
Benchmarking and Standardization of Evaluation Protocols: A Feedback-Driven Framework Using LLM Judges to Gatekeep and Iteratively Improve Synthetic Benchmarks

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 Most evaluation pipelines treat LLM judges as scorers or one-shot filters: models
2 generate items, a rubric assigns scores, and low-quality samples are discarded.
3 We take a different path. We position LLM judges as gatekeepers that actively
4 improve synthetic data through a nine-layer, iterative grading and feedback loop.
5 Each candidate prompt–response pair is scored against targeted rubrics (schema
6 conformity; BLUF/CTA quality; MECE structure; numeric/evidence consistency;
7 risk→mitigation→guardrail completeness; factuality; tone/audience fit; novelty/
8 contamination; CTA feasibility). When a layer fails, the judge emits machine-
9 actionable repair instructions; the item is revised or regenerated, re-evaluated, and
10 only admitted after passing all nine layers. Unlike prior paradigms that log evalua-
11 tions as by-products, we publish schema-based audit traces (per-layer scores, repair
12 histories, judge versions, similarity fingerprints) as first-class benchmark artifacts,
13 enabling contamination checks, reproducibility, and governance. Applied to six
14 structured genres, this closed-loop gatekeeping produces higher-quality synthetic
15 datasets that better align with human raters and yield more stable model deltas
16 than ungated or one-pass filtered baselines. We release rubric prompts, repair
17 templates, audit schemas, and evaluation scripts to support standardized, auditable
18 benchmarking.

19 1 Introduction

20 Reliable evaluation and fine-tuning of large language models (LLMs) for structured, high-stakes
21 genres—executive briefs, strategy memos, investment analyses, launch decisions, legal cases, and
22 policy memos—run into a persistent bottleneck: high-quality, balanced, and compliant training/e-
23 valuation data is scarce, sensitive, and heterogeneous. Naïve synthetic generation offers scale but
24 routinely injects noise, factual drift, stylistic inconsistency, and contamination risks. In domains
25 where errors carry real cost, “more data” is not a solution if quality is uncontrolled.

26 Much of today’s evaluation stack reflects two paradigms. (i) *Teacher–student* and related synthetic-
27 generation schemes create large corpora and then apply generic cleaning or reward models. (ii)
28 *LLM-as-judge* systems score outputs with multidimensional rubrics, sometimes using ensembles, and
29 often filter once: accept high scores, discard the rest. These paradigms help, but embed two gaps we
30 target directly: (1) judging is treated as *measurement*, not *control*; and (2) failures are observed but
31 rarely repaired before data enters a benchmark or fine-tuning set.

We treat LLM judges as **gatekeepers and corrective agents** in a nine-layer, iterative grading and feedback loop. Every candidate prompt–response pair is routed through targeted, domain-aware rubrics: schema conformity, BLUF/CTA quality, MECE organization, numeric and evidence consistency, risk→mitigation→guardrail completeness, factuality, tone/audience fit, novelty/contamination, and CTA feasibility. Failures yield *prescriptive* repairs; items are revised and re-evaluated until they pass or are rejected after bounded retries. Accepted items come with per-item audit trails.

Contributions.

1. A standardized, feedback-driven protocol that routes every candidate through nine rubric layers with ensemble judging, explicit thresholds, and bounded retries; failures trigger machine-actionable repair.
2. A schema-based audit specification that captures prompts, per-layer scores, failure modes, repair traces, judge identities/versions, and similarity fingerprints for contamination analysis.
3. An evaluation program contrasting this closed-loop approach against ungated and one-pass filtered baselines, including ablations by layer, ensemble size, thresholds, and model scale.
4. Demonstrated generality. Applied across six structured genres (executive briefs, strategy memos, investment briefs, launch decisions, legal cases, policy memos), showing how the same evaluation standard can span diverse domains.

2 Methodology

We target six structured genres: *executive brief*, *strategy memo*, *investment brief*, *launch decision*, *legal case*, *policy memo*. The pipeline operates in two stages under one principle: *immediate judging with prescriptive repair*.

2.1 Stage 1: Prompt synthesis with in-loop grading/gating

Modes. Two generation systems:

- **Strict (logic-based):** requires one-sentence BLUF and one-line CTA with mirror-rule enforcement.
- **Narrative-based:** forbids BLUF/CTA; requires 2–3 sentence overview with narrative anchors (benchmarks, precedents, strategy fit).

Both enforce: (i) explicit length band (e.g., 650–900 words); (ii) “return only the document text”; (iii) numbers/units policy; (iv) at least one Risk→Mitigation→Guardrail triplet. Category-specific hint banks encode scaffolds. Per-category strictness probabilities (e.g., legal case 55% strict; investment brief/launch decision 70%) choose modes.

Grading. Each prompt is graded by a category- and mode-specific judge returning JSON:

```
{"score": 1-5, "reason": "<short>"}
```

Must-haves are enforced per mode; scores: 1 (unstable) ... 5 (excellent).

Gatekeeping, retries, audits. Accept if $\text{score} \geq 4$ (overrides per category), retry up to MAX_ATTEMPTS, de-duplicate, and log to CSV/JSONL with `id`, `category`, `text`, `score`, `reason`, `attempts`, `mode`. Provenance includes generation mode, grading result, and attempts.

2.2 Stage 2: Document drafting with judge-gated evaluation

For each accepted prompt, we draft a document then evaluate against rubric checks.

Ingestion & mode. Prompts are loaded from `good_prompts.jsonl`, preserving the strict/narrative flag. `buff=True` implies BLUF/CTA required; `buff=False` forbids them.

Grader schema. The document grader returns JSON with:

- overall score,

- structured checks map (length band, BLUF/Overview/CTA flags, mirror rule, risk triplet, sections, units, acronyms, assumptions),
- sub-scores (structure, constraints, clarity, compliance).

Strict applies the mirror rule; narrative penalizes BLUF/CTA if present. Token ceilings and per-item logs are enforced.

2.3 Rubric suite (A1–A9) and global checks

2.4 Gating logic and iterative repair

On failure, graders return per-check flags (e.g., "mirror_rule_ok": false) with reasons ("BLUF missing numeric anchor"). These reasons are converted into explicit repair directives for targeted regeneration.

2.5 Example run (strict mode)

Step 1: Candidate prompt. “Create a Strategy Memo (700–900 words). Return only the document text. Include sections: Situation Overview; Options; Risks; Recommendation. Start with a one-sentence BLUF. Add a Call to Action with owner, budget, fewer than 2 milestones, and a success metric. Include at least one Risk→Mitigation→Guardrail triplet. Ensure all units are consistent.”

Step 2: Grading.

```
{"score": 3, "reason": "BLUF missing numeric anchor; CTA milestones lack dates"}
```

Step 3: Iterative repair.

```
Repair: Add numeric anchor to BLUF |
Repair: Add explicit milestone dates (e.g., Q2 2026)
```

Regenerated prompt: “BLUF: Decide immediately to allocate \$2M over 18 months... Milestones: complete pilot by June 2026; rollout by Dec 2026...”

Step 4: Acceptance.

```
{"score": 5, "reason": "All constraints satisfied"}
```

Logged to good_prompts.jsonl with {id, category=strategy_memo, mode=strict, buff=True, score=5, reason, attempts=2}.

Step 5: Document drafting.

```
{"overall": 4.5,
 "checks": {"BLUF_ok": true, "CTA_ok": true, "risk_triplet_ok": true,
            "mirror_rule_ok": true, "units_ok": true, "assumptions_ok": true},
 "reason": "Minor clarity issues in Options section"}
```

3 Results

3.1 From A1 to A9 — Closing the loop

A1 BLUF discipline. Vague openings became urgent, time-bound overviews.

```
Before: Our company faces a critical decision regarding the adoption of a new CRM
system...
After: By the end of Q2, prompted by a 20% decline in CSAT, we must evaluate
options...
```

A2 Section structure. Documents missing assumptions or duplicating content were reorganized into: Executive Summary → Background → Analysis → Stakeholders → Risks&Guardrails → Recommendations → Implementation → Assumptions → Acronyms → Units → Guardrail Enforcement.

127 **A3 Anchors & evidence.** General trade-offs were made specific with quantitative anchors and
128 sources.

129 Before: Option A ... may require significant upfront investment and training.
130 After: Option A ... may require significant upfront investment of \$100,000 USD.
131 Comparable case: Salesforce (McKinsey, 2020).
132

134 **A4 CTA completeness.** Weak CTAs were hardened into measurable directives.

135 Before: The organization should consider implementing a new CRM system.
136 After: Decide within 6 weeks; success measured by 20% CSAT, 10% revenue growth,
137 80% adoption.
138

140 **A5 Consistency & assumptions.** Hidden leaps of logic were surfaced as tagged assumptions.

141 Before: The new system will likely improve sales productivity and customer
142 satisfaction.
143 After: Assumption: Sales productivity will increase by 10% and satisfaction by 20%
144 (industry benchmarks).
145

147 **A6 Risk triplets.** Narrative risks were restructured into enforceable control logic.

148 Risk: Disruption of sales | Mitigation: Phased rollout | Guardrail: Pivot if sales
149 drop >5%
150 Risk: Data breaches | Mitigation: Encryption+controls | Guardrail: Quarterly
151 audits; escalate on breach
152
153

154 **A7 Completeness.** Acronyms expanded (CRM, IT, GDPR, CCPA); units normalized; guardrail
155 enforcement sections added with triggers/cadence.

156 **A8 Factuality & legal/source anchors.** Generic “security” became jurisdiction-anchored obligations.

157 Before: Robust security measures such as encryption and access controls...
158 After: Implementation must comply with GDPR and CCPA; conduct quarterly audits,
159 document DPIAs; enforce access controls and encryption at rest/in transit.
160

162 **A9 Tone/audience & units normalization.** Plain headings and consistent units improved readability.

163 Before: Option A requires investment of \$100,000
164 After: Option A requires investment of \$100,000 USD
165 Units of Measurement: USD; Percentage (%)
166
167

168 **Final gates.** CTA feasibility (owner/timing/budget coherence) validated that the plan is time-boxed,
169 budget-bounded, and measurable. Novelty/contamination scans passed for the CRM example.

170 **Net impact.** By A1–A7, drafts met structural discipline, CTA completeness, risk hygiene, and
171 assumptions transparency. Post-A7 gates added legal anchors, units normalization, and feasibility/-
172 contamination checks. Compared to ungated generation, discard rates dropped (more repairs, fewer
173 wasted generations), agreement with human raters increased, and model deltas stabilized. Across 8k
174 items, the protocol reduced discard rates by roughly 40

175 4 Discussion

176 **Limitations.** Judge bias can homogenize style; prescriptive repair may suppress creativity. Costs rise
177 with retries. Mitigations include diverse judge ensembles (different model families), periodic human
178 calibration against gold samples, and retry budgets. Subjective qualities (tone originality) remain
179 challenging.

180 **Broader implications.** Elevating per-item audits (scores, repairs, judge versions, fingerprints) to
181 first-class artifacts enables reproducibility, contamination checks, and governance. The framework
182 extends to code (tests/lint/security gates), science (citation/data anchors), and education (rubrics for
183 fairness).

References

- [1] Y. Bai et al. Constitutional AI: Harmlessness from AI Feedback. *arXiv:2212.08073*, 2023.
- [2] W.-L. Chiang et al. Chatbot Arena: An Open Platform for Evaluating LLMs by Human Preference. *arXiv:2403.04132*, 2024.
- [3] S. Duan et al. MDBench: A Large-scale Multi-document Benchmark for Long-context LLMs. *arXiv:2506.07257*, 2025.
- [4] P. Ganesh et al. LAB: Large-scale Alignment for ChatBots. *arXiv:2403.01081*, 2024.
- [5] S. Kim et al. AgoraBench: Dynamically Generated Data for Community-driven Evaluation of LLMs. In *Proc. ACL*, 2025.
- [6] LangChain. LangSmith Documentation. 2024. <https://www.langchain.com/langsmith>.
- [7] Langfuse. Langfuse Documentation. 2024. <https://langfuse.com/docs>.
- [8] A. Lee et al. Scaling Reinforcement Learning from AI Feedback. In *ICLR*, 2025.
- [9] Imarena. Arena-Hard-Auto: An Automatic LLM Benchmark. GitHub, 2024.
- [10] LMSYS / Emergent Mind. Arena-Hard Benchmarking Standard. 2024.
- [11] NVIDIA. Nemotron-4 340B Technical Report. *arXiv:2406.11704*, 2024.
- [12] H. Rahmani et al. JudgeBlender: Teaching Models to Blend their Judges. *arXiv:2411.02101*, 2024.
- [13] Z. Tan et al. JudgeBench: A Benchmark for Evaluating LLM-Based Judges. *OpenReview*, 2025.
- [14] P. Verga et al. Replacing Judges with Juries: Evaluating LLM Generations with a Panel of Diverse Models. *arXiv:2404.18796*, 2024.
- [15] P. Yu et al. R.I.P.: Better Models by Survival of the Fittest Prompts. *arXiv:2501.18578*, 2025.
- [16] L. Zheng et al. Judging LLM-as-a-Judge with MT-Bench and Chatbot Arena. *arXiv:2306.05685*, 2023.