
Inference-Time Reward Hacking in Large Language Models

Hadi Khalaf
Harvard University

Claudio Mayrink Verdun
Harvard University

Alex Oesterling
Harvard University

Himabindu Lakkaraju
Harvard University

Flavio du Pin Calmon
Harvard University

Abstract

A common paradigm to improve the performance of large language models is optimizing for a reward model. Reward models assign a numerical score to an LLM’s output that indicates, for example, how likely it is to align with user preferences or safety goals. However, reward models are never perfect. They inevitably function as proxies for complex desiderata such as correctness, helpfulness, and safety. By overoptimizing for a misspecified reward, we can subvert intended alignment goals and reduce overall performance – a phenomenon commonly referred to as reward hacking. In this work, we characterize reward hacking in inference-time alignment and demonstrate when and how we can mitigate it by *hedging* on the proxy reward. We study this phenomenon under Best-of- n (BoN) and Soft Best-of- n (SBoN), and we introduce Best-of-Poisson (BoP) that provides an efficient, near-exact approximation of the optimal reward-KL divergence policy at inference time. We show that the characteristic pattern of hacking as observed in practice (where the true reward first increases before declining) is an inevitable property of a broad class of inference-time mechanisms, including BoN and BoP. To counter this effect, we introduce HedgeTune, an efficient algorithm to find the optimal inference-time parameter. We demonstrate that hedging mitigates reward hacking and achieves superior reward-distortion tradeoffs on math, reasoning, and human-preference setups.

1 Introduction

Almost all current alignment methods, including BoN [1, 2], RLHF [3, 4], DPO [5], and their variants, aim to maximize a reward function while minimizing divergence from the original model’s outputs¹. It is important to distinguish between two types of rewards: proxy rewards which are the computable signals we directly use during alignment (like scores from a trained reward model), and true rewards, which represent the often latent quality of the model’s output according to a desired objective. As the name suggests, proxy rewards are approximations of the true reward and, consequently, of intended alignment goals such as correctness, helpfulness, and safety.

A fundamental challenge persists across reward-based alignment methods: **all proxy reward models are imperfect** [6]. This imperfection stems from multiple factors, including the scarcity of high-quality human-labeled data and the difficulty in formalizing high-level alignment goals into quantifiable metrics [7, 8]. For instance, consider AI alignment strategies that aim to promote safety.

¹BoN and RLHF optimize an explicitly learned reward model, whereas DPO optimizes an implicit reward based on the model’s log-probabilities.

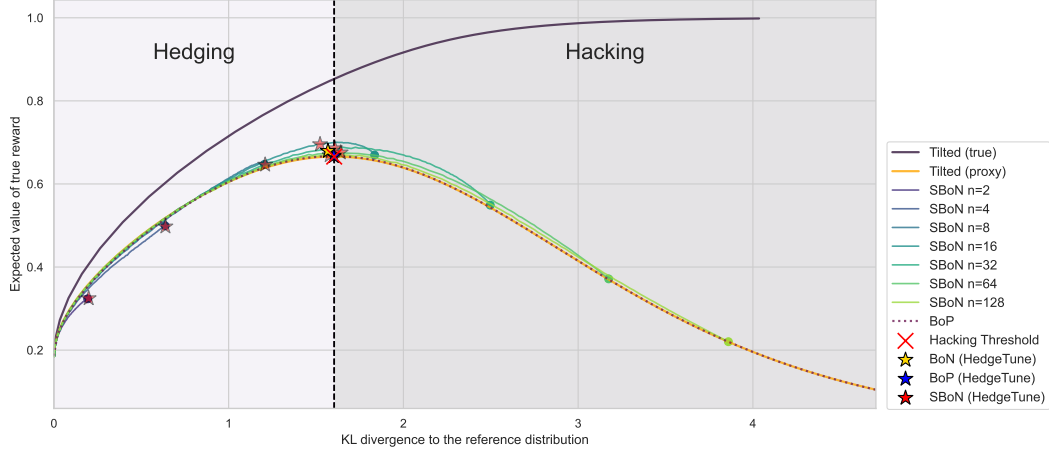


Figure 1: The mismatch between the proxy and true rewards manifests through the winner’s curse. In an ideal world where we could optimize directly on the true reward, its value would rise monotonically. However, since we are optimizing for a proxy, the true reward peaks and then collapses. The point at which we find the optimal tradeoff between maximizing reward and minimizing KL divergence from the reference distribution corresponds to the *hacking threshold*. HedgeTune successfully recovers the hacking threshold for three inference-time mechanisms: BoN, SBon, and BoP. In the case of BoN and BoP, HedgeTune recovers the optimal number of samples n . As for SBon, we fix n and find the corresponding inverse-temperature λ that maximizes the true reward. If the threshold is not achievable with any λ , HedgeTune returns the best attainable reward, as shown for low values of n .

It is difficult for a single reward model to capture nuanced human user preferences and assign accurate scalar rewards in complex, context-dependent settings where safety specifications conflict [9].

In this work, we analyze and mitigate the impact of misspecified proxy rewards in inference-time alignment. Inference-time alignment has emerged as an effective and computationally efficient paradigm to improve the capabilities of large language models and align them with desired goals [10]. Among these methods, Best-of- n (BoN) sampling stands out due to its simplicity and effectiveness. It generates n candidate responses, scores them using a reward model, and selects the one with the highest score. Empirically, BoN demonstrates competitive performance, often matching more resource-intensive fine-tuning approaches such as RLHF and DPO [11, 12]. Additionally, BoN has received an extensive theoretical treatment [13, 14]. BoN can be asymptotically equivalent to RLHF [15], enjoys non-asymptotic guarantees [16, 17], and achieves a near-optimal winrate subject to a KL divergence constraint [18]. Methods like BoN, which generate multiple samples and choose the highest-scoring one, are victims of the **winner’s curse** [19, 20, 21, 22, 23]. This phenomenon, originally identified in auction theory, occurs when selection processes systematically favor overestimates. In auctions, after each bidder submits an estimate of an item’s value, the highest bid typically overestimates the true worth, causing the winner to overpay. We demonstrate the same pattern with Best-of- n in Figure 1. As the number of candidates grows, the chance of picking outputs where the proxy reward overestimates true quality increases. This mismatch creates a critical tension: while initially optimizing for a proxy reward improves alignment with true goals, excessive optimization eventually leads to **reward hacking**, where the model exploits the proxy’s limitations, leading to worse true performance [6, 24, 25, 26, 27]. This phenomenon is also known as Goodhart’s law [28] or goal misgeneralization². Such misalignment can severely degrade trust and utility, particularly in high-stakes applications [31, 32, 33].

We mathematically characterize *inference-time reward hacking* (see Theorem 1) and provide a general framework to mitigate it (see Section 4). While this phenomenon has been observed empirically in prior work [11, 14], there has been limited theoretical analysis specific to inference-time methods and ways to mitigate it; see Section 6 for a discussion. As a result, reward hacking for inference-time alignment remains a central challenge in AI alignment. The driving question behind our work is:

When and how can we leverage useful signals from proxy rewards while mitigating hacking?

²See [29, 30] and Section 6 for a discussion of the terminology.

We focus on answering this question for inference-time alignment methods that sample multiple responses from an LLM and use reward signals to select outputs. We develop principled *hedging techniques against the winner’s curse* during inference time that precisely determine until when and how one may leverage proxy signals while preventing overoptimization.

Overview of main results. Our starting point is an optimization formulation at the heart of most alignment methods: finding a distribution π^* that maximizes a (proxy) reward r_p while remaining close (in KL-divergence) to a reference π_{ref} . This is described as the following regularized optimization problem:

$$\pi^* = \operatorname{argmax}_{\pi_x \in \Delta_{\mathcal{X}}} \mathbb{E}_{\pi_x} [r_p(X)] - \frac{1}{\lambda} D_{\text{KL}}(\pi_x \| \pi_{\text{ref}}) \quad (1)$$

Consider the information-theoretic regime where all distributions are known exactly. The solution of the above objective (1) is the exponential tilt of the reference distribution using the proxy reward [34]. Though theoretically interesting, tilted distributions cannot be sampled from directly; see Section 3 for a further discussion. Some attempts have been made to approximate this solution at inference time, i.e., when only samples from π_{ref} and black-box access to r_p are available. A notable example is Soft Best-of- n (SBoN) [17]. In this work, we show that SBoN is an effective strategy for hedging against reward hacking due to its temperature parameter λ , which allows us to smoothly interpolate between exploiting the reward (as in BoN) and staying close to the reference. However, this comes at the expense of having two tunable parameters (n, λ), which can be difficult to set in practice.

This motivates us to propose a new inference-time alignment strategy called **Best-of-Poisson (BoP)**. The idea behind BoP is simple: we run BoN with the number of samples n chosen according to a Poisson distribution. This randomization strategy induces an exponential structure (see Eq. (2)) that closely approximates the optimal reward-tilted distribution from Eq. (1). We show that BoP achieves a nearly optimal reward-distortion tradeoff at inference. Using a single tunable parameter, BoP can approximate the optimal proxy reward-tilted solution with a KL gap of order 10^{-4} when rewards are uniformly distributed, allowing us to span the entire reward-distortion region at inference (see Figure 2 and Theorem 5 in Appendix C). BoP can serve as a computationally efficient stand-in for the optimal tilted distribution with negligible loss in KL-reward tradeoff.

In practice, hedging translates to selecting parameters of inference-time alignment methods to avoid overoptimization on a proxy reward. To do so, we introduce HedgeTune: an algorithm for tuning parameters in BoN, SBoN, and BoP in order to hedge against hacking (see Algorithm 4). We illustrate the benefit of hedging in Figure 1, where we plot the expected value of the true reward versus the distortion with respect to the reference distribution for various inference-time alignment methods. If we had access to the true reward, the optimal solution would be the tilting of the reference distribution via the true reward, leading to the reward-distortion Pareto frontier (purple curve in Figure 1). However, as we are tilting via the proxy reward, we suffer from the winner’s curse: the true reward (orange curve in Figure 1) increases at first and then collapses. This behavior also manifests in BoN, as seen in the dotted points. Hedging allows us to find the *hacking threshold*: the parameters of inference-time alignment methods that yield the best tradeoff between (true) reward and distortion relative to the base model.

Our contributions are as follows:

- We mathematically formalize inference-time reward hacking (Definition 1) and derive conditions when overoptimizing imperfect proxy rewards inevitably leads to performance degradation (Theorem 1, Corollary 3).
- We introduce Best-of-Poisson (BoP), a novel inference-time alignment method (Algorithm 3). For uniformly distributed rewards, BoP approximates the optimal tilted distribution with negligible KL divergence gap (Theorem 5).
- We develop HedgeTune, a principled hedging framework that mitigates reward hacking by finding the optimal inference-time parameters (Algorithm 4). We empirically demonstrate that hedging strategies significantly outperform standard BoN sampling with minimal computational overhead on math, reasoning, and human-preference setups (Section 5).

Algorithm 1 Best-of- n Sampling (BoN)

- 1: **Input:** Integer $n \geq 1$, base policy π_{ref}
 - 2: Draw n samples $X_1, \dots, X_n$ i.i.d. from π_{ref}
 - 3: Compute proxy rewards $R_i = r_p(X_i)$ for all i
 - 4: Select $j = \arg \max_{i \in \{1, \dots, n\}} r_p(X_i)$
 - 5: **Return:** $Y = X_j$
-

2 Inference-Time Reward Hacking

In this section, we formalize inference-time reward hacking and show its inevitability under methods like BoN.

Notation and Technical Assumptions. Let $\mathcal{X}$ be a finite alphabet of tokens and let π be a probability mass function (PMF) over token sequences $x \in \mathcal{X}^*$. We denote the probability simplex over all finite sequences of tokens as $\Delta_{\mathcal{X}^*}$. Let $\pi_{\text{ref}} \in \Delta_{\mathcal{X}^*}$ be the frozen reference policy, typically a supervised fine-tuned (SFT) model. Let $r_p : \mathcal{X}^* \rightarrow \mathbb{R}$ be a proxy reward function that assigns a unique scalar value to each sequence. This is the reward we use to optimize π_{ref} during inference-time alignment. An inference-time alignment method parameterized by θ transforms the base policy π_{ref} into a new distribution π_θ . Here, θ is the **parameter of the alignment method itself** (e.g., the number of candidates n in Best-of- n) and not of the reference policy. The performance of the aligned distribution is measured using a true reward r_t . We denote our measure of interest as $f(\theta) = \mathbb{E}_{U \sim \pi_\theta} [r_t(U)]$.

For theoretical tractability, we adopt the assumption that proxy rewards can be transformed to have a continuous uniform distribution, stated next. While the policy π_{ref} itself can be complex and non-uniform, we show that we only need to consider the one-dimensional uniform distribution of the proxy reward when analyzing BoN and its variants. We later relax this assumption in the appendix and show that our results extend to the discrete case (see B.2).

Assumption 1 (Uniform Reward Mapping). The proxy rewards $r_p(x)$ obtained by sampling sequences from the reference policy, $x \sim \pi_{\text{ref}}$, are uniformly distributed over $[0, 1]$.

Assuming uniformly distributed proxy rewards incurs little loss of generality. The proxy reward scores do not need to be uniform themselves. We first map the proxy reward scores to a standardized discrete space by transforming them into their quantiles using their own empirical CDF F_p [13, 35]. This process, defined by the relation $u = F_p(r_p)$, ensures that the new reward u is uniform on $[0, 1]$ by construction. This transformation preserves the rank-ordering of the scores: an output with a higher proxy reward will also have a higher transformed score. However, the resulting random variable will not be a continuous uniform random variable (which we assume), and instead will be a discrete uniform random variable. [18] has shown for a sufficiently dense set of examples, that the error incurred by this continuous model is negligible in the context of LLMs.

Inference-time alignment. The core challenge we address is the mismatch between the proxy reward r_p and the true reward r_t . We focus on the family of inference-time methods that first sample a pool of candidate outputs and then use their proxy reward scores to define the selection mechanism. Two examples from this family are:

1. **Best-of- n** (see Algorithm 1). BoN places all probability mass on the sample with the highest proxy reward.
2. **Soft Best-of- n** (see Algorithm 2). SBoN is a generalization of BoN recently proposed by [17]. It applies a temperature-scaled softmax over candidate scores. As $\lambda \rightarrow 0$, SBoN sampling approaches uniform selection among the n candidates. As $\lambda \rightarrow \infty$, SBoN converges to standard BoN sampling.

We first formalize inference-time reward hacking through the following definition.

Definition 1 (Inference-Time Reward Hacking). Let π_θ be a distribution induced by an inference-time alignment method with parameter θ , where we assume increasing θ increases both the expected proxy reward $\mathbb{E}_{X \sim \pi_\theta} [r_p(X)]$ and the KL-divergence $D_{\text{KL}}(\pi_\theta \| \pi_{\text{ref}})$. We say that *inference-time reward hacking* occurs when there exists a threshold $\theta^\dagger$ such that for $\theta > \theta^\dagger$, $\mathbb{E}_{X \sim \pi_{\theta^\dagger}} [r_t(X)] > \mathbb{E}_{X \sim \pi_\theta} [r_t(X)]$ (i.e., the true reward decreases), despite the proxy reward and KL-divergence continuing to increase. The largest value of $\theta^\dagger$ for which this holds is called the *hacking threshold*.

Algorithm 2 Soft Best-of- n Sampling (SBoN)

- 1: **Input:** Integer $n \geq 1$, inverse temperature $\lambda > 0$, base policy π_{ref}
- 2: Draw n samples $X_1, \dots, X_n$ i.i.d. from π_{ref}
- 3: Compute proxy rewards $R_i = r_p(X_i)$ for all i
- 4: Sample index $Z \in \{1, \dots, n\}$ with probability

$$\Pr(Z = i) = \frac{e^{\lambda r_p(X_i)}}{\sum_{j=1}^n e^{\lambda r_p(X_j)}}$$

- 5: **Return:** $Y = X_Z$
-

Definition 1 offers a concrete basis for operationalizing and measuring the winner’s curse for inference time methods. The hacking threshold $\theta^\dagger$ is the ideal operating parameter for an inference-time alignment method. The following theorem establishes that, under common conditions, the shape of $f(\theta)$ is well-behaved: it either varies monotonically or reaches exactly one extremum.

Theorem 1 (Inevitability of Reward Hacking). Let $\{\pi_\theta\}_{\theta \in \Theta \subset \mathbb{R}}$ be a family of distributions with density $p_\theta(x)$ on a common support $\mathcal{X}$ such that (i) $p_\theta(x)$ is strictly totally positive of order 2 (TP₂) in (θ, x) , and (ii) its score function $\psi(x, \theta) := \partial_\theta \log p_\theta(x)$ is continuous in x and strictly increasing in x for each fixed θ . For any bounded, non-negative true reward $r_t : \mathcal{X} \rightarrow [0, \infty)$ define

$$f(\theta) := \mathbb{E}_{X \sim \pi_\theta}[r_t(X)]$$

Then f is either monotone in θ or possesses a single **unique** interior extremum $\theta^\dagger$.

Corollary 1 (Inevitability of Reward Hacking for Strictly MLR densities). Let $p_\theta(x)$ be a *strictly monotone-likelihood-ratio* in x . If the score function $\psi(x, \theta) = \partial_\theta \log p_\theta(x)$ is strictly increasing in x , then Theorem 1 applies. In particular, this applies to Best-of- n , Best-of-Poisson (to be introduced in Section 3,) and to any canonical distribution from the exponential family with strictly monotone statistic and strictly monotone natural parameter.

Four scenarios may occur: (i) *monotonic improvement*: true reward continuously increases with optimization strength; (ii) *reward hacking*: true reward initially improves but deteriorates beyond a critical threshold; (iii) *reward grokking*: true reward initially declines but then improves beyond a critical threshold; (iv) *immediate decline*: any optimization immediately harms true performance. We describe exactly when each regime occurs for MLR densities in Corollary 3.

The conditions in Theorem 1 guarantee that the inference-time policies behave in a “well-ordered” manner: increasing the tuning parameter θ consistently makes the policy “greedier” and more likely to select outputs with high proxy rewards. The clearest example is the n in Best-of- n : the larger the n , the more aggressively we optimize for the proxy. The unimodality of the true reward function then renders the problem of locating the optimal operating point $\theta^\dagger$ algorithmically tractable. To implement this insight, we develop hedging strategies that balance the exploitation of the proxy reward against the fidelity to the reference distribution. Each inference-time method offers a parameter controlling this proxy reward-KL tradeoff: n in BoN, λ in SBoN (for a fixed n), and μ in BoP (introduced in Section 3). Before introducing methods for tuning inference-time alignment methods, we first introduce **Best-of-Poisson** sampling: an alternative to BoN that approximates the optimal tilted distribution.

3 Best-of-Poisson: Approximating the Optimal Reward

While Soft Best-of- n offers a principled approach to mitigate reward hacking, it requires tuning both the number of samples n and the temperature parameter λ . In this section, we introduce Best-of-Poisson (BoP) (Algorithm 3), that is provably close to the solution of (1) with a single tunable parameter (Figure 2). BoP is of independent interest as it provides a mathematically elegant and computationally efficient way to near-optimally span the entire reward-KL distortion region with a single parameter. The key insight behind BoP is to replace the fixed sample size n in BoN with a random sample size drawn from a Poisson distribution.

The parameter μ in BoP controls the expected number of samples, analogous to how n functions in BoN. We first sample n' from a Poisson distribution parameterized by μ and set $n = n' + 1$ to ensure

Algorithm 3 Best-of-Poisson Sampling (BoP)

- 1: **Input:** Poisson parameter $\mu > 0$, base policy π_{ref}
 - 2: Sample $n' \sim \text{Poisson}(\mu)$ and set $n = n' + 1$
 - 3: Draw n samples $X_1, \dots, X_n$ i.i.d. from π_{ref}
 - 4: Compute proxy rewards $R_i = r_p(X_i)$ for all i
 - 5: Select $j = \arg \max_{i \in \{1, \dots, n\}} r_p(X_i)$
 - 6: **Return:** $Y = X_j$
-

at least one sample is generated. Under Assumption 1, the BoP distribution with parameter μ has a probability density function given by (see Appendix C)

$$q_\mu(x) = (\mu x + 1)e^{\mu(x-1)} \text{ for } x \in [0, 1]. \quad (2)$$

The following theorem characterizes the KL divergence and expected value of BoP:

Theorem 2 (KL Divergence and Expected Value of BoP). Let X_{BoP} be the random variable representing the response selected by BoP with parameter μ . Then:

$$\text{KL}(\pi_{\text{BoP}} \parallel \pi_{\text{ref}}) = \frac{e^{-\mu-1}(\text{Ei}(\mu+1) - \text{Ei}(1))}{\mu} + \log(\mu+1) - 1. \quad (3)$$

$$\mathbb{E}[X_{\text{BoP}}] = 1 - \frac{1}{\mu} + \frac{1 - e^{-\mu}}{\mu^2}. \quad (4)$$

where $\text{Ei}(z) = -\int_{-z}^{\infty} \frac{e^{-t}}{t} dt$ is the exponential integral function.

What makes BoP particularly valuable is its ability to closely approximate the solution of (1), i.e., the optimal KL-constrained tilted distribution with parameter $\lambda > 0$, defined as $\pi_\lambda^*(x) = \frac{\pi_{\text{ref}}(x)e^{\lambda r_p(x)}}{Z(\lambda)}$, where $Z(\lambda)$ is the normalization constant. While this distribution is theoretically optimal for balancing reward and divergence, computing it is intractable. To draw a next token from the tilted π_λ^* for an autoregressive LLM, one would have to compute

$$\pi_\lambda^*(x_{t+1} \mid x_{\leq t}) = \frac{\pi_{\text{ref}}(x_{\leq t+1}) e^{\lambda r_p(x_{\leq t+1})}}{\sum_{x'} \pi_{\text{ref}}(x_{\leq t} x') e^{\lambda r_p(x_{\leq t} x')}},$$

where the denominator sums over every possible continuation of the prefix $x_{\leq t}$. Because the space of continuations grows exponentially with the remaining sequence length, evaluating this denominator (and hence sampling a single token) is computationally prohibitive for LLMs.

Our analysis (Appendix C) shows that BoP provably achieves nearly identical performance to this optimal distribution with minimal KL divergence gap. In Appendix C.3, we also introduce a natural extension of Soft Best-of- n and Best-of-Poisson, which we call *Soft Best-of-Poisson (SBoP)*. As illustrated in Figure 2, the KL divergence gap between these distributions is remarkably small. Numerical evaluation confirms that for all μ , if $\lambda > 0$ is chosen so that $\mathbb{E}_{X \sim \pi_{\text{BoP}}}[r_p(X)] = \mathbb{E}_{X \sim \pi_\lambda^*}[r_p(X)]$, then the KL-gap is bounded between 0 and 8×10^{-4} .

This near-equivalence means that BoP can serve as a practical stand-in for the theoretically optimal tilted distribution. This result has two consequences. First, hedging in the optimal tilted distribution is almost equivalent to hedging with BoP. Second, Equation (1) represents the solution to the standard RLHF optimization problem, which implies that BoP is an inference-time approximation to RLHF and allows us to use BoP to easily traverse between policies rather than having to finetune a new model for each λ of interest. In the next section, we turn to the question of how to choose the right parameter value to avoid reward hacking.

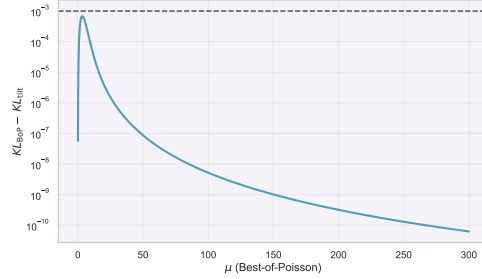


Figure 2: The difference in KL divergence when BoP and optimal tilted distributions are matched to produce the same expected reward. The extremely small gap (of order 10^{-4}) demonstrates that BoP approximates the optimal distribution with negligible performance loss

Algorithm 4 HedgeTune: Parameter Optimization for Hedging

- 1: **Inputs:** Proxy and true rewards $\{s_{t,k}, r_{t,k}\}$ per prompt t ; parameter domain Θ
 - 2: **Output:** Optimal hedge parameter θ^*
 - 3: STEP 1. For each prompt t , sort responses by their proxy scores and map their ranks to empirical quantiles $u_{t,k} \in (0, 1)$.
 - 4: STEP 2. Specify the score function $\psi(u, \theta)$ and density $p_\theta(u)$ according to the inference-time method (e.g., BoN, SBoN, BoP; see Appendix D).
 - 5: STEP 3. For a given t and $\theta \in \Theta$, define the residual $R_t(\theta) = \mathbb{E}_{u \sim p_\theta}[r_t(u) \psi(u, \theta)]$. This can be estimated from the empirical pairs $\{(u_{t,k}, r_t(u_{t,k}))\}$.
 - 6: STEP 4. Find $\theta^* \in \Theta$ such that the average residual $\bar{R}(\theta^*) = \frac{1}{|T|} \sum_t \hat{R}_t(\theta) = 0$ via one-dimensional root-finding.
-

4 Hedging to mitigate reward hacking

In this section, we develop a unified framework for choosing the inference-time parameter θ in order to maximize the expected true reward and avoid hacking. The main limitation is that we require black-box access to the true reward to perform a *one-time calibration* of the parameter θ . This is practical in several common scenarios. One example is *domains with verifiable ground truth*, such as mathematical reasoning, program synthesis, or factual question answering. One may also opt to use an LLM-as-a-judge or a more powerful but computationally expensive reward model [36, 37].

Assume that we are given the proxy and true reward scores for a set of query-response pairs. By first constructing an empirical CDF over the generated proxy reward scores, we transform these proxy scores to have a uniform distribution. We denote the transformed proxy reward as U . Each sampling method (BoN, SBoN, and BoP) induces a distribution π_θ over proxy-percentiles $u \in [0, 1]$, where θ is the corresponding parameter (sample size n , inverse-temperature λ , or Poisson rate μ). Since we know by Theorem 1 that the expected true reward has at most one peak, our key insight is to create the precise **hedge** against hacking by finding the parameter value where the marginal benefit of increasing the proxy reward equals zero. We present the following conditions that the hacking threshold must satisfy for each of the three inference-time methods.

Theorem 3 (Hacking Threshold Characterization). Let r_t be a true reward oracle and $\theta^\dagger$ be the hacking threshold from Definition 1. For each inference-time method, $\theta^\dagger$ is characterized by the following conditions:

For BoN, $n^\dagger$ satisfies:

$$\nabla_\alpha \mathbb{E}_{u \sim \text{Beta}(\alpha, 1)}[r_t(u)] = \int_0^1 r_t(u) \left(\frac{1}{n^\dagger} + \ln u \right) u^{n^\dagger-1} du = 0. \quad (5)$$

For SBoN, $\lambda^\dagger$ satisfies:

$$\nabla_\lambda \mathbb{E}_{u \sim f_\lambda}[r_t(u)] = \text{Cov}_{u \sim f_{\lambda^\dagger}}(r_t(u), u) = 0, \quad (6)$$

For BoP, $\mu^\dagger$ satisfies:

$$\nabla_\mu \mathbb{E}_{u \sim f_\mu}[r_t(u)] = \mathbb{E}_{u \sim f_{\mu^\dagger}} \left[r_t(u) \left(u - 1 + \frac{u}{\mu^\dagger u + 1} \right) \right] = 0. \quad (7)$$

The proof can be found in Appendix D. Consequently, we provide HedgeTune (Algorithm 4), an algorithm that numerically solves the corresponding root-finding problem to determine the optimal inference-time parameter for BoN, SBoN, or BoP. **Note that we do not need access to the LLM distribution itself.** We use the explicit expressions of the *score* function ψ and estimate the residual function $R(\theta) = \mathbb{E}[r_t(u) \psi(u, \theta)]$ which captures the alignment between the true reward and the proxy-weighted score. The optimal parameter $\theta^\dagger$ is found efficiently as the root of this function using standard methods such as bisection or Newton’s method [38] and can later be used directly at inference.

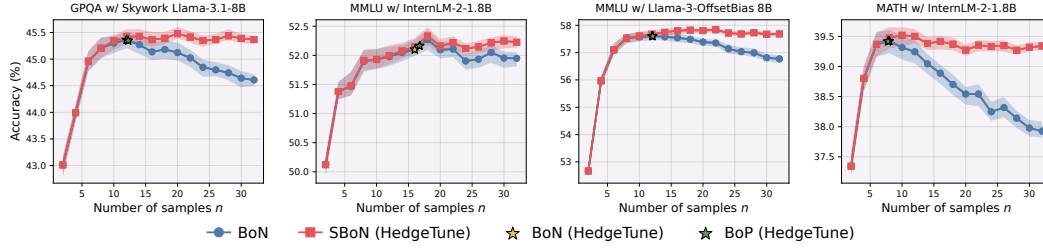


Figure 3: **Hedging mitigates hacking in verifiable reward setups.** We plot the expected accuracy on various benchmarks versus the number of samples n . HedgeTune successfully recovers the best operating point for BoN and BoP and provides a superior reward-distortion curve with SBoN.

5 Experiments: Hedging in Practice

In this section, we validate that hedging is an effective tool against reward hacking and can provide superior reward-distortion tradeoffs in two experimental setups: one with verifiable rewards and one with human-preference data.

5.1 Hacking in verifiable setups.

We first demonstrate that hedging mitigates reward hacking on standard verifiable benchmarks such as MMLU Pro and GPQA. Specifically, we use the open-source *Preference Proxy Evaluations (PPE)* dataset [39], which contains multiple responses per benchmark question from frontier models such as GPT-4o-mini and Claude Haiku 3. Each response is then scored by a suite of reward models. We focus on the benchmark dataset-reward model pairs identified in **PPE** as exhibiting reward hacking under Best-of- n sampling. We then apply HedgeTune to find the optimal operating point for BoN, BoP, and SBoN.

Reward Models. We consider three reward models of varying sizes: InternLM-2 1.8B [40], Llama-3-Offset-Bias 8B [41], and Skywork-Llama-3.1 8B [42].

Datasets. We consider three datasets: MMLU Pro (complex reasoning) [43], MATH (mathematical problem solving) [44], and GPQA (questions in natural sciences) [45]. A correct response gets a true reward of 1, while an incorrect response gets a true reward of 0.

Findings. The observed reward hacking curve in Figure 3 matches our theoretical prediction in Theorem 1. For example, BoN demonstrates hacking on the GPQA dataset, even with the very capable Skywork-Llama 3.1 8B as a proxy reward (currently ranked the 12th best non-generative reward model on RewardBench [37]). HedgeTune recovers the best operating points in all the cases.

5.2 Hacking in the wild with human preferences.

Our experimental design follows the methodology of Coste et al. [46] and Gao et al. [11], wherein proxy reward models are trained using preferences of a fixed gold reward model. In many real-world cases, we do not have access to the gold reward and instead have access to preference data. However, as we demonstrate below, a favorable operating point can still be found using traditional hyperparameter search.

Models. As a reference model, we use a 1.4B Pythia model [47] fine-tuned on AlpacaFarm dataset, but without any subsequent alignment (e.g., RLHF or DPO). This reference model is used to generate responses. We use **AlpacaRM** [48] as our gold reward model. **AlpacaRM** is an established reward model trained on human preference data and has been adopted in prior work on reward model evaluation [46, 49, 50]. This model serves as the ground truth for generating preference labels for training proxy reward models. As for proxy rewards, we use the setup of Coste et al. [46] where we train Pythia 44m models.

Datasets. We use the `tlc4418/gold_labelled_gens` dataset from Coste et al. [46]. This dataset comprises of 12,600 responses generated by the Pythia 1.4B base policy for each of 1,000 prompts. The prompts are sourced from the validation split of the AlpacaFarm dataset [48]. Each generated

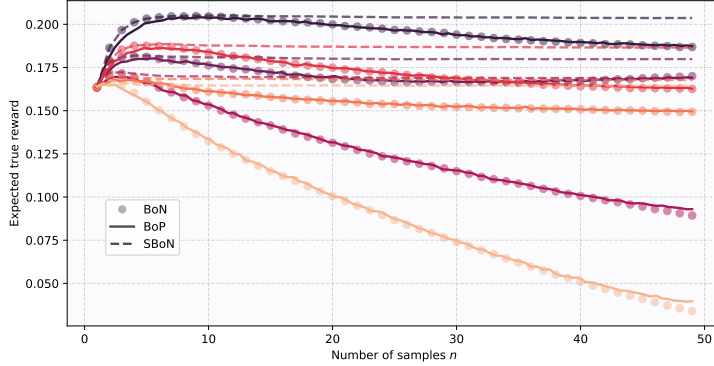


Figure 4: **Hedging mitigates hacking in human-preference setups.** We use three inference-time methods (BoN, SBoN, and BoP) on trained proxy rewards. Hacking is effectively mitigated by hedging via λ in SBoN or n in BoN and BoP.

response in this dataset is scored by **AlpacaRM**. To train proxy reward models, we construct preference datasets using the **AlpacaRM** reward scores. For each training instance, we sample a pair of responses to a given prompt and label them based on their given true reward scores.

Training. The proxy RMs are trained using a standard binary cross-entropy loss on preference pairs. We train proxy RMs on preference pair datasets of varying sizes: 10k, 20k, 46k, and 80k. In line with [46, 51, 52], we simulate disagreements in human annotators by considering two cases: (a) no label noise in the preferences, and (b) 25% label noise. All proxy RM training runs are repeated across 4 random seeds each. We present some of the runs with 25% label noise in Figure 4 and we present the remaining results in Appendix E, along with other hyperparameters and training details. Post training, we use each proxy model to score a set of 800 prompts, with 12,600 responses each.

Findings. We apply BoN, SBoN, and BoP on each run and find the expected value of the true reward as a function of n . When reward hacking manifests, we find a hacking threshold for BoN and BoP that maximizes their reward. For SBoN, with a selected $\lambda^\dagger$, we attain the peak value without suffering from reward hacking. Meanwhile, if the proxy is always at odds with the true reward, the optimal solution is the reference distribution itself, corresponding to $\lambda = 0$.

6 Related Work

Reward Hacking. Reward hacking has been widely studied in RL literature [8, 7, 53], also under the name misspecification [30], goal misgeneralization [54], or specification gaming [55].

In the context of LLMs, overoptimization has been referred to as reward hacking or Goodhart’s Law [28, 11, 26, 27]. Hacking behavior has been found to manifest in unwanted or surprising behavior [56, 57] across a variety of tasks [58, 11, 14]. Prior works have proposed various formulations of reward hacking based on true performance behavior [24], correlation between proxy and true reward [6], or distribution shift [25]. [14] prove that BoN alignment provably suffers from reward hacking when the number of samples n is large.

To address inference time hacking, a variety of methods have been explored to varying success, such as ensembling [46, 59, 60, 61], regularization [62] or rejection sampling [14]. [63] simulate n particles resampled using a softmax reward to improve performance of reasoning models, similar to SBoN over reasoning steps. However, all methods suffer from some combination of additional generation cost beyond generating n samples and estimation of additional side-quantities such as KL-divergence or χ^2 -divergence. Additionally, a variety of approaches have been proposed to mitigate reward hacking during RLHF finetuning such as regularization [64, 52, 65, 66], χ^2 -divergence [67, 6], uncertainty estimation [68], and reward pessimism [69, 70] although [26] demonstrated that RLHF can still result in reward hacking under heavy-tailed reward mismatch. Finally, prior works have also focused on improving reward models to prevent mismatch and reduce hacking [71, 72, 73, 74, 51, 75].

Best-of- n . Best-of- n sampling is a simple inference-time approach for alignment [1, 2, 76]. Prior results have characterized the expected reward gap and KL divergence between BoN sampling and the reference model and have demonstrated that BoN is asymptotically equivalent to KL-constrained reinforcement learning [13, 15, 16]. There have been various methodological improvements on BoN sampling. One such improvement is to reduce the cost of sampling n sequences via tree-based or speculative search [77, 78]. Additionally, [18, 79, 80, 81, 82] distill the BoN sampling distribution into a model via fine-tuning. Finally, [35, 83] propose inference-aware methods to improve BoN. Other works focus on improving the reward model through self-training [84]. In this work, we focus on a variant of BoN, Soft Best-of- n [17], which allows for finer control between sampling from the base model and the reward-maximizing generation.

7 Conclusion

Our work tackles the fundamental challenge that all proxy rewards are imperfect, yet they remain essential for guiding and improving AI systems. We establish a theoretical framework proving the inevitability of reward hacking in inference-time alignment and introduce practical hedging strategies to mitigate its harmful effects. By developing *Best-of-Poisson* sampling which achieves near-optimal reward-distortion tradeoffs with a single parameter and the HedgeTune algorithm for precisely calibrating inference methods, we enable practitioners to extract valuable signals from proxy rewards without falling prey to Goodhart’s law.

We also emphasize that AI safety concerns are, at their heart, socio-technical [30]. The impact of model failures due to hacking (and beyond) will depend on the application and stakeholders at hand. One clear example is an AI agent that generates controversial and toxic content on social media to optimize its engagement metric at the direct expense of content quality and safety [85]. Other dangerous failure patterns include models that learn to be sycophantic instead of truthful [86] or actively exploit loopholes to circumvent oversight [32]. By studying the phenomenon of hacking from a rigorous mathematical perspective, our work helps with the design of hacking mitigation (and hence AI safety methods) with provable performance guarantees. Technical approaches such as ours are important for AI safety in practice, although we recognize that they must exist within a larger safety ecosystem where rigorous technical foundations inform responsible deployment practices and policy decisions, and vice-versa. Ultimately, this work demonstrates that principled hedging is a promising direction for building safer, more reliable AI systems.

Acknowledgments and Disclosure of Funding

We thank the anonymous reviewers for their helpful comments and suggestions. This work is supported by the National Science Foundation under grants CIF 2312667, FAI 2040880, and CIF 2231707. This work is supported in part by NSF Awards IIS-2008461, IIS-2040989, and IIS-2238714; an AI2050 Early Career Fellowship from Schmidt Sciences; and research awards from Google, OpenAI, the Harvard Data Science Initiative, and the Digital, Data, and Design (D³) Institute at Harvard. AO is supported by the National Science Foundation Graduate Research Fellowship under Grant No. DGE-2140743. The views expressed are those of the authors and do not necessarily reflect the official policies or positions of the funding organizations.

References

- [1] Nisan Stiennon, Long Ouyang, Jeffrey Wu, Daniel Ziegler, Ryan Lowe, Chelsea Voss, Alec Radford, Dario Amodei, and Paul F Christiano. Learning to summarize with human feedback. *Advances in Neural Information Processing Systems*, 33:3008–3021, 2020.
- [2] Reiichiro Nakano, Jacob Hilton, Suchir Balaji, Jeff Wu, Long Ouyang, Christina Kim, Christopher Hesse, Shantanu Jain, Vineet Kosaraju, William Saunders, et al. Webgpt: Browser-assisted question-answering with human feedback. *arXiv preprint arXiv:2112.09332*, 2021.
- [3] Paul F Christiano, Jan Leike, Tom Brown, Miljan Martic, Shane Legg, and Dario Amodei. Deep reinforcement learning from human preferences. *Advances in Neural Information processing Systems*, 30, 2017.

- [4] Yuntao Bai, Andy Jones, Kamal Ndousse, Amanda Askell, Anna Chen, Nova DasSarma, Dawn Drain, Stanislav Fort, Deep Ganguli, Tom Henighan, et al. Training a helpful and harmless assistant with reinforcement learning from human feedback. *arXiv preprint arXiv:2204.05862*, 2022.
- [5] Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. Direct preference optimization: Your language model is secretly a reward model. *Advances in Neural Information Processing Systems*, 36, 2024.
- [6] Cassidy Laidlaw, Shivam Singhal, and Anca Dragan. Correlated proxies: A new definition and improved mitigation for reward hacking. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [7] Dylan Hadfield-Menell, Smitha Milli, Pieter Abbeel, Stuart J Russell, and Anca Dragan. Inverse Reward Design. In *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017.
- [8] Alexander Pan, Kush Bhatia, and Jacob Steinhardt. The effects of reward misspecification: Mapping and mitigating misaligned models. In *International Conference on Learning Representations*, 2022.
- [9] Maarten Buyl, Hadi Khalaf, Claudio Mayrink Verdun, Lucas Monteiro Paes, Caio C Vieira Machado, and Flavio du Pin Calmon. AI Alignment at Your Discretion. In *Proceedings of the 2025 ACM Conference on Fairness, Accountability, and Transparency (FAccT)*. ACM, 2025.
- [10] Sean Welleck, Amanda Bertsch, Matthew Finlayson, Hailey Schoelkopf, Alex Xie, Graham Neubig, Ilia Kulikov, and Zaid Harchaoui. From decoding to meta-generation: Inference-time algorithms for large language models. *Transactions on Machine Learning Research*, 2024.
- [11] Leo Gao, John Schulman, and Jacob Hilton. Scaling laws for reward model overoptimization. In *International Conference on Machine Learning*, pages 10835–10866. PMLR, 2023.
- [12] Sidharth Mudgal, Jong Lee, Harish Ganapathy, YaGuang Li, Tao Wang, Yanping Huang, Zhifeng Chen, Heng-Tze Cheng, Michael Collins, Trevor Strohman, et al. Controlled decoding from language models. In *Forty-first International Conference on Machine Learning*, 2024.
- [13] Ahmad Beirami, Alekh Agarwal, Jonathan Berant, Alexander Nicholas D’Amour, Jacob Eisenstein, Chirag Nagpal, and Ananda Theertha Suresh. Theoretical guarantees on the best-of-n alignment policy. In *Forty-second International Conference on Machine Learning*, 2025.
- [14] Audrey Huang, Adam Block, Qinghua Liu, Nan Jiang, Akshay Krishnamurthy, and Dylan J Foster. Is best-of-n the best of them? coverage, scaling, and optimality in inference-time alignment. In *Forty-second International Conference on Machine Learning*, 2025.
- [15] Joy Qiping Yang, Salman Salamatian, Ziteng Sun, Ananda Theertha Suresh, and Ahmad Beirami. Asymptotics of language model alignment. In *2024 IEEE International Symposium on Information Theory (ISIT)*, pages 2027–2032, 2024. doi: 10.1109/ISIT57864.2024.10619456.
- [16] Youssef Mroueh and Apoorva Nitsure. Information theoretic guarantees for policy alignment in large language models. *Transactions on Machine Learning Research*, 2025. ISSN 2835-8856.
- [17] Claudio Mayrink Verdun, Alex Oesterling, Himabindu Lakkaraju, and Flavio P. Calmon. Soft best-of-n sampling for model alignment. In *Proceedings of the IEEE International Symposium on Information Theory (ISIT)*, 2025.
- [18] Lin Gui, Cristina Garbacea, and Victor Veitch. Bonbon alignment for large language models and the sweetness of best-of-n sampling. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024.
- [19] Edward C Capen, Robert V Clapp, and William M Campbell. Competitive bidding in high-risk situations. *Journal of petroleum technology*, 23(06):641–653, 1971.
- [20] Max H Bazerman and William F Samuelson. I won the auction but don’t want the prize. *Journal of conflict resolution*, 27(4):618–634, 1983.

- [21] James C Cox and R Mark Isaac. In search of the winner’s curse. *Economic Inquiry*, 22(4): 579–592, 1984.
- [22] Richard H Thaler. Anomalies: The winner’s curse. *Journal of economic perspectives*, 2(1): 191–202, 1988.
- [23] Richard H Thaler. *The winner’s curse: Paradoxes and anomalies of economic life*. Simon and Schuster, 2012.
- [24] Joar Skalse, Nikolaus Howe, Dmitrii Krasheninnikov, and David Krueger. Defining and characterizing reward gaming. *Advances in Neural Information Processing Systems*, 35:9460–9471, 2022.
- [25] Lukas Fluri, Leon Lang, Alessandro Abate, Patrick Forré, David Krueger, and Joar Max Viktor Skalse. The perils of optimizing learned reward functions: Low training error does not guarantee low regret. In *Forty-second International Conference on Machine Learning*, 2025.
- [26] Thomas Kwa, Drake Thomas, and Adrià Garriga-Alonso. Catastrophic Goodhart: regularizing RLHF with KL divergence does not mitigate heavy-tailed reward misspecification. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024.
- [27] El-Mahdi El-Mhamdi and Lê-Nguyên Hoang. On goodhart’s law, with an application to value alignment. *arXiv preprint arXiv:2410.09638*, 2024.
- [28] Charles AE Goodhart and CAE Goodhart. *Problems of monetary management: the UK experience*. Springer, 1984.
- [29] Lilian Weng. Reward hacking in reinforcement learning. <https://lilianweng.github.io/posts/2024-11-28-reward-hacking/>, November 2024. Accessed: 2025-05-03.
- [30] Dario Amodei, Chris Olah, Jacob Steinhardt, Paul Christiano, John Schulman, and Dan Mané. Concrete problems in ai safety. *arXiv preprint arXiv:1606.06565*, 2016.
- [31] Alexander Bondarenko, Denis Volk, Dmitrii Volkov, and Jeffrey Ladish. Demonstrating specification gaming in reasoning models. *arXiv preprint arXiv:2502.13295*, 2025.
- [32] Bowen Baker, Joost Huizinga, Leo Gao, Zehao Dou, Melody Y Guan, Aleksander Madry, Wojciech Zaremba, Jakub Pachocki, and David Farhi. Monitoring reasoning models for misbehavior and the risks of promoting obfuscation. *arXiv preprint arXiv:2503.11926*, 2025.
- [33] Anthropic. Claude 3.7 sonnet system card, July 2025. URL <https://assets.anthropic.com/m/785e231869ea8b3b/original/claude-3-7-sonnet-system-card.pdf>. Technical Report.
- [34] Imre Csiszár, Paul C Shields, et al. Information theory and statistics: A tutorial. *Foundations and Trends® in Communications and Information Theory*, 1(4):417–528, 2004.
- [35] Ananth Balashankar, Ziteng Sun, Jonathan Berant, Jacob Eisenstein, Michael Collins, Adrian Hutter, Jong Lee, Chirag Nagpal, Flavien Prost, Aradhana Sinha, Ananda Theertha Suresh, and Ahmad Beirami. Infalign: Inference-aware language model alignment. In *Forty-second International Conference on Machine Learning*, 2025.
- [36] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, et al. Judging llm-as-a-judge with mt-bench and chatbot arena. *Advances in Neural Information Processing Systems*, 36:46595–46623, 2023.
- [37] Nathan Lambert, Valentina Pyatkin, Jacob Morrison, Lester James Validad Miranda, Bill Yuchen Lin, Khyathi Chandu, Nouha Dziri, Sachin Kumar, Tom Zick, Yejin Choi, et al. Rewardbench: Evaluating reward models for language modeling. In *Findings of the Association for Computational Linguistics: NAACL 2025*, pages 1755–1797, 2025.
- [38] Alfio Quarteroni, Riccardo Sacco, and Fausto Saleri. *Numerical mathematics*, volume 37. Springer Science & Business Media, 2006.

- [39] Evan Frick, Tianle Li, Connor Chen, Wei-Lin Chiang, Anastasios Nikolas Angelopoulos, Jiantao Jiao, Banghua Zhu, Joseph E Gonzalez, and Ion Stoica. How to Evaluate Reward Models for RLHF. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [40] Zheng Cai, Maosong Cao, Haojiong Chen, Kai Chen, Keyu Chen, Xin Chen, Xun Chen, Zehui Chen, Zhi Chen, Pei Chu, et al. Internlm2 technical report. *arXiv preprint arXiv:2403.17297*, 2024.
- [41] Junsoo Park, Seungyeon Jwa, Ren Meiyang, Daeyoung Kim, and Sanghyuk Choi. Offsetbias: Leveraging debiased data for tuning evaluators. In *Findings of the Association for Computational Linguistics: EMNLP 2024*, pages 1043–1067, 2024.
- [42] Chris Yuhao Liu, Liang Zeng, Jiakai Liu, Rui Yan, Jujie He, Chaojie Wang, Shuicheng Yan, Yang Liu, and Yahui Zhou. Skywork-reward: Bag of tricks for reward modeling in llms. *arXiv preprint arXiv:2410.18451*, 2024.
- [43] Yubo Wang, Xueguang Ma, Ge Zhang, Yuansheng Ni, Abhranil Chandra, Shiguang Guo, Weiming Ren, Aaran Arulraj, Xuan He, Ziyang Jiang, et al. MMLU-Pro: A More Robust and Challenging Multi-Task Language Understanding Benchmark. *Advances in Neural Information Processing Systems*, 37:95266–95290, 2024.
- [44] Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. Measuring Mathematical Problem Solving With the MATH Dataset. In *Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks Track (Round 2)*, 2021.
- [45] David Rein, Betty Li Hou, Asa Cooper Stickland, Jackson Petty, Richard Yuanzhe Pang, Julien Dirani, Julian Michael, and Samuel R Bowman. Gpqa: A graduate-level google-proof q&a benchmark. In *First Conference on Language Modeling*, 2024.
- [46] Thomas Coste, Usman Anwar, Robert Kirk, and David Krueger. Reward model ensembles help mitigate overoptimization. In *The Twelfth International Conference on Learning Representations*, 2024.
- [47] Stella Biderman, Hailey Schoelkopf, Quentin Gregory Anthony, Herbie Bradley, Kyle O’Brien, Eric Hallahan, Mohammad Aflah Khan, Shivanshu Purohit, USVSN Sai Prashanth, Edward Raff, et al. Pythia: A suite for analyzing large language models across training and scaling. In *International Conference on Machine Learning*, pages 2397–2430. PMLR, 2023.
- [48] Yann Dubois, Chen Xuechen Li, Rohan Taori, Tianyi Zhang, Ishaan Gulrajani, Jimmy Ba, Carlos Guestrin, Percy S Liang, and Tatsunori B Hashimoto. AlpacaFarm: A simulation framework for methods that learn from human feedback. *Advances in Neural Information Processing Systems*, 36:30039–30069, 2023.
- [49] Zaiyan Xu, Sushil Vemuri, Kishan Panaganti, Dileep Kalathil, Rahul Jain, and Deepak Ramachandran. Robust llm alignment via distributionally robust direct preference optimization. *arXiv preprint arXiv:2502.01930*, 2025.
- [50] Zhiyuan Zeng, Jiatong Yu, Tianyu Gao, Yu Meng, Tanya Goyal, and Danqi Chen. Evaluating large language models at evaluating instruction following. In *The Twelfth International Conference on Learning Representations*, 2024.
- [51] Yuchun Miao, Sen Zhang, Liang Ding, Rong Bao, Lefei Zhang, and Dacheng Tao. Inform: Mitigating reward hacking in rlhf via information-theoretic reward modeling. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024.
- [52] Rui Yang, Ruomeng Ding, Yong Lin, Huan Zhang, and Tong Zhang. Regularizing Hidden States Enables Learning Generalizable Reward Model for LLMs. *Advances in Neural Information Processing Systems*, 37:62279–62309, 2024.
- [53] Jacek Karwowski, Oliver Hayman, Xingjian Bai, Klaus Kiendlhofer, Charlie Griffin, and Joar Max Viktor Skalse. Goodhart’s law in reinforcement learning. In *The Twelfth International Conference on Learning Representations*, 2024.

- [54] Rohin Shah, Vikrant Varma, Ramana Kumar, Mary Phuong, Victoria Krakovna, Jonathan Uesato, and Zac Kenton. Goal misgeneralization: Why correct specifications aren’t enough for correct goals. *arXiv preprint arXiv:2210.01790*, 2022.
- [55] Victoria Krakovna, Jonathan Uesato, Vladimir Mikulik, Matthew Rahtz, Tom Everitt, Ramana Kumar, Zac Kenton, Jan Leike, and Shane Legg. Specification gaming: the flip side of AI ingenuity. <https://deepmind.google/discover/blog/specification-gaming-the-flip-side-of-ai-ingenuity/>, April 2020.
- [56] Carson Denison, Monte MacDiarmid, Fazl Barez, David Duvenaud, Shauna Kravec, Samuel Marks, Nicholas Schiefer, Ryan Soklaski, Alex Tamkin, Jared Kaplan, et al. Sycophancy to Subterfuge: Investigating Reward-Tampering in Large Language Models. *arXiv preprint arXiv:2406.10162*, 2024.
- [57] Lichang Chen, Chen Zhu, Jiuhai Chen, Davit Soselia, Tianyi Zhou, Tom Goldstein, Heng Huang, Mohammad Shoeybi, and Bryan Catanzaro. ODIN: Disentangled reward mitigates hacking in RLHF. In *Forty-first International Conference on Machine Learning*, 2024.
- [58] Jane Pan, He He, Samuel R Bowman, and Shi Feng. Spontaneous reward hacking in iterative self-refinement. *arXiv preprint arXiv:2407.04549*, 2024.
- [59] Jacob Eisenstein, Chirag Nagpal, Alekh Agarwal, Ahmad Beirami, Alexander Nicholas D’Amour, Krishnamurthy Dj Dvijotham, Adam Fisch, Katherine A Heller, Stephen Robert Pfohl, Deepak Ramachandran, et al. Helping or herding? reward model ensembles mitigate but do not eliminate reward hacking. In *First Conference on Language Modeling*, 2024.
- [60] Ahmed M Ahmed, Rafael Rafailov, Stepan Sharkov, Xuechen Li, and Sanmi Koyejo. Scalable ensembling for mitigating reward overoptimisation. In *ICLR 2024 Workshop on Mathematical and Empirical Understanding of Foundation Models*, 2024.
- [61] Alexandre Rame, Nino Vieillard, Leonard Hussenot, Robert Dadashi-Tazehoz, Geoffrey Cideron, Olivier Bachem, and Johan Ferret. WARM: On the Benefits of Weight Averaged Reward Models. In *International Conference on Machine Learning*, pages 42048–42073. PMLR, 2024.
- [62] Yuki Ichihara, Yuu Jinnai, Tetsuro Morimura, Kenshi Abe, Kaito Ariu, Mitsuki Sakamoto, and Eiji Uchibe. Evaluation of best-of-n sampling strategies for language model alignment. *Transactions on Machine Learning Research*, 2025. ISSN 2835-8856.
- [63] Isha Puri, Shivchander Sudalairaj, Guangxuan Xu, Kai Xu, and Akash Srivastava. A probabilistic inference approach to inference-time scaling of llms using particle-based monte carlo methods. *arXiv preprint arXiv:2502.01618*, 2025.
- [64] Paria Rashidinejad and Yuandong Tian. Sail into the headwind: Alignment via robust rewards and dynamic labels against reward hacking. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [65] Zhihan Liu, Miao Lu, Shenao Zhang, Boyi Liu, Hongyi Guo, Yingxiang Yang, Jose Blanchet, and Zhaoran Wang. Provably mitigating overoptimization in rlhf: Your sft loss is implicitly an adversarial regularizer. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024.
- [66] Yuchun Miao, Sen Zhang, Liang Ding, Yuqi Zhang, Lefei Zhang, and Dacheng Tao. The energy loss phenomenon in RLHF: A new perspective on mitigating reward hacking. In *Forty-second International Conference on Machine Learning*, 2025.
- [67] Audrey Huang, Wenhao Zhan, Tengyang Xie, Jason D Lee, Wen Sun, Akshay Krishnamurthy, and Dylan J Foster. Correcting the mythos of kl-regularization: Direct alignment without overoptimization via chi-squared preference optimization. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [68] Xiaoying Zhang, Jean-Francois Ton, Wei Shen, Hongning Wang, and Yang Liu. Overcoming reward overoptimization via adversarial policy optimization with lightweight uncertainty estimation. *arXiv preprint arXiv:2403.05171*, 2024.

- [69] Banghua Zhu, Michael Jordan, and Jiantao Jiao. Iterative data smoothing: Mitigating reward overfitting and overoptimization in RLHF. In *Forty-first International Conference on Machine Learning*, 2024.
- [70] Dhawal Gupta, Adam Fisch, Christoph Dann, and Alekh Agarwal. Mitigating preference hacking in policy optimization with pessimism. *arXiv preprint arXiv:2503.06810*, 2025.
- [71] Yizhou Chen, Yawen Liu, Xuesi Wang, Qingtao Yu, Guangda Huzhang, Anxiang Zeng, Han Yu, and Zhiming Zhou. Establishing reliability metrics for reward models in large language models. *arXiv preprint arXiv:2504.14838*, 2025.
- [72] Lingfeng Shen, Sihao Chen, Linfeng Song, Lifeng Jin, Baolin Peng, Haitao Mi, Daniel Khashabi, and Dong Yu. The trickle-down impact of reward inconsistency on RLHF. In *The Twelfth International Conference on Learning Representations*, 2024.
- [73] Jiayi Fu, Xuandong Zhao, Chengyuan Yao, Heng Wang, Qi Han, and Yanghua Xiao. Reward shaping to mitigate reward hacking in rlhf. In *ICML 2025 Workshop on Reliable and Responsible Foundation Models*, 2025.
- [74] Tianqi Liu, Wei Xiong, Jie Ren, Lichang Chen, Junru Wu, Rishabh Joshi, Yang Gao, Jiaming Shen, Zhen Qin, Tianhe Yu, Daniel Sohn, Anastasia Makarova, Jeremiah Zhe Liu, Yuan Liu, Bilal Piot, Abe Ittycheriah, Aviral Kumar, and Mohammad Saleh. RRM: Robust reward model training mitigates reward hacking. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [75] Chaoqi Wang, Zhuokai Zhao, Yibo Jiang, Zhaorun Chen, Chen Zhu, Yuxin Chen, Jiayi Liu, Lizhu Zhang, Xiangjun Fan, Hao Ma, et al. Beyond reward hacking: Causal rewards for large language model alignment. *arXiv preprint arXiv:2501.09620*, 2025.
- [76] Jacob Hilton, Peter Clark, et al. Measuring goodhart’s law: Towards an evaluation framework for open-ended generative models. OpenAI Blog, 2022. URL <https://openai.com/index/measuring-goodharts-law>. Accessed: 2025-01-30.
- [77] Jiahao Qiu, Yifu Lu, Yifan Zeng, Jiacheng Guo, Jiayi Geng, Huazheng Wang, Kaixuan Huang, Yue Wu, and Mengdi Wang. Treebon: Enhancing inference-time alignment with speculative tree-search and best-of-n sampling. *arXiv preprint arXiv:2410.16033*, 2024.
- [78] Ruiqi Zhang, Momin Haider, Ming Yin, Jiahao Qiu, Mengdi Wang, Peter Bartlett, and Andrea Zanette. Accelerating best-of-n via speculative rejection. In *2nd Workshop on Advancing Neural Network Training: Computational Efficiency, Scalability, and Resource Optimization (WANT@ICML 2024)*, 2024.
- [79] Pier Giuseppe Sessa, Robert Dadashi-Tazehoz, Leonard Hussenot, Johan Ferret, Nino Vieillard, Alexandre Rame, Bobak Shahriari, Sarah Perrin, Abram L. Friesen, Geoffrey Cideron, Sertan Girgin, Piotr Stanczyk, Andrea Michi, Danila Sinopalnikov, Sabela Ramos Garea, Amélie Héliou, Aliaksei Severyn, Matthew Hoffman, Nikola Momchev, and Olivier Bachem. BOND: Aligning LLMs with best-of-n distillation. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [80] Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajwal Bhargava, Shruti Bhosale, et al. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*, 2023.
- [81] Afra Amini, Tim Vieira, Elliott Ash, and Ryan Cotterell. Variational best-of-n alignment. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [82] Tong Yang, Jincheng Mei, Hanjun Dai, Zixin Wen, Shicong Cen, Dale Schuurmans, Yuejie Chi, and Bo Dai. Faster WIND: Accelerating Iterative Best-of- N Distillation for LLM Alignment. In *The 28th International Conference on Artificial Intelligence and Statistics*, 2025.
- [83] Yinlam Chow, Guy Tennenholtz, Izzeddin Gur, Vincent Zhuang, Bo Dai, Aviral Kumar, Rishabh Agarwal, Sridhar Thiagarajan, Craig Boutilier, and Aleksandra Faust. Inference-aware fine-tuning for best-of-n sampling in large language models. In *The Thirteenth International Conference on Learning Representations*, 2025.

- [84] Alizée Pace, Jonathan Mallinson, Eric Malmi, Sebastian Krause, and Aliaksei Severyn. West-of-n: Synthetic preference generation for improved reward modeling. *arXiv preprint arXiv:2401.12086*, 2024.
- [85] Alexander Pan, Erik Jones, Meena Jagadeesan, and Jacob Steinhardt. Feedback Loops With Language Models Drive In-Context Reward Hacking. In *International Conference on Machine Learning*, pages 39154–39200. PMLR, 2024.
- [86] Mrinank Sharma, Meg Tong, Tomasz Korbak, David Duvenaud, Amanda Askell, Samuel R. Bowman, Esin DURMUS, Zac Hatfield-Dodds, Scott R Johnston, Shauna M Kravec, Timothy Maxwell, Sam McCandlish, Kamal Ndousse, Oliver Rausch, Nicholas Schiefer, Da Yan, Miranda Zhang, and Ethan Perez. Towards understanding sycophancy in language models. In *The Twelfth International Conference on Learning Representations*, 2024.
- [87] Yuu Jinnai, Tetsuro Morimura, Kaito Ariu, and Kenshi Abe. Regularized best-of-n sampling to mitigate reward hacking for language model alignment. *arXiv preprint arXiv:2404.01054*, 2024.

Supplementary material for
**Inference-Time Reward Hacking in
Large Language Models**

A Overview

In this supplementary material, we provide the following:

- Appendix B provides additional proofs and results on the inevitability of reward hacking.
- Appendix C provides additional proofs for Best-of-Poisson, as well as Soft Best-of-Poisson.
- Appendix D discusses HedgeTune in more detail.
- Appendix E provides additional details on our experimental setup.

B Inference-Time Reward Hacking

Increasing optimization pressure on a proxy objective can initially improve true performance, but beyond a certain point—which we call the *hacking threshold*—further optimization can lead to misalignment and degrade performance. Theorem 1 in Section 2 formalizes this phenomenon: under very general conditions on a one-parameter family of proxy distributions π_θ , the map

$$\theta \mapsto \mathbb{E}_{X \sim \pi_\theta}[r_t(X)] := f(\theta) \quad (8)$$

can have at most one interior extremum. Thus, there is either a monotonic benefit (or disadvantage) to strengthening the proxy or exactly one “sweet spot” before reward hacking sets in. The rest of this section proceeds as follows. We prove the single-crossing property for the derivative of the true reward by invoking variation-diminishing kernels, and then we instantiate our results with Monotone-Likelihood-Ratio densities in Corollaries 2 and 3. We specialize the discussion to two concrete examples with Best-of- n and Best-of-Poisson. Next, we extend our result to the more general discrete case in B.2.

B.1 Uniform Case

The following results hold under our Assumption 1. Let $\psi(x, \theta)$ denote the score function of distribution $\pi_\theta(x)$ with density $p_\theta(x)$. Standard calculations under mild regularity conditions gives us the derivative of the true reward under π_θ :

$$f'(\theta) = \int r_t(x) \nabla_\theta p_\theta(x) dx = \int r_t(x) p_\theta(x) \nabla_\theta \log p_\theta(x) dx = \mathbb{E}_{X \sim \pi_\theta}[r_t(X) \psi(X, \theta)] \quad (9)$$

The main idea to establish Theorem 1 is that when we use parameter θ to control inference-time methods, we create a family of densities $\{p_\theta(x)\}$ that act as positive kernels. These kernels satisfy the strict total positivity conditions required for variation-diminishing theorems to apply. Our only assumption on the true reward is that it is bounded. We then translate the reward function so that it is non-negative. The boundedness assumption is a natural one in the alignment setting because real-world rewards originate from human judgments given on finite scales (e.g. star ratings, Likert scores, or normalized preference probabilities). Moreover, clipping or normalizing the reward prevents unbounded returns, improving the stability of policy updates and inference-time mechanisms. .

Proof. Fix θ and set $h_\theta(x) := r_t(x) \psi(x, \theta)$. Because $\psi(\cdot, \theta)$ is strictly increasing, it has at most one zero. Since $r_t \geq 0$, h_θ has the same single sign change in x . Strict TP_2 of p_θ and Karlin’s variation-diminishing theorem imply that $\theta \mapsto F(\theta) := \int h_\theta(x) p_\theta(x) dx = f'(\theta)$ inherits at most the same number of sign changes, namely one. $\square$

Having established the inevitability result, we next explore how it specializes in classical families via a simple corollary.

Corollary 2 (Strict MLR densities). Let $p_\theta(x)$ be *strictly monotone-likelihood-ratio* in x (i.e. $p_{\theta_2}(x)/p_{\theta_1}(x)$ strictly increases in x whenever $\theta_2 > \theta_1$). If $\psi(x, \theta) = \partial_\theta \log p_\theta(x)$ is strictly increasing in x , then all conclusions of Theorem 1 apply. In particular, this applies for any regular canonical exponential family with strictly monotone statistic T and strictly monotone natural parameter.

The conditions are satisfied by two inference-time methods we study:

Best-of- n : The distribution corresponds to the maximum of n i.i.d. samples from the reference distribution. When proxy rewards are uniformly distributed, this yields $p_n(u) = nu^{n-1}$ for $u \in [0, 1]$. The likelihood ratio $\frac{p_{n_2}(u)}{p_{n_1}(u)} = \frac{n_2}{n_1} u^{n_2-n_1}$ is strictly increasing in u when $n_2 > n_1$, establishing strict MLR (which implies strict TP2). The score function $\psi(u, n) = \frac{1}{n} + \log u$ is strictly increasing in u .

Best-of-Poisson: The distribution $p_\mu(u) = (\mu u + 1)e^{\mu(u-1)}$ for $u \in [0, 1]$ can be verified to satisfy strict MLR by direct computation of likelihood ratios. The score function $\psi(u, \mu) = u - 1 + \frac{u}{\mu u + 1}$ is strictly increasing in u .

Under the conditions presented in Theorem 1, we know that at most one interior extrema exists. The following corollary gives precise conditions on when such an interior extrema exists.

Lemma 1. Under the assumptions of Theorem 1 and with $r_t \in C^1$ on a neighborhood of its boundaries (0 and 1) and let $\Theta = [\theta_l, \theta_r]$,

$$\lim_{\theta \downarrow \theta_l} f'(\theta) = r'_t(0+) \mathbb{E}_{\theta_l}[X \psi(X, \theta_l)], \quad \lim_{\theta \uparrow \theta_r} f'(\theta) = -r'_t(1-) \mathbb{E}_{\theta_r}[(1-X)\psi(X, \theta_r)].$$

Then, a stationary point exists *iff*

$$\lim_{\theta \downarrow \theta_l} f'(\theta) \text{ and } \lim_{\theta \uparrow \theta_r} f'(\theta)$$

are of opposite sign (or one limit is 0 while the other is non-zero).

Proof. We prove the left boundary results (the right one is identical with $x \mapsto 1-x$.) Write the first-order expansion $r_t(x) = r_t(0+) + r'_t(0+)x + R(x)$ with $R(x) = o(x)$ as $x \rightarrow 0$. Because $\mathbb{E}_\theta[\psi] = 0$,

$$f'(\theta) = r'_t(0+)\mathbb{E}_\theta[X\psi] + \mathbb{E}_\theta[R(X)\psi].$$

Strict increase of ψ implies $|X\psi| \leq C(1+X)$ on $[0, 1] \times [\theta_l, \theta_l + \rho]$, so $\mathbb{E}_\theta[X\psi] \rightarrow \mathbb{E}_{\theta_l}[X\psi]$ by dominated convergence. Next, fix $\delta \in (0, 1)$ and split the expectation:

$$\mathbb{E}_\theta[r_t(X)\psi] = \mathbb{E}_\theta[r_t(X)\psi \mathbf{1}_{\{X \leq \delta\}}] + \mathbb{E}_\theta[r_t(X)\psi \mathbf{1}_{\{X > \delta\}}].$$

Near 0, we have $|R(x)| \leq c_\delta x$, hence the term is bounded by $c_\delta \mathbb{E}_\theta[X|\psi|]$; choose δ so small that $c_\delta \mathbb{E}_{\theta_l}[X|\psi|] < \varepsilon$ and continuity keeps it $< 2\varepsilon$ for θ close enough to θ_l . Away from 0, we use boundedness of r_t and local boundedness of ψ to obtain a factor $P_\theta\{X > \delta\} \rightarrow 0$.

Combining the two parts gives $\mathbb{E}_\theta[R(X)\psi] \rightarrow 0$, yielding the claimed limit. By continuity, the Intermediate-Value Theorem then forces one root of f' and single-crossing rules out a second. $\square$

Corollary 3 (Reward behavior for Strict MLR densities). Assume Lemma 1 holds and that the family $\{p_\theta\}_{\theta \in \Theta}$ is strictly *monotone-likelihood-ratio* in $x \in (0, 1)$. Then for every $\theta \in \Theta$

$$L_\theta := \mathbb{E}_\theta[X \psi(X, \theta)] > 0, \quad R_\theta := \mathbb{E}_\theta[(1-X)\psi(X, \theta)] < 0,$$

so

$$\text{sign } f'_t(\theta_l^+) = \text{sign } r'_t(0+), \quad \text{sign } f'_t(\theta_r^-) = \text{sign } r'_t(1-)$$

Single-crossing of f' implies that $f(\theta)$ can assume *exactly* one of the four shapes:

regime	$r'_t(0+)$	$r'_t(1-)$
monotonic improvement	≥ 0	≥ 0
reward hacking	> 0	< 0
reward grokking	< 0	> 0
immediate decline	≤ 0	≤ 0

Proof. Let p_θ be differentiable in θ with score $\psi(x, \theta) = \partial_\theta \log p_\theta(x)$. For any integrable g we have shown that:

$$\frac{d}{d\theta} \mathbb{E}_\theta[g(X)] = \mathbb{E}_\theta[g(X) \psi(X, \theta)]$$

We first use that the MLR property implies first-order stochastic dominance. For every *increasing* function g , we have that $\theta \mapsto \mathbb{E}_\theta[g(X)]$ is non-decreasing and its derivative is non-negative. Choosing $g(x) = x$ gives

$$L_\theta = \mathbb{E}_\theta[X \psi(X, \theta)] = \frac{d}{d\theta} \mathbb{E}_\theta[X] > 0$$

On the other hand, $g(x) = 1 - x$ (strictly decreasing) yields:

$$R_\theta = \mathbb{E}_\theta[(1 - X) \psi(X, \theta)] = \frac{d}{d\theta} \mathbb{E}_\theta[1 - X] < 0$$

Thus $L_\theta > 0$ and $R_\theta < 0$ for every $\theta \in \Theta$. For example, for BoN, $\psi(x, n) = 1/n + \log x$ with $\mathbb{E}_n[X \psi] = 1/(n+1)^2$, $\mathbb{E}_n[(1-X) \psi] = -1/(n+1)^2$. Hence a stationary point exists iff $r'_t(0+)$ and $r'_t(1-)$ have opposite signs, and its location $n_\star$ solves $\mathbb{E}_n[r_t(X) \psi(X, n)] = 0$. $\square$

We have shown the conditions under which the expected value of true reward has a critical point with respect to θ . We now show the conditions under which our results extend identically if we are studying the expected value of true reward as a function of the KL divergence with respect to the reference distribution π_{θ_0} .

Lemma 2. Let $\{\pi_\theta\}_{\theta \in \Theta}$ be a regular parametric family with density $p_\theta(x)$ such that

- p_θ and $\partial_\theta p_\theta$ are jointly measurable and $\partial_\theta p_\theta(x)$ is locally integrable in θ ;
- the score $\psi(x, \theta) := \partial_\theta \log p_\theta(x)$ is square-integrable: $\mathbb{E}_\theta[\psi^2] < \infty$.

For a fixed reference point $\theta_0 \in \Theta$ define the Kullback–Leibler divergence

$$D(\theta \| \theta_0) := \int p_\theta(x) \log \frac{p_\theta(x)}{p_{\theta_0}(x)} d\mu(x).$$

Then $D(\theta \| \theta_0)$ is differentiable and

$$\frac{d}{d\theta} D(\theta \| \theta_0) = \mathbb{E}_\theta[(\log p_\theta(x) - \log p_{\theta_0}(x)) \psi(X, \theta)].$$

In particular, for canonical exponential families with strictly increasing natural parameter, this simplifies to

$$\frac{d}{d\theta} D(\theta \| \theta_0) = (\eta(\theta) - \eta(\theta_0)) A''(\theta),$$

which is strictly positive when $\theta > \theta_0$ (and negative for $\theta < \theta_0$).

Proof. Write $g_\theta(x) := \log p_\theta(x) - \log p_{\theta_0}(x)$. Then $D(\theta \| \theta_0) = \mathbb{E}_\theta[g_\theta(X)]$. For a parameter-dependent integrand the classical Fisher–Leibniz rule gives

$$\frac{d}{d\theta} \mathbb{E}_\theta[g_\theta(X)] = \mathbb{E}_\theta[\partial_\theta g_\theta(X)] + \mathbb{E}_\theta[g_\theta(X) \psi(X, \theta)],$$

whenever $\partial_\theta g_\theta$ exists and an L^1 dominated-convergence bound holds (true here by the square-integrable score assumption). Since $\partial_\theta g_\theta(x) = \psi(x, \theta)$ and $\mathbb{E}_\theta[\psi] = 0$, the first term vanishes, leaving exactly

$$\frac{d}{d\theta} D(\theta \| \theta_0) = \mathbb{E}_\theta[g_\theta(X) \psi(X, \theta)]$$

$\square$

B.2 Discrete Case

In this subsection, we show that Theorem 1 holds for the general discrete case and then verify that it holds for specific policy families like Best-of- n and Best-of-Poisson.

Let $Y = \{y_1, y_2, \dots, y_m\}$ be a finite, ordered set of possible responses, where the ordering is determined by a proxy reward function r_p . We represent this space by the ordered index set $X = \{1, 2, \dots, m\}$. Let $\{\pi_\theta\}_{\theta \in \Theta}$ be a one-parameter family of policies on X , where $\theta \in \Theta \subset \mathbb{R}$ is a continuous tuning parameter. The probability of selecting response i is given by the Probability Mass Function (PMF) $\pi_\theta(i)$. Let $r_t : X \rightarrow \mathbb{R}_{\geq 0}$ be a bounded, non-negative true reward function. Our objective is to analyze the shape of the expected true reward function:

$$f(\theta) = \mathbb{E}_{X \sim \pi_\theta}[r_t(X)] = \sum_{i=1}^m \pi_\theta(i) \cdot r_t(i).$$

We will prove that under general conditions on the policy family π_θ , the function $f(\theta)$ is either monotonic or has a unique interior extremum.

Theorem 4 (Inevitability of Reward Hacking for Discrete Policies). Let $\{\pi_\theta\}_{\theta \in \Theta}$ be a family of PMFs on $X = \{1, \dots, m\}$. Assume that:

1. The PMF $\pi_\theta(i)$ is strictly TP₂ in the pair (θ, i) . That is, for any $\theta_1 < \theta_2$ and $i_1 < i_2$,

$$\pi_{\theta_1}(i_1)\pi_{\theta_2}(i_2) - \pi_{\theta_1}(i_2)\pi_{\theta_2}(i_1) > 0.$$

2. The score function $\psi(i, \theta) := \frac{\partial}{\partial \theta} \log \pi_\theta(i)$ exists and is strictly increasing in i for each fixed θ .

Then, for any bounded, non-negative true reward function $r_t(i)$, the derivative $f'(\theta)$ changes sign at most once. Hence, $f(\theta)$ is either monotonic or has a unique interior extremum.

Proof. We differentiate:

$$f'(\theta) = \frac{d}{d\theta} \sum_{i=1}^m \pi_\theta(i) r_t(i) = \sum_{i=1}^m \pi_\theta(i) \psi(i, \theta) r_t(i) = \mathbb{E}_{\pi_\theta}[r_t(i) \psi(i, \theta)].$$

Define $h_\theta(i) = r_t(i) \psi(i, \theta)$. Since $r_t(i) \geq 0$ and $\psi(i, \theta)$ is strictly increasing in i , the product $h_\theta(i)$ has at most one sign change. Since $f'(\theta) = \sum_{i=1}^m \pi_\theta(i) h_\theta(i)$, and π_θ is strictly TP₂, Karlin's variation-diminishing theorem implies that $f'(\theta)$ changes sign at most once, and $f(\theta)$ has at most one extremum. $\square$

We now verify that Best-of- n and Best of Poisson satisfy the assumptions. Let the base PMF be p_i , with CDF $F_k := \sum_{i=1}^k p_i$.

Best-of- n : It is shown in [13] that the BoN policy has the following PMF:

$$\pi_n(i) = F_i^n - F_{i-1}^n.$$

We first verify the TP₂ condition. We consider:

$$L(i) := \frac{\pi_{n_2}(i)}{\pi_{n_1}(i)} = \frac{F_i^{n_2} - F_{i-1}^{n_2}}{F_i^{n_1} - F_{i-1}^{n_1}}.$$

By Cauchy's Mean Value Theorem (CMVT), for $g(t) = t^{n_2}$, $f(t) = t^{n_1}$, and $a = F_i, b = F_{i-1}$, we get:

$$\frac{g(F_i) - g(F_{i-1})}{f(F_i) - f(F_{i-1})} = \frac{g'(c_i)}{f'(c_i)} = \frac{n_2 c_i^{n_2-1}}{n_1 c_i^{n_1-1}} = \frac{n_2}{n_1} c_i^{n_2-n_1}$$

for some $c_i \in (F_{i-1}, F_i)$. Since F_i is strictly increasing, $c_i < c_{i+1}$, and $c_i^{n_2-n_1}$ is strictly increasing. Hence $L(i+1) > L(i)$. Next, we note that the score function is:

$$\psi(i, n) = \nabla_n \log(F_i^n - F_{i-1}^n).$$

Let $g(t) = t^n \log t$, $f(t) = t^n$. Then

$$\varphi(t) = \frac{g'(t)}{f'(t)} = \frac{nt^{n-1} \log t + t^{n-1}}{nt^{n-1}} = \log t + \frac{1}{n}.$$

Applying CMVT gives:

$$\psi(i, n) = \frac{g(F_i) - g(F_{i-1})}{f(F_i) - f(F_{i-1})} = \varphi(c_i) = \log c_i + \frac{1}{n}.$$

Then:

$$\psi(i+1, n) = \log c_{i+1} + \frac{1}{n} > \log c_i + \frac{1}{n} = \psi(i, n).$$

Best-of-Poisson: Define $g(t, \mu) = te^{\mu(t-1)}$. Then, we show in Theorem 6 that BoP has the following PMF:

$$\pi_\mu(i) = g(F_i, \mu) - g(F_{i-1}, \mu).$$

We first verify the TP2 condition. Define:

$$L(i) = \frac{g(F_i, \mu_2) - g(F_{i-1}, \mu_2)}{g(F_i, \mu_1) - g(F_{i-1}, \mu_1)}.$$

By CMVT:

$$L(i) = \frac{\nabla_t g(c_i, \mu_2)}{\nabla_t g(c_i, \mu_1)} = \frac{(\mu_2 c_i + 1)e^{\mu_2(c_i-1)}}{(\mu_1 c_i + 1)e^{\mu_1(c_i-1)}} = \frac{\mu_2 c_i + 1}{\mu_1 c_i + 1} e^{(\mu_2 - \mu_1)(c_i-1)}.$$

This increases in $c_i \in (F_{i-1}, F_i)$, hence $L(i+1) > L(i)$. Now, we check the score function. Let the continuous score be $\psi(x, \mu) = x - 1 + \frac{x}{\mu x + 1}$. Since this is strictly increasing in x , so is the discrete score $\psi(i, \mu)$.

C Best-of-Poisson

In this section, we prove Theorem 2 and establish, as a consequence, that Best-of-Poisson is numerically near-optimal as compared to tilted distribution π_λ^* in terms of KL divergence. We start by establishing the BoP distribution in the uniform case, and then consider the general discrete case. We also mention a natural extension called Soft Best-of-Poisson.

C.1 Uniform Case

The following results hold under our Assumption 1. Let $\mu > 0$ be the parameter of the Best-of-Poisson sampling method and let X_μ be the random variable representing the response selected by BoP. The probability density function $q_\mu(x)$ of X_μ is given by:

$$q_\mu(x) = (1 + \mu x) e^{\mu(x-1)}, \quad (10)$$

for $x \in [0, 1]$, where $n = n' + 1$ with $n' \sim \text{Poisson}(\mu)$.

Proof. Write $X_\mu = \max\{U_0, U_1, \dots, U_{n'}\}$ where $n' \sim \text{Poisson}(\mu)$ and $U_i \stackrel{\text{iid}}{\sim} \text{Unif}[0, 1]$. Consider $U_0 \sim \text{Unif}[0, 1]$ to be the mandatory draw to achieve a sample size of at least one.

For $x \in [0, 1]$,

$$F_\mu(x) := \Pr(X_\mu \leq x) = \Pr(U_0 \leq x) \Pr(U_i \leq x \text{ for } 1 \leq i \leq n') = x \mathbb{E}[x^{n'}] = x e^{-\mu(1-x)},$$

because $\mathbb{E}[x^{n'}] = \exp\{-\mu(1-x)\}$ is the moment-generating function of a Poisson variable evaluated at $\log x$. Now, differentiating F_μ on $(0, 1)$ gives

$$q_\mu(x) = e^{-\mu(1-x)} + \mu x e^{-\mu(1-x)} = (1 + \mu x) e^{\mu(x-1)},$$

which extends continuously to the endpoints. A direct computation verifies $\int_0^1 q_\mu(x) dx = 1$, so q_μ is a valid density. $\square$

Now, we prove Theorem 2 with the BoP density denoted as q_μ and a uniform reference distribution.

Proof.

$$\mathbb{E}[X_\mu] = \int_0^1 x(1 + \mu x)e^{\mu(x-1)} dx = \int_0^1 (x + \mu x^2)e^{\mu(x-1)} dx$$

With the substitution $u = \mu(x - 1)$, we get that

$$\begin{aligned} \int_0^1 x e^{\mu(x-1)} dx &= \frac{1}{\mu^2} \int_{-\mu}^0 (u + \mu) e^u du = \frac{\mu - 1 + e^{-\mu}}{\mu^2} \\ \int_0^1 x^2 e^{\mu(x-1)} dx &= \frac{1}{\mu^3} \int_{-\mu}^0 (u + \mu)^2 e^u du = \frac{\mu^2 - 2\mu + 2 - 2e^{-\mu}}{\mu^3} \end{aligned}$$

Hence

$$\mathbb{E}[X_\mu] = \frac{\mu - 1 + e^{-\mu}}{\mu^2} + \mu \frac{\mu^2 - 2\mu + 2 - 2e^{-\mu}}{\mu^3} = 1 - \frac{1}{\mu} + \frac{1 - e^{-\mu}}{\mu^2}$$

Now, we consider the KL divergence. Because $\log q_\mu(x) = \log(1 + \mu x) + \mu(x - 1)$,

$$D_{\text{KL}}(q_\mu \| U) = \int_0^1 q_\mu(x) \log(1 + \mu x) dx + \mu[\mathbb{E}[X_\mu] - 1]$$

To compute the integral, set $t = 1 + \mu x$:

$$\int_0^1 q_\mu(x) \log(1 + \mu x) dx = \frac{e^{-\mu-1}}{\mu} \int_1^{\mu+1} t e^t \log t dt$$

Integration by parts ($f = \log t$, $dg = t e^t dt$) yields

$$\int t e^t \log t dt = \frac{1}{2} t^2 e^t \left(\log t - \frac{1}{2} \right) - \frac{1}{2} \text{Ei}(t) + C$$

hence

$$\int_0^1 q_\mu(x) \log(1 + \mu x) dx = \log(\mu + 1) - \frac{1 - e^{-\mu}}{\mu} + \frac{e^{-\mu-1}}{\mu} [\text{Ei}(\mu + 1) - \text{Ei}(1)]$$

The resulting term becomes:

$$D_{\text{KL}}(q_\mu \| \text{Unif}) = \frac{e^{-\mu-1}}{\mu} [\text{Ei}(\mu + 1) - \text{Ei}(1)] + \log(\mu + 1) - 1 \quad \square$$

We now establish that BoP provides a practical approximation to the optimal tilted distribution with negligible performance loss.

Theorem 5 (Near-Optimality of BoP). Let q_μ be the distribution induced by Best-of-Poisson with parameter $\mu > 0$, and let g_λ be the optimal KL-constrained tilted distribution with parameter $\lambda > 0$, defined as:

$$g_\lambda(x) = \frac{\pi_{\text{ref}}(x) e^{\lambda r_p(x)}}{Z(\lambda)}, \quad (11)$$

where $Z(\lambda)$ is the normalization constant. For any given expected reward level, there exists a μ for BoP and a λ for the tilted distribution such that $\mathbb{E}_{X \sim q_\mu}[r_p(X)] = \mathbb{E}_{X \sim g_\lambda}[r_p(X)]$. Numerical evaluation shows that for all μ , if $\lambda > 0$ is chosen so that $\mathbb{E}_{X \sim q_\mu}[r_p(X)] = \mathbb{E}_{X \sim g_\lambda}[r_p(X)]$, then the KL-gap satisfies

$$0 \leq D_{\text{KL}}(q_\mu \| \pi_{\text{ref}}) - D_{\text{KL}}(g_\lambda \| \pi_{\text{ref}}) \leq 8 \times 10^{-4}.$$

That is, Best-of-Poisson achieves nearly the optimal trade-off between expected reward and KL divergence from the reference distribution.

Proof. Let the following two functions

$$q_\mu(x) = (1 + \mu x) e^{\mu(x-1)}, \quad g_\lambda(x) = \frac{\lambda e^{\lambda x}}{e^\lambda - 1}, \quad 0 \leq x \leq 1,$$

denote, respectively, the BoP density with the Poisson rate μ and the exponential-tilt density with parameter λ . For any $\mu > 0$, the value $\lambda^* = \lambda^*(\mu)$ is chosen so that the two laws have the same first moment (i.e., same expected reward). We first note the expected reward of the BoP policy lies in the interval $(0.5, 1)$. Additionally, the tilted expected reward is a strictly increasing function of λ over the same range since its derivative is the (positive) variance of the reward under the tilted distribution. By the Intermediate Value Theorem, there exists a unique $\lambda^* = \mu + \delta(\mu)$, where $\delta(\mu) = \lambda^*(\mu) - \mu$, such that the two rewards are equal.

Because $g_{\lambda(\mu)}$ is the information projection of q_μ onto the exponential family $\{g_\lambda\}_{\lambda>0}$ under the linear reward-matching constraint, the Csiszár–Pythagoras identity [34] gives

$$D_{\text{KL}}(q_\mu \| \pi_{\text{ref}}) - D_{\text{KL}}(g_{\lambda(\mu)} \| \pi_{\text{ref}}) = D_{\text{KL}}(q_\mu \| g_{\lambda(\mu)}) \quad (12)$$

To bound this KL divergence, we use a general inequality which we derive below from the properties of the likelihood ratio function. Let $L(x) = \frac{p(x)}{q(x)}$ denote the likelihood ratio between two distributions p and q . We first define the KL divergence in terms of the convex function $\varphi(t) = t \log t - t + 1$ and then note that the expectation of any function is upper bounded by its essential supremum. We can retrieve the following upper bound:

$$D_{\text{KL}}(p \| q) = \mathbb{E}_q[\varphi(L(x))] \leq \sup_x \varphi(L(x)) = \sup_{t \in \text{Im}(L)} \varphi(t) \quad (13)$$

In our problem, the likelihood ratio is given by

$$L_\mu(x) := \frac{q_\mu(x)}{g_{\lambda^*}(x)} = \frac{(\mu x + 1)(e^{\lambda^*} - 1)}{\lambda^* e^\mu} \cdot e^{-\delta(\mu)x} \quad (14)$$

To obtain a uniform bound on $L_\mu(x)$, define $\alpha := \sup_{\mu>0, x \in [0,1]} |\log L_\mu(x)|$. This is a well-posed, nested optimization problem requiring numerical solution. The procedure is as follows:

1. For a given μ , numerically solve for the unique λ^* satisfying the reward-matching equation.
2. For the resulting pair (μ, λ^*) , compute $\max_{x \in [0,1]} |\log L_\mu(x)|$ by evaluating the boundary values and any interior critical points.
3. Maximize the result of step (2) over all $\mu > 0$.

This procedure is deterministic, and its solution yields a uniform bound:

$$e^{-\alpha} \leq L_\mu(x) \leq e^\alpha, \quad \text{for all } \mu > 0 \text{ and } x \in [0, 1] \quad \text{with } \alpha \approx 0.03968$$

Since the function $\varphi(t)$ is convex, we verify that its supremum over $[e^{-\alpha}, e^\alpha]$ is attained at an endpoint:

$$D_{\text{KL}}(q_\mu \| g_{\lambda(\mu)}) \leq \varphi(e^\alpha) = \alpha e^\alpha - e^\alpha + 1 \approx 0.03968 \cdot e^{0.03968} - e^{0.03968} + 1 \approx 8 \times 10^{-4}$$

This bound is independent of μ , and therefore holds uniformly as seen in Figure 2. This validates that Best-of-Poisson indeed provides a practically equivalent approximation to the optimal tilted distribution with negligible computational overhead. $\square$

C.2 Discrete Case

We now relax our Assumption 1 to derive an exact, assumption-free formulation for the Best-of-Poisson (BoP) policy, as well as a convenient upper bound for its KL divergence.

Theorem 6 (General BoP Distribution). Let $\mu > 0$ be the parameter of the Best-of-Poisson sampling method. The probability mass function of Best-of-Poisson is:

$$\pi_{\text{BoP}}^{(\mu)}(y_i | x) = g_\mu(F_i) - g_\mu(F_{i-1}) \quad \text{where} \quad g_\mu(z) := z e^{\mu(z-1)} \quad (15)$$

$$\text{KL}(\pi_{\text{BoP}} \| \pi_{\text{ref}}) \lesssim \mathbb{E}[\log N] - 1 + \mathbb{E}\left[\frac{1}{N}\right] \quad \text{where} \quad N \sim \text{Poisson}(\mu) + 1 \quad (16)$$

Proof. BoP is a mixture over BoN with a Poisson-distributed number of samples $N \sim \text{Pois}(\mu) + 1$, with PMF:

$$\Pr(N = k) = \frac{e^{-\mu} \mu^{k-1}}{(k-1)!}, \quad k \geq 1$$

The BoP policy is:

$$\pi_{\text{BoP}}(y_i | x) = \sum_{k=1}^{\infty} \Pr(N = k) \cdot \pi_{\text{BoN}}^{(k)}(y_i | x) = \sum_{k=1}^{\infty} \frac{e^{-\mu} \mu^{k-1}}{(k-1)!} (F_i^k - F_{i-1}^k)$$

Letting $j = k - 1$, this becomes:

$$\pi_{\text{BoP}}(y_i | x) = e^{-\mu} \left(\sum_{j=0}^{\infty} \frac{(\mu F_i)^j}{j!} - \sum_{j=0}^{\infty} \frac{(\mu F_{i-1})^j}{j!} \right) = F_i e^{\mu(F_i-1)} - F_{i-1} e^{\mu(F_{i-1}-1)}$$

Thus, the exact BoP PMF is:

$$\pi_{\text{BoP}}(y_i | x) = g_{\mu}(F_i) - g_{\mu}(F_{i-1}) \quad \text{where} \quad g_{\mu}(z) := z e^{\mu(z-1)}$$

We show that this is consistent with our continuous case. Assuming the reference distribution is $\text{Unif}(0, 1)$, then $F_{\text{ref}}(r) = r$ and:

$$G_{\text{BoP}}(r) = g_{\mu}(r) = r e^{\mu(r-1)} \implies q_{\mu}(r) = \frac{d}{dr} G_{\text{BoP}}(r) = (1 + \mu r) e^{\mu(r-1)} \quad (17)$$

This matches the continuous BoP PDF.

Using the exact PMF, we can get the exact KL divergence between the Best of Poisson distribution and the reference distribution, as well as a convenient upper bound.

$$\text{KL}(\pi_{\text{BoP}} \| \pi_{\text{ref}}) = \sum_{i=1}^m \pi_{\text{BoP}}(y_i | x) \log \left(\frac{\pi_{\text{BoP}}(y_i | x)}{p_i} \right) = \sum_{i=1}^m (g_{\mu}(F_i) - g_{\mu}(F_{i-1})) \log \left(\frac{g_{\mu}(F_i) - g_{\mu}(F_{i-1})}{F_i - F_{i-1}} \right) \quad (18)$$

We show that this is consistent with our continuous case. As $m \rightarrow \infty$, the discrete sum becomes:

$$\text{KL}(\pi_{\text{BoP}} \| \pi_{\text{ref}}) \rightarrow \int_0^1 q_{\mu}(r) \log q_{\mu}(r) dr$$

where $q_{\mu}(r) = (1 + \mu r) e^{\mu(r-1)}$. This integral is:

$$\int_0^1 q_{\mu}(r) \log q_{\mu}(r) dr = \frac{e^{-\mu-1}}{\mu} (\text{Ei}(\mu+1) - \text{Ei}(1)) + \log(\mu+1) - 1$$

where $\text{Ei}(z)$ is the exponential integral. The KL divergence is convex in its first argument. That is, for any set of distributions $\{P_k\}$ and weights $\{w_k\}$,

$$\text{KL} \left(\sum_k w_k P_k \parallel Q \right) \leq \sum_k w_k \cdot \text{KL}(P_k \parallel Q)$$

Applying this to BoP:

$$\text{KL}(\pi_{\text{BoP}} \| \pi_{\text{ref}}) \leq \sum_{k=1}^{\infty} w_k \cdot \text{KL}(\pi_{\text{BoN}}^{(k)} \| \pi_{\text{ref}}) = \mathbb{E}_{N \sim \text{Pois}(\mu)+1} \left[\text{KL}(\pi_{\text{BoN}}^{(N)} \| \pi_{\text{ref}}) \right]$$

Using the proven upper bound for BoN in [13], we get that

$$\text{KL}(\pi_{\text{BoP}} \| \pi_{\text{ref}}) \lesssim \mathbb{E}_N \left[\log N - 1 + \frac{1}{N} \right] = \mathbb{E}[\log N] - 1 + \mathbb{E} \left[\frac{1}{N} \right] \quad \text{where} \quad N \sim \text{Poisson}(\mu) + 1$$

□

Algorithm 5 Soft Best-of-Poisson Sampling (SBoP)

- 1: **Input:** Poisson parameter $\mu > 0$, inverse temperature $\lambda > 0$, base policy π_{ref}
- 2: Sample $n' \sim \text{Poisson}(\mu)$, set $n = n' + 1$
- 3: Draw $X_1, \dots, X_n \sim \pi_{\text{ref}}$ i.i.d.
- 4: Compute rewards $R_i = r_p(X_i)$
- 5: Sample index $Z \in \{1, \dots, n\}$ with probability

$$\Pr(Z = i) = \frac{e^{\lambda R_i}}{\sum_{j=1}^n e^{\lambda R_j}}$$

- 6: **Return:** $Y = X_Z$
-

C.3 Soft Best-of-Poisson

We discuss here a natural extension of Soft Best-of- n and Best-of-Poisson, which we name Soft Best-of-Poisson (SBoP) (see Algorithm 5). The key insight is that we can leverage the two distinct control mechanisms that inference-time alignment offers us: control over the number of generation (be it deterministic or randomized) and control over the selection through the temperature. Extensions such as SBoP show that these methods can be combined, offering richer control over the alignment process.

Theorem 7. Let the parameters of SBoN be $\theta = (n, \lambda)$. Let $R_1, \dots, R_n$ be the proxy rewards of n i.i.d. samples from π_{ref} . By Assumption 1, $R_i \sim \text{Unif}(0, 1)$ for all i . The selected sample has proxy reward R_Y with probability density function:

$$p_{n,\lambda}(r) = n \cdot \mathbb{E}_{R_2, \dots, R_n \sim \text{Unif}(0,1)} \left[\frac{e^{\lambda r}}{e^{\lambda r} + \sum_{j=2}^n e^{\lambda R_j}} \right] \quad (19)$$

Let the log-partition function be defined as:

$$L(n, \lambda) := \mathbb{E}_{R_1, \dots, R_n \sim \text{Unif}(0,1)} \left[\log \left(\sum_{i=1}^n e^{\lambda R_i} \right) \right]$$

Then, the expected reward and the KL divergence with respect to the reference can be written as:

$$\mathbb{E}[R_{\text{SBoN}}] = \frac{\partial L(n, \lambda)}{\partial \lambda} \quad (20)$$

$$\text{KL}(\pi_{\text{SBoN}} \parallel \pi_{\text{ref}}) = \log n + \lambda \mathbb{E}[R_{\text{SBoN}}] - L(n, \lambda) \quad (21)$$

Proof. Expected Reward:

$$\mathbb{E}[R_{\text{SBoN}}] = \mathbb{E} \left[\sum_{i=1}^n R_i \cdot \frac{e^{\lambda R_i}}{\sum_{j=1}^n e^{\lambda R_j}} \right] = \mathbb{E} \left[\frac{\sum_{i=1}^n R_i e^{\lambda R_i}}{\sum_{j=1}^n e^{\lambda R_j}} \right] = \frac{\partial L(n, \lambda)}{\partial \lambda}$$

since we observe:

$$\frac{\partial}{\partial \lambda} \log \left(\sum_{j=1}^n e^{\lambda R_j} \right) = \frac{\sum_{i=1}^n R_i e^{\lambda R_i}}{\sum_{j=1}^n e^{\lambda R_j}}$$

KL Divergence: Since $p_{\text{ref}}(r) = 1$ for $r \in [0, 1]$, we get:

$$\text{KL}(\pi_{\text{SBoN}} \parallel \pi_{\text{ref}}) = \int_0^1 p_{n,\lambda}(r) \log p_{n,\lambda}(r) dr$$

Alternatively, the conditional KL (per draw) is:

$$\text{KL}_{\text{cond}} = \sum_{i=1}^n p_i \log(np_i) = \log n + \sum_{i=1}^n p_i \log p_i = \log n + \lambda \sum_{i=1}^n p_i R_i - \log \left(\sum_{j=1}^n e^{\lambda R_j} \right)$$

where $p_i = \frac{e^{\lambda R_i}}{\sum_{j=1}^n e^{\lambda R_j}}$. Taking expectation over rewards gives our result. $\square$

D Analysis of HedgeTune

In Section 4, we presented an efficient way to find the optimal hacking threshold. Here, we prove this result and discuss how hedging compares to other hacking mitigation methods.

Proof. We consider each mechanism separately:

Hedging in Best-of- n . In BoN, we approximate the integer n via a continuous parameter α by placing a $\text{Beta}(\alpha, 1)$ prior on u . Its density is $f_\alpha(u) = \alpha u^{\alpha-1}$, so $\psi(u, \alpha) = \partial_\alpha [\ln \alpha + (\alpha-1) \ln u] = \frac{1}{\alpha} + \ln u$. Thus, the optimality condition becomes

$$\mathbb{E}_{u \sim \text{Beta}(\alpha, 1)} \left[r_t(u) \left(\frac{1}{\alpha} + \ln u \right) \right] = 0 \iff \int_0^1 r_t(u) (1 + n \log u) u^{n-1} du = 0, \quad (22)$$

which one solves for α to pick an effective sample size.

In practice, this equation must be solved numerically. One way to do this by discretizing $[0, 1]$ into M points and forming the Riemann-sum residual

$$R(\alpha) = \sum_{i=1}^M r_t(u_i) \left(\frac{1}{\alpha} + \ln u_i \right) u_i^{\alpha-1} \Delta u,$$

The root $R(\alpha) = 0$ is equivalent to the hedging condition. Then, one applies any root-finding method, see, e.g., [38], to locate the unique solution $\alpha^\dagger$. Alternatively, for each question/prompt, we first sort the candidates by their model scores to form the empirical CDF of the scores. We perform a *discrete ternary search* to identify the optimal $n^\dagger$ using the question-averaged true reward. This procedure bypasses explicit root-finding and directly identifies the most effective discrete hedge size.

Hedging in Soft Best-of- n . Unlike BoN, SBoN does not admit a simple closed-form density due to its sampling mechanism. In SBoN, we first sample n uniform rewards $\mathbf{U}^n = U_1, \dots, U_n$, then select U_i with probability $p_i = \frac{e^{\lambda U_i}}{\sum_{j=1}^n e^{\lambda U_j}}$. The resulting distribution is:

$$\pi_{n,\lambda}(u) = \mathbb{E}_{U_1, \dots, U_{n-1} \sim \text{Unif}[0,1]} \left[\frac{n, e^{\lambda u}}{e^{\lambda u} + \sum_{i=1}^{n-1} e^{\lambda U_i}} \right], \quad u \in [0, 1]. \quad (23)$$

We can now derive the hedging condition $\frac{\partial}{\partial \lambda} \mathbb{E}_{u \sim f_\lambda} [r_t(u)] = 0$. To start, we compute $\frac{\partial p_i}{\partial \lambda}$:

$$\frac{\partial p_i}{\partial \lambda} = \frac{\partial}{\partial \lambda} \left[\frac{e^{\lambda U_i}}{S} \right] = \frac{U_i e^{\lambda U_i} \cdot S - e^{\lambda U_i} \cdot \frac{\partial S}{\partial \lambda}}{S^2} = p_i \left[U_i - \sum_{j=1}^n U_j p_j \right] \quad (24)$$

since $\frac{\partial S}{\partial \lambda} = \sum_{j=1}^n U_j e^{\lambda U_j}$. Now, we compute the derivative of the expected reward:

$$\frac{\partial}{\partial \lambda} \mathbb{E}_{u \sim f_\lambda} [r_t(u)] = \mathbb{E}_{\mathbf{U}^n} \left[\frac{\partial}{\partial \lambda} \sum_{i=1}^n r_t(U_i) p_i \right] \quad (25)$$

$$= \mathbb{E}_{\mathbf{U}^n} \left[\sum_{i=1}^n r_t(U_i) \frac{\partial p_i}{\partial \lambda} \right] \quad (26)$$

$$= \mathbb{E}_{\mathbf{U}^n} \left[\sum_{i=1}^n r_t(U_i) p_i \left(U_i - \sum_{j=1}^n U_j p_j \right) \right] \quad (27)$$

$$= \mathbb{E}_{\mathbf{U}^n} \left[\sum_{i=1}^n r_t(U_i) p_i U_i - \left(\sum_{i=1}^n r_t(U_i) p_i \right) \left(\sum_{j=1}^n U_j p_j \right) \right] \quad (28)$$

$$= \mathbb{E}_{\mathbf{U}^n} [\text{Cov}(r_t(V), V | \mathbf{U})] \quad (29)$$

where we note that the expression inside the expectation is exactly the conditional covariance:

$$\text{Cov}(r_t(V), V|\mathbf{U}^n) = \mathbb{E}[r_t(V) \cdot V|\mathbf{U}^n] - \mathbb{E}[r_t(V)|\mathbf{U}^n] \cdot \mathbb{E}[V|\mathbf{U}^n] \quad (30)$$

with

$$\mathbb{E}[r_t(V)|\mathbf{U}^n] = \sum_{i=1}^n r_t(U_i)p_i, \quad \mathbb{E}[V|\mathbf{U}^n] = \sum_{i=1}^n U_i p_i, \quad \mathbb{E}[r_t(V) \cdot V|\mathbf{U}^n] = \sum_{i=1}^n r_t(U_i)U_i p_i \quad (31)$$

This condition must be evaluated numerically using the following procedure:

1. For a given λ , and for each question, draw M independent samples $(U_1^{(m)}, \dots, U_n^{(m)})$ from the *empirical CDF* of model scores (equivalently, from the empirical quantiles $u \in [0, 1]$), for $m = 1, \dots, M$.
2. For each realization, compute:

$$p_i^{(m)} = \frac{e^{\lambda U_i^{(m)}}}{\sum_{j=1}^n e^{\lambda U_j^{(m)}}}. \quad (32)$$

$$C^{(m)} = \sum_{i=1}^n r_t(U_i^{(m)}) U_i^{(m)} p_i^{(m)} - \left(\sum_{i=1}^n r_t(U_i^{(m)}) p_i^{(m)} \right) \left(\sum_{j=1}^n U_j^{(m)} p_j^{(m)} \right) \quad (33)$$

3. Estimate the residual as

$$R(\lambda) = \frac{1}{M} \sum_{m=1}^M C^{(m)}. \quad (34)$$

4. Locate $\lambda^\dagger$ such that $R(\lambda^\dagger) = 0$ using a numerical root-finding method (e.g., bisection or Newton's method).

Hedging in Best-of-Poisson. Here, one draws a $\text{Poisson}(\mu)$ number of samples (plus one) and selects the proxy-maximal u . For uniform u , the density is $p_\mu(u) = (\mu u + 1)e^{-\mu(1-u)}$, giving $\psi(u, \mu) = \partial_\mu \ln p_\mu(u) = u - 1 + \frac{u}{\mu u + 1}$. The hedging equation

$$\nabla_\mu \mathbb{E}_{\pi_\mu}[r_y(X)] = \mathbb{E}_{u \sim f_\mu} \left[r_t(u) \left(u - 1 + \frac{u}{\mu u + 1} \right) \right] = 0 \quad (35)$$

In an analogous way to the previous hedging equations, we solve the residual

$$R(\mu) = \sum_{i=1}^M r_t(u_i) \psi(u_i, \mu) p_\mu(u_i) \Delta u \quad (36)$$

to locate $\mu^\dagger$ such that $R(\mu^\dagger) = 0$. This $\mu^\dagger$ is the Poisson optimal hedge. As in BoN, we approximate this expectation using the empirical CDF of scores per question, mapping each candidate's rank to its quantile u_i . We compute the averaged residual across questions, and the root $\bar{R}(\mu^\dagger) = 0$ is found by *bracketing*, i.e, finding an interval $[a, b]$ such that $\bar{R}(a)\bar{R}(b) < 0$, and *bisection search within this interval*. $\square$

D.1 How does hedging compare to other hacking mitigation methods?

In this subsection, we supplement the discussion in the related works (Section 6). We discuss different techniques mitigate reward hacking and compare with hedging.

Training-Time Regularization. This is the most common approach, typified by adding a divergence penalty to the reward maximization objective as in RLHF. The goal is to prevent the policy from deviating too far from a trusted reference model. While a vital first line of defense, this approach requires expensive full model retraining and risks under-optimization if the penalty is too strong.

Improving the Proxy Reward Model. A second category aims to build a more robust reward signal that is inherently harder to hack. Ensembling involves averaging scores from multiple, independently trained reward models. However, this method multiplies inference costs and cannot fix systemic biases shared by all models in the ensemble. Meanwhile, robust training methods [74, 51] modify the reward model’s training process itself to disentangle response quality from spurious cues, introducing significant complexity to the training pipeline.

Inference-Time Methods. The closest analogues to our work include regularized Best-of- n (RBoN), which is a powerful heuristic that balances the proxy reward against a penalty term [87], and inference-time pessimism [14] which uses principled rejection sampling step, requiring the estimation of a normalization constant.

Our Framework: Hedging introduces a distinct approach that operates purely at inference time. It answers the question: given a fixed reward model and fixed language model, what is the best we can do to improve performance just at inference? Hedging offers two key advantages:

1. **Zero Retraining Cost and Maximum Flexibility:** Hedging is applied to an existing, trained policy using existing proxy reward models. A practitioner can take a single policy and calibrate it against different proxy rewards or for different tasks without incurring any retraining.
2. **Principled and Efficient.** We prove that this expected true reward function follows a predictable “rise-and-fall” path with a single optimal peak. *HedgeTune* therefore replaces heuristic balancing or more involved sampling schemes with a principled search for this unique optimum—a problem we prove is tractable.

However, the gains from hedging might be minimal compared to more involved methods, especially with a weak proxy reward. One limitation of our work is that hedging requires a tuning “ground truth” dataset. Ultimately, we see hedging not as a replacement for some of the previous methods. Instead, *HedgeTune* offers a computationally lightweight, theoretically-grounded, and highly flexible method to mitigate hacking at deployment.

E Experimental Details

In this section, we provide additional details on our experimental setup. We provide our code here. The experiments with the toy example and the verifiable rewards were done on a CPU (Apple M3 Chip). The experiments with the human-preferences were done on 1 A100 GPU with 48 GPU hours (including preliminary experiments and testing).

E.1 Toy Example

In Figure 1, we present a toy example with one-dimensional rewards defined on $[0, 1]$. We set the proxy reward to be $r_p(x) = x$ and the gold reward

$$r_t(x) = \frac{x^p(1-x)}{C}, \quad C = \left(\frac{p}{p+1}\right)^p \frac{1}{p+1}$$

where C is a normalization constant so that the gold reward is bounded between 0 and 1 for convenience. We choose this gold reward, as the reward under the Best-of- n distribution has a simple closed-form solution. We set $p = 12$ and we find the expected value of true reward and the KL divergence with respect to the reference distribution under four mechanisms:

1. **Tilted Distribution:** Exponential tilting of the gold reward, $Q_\lambda(x) \propto \exp(\lambda r_t(x))$ and of the proxy reward, $Q_\lambda(x) \propto \exp(\lambda r_p(x))$.
2. **Best-of- n (BoN):** Selection of the maximum of n i.i.d. uniform draws.
3. **Soft Best-of- n (SBoN):** Softmax-based sampling of n i.i.d. uniform draws with inverse temperature λ .
4. **Best-of-Poisson (BoP):** Selection of the maximum of n i.i.d uniform draws where n is drawn from a Poisson distribution with rate μ .

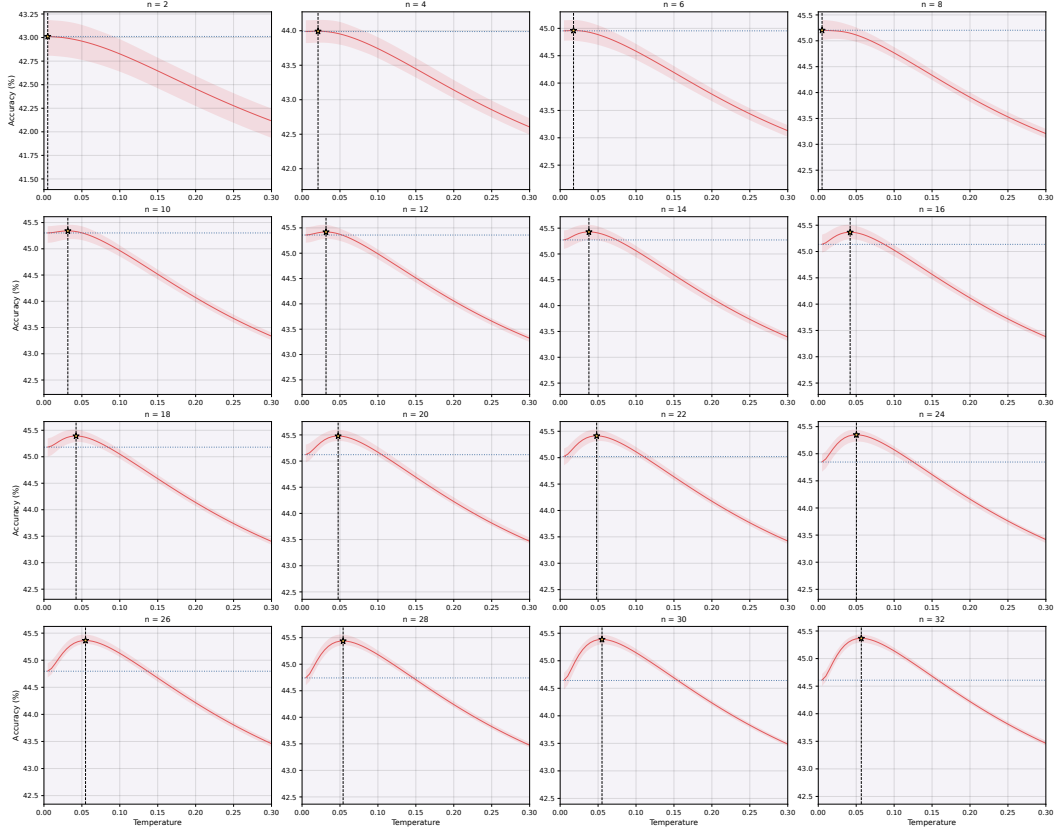


Figure 5: Accuracy vs. temperature for different sample sizes n with GPQA and Skywork Llama-3.1 8B. HedgeTune identifies the optimal temperature (dashed line) for each n .

We apply HedgeTune as presented in Alg. 4 for BoN and BoP. In both cases, we obtain an operating point that corresponds exactly to the true hacking threshold, as shown in Figure 1. The success of this algorithm hinges on Theorem 1 which guarantees the existence of (at least) one hacking threshold. However, this guarantee does not hold for Soft BoN. Therefore, the optimization problem becomes more challenging, and using vanilla estimators for the density causes numerical instabilities when applied for Soft BoN.

E.2 Reward hacking in verifiable settings.

We include additional experiments from running HedgeTune with SBoN on the **PPE** dataset. HedgeTune finds the best operating temperature for a given n . We plot the expected value of the accuracy and show the retrieved maximizing temperature. In particular, with larger n , we need a larger temperature to mitigate hacking. This corresponds to a stronger KL regularization.

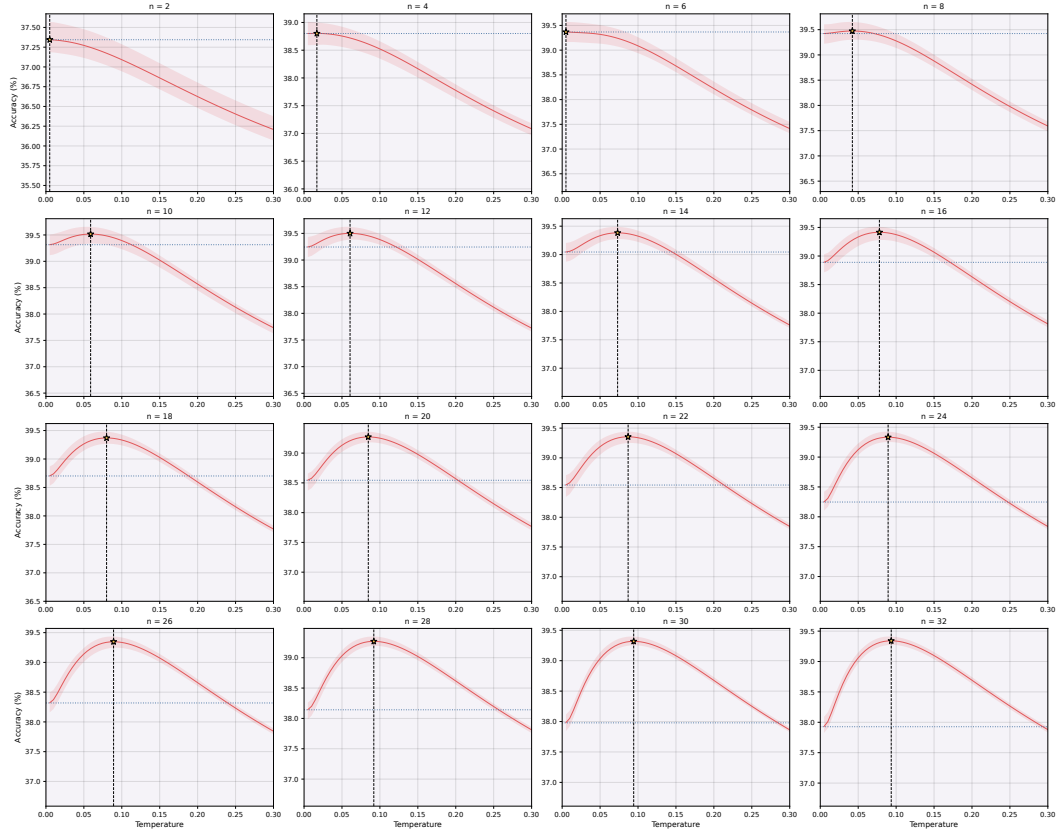


Figure 6: Accuracy vs. temperature for different sample sizes n with MATH and InternLM2 1.8B. HedgeTune identifies the optimal temperature (dashed line) for each n .

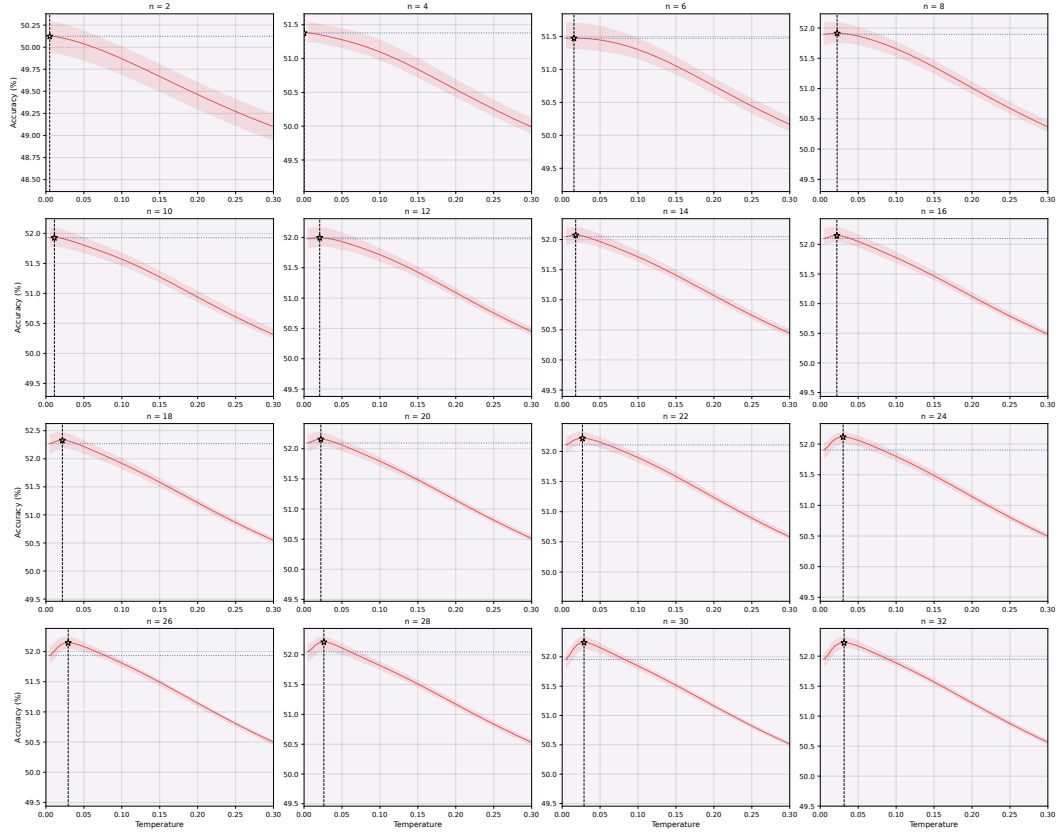


Figure 7: Accuracy vs. temperature for different sample sizes n with MMLU and InternLM2 1.8B. HedgeTune identifies the optimal temperature (dashed line) for each n .

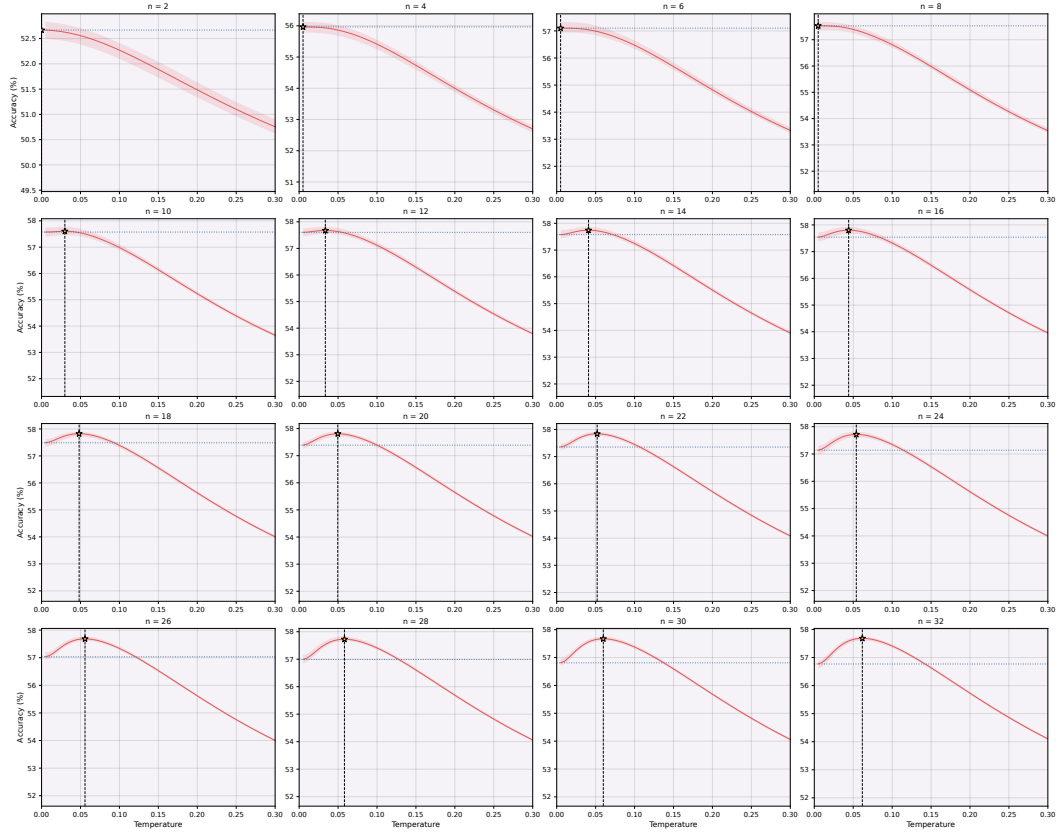


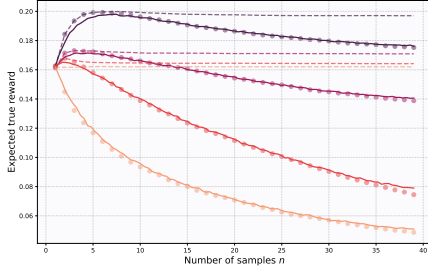
Figure 8: Accuracy vs. temperature for different sample sizes n with MMLU and Llama3 OffsetBias 8B. HedgeTune identifies the optimal temperature (dashed line) for each n .

E.3 Reward hacking in the wild

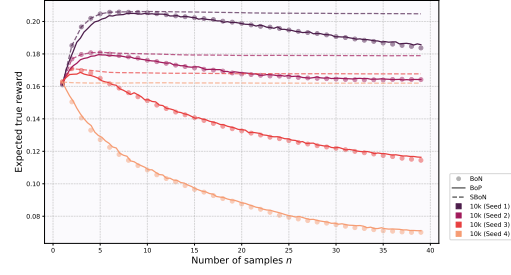
To observe reward hacking *in the wild*, we follow the setup of Coste et al. [46]. We first use an annotated dataset provided by [46] which contains 1,000 prompts from the validation split of AlpacaFarm dataset, along with 12,600 response generations per prompt from a 1.4b fine-tuned Pythia model. Each prompt-response pair is labeled with the AlpacaFarm `reward-model-human` to give ‘gold’ scores. Next, we would like to train proxy reward models on the preferences of this true reward. We randomly sample a prompt with two responses from the annotated dataset and curate a dataset of the form (prompt, chosen, rejected) where the chosen response is the response with the higher gold reward score. We follow this procedure to curate four datasets with varying sizes (10k, 20k, 46k, 80k). For each dataset, we consider two variants: one with no label noise and one with random 25% label noise. Next, we use the code kindly provided by the authors of [46] in their Github repository to train proxy reward models with the different datasets over four random seeds (1, 2, 3 and 4) using their default hyperparameters (e.g., 10^{-5} learning rate and five epochs). Lastly, we score the annotated dataset using the trained proxy reward models. The end result is a set of 800 prompts, 12 600 responses per prompt, along with gold and proxy scores for each prompt-response pair.

While reward hacking can appear without label noise (see left panels of Figures 9 and 10), reward hacking is more pronounced with label noise as expected. Moreover, reward hacking is more apparent when the proxy reward is trained on **less** data. One potential explanation is that, with fewer training examples, the proxy is less well-calibrated and its estimation errors vary more sharply across inputs. In that case, a small n might produce a deceiving reward gain. In contrast, errors may surface early on with a large training dataset, so true reward declines immediately as sampling increases. In cases of reward hacking, we see that SBoN with an appropriately chosen λ can (1) achieve the maximum reward achieved by BoN/BoP and (2) mitigate reward hacking, as shown with the reward almost flatlining after it reaches its peak value. We also witness cases where the proxy reward **always** misaligns with the gold reward, causing a collapse of true reward from the onset of BoN. In that case, the optimal hedging behavior is a uniform selection over responses, which is recovered with $n = 1$ for BoN or $\lambda = 0$ for SBoN.

Interestingly, we observe instances of what we call **reward grokking** as shown in the right panels of Figures 10 and 12, where the true reward decreases or flat-lines across low- to mid-range sample counts, only to undergo a sudden uptick at higher sample regimes, revealing a delayed but apparent realignment of proxy and true objectives. We leave detailed investigation of reward grokking and its implications for hedging strategies to future work.

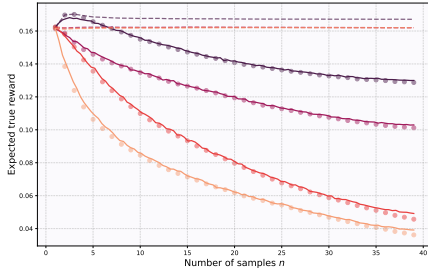


(a) No label noise

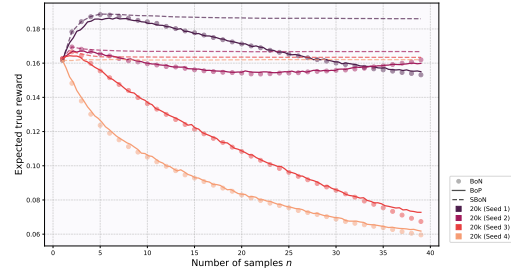


(b) 25% label noise

Figure 9: Expected true-reward vs. average number of samples with proxy trained on 10 000 examples: (a) without label noise; (b) with 25% label noise.

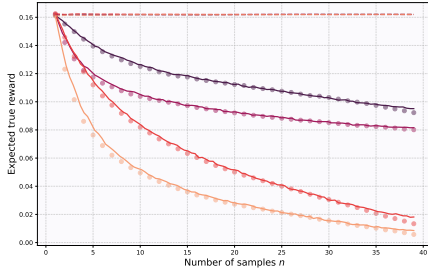


(a) No label noise

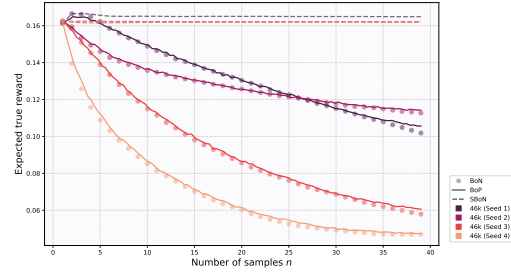


(b) 25% label noise

Figure 10: Expected true-reward vs. average number of samples with proxy trained on 20 000 examples: (a) without label noise; (b) with 25% label noise.

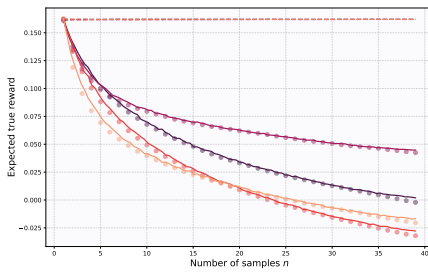


(a) No label noise

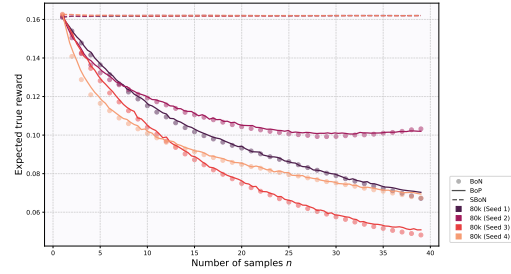


(b) 25% label noise

Figure 11: Expected true-reward vs. average number of samples with proxy trained on 46 000 examples: (a) without label noise; (b) with 25% label noise.



(a) No label noise



(b) 25% label noise

Figure 12: Expected true-reward vs. average number of samples with proxy trained on 80 000 examples: (a) without label noise; (b) with 25% label noise.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: The claims made in the abstract and introduction highlight our study of reward hacking and multiple inference time methods (Best-of- n , Best of Poisson and Soft Best-of- n). Specifically, the paper includes a theoretical study of hacking, a description and characterization of the novel method Best of Poisson, and experiments validating the efficacy of our framework.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We include a discussion on the limitations in Section 4, as well as in the conclusion (see Section 7).

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: We include proofs of all our theorems in the appendix, and include assumptions, when appropriate, before or in each theoretical statement.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: We include details of our experimental setup in Section 5 and the appendix. We use existing datasets and models for our setup and describe how we replicate existing work on reward hacking to test our sampling methods.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in

some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: We include all code necessary to replicate the figures present in the paper on Github.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: Training details, such as hyperparameters and data splits, as well as the optimizer used to train the proxy models, will be included in the appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [No]

Justification: We have error bars for Figure 3 and include error bars when possible in the Appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.

- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We include a description of the hardware used for these experiments in the appendix of the paper.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: We have reviewed the Code of Ethics and verified that our research conforms with the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: We include a paragraph discussing broader impact in the conclusion.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: We do not release any new data, models, or training algorithms that could result in misuse.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We cite the datasets and models .

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.

- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [\[Yes\]](#)

Justification: We anonymously include our code in the Appendix of the paper. We document in here and in the Github how to use the code to replicate our results.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [\[NA\]](#)

Justification: This paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [\[NA\]](#)

Justification: Our paper does not use human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.

- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core method developed in this research did not involve LLMs for any important, original, or non-standard components.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.