

LLM-as-a-Verifier: A General-Purpose Verification Framework

Jacky Kwok¹, Shulu Li², Pranav Atreya², Yuejiang Liu¹, Yixing Jiang¹

Chelsea Finn¹, Marco Pavone^{1,3}, Ion Stoica², Azalia Mirhoseini¹

¹Stanford University ²UC Berkeley ³NVIDIA Research

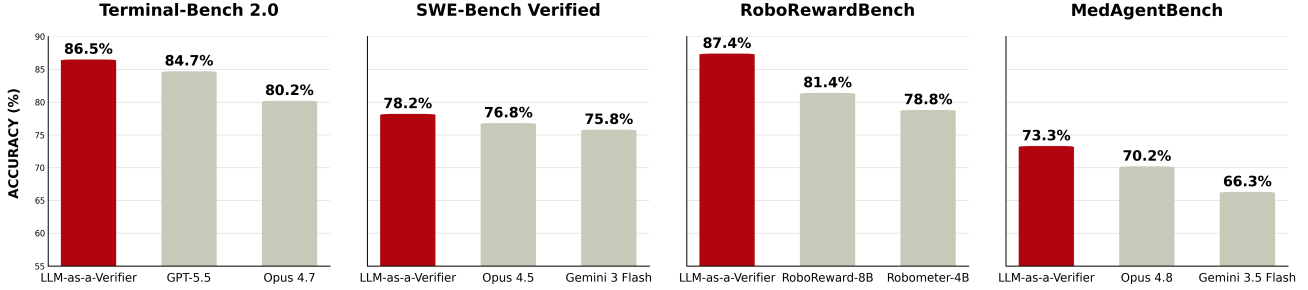


Figure 1: Overall Performance Results. Our proposed framework, **LLM-as-a-Verifier**, achieves state-of-the-art performance across coding, robotics, and medical domains: Terminal-Bench V2 (86.5%), SWE-Bench Verified (78.2%), RoboRewardBench (87.4%), and MedAgentBench (73.3%).

Abstract: Scaling pre-training, post-training, and test-time compute have become the central paradigms for improving the capabilities of large language models (LLMs). In this work, we identify verification—the ability to determine the correctness of a solution—as a new scaling axis. To unlock this and demonstrate its effectiveness, we introduce LLM-as-a-Verifier, a general-purpose verification framework that provides fine-grained feedback for agentic tasks without requiring additional training. Unlike standard LM judges that prompt LLMs to produce discrete scores for candidate solutions, LLM-as-a-Verifier computes the expectation over the distribution of scoring token logits to generate continuous scores. This probabilistic formulation substantially reduces tie rates when comparing complex solutions and enables verification to scale along multiple dimensions: (1) score granularity, (2) repeated evaluation, and (3) criteria decomposition. In particular, we show that scaling the scoring granularity leads to better separation between positive and negative solutions, resulting in more calibrated comparisons. Moreover, scaling repeated evaluation and criteria decomposition consistently leads to additional gains in verification accuracy through variance and complexity reduction. To make verification scaling practical, we further introduce a cost-efficient ranking algorithm for selecting the best solution among candidates using the preference probabilities derived from the verifier’s continuous scores. LLM-as-a-Verifier is effective across coding, robotics, and medical domains. It achieves state-of-the-art performance on Terminal-Bench V2 (86.5%), SWE-Bench Verified (78.2%), RoboRewardBench (87.4%), and MedAgentBench (73.3%). Beyond verification, the fine-grained signals from LLM-as-a-Verifier can also serve as a proxy for estimating task progress. We build extensions for Claude Code and Codex, enabling developers to monitor and improve their own agentic systems. Finally, we show that LLM-as-a-Verifier can be used as a dense reward signal for RL, improving the sample efficiency of SAC and GRPO on robotics and mathematical reasoning benchmarks.

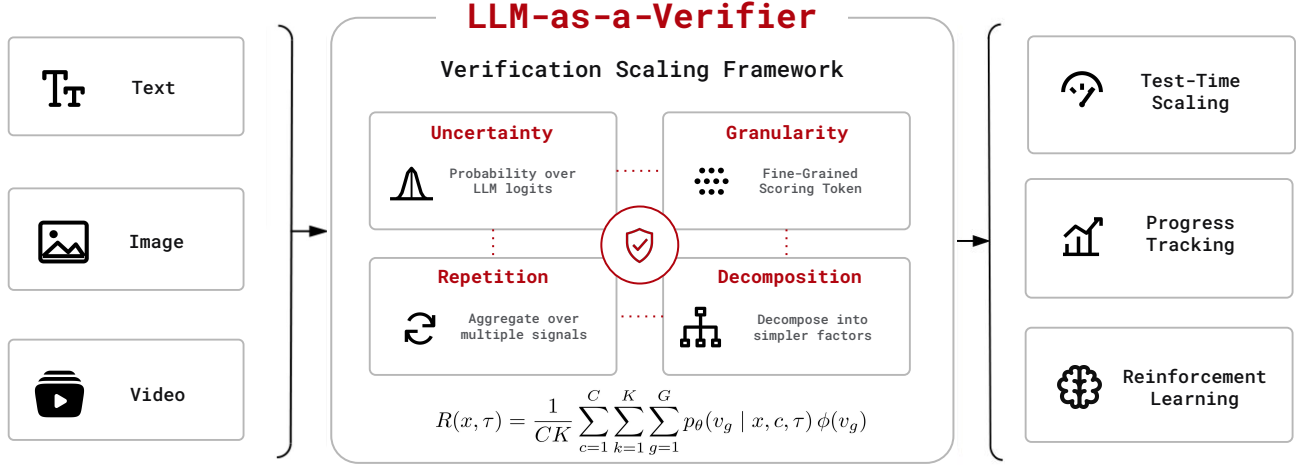


Figure 2: Multiple modalities, many applications, one unified verification framework. We present LLM-as-a-Verifier, a general-purpose framework that provides fine-grained feedback for any modality without requiring additional training. By leveraging the full distribution of scoring-token logits, our method captures evaluation uncertainty and enables verification to scale along three dimensions: score granularity, repeated evaluation, and criteria decomposition. The resulting fine-grained feedback can be used for test-time scaling, progress tracking, and reinforcement learning.

1. Introduction

Recent advances in large language models (LLMs) have established scaling as a central paradigm for improving their capabilities. Performance has been driven by scaling along multiple axes, including pre-training data and compute, post-training optimization, and test-time inference [1–3]. However, while generation has benefited significantly from these scaling paradigms, verification—the ability to determine the quality or correctness of a solution—has not seen the same degree of scaling. In this work, we argue that verification itself constitutes a distinct and underexplored scaling axis. Unlike generation, which benefits from well-established scaling laws, verification in current systems remains fundamentally limited. In particular, standard LM judges collapse scoring distributions into coarse discrete scores [4, 5], leading to ties and poor discrimination, while learned reward models are constrained by training data and often fail to generalize across domains [6, 7]. These limitations hinder the scalability of verification, preventing further performance improvements.

To this end, we introduce LLM-as-a-Verifier, a general-purpose verification framework that provides dense and fine-grained feedback without requiring additional training. Unlike traditional approaches that prompt LLMs to produce discrete scores within the language space [4], LLM-as-a-Verifier estimates the quality of candidate solutions by computing the expectation over the distribution of scoring token logits. In Fig. 4, we show that this probabilistic formulation unlocks multiple axes of scaling for verification. We first demonstrate that scaling the number of extracted token logits consistently reduces the tie rate when comparing complex solutions and improves the separation between positive and negative solutions. We observe that an individual evaluation or a single criterion

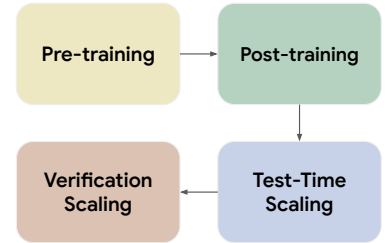


Figure 3: Scaling paradigms for large language models.

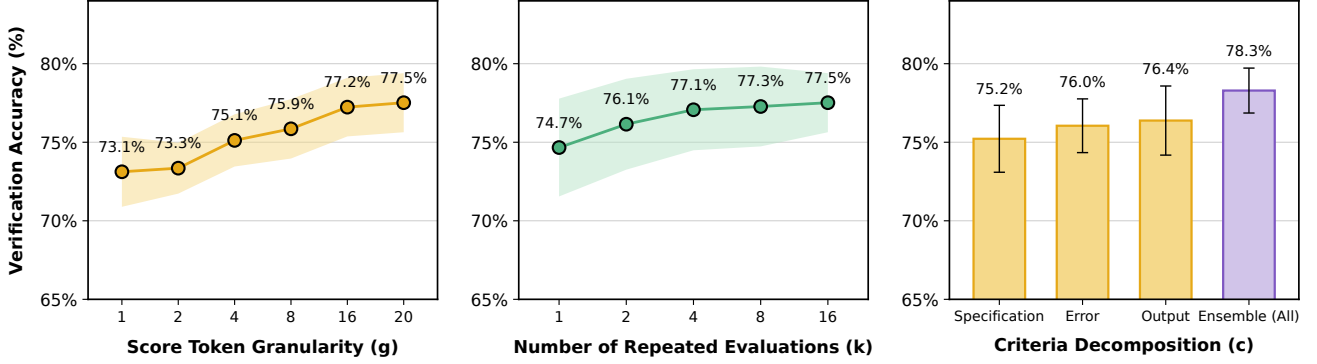


Figure 4: Verification Scaling. We find that verification accuracy consistently improves as we scale across multiple dimensions: (1) the granularity of score tokens, (2) the number of repeated evaluations, and (3) the decomposition of evaluation criteria. Verification accuracy is measured as the pairwise accuracy of the verifier in assigning a higher score to the ground-truth successful solution than to failed solutions for the same task on Terminal-Bench V2.

can be biased or noisy. To mitigate this, we scale verification along two additional dimensions, repeated evaluations (which reduces variance) and criteria decomposition (which reduces prompt bias), leading to higher verification accuracy. We quantify these scaling benefits under controlled budgets, comparing LLM-as-a-Verifier against a discrete LM judge baseline in Sec. 4. To make verification scaling practical, we further introduce a cost-efficient ranking algorithm for selecting the best solution among candidates using the preference probabilities derived from the verifier’s continuous scores.

Interestingly, we find that the fine-grained signals produced by LLM-as-a-Verifier enable the evaluation of entire interaction trajectories rather than only intermediate steps or final outcomes as in PRMs and ORMs [7, 8] for agentic tasks. When used as a trajectory reward model with our cost-efficient ranking algorithm, LLM-as-a-Verifier outperforms frontier models on challenging benchmarks across coding, robotics, and medical domains. It achieves state-of-the-art performance on Terminal-Bench V2 (86.5%), SWE-Bench Verified (78.2%), RoboRewardBench (87.4% Trajectory Preference Accuracy), and MedAgentBench (73.3%).

Beyond its role as a verifier, our approach can also serve as a proxy for estimating task progress. Notably, we observe a strong correlation between the chronological order of steps and the verifier score (Fig. 8). To instantiate these capabilities, we provide extensions for Claude Code and Codex, enabling users to monitor task progress and harness the benefits of LLM-as-a-Verifier to improve their own agentic systems. In robotics, our approach outperforms state-of-the-art reward models, including Robometer [9], TOPReward [10], and RoboReward [11], achieving a mean Value-Order Correlation (VOC) of 0.966. Overall, LLM-as-a-Verifier provides a scalable mechanism for improving the evaluation and monitoring of autonomous agents and robots in real-world environments.

Additionally, we demonstrate that using LLM-as-a-Verifier as a dense reward signal improves the sample efficiency of both off-policy and on-policy reinforcement learning algorithms. On LIBERO [12], LLM-as-a-Verifier achieves $\approx 1.8\times$ higher sample efficiency than sparse reward baselines when fine-tuning a π_0 policy with DSRL-SAC [13], while also reaching a higher final success rate. On the MATH reasoning benchmark, it achieves $\approx 1.1\times$ higher sample efficiency when fine-tuning Qwen3-8B with GRPO [14].

In summary, our contributions are as follows:

1. We introduce LLM-as-a-Verifier, a probabilistic verification framework that leverages the full distribution of scoring token logits to produce fine-grained feedback and characterize three key axes of verification scaling: (1) score granularity, (2) repeated evaluation, and (3) criteria decomposition.
2. We propose a cost-efficient algorithm for ranking candidates and demonstrate that, when combined with verification scaling, LLM-as-a-Verifier achieves state-of-the-art performance across coding, robotics, and medical benchmarks without requiring additional training.
3. We show that the fine-grained verifier score correlates with an agent’s task progress and can be used to monitor the behavior of agents and robots.
4. We demonstrate that LLM-as-a-Verifier can provide dense feedback for reinforcement learning, improving the sample efficiency of both on-policy and off-policy algorithms across robotics and mathematical reasoning benchmarks.

2. Preliminaries

We model an agent interacting with an environment as a finite-horizon Markov Decision Process (MDP) $\mathcal{M} = (\mathcal{C}, \mathcal{S}, \mathcal{A}, P, R, H)$, where $\mathcal{C}$ denotes the space of contexts, $\mathcal{S}$ the state space, $\mathcal{A}$ the action space, $P : \mathcal{C} \times \mathcal{S} \times \mathcal{A} \rightarrow \Delta(\mathcal{S})$ the transition dynamics, $R : \mathcal{C} \times \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$ the reward function, and $H \in \mathbb{N}^+$ the horizon. At the beginning of each episode, a task prompt $x \in \mathcal{C}$ is sampled, and the agent begins in an initial state $s_1 \in \mathcal{S}$. At each timestep $t \in [1, H]$, the agent observes the current state s_t , selects an action $a_t \in \mathcal{A}$, and transitions to the next state $s_{t+1} \sim P(\cdot \mid x, s_t, a_t)$. In LLM-based agents, states correspond to prior interaction histories, and actions correspond to token sequences, such as natural language responses, code edits, and tool calls. A trajectory is defined as $\tau = (s_1, a_1, s_2, a_2, \dots, s_H, a_H)$. We assume access to a language model $\pi_\theta : \mathcal{C} \times \mathcal{S} \rightarrow \Delta(\mathcal{A})$, parameterized by θ , from which actions are sampled autoregressively. A reward model assigns a scalar score to actions or trajectories. Conventional approaches rely on prompting LLMs to produce discrete scores in the language space. Formally, such reward models can be written as $R_{\text{LM}}(x, \tau) \in \{1, \dots, G\}$, where the score is the generated token.

3. Proposed Approach: LLM-as-a-Verifier

3.1. Motivation

Most models already possess the capability to solve many tasks: when executed repeatedly, they often produce a correct solution at least once. As shown in Fig. 5 (left), the fraction of solved tasks increases consistently as we scale the number of sampled trajectories on Terminal-Bench, assuming access to an oracle verifier that always picks the optimal trajectory. Under this setting, the success rate reaches 98.9% when pooling trajectories across the full Terminal-Bench V2 leaderboard, effectively solving nearly the entire benchmark. However, capturing this headroom requires a verifier that can reliably distinguish correct trajectories from incorrect ones. While standard LM judges [4] can be used as verifiers, they fail to provide sufficiently fine-grained feedback. Specifically, they prompt

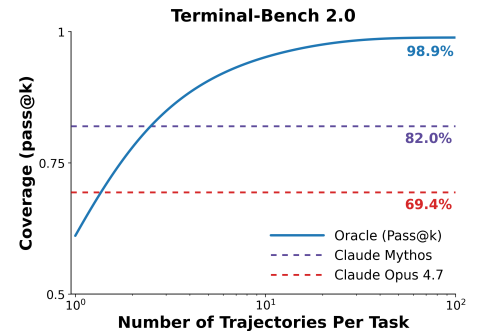


Figure 5: Oracle Pass@ K reaches 98.9% on Terminal-Bench V2.

the model to output a discrete score token and select the highest-probability token as the final score, collapsing the full scoring distribution into a single value. This leads to inherently coarse evaluations. When comparing complex solutions, standard LM judges often assign the same score, resulting in ties and failing to discriminate between them. As a result, coarse scoring induces a high tie rate (27%) on Terminal-Bench, with distinct trajectories often collapsing to the same score, as illustrated in Figure 7. One could instead train a reward model [15], but such methods are constrained by their training data and often fail to generalize across domains. These limitations motivate the need for a generalizable framework that can provide fine-grained verification signals.

3.2. Methodology

Fine-Grained Reward Estimation. By definition, a judge is one who forms an overall opinion and assigns a decision, whereas a verifier is one who confirms the truth or correctness of something and requires more detailed evaluations. To this end, we introduce LLM-as-a-Verifier, a probabilistic verification framework that provides fine-grained feedback by scaling scoring granularity, repeated evaluation, and criteria decomposition.

Let $V_{\text{score}} = \{v_1, \dots, v_G\}$ denote an ordered set of tokens representing discrete score levels. Given a task prompt x , a language model p_θ , a criterion c , and two candidate trajectories τ_i and τ_j , we construct scoring prompts and obtain their conditional distributions $p_\theta(v \mid x, c, \tau_i)$ and $p_\theta(v \mid x, c, \tau_j)$ by extracting the logprobs from $\langle \text{score}_A \rangle$ and $\langle \text{score}_B \rangle$ tags using the following prompt:

```
You are an expert [domain] reviewer. You will see a task description and two trajectories.
Evaluation Criteria: [domain specific criteria]
Task: {task prompt}
Trajectory A: {A} Trajectory B: {B}
Carefully analyze each trajectory, then provide your final scores:

<score_A> INTEGER_1_TO_20 </score_A>
<score_B> INTEGER_1_TO_20 </score_B>

Rating Rules: Rate correctness on a 1-20 scale based on evaluation criteria (1 = incorrect,
10 = borderline, 20 = correct)

Note: We use a letter-based scale instead of digits to enable logprob extraction for
granularity scaling.
```

Rather than collapsing each distribution to a single discrete score, we approximate the reward of a trajectory as:

$$R(x, \tau) = \frac{1}{CK} \sum_{c=1}^C \sum_{k=1}^K \sum_{g=1}^G p_\theta(v_g \mid x, c, \tau) \phi(v_g) \quad (3.1)$$

where C is the number of evaluation criteria, K is the number of repeated verifications, G is the number of score tokens (granularity level), $p_\theta(v_g \mid x, c, \tau)$ is the probability assigned by model θ to score token v_g , and $\phi(v_g)$ maps each score token to a scalar value.

We first normalize $R(x, \tau) \in [0, 1]$ by the linear map $R \mapsto (R - \phi_{\min})/(\phi_{\max} - \phi_{\min})$. Then, we convert these continuous rewards into a pairwise preference using the Bradley–Terry model, treating $R(x, \tau)$ as

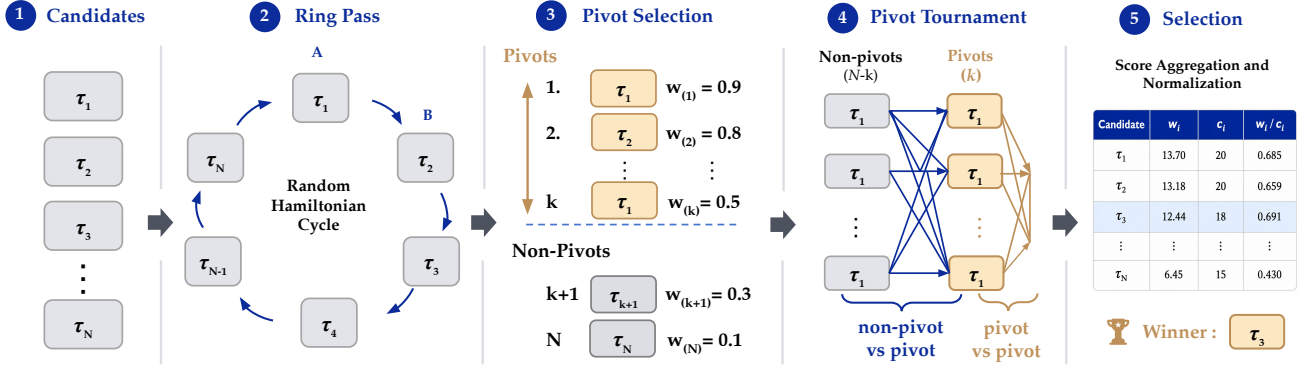


Figure 6: Probabilistic Pivot Tournament. A five-stage pipeline for selecting the best of N candidates under a constrained verification budget. **(1) Candidates:** the pool $\{\tau_1, \dots, \tau_N\}$ to be ranked. **(2) Ring pass:** a random Hamiltonian cycle scores the N adjacent pairs so every candidate appears once in the “A” slot and once in “B”, canceling the model’s positional bias. **(3) Pivot selection:** candidates are ranked by their ring-pass scores $w_{(i)}$, and the top- k candidates form the pivot set $\mathcal{P}$. **(4) Pivot tournament:** every *non-pivot-vs-pivot* and *pivot-vs-pivot* pair is scored via Eq. 3.2, concentrating the budget on uncertain top candidates and cutting cost from $\mathcal{O}(N^2)$ to $\mathcal{O}(Nk)$. **(5) Selection:** comparisons are aggregated into win mass w_i and count c_i , and the candidate with the highest normalized w_i/c_i is returned.

the latent strength of trajectory τ :

$$P(\tau_i \succ \tau_j \mid x) = \frac{1}{1 + \exp(-(R(x, \tau_i) - R(x, \tau_j)))}, \quad (3.2)$$

Probabilistic Pivot Tournament. To pick the best trajectory among N candidates, we can run a round-robin tournament that scores all $\binom{N}{2}$ pairs and accumulates wins

$$w_i += P(\tau_i \succ \tau_j \mid x), \quad w_j += 1 - P(\tau_i \succ \tau_j \mid x),$$

using the preference probability of Eq. 3.2. However, such a schedule scales as $\mathcal{O}(N^2)$ pairwise verifications and quickly dominates verifier cost as N grows. We propose a budget-efficient alternative, Probabilistic Pivot Tournament (PPT), illustrated in Fig. 6, in which every candidate is compared only against a small set of $k \ll N$ pivots, reducing the budget from $\mathcal{O}(N^2)$ to $\mathcal{O}(Nk)$. Critically, the choice of pivots determines whether the saved budget is well spent: arbitrary anchors waste verifications on candidates that are clearly weak. We therefore introduce a *ring-based pivot selection* step that both removes the verifier’s positional bias and concentrates the remaining budget on uncertain top candidates. PPT proceeds in three steps:

1) *Ring pass.* We sample a uniformly random Hamiltonian cycle γ over $\{1, \dots, N\}$ and score the N adjacent pairs $\{(\gamma_t, \gamma_{t+1 \bmod N})\}_{t=1}^N$. By the cyclic structure, every candidate appears *exactly once in the “A” position and once in the “B” position* of the verifier prompt, so any systematic preference of the language models for one slot over the other cancels in expectation across the ring.

2) *Pivot selection.* We rank candidates by their ring-pass mean preference w_i/c_i and choose the top- k as the pivot set $\mathcal{P}$. Selecting pivots from the empirical leaders allocates the remaining verification budget

to the candidates most likely to be correct, so the subsequent pairwise comparisons distinguish among uncertain top candidates rather than spending queries on weak anchors.

3) *Pivot rounds.* With the pivot set fixed, we score (i) every *non-pivot vs. pivot* pair (i, p) with $i \notin \mathcal{P}$, $p \in \mathcal{P}$, and (ii) every *pivot vs. pivot* pair within $\binom{\mathcal{P}}{2}$. All ring and pivot-round comparisons are aggregated into the same w_i , c_i , and we select $i^* \in \arg \max_i w_i/c_i$. Normalizing by c_i removes the bias that pivots participate in more comparisons than non-pivots. The total number of pairwise verifications is $N + k(N - k) + \binom{k}{2}$ which scales as $\mathcal{O}(Nk)$ where $k \ll N$. The full generation and verification pipeline is given in Algorithm 1 (Appendix B.2).

To rigorously evaluate the ranking algorithms and assess performance on large candidate pools, we curate 20 trajectories per task using the Terminus-2 harness, and benchmark all methods in this setting. Table 9 characterizes the budget–accuracy trade-off of PPT, showing that our method outperforms prior approaches (e.g., V1 [5]) while requiring fewer comparisons. Notably, performance improves consistently as the number of pivots increases. Further ablations in Appendix B.2.

4. Verification Scaling

Equation 3.1 illustrates three independent axes along which verification can be scaled: the granularity of score tokens G , the number of repeated evaluations K , and the number of evaluation criteria C . Each axis targets a different source of error in the reward estimate, and we find that the three act as complementary levers: increasing granularity improves score separation between candidate solutions, repeated evaluation averages out biases from individual verification passes, and criteria decomposition captures complementary aspects of trajectory quality. For all scaling experiments, we use Gemini 2.5 Flash [16] as the verifier, which allows us to extract up to 20 top logprobs per scoring token. In Fig. 4, we show that verification accuracy on Terminal-Bench 2.0 improves along all three dimensions, rising from 73.1% at $G=1$ to 77.5% at $G=20$, from 74.7% at $K=1$ to 77.4% at $K=16$, and from 75.2%–76.4% for any single criterion to 78.3% when the three criteria are ensembled. We measure the pairwise verification accuracy over 200 randomly sampled trajectories from Terminal-Bench, spanning multiple agent harnesses. Each axis is a knob that the practitioner can tune depending on the latency budget of the downstream application. While our primary experiments use a logprob-accessible model, Appendix B.6 demonstrates that our framework is also compatible with frontier models that do not expose token-level log-probabilities via a simple two-stage workaround.

4.1. Scoring Token Granularity

Standard LM judges collapse the scoring distribution to the single highest-probability token, yielding a discrete reward $R_{\text{LM}}(t, \tau) \in \{1, \dots, G\}$ with resolution $1/G$. Intuitively, enlarging the ordered token set V_{score} does not grant the verifier any new information about the trajectory. Yet, it grants the decoder a finer space in which to project the model’s internal belief, so that nearby beliefs that would have been rounded to the same integer are now mapped to continuous rewards.

Signal-to-Noise Ratio. To isolate why finer granularity improves verification, we decompose the pairwise score gap $\Delta = s_c - s_i$ between correct (s_c) and incorrect (s_i) trajectories into a signal and a noise component. We define the signal-to-noise ratio as in Eq. 4.1 (Table 1, left), where $\mathbb{E}(s_c - s_i)$ captures how strongly the verifier prefers the correct trajectory over the incorrect one (*signal strength*),

$$\text{SNR}(G) = \frac{\mathbb{E}[s_c - s_i]}{\sqrt{\text{Var}(s_c - s_i)}} \quad (4.1)$$

Granularity G	1	4	16	20
SNR ($k=16$)	0.775	0.786	0.797	0.799

Table 1: **Signal-to-noise ratio (SNR).** (Left) The SNR measures how reliably the verifier separates correct (s_c) from incorrect (s_i) trajectories (Eq. 4.1). (Right) As the number of scoring tokens G increases, the SNR grows, indicating better-calibrated score separation.

and the denominator, $\text{Var}(s_c - s_i)$, captures how inconsistent that preference is across pairs (*noise*). Pairwise verification accuracy is a monotonic function of $\text{SNR}(G)$: holding sample size fixed, a larger standardized gap implies a higher probability that $s_c > s_i$. Empirically, we find that $\text{SNR}(G)$ increases from 0.775 at $G=1$ to 0.799 at $G=20$ on Terminal-Bench (Table 1). Finer-grained tokens therefore produce better-calibrated scores that more reliably separate correct from incorrect trajectories, which in turn improves the pairwise accuracy from 73.1% to 77.5%.

Case Study: query-optimize. To concretely illustrate how scaling granularity to $G=20$ and our probabilistic formulation sharpen the verifier’s signal, we analyze a representative trajectory pair from the query-optimize task on Terminal-Bench V2, generated by Claude Opus 4.5 under the OpenHands harness and scored by Gemini 2.5 Flash. Here the agent is given a slow SQL query over a database and asked to produce an equivalent optimized version. Both candidate trajectories generate queries that execute faster, but they differ critically in their verification procedures. The correct trajectory waits the full 5 minutes for the original query to complete on the canonical database and performs a direct diff against the optimized output. In contrast, the failing trajectory never validates equivalence on the database and instead creates a new database. As shown in the reasoning traces in Appendix B.4, Gemini 2.5 Flash reliably identifies this failure mode, but expresses it in graded, hedged language (e.g., “*slightly cleaner*,” “*marginally more direct*”), as if the discrepancy were minor. When evaluated over 100 repetitions, a standard LM judge on a 1–5 scale collapses these nuanced assessments into discrete scores (Table 2), producing ties (e.g., 5 vs. 5) in 88 out of 100 runs, thus failing to meaningfully discriminate between the candidates. Taking the expectation over the *same* 5-point distribution eliminates ties entirely—ranking the correct trajectory higher in 69 runs—and scaling the granularity to $G=20$ sharpens the signal further, letting LLM-as-a-Verifier rank the correct trajectory strictly higher in 77 out of 100 runs.

Table 2: **Judges vs. Verifiers on query-optimize.** Over 100 repeated evaluations, we count how often the correct trajectory is scored higher than ($s_c > s_i$), tied with ($s_c = s_i$), or lower than ($s_c < s_i$) the incorrect one. The discrete 1–5 judge produces ties in 88/100 evaluations. Taking the expectation over the same 1–5 scale eliminates ties and correctly ranks the trajectory in 69/100 evaluations. Increasing the score granularity to $G=20$ further improves discrimination, correctly ranking the trajectory in 77/100 evaluations.

Method	Reward $R(x, \tau)$	$\#(s_c > s_i)$ ✓	$\#(s_c = s_i)$ (tie)	$\#(s_c < s_i)$ ✗
Judge (discrete, $G=5$)	$\phi(\arg \max_g p_\theta(v_g))$	12/100	88/100	0/100
Verifier (continuous, $G=5$)	$\sum_{g=1}^5 p_\theta(v_g) \phi(v_g)$	69/100	0/100	31/100
Verifier (continuous, $G=20$)	$\sum_{g=1}^{20} p_\theta(v_g) \phi(v_g)$	77/100	0/100	23/100

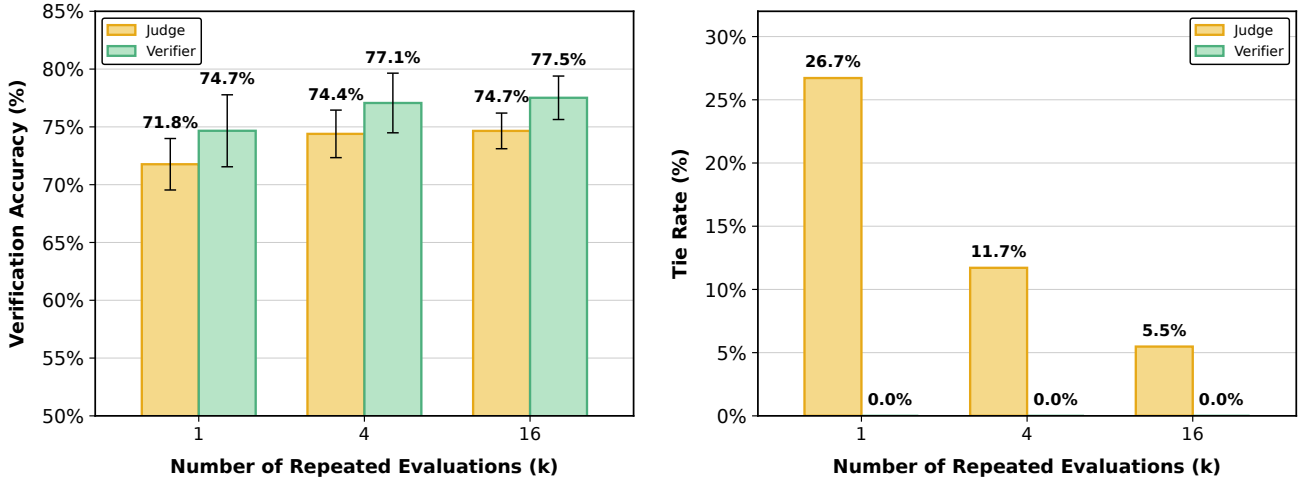


Figure 7: Verifier (continuous) vs. Judge (discrete) on Terminal-Bench V2 across $k \in \{1, 4, 16\}$ repeated evaluations. **Left: Pairwise verification accuracy. The verifier achieves 74.7% at $k=1$ and improves to 77.5% at $k=16$, consistently outperforming the judge across all evaluation budgets. **Right:** Tie rate. The judge produces ties in 26.7% of comparisons at $k=1$ due to coarse discrete scoring, decreasing to 5.5% at $k=16$ as averaging breaks ties. In contrast, the verifier yields zero ties.**

4.2. Repeated Evaluation

While granularity improves score calibration within a single forward pass, it does not address a second source of error: the verifier’s variance on one evaluation. Even at high G , a single evaluation $R^{(k)}(x, \tau)$ can be skewed by spurious features of the prompt or failure modes of the verifier on a particular trajectory. Averaging K independent evaluations $\frac{1}{K} \sum_{k=1}^K R^{(k)}(x, \tau)$ is a Monte Carlo estimator of the underlying expected reward; its variance shrinks as $\mathcal{O}(1/K)$ while its bias is unchanged. This complements granularity rather than duplicating it: granularity sharpens each individual estimator, while repeated evaluation averages out the noise that granularity cannot remove. Fig. 4 (middle) shows that the accuracy increases from 74.7% at $K=1$ to 77.5% at $K=16$. However, gains diminish with larger K : early improvements arise from variance reduction, while additional evaluations contribute diminishing returns due to correlated biases on harder examples. Importantly, repeated evaluation benefits discrete judges, whose coarse scoring induces high tie rates at low K . While increasing K helps break these ties through averaging, this mechanism is fundamentally limited for discrete judges. We show that a single-pass verifier ($K=1$) already matches a heavily ensembled judge ($K=16$), highlighting that fine-grained probabilistic scoring provides a stronger signal.

4.3. Criteria Decomposition

Granularity and repeated evaluation both assume that the rubric itself is adequate; neither helps if a single monolithic criterion is a poor proxy for trajectory quality. In long-horizon agentic tasks, judgments like “is this trajectory correct?” conflate several logically distinct factors, and verifiers asked a compound question often latch onto whichever factor is most salient in the prompt. Therefore, we replace a single monolithic rubric with an ensemble over C simpler sub-criteria. Concretely, for code-agent trajectories we decompose correctness into three factors that are individually easier to verify: Specification (whether the trajectory satisfies all task requirements), Output (whether the final output format matches the expected

result), and Errors (whether the trajectory is free of failure signals in logs and tool outputs). The final reward averages the expected scores across criteria, as in the outer sum of Eq. 1. In Fig. 4 (right), any one criterion alone achieves 75.2%–76.4% accuracy, and their ensemble reaches 78.3%.

5. Experiments

We evaluate LLM-as-a-Verifier as a trajectory reward model (TRM) for test-time scaling across four benchmarks that span three domains: coding (Terminal-Bench V2 [17], SWE-Bench Verified [18]), robotics (RoboRewardBench [11]), and medical (MedAgentBench [19]). Across all four, we use the same protocol: a generation policy π_θ produces N candidate trajectories per task, the verifier scores every pair using the probabilistic pivot tournament as described in Algorithm 1, and the trajectory with the highest normalized score is submitted. Unless otherwise noted, the verifier is run with granularity $G=20$, repeated evaluations $K=8$, and the three-criterion decomposition described in Section 4.3. Our method is training-free and plug-and-play: the same verification framework is applied across all four benchmarks without any per-domain fine-tuning. Overall results are summarized in Fig. 1, and per-benchmark headline numbers, including baseline accuracies, Pass@1, and oracle Pass@ N , are reported in Table 3.

Table 3: **Per-benchmark performance and gains from verification.** Baseline accuracies (left) are obtained under a fixed agent harness. On the same candidate pools, we report Pass@1, the oracle Pass@ N upper bound, and the accuracy achieved by LLM-as-a-Verifier (right). Our method consistently improves over Pass@1 and recovers a large portion of the oracle headroom, achieving state-of-the-art performance on each.

Benchmark	Baseline models (accuracy)			LLM-as-a-Verifier		
	#1	#2	#3	Pass@1	Oracle	Ours
Terminal-Bench V2	GPT-5.5 (84.7%)	Opus 4.7 (80.2%)	G3.1 Pro (80.2%)	83.1%	92.1%	86.5%
SWE-Bench Verified	Opus 4.5 (76.8%)	G3 Flash (75.8%)	M2.5 (75.8%)	76.1%	84.4%	78.2%
MedAgentBench	Opus 4.8 (70.2%)	G3.5 Flash (66.3%)	GPT-5.5 (65.1%)	70.2%	75.0%	73.3%

5.1. Terminal-Bench V2

Terminal-Bench V2 [17] measures an agent’s proficiency in shell-based environments across long-horizon tasks that require multi-step reasoning, file manipulation, and recovery from failed tool calls. The benchmark is particularly difficult for verifiers because many trajectories produce syntactically plausible but incorrect terminal states. We use Capy [20] as the scaffold and sample $N=5$ trajectories per task from GPT-5.5; Gemini 2.5 Flash serves as the verifier. The Pass@1 of GPT-5.5 under Capy is 83.1%, and the oracle Pass@5 upper bound on this candidate pool is 92.1%. LLM-as-a-Verifier improves the accuracy from 83.1% to **86.5%**, surpassing Claude Mythos + Terminus-2 [21] (82.0%), GPT-5.5 + NexAU-AHE (84.7%), Claude Opus 4.7 + WOZCODE (80.2%), and Gemini 3.1 Pro + TongAgents (80.2%) and setting a new state of the art on Terminal-Bench V2.¹ We further show that these gains are not tied to a specific harness. For additional generalization results on Terminus-2 and Terminus-Kira, refer to Appendix B.1.

¹Baseline accuracies are extracted from the official Terminal-Bench V2 and SWE-Bench Verified leaderboard.

5.2. SWE-Bench Verified

SWE-Bench Verified [18] is a human-curated subset of 500 real-world GitHub issues where each task requires an agent to produce a patch that resolves the issue and passes the maintainer’s hidden test suite. It stresses long-context reasoning, cross-file edits, and compliance with an existing codebase. We use mini-swe-agent as the scaffold and, in contrast to the homogeneous proposal pool used on Terminal-Bench, draw a heterogeneous pool of $N=3$ candidates per task by sampling one trajectory each from Claude Opus 4.5, Gemini 3 Flash, and MiniMax M2.5. Gemini 2.5 Flash again serves as the verifier. The mean Pass@1 across this candidate pool is 76.1% and the oracle Pass@3 upper bound is 84.4%. LLM-as-a-Verifier achieves **78.2%** on SWE-Bench Verified, outperforming Claude Opus 4.5 (76.8%), Gemini 3 Flash (75.8%), and MiniMax M2.5 (75.8%). These results highlight the verifier’s ability to select the strongest trajectory from a diverse set of candidates produced by different model families.

5.3. RoboRewardBench

RoboRewardBench [11] evaluates reward models on robotic manipulation trajectories. Following Liang et al. [9], we curate pairs of rollout videos that follow the same natural-language instruction but make different amounts of progress; the reward model must output a preference indicating which rollout makes more progress. Unlike the coding and clinical benchmarks, inputs here are multi-frame videos, so the verifier must integrate visual context across frames to reason about physical progress toward the goal. We use Qwen 3.6 35B as the base VLM verifier and apply the same probabilistic formulation (Eq. 1) over scoring tokens extracted from the VLM’s logits, with granularity $G=20$ and $K=8$ repeated verifications. We evaluate on 500 randomly sampled trajectory pairs from the RoboRewardBench and compare against (i) a discrete LLM-as-a-Judge baseline using the same VLM, (ii) reward models specifically trained on robotics data—RoboReward-8B (trained on $\sim 45k$ episodes) and Robometer-4B (trained on $\sim 1M$ comparisons), and (iii) TOPReward [10] (Qwen 3.6). As shown in Table 4 and Appendix B.5, LLM-as-a-Verifier achieves **87.4%** preference accuracy, outperforming the discrete LLM-as-a-Judge baseline (70.8%), RoboReward-8B (81.4%), Robometer-4B [9] (78.8%), and TOPReward (74.7%), despite being applied zero-shot and without any fine-tuning. We also evaluate on RoboRewardBench by measuring the Mean Absolute Error (MAE) between predicted rewards and human annotations. As shown in Table 5, using the continuous reward formulation in Eq. 3.1 together with $K=8$ repeated evaluations substantially improves alignment with human judgments, reducing the MAE from 1.11 to **0.72**.

Table 4: **Preference accuracy on RoboRewardBench.** LLM-as-a-Verifier outperforms trained robotics reward models.

Method	Accuracy (%)
TOPReward	74.7
Robometer-4B	78.8
RoboReward-8B	81.4
LLM-as-a-Judge (Discrete)	70.8
LLM-as-a-Verifier (Ours)	87.4

Table 5: **Evaluation on RoboRewardBench against human annotations.** We report Mean Absolute Error (MAE; lower is better) between human labels and predicted rewards. LLM-as-a-Verifier uses continuous rewards ($K=8$), whereas the baseline extracts discrete scores from the model.

Models	RoboRewardBench MAE ($\downarrow$)
RoboReward 8B	1.11
RoboReward 8B + LLM-as-a-Verifier	0.72

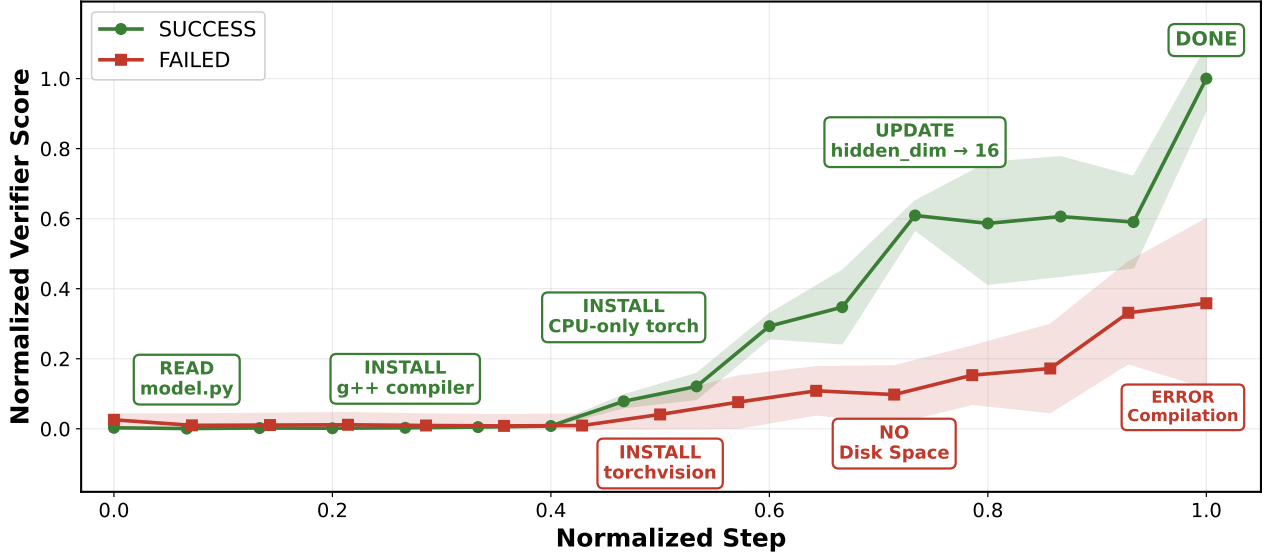


Figure 8: We observe a strong correlation between the chronological progression of code generation steps and the scores from LLM-as-a-Verifier. The example task above requires the agent to run MNIST inference. The successful trajectory follows a coherent sequence of events—*Read model.py* → *Install g++ compiler* → *Install CPU-only torch* → *Update hidden_dim* → *DONE* and exhibits consistently increasing verifier scores. In contrast, the failed trajectory is characterized by erroneous behaviors—it unnecessarily installs the large *torchvision* package, which exhausts the available disk space and hits a compilation error—resulting in significantly lower scores. Results are shown for the *pytorch-model-cli* task from Terminal-Bench V2, using Gemini 2.5 Pro with Terminus 2 and Gemini 2.5 Flash as the verifier.

5.4. MedAgentBench

MedAgentBench [19] evaluates LLM agents on medical tasks that involve patient information retrieval, guideline lookup, and multi-step tool use in a simulated electronic health records (EHR) environment. It covers a regime where ground-truth trajectory checkers are expensive to construct and where verification errors carry real safety consequences, making it a natural stress test for general-purpose verifiers. We use the AgentBench harness and sample $N=5$ trajectories per task from Claude Opus 4.8, then apply the same verification procedure. The Pass@1 of Claude Opus 4.8 on this pool is 70.2% and LLM-as-a-Verifier achieves **73.3%**, outperforming Opus 4.8 (70.2%), Gemini 3.5 Flash (66.3%), and GPT-5.5 (65.1%).

6. Fine-grained Verifier Signals as a Proxy for Task Progress

Beyond selecting the best trajectory, the fine-grained signal produced by LLM-as-a-Verifier can serve as a scalar proxy for how far an agent has progressed through a task. We quantify this using the *Value-Order Correlation* (VOC), the Spearman rank correlation between the chronological index of a step and the verifier’s predicted value for the prefix ending at that step, following Ma et al. [22]. Intuitively, a verifier that tracks task progress should assign monotonically higher scores to later prefixes of a successful rollout, yielding $\text{VOC} \rightarrow 1$, and should remain robust to failure modes such as getting stuck or regressing.

$$\text{VOC} = \text{rank-correlation}(\text{argsort}(s_{t_1}, s_{t_2}, \dots, s_{t_K}), (t_1, t_2, \dots, t_K)). \quad (6.1)$$

VOC on code generation. On Terminal-Bench V2, we measure VOC between the chronological step of each agent action and the verifier’s score on the corresponding trajectory prefix. LLM-as-a-Verifier produces consistently increasing scores on successful rollouts while remaining largely flat on trajectories that stall or drift toward failure, allowing the same scalar to serve as both a progress measure and an early-warning signal. Figure 8 illustrates this on the `pytorch-model-cli` task, where the successful run’s score rises monotonically while the failed run’s stays low. This dual use motivates our Claude Code and Codex extensions, which surface the live verifier score to the user so that long-running agentic jobs can be monitored, paused, or rolled back before they commit broken state to disk. Across 500 (success, failure) pairs drawn from Terminal-Bench V2 runs, the verifier (Gemini 2.5 Flash, $G=20$) attains Spearman VOC 0.848 on successful trajectories and 0.769 on failed ones (full numbers reported in Table 6). The code-generation VOC numbers indicate that the fine-grained signal produced by LLM-as-a-Verifier is not merely a better ranker but a calibrated estimator of task progress, opening a path toward safer real-world deployment of autonomous agents.

Trajectory outcome	Spearman VOC (rank correlation)
Successful	0.848 ± 0.012
Failed	0.769 ± 0.016
Success – Failed (gap)	+0.079

Table 6: **Value-Order Correlation by trajectory outcome on Terminal-Bench V2.** Mean Spearman rank correlation between step index and verifier progress score, computed over 500 randomly sampled trajectories from Terminal-Bench V2. The verifier (Gemini 2.5 Flash, $G=20$) exhibits near-monotonic progress for successful trajectories, while failed rollouts show weaker correlation, indicating limited or inconsistent progress. We observe a 0.08 Spearman gap between successful and failed trajectories generated by the same agent backbone on the same task.

VOC on robotics. We compute VOC over 500 trajectories from the held-out RoboReward dataset. As shown in Table 7, LLM-as-a-Verifier (Qwen 3.6, $K=5$, $G=20$) attains **0.966**, substantially exceeding RoboReward-8B (0.877), Robometer-4B (0.780), and TOPReward (0.565). Qualitatively, TOPReward tends to saturate at $P(\text{True})=1.0$ almost immediately and therefore loses the ability to discriminate mid-trajectory progress when a rollout eventually fails, whereas our expectation over the full scoring distribution preserves a smooth, chronologically-aligned signal throughout the episode.

Method	Spearman VOC (rank correlation)
LLM-as-a-Verifier (Qwen 3.6 35B, 5 reps, 20 granularity)	0.966
RoboReward-8B	0.877
Robometer-4B	0.780
TOPReward (Qwen 3.6, $P(\text{true})$)	0.565

Table 7: **Value-Order Correlation on 500 trajectories from RoboRewardBench.** LLM-as-a-Verifier with $K=5$ repeated evaluations and $G=20$ scoring granularity attains the highest rank-correlation between the chronological step index and the verifier’s predicted progress score.

Coding Agent Extension To demonstrate the applicability of LLM-as-a-Verifier to real-world coding agents, we develop *TurboAgent*, a drop-in extension for Claude Code and other OpenAI-API compatible clients. TurboAgent operates as an inference-time proxy that transparently sits between the client and the LLM provider, requiring no modifications to either the underlying agent harness or the backend model. The proxy design also allows TurboAgent to be plugged transparently into existing benchmarks such as Terminal-Bench [17]. For each request, it dispatches N candidate trajectories to the backend model in parallel and selects the best response using the proposed *Probabilistic Pivot Tournament* (PPT). Beyond verification, TurboAgent also provides a web-based interface for visualizing verifier outputs and monitoring agent progress in real time.

7. Dense Reward for Reinforcement Learning

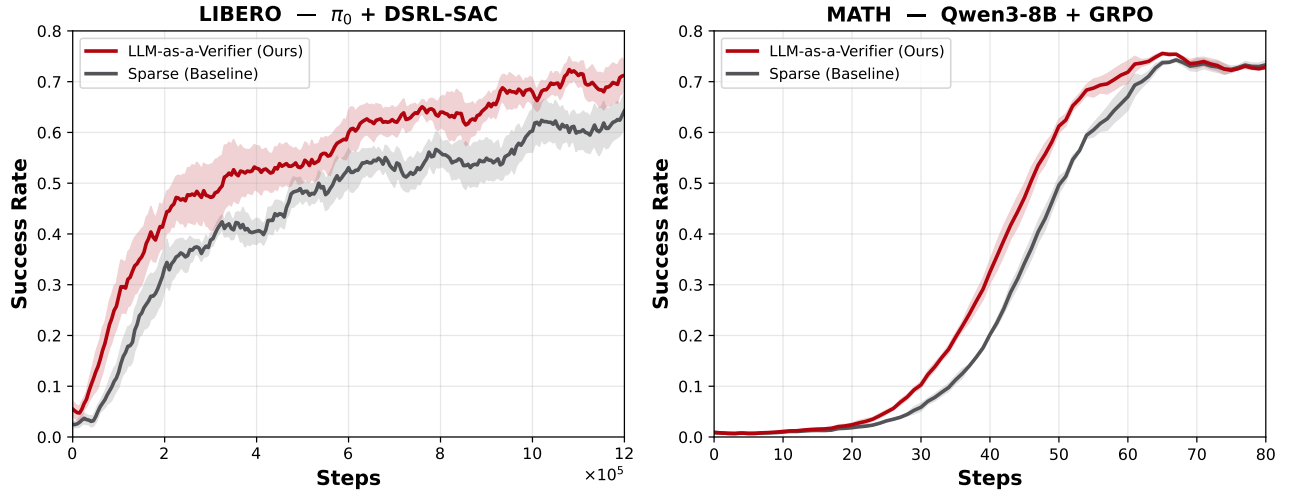


Figure 9: LLM-as-a-Verifier improves RL sample efficiency. Success rate versus training steps for off-policy (left) and on-policy (right) reinforcement learning, comparing sparse-reward baselines with dense rewards from LLM-as-a-Verifier. **Left:** A π_0 policy fine-tuned on the LIBERO ketchup task with DSRL-SAC. The verifier progress reward (Eq. 7.1) achieves the same success rate using $\approx 1.8\times$ fewer environment steps and reaches a higher final success rate (0.76 vs. 0.69). **Right:** Qwen3-8B fine-tuned on MATH with GRPO. The verifier reasoning reward (Eq. 7.2) improves sample efficiency by $\approx 1.1\times$. Results are averaged over multiple seeds (LIBERO $n=5$, MATH $n=3$).

The progress signal from the previous section also helps remediate a long-standing difficulty in reinforcement learning (RL): the *credit assignment problem*. We show that the fine-grained score of LLM-as-a-Verifier (Eq. 3.1) is a drop-in dense reward for both off-policy and on-policy RL, improving sample efficiency without any reward-model training or environment-specific shaping.

Off-policy RL: dense progress rewards for DSRL-SAC. We fine-tune the π_0 [23] vision–language–action model on LIBERO with DSRL using Soft Actor–Critic (SAC). At the end of each rollout we query the VLM verifier with the task instruction x and a uniformly sub-sampled sequence of rendered frames, obtaining

a per-step progress curve $\rho_t = R(x, \tau_{1:t}) \in [0, 1]$. We then relabel the rollout with the shaped reward:

$$r_t = r_t^{\text{env}} + \lambda \rho_t, \quad (7.1)$$

store the relabeled transitions (s_t, a_t, r_t, s_{t+1}) in the replay buffer $\mathcal{D}$, and train the SAC critic on the relabeled returns sampled from $\mathcal{D}$. The coefficient λ trades off environment and verifier rewards. Because shaping is applied offline to stored trajectories and leaves the SAC objective untouched, it adds dense intermediate signals at no additional algorithmic cost.

On-policy RL: dense reasoning rewards for GRPO. We fine-tune Qwen3-8B [24] on MATH using Group Relative Policy Optimization (GRPO) [14], which samples a group of G responses $\{y_i\}_{i=1}^G$ for each prompt x and estimates each response’s advantage relative to the group. During the early stages of training, it is common for all sampled responses to produce incorrect final answers, causing the group-relative advantage to collapse to zero and yielding no gradient. LLM-as-a-Verifier mitigates this issue by evaluating the reasoning trace of each completion using the probabilistic pivot tournament (Eq. 3.2), assigning each response a normalized preference score $\bar{R}_i \in [0, 1]$ that captures fine-grained differences in reasoning quality even when the final answers are identical. We incorporate this verifier-derived score into the standard correctness and format reward with weight β :

$$r_i = r_{\text{correct},i} + r_{\text{format},i} + \beta r_{\text{reasoning},i}. \quad (7.2)$$

Empirical findings. Across both regimes, dense verifier rewards improve sample efficiency over the sparse baselines (Fig. 9). We quantify sample efficiency as the ratio of training steps the sparse baseline requires to reach a target success rate to the steps our dense reward requires. On LIBERO, shaping a π_0 policy trained with DSRL-SAC reaches a matched success rate in substantially fewer environment steps— $1.8\times$ higher sample efficiency across success-rate targets from 0.2 to 0.6—while also attaining a higher final success rate (0.76 vs. 0.69). On MATH, augmenting GRPO with the reasoning reward gives a smaller but consistent gain of $\approx 1.1\times$ (a $\sim 10\%$ reduction in the optimizer steps needed to reach a matched accuracy). We report reward-shaping hyperparameters and additional ablations in Appendix B.7.

8. Discussion

In this work, we argue that verification constitutes an underexplored axis of scaling. To realize this, we propose LLM-as-a-Verifier, a general-purpose framework that delivers fine-grained feedback for agentic tasks. Unlike standard LM judges that output a single discrete score, our approach computes a continuous reward by taking the expectation over the distribution of scoring-token logits and enables verification scaling across multiple dimensions, including (1) score granularity, (2) repeated evaluation, and (3) criteria decomposition. When used as a trajectory reward model for test-time scaling, it achieves state-of-the-art performance on Terminal-Bench V2, SWE-Bench Verified, RoboRewardBench, and MedAgentBench. Beyond ranking, the fine-grained verifier signal can be used as a progress estimator, opening a path toward safer real-world deployment of autonomous agents. Finally, we demonstrate that LLM-as-a-Verifier can be used as a dense reward signal for RL, improving the sample efficiency of SAC and GRPO on robotics and mathematical reasoning benchmarks.

9. Related Work

Test-Time Scaling. Test-time scaling improves model performance by spending additional inference compute on deliberation, search, or candidate generation. One line of work improves a single response by eliciting intermediate reasoning with chain-of-thought [25, 26], decomposing problems into simpler subproblems [27], or marginalizing over sampled reasoning paths [28]. Another line searches over intermediate thoughts [29, 30], actions [31, 32], or latent world states [33] based on test-time feedback [34–39]. Repeated sampling and best-of- N selection further scale candidate pools for code generation [40], general reasoning [3, 41], parallel self-verification [5], and inference-aware training [42]. These approaches can expose substantial oracle headroom, but realizing this headroom requires a reliable selector. Our LLM-as-a-Verifier instead shows that verifier quality can be improved by scaling score granularity, repeated evaluation, and criteria decomposition, and that better verification directly improves best-of- N selection for long-horizon decision making.

LLM-as-a-Judge. LLM-as-a-judge methods provide a scalable alternative to human evaluation by prompting large models to score or compare generated outputs. Probability- and form-filling-based evaluators extract richer scoring signals from LLMs [43, 44], while benchmark-style evaluators use LLM preferences to evaluate instruction-following systems [45–47]. A complementary line makes evaluation more fine-grained through skill decompositions [48], customized rubrics and specialized open judges [49–52], and hierarchical criteria [53]. Other work trains scalable judge models [54, 55] or combines multiple judges through debate and weak-verifier ensembling [56, 57]. Recent studies have also analyzed judge reliability, including fairness and position biases [58, 59], cognitive and self-enhancement biases [60, 61], general judge benchmarks [62, 63], as well as domain-specific judge evaluation [64]. Multimodal judge models extend this paradigm to image and vision-language evaluation [65–67]. Our work builds on these works but differs in setting, objective, and scaling characterization. Rather than evaluating isolated natural-language responses, LLM-as-a-Verifier verifies long-horizon agent trajectories involving tool use, code execution, robotics, and medical decision-making. Moreover, we systematically study how verification quality scales with score granularity, repeated evaluations, and criteria decomposition.

Verifiable Reward. Reward models convert candidate solutions, actions, or trajectories into scalar feedback for selection, monitoring, or policy optimization. In language reasoning, learned verifiers have been trained as outcome reward models for final-answer selection [7], process reward models for step-level supervision [8, 68], and generative verifiers that cast reward modeling as next-token prediction [6]. In robotics, reward signals have been derived from value-implicit visual representations [69], language-image reward representations [70], pretrained vision-language models [22, 71, 72], VLM feedback or preferences [73], LLM-generated reward code [74–76], token-probability progress [10], and trained general-purpose robotic reward models based on large-scale trajectory or preference data [9, 11, 77, 78]. Recent work has also developed action-level verifiers for guided sampling [79–81], runtime monitoring [82], and multimodal alignment [83]. A complementary line makes verification more reliable by factorizing holistic judgments into smaller checks [53, 84–88]. Orthogonal to these methods, our work studies how verifier quality scales with score granularity, repeated evaluation, and criteria decomposition across multiple domains in a general-purpose framework.

10. Acknowledgments

We thank the members of the UC Berkeley Sky Computing Lab, Stanford Scaling Intelligence Lab, IRIS Lab, and Autonomous Systems Lab for their constructive feedback and informative discussions. This work was supported by Google; Google DeepMind; Google Cloud; Stanford HAI; DARPA (HR00112520038, Fallingwater); NSF (24-554, AIMing); NASA ULI; Schmidt Sciences; and Lightspeed. We also acknowledge the support of IBM and Felicis as members of Stanford HAI’s Industry Affiliates Program.

References

- [1] Jared Kaplan, Sam McCandlish, Tom Henighan, Tom B. Brown, Benjamin Chess, Rewon Child, Scott Gray, Alec Radford, Jeffrey Wu, and Dario Amodei. Scaling Laws for Neural Language Models, January 2020. URL <http://arxiv.org/abs/2001.08361>. arXiv:2001.08361 [cs].
- [2] Leo Gao, John Schulman, and Jacob Hilton. Scaling Laws for Reward Model Overoptimization, October 2022. URL <http://arxiv.org/abs/2210.10760>. arXiv:2210.10760 [cs].
- [3] Charlie Snell, Jaehoon Lee, Kelvin Xu, and Aviral Kumar. Scaling LLM Test-Time Compute Optimally can be More Effective than Scaling Model Parameters, August 2024. URL <http://arxiv.org/abs/2408.03314>. arXiv:2408.03314 [cs].
- [4] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric P. Xing, Hao Zhang, Joseph E. Gonzalez, and Ion Stoica. Judging LLM-as-a-judge with MT-Bench and Chatbot Arena, 2023. URL <https://arxiv.org/abs/2306.05685>.
- [5] Harman Singh, Xiuyu Li, Kusha Sareen, Monishwaran Maheswaran, Sijun Tan, Xiaoxia Wu, Junxiong Wang, Alpay Ariyak, Qingyang Wu, Samir Khaki, Rishabh Tiwari, Long Lian, Yucheng Lu, Boyi Li, Alane Suhr, Ben Athiwaratkun, and Kurt Keutzer. V_1 : Unifying Generation and Self-Verification for Parallel Reasoners, 2026. URL <https://arxiv.org/abs/2603.04304>.
- [6] Lunjun Zhang, Arian Hosseini, Hritik Bansal, Mehran Kazemi, Aviral Kumar, and Rishabh Agarwal. Generative verifiers: Reward modeling as next-token prediction, 2025. URL <https://arxiv.org/abs/2408.15240>.
- [7] 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, November 2021. URL <http://arxiv.org/abs/2110.14168>. arXiv:2110.14168 [cs].
- [8] Hunter Lightman, Vineet Kosaraju, Yura Burda, Harri Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. Let’s Verify Step by Step, May 2023. URL <http://arxiv.org/abs/2305.20050>. arXiv:2305.20050 [cs].
- [9] Anthony Liang, Yigit Korkmaz, Jiahui Zhang, Minyoung Hwang, Abrar Anwar, Sidhant Kaushik, Aditya Shah, Alex S Huang, Luke Zettlemoyer, Dieter Fox, et al. Robometer: Scaling general-purpose robotic reward models via trajectory comparisons. *arXiv preprint arXiv:2603.02115*, 2026.

- [10] Shirui Chen, Cole Harrison, Ying-Chun Lee, Angela Jin Yang, Zhongzheng Ren, Lillian J Ratliff, Jiafei Duan, Dieter Fox, and Ranjay Krishna. Topreward: Token probabilities as hidden zero-shot rewards for robotics. *arXiv preprint arXiv:2602.19313*, 2026.
- [11] Tony Lee, Andrew Wagenmaker, Karl Pertsch, Percy Liang, Sergey Levine, and Chelsea Finn. RoboReward: General-purpose vision-language reward models for robotics, 2026. URL <https://arxiv.org/abs/2601.00675>.
- [12] Bo Liu, Yifeng Zhu, Chongkai Gao, Yihao Feng, Qiang Liu, Yuke Zhu, and Peter Stone. Libero: Benchmarking knowledge transfer for lifelong robot learning, 2023. URL <https://arxiv.org/abs/2306.03310>.
- [13] Andrew Wagenmaker, Mitsuhiko Nakamoto, Yunchu Zhang, Seohong Park, Waleed Yagoub, Anusha Nagabandi, Abhishek Gupta, and Sergey Levine. Steering your diffusion policy with latent space reinforcement learning, 2025. URL <https://arxiv.org/abs/2506.15799>.
- [14] DeepSeek-AI, Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Peiyi Wang, Qihao Zhu, Runxin Xu, Ruoyu Zhang, Shirong Ma, Xiao Bi, Xiaokang Zhang, Xingkai Yu, Yu Wu, Z. F. Wu, Zhibin Gou, Zhihong Shao, Zhuoshu Li, Ziyi Gao, Aixin Liu, Bing Xue, Bingxuan Wang, Bochao Wu, Bei Feng, Chengda Lu, Chenggang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, Damai Dai, Deli Chen, Dongjie Ji, Erhang Li, Fangyun Lin, Fucong Dai, Fuli Luo, Guangbo Hao, Guanting Chen, Guowei Li, H. Zhang, Han Bao, Hanwei Xu, Haocheng Wang, Honghui Ding, Huajian Xin, Huazuo Gao, Hui Qu, Hui Li, Jianzhong Guo, Jiashi Li, Jiawei Wang, Jingchang Chen, Jingyang Yuan, Junjie Qiu, Junlong Li, J. L. Cai, Jiaqi Ni, Jian Liang, Jin Chen, Kai Dong, Kai Hu, Kaige Gao, Kang Guan, Kexin Huang, Kuai Yu, Lean Wang, Lecong Zhang, Liang Zhao, Litong Wang, Liyue Zhang, Lei Xu, Leyi Xia, Mingchuan Zhang, Minghua Zhang, Minghui Tang, Meng Li, Miaoju Wang, Mingming Li, Ning Tian, Panpan Huang, Peng Zhang, Qiancheng Wang, Qinyu Chen, Qiushi Du, Ruiqi Ge, Ruisong Zhang, Ruizhe Pan, Runji Wang, R. J. Chen, R. L. Jin, Ruyi Chen, Shanghao Lu, Shangyan Zhou, Shanhuang Chen, Shengfeng Ye, Shiyu Wang, Shuiping Yu, Shunfeng Zhou, Shuting Pan, S. S. Li, Shuang Zhou, Shaoqing Wu, Shengfeng Ye, Tao Yun, Tian Pei, Tianyu Sun, T. Wang, Wangding Zeng, Wanbiao Zhao, Wen Liu, Wenfeng Liang, Wenjun Gao, Wenqin Yu, Wentao Zhang, W. L. Xiao, Wei An, Xiaodong Liu, Xiaohan Wang, Xiaokang Chen, Xiaotao Nie, Xin Cheng, Xin Liu, Xin Xie, Xingchao Liu, Xinyu Yang, Xinyuan Li, Xuecheng Su, Xuheng Lin, X. Q. Li, Xiangyue Jin, Xiaojin Shen, Xiaosha Chen, Xiaowen Sun, Xiaoxiang Wang, Xinnan Song, Xinyi Zhou, Xianzu Wang, Xinxia Shan, Y. K. Li, Y. Q. Wang, Y. X. Wei, Yang Zhang, Yanhong Xu, Yao Li, Yao Zhao, Yaofeng Sun, Yaohui Wang, Yi Yu, Yichao Zhang, Yifan Shi, Yiliang Xiong, Ying He, Yishi Piao, Yisong Wang, Yixuan Tan, Yiyang Ma, Yiyuan Liu, Yongqiang Guo, Yuan Ou, Yuduan Wang, Yue Gong, Yuheng Zou, Yujia He, Yunfan Xiong, Yuxiang Luo, Yuxiang You, Yuxuan Liu, Yuyang Zhou, Y. X. Zhu, Yanhong Xu, Yanping Huang, Yaohui Li, Yi Zheng, Yuchen Zhu, Yunxian Ma, Ying Tang, Yukun Zha, Yuting Yan, Z. Z. Ren, Zehui Ren, Zhangli Sha, Zhe Fu, Zhean Xu, Zhenda Xie, Zhengyan Zhang, Zhewen Hao, Zhicheng Ma, Zhigang Yan, Zhiyu Wu, Zihui Gu, Zijia Zhu, Zijun Liu, Zilin Li, Ziwei Xie, Ziyang Song, Zizheng Pan, Zhen Huang, Zhipeng Xu, Zhongyu Zhang, and Zhen Zhang. DeepSeek-r1: Incentivizing reasoning capability in LLMs via reinforcement learning. 645(8081):633–638. ISSN 0028-0836, 1476-4687. doi: 10.1038/s41586-025-09422-z. URL <http://arxiv.org/abs/2501.12948>.
- [15] Nisan Stiennon, Long Ouyang, Jeff Wu, Daniel M. Ziegler, Ryan Lowe, Chelsea Voss, Alec Radford,

- Dario Amodei, and Paul Christiano. Learning to summarize from human feedback, February 2022. URL <http://arxiv.org/abs/2009.01325>. arXiv:2009.01325 [cs].
- [16] Google. Gemini 2.5 Flash. <https://ai.google.dev/gemini-api/docs/models/gemini-2.5-flash>, 2026. Accessed: 2026-05-06.
- [17] Mike A. Merrill, Alexander G. Shaw, Nicholas Carlini, Boxuan Li, Harsh Raj, Ivan Bercovich, Lin Shi, Jeong Yeon Shin, Thomas Walshe, E. Kelly Buchanan, Junhong Shen, Guanghao Ye, Haowei Lin, Jason Poulos, Maoyu Wang, Marianna Nezhurina, Jenia Jitsev, Di Lu, Orfeas Menis Mastromichalakis, Zhiwei Xu, Zizhao Chen, Yue Liu, Robert Zhang, Leon Liangyu Chen, Anurag Kashyap, Jan-Lucas Uslu, Jeffrey Li, Jianbo Wu, Minghao Yan, Song Bian, Vedang Sharma, Ke Sun, Steven Dillmann, Akshay Anand, Andrew Lanpouthakoun, Bardia Koopah, Changran Hu, Etash Guha, Gabriel H. S. Dreiman, Jiacheng Zhu, Karl Krauth, Li Zhong, Niklas Muennighoff, Robert Amanfu, Shangyin Tan, Shreyas Pimpalgaonkar, Tushar Aggarwal, Xiangning Lin, Xin Lan, Xuandong Zhao, Yiqing Liang, Yuanli Wang, Zilong Wang, Changzhi Zhou, David Heineman, Hange Liu, Harsh Trivedi, John Yang, Junhong Lin, Manish Shetty, Michael Yang, Nabil Omi, Negin Raoof, Shanda Li, Terry Yue Zhuo, Wuwei Lin, Yiwei Dai, Yuxin Wang, Wenhao Chai, Shang Zhou, Dariush Wahdany, Ziyu She, Jiaming Hu, Zhikang Dong, Yuxuan Zhu, Sasha Cui, Ahson Saiyed, Arinbjörn Kolbeinsson, Jesse Hu, Christopher Michael Rytting, Ryan Marten, Yixin Wang, Alex Dimakis, Andy Konwinski, and Ludwig Schmidt. Terminal-Bench: Benchmarking Agents on Hard, Realistic Tasks in Command Line Interfaces, January 2026. URL <http://arxiv.org/abs/2601.11868>. arXiv:2601.11868 [cs].
- [18] Carlos E. Jimenez, John Yang, Alexander Wettig, Shunyu Yao, Kexin Pei, Ofir Press, and Karthik Narasimhan. SWE-bench: Can Language Models Resolve Real-World GitHub Issues?, November 2024. URL <http://arxiv.org/abs/2310.06770>. arXiv:2310.06770 [cs].
- [19] Yixing Jiang, Kameron C. Black, Gloria Geng, Danny Park, James Zou, Andrew Y. Ng, and Jonathan H. Chen. A virtual ehr environment to benchmark medical llm agents. *NEJM AI*, 2(9), 2025. doi: 10.1056/AIdbp2500144.
- [20] Capy. URL <https://www.tbench.ai/leaderboard/terminal-bench/2.0/capy-build/unknown/gpt-5.5%40openai>.
- [21] Terminus 2. URL https://github.com/harbor-framework/terminal-bench/tree/main/terminal_bench/agents/terminus_2.
- [22] Yecheng Jason Ma, Joey Hejna, Ayzaan Wahid, Chuyuan Fu, Dhruv Shah, Jacky Liang, Zhuo Xu, Sean Kirmani, Peng Xu, Danny Driess, Ted Xiao, Jonathan Tompson, Osbert Bastani, Dinesh Jayaraman, Wenhao Yu, Tingnan Zhang, Dorsa Sadigh, and Fei Xia. Vision language models are in-context value learners. In *International Conference on Learning Representations*, 2025. doi: 10.48550/arXiv.2411.04549. URL <https://arxiv.org/abs/2411.04549>.
- [23] Kevin Black, Noah Brown, Danny Driess, Adnan Esmail, Michael Equi, Chelsea Finn, Niccolo Fusai, Lachy Groom, Karol Hausman, Brian Ichter, Szymon Jakubczak, Tim Jones, Liyiming Ke, Sergey Levine, Adrian Li-Bell, Mohith Mothukuri, Suraj Nair, Karl Pertsch, Lucy Xiaoyang Shi, James Tanner, Quan Vuong, Anna Walling, Haohuan Wang, and Ury Zhilinsky. π_0 : A vision-language-action flow model for general robot control, 2026. URL <https://arxiv.org/abs/2410.24164>.
- [24] An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, et al. Qwen3 technical report. *arXiv preprint arXiv:2505.09388*, 2025.

- [25] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed H. Chi, Quoc V. Le, and Denny Zhou. Chain-of-thought prompting elicits reasoning in large language models, 2022. URL <https://arxiv.org/abs/2201.11903>.
- [26] Takeshi Kojima, Shixiang Shane Gu, Machel Reid, Yutaka Matsuo, and Yusuke Iwasawa. Large language models are zero-shot reasoners, 2022. URL <https://arxiv.org/abs/2205.11916>.
- [27] Denny Zhou, Nathanael Schärli, Le Hou, Jason Wei, Nathan Scales, Xuezhi Wang, Dale Schuurmans, Claire Cui, Olivier Bousquet, Quoc V. Le, and Ed H. Chi. Least-to-most prompting enables complex reasoning in large language models. In *International Conference on Learning Representations*, 2023. URL <https://openreview.net/forum?id=WZH7099tgfM>.
- [28] Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc Le, Ed Chi, Sharan Narang, Aakanksha Chowdhery, and Denny Zhou. Self-consistency improves chain of thought reasoning in language models. URL <http://arxiv.org/abs/2203.11171>.
- [29] Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran, Tom Griffiths, Yuan Cao, and Karthik Narasimhan. Tree of thoughts: Deliberate problem solving with large language models. *Advances in neural information processing systems*, 36:11809–11822, 2023.
- [30] Maciej Besta, Nils Blach, Ales Kubicek, Robert Gerstenberger, Michal Podstawski, Lukas Gianinazzi, Joanna Gajda, Tomasz Lehmann, Hubert Niewiadomski, Piotr Nyczyk, and Torsten Hoeffler. Graph of thoughts: Solving elaborate problems with large language models, 2023. URL <https://arxiv.org/abs/2308.09687>.
- [31] Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan Cao. ReAct: Synergizing reasoning and acting in language models. URL <http://arxiv.org/abs/2210.03629>.
- [32] Andy Zhou, Kai Yan, Michal Shlapentokh-Rothman, Haohan Wang, and Yu-Xiong Wang. Language agent tree search unifies reasoning acting and planning in language models. *arXiv preprint arXiv:2310.04406*, 2023.
- [33] Shibo Hao, Yi Gu, Haodi Ma, Joshua Jiahua Hong, Zhen Wang, Daisy Zhe Wang, and Zhiting Hu. Reasoning with language model is planning with world model. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, 2023. URL <https://arxiv.org/abs/2305.14992>.
- [34] Yixuan Weng, Minjun Zhu, Fei Xia, Bin Li, Shizhu He, Shengping Liu, Bin Sun, Kang Liu, and Jun Zhao. Large language models are better reasoners with self-verification, 2022. URL <https://arxiv.org/abs/2212.09561>.
- [35] Aman Madaan, Niket Tandon, Prakhar Gupta, Skyler Hallinan, Luyu Gao, Sarah Wiegrefe, Uri Alon, Nouha Dziri, Shrimai Prabhumoye, Yiming Yang, Shashank Gupta, Bodhisattwa Prasad Majumder, Katherine Hermann, Sean Welleck, Amir Yazdanbakhsh, and Peter Clark. Self-refine: Iterative refinement with self-feedback, 2023. URL <https://arxiv.org/abs/2303.17651>.
- [36] Noah Shinn, Federico Cassano, Edward Berman, Ashwin Gopinath, Karthik Narasimhan, and Shunyu Yao. Reflexion: Language agents with verbal reinforcement learning, 2023. URL <https://arxiv.org/abs/2303.11366>.

- [37] Zhibin Gou, Zhihong Shao, Yeyun Gong, Yelong Shen, Yujiu Yang, Nan Duan, and Weizhu Chen. CRITIC: Large language models can self-correct with tool-interactive critiquing. In *International Conference on Learning Representations*, 2024. URL <https://arxiv.org/abs/2305.11738>.
- [38] Lakshya A. Agrawal, Shangyin Tan, Dilara Soylu, Noah Ziemis, Rishi Khare, Krista Opsahl-Ong, Arnav Singhvi, Herumb Shandilya, Michael J. Ryan, Meng Jiang, Christopher Potts, Koushik Sen, Alexandros G. Dimakis, Ion Stoica, Dan Klein, Matei Zaharia, and Omar Khattab. GEPA: Reflective Prompt Evolution Can Outperform Reinforcement Learning, July 2025. URL <http://arxiv.org/abs/2507.19457>. arXiv:2507.19457 [cs].
- [39] Alexander Novikov, Ngân Vũ, Marvin Eisenberger, Emilien Dupont, Po-Sen Huang, Adam Zsolt Wagner, Sergey Shirobokov, Borislav Kozlovskii, Francisco J. R. Ruiz, Abbas Mehrabian, M. Pawan Kumar, Abigail See, Swarat Chaudhuri, George Holland, Alex Davies, Sebastian Nowozin, Pushmeet Kohli, and Matej Balog. AlphaEvolve: A coding agent for scientific and algorithmic discovery, June 2025. URL <http://arxiv.org/abs/2506.13131>. arXiv:2506.13131 [cs].
- [40] Yujia Li, David Choi, Junyoung Chung, Nate Kushman, Julian Schrittwieser, Rémi Leblond, Tom Eccles, James Keeling, Felix Gimeno, Agustin Dal Lago, Thomas Hubert, Peter Choy, Cyprien de Masson d’Autume, Igor Babuschkin, Xinyun Chen, Po-Sen Huang, Johannes Welbl, Sven Gowal, Alexey Cherepanov, James Molloy, Daniel J. Mankowitz, Esme Sutherland Robson, Pushmeet Kohli, Nando de Freitas, Koray Kavukcuoglu, and Oriol Vinyals. Competition-level code generation with AlphaCode. *Science*, 378(6624):1092–1097, 2022. doi: 10.1126/science.abq1158. URL <https://www.science.org/doi/10.1126/science.abq1158>.
- [41] Bradley Brown, Jordan Juravsky, Ryan Ehrlich, Ronald Clark, Quoc V. Le, Christopher Ré, and Azalia Mirhoseini. Large language monkeys: Scaling inference compute with repeated sampling, 2024. URL <https://arxiv.org/abs/2407.21787>.
- [42] Yinlam Chow, Guy Tennenholtz, Izzeddin Gur, Vincent Zhuang, Bo Dai, Sridhar Thiagarajan, Craig Boutilier, Rishabh Agarwal, Aviral Kumar, and Aleksandra Faust. Inference-aware fine-tuning for best-of-N sampling in large language models, 2024. URL <https://arxiv.org/abs/2412.15287>.
- [43] Jinlan Fu, See-Kiong Ng, Zhengbao Jiang, and Pengfei Liu. GPTScore: Evaluate as you desire, 2023. URL <https://arxiv.org/abs/2302.04166>.
- [44] Yang Liu, Dan Iter, Yichong Xu, Shuohang Wang, Ruochen Xu, and Chenguang Zhu. G-eval: NLG evaluation using GPT-4 with better human alignment. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 2511–2522, Singapore, December 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.emnlp-main.153. URL <https://aclanthology.org/2023.emnlp-main.153/>.
- [45] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric P. Xing, Hao Zhang, Joseph E. Gonzalez, and Ion Stoica. Judging LLM-as-a-judge with MT-bench and chatbot arena. In *Advances in Neural Information Processing Systems*, 2023. URL <https://openreview.net/forum?id=uccHPGDlao>.
- [46] Yann Dubois, Balázs Galambosi, Percy Liang, and Tatsunori B. Hashimoto. Length-controlled AlpacaEval: A simple way to debias automatic evaluators, 2024. URL <https://arxiv.org/abs/2404.04475>.

- [47] Tianle Li, Wei-Lin Chiang, Evan Frick, Lisa Dunlap, Tianhao Wu, Banghua Zhu, Joseph E. Gonzalez, and Ion Stoica. From crowdsourced data to high-quality benchmarks: Arena-Hard and BenchBuilder pipeline, 2024. URL <https://arxiv.org/abs/2406.11939>.
- [48] Seonghyeon Ye, Doyoung Kim, Sungdong Kim, Hyeonbin Hwang, Seungone Kim, Yongrae Jo, James Thorne, Juho Kim, and Minjoon Seo. Flask: Fine-grained language model evaluation based on alignment skill sets. *arXiv preprint arXiv:2307.10928*, 2023.
- [49] Seungone Kim, Jamin Shin, Yejin Cho, Joel Jang, Shayne Longpre, Hwaran Lee, Sangdoo Yun, Seongjin Shin, Sungdong Kim, James Thorne, and Minjoon Seo. Prometheus: Inducing fine-grained evaluation capability in language models. In *International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=8euJaTveKw>.
- [50] Seungone Kim, Juyoung Suk, Shayne Longpre, Bill Yuchen Lin, Jamin Shin, Sean Welleck, Graham Neubig, Moontae Lee, Kyungjae Lee, and Minjoon Seo. Prometheus 2: An open source language model specialized in evaluating other language models, 2024. URL <https://arxiv.org/abs/2405.01535>.
- [51] Junlong Li, Shichao Sun, Weizhe Yuan, Run-Ze Fan, Hai Zhao, and Pengfei Liu. Generative judge for evaluating alignment. In *International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=gtkFw6sZGS>.
- [52] Xinyu Hu, Li Lin, Mingqi Gao, Xunjian Yin, and Xiaojun Wan. Themis: A reference-free NLG evaluation language model with flexibility and interpretability, 2024. URL <https://arxiv.org/abs/2406.18365>.
- [53] Yuxuan Liu, Tianchi Yang, Shaohan Huang, Zihan Zhang, Haizhen Huang, Furu Wei, Weiwei Deng, Feng Sun, and Qi Zhang. HD-eval: Aligning large language model evaluators through hierarchical criteria decomposition. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 7641–7660, Bangkok, Thailand, August 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.acl-long.413. URL <https://aclanthology.org/2024.acl-long.413/>.
- [54] Yidong Wang, Zhuohao Yu, Wenjin Yao, Zhengran Zeng, Linyi Yang, Cunxiang Wang, Hao Chen, Chaoya Jiang, Rui Xie, Jindong Wang, Xing Xie, Wei Ye, Shikun Zhang, and Yue Zhang. PandaLM: An automatic evaluation benchmark for LLM instruction tuning optimization. In *International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=5Nn2BLV7SB>.
- [55] Lianghui Zhu, Xinggang Wang, and Xinlong Wang. JudgeLM: Fine-tuned large language models are scalable judges. In *International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=xsELpEPn4A>.
- [56] Chi-Min Chan, Weize Chen, Yusheng Su, Jianxuan Yu, Wei Xue, Shanghang Zhang, Jie Fu, and Zhiyuan Liu. Chateval: Towards better llm-based evaluators through multi-agent debate. *arXiv preprint arXiv:2308.07201*, 2023.
- [57] Jon Saad-Falcon, E Kelly Buchanan, Mayee F Chen, Tzu-Heng Huang, Brendan McLaughlin, Tanvir Bhathal, Shang Zhu, Ben Athiwaratkun, Frederic Sala, Scott Linderman, et al. Shrinking the generation-verification gap with weak verifiers. *arXiv preprint arXiv:2506.18203*, 2025.

- [58] Peiyi Wang, Lei Li, Liang Chen, Zefan Cai, Dawei Zhu, Binghuai Lin, Yunbo Cao, Lingpeng Kong, Qi Liu, Tianyu Liu, and Zhifang Sui. Large language models are not fair evaluators. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 9440–9450, Bangkok, Thailand, August 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.acl-long.511. URL <https://aclanthology.org/2024.acl-long.511/>.
- [59] Zhiyuan Zeng, Jiatong Yu, Tianyu Gao, Yu Meng, Tanya Goyal, and Danqi Chen. Evaluating large language models at evaluating instruction following. In *International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=tr0KidwPLc>.
- [60] Ryan Koo, Minhwa Lee, Vipul Raheja, Jong Inn Park, Zae Myung Kim, and Dongyeop Kang. Benchmarking cognitive biases in large language models as evaluators, 2023. URL <https://arxiv.org/abs/2309.17012>.
- [61] Yiqi Liu, Nafise Sadat Moosavi, and Chenghua Lin. LLMs as narcissistic evaluators: When ego inflates evaluation scores, 2023. URL <https://arxiv.org/abs/2311.09766>.
- [62] Sijun Tan, Siyuan Zhuang, Kyle Montgomery, William Yuan Tang, Alejandro Cuadron, Chenguang Wang, Raluca Ada Popa, and Ion Stoica. JudgeBench: A benchmark for evaluating LLM-based judges. In *International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=G0dksFayVq>.
- [63] Hui Huang, Xingyuan Bu, Hongli Zhou, Yingqi Qu, Jing Liu, Muyun Yang, Bing Xu, and Tiejun Zhao. An empirical study of LLM-as-a-judge for LLM evaluation: Fine-tuned judge model is not a general substitute for GPT-4. In *Findings of the Association for Computational Linguistics: ACL 2025*, pages 5880–5895, Vienna, Austria, July 2025. Association for Computational Linguistics. doi: 10.18653/v1/2025.findings-acl.306. URL <https://aclanthology.org/2025.findings-acl.306/>.
- [64] Hongchao Jiang, Yiming Chen, Yushi Cao, Hung-yi Lee, and Robby T. Tan. CodeJudgeBench: Benchmarking LLM-as-a-judge for coding tasks, 2025. URL <https://arxiv.org/abs/2507.10535>.
- [65] Dongping Chen, Ruoxi Chen, Shilin Zhang, Yinuo Liu, Yaochen Wang, Huichi Zhou, Qihui Zhang, Yao Wan, Pan Zhou, and Lichao Sun. MLLM-as-a-judge: Assessing multimodal LLM-as-a-judge with vision-language benchmark, 2024. URL <https://arxiv.org/abs/2402.04788>.
- [66] Seongyun Lee, Seungone Kim, Sue Hyun Park, Geewook Kim, and Minjoon Seo. Prometheus-vision: Vision-language model as a judge for fine-grained evaluation, 2024. URL <https://arxiv.org/abs/2401.06591>.
- [67] Tianyi Xiong, Xiyao Wang, Dong Guo, Qinghao Ye, Haoqi Fan, Quanquan Gu, Heng Huang, and Chunyuan Li. LLaVA-critic: Learning to evaluate multimodal models, 2024. URL <https://arxiv.org/abs/2410.02712>.
- [68] Jonathan Uesato, Nate Kushman, Ramana Kumar, Francis Song, Noah Siegel, Lisa Wang, Antonia Creswell, Geoffrey Irving, and Irina Higgins. Solving math word problems with process- and outcome-based feedback, 2022. URL <https://arxiv.org/abs/2211.14275>.
- [69] Yecheng Jason Ma, Shagun Sodhani, Dinesh Jayaraman, Osbert Bastani, Vikash Kumar, and Amy Zhang. VIP: Towards universal visual reward and representation via value-implicit pre-training, 2022. URL <https://arxiv.org/abs/2210.00030>.

- [70] Yecheng Jason Ma, William Liang, Vaidehi Som, Vikash Kumar, Amy Zhang, Osbert Bastani, and Dinesh Jayaraman. LIV: Language-image representations and rewards for robotic control, 2023. URL <https://arxiv.org/abs/2306.00958>.
- [71] Sumedh A. Sontakke, Jesse Zhang, Sébastien M. R. Arnold, Karl Pertsch, Erdem Biyik, Dorsa Sadigh, Chelsea Finn, and Laurent Itti. RoboCLIP: One demonstration is enough to learn robot policies, 2023. URL <https://arxiv.org/abs/2310.07899>.
- [72] Juan Rocamonde, Victoriano Montesinos, Elvis Nava, Ethan Perez, and David Lindner. Vision-language models are zero-shot reward models for reinforcement learning, 2023. URL <https://arxiv.org/abs/2310.12921>.
- [73] Yufei Wang, Zhanyi Sun, Jesse Zhang, Zhou Xian, Erdem Biyik, David Held, and Zackory Erickson. RL-VLM-F: Reinforcement learning from vision language foundation model feedback, 2024. URL <https://arxiv.org/abs/2402.03681>.
- [74] Wenhao Yu, Nimrod Gileadi, Chuyuan Fu, Sean Kirmani, Kuang-Huei Lee, Montse Gonzalez Arenas, Hao-Tien Lewis Chiang, Tom Erez, Leonard Hasenclever, Jan Humplik, Brian Ichter, Ted Xiao, Peng Xu, Andy Zeng, Tingnan Zhang, Nicolas Heess, Dorsa Sadigh, Jie Tan, Yuval Tassa, and Fei Xia. Language to rewards for robotic skill synthesis, 2023. URL <https://arxiv.org/abs/2306.08647>.
- [75] Tianbao Xie, Siheng Zhao, Chen Henry Wu, Yitao Liu, Qian Luo, Victor Zhong, Yanchao Yang, and Tao Yu. Text2Reward: Reward shaping with language models for reinforcement learning, 2023. URL <https://arxiv.org/abs/2309.11489>.
- [76] Yecheng Jason Ma, William Liang, Guanzhi Wang, De-An Huang, Osbert Bastani, Dinesh Jayaraman, Yuke Zhu, Linxi Fan, and Anima Anandkumar. Eureka: Human-level reward design via coding large language models, 2023. URL <https://arxiv.org/abs/2310.12931>.
- [77] Jiahui Zhang, Yusen Luo, Abrar Anwar, Sumedh Anand Sontakke, Joseph J. Lim, Jesse Thomason, Erdem Biyik, and Jesse Zhang. ReWiND: Language-guided rewards teach robot policies without new demonstrations, 2025. URL <https://arxiv.org/abs/2505.10911>.
- [78] Qianzhong Chen, Justin Yu, Mac Schwager, Pieter Abbeel, Yide Shentu, and Philipp Wu. SARM: Stage-aware reward modeling for long horizon robot manipulation, 2025. URL <https://arxiv.org/abs/2509.25358>.
- [79] Mitsuhiro Nakamoto, Oier Mees, Aviral Kumar, and Sergey Levine. Steering your generalists: Improving robotic foundation models via value guidance, 2024. URL <https://arxiv.org/abs/2410.13816>.
- [80] Yuejiang Liu, Jubayer Ibn Hamid, Annie Xie, Yoonho Lee, Maximilian Du, and Chelsea Finn. Bidirectional decoding: Improving action chunking via guided test-time sampling, 2024. URL <https://arxiv.org/abs/2408.17355>.
- [81] Jacky Kwok, Christopher Agia, Rohan Sinha, Matt Foutter, Shulu Li, Ion Stoica, Azalia Mirhoseini, and Marco Pavone. RoboMonkey: Scaling test-time sampling and verification for vision-language-action models. In Joseph Lim, Shuran Song, and Hae-Won Park, editors, *Proceedings of The 9th Conference on Robot Learning*, volume 305 of *Proceedings of Machine Learning Research*, pages 3200–3217. PMLR, 2025. URL <https://proceedings.mlr.press/v305/kwok25a.html>.

- [82] Christopher Agia, Rohan Sinha, Jingyun Yang, Zi-ang Cao, Rika Antonova, Marco Pavone, and Jeannette Bohg. Unpacking failure modes of generative policies: Runtime monitoring of consistency and progress, 2024. URL <https://arxiv.org/abs/2410.04640>.
- [83] Jacky Kwok, Xilun Zhang, Mengdi Xu, Yuejiang Liu, Azalia Mirhoseini, Chelsea Finn, and Marco Pavone. Scaling verification can be more effective than scaling policy learning for vision-language-action alignment, 2026. URL <https://arxiv.org/abs/2602.12281>.
- [84] Sewon Min, Kalpesh Krishna, Xinxi Lyu, Mike Lewis, Wen-tau Yih, Pang Wei Koh, Mohit Iyyer, Luke Zettlemoyer, and Hannaneh Hajishirzi. FActScore: Fine-grained atomic evaluation of factual precision in long form text generation, 2023. URL <https://arxiv.org/abs/2305.14251>.
- [85] Alexander R. Fabbri, Chien-Sheng Wu, Wenhao Liu, and Caiming Xiong. QAFactEval: Improved QA-based factual consistency evaluation for summarization, 2021. URL <https://arxiv.org/abs/2112.08542>.
- [86] Potsawee Manakul, Adian Liusie, and Mark J. F. Gales. SelfCheckGPT: Zero-resource black-box hallucination detection for generative large language models, 2023. URL <https://arxiv.org/abs/2303.08896>.
- [87] Yuejiang Liu, Fan Feng, Lingjing Kong, Weifeng Lu, Jinzhou Tang, Kun Zhang, Kevin Murphy, Chelsea Finn, and Yilun Du. World action verifier: Self-improving world models via forward-inverse asymmetry, 2026. URL <https://arxiv.org/abs/2604.01985>.
- [88] Wei-Cheng Tseng, Gashon Hussein, Yuzhu Dong, Allen Z. Ren, Lucy X. Shi, XuDong Wang, Sergey Levine, Zhaoshuo Li, Jinwei Gu, Florian Shkurti, Ming-Yu Liu, and Quan Vuong. SC3-eval: Evaluating robot foundation models via self-consistent video generation, 2026. URL <https://arxiv.org/abs/2606.18610>.

Appendix

A. Limitations and Future Work

The current framework has several limitations that suggest directions for future work. First, it assumes access to scoring-token logits, which excludes several frontier models available only through restricted APIs; in Appendix B.6 we describe a simple two-stage workaround that recovers most of the gain by routing the reasoning of a closed model through an open verifier whose logits are accessible. Second, the proposed scaling axes are not exhaustive: criteria decomposition could be learned or dynamically generated per domain rather than hand-designed, and repeated evaluation could be replaced with an adaptive compute allocation strategy guided by the verifier’s own uncertainty. Finally, while we already show that the verifier can serve as a dense reward for reinforcement learning, our experiments are limited to single-turn settings; extending it to multi-turn RL—where the verifier supplies per-step rewards over long-horizon agentic rollouts to shape credit assignment across many interdependent actions—is a promising direction for future work.

B. Additional Results and Analyses

B.1. Agent Harness Generalization on Terminal-Bench V2

To verify that the gains of LLM-as-a-Verifier are not tied to a specific agent scaffold, we repeat the Terminal-Bench V2 evaluation under two additional harnesses beyond the Capy scaffold used in our main results, each paired with the model its authors tuned for: Terminus-Kira with Claude Opus 4.6 and Terminus-2 with GPT-5.3-Codex. In every case we sample $N=5$ trajectories per task and apply the same Gemini 2.5 Flash verifier with $G=20$, $K=8$, and the three-criterion decomposition of Section 4.3; only the proposal generator and the harness change. Table 8 reports the resulting verifier accuracies alongside the official single-trajectory accuracies of Claude Opus 4.6 and Gemini 3.1 Pro under each harness.

Table 8: **Harness generalization on Terminal-Bench V2.** LLM-as-a-Verifier significantly boosts the accuracy on two additional harnesses, Terminus-2 (71.2%) and Terminus-Kira (79.4%).

Agent Harness	LLM-as-a-Verifier	Claude Opus 4.6	Gemini 3.1 Pro
Terminus-Kira (Opus 4.6)	79.4%	74.7%	74.8%
Terminus-2 (GPT-5.3-Codex)	71.2%	62.9%	68.5%

The verifier delivers the same qualitative gain across both harnesses despite their setups, observation formats, and models. Terminus-Kira (Opus 4.6 proposals) gains ~ 5 points over the strongest baseline, and Terminus-2 (GPT-5.3-Codex proposals)—the weaker harness in absolute terms—still gains 2.7 points over Gemini 3.1 Pro and 8.3 points over Claude Opus 4.6. The transfer indicates that the verifier reasons about *terminal state and task progress* rather than about scaffold-specific syntactic patterns: the same prompt template generalizes across stylistically different rollouts.

B.2. Probabilistic Pivot Tournament: Budget–Accuracy Trade-off

We provide the full pseudocode for the LLM-as-a-Verifier pipeline. Algorithm 1 embeds the fine-grained reward of Eq. 3.1—the expectation over the verifier’s scoring-token distribution, averaged across the

C criteria and K repeated evaluations—inside Probabilistic Pivot Tournament with ring-based pivot selection: a random Hamiltonian cycle gives every candidate one “A” position and one “B” position to cancel the verifier’s positional bias, the top- k candidates by ring mean preference become the pivot set $\mathcal{P}$, and each remaining candidate is compared only against $\mathcal{P}$ using the Bradley–Terry preference of Eq. 3.2. The trajectory with the highest count-normalized score is returned, reducing the budget from $\mathcal{O}(N^2)$ to $\mathcal{O}(Nk)$ while concentrating verifications on the candidates most likely to be correct.

Algorithm 1 Probabilistic Pivot Tournament with Ring-based Pivot Selection. A random Hamiltonian cycle is first scored to give every candidate one “A” position and one “B” position; the top- k candidates by ring mean preference become the pivots, and each remaining candidate is compared only against the pivot set using the soft preference probability of Eq. 3.2.

Require: Task x ; generation policy π_θ ; verifier LM p_θ ; score tokens V_{score} with map ϕ ; criteria $\mathcal{C}$; candidates N ; repetitions K ; pivots k

Ensure: Selected trajectory τ^*

- 1: **for** $i = 1, \dots, N$ **do** ▷ candidate generation
- 2: Sample $\tau_i \sim \pi_\theta(\cdot \mid x)$
- 3: **end for**
- 4: Initialize $w_i \leftarrow 0$, $c_i \leftarrow 0$ for $i \in \{1, \dots, N\}$
- 5: Sample random permutation γ of $\{1, \dots, N\}$ ▷ ring pass
- 6: $\mathcal{E}_{\text{ring}} \leftarrow \{(\gamma_t, \gamma_{t+1 \bmod N}) : t = 1, \dots, N\}$
- 7: **for each pair** $(i, j) \in \mathcal{E}_{\text{ring}}$ **do**
- 8: $(R_i, R_j) \leftarrow (R(x, \tau_i), R(x, \tau_j))$ ▷ reward via Eq. 3.1
- 9: $p \leftarrow \sigma(R_i - R_j)$ ▷ Eq. 3.2
- 10: $w_i += p$, $w_j += 1 - p$, $c_i += 1$, $c_j += 1$
- 11: **end for**
- 12: $\mathcal{P} \leftarrow$ top- k candidates by w_i/c_i ▷ ring-based pivot selection
- 13: $\mathcal{E}_{\text{piv}} \leftarrow (\{(i, p) : i \notin \mathcal{P}, p \in \mathcal{P}\} \cup \{(p_1, p_2) \in \mathcal{P}^2 : p_1 < p_2\}) \setminus \mathcal{E}_{\text{ring}}$
- 14: **for each pair** $(i, j) \in \mathcal{E}_{\text{piv}}$ **do** ▷ $\mathcal{O}(Nk)$ pivot rounds
- 15: $(R_i, R_j) \leftarrow (R(x, \tau_i), R(x, \tau_j))$ ▷ reward via Eq. 3.1
- 16: $p \leftarrow \sigma(R_i - R_j)$ ▷ Eq. 3.2
- 17: $w_i += p$, $w_j += 1 - p$, $c_i += 1$, $c_j += 1$
- 18: **end for**
- 19: **return** $\tau^* \leftarrow \tau_{i^*}$ where $i^* \in \arg \max_i w_i/c_i$

We characterize the budget–accuracy trade-off of Probabilistic Pivot Tournament (PPT, Algorithm 1) and compare it against the V1 baseline Singh et al. [5]. We curate $N = 20$ candidate trajectories per task on Terminal-Bench V2 (89 tasks, Terminus-2 harness) and report the total number of queried pairs and selection accuracy in Table 9. PPT improves steadily as the number of pivots increases, outperforming the V1 under comparable verification budgets. With only $k=3$, PPT already surpasses the best V1 result, achieving 66.17% accuracy with 4,723 queried pairs. Increasing the number of pivots further improves accuracy: $k=5$ reaches 66.27% accuracy using 6,609 pairs, while $k=9$ achieves 67.13% accuracy with 9,630 pairs, approaching full round-robin performance while using substantially fewer comparisons.

Table 9: Probabilistic Pivot Tournament (PPT) scales with the number of pivots, achieving higher accuracy as k increases while maintaining a low verification budget on Terminal-Bench V2.

Method	Pairs queried	Accuracy (%)
pass@1	—	52.64
V1 (1 N budget) [5]	1,400	64.64
V1 (3 N budget) [5]	4,200	65.62
V1 (5 N budget) [5]	7,000	65.85
V1 (7 N budget) [5]	9,800	65.53
PPT $k=1$	2,570	65.83
PPT $k=3$	4,723	66.17
PPT $k=5$ (ours)	6,609	66.27
PPT $k=7$ (ours)	8,242	66.67
PPT $k=9$ (ours)	9,630	67.13
Full Round-Robin	13,111	67.42

B.3. LLM-as-a-Verifier as a Process and Outcome Reward Model

We evaluate LLM-as-a-Verifier as a process reward model (PRM) for per-step verification on agentic tasks and as an outcome reward model (ORM) for Best-of- N on coding and mathematics benchmarks. Used as a PRM, pass@1 scales monotonically with the number of sampled actions per step k : from 48.7% to 55.7% on TauBench and from 49.8% to 54.3% on Terminal-Bench as k grows from 1 to 9 (Table 10), at $3\times$ less compute than V_1 [5]. Used as an ORM (Table 11), it improves pass@1 by 9.5% on SWE-Bench Lite, 18.5% on AIME, and 21.3% on HMMT over the base model and outperforms pointwise and pairwise baselines, while improving verification accuracy by +6.2 points over V_1 at $1.5\times$ less compute.

Table 10: LLM-as-a-Verifier as a PRM: pass@1 scales with the number of sampled actions per step. Per-step verification with LLM-as-a-Verifier increases pass@1 monotonically as the number of candidate actions k grows.

Sampled actions per step	$k=1$	$k=3$	$k=5$	$k=9$
TauBench (Gemini 2.5 Flash, pass@1)	48.7	54.0	55.1	55.7
Terminal-Bench (Gemini 3 Flash, pass@1)	49.8	51.4	52.0	54.3

Table 11: LLM-as-a-Verifier as an ORM improves pass@1 accuracy across coding and mathematics benchmarks. Under Best-of- N sampling, LLM-as-a-Verifier outperforms the base model and the pointwise/pairwise verifier baselines, with absolute improvements of 9.5% on SWE-Bench Lite, 18.5% on AIME, and 21.3% on HMMT over the base model, while using a smaller verification budget than V_1 [5].

Method	SWE-Bench Lite	AIME	HMMT
Base model	23.5	71.5	52.0
Pointwise (LM judge, N budget)	29.7	83.3	63.5
Pairwise (V_1 [5], $3N$ budget)	31.0	86.0	73.3
LLM-as-a-Verifier (ours, $2N$ budget)	33.0	90.0	73.3

B.4. Case Study: query-optimize

This appendix expands the query-optimize case study from Section 4.1 with the full task specification, ground-truth breakdown, and verifier reasoning trace. The trajectory pair is drawn under the OpenHands harness, with Claude Opus 4.5 as the proposal generator and Gemini 2.5 Flash as the verifier.

Task instruction.

You are given the Open English Wordnet (OEWn) database in SQLite format, located at `/app/oewn.sqlite`.

I implemented a sql query but it is not optimized. I have saved it in `/app/my-sql-query.sql`. Please make the query as efficient as possible while ensuring that the same output is produced.

Please save your solution in the file `/app/sol.sql`. This file must contain no comments, just one single sql query terminated by a semicolon.

Finally, please use sqlite syntax! Your code will not execute in sqlite if you use other dialects.

Ground-truth breakdown. Both candidate trajectories save an optimized SQL query and report a passing internal diff check, but Terminal-Bench’s hidden grader assigns reward 1 to one and reward 0 to the other. The failing trajectory’s verification step is methodologically unsound:

- The agent attempts to run the original query against `/app/oewn.sqlite` twice (60s timeout, then 5m02s) and aborts both times.
- To obtain a reference output, the agent runs `cp /app/oewn.sqlite /tmp/oewn_test.sqlite` and then issues `CREATE INDEX` on `senses(wordid)`, `senses(synsetid)`, and `synsets(synsetid)` on the copy.
- Subsequent steps compare (*original query on indexed copy*) to (*optimized query on canonical, unindexed database*)—two different physical access paths whose tied `ORDER BY` keys are not guaranteed to break the same way at the `LIMIT 500` boundary.
- All verification artifacts are then deleted (`rm /tmp/oewn_test.sqlite /tmp/original_output.txt /tmp/optimized_output.txt`), removing any evidence that could be re-checked.
- The optimized query is therefore never actually validated against the original on the canonical, unindexed database that Terminal-Bench’s grader uses, and `sol.sql` fails the hidden test.

The correct trajectory simply waits the full 5m03s for the original query to complete on the canonical database and runs a direct `diff` (exit code 0) before exiting.

Gemini 2.5 Flash reasoning trace. Across 16 reasoning traces with thinking enabled, the verifier reliably identifies the soundness issue. Excerpt:

“The agent was unable to execute the original query to completion on the provided `/app/oewn.sqlite` database within a reasonable timeframe (it was interrupted twice, once after 60s and once after 5m2s). To obtain a reference output for comparison, the agent copied the database (...) and then added indexes to this copied database. ... By modifying the database to obtain the reference output, the agent violated the implicit constraint of the task. Therefore, it did not correctly verify that its optimized query produces the same output as the original query running on the original, unindexed database.”

B.5. Scaling Repeated Evaluation and Visual Context on RoboRewardBench

Figure 10 extends the repeated-evaluation analysis of Section 4.2 to robotic manipulation. Trajectory-preference accuracy on RoboRewardBench rises from 81.5% at $K=1$ to 87.4% at $K=8$ and saturates at large K as the noise floor is reached. LLM-as-a-Verifier outperforms LLM-as-a-Judge and the trained baselines TOPReward, RoboReward-8B, and Robometer-4B at every budget, indicating that the gains of repeated evaluation transfer cleanly across modalities despite the change of input from text to multi-frame video.

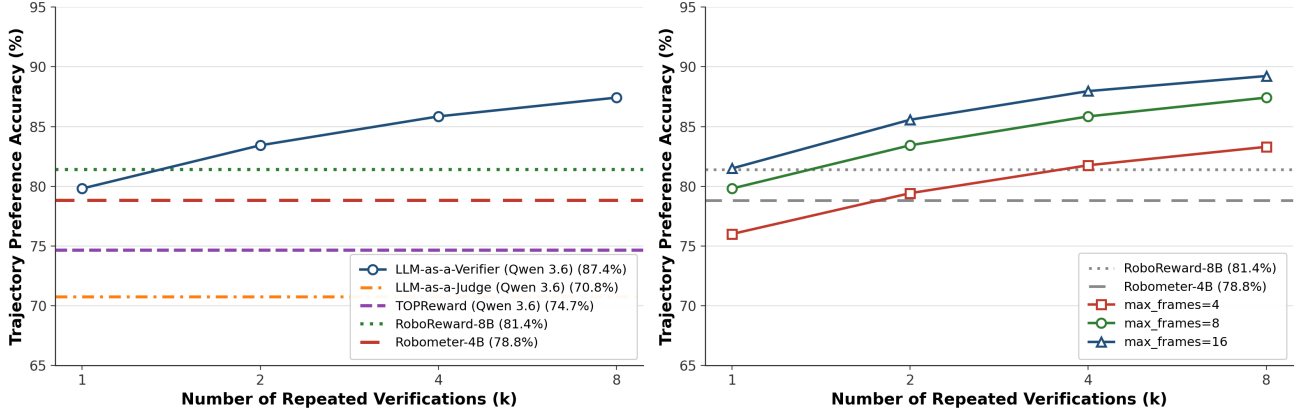


Figure 10: Trajectory-preference accuracy improves consistently as the number of repeated evaluations k increases, rising from 81.5% at $k=1$ to 87.4% at $k=8$. LLM-as-a-Verifier outperforms LLM-as-a-Judge and prior reward models (TOPReward, RoboReward-8B, Robometer-4B) across all budgets, with gains saturating at larger k .

B.6. Recovering Continuous Rewards for Logit-Restricted Frontier Models

LLM-as-a-Verifier requires access to the verifier’s scoring-token logits in order to evaluate Eq. 3.1. A growing class of frontier models, including GPT-5.5 and Claude Opus 4.7 exposes only sampled completions through their public APIs and does not return token-level logprobs, which precludes a direct in-place substitution of the verifier backbone. We describe a simple two-stage workaround that recovers most of the calibrated reward signal by decoupling reasoning from scoring.

Two-stage pipeline. For each pair (τ_i, τ_j) , we first prompt the closed model (GPT-5.5) with our standard pairwise template and require it to emit a free-form `<reasoning>...</reasoning>` block analyzing both trajectories before producing a discrete 1–10 score. We then forward the task, both trajectories, and the closed-model reasoning to an open verifier (Gemini 2.5 Flash, $G=20$) and read its logprobs at the `<score_A>` and `<score_B>` positions to compute the continuous reward of Eq. 3.1. Stage 1 contributes domain-specific reasoning quality from the frontier model; stage 2 supplies the calibrated probability distribution that the closed API withholds.

Setup and Findings. We evaluate on the same Terminal-Bench V2 swing-pair suite used throughout Section 4.2. For each pair we generate 16 independent reasoning traces with GPT-5.5 and 16 independent

stage-2 evaluations with Gemini 2.5 Flash, then average the $K \in \{1, 2, 4, 8, 16\}$ subsampled scores per trajectory and check whether $\bar{R}(x, \tau_{\text{correct}}) > \bar{R}(x, \tau_{\text{incorrect}})$. Mean accuracy and tie rate are estimated over 400 bootstrap subsamples per K . Table 12 shows that the workaround dominates the discrete baseline at every budget. At $K=1$, routing the reasoning through Gemini 2.5 Flash recovers a +5.2-point accuracy gain over directly using GPT-5.5’s integer scores (80.1% vs. 74.9%) and eliminates the 10.9% tie rate that the closed model’s coarse outputs incur. The continuous variant saturates almost immediately—accuracy moves only 1.1 points from $K=1$ to $K=16$ —whereas the discrete variant relies on heavy ensembling (+4.2 points from $K=1$ to $K=16$) primarily to break ties. Even at $K=16$, the continuous workaround leads by +2.1 points and retains zero ties.

Table 12: Accuracy and tie rate on Terminal-Bench V2 as the number of repeated evaluations K scales from 1 to 16. *GPT-5.5 (Discrete)* averages the integer 1–10 scores returned by GPT-5.5; *GPT-5.5 → Gemini 2.5 Flash (Continuous)* forwards the GPT-5.5 reasoning to Gemini 2.5 Flash and reads continuous rewards from its scoring-token logits.

K	GPT-5.5 (Discrete)		GPT-5.5 → Gemini 2.5 Flash (Continuous)	
	Accuracy (%)	Tie rate (%)	Accuracy (%)	Tie rate (%)
1	74.9	10.9	80.1	0.0
2	76.3	9.1	80.5	0.0
4	77.6	7.0	81.0	0.0
8	78.4	5.8	80.9	0.0
16	79.1	5.0	81.2	0.0

B.7. LLM-as-a-Verifier as a Dense Reward for RL

This appendix details the RL experiments of Section 7. In both settings the only difference between the baseline and our method is the reward: the policy, optimizer, hyperparameters, evaluation protocol, and random seeds are held fixed.

Off-policy RL (DSRL-SAC). We fine-tune a π_0 policy on the LIBERO-90 ketchup task with DSRL-SAC. The verifier is Qwen 3.6 35B served over SGLang and queried with the task instruction and a uniform sub-sample of $N_f=10$ rendered frames per rollout; each frame’s completion is scored on the granularity-20 scale of Eq. 3.1, decoded as the expectation over the top-20 scoring-token logprobs, and the evaluation is repeated $K = 3$ times with coefficient $\lambda=1$ (Eq. 7.1). Reward shaping is applied to successful rollouts by default; relabeling all rollouts (success and failure) is an optional variant in the code base. Relabeled transitions are written to the SAC replay buffer and the critic is trained on the shaped returns. We run 5 seeds per condition to 1.5M environment steps and report the cross-seed mean of the simulator success rate.

On-policy RL (GRPO). We fine-tune Qwen3-8B on Hendrycks MATH with GRPO using Tinker, with a group size of $M=16$, 64 groups per optimization batch, learning rate 2×10^{-5} , and a maximum generation length of 512 tokens. For each group of completions, Gemini 2.5 Flash scores the reasoning traces through the probabilistic pivot tournament (Eq. 3.2). This preference is standardized within the group and added to the correctness and format reward with weight $\beta = 0.1$ (Eq. 7.2).