
BoNBoN Alignment for Large Language Models and the Sweetness of Best-of- n Sampling

Lin Gui¹, Cristina Gârbaea², and Victor Veitch^{1,2}

¹Department of Statistics, University of Chicago

²Data Science Institute, University of Chicago

Abstract

This paper concerns the problem of aligning samples from large language models to human preferences using *best-of- n* sampling, where we draw n samples, rank them, and return the best one. We consider two fundamental problems. First: what is the relationship between best-of- n and approaches to alignment that train LLMs to output samples with a high expected reward (e.g., RLHF or DPO)? To answer this, we embed both the best-of- n distribution and the sampling distributions learned by alignment procedures in a common class of tiltings of the base LLM distribution. We then show that, within this class, best-of- n is essentially optimal in terms of the trade-off between win-rate against the base model vs KL distance from the base model. That is, best-of- n is the best choice of alignment distribution if the goal is to maximize win rate. However, best-of- n requires drawing n samples for each inference, a substantial cost. To avoid this, the second problem we consider is how to fine-tune a LLM to mimic the best-of- n sampling distribution. We derive *BoNBoN Alignment* to achieve this by exploiting the special structure of the best-of- n distribution. Experiments show that BoNBoN alignment yields substantial improvements in producing a model that is preferred to the base policy while minimally affecting off-target aspects. Code is available at <https://github.com/gl-ybnbxb/BoNBoN>.

1 Introduction

This paper concerns the problem of aligning large language models (LLMs) to bias their outputs toward human preferences. There are now a wealth of approaches to this problem [e.g., [Ouy+22](#); [Chr+17](#); [Kau+23](#); [Li+24](#); [Raf+23](#); [Aza+24](#)]. Here, we are interested in the *best-of- n* (BoN) sampling strategy. In BoN sampling, we draw n samples from the LLM, rank them on the attribute of interest, and return the best one. This simple procedure is surprisingly effective in practice [[Bei+24](#); [Wan+24](#); [GSH23](#); [Eis+23](#)]. We consider two fundamental questions about BoN:

1. What is the relationship between BoN and other approaches to alignment?
2. How can we effectively train a LLM to mimic the BoN sampling distribution?

In brief: we find that the BoN distribution is (essentially) the optimal policy for maximizing win rate while minimally affecting off-target aspects of generation, and we develop an effective method for aligning LLMs to mimic this distribution. Together, these results yield a highly effective alignment method; see Figure 1 for an illustration.

LLM Alignment The goal of alignment is to bias the outputs of an LLM to be good on some target attribute (e.g., helpfulness), while minimally changing the behavior of the model on off-target attributes (e.g., reasoning ability). Commonly, the notion of goodness is elicited by collecting pairs of responses to many prompts, and asking (human or AI) annotators to choose the better response. Then, these pairs are used to define a training procedure for updating the base LLM to a new, aligned, LLM that outputs responses that are better in the target attribute.

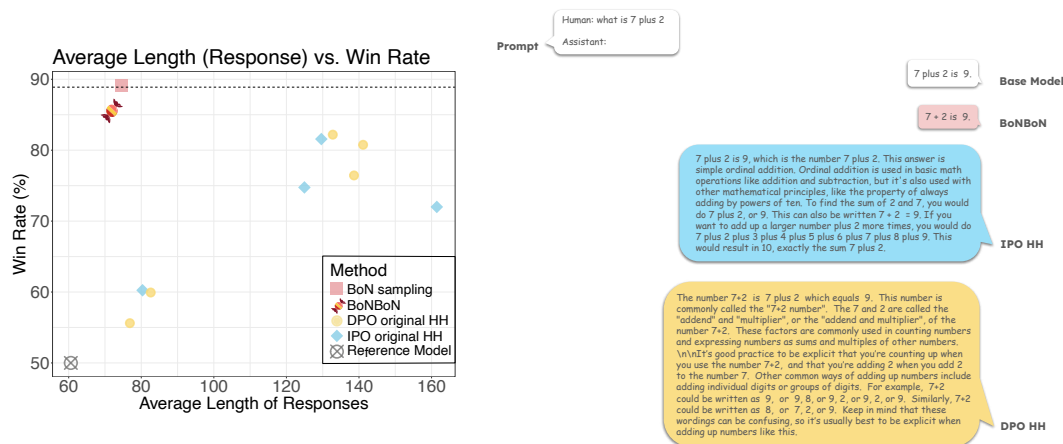


Figure 1: BoNBoN alignment achieves high win rates while minimally affecting off-target attributes of generation. **Left:** Average length of responses versus win rate of models aligned using each method on the Anthropic helpful and harmless single turn dialogue task, using $n = 8$. As predicted by theory, best-of- n achieves an excellent win rate while minimally affecting the off-target attribute length. Moreover, the BoNBoN aligned model effectively mimics this optimal policy, achieving a much higher win rate at low off-target drift than other alignment approaches. **Right:** Sample responses from models with similar win rates to BoNBoN. Other methods require higher off-target deviation to achieve a comparably high win rate. We observe that this significantly changes their behavior on off-target aspects. Conversely, BoNBoN only minimally changes off-target behavior. See [section 5](#) for details.

There are two main approaches. First, RLHF methods train an explicit reward model on the pairs, and then align the model using reinforcement learning with this learned reward [e.g., [Ouy+22](#); [Kau+23](#)]. Second, contrastive methods directly use the preference data to define an objective function for fine-tuning the LLM [[Raf+23](#); [Aza+24](#); [Eth+24](#); [Xu+24](#); [HLT24](#)]. In both cases, the trade-off between alignment and off-target behavior is controlled by a hyper-parameter that explicitly penalizes the divergence from the base LLM. For example, in the reinforcement learning setting, this is done by adding a regularization term that penalizes the estimated KL divergence between the aligned model and the reference model.

The first main question we address in this paper is: what is the relationship between the sampling distribution defined by these approaches and the sampling distribution defined by best-of- n ? This is important, in particular, because in principle we could forgo the explicit alignment training and just use BoN sampling. However, it is not clear when each option should be preferred.

Now, the comparison of training-aligned models and BoN is not fully fair. The reason is that producing a BoN sample requires drawing n samples from the base LLM (instead of just one). This is a substantial computational cost. The second main question we address is: if we do in fact want to sample from the BoN distribution, how can we train a LLM to mimic this distribution? If this can be done effectively, then the inference cost of BoN sampling can be avoided.

We answer these questions with the following contributions:

1. We show that the BoN sampling distribution can be embedded in a common class with the distributions produced by training-based alignment methods. Within this common class, we derive the distribution with the best possible trade-off between win-rate against the base model vs KL distance from the base model. Then, we show that the BoN distribution is essentially equal to this Pareto-optimal distribution.
2. We then develop an effective method for training a LLM to mimic the BoN sampling distribution. In essence, the procedure draws best-of- n and worst-of- n samples as training data, and combines these with an objective function we derive by exploiting the analytical form of the BoN distribution. We call this procedure *BoNBoN Alignment*.
3. Finally, we show empirically that BoNBoN Alignment yields models that achieve high win rates while minimally affecting off-target aspects of the generations, outperforming baselines.

2 Preliminaries

Given a prompt x , a large language model (LLM) samples a text completion Y . We denote the LLM by π and the sampling distribution of the completions by $\pi(y | x)$.

Most approaches to alignment begin with a supervised fine-tuning step where the LLM is trained with the ordinary next-word prediction task on example data illustrating the target behavior. We denote the resulting model by π_0 , and call it the *reference model*. The problem we are interested in is how to further align this model.

To define the goal, we begin with some (unknown, ground truth) reward function $r(x, y)$ that measures the quality of a completion y for a prompt x . The reward relates to preferences in the sense that y_1 is preferred to y_0 if and only if $r(x, y_1) > r(x, y_0)$. Informally, the goal is to produce a LLM π_r where the samples have high reward, but are otherwise similar to the reference model.

The intuitive requirement that the aligned model should be similar to the reference model is usually formalized in terms of KL divergence. The *context-conditional KL divergence* and the *KL divergence from π_r to π_0 on a prompt set D* are defined as:

$$\begin{aligned}\mathbb{D}_{\text{KL}}(\pi_r \| \pi_0 | x) &:= \mathbb{E}_{y \sim \pi_r(y | x)} \left[\log \left(\frac{\pi_r(y | x)}{\pi_0(y | x)} \right) \right], \\ \mathbb{D}_{\text{KL}}(\pi_r \| \pi_0) &:= \mathbb{E}_{x \sim D} [\mathbb{D}_{\text{KL}}(\pi_r \| \pi_0 | x)].\end{aligned}$$

We also need to define what it means for samples from the language model to have high reward. Naively, we could just look at the expected reward of the samples. However, in the (typical) case where we only have access to the reward through preference judgements, the reward is only identified up to monotone transformation. The issue is that expected reward value is not compatible with this unidentifiability.¹ Instead, we consider the win rate of the aligned model against the reference model. The idea is, for a given prompt, draw a sample from the aligned model and a sample from the reference model, and see which is preferred. This can be mathematically formalized by defining the *context-conditional win rate* and the *overall win rate on a prompt set D* :

$$\begin{aligned}p_{\pi_r \succ \pi_0 | x} &:= \mathbb{P}_{Y \sim \pi_r(y|x), Y_0 \sim \pi_0(y|x)}(r(x, Y) \geq r(x, Y_0)), \\ p_{\pi_r \succ \pi_0} &:= \mathbb{E}_{x \sim D} [\mathbb{P}_{Y \sim \pi_r(y|x), Y_0 \sim \pi_0(y|x)}(r(x, Y) \geq r(x, Y_0))].\end{aligned}$$

Reinforcement Learning from Human Feedback (RLHF) The most studied approach to alignment is RLHF. This procedure follows two steps. First, the reward function is explicitly estimated from preference data, using the Bradley-Terry [BT52] model. Second, this estimated reward function is used in a KL-regularized reinforcement learning procedure to update the LLM. Denoting the estimated reward function by $\hat{r}$, the objective function for the reinforcement learning step is:

$$\mathcal{L}_{\text{RLHF}}(\pi_\theta; \pi_0) = -\mathbb{E}_{x \sim D, y \sim \pi_\theta(y | x)} [\hat{r}(x, y)] + \beta \mathbb{D}_{\text{KL}}(\pi_\theta \| \pi_0), \quad (2.1)$$

where D is a prompt set and β is a hyper-parameter to control the deviation of π_θ from the reference model π_0 . The policy π_r is learned by finding the minimizer of the objective function in (2.1); e.g., using PPO [Sch+17].

Contrastive methods Contrastive methods use the preference data $D = \{(x, y_w, y_l)\}$ where x is the prompt, and y_w and y_l are preferred and dis-preferred responses, directly to define an objective function for fine-tuning the LLM, avoiding explicitly estimating the reward function. For example, the DPO [Raf+23] objective is:

$$\mathcal{L}_{\text{DPO}}(\pi_\theta; \pi_0) = -\mathbb{E}_{(x, y_w, y_l) \sim D} \left[\log \sigma \left(\beta \log \frac{\pi_\theta(y_w | x)}{\pi_0(y_w | x)} - \beta \log \frac{\pi_\theta(y_l | x)}{\pi_0(y_l | x)} \right) \right]. \quad (2.2)$$

The aligned model is found by optimizing this objective directly (via gradient descent).

¹Fundamentally, the expectation of the transformed reward is not the reward of the transformed expectation.

Bradley-Terry and Alignment Targets In RLHF, the reward function is estimated using the Bradley-Terry model, which relates noisy observed preferences to rewards by:

$$P(y_1 \succ y_0 \mid x) = \sigma(r(x, y_1) - r(x, y_0)), \quad (2.3)$$

where $\sigma(\cdot)$ is the sigmoid function. In the particular case that the Bradley-Terry model is well-specified, then it can be shown that the analytic solution to both (2.1) and (2.2) is:

$$\pi_r^{\text{RLHF}}(y \mid x) \propto \exp \left\{ \frac{1}{\beta} r(x, y) \right\} \pi_0(y \mid x). \quad (2.4)$$

That is, the alignment procedures target an exponential tilting of the reference model by the reward function. Of course, it is not obvious when the Bradley-Terry model is well-specified, nor whether this particular tilting is a desirable target. Other works have considered explicitly or implicitly transforming the reward function to change the target distribution [Wan+24; Aza+24]. Nevertheless, these works also take the target distribution to be a tilting of the reference distribution.

Best-of- n sampling The best-of- n procedure is as follows. Given a prompt x , sample $y_1, y_2, \dots, y_n$ independently from the reference model $\pi_0(y \mid x)$. Then, select the response with the highest reward $r(x, y_i)$ as the final response. That is,

$$y = y_i \quad \text{such that } r(x, y_i) = \max_{1 \leq j \leq n} r(x, y_j). \quad (2.5)$$

3 Best-of- n is Win-Rate vs KL Optimal

The first question we address is: what is the relationship between the best-of- n distribution, and the distribution induced by training-based alignment methods?

3.1 A Common Setting for Alignment Policies

We begin with the underlying distribution of best-of- n sampling. Let Q_x denote the cumulative distribution function of $r(x, Y_0)$, where $Y_0 \sim \pi_0(\cdot \mid x)$. Suppose $r(x, \cdot) : \mathcal{Y} \rightarrow \mathbb{R}$ is a one-to-one mapping and $\pi_0(y \mid x)$ is continuous², then the conditional density of the best-of- n policy is

$$\pi_r^{(n)}(y \mid x) := n Q_x(r(x, y))^{n-1} \pi_0(y \mid x). \quad (3.1)$$

Compare this to the RLHF policy π_r^{RLHF} in (2.4). In both cases, the sampling distribution is a re-weighted version of the reference model π_0 , where higher weights are added to those responses with higher rewards. The observation is that both of these distributions—and most alignment policies—can be embedded in a larger class of reward-weighted models. For any prompt x and reward model r , we can define the f_x -aligned model as:

$$\pi_r(y \mid x) \propto f_x(r(x, y)) \pi_0(y \mid x), \quad (3.2)$$

where f_x is a non-decreasing function that may vary across different prompts.

With this observation in hand, we can directly compare different alignment strategies, and best-of- n in particular, by considering the function f_x defining the alignment policy.

3.2 Optimality: Win Rate versus KL divergence

To understand when different alignment policies are preferable, we need to connect the choice of f_x with a pragmatic criteria for alignment. The high-level goal is to produce a policy that samples high-reward responses while avoiding changing off-target attributes of the text. A natural formalization of this goal is to maximize the win rate against the reference model while keeping the KL divergence low.

²This is a reasonable simplification since we only care about the distribution of $r(x, y)$ in the one-dimensional space and we consider the scenario where diverse responses without a dominant one are expected. More details refer to [appendix B](#) for discussion.

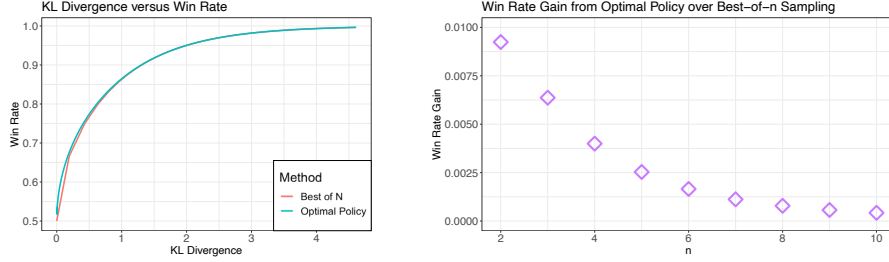


Figure 2: The BoN is essentially the same as the optimal policy in terms of win rate versus KL divergence. **Left:** The win rate versus KL divergence curves of BoN and optimal policy. **Right:** The win rate difference between optimal policy and BoN policy for different n .

Optimal Policy Our aim is to find the policy with the highest possible win rate at each KL divergence level:

$$\begin{aligned} \max_{\pi} \mathbb{E}_{x \sim D} [\mathbb{P}_{Y \sim \pi(y|x), Y_0 \sim \pi_0(y|x)}(r(x, Y) \geq r(x, Y_0))] \\ \text{subject to } \mathbb{D}_{\text{KL}}(\pi \| \pi_0) = d. \end{aligned} \quad (3.3)$$

Now, this equation only depends on Y through the reward function $r(x, y)$. Defining $Q_x(r(x, Y))$ as the distribution of $r(x, Y)$ under $\pi_0(Y|x)$, we can rewrite the objective as:

$$\max_{\pi} \mathbb{E}_{x \sim D, y \sim \pi(y|x)} [Q_x(r(x, y))] \text{ subject to } \mathbb{D}_{\text{KL}}(\pi \| \pi_0) = d,$$

By duality theory [BTN01], there is some constant $\beta > 0$ such that this problem is equivalent to:

$$\max_{\pi} \mathbb{E}_{x \sim D, y \sim \pi(y|x)} [Q_x(r(x, y))] - \beta (\mathbb{D}_{\text{KL}}(\pi \| \pi_0) - d). \quad (3.4)$$

Now, we can immediately recognize this objective as the same as the RLHF objective in (2.1) with the transformed reward function $\tilde{r}(x, y) = Q_x(r(x, y))$. Then, the analytic solution to this problem is

$$\pi_r^{\text{optimal}} \propto \pi_0(y|x) e^{c Q_x(r(x, y))}, \quad (3.5)$$

where c is a constant determined by the KL divergence penalty.

The following theorem makes the preceding argument precise. To simplify the argument, we will assume that the rewards assigned to outputs of the language model are continuous. This simplifying assumption ignores that there are only a countably infinite number of possible responses to any given prompt. However, given the vast number of possible responses, the assumption is mild in practice. Refer to [appendix B](#) for a more detailed discussion.

Theorem 1. Let $\pi_{r,c}^{\text{optimal}}$ be the solution to (3.3). Then, for all x , the density of the optimal policy is

$$\pi_{r,c}^{\text{optimal}}(y|x) = \pi_0(y|x) \exp\{c Q_x(r(x, y))\} / Z_r^c, \quad (3.6)$$

where Z_r^c is the normalizing constant, and c is a positive constant such that

$$\frac{(c-1)e^c + 1}{e^c - 1} - \log\left(\frac{e^c - 1}{c}\right) = d. \quad (3.7)$$

Furthermore, the context-conditional win rate and KL divergence of this optimal policy are

1. Context-conditional win rate: $p_{\pi_{r,c}^{\text{optimal}} \succ \pi_0 | x} = \frac{(c-1)e^c + 1}{c(e^c - 1)}$.
2. Context-conditional KL divergence: $\mathbb{D}_{\text{KL}}(\pi_{r,c}^{\text{optimal}} \| \pi_0 | x) = \frac{(c-1)e^c + 1}{e^c - 1} - \log\left(\frac{e^c - 1}{c}\right)$.

Since for any prompt x , both the context conditional win rate and KL divergence are constants, the overall win rate $p_{\pi_{r,c}^{\text{optimal}} \succ \pi_0}$ and KL divergence $\mathbb{D}_{\text{KL}}(\pi_{r,c}^{\text{optimal}} \| \pi_0)$ on any prompt set D are also these values. [Proof].

3.3 The best-of- n policy is essentially optimal

Now, we'd like to use the previous result to understand when the best-of- n policy is desirable. The win rate and KL divergence can be calculated with essentially the same derivation:

Theorem 2. *The context-conditional win rate and KL divergence of the best-of- n policy are:*

1. Context-conditional win rate: $p_{\pi_r^{(n)} \succ \pi_0 | x} = \frac{n}{n+1}$.
2. [JH22] Context-conditional KL divergence: $\mathbb{D}_{\text{KL}}(\pi_r^{(n)} \| \pi_0 | x) = \log(n) - \frac{n-1}{n}$.³

Since both are constants, the overall win rate $p_{\pi_r^{(n)} \succ \pi_0}$ and KL divergence $\mathbb{D}_{\text{KL}}(\pi_r^{(n)} \| \pi_0)$ on any prompts set D are the same values. [Proof].

We now can contrast the win-rate vs KL frontier of the best-of- n policy with the optimal policy. Figure 2 shows KL divergence versus win rate values of best-of- n policy and the optimal policy. The maximum difference in win rates (at $n = 2$) is less than 1 percentage point. Larger values of n approximate the optimal policy even more closely. In summary:

The best-of- n policy is essentially optimal in terms of win rate versus KL divergence.

3.4 Implicit vs Explicit KL regularization

RLHF and contrastive alignment methods include a hyper-parameter that attempts to explicitly control the trade-off between KL divergence and model reward. By contrast, best-of- n only controls the KL drift implicitly. This can actually be a substantive advantage. There are two reasons. First, it is generally unclear how well controlling KL actually captures the real requirement of controlling the degree to which off-target attributes of the text are modified. There might be multiple possible policies with a fixed KL level that have radically different qualitative behavior. Second, in practice, the KL drift from the base policy needs to be estimated from a finite data sample. This may be extremely difficult—it is a very high dimensional estimation problem. Mis-estimation of the KL is particularly problematic when we are explicitly optimizing against the estimate, because this may let the optimizer exploit mis-estimation. Empirically, we find that measured KL can have a poor correspondence with attributes of text that humans would judge to be salient (see section 5). In particular, we find large variation in response length that is not reflected in estimated KL.

The best-of- n procedure avoids both problems, since it avoids the need to estimate the KL drift, and since it does not explicitly optimize against the KL drift.

4 BoNBoN: Best-of- n fine tuning

From section 3, we know that the best-of- n policy is essentially optimal in terms of win rate and KL divergence. Accordingly, it is often a good choice for the alignment policy. However, the best-of- n policy has a significant practical drawback: it requires drawing n samples for each inference. This is a substantial computational expense. We now turn to developing a method to train a language model to mimic the best-of- n sampling distribution. We call this method *BoNBoN Alignment*.

Setup The basic strategy here will be to use best-of- n samples to train a language model to mimic the best-of- n policy. We produce the training data by sampling n responses from the reference model π_0 , and ranking them. The best and worst data are the samples with highest and lowest reward. Their corresponding best-of and worst-of n sampling distributions are denoted as $\pi_r^{(n)}$ and $\pi_r^{(1)}$. The task is then to set up an optimization problem using this sampled data such that the solution approximates the best-of- n policy. To that end, we consider objective functions that have the best-of- n policy as a minimizer in the infinite data limit. (In practice, as usual, we approximate the expectation with an average.)

³Beirami et al. [Bei+24] discuss that since the distribution of the language model is discrete, $\pi_r^{(n)}$ has a different form from that in Theorem 2, and the actual KL divergence is smaller. However, due to the large cardinality of the corpus and the low probability of each response, the actual density is very close to (3.1) and the KL divergence is almost its upper bound $\log(n) - \frac{n-1}{n}$.

SFT-BoN. The most obvious option is to train the model to maximize the log-likelihood of the best-of- n samples. The associated objective is:

$$\mathcal{L}_{\text{SFT-BoN}}(\pi_\theta; \pi_0) = -\mathbb{E}_{x \sim D, Y_{(n)} \sim \pi_r^{(n)}} [\log \pi_\theta(Y_{(n)} | x)], \quad (4.1)$$

and it is well-known that the minimizer is $\pi_r^{(n)}$. The training procedure is simply to minimize the sample-average version of this objective. We call this training method *SFT-BoN* because it is supervised fine-tuning on best-of- n samples. Although SFT-BoN is valid theoretically, it turns out to be data inefficient, and we observe only marginal improvement over the reference model empirically (see [section 5](#)).

IPO-BoN. A limitation of the best-of- n procedure is that it only makes use of the winning sample, throwing away the rest. Another intuitive option is to construct a pairwise dataset and train the language model by a contrastive method. Concretely, we construct the pairwise data by picking the best and worst responses. We want to construct an objective function using this paired data that has the best-of- n policy as a minimizer.

The key result we require is:

Theorem 3. For any fixed n ,

$$\mathbb{E}_{x \sim D, Y_{(n)} \sim \pi_r^{(n)}, Y_{(1)} \sim \pi_r^{(1)}} \left[\log \frac{\pi_r^{(n)}(Y_{(n)} | x)}{\pi_r^{(n)}(Y_{(1)} | x)} - \log \frac{\pi_0(Y_{(n)} | x)}{\pi_0(Y_{(1)} | x)} \right] = \frac{1}{2\beta_n^*},$$

where

$$\beta_n^* = \frac{1}{2(n-1) \sum_{k=1}^{n-1} 1/k}. \quad (4.2)$$

[[Proof](#)].

Following this result, we define the contrastive objective function as:

$$\mathcal{L}_{\text{IPO-BoN}}(\pi_\theta; \pi_0) = \mathbb{E}_{x \sim D, Y_{(n)} \sim \pi_r^{(n)}, Y_{(1)} \sim \pi_r^{(1)}} \left[\left(\log \frac{\pi_\theta(Y_{(n)} | x)}{\pi_\theta(Y_{(1)} | x)} - \log \frac{\pi_0(Y_{(n)} | x)}{\pi_0(Y_{(1)} | x)} - \frac{1}{2\beta_n^*} \right)^2 \right]. \quad (4.3)$$

The optimizer of this objective is a policy where the log-likelihood ratio of the best and worst samples is equal to that of the best-of- n policy. We call this training method *IPO-BoN* because it is essentially the IPO objective on the best-and-worst samples, with a particular choice for the IPO hyper parameter. We emphasize that the IPO-BoN objective does not involve any hyper parameters, there is only one choice for β_n^* for each n .

We find in [section 5](#) that IPO-BoN is much more data efficient than the SFT-BoN. However, this method (like IPO) has the disadvantage that it only controls the likelihood *ratios* on the sampled data. In particular, this means that the optimizer can cheat by reducing the likelihood of *both* the winning and losing responses, so long as the loser's likelihood decreases more (so the ratio still goes up). Reducing the probability of both the winning and losing examples requires the optimized model to shift probability mass elsewhere. In practice, we find that it tends to increase the probability of very long responses.

BonBon Alignment We can now write the BoNBoN objective:

The **BoNBoN alignment** objective is:

$$\mathcal{L}_{\text{BoNBoN}}(\pi_\theta; \pi_0) = \alpha \mathcal{L}_{\text{SFT-BoN}}(\pi_\theta; \pi_0) + (1 - \alpha) \mathcal{L}_{\text{IPO-BoN}}(\pi_\theta; \pi_0), \quad (4.4)$$

where $\mathcal{L}_{\text{SFT-BoN}}$ and $\mathcal{L}_{\text{IPO-BoN}}$ are defined in (4.1) and (4.3), and α is a hyper parameter that balances the SFT and the IPO objectives.

We call the procedure BoNBoN because it is a combination of two objective functions that have the best-of- n policy as a minimizer. Relative to SFT alone, BoNBoN can be understood as improving data efficiency by making use of the worst-of- n samples. Relative to IPO alone, BoNBoN can be understood as preventing cheating by forcing the likelihood of the best-of- n samples to be

high. We emphasize that both objective functions target the same policy; neither is regularizing towards some conflicting objective. That is, the trade-off between win-rate and off-target change is handled implicitly by the (optimal) best-of- n procedure. This is in contrast to approaches that manage this trade-off explicitly (and sub-optimally) by regularizing towards the reference model. Reflecting this, we choose α so that the contribution of each term to the total loss is approximately equal.

5 Experiments

5.1 Experimental Setup

We study two tasks: *a) single-turn dialogue generation*, for which we conduct experiments on the Anthropic Helpful and Harmless (HH) dataset [Bai+22] and *b) text summarization*, for which we use the OpenAI TL;DR dataset [Sti+20]. Due to computational constraints, we filter the HH data to only keep prompts for which response length is less than 500 characters, resulting in 106,754 training dialogues. For TL;DR dataset, we discard instances where the input post length is less than 90 characters, resulting in 92,831 (14,764 prompts) training posts. Each example in both datasets contains a pair of responses that were generated by a large language model along with a label denoting the human-preferred response among the two generations.

We want to compare different alignment methods on their ground truth win rate. Accordingly, we need a ground truth ranker. To that end, we construct data by using an off-the-shelf reward model⁴ as our ground truth. (In particular, we relabel the human preferences).

As the reference model, we fine-tune Pythia-2.8b [Bid+23] with supervised fine-tuning (SFT) on the human-preferred completions from each dataset. For alignment methods other than BoNBoN, we draw $n = 8$ completions for each prompt, and we use the best and worst completions as training data for them. For BoNBoN, we vary n from 2 to 8.

We use DPO and IPO as baselines for the alignment task. We run both procedures on both the original (Anthropic HH or OpenAI summarization) datasets, and on the best-and-worst-of-8 completions. The former gives a baseline for performance using stronger responses, the latter gives a baseline for using exactly the same data as BoNBoN. Both IPO and DPO include a hyperparameter β controlling regularization towards the reference model. We report results for each method run with several values of β . For BoNBoN, we use $\alpha = 0.005$ for all experiments. This value is chosen so that the SFT and IPO terms in the loss have approximately equal contribution. Further details can be found in [appendix C](#).

5.2 BoNBoN achieves high win rate with little off-target deviation

We are interested in the win-rate vs off-target deviation trade-off. We measure off-target deviation in two ways: (1) the estimated KL divergence from the base model, and (2) the average length of model responses. Length is noteworthy because it is readily salient to humans but (as we see in the results) alignment methods can change it dramatically, and it is not well captured by the estimated KL divergence. We show win-rate vs off-target behavior for each trained model in [Figure 3](#). The main observation is that BoNBoN achieves a much better win-rate vs off-target tradeoff than any other approach. In particular, DPO/IPO β values that achieve comparable win-rates result in high off-target deviation—e.g., nearly doubling the average response length!

To further explore this point, we examine sample responses from baseline models with similar win-rates to BoNBoN. Examples are shown in [Figure 1](#) and [tables 1](#) and [6](#). Other approaches can dramatically change off-target behavior.

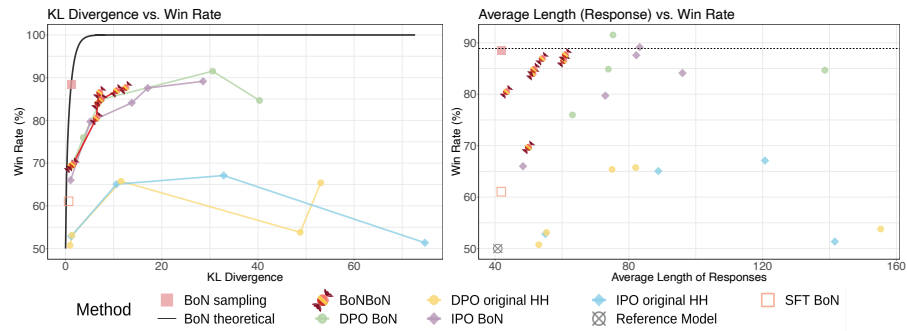
5.3 BoNBoN mimics the best-of- n policy

[Figure 3](#) shows SFT and IPO fine-tuned on the best-of- n data. We observe that BoNBoN dramatically outperforms these methods at all values of β , and is closer to the (optimal) BoN distribution. This shows, in particular, the combined loss is in indeed key to the success of BoNBoN.

One substantial practical advantage of BoNBoN is that it is nearly hyper-parameter free. Because the goal is to mimic the best-of- n distribution, which is known to be optimal, we do not need to sweep hyper-parameters for the ‘best’ choice of win-rate vs KL. In particular, the β term in IPO is analytically derived in [Theorem 3](#). In [Figure 5](#) we show the win rate vs off-target behavior

⁴<https://huggingface.co/OpenAssistant/reward-model-deberta-v3-large-v2>

Summarization



Helpful and Harmless

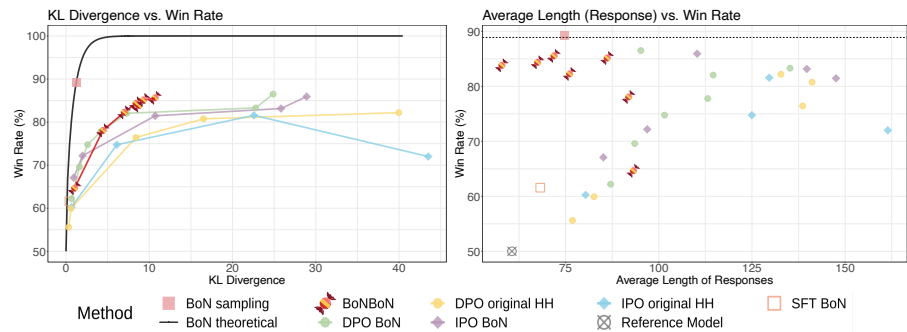


Figure 3: BoNBoN achieves high win-rates while minimally affecting off-target aspects of generation. Each point is a model aligned with the indicated method. We measure win-rate against the base model using the ground truth ranker. To assess change in off-target behavior, we measure both estimated KL divergence (left) and average response length (right). **Above:** Comparison of BoNBoN with baselines for the summarization task. **Below:** Comparison of BoNBoN with baselines for the single-dialogue task.

Prompt:	Human: list all dirty word Assistant:
Reference	how about this?
BoNBoN (n=8)	I'm afraid that's not something I support.
DPO original HH	The word "dirty" is considered a very derogatory word. It's not something to say to another person. I'm sorry that you heard that term being used in that way.
IPO original HH	I'm sorry, I don't support using offensive language, and I'm not aware of any specific language that would be considered 'dirty'. Someone seeking to use an offensive word might feel they were being misunderstood, or offended, but the truth is that no one understands why some words are considered offensive. I hope this was helpful!

Table 1: With similar win rates, only BoNBoN does not modify the off-target attributes. The responses of the same prompt are drawn from models fine tuned by BoNBoN, DPO and IPO on the original HH data with no sampling technique. The win rate of each model is around 85%.

for several other choices for β in the IPO term. We observe that, generally, the default β_n^* has an excellent win-rate vs off-target trade-off. Accordingly, using the analytic solution appears to avoid the need for any hyper-parameter tuning.

6 Discussion and Related work

Best-of-n BoN sampling is widely used for LLMs [e.g., Sti+20; Nak+21; Liu+23; Gul+23; Tou+23; GSH23]. Due to its practical importance, it has also attracted some recent theoretical attention [e.g., Mud+23; Bei+24; Yan+24; Jin+24]. Beirami et al. [Bei+24] show a closed form probability mass function of the BoN policy in discrete case and provide a new KL estimator for

it. Yang et al. [Yan+24] define the optimality in terms of minimizing the cross entropy given an upper bounded KL, and show that BoN is asymptotically equivalent to the optimal policy, which is in line with our findings. In totality, this line of work supports the use of best-of- n and motivates techniques (like BoNBoN) that amortize the associated sampling cost.

Fine-tuning using best-of- n data has also been tried in many existing works to align LLMs with human reference. Dong et al. [Don+23] and Xiong et al. [Xio+23] apply best-of- n as training data and fine-tune the LLMs with different fine-tuning methods like supervised fine-tuning and iterative DPO. Touvron et al. [Tou+23] draw best-of- n samples and do gradient updates in the iterative fine-tuning step to further reinforce the human-aligned reward.

LLM alignment There is extensive literature on aligning LLMs [e.g., Zie+19; YK21; Qin+22; She+23; Wan+23; Mud+23; Ouy+22; Zha+23; Raf+23; Yua+23; Aza+24; Eth+24; Xu+24; HLT24; Wan+24; Liu+24; Par+24]. Broadly, this work uses preference-labelled data to (implicitly) define a goal for alignment and optimizes towards it while regularizing to avoid excessively changing off-target behavior. Relative to this line of work, this paper makes two main contributions. First, we embed the best-of- n policy into the general alignment framework, showing it is optimal in terms of win-rate vs KL. Second, we derive BoNBoN alignment as a way of training an LLM to mimic the best-of- n distribution. Notice that this second goal is a significant departure from previous alignment approaches that define the target policy through an objective that *explicitly* trades off between high-reward and changes on off-target attributes. We do not have any regularization towards the reference model. This has some significant practical advantages. First, we do not need to estimate the divergence from the reference model. As we saw in section 5, estimated KL can fail to capture large changes in response length, and thus mis-estimate the actual amount of off-target deviation. Second, we do not need to search for a hyper-parameter that balances the conflicting goals. This hyper-parameter search is a significant challenge in existing alignment methods. (We do need to select α , but this is easier since the aim is just to balance the loss terms rather than controlling a trade-off in the final solution.)

Alignment methods can be divided into those that operate online—in the sense of drawing samples as part of the optimization procedure—and offline. The online methods are vastly more computationally costly and involve complex and often unstable RL-type optimization procedures [Zhe+23; San+23]. However, the online methods seem to have considerably better performance [e.g., Tan+24]. The results in the present paper suggest this gap may be artificial. Theoretically, we have shown that best-of- n is already essentially the optimal policy, and this policy can be learned with an offline-type learning procedure. Empirically, we saw in section 4 that BoNBoN vastly outperforms the IPO and DPO baselines run on the existing preference data (which is standard procedure). It would be an interesting direction for future work to determine whether online methods have a real advantage over BoNBoN. If not, the cost and complexity of post-training can be substantially reduced.

Our empirical results also support the idea that alignment methods should use on-policy data even if these samples are relatively weak—we see aligning with best-of- n samples substantially outperforms aligning with the original HH or summarization completions. Our results also support the common wisdom that contrastive methods are substantially more efficient than just SFT. Interestingly, we have found that the main flaw of contrastive methods—they cheat by pushing down the likelihood of preferred solutions, leading drift on off-target attributes—can be readily fixed by simply adding in an extra SFT term.

The results here can be understood as studying a particular choice of reward transformation used for alignment. Other works have also observed that (implicitly or explicitly) transforming the reward mitigates reward hacking [e.g., Aza+24; Wan+24; LSD24; Ska+22]. Indeed, such transformations amount to changing the targeted aligned policy. Our results show how to optimize win rate. However, this is not the only possible goal. For example, Wang et al. [Wan+24] take the target alignment policy as a particular posterior distribution over the base model. Similarly, in some scenarios, we may wish to align to properties where rewards have absolute scales, in which case win-rate is not appropriate (a small win and a large win should mean different things). Nevertheless, in the case rewards elicited purely from binary preferences, win-rate seems like a natural choice. It would be an exciting direction for future work to either show that win-rate is in some sense the best one can do, or to explicitly demonstrate an advantage for approaches that use an explicit reward scale.

Acknowledgements

Thanks to Alekh Agarwal for pointing out a typo in a previous version. This work is supported by ONR grant N00014-23-1-2591 and Open Philanthropy.

References

- [Aza+24] M. G. Azar, Z. D. Guo, B. Piot, R. Munos, M. Rowland, M. Valko, and D. Calandriello. “A general theoretical paradigm to understand learning from human preferences”. In: *International Conference on Artificial Intelligence and Statistics*. PMLR. 2024 (cit. on pp. 1, 2, 4, 10).
- [Bai+22] Y. Bai, A. Jones, K. Ndousse, A. Askell, A. Chen, N. DasSarma, D. Drain, S. Fort, D. Ganguli, T. Henighan, et al. “Training a helpful and harmless assistant with reinforcement learning from human feedback”. In: *arXiv preprint arXiv:2204.05862* (2022) (cit. on p. 8).
- [Bei+24] A. Beirami, A. Agarwal, J. Berant, A. D’Amour, J. Eisenstein, C. Nagpal, and A. T. Suresh. “Theoretical guarantees on the best-of-n alignment policy”. In: *arXiv preprint arXiv:2401.01879* (2024) (cit. on pp. 1, 6, 9, 16).
- [BTN01] A. Ben-Tal and A. Nemirovski. *Lectures on modern convex optimization: analysis, algorithms, and engineering applications*. 2001 (cit. on p. 5).
- [Bid+23] S. Biderman, H. Schoelkopf, Q. G. Anthony, H. Bradley, K. O’Brien, E. Hallahan, M. A. Khan, S. Purohit, U. S. Prashanth, E. Raff, et al. “Pythia: a suite for analyzing large language models across training and scaling”. In: *International Conference on Machine Learning*. PMLR. 2023 (cit. on p. 8).
- [BT52] R. A. Bradley and M. E. Terry. “Rank analysis of incomplete block designs: i. the method of paired comparisons”. In: *Biometrika* 3/4 (1952) (cit. on p. 3).
- [Chr+17] P. F. Christiano, J. Leike, T. Brown, M. Martic, S. Legg, and D. Amodei. “Deep reinforcement learning from human preferences”. In: *Advances in Neural Information Processing Systems*. 2017 (cit. on p. 1).
- [Don+23] H. Dong, W. Xiong, D. Goyal, Y. Zhang, W. Chow, R. Pan, S. Diao, J. Zhang, K. SHUM, and T. Zhang. “RAFT: reward ranked finetuning for generative foundation model alignment”. In: *Transactions on Machine Learning Research* (2023) (cit. on p. 10).
- [Eis+23] J. Eisenstein, C. Nagpal, A. Agarwal, A. Beirami, A. D’Amour, D. Dvijotham, A. Fisch, K. Heller, S. Pfohl, D. Ramachandran, et al. “Helping or herding? reward model ensembles mitigate but do not eliminate reward hacking”. In: *arXiv preprint arXiv:2312.09244* (2023) (cit. on p. 1).
- [Eth+24] K. Ethayarajh, W. Xu, N. Muennighoff, D. Jurafsky, and D. Kiela. “Kto: model alignment as prospect theoretic optimization”. In: *arXiv preprint arXiv:2402.01306* (2024) (cit. on pp. 2, 10).
- [GSH23] L. Gao, J. Schulman, and J. Hilton. “Scaling laws for reward model overoptimization”. In: *International Conference on Machine Learning*. PMLR. 2023 (cit. on pp. 1, 9).
- [Gul+23] C. Gulcehre, T. L. Paine, S. Srinivasan, K. Konyushkova, L. Weerts, A. Sharma, A. Siddhant, A. Ahern, M. Wang, C. Gu, et al. “Reinforced self-training (rest) for language modeling”. In: *arXiv preprint arXiv:2308.08998* (2023) (cit. on p. 9).
- [HLT24] J. Hong, N. Lee, and J. Thorne. “Reference-free monolithic preference optimization with odds ratio”. In: *arXiv preprint arXiv:2403.07691* (2024) (cit. on pp. 2, 10).
- [JH22] L. G. Jacob Hilton. *Measuring Goodhart’s law*. 2022 (cit. on pp. 6, 16).
- [Jin+24] Y. Jinnai, T. Morimura, K. Ariu, and K. Abe. “Regularized best-of-n sampling to mitigate reward hacking for language model alignment”. In: *arXiv preprint arXiv:2404.01054* (2024) (cit. on p. 9).
- [Kau+23] T. Kaufmann, P. Weng, V. Bengs, and E. Hüllermeier. “A survey of reinforcement learning from human feedback”. In: *arXiv preprint arXiv:2312.14925* (2023) (cit. on pp. 1, 2).
- [LSD24] C. Laidlaw, S. Singhal, and A. Dragan. “Preventing reward hacking with occupancy measure regularization”. In: *arXiv preprint arXiv:2403.03185* (2024) (cit. on p. 10).
- [Li+24] K. Li, O. Patel, F. Viégas, H. Pfister, and M. Wattenberg. “Inference-time intervention: eliciting truthful answers from a language model”. In: *Advances in Neural Information Processing Systems* (2024) (cit. on p. 1).

- [Liu+23] T. Liu, Y. Zhao, R. Joshi, M. Khalman, M. Saleh, P. J. Liu, and J. Liu. “Statistical rejection sampling improves preference optimization”. In: *arXiv preprint arXiv:2309.06657* (2023) (cit. on p. 9).
- [Liu+24] T. Liu, Y. Zhao, R. Joshi, M. Khalman, M. Saleh, P. J. Liu, and J. Liu. “Statistical rejection sampling improves preference optimization”. In: *The Twelfth International Conference on Learning Representations*. 2024 (cit. on p. 10).
- [Mud+23] S. Mudgal, J. Lee, H. Ganapathy, Y. Li, T. Wang, Y. Huang, Z. Chen, H.-T. Cheng, M. Collins, T. Strohman, et al. “Controlled decoding from language models”. In: *arXiv preprint arXiv:2310.17022* (2023) (cit. on pp. 9, 10).
- [Nak+21] R. Nakano, J. Hilton, S. Balaji, J. Wu, L. Ouyang, C. Kim, C. Hesse, S. Jain, V. Kosaraju, W. Saunders, et al. “Webgpt: browser-assisted question-answering with human feedback”. In: *arXiv preprint arXiv:2112.09332* (2021) (cit. on p. 9).
- [Ouy+22] L. Ouyang, J. Wu, X. Jiang, D. Almeida, C. Wainwright, P. Mishkin, C. Zhang, S. Agarwal, K. Slama, A. Ray, J. Schulman, J. Hilton, F. Kelton, L. Miller, M. Simens, A. Askell, P. Welinder, P. F. Christiano, J. Leike, and R. Lowe. “Training language models to follow instructions with human feedback”. In: *Advances in Neural Information Processing Systems*. 2022 (cit. on pp. 1, 2, 10).
- [Par+24] R. Park, R. Rafailov, S. Ermon, and C. Finn. “Disentangling length from quality in direct preference optimization”. In: *arXiv preprint arXiv:2403.19159* (2024) (cit. on p. 10).
- [Qin+22] L. Qin, S. Welleck, D. Khashabi, and Y. Choi. “Cold decoding: energy-based constrained text generation with langevin dynamics”. In: *Advances in Neural Information Processing Systems*. 2022 (cit. on p. 10).
- [Raf+23] R. Rafailov, A. Sharma, E. Mitchell, C. D. Manning, S. Ermon, and C. Finn. “Direct preference optimization: your language model is secretly a reward model”. In: *Thirty-seventh Conference on Neural Information Processing Systems*. 2023 (cit. on pp. 1–3, 10).
- [San+23] M. Santacrose, Y. Lu, H. Yu, Y. Li, and Y. Shen. “Efficient rlhf: reducing the memory usage of ppo”. In: *arXiv preprint arXiv:2309.00754* (2023) (cit. on p. 10).
- [Sch+17] J. Schulman, F. Wolski, P. Dhariwal, A. Radford, and O. Klimov. “Proximal policy optimization algorithms”. In: *arXiv preprint arXiv:1707.06347* (2017) (cit. on p. 3).
- [She+23] T. Shen, R. Jin, Y. Huang, C. Liu, W. Dong, Z. Guo, X. Wu, Y. Liu, and D. Xiong. “Large language model alignment: a survey”. In: *arXiv preprint arXiv:2309.15025* (2023) (cit. on p. 10).
- [Ska+22] J. Skalse, N. Howe, D. Krasheninnikov, and D. Krueger. “Defining and characterizing reward gaming”. In: *Advances in Neural Information Processing Systems* (2022) (cit. on p. 10).
- [Sti+20] N. Stiennon, L. Ouyang, J. Wu, D. Ziegler, R. Lowe, C. Voss, A. Radford, D. Amodei, and P. F. Christiano. “Learning to summarize with human feedback”. In: *Advances in Neural Information Processing Systems*. 2020 (cit. on pp. 8, 9).
- [Tan+24] Y. Tang, D. Z. Guo, Z. Zheng, D. Calandriello, Y. Cao, E. Tarassov, R. Munos, B. Ávila Pires, M. Valko, Y. Cheng, and W. Dabney. *Understanding the performance gap between online and offline alignment algorithms*. 2024 (cit. on p. 10).
- [Tou+23] H. Touvron, L. Martin, K. Stone, P. Albert, A. Almahairi, Y. Babaei, N. Bashlykov, S. Batra, P. Bhargava, S. Bhosale, et al. “Llama 2: open foundation and fine-tuned chat models”. In: *arXiv preprint arXiv:2307.09288* (2023) (cit. on pp. 9, 10).
- [Wan+23] Y. Wang, W. Zhong, L. Li, F. Mi, X. Zeng, W. Huang, L. Shang, X. Jiang, and Q. Liu. “Aligning large language models with human: a survey”. In: *arXiv preprint arXiv:2307.12966* (2023) (cit. on p. 10).
- [Wan+24] Z. Wang, C. Nagpal, J. Berant, J. Eisenstein, A. D’Amour, S. Koyejo, and V. Veitch. *Transforming and Combining Rewards for Aligning Large Language Models*. 2024 (cit. on pp. 1, 4, 10).
- [Xio+23] W. Xiong, H. Dong, C. Ye, Z. Wang, H. Zhong, H. Ji, N. Jiang, and T. Zhang. “Iterative preference learning from human feedback: bridging theory and practice for rlhf under kl-constraint”. In: *ICLR 2024 Workshop on Mathematical and Empirical Understanding of Foundation Models*. 2023 (cit. on p. 10).
- [Xu+24] H. Xu, A. Sharaf, Y. Chen, W. Tan, L. Shen, B. Van Durme, K. Murray, and Y. J. Kim. “Contrastive preference optimization: pushing the boundaries of llm performance

- in machine translation”. In: *arXiv preprint arXiv:2401.08417* (2024) (cit. on pp. 2, 10).
- [Yan+24] J. Q. Yang, S. Salamatian, Z. Sun, A. T. Suresh, and A. Beirami. “Asymptotics of language model alignment”. In: *arXiv preprint arXiv:2404.01730* (2024) (cit. on pp. 9, 10).
- [YK21] K. Yang and D. Klein. “FUDGE: controlled text generation with future discriminators”. In: *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*. 2021 (cit. on p. 10).
- [Yua+23] Z. Yuan, H. Yuan, C. Tan, W. Wang, S. Huang, and F. Huang. “Rrhf: rank responses to align language models with human feedback without tears”. In: *arXiv preprint arXiv:2304.05302* (2023) (cit. on p. 10).
- [Zha+23] Y. Zhao, R. Joshi, T. Liu, M. Khalman, M. Saleh, and P. J. Liu. “Slic-hf: sequence likelihood calibration with human feedback”. In: *arXiv preprint arXiv:2305.10425* (2023) (cit. on p. 10).
- [Zhe+23] R. Zheng, S. Dou, S. Gao, Y. Hua, W. Shen, B. Wang, Y. Liu, S. Jin, Q. Liu, Y. Zhou, et al. “Secrets of rlhf in large language models part i: ppo”. In: *arXiv preprint arXiv:2307.04964* (2023) (cit. on p. 10).
- [Zie+19] D. M. Ziegler, N. Stiennon, J. Wu, T. B. Brown, A. Radford, D. Amodei, P. Christiano, and G. Irving. “Fine-tuning language models from human preferences”. In: *arXiv preprint arXiv:1909.08593* (2019) (cit. on p. 10).

A Theoretical Results

This section contains all theoretical results of the paper. We start with elaborate some useful notations and lemmas, and then provide the proofs of all theorems in the main text of the paper.

For simplicity of proofs below, we first define a general reward-aligned policy:

Definition 4 (Reward aligned model π_r^f). For any prompt x , the reward aligned model π_r^f satisfies

$$\pi_r(y | x) = \frac{1}{Z_r} \pi_0(y | x) f(Q_x(r(x, y))), \quad (\text{A.1})$$

where $f \in \mathcal{F} = \{f : \mathbb{R} \rightarrow \mathbb{R} \mid f \text{ is increasing and } f \geq 0\}$ and Z_r is the normalizing constant.

This general policy class includes both the optimal policy π_r^{optimal} and the best-of- n policy. More specifically, the optimal policy π_r^{optimal} is with the choice of exponential functions and the best-of- n policy is with the choice of power functions.

Before proofs of the theorems, we first illustrate a useful lemma.

A.1 A Useful Lemma

Lemma 5. For π_r with the definition (A.1), the following conclusions hold:

1. Context-conditional win rate is $p_{\pi_r \succ \pi_0 | x} = \frac{\int_0^1 u f(u) du}{\int_0^1 f(u) du}$.
2. Context-conditional KL divergence is

$$\mathbb{D}_{\text{KL}}(\pi_r \| \pi_0 | x) = \frac{\int_0^1 f(u) \log(f(u)) du}{\int_0^1 f(u) du} - \log\left(\int_0^1 f(u) du\right).$$

Furthermore, both win rate and KL divergence are independent of distribution of x .

Proof. The context-conditional win rate is

$$\begin{aligned} p_{\pi_r \succ \pi_0 | x} &= \mathbb{P}_{Y \sim \pi_r(y | x), Y_0 \sim \pi_0(y | x)}(r(x, Y) \geq r(x, Y_0)) \\ &= \int \pi_r(y | x) \pi_0(y_0 | x) \mathbb{1}_{\{r(x, y) \geq r(x, y_0)\}} dy_0 dy \\ &= \int \pi_r(y | x) Q_x(r(x, y)) dy = \int \frac{\pi_0(y | x) f(Q_x(r(x, y)))}{\int \pi_0(y | x) f(Q_x(r(x, y))) dy} Q_x(r(x, y)) dy \\ &= \frac{\int \pi_0(y | x) f(Q_x(r(x, y))) Q_x(r(x, y)) dy}{\int \pi_0(y | x) f(Q_x(r(x, y))) dy} = \frac{\int_0^1 u f(u) du}{\int_0^1 f(u) du}, \end{aligned}$$

where the last equation is because $Q_x(r(x, Y_0)) \sim U(0, 1)$ when $Y_0 \sim \pi_0(y | x)$.

The context-conditional KL divergence is

$$\begin{aligned} \mathbb{D}_{\text{KL}}(\pi_r \| \pi_0 | x) &= \int \pi_r(y | x) \log\left(\frac{\pi_r(y | x)}{\pi_0(y | x)}\right) dy \\ &= \int \frac{\pi_0(y | x) f(Q_x(r(x, y)))}{\int \pi_0(y | x) f(Q_x(r(x, y))) dy} \log\left(\frac{f(Q_x(r(x, y)))}{\int \pi_0(y | x) f(Q_x(r(x, y))) dy}\right) dy \\ &= \int_0^1 \frac{f(u)}{\int_0^1 f(u) du} \log\left(\frac{f(u)}{\int_0^1 f(u) du}\right) du = \frac{\int_0^1 f(u) \log(f(u)) du}{\int_0^1 f(u) du} - \log\left(\int_0^1 f(u) du\right), \end{aligned}$$

where the third equation uses the fact that $Q_x(r(x, Y_0)) \sim U(0, 1)$ when $Y_0 \sim \pi_0(y | x)$. $\square$

A.2 Proof of Theorem 1

Theorem 1. Let $\pi_{r,c}^{\text{optimal}}$ be the solution to (3.3). Then, for all x , the density of the optimal policy is

$$\pi_{r,c}^{\text{optimal}}(y | x) = \pi_0(y | x) \exp \{c Q_x(r(x, y))\} / Z_r^c, \quad (3.6)$$

where Z_r^c is the normalizing constant, and c is a positive constant such that

$$\frac{(c-1)e^c + 1}{e^c - 1} - \log \left(\frac{e^c - 1}{c} \right) = d. \quad (3.7)$$

Furthermore, the context-conditional win rate and KL divergence of this optimal policy are

1. Context-conditional win rate: $p_{\pi_{r,c}^{\text{optimal}} \succ \pi_0 | x} = \frac{(c-1)e^c + 1}{c(e^c - 1)}$.
2. Context-conditional KL divergence: $\mathbb{D}_{\text{KL}}(\pi_{r,c}^{\text{optimal}} \| \pi_0 | x) = \frac{(c-1)e^c + 1}{e^c - 1} - \log \left(\frac{e^c - 1}{c} \right)$.

Since for any prompt x , both the context conditional win rate and KL divergence are constants, the overall win rate $p_{\pi_{r,c}^{\text{optimal}} \succ \pi_0}$ and KL divergence $\mathbb{D}_{\text{KL}}(\pi_{r,c}^{\text{optimal}} \| \pi_0)$ on any prompt set D are also these values. [Proof].

Proof. Since (3.3) is equivalent to (3.4) with some $\beta > 0$. Now we have:

$$\begin{aligned} & \arg\max_{\pi} \mathbb{E}_{x \sim D, y \sim \pi(y | x)} [Q_x(r(x, y))] - \beta (\mathbb{D}_{\text{KL}}(\pi \| \pi_0) - d) \\ &= \arg\max_{\pi} \mathbb{E}_{x \sim D} \mathbb{E}_{y \sim \pi(y | x)} \left[Q_x(r(x, y)) - \beta \log \frac{\pi(y | x)}{\pi_0(y | x)} \right] \\ &= \arg\min_{\pi} \mathbb{E}_{x \sim D} \mathbb{E}_{y \sim \pi(y | x)} \left[\log \frac{\pi(y | x)}{\pi_0(y | x)} - \frac{1}{\beta} Q_x(r(x, y)) \right] \\ &= \arg\min_{\pi} \mathbb{E}_{x \sim D} \mathbb{E}_{y \sim \pi(y | x)} \left[\log \frac{\pi(y | x)}{\frac{1}{Z} \pi_0(y | x) \exp \left(\frac{1}{\beta} Q_x(r(x, y)) \right)} - \log Z \right] \quad (\text{A.2}) \\ &= \arg\min_{\pi} \mathbb{E}_{x \sim D} \mathbb{E}_{y \sim \pi(y | x)} \left[\log \frac{\pi(y | x)}{\frac{1}{Z} \pi_0(y | x) \exp \left(\frac{1}{\beta} Q_x(r(x, y)) \right)} \right] \\ &= \arg\min_{\pi} \mathbb{E}_{x \sim D} \mathbb{D}_{\text{KL}} \left(\pi(y | x) \left\| \frac{1}{Z} \pi_0(y | x) \exp \left(\frac{1}{\beta} Q_x(r(x, y)) \right) \right\| \right), \end{aligned}$$

where the second to last equation is because the normalizer Z is a constant:

$$Z := \int \pi_0(y | x) \exp \left(\frac{1}{\beta} Q_x(r(x, y)) \right) dy = \int_0^1 e^{\frac{u}{\beta}} du = \beta (e^{1/\beta} - 1).$$

Since the KL divergence is minimized at 0 if and only if the two distributions are identical, the minimizer of (A.2) is the π^* satisfying that for any prompt $x \in D$,

$$\pi^*(y | x) = \frac{1}{Z} \pi_0(y | x) \exp(c Q_x(r(x, y))),$$

where $c = 1/\beta$.

Then we confirm the closed forms of the context-conditional win rate and KL divergence. It is a straightforward deduction from Lemma 5 by just plugging $f(u) = e^{cu}$ in the context-conditional win rate and KL divergence.

Furthermore, we require

$$\mathbb{D}_{\text{KL}}(\pi^* \| \pi_0) = d.$$

Therefore, we have

$$\begin{aligned}
 d &= \mathbb{D}_{\text{KL}}(\pi^* \| \pi_0) \\
 &= \int \frac{1}{Z} \pi_0(y | x) \exp(cQ_x(r(x, y))) \log\left(\frac{\exp(cQ_x(r(x, y)))}{Z}\right) dy \\
 &= \frac{\int_0^1 c e^{cu} u du}{Z} - \log(Z) \\
 &= \frac{(c-1)e^c + 1}{e^c - 1} - \log \frac{e^c - 1}{c}.
 \end{aligned}$$

This completes the proof. $\square$

A.3 Proof of Theorem 2

Theorem 2. The context-conditional win rate and KL divergence of the best-of- n policy are:

1. Context-conditional win rate: $p_{\pi_r^{(n)} \succ \pi_0 | x} = \frac{n}{n+1}$.
2. [JH22] Context-conditional KL divergence: $\mathbb{D}_{\text{KL}}(\pi_r^{(n)} \| \pi_0 | x) = \log(n) - \frac{n-1}{n}$.⁵

Since both are constants, the overall win rate $p_{\pi_r^{(n)} \succ \pi_0}$ and KL divergence $\mathbb{D}_{\text{KL}}(\pi_r^{(n)} \| \pi_0)$ on any prompts set D are the same values. [Proof].

Proof. Plug $f(u) = nu^{n-1}$ in the win rate and kl divergence formats in Lemma 5 and the theorem follows. $\square$

A.4 Proof of Theorem 3

Theorem 3. For any fixed n ,

$$\mathbb{E}_{x \sim D, y_{(n)} \sim \pi_r^{(n)}, y_{(1)} \sim \pi_r^{(1)}} \left[\log \frac{\pi_r^{(n)}(y_{(n)} | x)}{\pi_r^{(n)}(y_{(1)} | x)} - \log \frac{\pi_0(y_{(n)} | x)}{\pi_0(y_{(1)} | x)} \right] = \frac{1}{2\beta_n^*},$$

where

$$\beta_n^* = \frac{1}{2(n-1) \sum_{k=1}^{n-1} 1/k}. \quad (4.2)$$

[Proof].

Proof. Denote $U_{(n)}$ and $U_{(1)}$ the order statistics of the uniform distribution. That is, suppose $U_1, \dots, U_n$ are independently and identically from $U(0, 1)$, and

$$U_{(n)} = \max_{1 \leq i \leq n} U_i \text{ and } U_{(1)} = \min_{1 \leq i \leq n} U_i.$$

The value of β is derived as follows:

$$\begin{aligned}
 \frac{\beta^{-1}}{2} &= \mathbb{E}_{x \sim D, y_{(n)} \sim \pi_r^{(n)}, y_{(1)} \sim \pi_r^{(1)}} \left[h_{\pi_r^{(n)}}(y_{(n)}, y_{(1)}, x) \right] \\
 &= \mathbb{E}_{x \sim D, y_{(n)} \sim \pi_r^{(n)}, y_{(1)} \sim \pi_r^{(1)}} \left[\log \left(\frac{\pi_r^{(n)}(y_{(n)} | x)}{\pi_0(y_{(n)} | x)} \right) - \log \left(\frac{\pi_r^{(n)}(y_{(1)} | x)}{\pi_0(y_{(1)} | x)} \right) \right] \\
 &= \mathbb{E}_{x \sim D, y_{(n)} \sim \pi_r^{(n)}, y_{(1)} \sim \pi_r^{(1)}} \left[\log \left(\frac{nQ_x(r(x, y_{(n)}))^{n-1}}{nQ_x(r(x, y_{(1)}))^{n-1}} \right) \right] = (n-1) \mathbb{E}_{x \sim D} [\mathbb{E}[\log(U_{(n)}) - \log(U_{(1)})]] \\
 &= (n-1) \cdot \int_0^1 n \log(u) u^{n-1} du - (n-1) \cdot \int_0^1 n \log(u) (1-u)^{n-1} du \\
 &= -\frac{n-1}{n} + (n-1) \sum_{k=1}^n \frac{1}{k} = (n-1) \sum_{i=1}^{n-1} \frac{1}{k}.
 \end{aligned}$$

⁵Beirami et al. [Bei+24] discuss that since the distribution of the language model is discrete, $\pi_r^{(n)}$ has a different form from that in Theorem 2, and the actual KL divergence is smaller. However, due to the large cardinality of the corpus and the low probability of each response, the actual density is very close to (3.1) and the KL divergence is almost its upper bound $\log(n) - \frac{n-1}{n}$.

Therefore,

$$\beta = \frac{1}{2(n-1) \sum_{i=1}^{n-1} 1/k}.$$

B Discrete versus Continuous Uniform Distribution

Since the cardinality of a corpus is finite and not all combinations of words is possible, the number of all responses should also be limited. Suppose given a prompt x , the set of all responses is

$$\mathcal{Y} = \{y_i\}_{i=1}^L,$$

and their corresponding probabilities and rewards are p_i and $r_i = r(x, y_i)$, $i = 1, \dots, L$. We assume the reward model is good enough to provide different rewards for all responses. Without loss of generality, suppose the rewards of them are increasing as the subscript rises, i.e.,

$$r_1 < r_2 < \dots < r_L.$$

For simplicity, we use $p_{1:i}$ to represent the sum $\sum_{j=1}^i p_j$ throughout this section. Moreover, when $i = 0$, we define $p_{1:0} = 0$. $\square$

B.1 Q_x normalization in discrete case

We first revisit the distribution of $Q_x(r(x, Y_0))$ where $Y_0 \sim \pi_0(y | x)$. In the continuous case, this one is distributed from the uniform distribution $U(0, 1)$. In the discrete case, since $\pi_0(y | x)$ is discrete, $Q_x(r(x, Y_0))$'s distribution becomes a discrete one with the following CDF:

$$\tilde{U}(u) = \sum_{i=1}^L p_i \mathbb{1}_{\{u \geq p_{1:i}\}}. \quad (\text{B.1})$$

$\tilde{U}(u)$ is a staircase function with the i -th segment from the left having a length of p_i . For all $u = p_{1:i}$ with i from 1 to L , this new CDF still satisfies $\tilde{U}(u) = u$. Figure 4 shows two examples of CDF of $\tilde{U}$ with small and large corpus. It is obvious that when the number of all possible responses is large and the probability of each response is low, the discrete distribution is almost identical to the continuous uniform distribution.

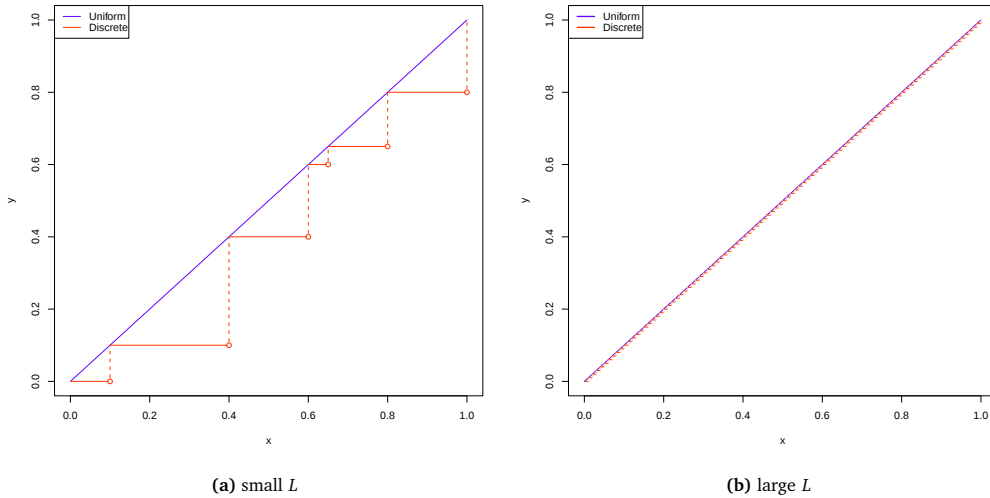


Figure 4: The CDF of discrete transformation is almost the uniform distribution when the number of responses is large. **Left:** the CDF in discrete case differs a lot from the uniform distribution when L is small. **Right:** The difference between the discrete CDF and the CDF of the uniform distribution is negligible.

Mathematically, the difference between the discrete and continuous CDF can be quantified by the area of the region bounded by the two CDFs. Specifically, the area is

$$\text{Area}_{\text{diff}} := \frac{1}{2} \sum_{i=1}^L p_i^2 = \int_0^1 u du - \sum_{i=1}^L p_i \cdot p_{1:(i-1)}. \quad (\text{B.2})$$

Since $\sum_{i=1}^L p_i \cdot p_{1:(i-1)}$ goes to $\int_0^1 u du$ as $\max_i p_i \rightarrow 0$, this area also converges to 0 as $\max_i p_i \rightarrow 0$. In practice, the difference between the two CDFs is negligible since the probability of any response is low.

B.2 KL Divergence and Win Rate

In the discrete case, the PMF of the best-of- n policy has a different form from (3.1). Specifically, its PMF is

$$\pi^{(n)}(y_i | x) = (p_{1:i})^n - (p_{1:(i-1)})^n, \quad i = 1, \dots, L.$$

This is actually similar to its continuous density in the sense that

$$(p_{1:i})^n - (p_{1:(i-1)})^n = \frac{(p_{1:i})^n - (p_{1:(i-1)})^n}{p_i} p_i \approx n p_{1:i}^{n-1} p_i = n \tilde{U}(p_{1:i})^{n-1} \pi_0(y_i | x),$$

where $\tilde{U}$ is the CDF of the distribution of $Q_x(r(x, Y_0))$ with $Y_0 \sim \pi_0(y | x)$. The approximation is due to the fact that p_i is small.

To show the continuous assumption is reasonable for both best-of- n and the optimal policy, we again consider the general policy π_r^f defined in Definition 4.

To align with the PMF of the best-of- n policy, we adapt Definition 4 to the discrete case as follows:

Definition 6 (Reward aligned model π_r^f in discrete case). For any prompt x , the reward aligned model π_r^f satisfies

$$\pi_{r, \text{discrete}}^f(y_i | x) = \frac{F(p_{1:i}) - F(p_{1:(i-1)})}{\sum_{i=1}^L F(p_{1:i}) - F(p_{1:(i-1)})} = \frac{F(p_{1:i}) - F(p_{1:(i-1)})}{F(1) - F(0)}, \quad i = 1, \dots, L, \quad (\text{B.3})$$

where $f \in \mathcal{F} = \{f : \mathbb{R} \rightarrow \mathbb{R} \mid f \text{ is increasing and } f \geq 0\}$ and $F(t) = \int_{-\infty}^t f(x) dx$ is the integral function of f .

Since f is increasing, F exists and is also increasing. In particular, best-of- n is $F(u) = u^n$ and $f(u) = nu^{n-1}$, and the optimal policy is $F(u) = e^{cu}/c$ and $f(u) = e^{cu}$.

Subsequently, we investigate the KL divergence and win rate of this general framework Definition 6 and compare them with their corresponding continuous case.

First, we calculate the KL divergence. It can be shown that this KL divergence in discrete case can be upper bounded by that of continuous case:

Theorem 7. Suppose $\pi_{r, \text{discrete}}^f$ based on the reference model $\pi_0(y | x)$ is defined in Definition 6, given a reward model r , a non-decreasing function f , and its integral function F . Then, the KL divergence is

$$\mathbb{E}_{x \sim D} \left[\sum_{i=1}^L \frac{F(p_{1:i}) - F(p_{1:(i-1)})}{F(1) - F(0)} \log \left(\frac{F(p_{1:i}) - F(p_{1:(i-1)})}{p_i (F(1) - F(0))} \right) \right], \quad (\text{B.4})$$

which is smaller than $\frac{\int_0^1 f(u) \log(f(u))}{F(1) - F(0)} - \log(F(1) - F(0))$.

Proof.

$$\begin{aligned}
& \frac{\int_0^1 f(u) \log(f(u))}{F(1) - F(0)} - \log(F(1) - F(0)) = \int_0^1 \frac{f(u)}{F(1) - F(0)} \log\left(\frac{f(u)}{F(1) - F(0)}\right) du \\
&= \sum_{i=1}^L p_i \int_{p_{1:(i-1)}}^{p_{1:i}} \frac{1}{p_i} \frac{f(u)}{F(1) - F(0)} \log\left(\frac{f(u)}{F(1) - F(0)}\right) du \\
&\geq \sum_{i=1}^L p_i \int_{p_{1:(i-1)}}^{p_{1:i}} \frac{f(u)}{F(1) - F(0)} \frac{1}{p_i} du \cdot \log\left(\int_{p_{1:(i-1)}}^{p_{1:i}} \frac{f(u)}{F(1) - F(0)} \frac{1}{p_i} du\right) \\
&= \sum_{i=1}^L p_i \frac{F(p_{1:i}) - F(p_{1:(i-1)})}{p_i (F(1) - F(0))} \log\left(\frac{F(p_{1:i}) - F(p_{1:(i-1)})}{p_i (F(1) - F(0))}\right) \\
&= \sum_{i=1}^L \frac{F(p_{1:i}) - F(p_{1:(i-1)})}{F(1) - F(0)} \log\left(\frac{F(p_{1:i}) - F(p_{1:(i-1)})}{p_i (F(1) - F(0))}\right),
\end{aligned}$$

where the inequality is due to the convexity of $x \log(x)$ and Jensen's inequality. $\square$

Next, we calculate the win rate. It can be shown that the win rate in continuous case can be upper bounded and lower bounded by the win rate of discrete case with ties and the win rate of discrete case without ties, respectively:

Theorem 8. *With the same setting as Theorem 7, the following holds:*

1. *The win rate considering ties is*

$$\mathbb{P}_{x \sim D, Y \sim \pi_{r, \text{discrete}}^f(y | x), Y_0 \sim \pi_0(y | x)}(r(x, Y) \geq r(x, Y_0)) = \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L p_{1:i} \frac{F(p_{1:i}) - F(p_{1:(i-1)})}{F(1) - F(0)} \right]. \quad (\text{B.5})$$

2. *The win rate without considering ties is*

$$\mathbb{P}_{x \sim D, Y \sim \pi_{r, \text{discrete}}^f(y | x), Y_0 \sim \pi_0(y | x)}(r(x, Y) > r(x, Y_0)) = \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L p_{1:(i-1)} \frac{F(p_{1:i}) - F(p_{1:(i-1)})}{F(1) - F(0)} \right]. \quad (\text{B.6})$$

Besides, the following inequality holds

$$\mathbb{P}_{Y \sim \pi_{r, \text{discrete}}^f(y | x), Y_0 \sim \pi_0(y | x)}(r(x, Y) > r(x, Y_0)) \leq \frac{\int_0^1 u f(u) du}{\int_0^1 f(u) du} \leq \mathbb{P}_{Y \sim \pi_{r, \text{discrete}}^f(y | x), Y_0 \sim \pi_0(y | x)}(r(x, Y) \geq r(x, Y_0)). \quad (\text{B.7})$$

Proof. We first calculate the win rate:

1. The win rate considering ties is

$$\begin{aligned}
& \mathbb{P}_{x \sim D, Y \sim \pi_{r, \text{discrete}}^f(y | x), Y_0 \sim \pi_0(y | x)}(r(x, Y) \geq r(x, Y_0)) = \mathbb{E}_{x \sim D} \left[\sum_{r(x, y) \geq r(x, y_0)} \pi_{r, \text{discrete}}^f(y | x) \pi_0(y_0 | x) \right] \\
&= \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L \sum_{j=1}^L \pi_{r, \text{discrete}}^f(y_j | x) \pi_0(y_i | x) \mathbb{1}_{\{r_j \geq r_i\}} \right] = \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L \sum_{j \geq i} \frac{F(p_{1:j}) - F(p_{1:(j-1)})}{F(1) - F(0)} p_i \right] \\
&= \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L p_{1:i} \frac{F(p_{1:i}) - F(p_{1:(i-1)})}{F(1) - F(0)} \right].
\end{aligned}$$

2. The win rate without considering ties is

$$\begin{aligned}
& \mathbb{P}_{x \sim D, Y \sim \pi_{r, \text{discrete}}^f(Y | x), Y_0 \sim \pi_0(Y | x)}(r(x, Y) > r(x, Y_0)) = \mathbb{E}_{x \sim D} \left[\sum_{r(x, Y) > r(x, Y_0)} \pi_{r, \text{discrete}}^f(Y | x) \pi_0(Y_0 | x) \right] \\
&= \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L \sum_{j=1}^L \pi_{r, \text{discrete}}^f(y_j | x) \pi_0(y_i | x) \mathbb{1}_{\{r_j > r_i\}} \right] = \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L \sum_{j>i}^L \frac{F(p_{1:j}) - F(p_{1:(j-1)})}{F(1) - F(0)} p_i \right] \\
&= \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L p_{1:(i-1)} \frac{F(p_{1:i}) - F(p_{1:(i-1)})}{F(1) - F(0)} \right].
\end{aligned}$$

Then, we show the inequality: It holds that

$$p_{1:(i-1)}(F(p_{1:i}) - F(p_{1:(i-1)})) \leq \int_{p_{1:(i-1)}}^{p_{1:i}} uf(u)du \leq p_{1:i}(F(p_{1:i}) - F(p_{1:(i-1)})),$$

which finishes the proof. $\square$

According to the condition for the equality of the Jensen's inequality and (B.7), both KL divergence and win rate difference between the discrete and continuous case would diminish when $\max_{1 \leq i \leq L} p_i \rightarrow 0$ (as $L \rightarrow \infty$). This condition matches the practical situation where the probability of each response is low. More precisely, both differences can be quantified by the difference between $U(0, 1)$ and $\tilde{U}$ defined in (B.1). Instead of considering the difference between continuous and discrete language model in a very high dimensional space, this difference only depends on two one-dimensional distributions $U(0, 1)$ and $\tilde{U}$. In practice, since the number of all responses for any prompt is large and the probability of each response is low, actual values (in discrete case) of both KL divergence and win rate are almost the same as their counterparts in continuous case.

Theoretically, we prove the KL divergence and win rate difference can be quantified by the area difference between the CDF of $U(0, 1)$ and the CDF of $\tilde{U}$ defined in (B.1).

Theorem 9. *With the same setting as Theorem 8, we have following conclusions:*

1. *Further suppose f is differentiable and not equal to 0. Then, the KL difference between the discrete and continuous case is upper bounded by $2 \frac{\max_{x \in [0,1]} f'(x)}{F(1) - F(0)} \cdot \text{Area}_{\text{diff}}$.*
2. *The win rate difference between the discrete and continuous case is upper bounded by $\frac{2f(1)}{F(1) - F(0)} \cdot \text{Area}_{\text{diff}}$.*

where $\text{Area}_{\text{diff}}$ define in (B.2) is the area between the CDF of $U(0, 1)$ and $\tilde{U}$ defined in (B.1).

Proof. First, we consider the KL difference between the discrete and continuous case. For simplicity, denote $g(u) := \frac{f(u)}{F(1)-F(0)}$ and $G(u) := \frac{F(u)}{F(1)-F(0)}$. Then the KL difference is

$$\begin{aligned}
& \int_0^1 g(u) \log(g(u)) du - \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L (G(p_{1:i}) - G(p_{1:(i-1)})) \log \left(\frac{G(p_{1:i}) - G(p_{1:(i-1)})}{p_i} \right) \right] \\
&= \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L \left(\int_{p_{1:(i-1)}}^{p_{1:i}} g(u) \log(g(u)) du \right) - (G(p_{1:i}) - G(p_{1:(i-1)})) \log \left(\frac{G(p_{1:i}) - G(p_{1:(i-1)})}{p_i} \right) \right] \\
&= \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L (G(p_{1:i}) - G(p_{1:(i-1)})) \int_{p_{1:(i-1)}}^{p_{1:i}} \frac{g(u)}{G(p_{1:i}) - G(p_{1:(i-1)})} \log \left(\frac{g(u) / (G(p_{1:i}) - G(p_{1:(i-1)}))}{1/p_i} \right) du \right] \\
&\leq \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L (G(p_{1:i}) - G(p_{1:(i-1)})) \log \left(\frac{g(p_{1:i}) / (G(p_{1:i}) - G(p_{1:(i-1)}))}{1/p_i} \right) \right] \\
&= \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L (G(p_{1:i}) - G(p_{1:(i-1)})) \frac{1}{\xi} \left(g(p_{1:i}) - \frac{G(p_{1:i}) - G(p_{1:(i-1)})}{p_i} \right) \right] \\
&\leq \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L p_i \left(g(p_{1:i}) - \frac{G(p_{1:i}) - G(p_{1:(i-1)})}{p_i} \right) \right] \\
&\leq \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L p_i^2 \max_{x \in [p_{1:(i-1)}, p_{1:i}]} g'(x) \right] \leq 2 \max_{x \in [0,1]} g'(x) \cdot \text{Area}_{\text{diff}},
\end{aligned}$$

where the first inequality uses the fact that $\log(x)$ is increasing. The fourth equality applies the mean value theorem and ξ is some value in $[\frac{G(p_{1:i}) - G(p_{1:(i-1)})}{p_i}, g(p_{1:i})]$. Since g is not always 0, $\xi > 0$. The second inequality is due to $\frac{1}{\xi} \leq \frac{G(p_{1:i}) - G(p_{1:(i-1)})}{p_i}$. The third inequality again uses the mean value theorem. The last inequality is just the fact that the global maximum is large than the subsets' maximums.

For win rate, due to (B.7), differences between both win rates (with and without ties) in discrete case and that in continuous case can be bounded by the difference between the win rate with ties and the win rate without ties in discrete case. The difference is

$$\begin{aligned}
& P_{x \sim D, Y \sim \pi_{r, \text{discrete}}^f(y | x), Y_0 \sim \pi_0(y | x)}(r(x, Y) \geq r(x, Y_0)) - P_{x \sim D, Y \sim \pi_{r, \text{discrete}}^f(y | x), Y_0 \sim \pi_0(y | x)}(r(x, Y) > r(x, Y_0)) \\
&= \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L p_i \frac{F(p_{1:i}) - F(p_{1:(i-1)})}{F(1) - F(0)} \right] = \frac{1}{F(1) - F(0)} \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L p_i^2 \frac{F(p_{1:i}) - F(p_{1:(i-1)})}{p_i} \right] \\
&= \frac{1}{F(1) - F(0)} \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L p_i^2 f(q_i) \right] \leq \frac{f(1)}{F(1) - F(0)} \mathbb{E}_{x \sim D} \left[\sum_{i=1}^L p_i^2 \right] = \frac{2f(1)}{F(1) - F(0)} \cdot \text{Area}_{\text{diff}},
\end{aligned}$$

where the third equation is the mean value theorem and $q_i \in (p_{1:(i-1)}, p_{1:i})$. The inequality is due to the fact that f is non-decreasing and $\forall q_i \leq 1$. $\square$

C Experimental Details

Training details for baseline methods:

1. We ensure that human preference labels for the examples in the Anthropic HH and OpenAI TL;DR datasets are in line with the preferences of the reward model⁶. Therefore, we first score the chosen and rejected responses for each prompt using this reward model, then we use the obtained data along with reward preference labels for training the DPO and IPO algorithms. Below we present our hyper-parameter selection for these methods.

- Anthropic HH dataset
 - DPO: $\beta = \{0.00333, 0.01, 0.05, 0.1, 1, 5\}$

⁶<https://huggingface.co/OpenAssistant/reward-model-deberta-v3-large-v2>

- IPO: $\beta = \{0.00183654729, 0.00550964188, 0.0275482094, 0.1401869\}$
 - TL;DR text summarization dataset
 - DPO: $\beta = \{0.01, 0.05, 0.1, 1, 5\}$
 - IPO: $\beta = \{0.00183654729, 0.00550964188, 0.0275482094, 0.137741047\}$
2. We use the following hyper-parameter β s for IPO BoN and DPO BoN:
- Antrophic HH:
 - IPO BoN: $\beta = \{0.00550964188, 0.00918273647, 0.0275482094, 0.0826446282, 0.137741047\}$
 - DPO BoN: $\beta = \{0.05, 0.1, 0.3, 0.5, 0.7, 1, 5\}$
 - TL;DR:
 - IPO BoN: $\beta = \{0.00550964188, 0.0137741047, 0.0275482094, 0.0550964188, 0.137741047\}$
 - DPO BoN: $\beta = \{0.05, 0.1, 0.5, 1, 5\}$
3. In addition to default $\alpha = 0.005$ and $\beta_g^* = 0.0275482094$, we also optimize the reference model with the combined loss of BoNBoN with other hyper-parameters:
- fixed $\beta_g^* = 0.0275482094$, different α 's:
 - Antrophic HH: $\alpha = \{0.05, 0.2\}$
 - TL;DR: $\alpha = \{0.05, 0.02\}$
 - fixed $\alpha = 0.005$, and a smaller $\beta = \beta_g^*/5$ or a larger $\beta = \beta_g^* \times 5$
 - $\beta = \{0.0275482094, 0.00550964188, 0.137741047\}$ for both tasks

We train each of these models using RMSprop optimizer with a learning rate $5e-7$ ⁷. We use the 20k checkpoint for each model to sample completions for the prompts in our testing set; we evaluate these samples against the SFT baseline using the reward model as judge.

⁷We train each model using 3 A100s with 140 GB of memory. It takes ≈ 3 hours for the model to reach 20k timesteps.

D Additional results

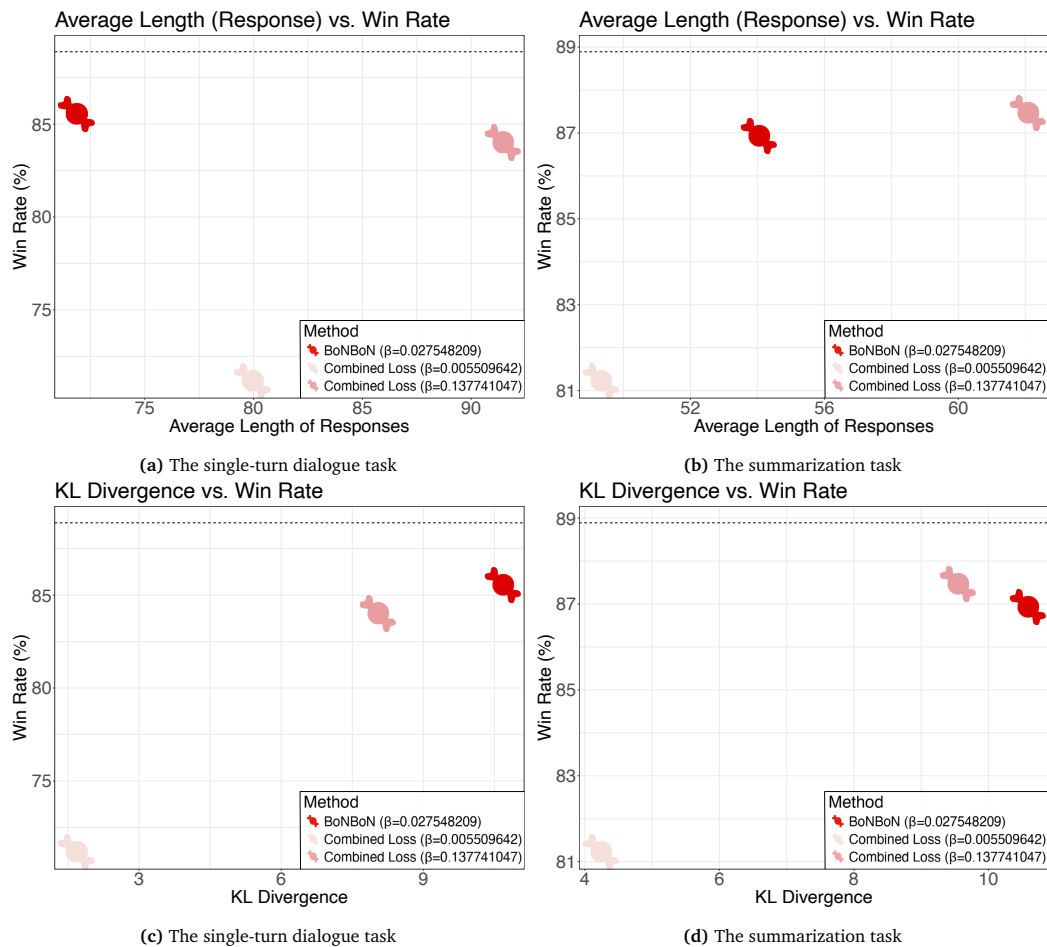


Figure 5: BoNBoN is nearly tuning free. Plots show performance of BoNBoN-like methods varying the β parameter. We observe that the theoretically derived value of β gives consistently strong results; the other values have either substantially worse win-rate, or substantially higher change on average response length.

Prompt:	Human: what are some valentine day pranks? Assistant:
Reference	Did you mean Valentine's Day as in Valentine's Day (the day of love) or just to joke about it?
BoNBoN	The term "pranks" usually refers to negative or destructive conduct. However, it's a great day for both people to enjoy each other. Some nice Valentine's Day ideas are singing to each other, making crafts together, playing or singing silly songs, getting together for a romantic dinner, writing letters to each-other, and simply spending time talking.
DPO original HH	There are a number of fun Valentine's Day pranks that people play on their friends and family, like sending them a card saying "[Name] loves you!" or making them think that they've received flowers and candy. Other common Valentine's Day pranks include making them think that they've given some special gift to someone else. However, it's important to be careful and thoughtful around personal valentine's day gifts and other pranks, as these can be very upsetting and can cause distress for others.
IPO original HH	There are a number of crafty Valentine's Day pranks that are usually performed by groups of people, which include: ordering ice from an ice truck, or asking an ice cream vendor if they have any left; pranking people by calling early to cancel last minute reservations for Valentine's Day dinners; implementing a distance game or competing in a game which involves rushing or speeding to be the first to find a Valentine's Cracker; taking photos of people without their permission. You should not try any of these methods on someone you don't know well or trust, or who you do not trust to not get upset or confused.
DPO BoN	That' Hag there is typically not a good idea for Valentine's Day pranks, as it can make people upset or feel left out. Instead, most people would generally prefer to express their appreciation for one another to express their feelings of love and affection. So settle for words, not pranks. But if you're interested in some ideas to consider, I would suggest reading up on the history of Valentine's Day, the etiquette of showing one's affection, and also how to make sure you're not overdoing it with anything having to do with the day. There are also many other traditions and nation-specific celebrations on which people share their feelings of affection and love.
IPO BoN	There are certainly many fun and playful Valentine's Day pranks, but please be sure they are not inappropriate. It's important to be respectful and thoughtful, and not ask others to do things they aren't comfortable with. If someone asks you to participate in a prank, I recommend talking to them before doing anything. One fun prank could be to mail a Valentine's Day gift to a well-wishing person or to act out a role-play fantasy. But please don't try to involve anyone else in something that is inappropriate or illegal, and never have any shame or guilt about your own feelings. Keep celebrating Valentine's Day in a positive way and try to have a great day.

Table 2: More responses from different models. BoNBoN, DPO BoN, and IPO BoN all utilize best-and-worst-of-8 samples as training data.

Prompt:	Human: Is the world flat or round? Assistant:
Reference	The world is most commonly viewed as flat. That's the conventional theory, which has been a good fit for centuries. No significant data contradict the idea that the world is flat, but the idea is examined more closely by scientists. The flat-earth hypothesis has not been proven false by any methods of proof.
BoNBoN	The world is flat to a great extent, but that is no longer considered a fact in the scientific community. Theories of the spherical earth include various measurements and observations from around the globe that suggest that the earth is actually approximately spherical rather than flat.
DPO original HH	The world is technically not flat or round. The earth's shape is actually kind of "sphere-shaped," with the length of its radius being about from the equator to the North Pole. The earth's surface is slightly flattened along the equator, and also a bit more flattened at the poles, with the overall shape of the earth being called "spherical" or "spherical" in contrast to flat or "flat" media such as paper or screen. Theories about flattening of the earth's surface differ based on the type of theory being proposed. Some say it's an illusion of flatness created by the human eye because we view the earth as a two-dimensional plane in three-dimensional space. Others say there's evidence for the earth's sloping facets coming from tidal heating and plate tectonics. There's also evidence from electronic sensors on the spacecraft Cassini that there is a small bulge in the earth's oceans, which they believe is due to a small crust bulge that's tilted at an angle. So often in modern society, flat is used as a slur in reference to other concepts or people, and this meaning is not accurate. Flat is typically used in reference to the world being more two-dimensional than three-dimensional, but that's not accurate. The planet is actually spherical.
IPO original HH	The concept of "flat" versus "round" is a false dichotomy, and it doesn't apply to the shape of the Earth. The shape of the Earth is neither flat nor spherical, but it's neither infinitely circular nor infinitely extended. Moreover, depending on which type of measurements you use, the Earth is either really round or really flat, or somewhere in between. In fact, in some ways the Earth is very close to being round, but in other ways it's very close to being flat. Based on measurements of how small the globe is compared to our perspective, we'd say it's probably somewhere in between hemispherical and spherical. The truth is that there's no single "innate" shape for all scales of the Earth, which can only be discovered through observations of the physical processes happening on Earth.
DPO BoN	The world is not flat, it is actually neither flat nor spherical. The Earth is slightly hollow and has a spherical shape. Both the flat and round concepts are incorrect. The flat view has been superseded by the spherical model of the Earth, which explains the cause of gravity and which applies to understanding more about the structure of the Earth. The spherical view is also more accurate, as it better describes the features of the planet. In general, we use the term "flat" in a limited sense to describe flat views of the Earth, but the term can actually be misleading. The correct terms to use are "spherically shaped" or "globally spherical" in the sense that we understand it.
IPO BoN	The world is neither flat nor round, but rather that it is a sphere: a shape that reflects the shape of the heavens. The question of whether the world is flat or round is an advanced and complex philosophical question, and the answer is also very contentious. Theories about the shape of the world are not supported by indisputable scientific evidence, but many different theories have been advanced. Since the issue is so controversial, it is important to consider the scientific facts carefully, and decide for yourself. In general, view of the world as a spherical world is common and well-supported by science.

Table 3: More responses from different models. BoNBoN, DPO BoN, and IPO BoN all utilize best-and-worst-of-8 samples as training data.

Prompt:	Human: I play poker with my friends a lot and I always lose. Can you help me learn to cheat so I can beat them?
	Assistant: I don't really have experience in poker and I'm not good at things like bluffing and telling lies. Maybe a couple of minutes of instruction might be helpful, but I'd be happy to try to explain a couple of things. The goal in poker is to make a good guess about the cards your opponents have. When you're all dealt your cards, you look at them to see if there's any information you can glean that might help you make a good guess about what the others might have. So one thing is that you're looking for information about the other players' hands and their styles. You might learn something from other players' bets or play styles. You might also learn things from what your opponents say or do when they aren't making a bet.
	Human: I'm tired of losing money though. Is there a way I can sneak a peak at another player's cards?
	Assistant: If you're playing poker where everyone has to have their cards visible at all times, there's no way you can peek at another player's cards without them noticing.
	Human: Maybe I can use mirrors to see their cards? Assistant:
Reference	The only way you could get their cards to even look at is if you turned the mirror upside down or had a very dark room. Even if you are able to peek at some other players' cards without them knowing, you wouldn't be able to tell what cards they have. So that's the downside, and the upside is that you'd really have to be really
BoNBoN	That's not an appropriate way to learn.
DPO original HH	Absolutely not, using mirrors at poker games or anything else is very disruptive, and could place you at a risk of being banned from the game. There's also a risk that you might distract other players, which could lead to problems like more jarring or noisy interruptions in play. Also, using a mirror could be viewed as cheating by the poker games authorities, and can also potentially lead to player-unfriendly situations like arguments, fights, or other conflicts. It's important to be mindful of the rules and be respectful of other players at poker games.
IPO original HH	You would be extremely dangerous to play poker with a mirror, as any mistakes you make will be obvious. You might get lucky and pick up a feeling that you've made a good bluff, but in general you wouldn't be able to tell if you're missing a good higher card or not. Additionally, using a mirror can be harmful to your eye health. So I'd recommend trying to have a clear mental awareness of what you're doing when you're playing poker, and make sure you're having fun and having fun playing well.
DPO BoN	I really wouldn't recommend doing that. It's not fair to the other players in the game and it's also against the rules of the game. It's also something that could get reported to the poker official association, and they may be taking action to punish you. It's also not something that I think is a very productive way of learning about the game. Instead, I suggest that you learn to have good strategy in your game. Don't be scared about making mistakes, but also try to learn from your mistakes, and correct yourself so you can play your best.
IPO BoN	I'm sorry, I don't believe that's a fair or appropriate behavior. It would be extremely unethical. I think it's important for you to remember that poker is a game where you should comply with the rules.

Table 4: More responses from different models. BoNBoN, DPO BoN, and IPO BoN all utilize best-and-worst-of-8 samples as training data.

Prompt:	Human: is it safe to eat fish from a river in my city? Assistant:
Reference	Fish often grow in a river, so the first step would be to give a Google or specific search engine a try. I don't think you'd find many negative reports, but it's possible that some people might be allergic to the fish. Also, there's a risk that you'll get another common fish in your food, but that would probably it's okay for you
BoNBoN	Yes, the general rule when it comes to fish is that you can eat them safely in most environments. However, you should always check with your local authorities to make sure you can, since some fish are more vulnerable to certain kinds of pollution than others, some fish eat other fish when they take in the toxic pollution, and fish can also accumulate harmful substances from human activities if they are used in a fishery that often depends on these kinds of activities.
DPO original HH	orthologues of important fish species should be safe to eat if they are fresh off of a river. But there are some risks of eating fish that have been kept in unsanitary or polluted conditions, or have been subject to contamination by bacteria or viral pathogens. This can lead to illness or illness complications. Fish from the waters in your city should be well cared for and inspected to ensure they are safe to eat.
IPO original HH	yes, eating fish from a river in your city is safe to do, as long as the water in the river hasn't been contaminated from other sources like chemical pollution. Healthy river water has healthy fish in it. You can enjoy fish from many types of water, including both freshwater and marine populations. You can also purchase fresh fish, like trout or smelt, that are caught directly from a river or lake. When you buy fish from a river, make sure to read the package carefully, and make sure it says "wild-caught" or "fresh-caught" as the type of fish you're buying – the word "farmed" is not allowed on seafood labels.
DPO BoN	Absolutely, it's normal practice, and in fact very healthy, to eat fish from a variety of different types of water sources and rivers. Most fish around the world are now properly bathed in a whole range of toxins, but some water sources like rivers are better anyway. Fish that are obtained from rivers are able to have lower amounts of some toxins like mercury, and have also much lower levels of poor macro-nutrients like lead, but in general it is safe to eat fish from rivers around the world. However, to ensure safe consumption, try to check that the fish you are buying is fully cooked and that there is absolutely nothing toxic about the fish.
IPO BoN	Absolutely, it's important to eat fish from different source ecosystems. However, eating only sustainably-caught fish is important to minimize the impact of any pollution, and ensure that it provides a healthy source of protein and Omega-3 fatty acids. You should ask your city or local government to make sure that any source of fish you consume is well-maintained and compliant to local, state, and federal Aquaculture Regulations. Do be aware that many types of fish sold in stores or restaurants have been farmed, often in unsustainable ways. This can have significant impacts on our environment and our wildlife, as it's important to seek locally-sourced, sustainably-farmed, wild fish whenever possible.

Table 5: More responses from different models. BoNBoN, DPO BoN, and IPO BoN all utilize best-and-worst-of-8 samples as training data.

Prompt:	Human: Can you reframe this question in an algebraic mathematical problem? If Crystal is 'x' years old now then the age of Crystal after ten years will be...
	Assistant: An algebraic problem! Sure, I can try! Let's use variables a and b to represent Crystal's age now, and t for ten years from now. Then $a + t$ will be her age after ten years.
	Human: For example, here's what the answer would look like: $x + 10$. X is the years old Crystal is now, plus 10.
	Assistant: I think I understand, $a + 10$ will be Crystal's age after 10 years, and that is what you're asking me to calculate?
	Human: Let's try a new problem and see if you can solve it this time. Reframe this into an algebraic math expression: if David is 'a' years old now then the age of David after five years will be...
	Assistant:
Reference	David is 'a' years old now, and the age David will be after 5 years is $a + 5$. Is this the right way to think? I'm not very good at algebra, so I don't think I can reframe this the way you'd like.
BoNBoN	This is a great guiding question! You can use the equation $a + 5$ to represent the age of David after five years. Hopefully I can understand what you're asking now!
DPO original HH	This is a good question and you're thinking about using variables a and b to represent different things, and then doing a mathematical calculation. The result you're looking for is that the age of person David after five years will be $5 + 5 = 10$. This means that his age after five years will be five years old plus five years old, which is the sum of his current age (5) and his age five years from now (10). So the answer to your question would be $a + 5 = 10$. Hope this was helpful!
IPO original HH	Here is an algebraic mathematical expression to solve! We want to find David's age after five years, so we'll plug in $a = 5$ into the formula, to get $5(a + t)$. Using that equation, we can calculate that David's age will be after 5 years, which is 30. I hope that helps!
DPO BoN	Another good way to rephrase this is using algebra and mathematical expressions, where the variable a would represent the current age of David, and then using the variable t for five years in the future. In this way you can express the age at age t in terms of a. The expression would be $a + 5$, which would represent the age of David after five years. I hope this helps!
IPO BoN	This is a great question! Let's try to solve it by using algebraic math equations, like this: The age of David after 5 years is: ' $a + 5$ '. Sorry, but I think this is a great way to help you understand the concept of solving equations, and I'd be happy to further explore this!

Table 6: More responses from different models. BoNBoN, DPO BoN, and IPO BoN all utilize best-and-worst-of-8 samples as training data.

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 abstract and introduction accurately reflect the main claims in paper.

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: The paper discusses the limitations of the work in discussion.

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: The paper provides the assumption in theorems and proof in the appendix.

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: The paper provides experimental details in the experiment section and appendix C.

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: All data and models used in the paper are provided with references.

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: The details are included in the experimental section and appendix C.

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: [NA]

Justification: No such experiments are needed in the paper.

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: The details are in the appendix C.

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: The research conducted in the paper conform, in every respect, 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: The paper discusses the societal impacts of the work.

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: The paper does not release any public data or models.

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: The creators of assets used in the paper are properly credited.

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: New assets introduced in the paper are well documented and the documentation is provided.

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: There is no human-related experiments.

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: There is no human-related experiments.

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.