

Synthetic Preference Interpolation for Language Model Alignment

Anonymous ACL submission

Abstract

Ensuring alignment with human preferences is a critical and challenging aspect of large language models (LLMs). Currently, the most widely adopted alignment methods, such as those based on Direct Preference Optimization (DPO), leverage pairwise preference data for training and have demonstrated promising results. However, these methods face limitations, as they cannot fully exploit the rich information inherent in preference data, such as intermediate quality levels between chosen and rejected samples. Motivated by this insight, we propose **Synthetic Preference Interpolation Alignment (SPIA)**, a novel alignment algorithm that introduces interpolated synthetic preferences to better capture the nuances between samples of different quality levels. By constructing synthetic preference data that reflects intermediate quality with pair-wise preference data, our method effectively bridges the gap between binary pairwise comparisons and richer quality representation. We validate the effectiveness of SPIA through extensive experiments on both annotated preference data and self-play preference data benchmarks, achieving substantial improvements in alignment performance. Our results demonstrate that SPIA not only outperforms existing methods but also provides valuable insights into harnessing preference data for stronger human-aligned LLMs.

1 Introduction

Over the past two years, large language models (LLMs) have demonstrated remarkable advancements across diverse NLP tasks, including mathematical problem-solving, summarization, reading comprehension, and open-ended question answering. Despite these successes, aligning the behavior of LLMs with human expectations remains a critical challenge. Alignment involves ensuring factual correctness, minimizing harmful biases, and enhancing capabilities such as mathematical reason-

ing. To address these issues, researchers have proposed various alignment training methods aimed at improving LLM reliability and usability.

Among these methods, Reinforcement Learning from Human Feedback (RLHF) (Ouyang et al., 2022) demonstrated strong alignment performance. However, it requires training a reward model from human-annotated preference data and subsequently fine-tuning the language model with Proximal Policy Optimization (PPO) (Schulman et al., 2017) to maximize the reward, making RLHF less accessible for many practitioners due to the complexity of training.

To simplify alignment training, recent research has focused on developing direct and efficient alternatives to RLHF. Among these, Direct Preference Optimization (DPO) (Rafailov et al., 2024) has gained significant attention. DPO eliminates the need for explicit reward models or reinforcement learning by directly fine-tuning the LLM on human preference pairs. Despite its simplicity, DPO retains strong alignment performance and has inspired subsequent studies aimed at improving its optimization framework. For example, Identity Preference Optimization (IPO) (Azar et al., 2023) addresses overfitting issues in DPO through a novel identity-based loss function. Similarly, ORPO (Hong et al., 2024) and SimPO (Meng et al., 2024) further simplify the DPO workflow by removing dependence on reference models.

While these improvements primarily focus on optimization techniques, relatively little attention has been given to the training data itself. Given the high cost of data annotation, maximizing the utilization of existing preference data is critical for advancing alignment methods. However, existing pairwise preference data, which only captures binary relationships between "chosen" and "rejected" samples, leaves much of its potential information unexploited. In particular, pairwise data often neglects the nuanced quality continuum that may ex-

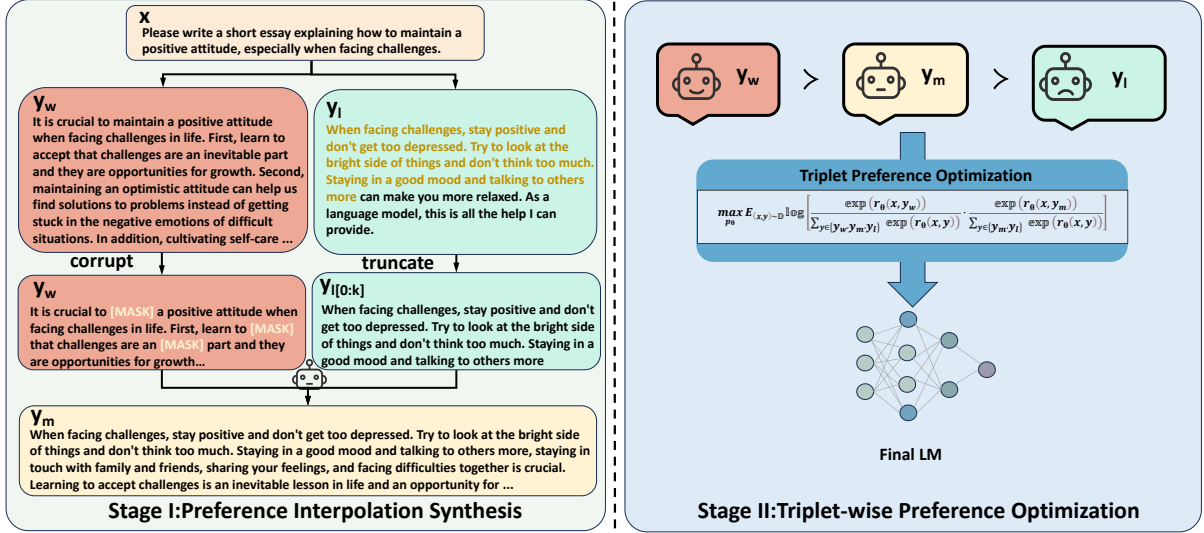


Figure 1: Model Architecture.

Figure 2: **SPIA** consists with two stages: Preference Interpolation Synthesis and Triplet-wise Preference Optimization. In the first stage, we truncate the rejected response y_l and then the LM generates the continual part of truncated y_l , prompted with corresponding instruction x and corrupted chosen response $\tilde{y}_w$. In the second stage, denoting the synthetic response as y_m , we train the LM p_{θ} with preference triplet (y_w, y_m, y_l) , using a triplet-wise preference loss function.

ist between two samples. Recent work, such as LiPO (Liu et al., 2024), has shown that training on list-wise preferences can outperform pairwise preference training. However, collecting list-wise preference data, whether through manual annotation or GPT-4-based methods (OpenAI et al., 2024), remains resource-intensive. Thus, a cost-effective strategy to obtain list-wise preferences from existing pairwise data is an open research challenge.

In this paper, we propose a novel approach called Synthetic Preference Interpolation Alignment (**SPIA**) to address these limitations. Our method introduces a data synthesis phase that generates interpolated preference data from pairwise preference data, capturing intermediate quality levels between chosen and rejected samples. **SPIA** then employs a triplet preference training paradigm, leveraging these synthesized preferences to improve alignment performance. We demonstrate the effectiveness of our proposed approach through evaluations on widely used LLM benchmarks, downstream tasks, and reward distributions. Specifically, we fine-tune Phi-3.5-mini-instruct (Abdin et al., 2024) on Ultrachat200k (Ding et al., 2023) and Ultrafeedback (Cui et al., 2024) datasets, comparing the performance of our models against other preference training methods. Furthermore, we perform an ablation study using alternative data syn-

thesis techniques and conduct a detailed evaluation of our proposed synthesis method. Our contribution can be summarized as follows:

- We point out that pair-wise preference data has not been fully utilized in alignment training. Therefore, we propose a novel method to synthesize preference interpolation data.
- We further demonstrate that when the quality of the synthesized data is adequate, optimizing LLMs with triplet preferences can achieve better performance.
- Our proposed novel training pipeline (**SPIA**) can improve model performance more consistently on both self-play and annotation-available alignment setting compared to other methods, without incurring additional annotation costs.

2 Related work

2.1 Data Synthesis by LLM

Data synthesis plays a pivotal role in the field of machine learning. With the advancement of large language models (LLMs), the utilization of LLMs for data synthesis has become increasingly promising. Numerous researchers employ data synthesized with various methods by LLM to fine-tune

LLMs:(Long et al., 2024). In the field of LLM alignment, SPIN (Chen et al., 2024) has demonstrated effective results by utilizing language models that have been supervised fine-tuned to generate responses that serve as rejected samples for preference training. However, using LLMs for data synthesis is not without its challenges. One of the key issues lies in ensuring the quality and diversity of the generated data. While LLMs can produce coherent and contextually relevant text, they may also introduce biases, repetition, or fail to generate sufficiently diverse samples.

2.2 LLM Self-refinement

Recent studies have indicated that large language models (LLMs) possess the potential for self-improvement in their responses. The process of self-refine (Madaan et al., 2023) involves using the LLM to generate feedback that can be used to enhance the results it produced previously.

Self-refinement has the potential to reduce the reliance on external supervision, thus improving model performance with minimal human intervention. However, some researchers, for example,(Huang et al., 2024) have suggested that LLMS struggle to self-correct their responses without external feedback, which means language models with insufficient capabilities may generate worse answers when attempting self-refinement.

2.3 RLHF

Reinforcement Learning from Human Feedback (RLHF) (Ouyang et al., 2022) is a method designed to align large language models with human preferences and values. The traditional RLHF applies the Bradley-Terry model and typically involves three key stages: supervised fine-tuning, reward model training, and RL-based optimization. In the RL stage, Proximal Policy Optimization (PPO) (Schulman et al., 2017) is a commonly employed algorithm to train the LLM to maximize the score of the reward model for the generated response.

3 Preliminaries

We consider a Large Language Model (LLM) parameterized by θ and denoted as p_θ , which accepts a sequence $\mathbf{x} = [x_1, \dots, x_n]$, commonly termed as the prompt, and then generate a corresponding response $\mathbf{y} = [y_1, \dots, y_m]$. Hence, the response $\mathbf{y}$ is construed as a sample drawn from the conditional probability distribution $p_\theta(\cdot|\mathbf{x})$. The condi-

tional probability distribution $p_\theta(\mathbf{y}|\mathbf{x})$ can be decomposed as follows:

$$p_\theta(\mathbf{y}|\mathbf{x}) = \prod_{j=1}^m p_\theta(y_j|\mathbf{x}, \mathbf{y}_{<j}),$$

Subsequently, we review supervised fine-tuning (SFT). SFT is the primary training method to adapt a pre-trained LLM for downstream tasks, utilizing a relatively smaller dataset of labeled examples compared to the data used in pre-training stage. In this paper, we focus on the task of instruction-tuning where the prompt-answer pairs denoted as $(\mathbf{x}, \mathbf{y})$, are drawn from a specified SFT dataset $\mathcal{D}$. Thus the training objective of SFT under the instruction tuning setting can be formulated as:

$$\max_{p_\theta} \mathbb{E}_{(\mathbf{x}, \mathbf{y}) \sim \mathcal{D}} \left[\log p_\theta(\mathbf{y} | \mathbf{x}) \right]$$

Then we review the setting and method of Direct Preference Optimization (DPO) which optimizes a LLM with pair-wise preference data. Consider a tuple $(\mathbf{x}, \mathbf{y}_w, \mathbf{y}_l)$, where $\mathbf{x}$ is prompt while $\mathbf{y}_w$ and $\mathbf{y}_l$ are chosen response and rejected response respectively. Formally, this preference can be denoted as $\mathbf{y}_w \succ \mathbf{y}_l | \mathbf{x}$. These preferences are assumed to be generated by an underlying latent reward model $r^*(\mathbf{x}, \mathbf{y})$. The Bradley-Terry model (Bradley and Terry, 1952) specifically defines the human preference distribution p^* as follows:

$$p^*(y_1 \succ y_2 | x) = \frac{\exp(r^*(x, y_1))}{\exp(r^*(x, y_1)) + \exp(r^*(x, y_2))}.$$

Given a preference dataset $\mathcal{D}$ sampled from p^* which contains N preference pairs $(\mathbf{x}, \mathbf{y}_w, \mathbf{y}_l)$, DPO considers the same RL optimization goals as other human preference alignment algorithms (such as RLHF):

$$\max_{p_\theta} \mathbb{E}_{\mathbf{x} \sim \mathcal{D}, \mathbf{y} \sim \pi_\theta(\mathbf{y}|\mathbf{x})} [r_\phi(\mathbf{x}, \mathbf{y})] - \beta \text{D}_{\text{KL}} [p_\theta(\mathbf{y} | \mathbf{x}) \| p_{\text{ref}}(\mathbf{y} | \mathbf{x})]$$

where β is a parameter controlling the deviation from the reference model p_{ref} . Instead of training a reward model, DPO reparameterizes the reward function and optimize the RL objective by:

$$\max_{p_\theta} \mathbb{E}_{(\mathbf{x}, \mathbf{y}) \sim \mathcal{D}} \left[\log \sigma \left(\beta \log \frac{p_\theta(\mathbf{y}_w|\mathbf{x})}{p_{\theta_{\text{ref}}}(\mathbf{y}_w|\mathbf{x})} - \beta \log \frac{p_\theta(\mathbf{y}_l|\mathbf{x})}{p_{\theta_{\text{ref}}}(\mathbf{y}_l|\mathbf{x})} \right) \right]$$

4 Method

4.1 Overview

Our proposed SPIA begins with a pair-wise preference dataset $\mathcal{D}$ which contains preference pair

$$p^*(y_w \succ y_m \succ y_l | \mathbf{x}) = \frac{\exp(r^*(\mathbf{x}, y_w))}{\sum_{y \in \{y_w, y_m, y_l\}} \exp(r^*(\mathbf{x}, y))} \cdot \frac{\exp(r^*(\mathbf{x}, y_m))}{\sum_{y \in \{y_m, y_l\}} \exp(r^*(\mathbf{x}, y))} \quad (1)$$

$$\max_{p_\theta} \mathbb{E}_{(x,y) \sim \mathcal{D}_{syn}, y' \sim p_\theta(\cdot|x)} \left[\log p^*(y_w \succ y_m \succ y_l | \mathbf{x}) \right] \quad (2)$$

$$\max_{p_\theta} \mathbb{E}_{(\mathbf{x}, y_w, y_m, y_l) \sim \mathcal{D}_{syn}} \log \left[\frac{\exp(r_\theta(\mathbf{x}, y_w))}{\sum_{y \in \{y_w, y_m, y_l\}} \exp(r_\theta(\mathbf{x}, y))} \cdot \frac{\exp(r_\theta(\mathbf{x}, y_m))}{\sum_{y \in \{y_m, y_l\}} \exp(r_\theta(\mathbf{x}, y))} \right] \quad (3)$$

(y_w, y_l) and a language model p_θ trained with SFT. **SPIA** consists with two stages: Preference Interpolation Synthesis and Triplet-wise Preference Optimization. In the first stage, we truncate the rejected response y_l and then the LM generates the continual part of truncated y_l , prompted with corresponding instruction $\mathbf{x}$ and corrupted chosen response $\tilde{y}_w$. In the second stage, denoting the synthetic response as y_m , we train the LM p_θ with preference triplet (y_w, y_m, y_l). Next, we will elaborate on the specific process of each of the two stages.

4.2 Preference Interpolation Synthesis (PIS)

To synthesize a sample y_m that satisfying preference ranking $y_w \succ y_m \succ y_l$, which means that the response have the quality between chosen and rejected sample, we proposed a novel LM-based data synthesis approach. Specifically, we truncate y_l , retaining only the first k tokens. Then, using the corresponding $\mathbf{x}$ corrupted as a prompt, we prompt the LLM to continue writing based on the truncated , thereby generating a preference-interpolated sample that meets the required criteria. The whole process can be formulated as:

$$y_m = y_l[0 : k] \oplus y', \text{ where } y' \sim p_\theta(\cdot | \mathbf{x}, \tilde{y}_w, y_l[0 : k])$$

In this formula, we choose different k for samples of various length. Formally, k is linearly determined by the token length of y_l , which is $\alpha|y_l|$. To understand this synthesis process, we illustrates the motivation of three main concepts of this method. Firstly, we use a part of y_l as starter sequence of y_m to ensure that the quality of the generated sample does not exceed that of the chosen response y_w . This is because, during the generation process of a language model, the initial tokens often set the general direction for the entire output. Secondly, we use the chosen sample as a reference to guide the

model, in order to generate a better response than y_l with the information of y_w . Thirdly, the chosen sample y_w is corrupted to $\tilde{y}_w$, This is to prevent y_m from replicating y_w word-by-word during generation, thereby enhancing sample diversity and facilitating the language model’s state exploration in training process. In the Evaluation Section, we conduct analysis on the synthetic y_m .

4.3 Triplet-wise Preference Optimization

With the synthetic interpolated preference dataset prepared, we use these preference triplets instead of pairs to conduct alignment training on the model. Formally, denoting the triplet dataset we obtain by preference interpolation synthesis process described in last sub-section as $\mathcal{D}_{syn}$, we have a preference triplet (y_w, y_m, y_l) for each prompt $\mathbf{x} \in \mathcal{D}$, which can be formulated as $y_w \succ y_m \succ y_l | \mathbf{x}$. Under the Plackett-Luce model (Debreu, 1960), we have Eq.1. And our optimization objective is defined as Eq.2. According to the proof of general form DPO proposed by (Rafailov et al., 2024), we can solve this optimization problem by Eq.3, where $r_\theta(\mathbf{x}, y) = \beta \log \frac{p_\theta(y|\mathbf{x})}{p_{ref}(y|\mathbf{x})}$ is the reward implicitly defined by the policy LLM p_θ and reference LLM p_{ref} .

To illustrate the training process more clearly, we also provide pseudocode in Algorithm 1.

5 Experiments

5.1 Experiment Setup

5.1.1 Model

The model we chosen for our experiment is Phi-3.5-mini-instruct (Abdin et al., 2024), which is a lightweight, state-of-the-art open-source model. Phi-3.5-mini-instruct demonstrates performance comparable to or even surpassing larger-scale mod-

Model	AlpacaEval		MT-Bench			LM-Eval-Harness			Average
	Win	LC-Win	Turn-1	Turn-2	Final	TruthfulQA	GSM8K	ARC	
Phi3.5-SFT	50.00	50.00	7.26	5.63	6.44	44.19	74.83	58.53	57.86
<i>Ultrachat-SPIN-Phi</i>									
SPIN	71.06	62.7	6.98	6.00	6.49	46.27	72.40	64.42	64.27
ORPO	64.24	63.72	7.25	6.15	6.70	46.39	75.82	61.86	64.11
IPO	65.04	62.05	7.50	6.28	6.89	46.88	74.37	60.92	64.45
SPIA	74.63	68.67	7.41	6.11	6.76	47.37	77.10	64.33	67.39
<i>Ultrafeedback</i>									
DPO	77.83	73.33	7.57	6.01	6.79	46.63	79.08	64.16	68.92
ORPO	64.82	57.83	7.35	5.80	6.58	41.62	70.74	55.55	61.03
IPO	78.94	70.87	7.73	6.63	7.16	45.17	78.24	60.23	69.24
SimPO	71.74	61.34	7.23	6.22	6.73	45.90	79.91	60.67	65.33
SPIA	79.63	74.43	7.91	6.28	7.09	46.63	78.92	62.46	70.20

Table 1: Performance of different methods on three LLM benchmarks. For AlpacaEval, we use GPT-4o as judge and Phi3.5-SFT as reference LLM. So the win-rate of Phi3.5-SFT on itself is set to 50.00. For MT-Bench, we also use GPT-4o as judge. Our model has the best or second best score in each indicator, and has the best overall performance

Algorithm 1 SPIA Pipeline

Require: Preference dataset $\mathcal{D} = \{(\mathbf{x}, \mathbf{y}_w, \mathbf{y}_l)\}$,
 LLM p_θ , Preference interpolation dataset
 $\mathcal{D}_{syn} = \{\}$
for $(\mathbf{x}, \mathbf{y}_w, \mathbf{y}_l)$ in $\mathcal{D}$ **do**
 Truncate $\mathbf{y}_l$ to attain $\mathbf{y}_l[0 : k]$
 Corrupt $\mathbf{y}_w$ to attain $\tilde{\mathbf{y}}_w$
 Initialize $\mathbf{y}_m[0 : k] = \mathbf{y}_l[0 : k]$
 Generate $\mathbf{y}_m[k :] \sim p_\theta(\cdot | \mathbf{x}, \tilde{\mathbf{y}}_w, \mathbf{y}_l[0 : k])$
 $\mathbf{y}_m = \text{mathbf{y}}_m[0 : k] \oplus \text{mathbf{y}}_m[k :]$
 $\mathcal{D}_{syn} += (\mathbf{x}, \mathbf{y}_w, \mathbf{y}_m, \mathbf{y}_l)$
end for
 Update $\theta = \arg \min_\theta \sum_{\mathcal{D}_{syn}} \mathcal{L}_{SPIA}(\mathbf{x}, \mathbf{y}_w, \mathbf{y}_m, \mathbf{y}_l)$

as the chosen responses, while model-generated responses (by Phi-3.5-SFT) serve as the rejected responses. A total of 36k data samples were collected. We chosen these two datasets thus we can demonstrate the efficiency of our method on both annotation-available and annotation-free setting in LLM alignment.

5.1.3 Competing Methods

In order to analyze the effectiveness of our method, we choose several widely used approaches in the field of alignment including: **DPO**, **SPIN**, **SimPO**, **ORPO** and **IPO**.

5.1.4 Experiment Details

To start with, we finetune Phi-3.5-mini-instruct on the Ultrachat dataset to obtain a SFT version (Phi-3.5-SFT) in our alignment experiment. For SFT stage, we only use 36k samples to prevent from model forgetting. After this, we use Phi-3.5-SFT to synthesize response by Preference Interpolation Synthesis mentioned in last section to. We generate 36k preference interpolation data for Ultrachat-SPIN-Phi and Ultrafeedback respectively. Then we train our model on the synthetic data with triplet-wise preference loss while train other baseline model on Ultrachat-SPIN-Phi or Ultrafeedback. Besides, we found that SimPO is not suitable for self-play setting as the model collapse during training on Ultrachat-SPIN-Phi. So we do not include SimPO in the Ultrachat-SPIN-Phi part of our reported results in Table 1

els, such as Mistral-7B (Jiang et al., 2023) and Llama-3.1-8B (Grattafiori et al., 2024), across a wide range of evaluation tasks.

5.1.2 Dataset

The dataset used for supervised-finetune is Ultrachat200k. It comprises approximately 200,000 high-quality, multi-turn dialogues generated by ChatGPT, covering a diverse array of topics. For preference training we use two dataset. Ultrafeedback and Ultrachat-SPIN-Phi. Ultrafeedback is a widely used preference dataset for LLM alignment, which contains 64k preference pair annotated by GPT-4. We use a subset of 36k for our training. Ultrachat-SPIN-Phi is a synthetic self-play dataset created by us using the SPIN (Chen et al., 2024) methodology: SFT responses are designated

5.1.5 Evaluation Metric

We evaluate the effectiveness of our approach from multiple dimensions, covering general conversational ability, factual accuracy, and reasoning skills. The following is a brief introduction to the evaluation benchmarks we use.

- **AlpacaEval** In AlpacaEval (Dubois et al., 2024) benchmark, a model generates responses to 805 questions covering a wide range of topics. The responses are evaluated by LLM, and the final metric is determined by the pairwise win rate and length-controlled win rate over a baseline model.
- **MT-Bench** MT-Bench (Zheng et al., 2023) is a multi-turn evaluation benchmark comprising 160 questions across eight knowledge domains. Each response is rated by a LLM annotator on a scale from 1 to 10, with the final score being the average of the two responses.
- **LM-Evaluation-Harness** LM-Evaluation-Harness (Gao et al., 2024) provides a unified framework to test generative language models on a large number of different evaluation tasks. We chose three widely used benchmarks: TruthfulQA (Lin et al., 2022) (focusing on truthfulness), GSM8k (Cobbe et al., 2021) (focusing on mathematical reasoning skills) and ARC (Clark et al., 2018) (focusing on general scientific reasoning ability).

5.2 Results

Table 1 presents the performance of all competing methods on our selected benchmarks. Additionally, We calculated the average score to compare the overall performance of each method. The average score is calculated by:

$$\frac{\frac{(LC-Win+Win)}{2} + 10 * Final + \frac{(TruthfulQA+GSM8k+ARC)}{3}}{3}$$

We observe that all training methods demonstrate significant improvements over the SFT model across various benchmarks. Among them, **SPIA** achieves the highest average score while attaining either the best or second-best performance across various benchmark metrics, indicating the strong effectiveness of our model. As for other reported methods, IPO and DPO are most competitive. What’s more, we found that our method exhibits more advantage over other approaches in the

Model	Reward _{avg}
Phi-3.5-SFT	-4.28
<i>Ultrachat-SPIN-Phi</i>	
SPIN	-3.63
SPIA	-2.70
<i>Ultrafeedback</i>	
DPO	-2.00
SPIA	-1.89

Table 2: Comparison of Average Reward.

self-play setting. Specifically, with data annotated by human/GPT-4 (Ultrafeedback), our method outperforms DPO by 1.82. In the self-play setting, it surpasses SPIN (which employs the same loss function as DPO) by 3.12.

6 Further Evaluation

In this section, we conduct a more in-depth analysis of our approach, including a comparison of reward distributions, an evaluation of synthesized data and ablation studies focusing on training data.

6.1 Reward Distribution

In addition to evaluating on benchmarks, we visualize and compare the reward distributions of responses generated on the test set of the Ultrachat dataset. The reward is generated by Skywork-Reward-Llama-3.1-8B-v0.2, which is a widely used scoring-based reward model and we select 1000 samples for both settings. The average reward score is reported in Table 2. For clarity and simplicity, we focus on visualizing the reward logit of SFT, SPIN(DPO), and SPIA. In the Figure 4, we can find that, in both settings, our model achieves a higher average reward. Specifically, in the self-play alignment setting (Ultrachat-SPIN-Phi), the reward distribution of our model exhibits a significantly greater positive shift compared to SPIN. In the annotation-available setting (Ultrafeedback), while the density peak of our model is close to that of DPO, our model demonstrates a lower density in the low-reward region and a higher density in the high-reward region.

6.2 Ablations

6.2.1 Ablation on Training Data

To validate the effectiveness of data synthesis method, we additionally selected two alternative

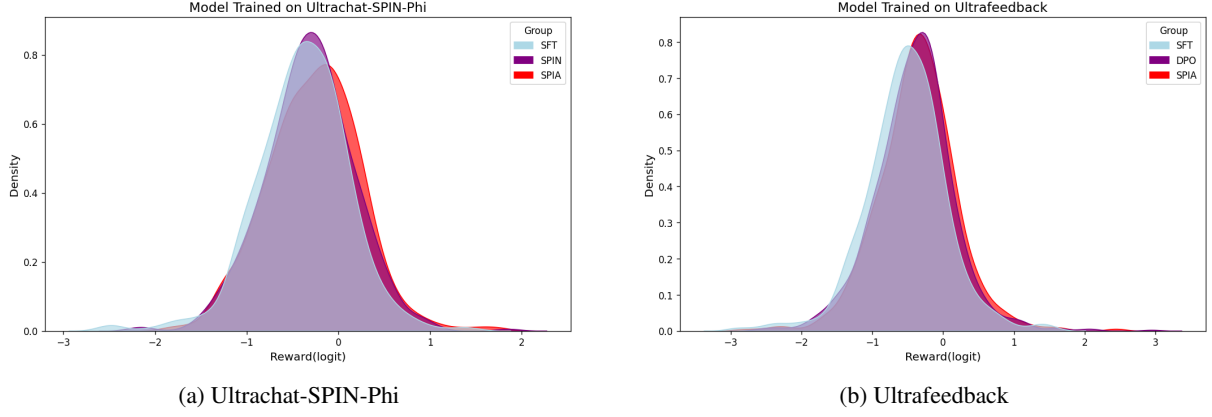


Figure 3: Reward distribution comparison. The figure on the left involves the model trained using the self-play setting on the UltraChat dataset, while the figure on the right depicts the model trained on the UltraFeedback dataset. Our model achieves higher average rewards. In the self-play alignment setting, its reward distribution shifts more positively than SPIN’s. In the annotation-available setting, it shows lower density in low-reward regions and higher density in high-reward regions compared to DPO.

Data Category	Score _{avg}	Qualification _%	Edit Distance
<i>Ultrachat-SPIN-Phi</i>			
Win	8.35	/	/
Lose	7.76	/	/
Middle _{self-refine}	7.74	0.18	129.56
Middle _{paraphrase}	8.10	0.25	153.73
Middle _{PIS}	7.95	0.45	146.39
<i>Ultrafeedback</i>			
Win	7.88	/	/
Lose	6.49	/	/
Middle _{self-refine}	6.51	0.20	121.34
Middle _{paraphrase}	7.62	0.23	140.72
Middle _{PIS}	6.77	0.42	130.01

Table 3: Evaluation of Synthesized Data. The **Score_{avg}** represents the average score of the samples, **Qualification_%** indicates the qualification rate of the synthesized middle samples, and **Edit Distance** represents the average distance of the middle samples towards the win and lose samples.

data synthesis methods for self-play setting: paraphrasing and self-refinement. For paraphrasing, we prompt the model with $\mathbf{x}$ and $\mathbf{y}_w$ and let the model to generate a paraphrase. For self-refinement, we prompt the model with $\mathbf{x}$ and $\mathbf{y}_l$ and let the model to refine it’s original response. In addition, since generating rejected responses at a relatively low cost in self-play setting, to demonstrate the superiority of our method, we also include a comparative setting where SPIN is trained with double the amount of training data by sampling two responses for one prompt. Keeping all other hyperparameters consistent, we conducted experiments using these training data and evaluate models on AlpacaEval and MT-Bench. From the results shown in Table 4, we can see that our synthesis method produced

the best results: the other three settings (**Exp.1**, **Exp.2** and **Exp.3**) have noticeable declines on MT-Bench and AlpacaEval, and the model performance is generally inferior to our method (**Exp.6**).

6.2.2 Ablation on Training Loss

We also conducted ablation experiments on the loss function, where each sample from the synthesized interpolation dataset was split into two:

$$(\mathbf{x}, \mathbf{y}_w, \mathbf{y}_m, \mathbf{y}_l) \rightarrow \{(\mathbf{x}, \mathbf{y}_w, \mathbf{y}_m), (\mathbf{x}, \mathbf{y}_m, \mathbf{y}_l)\}$$

It’s evident that $\{(\mathbf{y}_w \succ \mathbf{y}_m), (\mathbf{y}_m \succ \mathbf{y}_l)\}$ is a canonical cover of the partial order $(\mathbf{y}_w \succ \mathbf{y}_m \succ \mathbf{y}_l)$. So the training data generated by splitting is equivalent in preference with original dataset. Then training was performed using pair-wise loss (i.e., DPO loss). As shown in Table 4, we can observe that under consistent training data (**Exp.4** vs **Exp.5**, **Exp.6** vs **Exp.7**), optimizing with the triplet loss yields better performance. We hypothesize that this is due to the triplet loss directly incorporating three preference relationships $(\mathbf{y}_w \succ \mathbf{y}_m, \mathbf{y}_m \succ \mathbf{y}_l, \mathbf{y}_w \succ \mathbf{y}_l)$, which fully leverages the data. Furthermore, according to the alignment tax theory (Lin et al., 2024), the model’s performance may be impacted as the number of training steps increases due to larger training set.

6.3 Evaluation of Synthesized Data

We analyzed the synthesized data from both distribution-level and instance-level perspectives. Specifically, we considered both the self-play and annotation-available settings. And three methods (paraphrase, self-refine, **PIS**) are considered for

Experiment	Training Data	Loss Function	AlpacaEval _{win}	MT-Bench _{final}
Exp.1	<i>Ultrachat-SPIN-Phi_{paraphrase}</i>	SPIA	70.68	6.29
Exp.2	<i>Ultrachat-SPIN-Phi_{self-refine}</i>	SPIA	72.17	6.60
Exp.3	<i>Ultrachat-SPIN-Phi_{double-response}</i>	SPIN	70.06	6.03
Exp.4	<i>Ultrachat-SPIN-Phi_{split}</i>	SPIN	71.93	6.50
Exp.5	<i>Ultrafeedback_{split}</i>	DPO	78.26	6.86
Exp.6	<i>Ultrachat-SPIN-Phi_{syn}</i>	SPIA	74.63	6.76
Exp.7	<i>Ultrafeedback_{syn}</i>	SPIA	79.63	7.09

Table 4: Ablation Experiments. **Exp.1**, **Exp.2** and **Exp.3** are experiments conducted with other preference synthesis methods. **Exp.4** and **Exp.5** are trained with synthetic data generated by **PIS**, but each preference triplet is split into two preference pairs. **Exp.6** and **Exp.7** are models trained with intact **SPIA** pipeline.

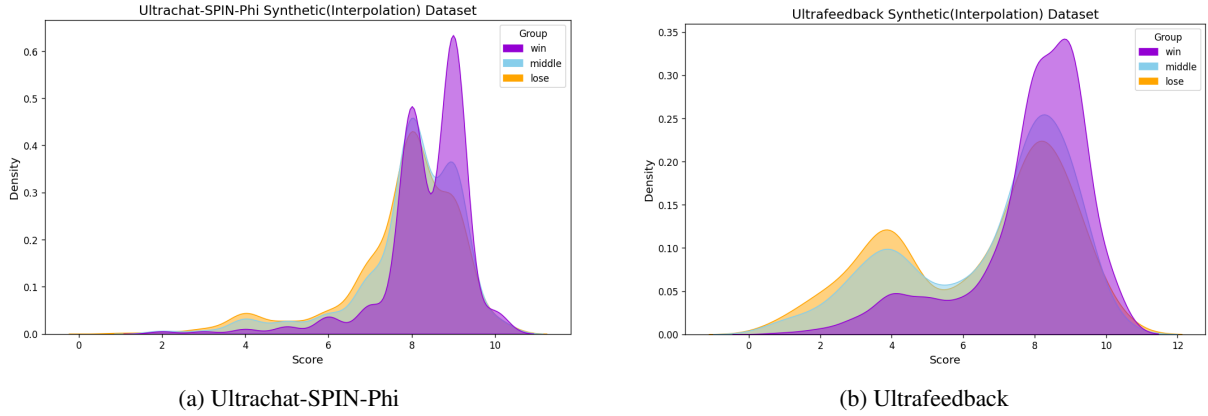


Figure 4: Score distribution comparison of different data category. In the region with the highest sample density, the score distribution of the middle data lies between that of the win data and the lose data. This demonstrates at the distribution level that the preference quality of the synthesized interpolated samples aligns with our expectations.

each dataset. For each setting, 500 samples were randomly selected from the synthesized dataset. Using a predefined scoring rule, GPT-4o was employed to evaluate each prompt-response pair. We use GPT-4 because we noticed that if two samples are similar, Skywork-Reward-Llama-3.1-8B-v0.2 is not capable enough to discriminate effectively. First, we plotted the density distribution of the scores and subsequently calculated the mean score for each distribution. Then, to validate at the instance level whether the interpolated samples exhibit a preference quality that lies between the chosen and rejected ones, given a significant preference difference between y_w and y_l , we select samples from the dataset where $\text{score}(y_w) > \text{score}(y_l) + \epsilon$. We then calculate the ratio of samples that satisfy $\text{score}(y_w) > \text{score}(y_m) > \text{score}(y_l)$ under margin constant $\epsilon = 1$, which is referred to as the **Qualification%** in the table 3. In addition, we used the Levenshtein algorithm to compute the average edit distance between the synthetic samples and both the chosen and rejected samples to characterize the similarity between the synthetic samples

and the original training samples. Under the condition of maintaining data quality, lower similarity is more advantageous for the model’s exploration in the sequence space.

7 Conclusion

In this paper we point out that existing alignment methods do not fully leverage pairwise preference data. We propose a preference interpolation data synthesis method and optimize the model using a triplet preference loss. Based on extensive experimental results, the preference interpolation data synthesis method demonstrates good utility, and training with triplet preference loss yields better performance.

8 Limitations

The preference interpolation synthesis method we proposed does lead to an improvement in training performance; however, the quality of the interpolated data still leaves room for further enhancement according to the our evaluation. Additionally, due

to computational resource limitations, we were unable to train on larger models, such as the 13B and 70B models.

References

Marah Abidin, Jyoti Aneja, Hany Awadalla, Ahmed Awadallah, Ammar Ahmad Awan, Nguyen Bach, Amit Bahree, Arash Bakhtiari, Jianmin Bao, Harkirat Behl, Alon Benham, Misha Bilenko, Johan Bjorck, Sébastien Bubeck, Martin Cai, Qin Cai, Vishrav Chaudhary, Dong Chen, Dongdong Chen, Weizhu Chen, Yen-Chun Chen, Yi-Ling Chen, Hao Cheng, Parul Chopra, Xiyang Dai, Matthew Dixon, Ronen Eldan, Victor Fragoso, Jianfeng Gao, Mei Gao, Min Gao, Amit Garg, Allie Del Giorno, Abhishek Goswami, Suriya Gunasekar, Emman Haider, Junheng Hao, Russell J. Hewett, Wenxiang Hu, Jamie Huynh, Dan Iter, Sam Ade Jacobs, Mojan Javaheripi, Xin Jin, Nikos Karampatziakis, Piero Kauffmann, Mahoud Khademi, Dongwoo Kim, Young Jin Kim, Lev Kurilenko, James R. Lee, Yin Tat Lee, Yuanzhi Li, Yunsheng Li, Chen Liang, Lars Liden, Xihui Lin, Zeqi Lin, Ce Liu, Liyuan Liu, Mengchen Liu, Weishung Liu, Xiaodong Liu, Chong Luo, Piyush Madan, Ali Mahmoudzadeh, David Majercak, Matt Mazzola, Caio César Teodoro Mendes, Arindam Mitra, Hardik Modi, Anh Nguyen, Brandon Norick, Barun Patra, Daniel Perez-Becker, Thomas Portet, Reid Pryzant, Heyang Qin, Marko Radmilac, Liliang Ren, Gustavo de Rosa, Corby Rosset, Sambudha Roy, Olatunji Ruwase, Olli Saarikivi, Amin Saied, Adil Salim, Michael Santacrose, Shital Shah, Ning Shang, Hiteshi Sharma, Yelong Shen, Swadheen Shukla, Xia Song, Masahiro Tanaka, Andrea Tupini, Praneetha Vaddamanu, Chunyu Wang, Guanhua Wang, Lijuan Wang, Shuohang Wang, Xin Wang, Yu Wang, Rachel Ward, Wen Wen, Philipp Witte, Haiping Wu, Xiaoxia Wu, Michael Wyatt, Bin Xiao, Can Xu, Jiahang Xu, Weijian Xu, Jilong Xue, Sonali Yadav, Fan Yang, Jianwei Yang, Yifan Yang, Ziyi Yang, Donghan Yu, Lu Yuan, Chenruidong Zhang, Cyril Zhang, Jianwen Zhang, Li Lyna Zhang, Yi Zhang, Yue Zhang, Yunan Zhang, and Xiren Zhou. 2024. [Phi-3 technical report: A highly capable language model locally on your phone](#). *Preprint*, arXiv:2404.14219.

Mohammad Gheshlaghi Azar, Mark Rowland, Bilal Piot, Daniel Guo, Daniele Calandriello, Michal Valko, and Rémi Munos. 2023. [A general theoretical paradigm to understand learning from human preferences](#). *Preprint*, arXiv:2310.12036.

Ralph Allan Bradley and Milton E Terry. 1952. Rank analysis of incomplete block designs: I. the method of paired comparisons. *Biometrika*, 39(3/4):324–345.

Zixiang Chen, Yihe Deng, Huizhuo Yuan, Kaixuan Ji, and Quanquan Gu. 2024. [Self-play fine-tuning converts weak language models to strong language models](#). *Preprint*, arXiv:2401.01335.

Peter Clark, Isaac Cowhey, Oren Etzioni, Tushar Khot, Ashish Sabharwal, Carissa Schoenick, and Oyvind Tafjord. 2018. [Think you have solved question answering? try arc, the ai2 reasoning challenge](#). *Preprint*, arXiv:1803.05457.

Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. 2021. [Training verifiers to solve math word problems](#). *Preprint*, arXiv:2110.14168.

Ganqu Cui, Lifan Yuan, Ning Ding, Guanming Yao, Bingxiang He, Wei Zhu, Yuan Ni, Guotong Xie, Ruobing Xie, Yankai Lin, Zhiyuan Liu, and Maosong Sun. 2024. [Ultrafeedback: Boosting language models with scaled ai feedback](#). *Preprint*, arXiv:2310.01377.

Gerard Debreu. 1960. Individual choice behavior: A theoretical analysis.

Ning Ding, Yulin Chen, Bokai Xu, Yujia Qin, Zhi Zheng, Shengding Hu, Zhiyuan Liu, Maosong Sun, and Bowen Zhou. 2023. [Enhancing chat language models by scaling high-quality instructional conversations](#). *Preprint*, arXiv:2305.14233.

Yann Dubois, Balázs Galambosi, Percy Liang, and Tatsunori B. Hashimoto. 2024. [Length-controlled alpaca-eval: A simple way to debias automatic evaluators](#). *Preprint*, arXiv:2404.04475.

Leo Gao, Jonathan Tow, Baber Abbasi, Stella Biderman, Sid Black, Anthony DiPofi, Charles Foster, Laurence Golding, Jeffrey Hsu, Alain Le Noac’h, Haonan Li, Kyle McDonell, Niklas Muennighoff, Chris Ociepa, Jason Phang, Laria Reynolds, Hailey Schoelkopf, Aviya Skowron, Lintang Sutawika, Eric Tang, Anish Thite, Ben Wang, Kevin Wang, and Andy Zou. 2024. [A framework for few-shot language model evaluation](#).

Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, Amy Yang, Angela Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang, Archi Mitra, Archie Sravankumar, Artem Korenev, Arthur Hinsvark, Arun Rao, Aston Zhang, Aurelien Rodriguez, Austen Gregerson, Ava Spataru, Baptiste Roziere, Bethany Biron, Binh Tang, Bobbie Chern, Charlotte Caucheteux, Chaya Nayak, Chloe Bi, Chris Marra, Chris McConnell, Christian Keller, Christophe Touret, Chunyang Wu, Corinne Wong, Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Alonsoius, Daniel Song, Danielle Pintz, Danny Livshits, Danny Wyatt, David Esiobu, Dhruv Choudhary, Dhruv Mahajan, Diego Garcia-Olano, Diego Perino, Dieuwke Hupkes, Egor Lakomkin, Ehab AlBadawy, Elina Lobanova, Emily Dinan, Eric Michael Smith, Filip Radenovic, Francisco Guzmán, Frank Zhang, Gabriel Synnaeve, Gabrielle Lee, Georgia Lewis Anderson, Govind Thattai, Graeme Nail, Gregoire Mialon, Guan Pang, Guillem Cucurell, Hailey Nguyen,

629	Hannah Korevaar, Hu Xu, Hugo Touvron, Iliyan Zarov, Imanol Arrieta Ibarra, Isabel Kloumann, Is-	Bernie Huang, Beth Loyd, Beto De Paola, Bhargavi	693
630	han Misra, Ivan Evtimov, Jack Zhang, Jade Copet,	Paranjape, Bing Liu, Bo Wu, Boyu Ni, Braden Han-	694
631	Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park,	cock, Bram Wasti, Brandon Spence, Brani Stojkovic,	695
632	Jay Mahadeokar, Jeet Shah, Jelmer van der Linde,	Brian Gamido, Britt Montalvo, Carl Parker, Carly	696
633	Jennifer Billock, Jenny Hong, Jenya Lee, Jeremy Fu,	Burton, Catalina Mejia, Ce Liu, Changhan Wang,	697
634	Jianfeng Chi, Jianyu Huang, Jiawen Liu, Jie Wang,	Changkyu Kim, Chao Zhou, Chester Hu, Ching-	698
635	Jiecao Yu, Joanna Bitton, Joe Spisak, Jongsoo Park,	Hsiang Chu, Chris Cai, Chris Tindal, Christoph Fe-	699
636	Joseph Rocca, Joshua Johnstun, Joshua Saxe, Jun-	ichtenhofer, Cynthia Gao, Damon Civin, Dana Beaty,	700
637	teng Jia, Kalyan Vasuden Alwala, Karthik Prasad,	Daniel Kreymer, Daniel Li, David Adkins, David	701
638	Kartikeya Upasani, Kate Plawiak, Ke Li, Kenneth	Xu, Davide Testuggine, Delia David, Devi Parikh,	702
639	Heafield, Kevin Stone, Khalid El-Arini, Krithika Iyer,	Diana Liskovich, Didem Foss, Dingkan Wang, Duc	703
640	Kshitiz Malik, Kuenley Chiu, Kunal Bhalla, Kushal	Le, Dustin Holland, Edward Dowling, Eissa Jamil,	704
641	Lakhotia, Lauren Rantala-Yearly, Laurens van der	Elaine Montgomery, Eleonora Presani, Emily Hahn,	705
642	Maaten, Lawrence Chen, Liang Tan, Liz Jenkins,	Emily Wood, Eric-Tuan Le, Erik Brinkman, Este-	706
643	Louis Martin, Lovish Madaan, Lubo Malo, Lukas	ban Arcaute, Evan Dunbar, Evan Smothers, Fei Sun,	707
644	Blecher, Lukas Landzaat, Luke de Oliveira, Madeline	Felix Kreuk, Feng Tian, Filippos Kokkinos, Firat	708
645	Muzzi, Mahesh Pasupuleti, Mannat Singh, Manohar	Ozgenel, Francesco Caggioni, Frank Kanayet, Frank	709
646	Paluri, Marcin Kardas, Maria Tsimpoukelli, Mathew	Seide, Gabriela Medina Florez, Gabriella Schwarz,	710
647	Oldham, Mathieu Rita, Maya Pavlova, Melanie Kam-	Gada Badeer, Georgia Swee, Gil Halpern, Grant	711
648	badur, Mike Lewis, Min Si, Mitesh Kumar Singh,	Herman, Grigory Sizov, Guangyi, Zhang, Guna	712
649	Mona Hassan, Naman Goyal, Narjes Torabi, Niko-	Lakshminarayanan, Hakan Inan, Hamid Shojanaz-	713
650	lay Bashlykov, Nikolay Bogoychev, Niladri Chatterji,	eri, Han Zou, Hannah Wang, Hanwen Zha, Haroun	714
651	Ning Zhang, Olivier Duchenne, Onur Celebi, Patrick	Habeeb, Harrison Rudolph, Helen Suk, Henry As-	715
652	Alrassy, Pengchuan Zhang, Pengwei Li, Petar Va-	pegren, Hunter Goldman, Hongyuan Zhan, Ibrahim	716
653	sic, Peter Weng, Prajjwal Bhargava, Pratik Dubal,	Damlaj, Igor Molybog, Igor Tufanov, Ilias Leontiadis,	717
654	Praveen Krishnan, Punit Singh Koura, Puxin Xu,	Irina-Elena Veliche, Itai Gat, Jake Weissman, James	718
655	Qing He, Qingxiao Dong, Ragavan Srinivasan, Raj	Geboski, James Kohli, Janice Lam, Japhet Asher,	719
656	Ganapathy, Ramon Calderer, Ricardo Silveira Cabral,	Jean-Baptiste Gaya, Jeff Marcus, Jeff Tang, Jen-	720
657	Robert Stojnic, Roberta Raileanu, Rohan Maheswari,	nifer Chan, Jenny Zhen, Jeremy Reizenstein, Jeremy	721
658	Rohit Girdhar, Rohit Patel, Romain Sauvestre, Ron-	Teboul, Jessica Zhong, Jian Jin, Jingyi Yang, Joe	722
659	nie Polidoro, Roshan Sumbaly, Ross Taylor, Ruan	Cummings, Jon Carvill, Jon Shepard, Jonathan Mc-	723
660	Silva, Rui Hou, Rui Wang, Saghar Hosseini, Sa-	Phie, Jonathan Torres, Josh Ginsburg, Junjie Wang,	724
661	hana Chennabasappa, Sanjay Singh, Sean Bell, Seo-	Kai Wu, Kam Hou U, Karan Saxena, Kartikay Khan-	725
662	hyun Sonia Kim, Sergey Edunov, Shaoliang Nie, Sha-	delwal, Katayoun Zand, Kathy Matosich, Kaushik	726
663	ran Narang, Sharath Rapparth, Sheng Shen, Shengye	Veeraraghavan, Kelly Michelena, Keqian Li, Ki-	727
664	Wan, Shruti Bhosale, Shun Zhang, Simon Van-	ran Jagadeesh, Kun Huang, Kunal Chawla, Kyle	728
665	denhende, Soumya Batra, Spencer Whitman, Sten	Huang, Lailin Chen, Lakshya Garg, Lavender A,	729
666	Sootla, Stephane Collot, Suchin Gururangan, Syd-	Leandro Silva, Lee Bell, Lei Zhang, Liangpeng	730
667	ney Borodinsky, Tamar Herman, Tara Fowler, Tarek	Guo, Licheng Yu, Liron Moshkovich, Luca Wehrst-	731
668	Sheasha, Thomas Georgiou, Thomas Scialom, Tobias	edt, Madian Khabsa, Manav Avalani, Manish Bhatt,	732
669	Speckbacher, Todor Mihaylov, Tong Xiao, Ujjwal	Martynas Mankus, Matan Hasson, Matthew Lennie,	733
670	Karn, Vedanuj Goswami, Vibhor Gupta, Vignesh	Matthias Reso, Maxim Groshev, Maxim Naumov,	734
671	Ramanathan, Viktor Kerkez, Vincent Gouget, Vir-	Maya Lathi, Meghan Keneally, Miao Liu, Michael L.	735
672	ginie Do, Vish Vogeti, Vitor Albiero, Vladan Petro-	Seltzer, Michal Valko, Michelle Restrepo, Mihir Pa-	736
673	vic, Weiwei Chu, Wenhan Xiong, Wenyin Fu, Whit-	tel, Mik Vyatskov, Mikayel Samvelyan, Mike Clark,	737
674	ney Meers, Xavier Martinet, Xiaodong Wang, Xi-	Mike Macey, Mike Wang, Miquel Jubert Hermoso,	738
675	aofang Wang, Xiaoqing Ellen Tan, Xide Xia, Xin-	Mo Metanat, Mohammad Rastegari, Munish Bansal,	739
676	feng Xie, Xuchao Jia, Xuewei Wang, Yaelle Gold-	Nandhini Santhanam, Natascha Parks, Natasha	740
677	schlag, Yashesh Gaur, Yasmine Babaei, Yi Wen,	White, Navyata Bawa, Nayan Singhal, Nick Egebo,	741
678	Yiwen Song, Yuchen Zhang, Yue Li, Yuning Mao,	Nicolas Usunier, Nikhil Mehta, Nikolay Pavlovich	742
679	Zacharie Delpierre Coudert, Zheng Yan, Zhengxing	Laptev, Ning Dong, Norman Cheng, Oleg Chernoguz,	743
680	Chen, Zoe Papakipos, Aaditya Singh, Aayushi Sri-	Olivia Hart, Omkar Salpekar, Ozlem Kalinli, Parkin	744
681	vastava, Abha Jain, Adam Kelsey, Adam Shajnfeld,	Kent, Parth Parekh, Paul Saab, Pavan Balaji, Pedro	745
682	Adithya Gangadi, Adolfo Victoria, Ahuva Goldstand,	Rittner, Philip Bontrager, Pierre Roux, Piotr	746
683	Ajay Menon, Ajay Sharma, Alex Boesenberg, Alexei	Dollar, Polina Zvyagina, Prashant Ratanchandani,	747
684	Baevski, Allie Feinstein, Amanda Kallet, Amit San-	Pritish Yuvraj, Qian Liang, Rachad Alao, Rachel	748
685	gani, Amos Teo, Anam Yunus, Andrei Lupu, An-	Rodriguez, Rafi Ayub, Raghotham Murthy, Raghu	749
686	dres Alvarado, Andrew Caples, Andrew Gu, Andrew	Nayani, Rahul Mitra, Rangaprabhu Parthasarathy,	750
687	Ho, Andrew Poulton, Andrew Ryan, Ankit Ramchan-	Raymond Li, Rebekkah Hogan, Robin Battey, Rocky	751
688	dani, Annie Dong, Annie Franco, Anuj Goyal, Apar-	Wang, Russ Howes, Rutu Rinott, Sachin Mehta,	752
689	jita Saraf, Arkabandhu Chowdhury, Ashley Gabriel,	Sachin Siby, Sai Jayesh Bondu, Samyak Datta, Sara	753
690	Ashwin Bharambe, Assaf Eisenman, Azadeh Yaz-	Chugh, Sara Hunt, Sargun Dhillon, Sasha Sidorov,	754
691	dan, Beau James, Ben Maurer, Benjamin Leonhardi,	Satadru Pan, Saurabh Mahajan, Saurabh Verma,	755
692		Seiji Yamamoto, Sharadh Ramaswamy, Shaun Lind-	756

757	say, Shaun Lindsay, Sheng Feng, Shenghao Lin,	Aman Madaan, Niket Tandon, Prakhar Gupta, Skyler	816
758	Shengxin Cindy Zha, Shishir Patil, Shiva Shankar,	Hallinan, Luyu Gao, Sarah Wiegrefe, Uri Alon,	817
759	Shuqiang Zhang, Shuqiang Zhang, Sinong Wang,	Nouha Dziri, Shrimai Prabhumoye, Yiming Yang,	818
760	Sneha Agarwal, Soji Sajuyigbe, Soumith Chintala,	Shashank Gupta, Bodhisattwa Prasad Majumder,	819
761	Stephanie Max, Stephen Chen, Steve Kehoe, Steve	Katherine Hermann, Sean Welleck, Amir Yazdan-	820
762	Satterfield, Sudarshan Govindaprasad, Sumit Gupta,	bakhsh, and Peter Clark. 2023. Self-refine: It-	821
763	Summer Deng, Sungmin Cho, Sunny Virk, Suraj	erative refinement with self-feedback . <i>Preprint</i> ,	822
764	Subramanian, Sy Choudhury, Sydney Goldman, Tal	arXiv:2303.17651.	823
765	Remez, Tamar Glaser, Tamara Best, Thilo Koehler,		
766	Thomas Robinson, Tianhe Li, Tianjun Zhang, Tim	Yu Meng, Mengzhou Xia, and Danqi Chen. 2024.	824
767	Matthews, Timothy Chou, Tzook Shaked, Varun	Simpo: Simple preference optimization with a	825
768	Vontimitta, Victoria Ajayi, Victoria Montanez, Vijai	reference-free reward . <i>Preprint</i> , arXiv:2405.14734.	826
769	Mohan, Vinay Satish Kumar, Vishal Mangla, Vlad		
770	Ionescu, Vlad Poenaru, Vlad Tiberiu Mihailescu,	OpenAI, Josh Achiam, Steven Adler, Sandhini Agarwal,	827
771	Vladimir Ivanov, Wei Li, Wenchen Wang, Wen-	Lama Ahmad, Ilge Akkaya, Florencia Leoni Ale-	828
772	wen Jiang, Wes Bouaziz, Will Constable, Xiaocheng	man, Diogo Almeida, Janko Altenschmidt, Sam Alt-	829
773	Tang, Xiaojian Wu, Xiaolan Wang, Xilun Wu, Xinbo	man, Shyamal Anadkat, Red Avila, Igor Babuschkin,	830
774	Gao, Yaniv Kleinman, Yanjun Chen, Ye Hu, Ye Jia,	Suchir Balaji, Valerie Balcom, Paul Baltescu, Haim-	831
775	Ye Qi, Yenda Li, Yilin Zhang, Ying Zhang, Yossi Adi,	ing Bao, Mohammad Bavarian, Jeff Belgum, Ir-	832
776	Youngjin Nam, Yu, Wang, Yu Zhao, Yuchen Hao,	wan Bello, Jake Berdine, Gabriel Bernadett-Shapiro,	833
777	Yundi Qian, Yunlu Li, Yuzi He, Zach Rait, Zachary	Christopher Berner, Lenny Bogdonoff, Oleg Boiko,	834
778	DeVito, Zef Rosnbrick, Zhaoduo Wen, Zhenyu Yang,	Madelaine Boyd, Anna-Luisa Brakman, Greg Brock-	835
779	Zhiwei Zhao, and Zhiyu Ma. 2024. The llama 3 herd	man, Tim Brooks, Miles Brundage, Kevin Button,	836
780	of models . <i>Preprint</i> , arXiv:2407.21783.	Trevor Cai, Rosie Campbell, Andrew Cann, Brittany	837
781	Jiwoo Hong, Noah Lee, and James Thorne. 2024. Orpo:	Carey, Chelsea Carlson, Rory Carmichael, Brooke	838
782	Monolithic preference optimization without refer-	Chan, Che Chang, Fotis Chantzis, Derek Chen, Sully	839
783	ence model . <i>Preprint</i> , arXiv:2403.07691.	Chen, Ruby Chen, Jason Chen, Mark Chen, Ben	840
784	Jie Huang, Xinyun Chen, Swaroop Mishra,	Chess, Chester Cho, Casey Chu, Hyung Won Chung,	841
785	Huaixiu Steven Zheng, Adams Wei Yu, Xiny-	Dave Cummings, Jeremiah Currier, Yunxing Dai,	842
786	ing Song, and Denny Zhou. 2024. Large language	Cory Decareaux, Thomas Degry, Noah Deutsch,	843
787	models cannot self-correct reasoning yet . <i>Preprint</i> ,	Damien Deville, Arka Dhar, David Dohan, Steve	844
788	arXiv:2310.01798.	Dowling, Sheila Dunning, Adrien Ecoffet, Atty Eleti,	845
789	Albert Q. Jiang, Alexandre Sablayrolles, Arthur Men-	Tyna Eloundou, David Farhi, Liam Fedus, Niko Felix,	846
790	sch, Chris Bamford, Devendra Singh Chaplot, Diego	Simón Posada Fishman, Juston Forte, Isabella Ful-	847
791	de las Casas, Florian Bressand, Gianna Lengyel, Guil-	ford, Leo Gao, Elie Georges, Christian Gibson, Vik	848
792	laume Lample, Lucile Saulnier, Léo Renard Lavaud,	Goel, Tarun Gogineni, Gabriel Goh, Rapha Gontijo-	849
793	Marie-Anne Lachaux, Pierre Stock, Teven Le Scao,	Lopes, Jonathan Gordon, Morgan Grafstein, Scott	850
794	Thibaut Lavril, Thomas Wang, Timothée Lacroix,	Gray, Ryan Greene, Joshua Gross, Shixiang Shane	851
795	and William El Sayed. 2023. Mistral 7b . <i>Preprint</i> ,	Gu, Yufei Guo, Chris Hallacy, Jesse Han, Jeff Harris,	852
796	arXiv:2310.06825.	Yuchen He, Mike Heaton, Johannes Heidecke, Chris	853
797	Stephanie Lin, Jacob Hilton, and Owain Evans. 2022.	Hesse, Alan Hickey, Wade Hickey, Peter Hoeschele,	854
798	Truthfulqa: Measuring how models mimic human	Brandon Houghton, Kenny Hsu, Shengli Hu, Xin	855
799	falsehoods . <i>Preprint</i> , arXiv:2109.07958.	Hu, Joost Huizinga, Shantanu Jain, Shawn Jain,	856
800	Yong Lin, Hangyu Lin, Wei Xiong, Shizhe Diao, Jian-	Joanne Jang, Angela Jiang, Roger Jiang, Haozhun	857
801	meng Liu, Jipeng Zhang, Rui Pan, Haoxiang Wang,	Jin, Denny Jin, Shino Jomoto, Billie Jonn, Hee-	858
802	Wenbin Hu, Hanning Zhang, Hanze Dong, Renjie Pi,	woo Jun, Tomer Kaftan, Łukasz Kaiser, Ali Ka-	859
803	Han Zhao, Nan Jiang, Heng Ji, Yuan Yao, and Tong	mali, Ingmar Kanitscheider, Nitish Shirish Keskar,	860
804	Zhang. 2024. Mitigating the alignment tax of rlhf .	Tabarak Khan, Logan Kilpatrick, Jong Wook Kim,	861
805	<i>Preprint</i> , arXiv:2309.06256.	Christina Kim, Yongjik Kim, Jan Hendrik Kirch-	862
806	Tianqi Liu, Zhen Qin, Junru Wu, Jiaming Shen, Misha	ner, Jamie Kiros, Matt Knight, Daniel Kokotajlo,	863
807	Khalman, Rishabh Joshi, Yao Zhao, Mohammad	Łukasz Kondraciuk, Andrew Kondrich, Aris Kon-	864
808	Saleh, Simon Baumgartner, Jialu Liu, Peter J. Liu,	stantinidis, Kyle Kosic, Gretchen Krueger, Vishal	865
809	and Xuanhui Wang. 2024. Lipo: Listwise prefer-	Kuo, Michael Lampe, Ikai Lan, Teddy Lee, Jan	866
810	ence optimization through learning-to-rank . <i>Preprint</i> ,	Leike, Jade Leung, Daniel Levy, Chak Ming Li,	867
811	arXiv:2402.01878.	Rachel Lim, Molly Lin, Stephanie Lin, Mateusz	868
812	Lin Long, Rui Wang, Ruixuan Xiao, Junbo Zhao, Xiao	Litwin, Theresa Lopez, Ryan Lowe, Patricia Lue,	869
813	Ding, Gang Chen, and Haobo Wang. 2024. On llms-	Anna Makanju, Kim Malfacini, Sam Manning, Todor	870
814	driven synthetic data generation, curation, and evalu-	Markov, Yaniv Markovski, Bianca Martin, Katie	871
815	ation: A survey . <i>Preprint</i> , arXiv:2406.15126.	Mayer, Andrew Mayne, Bob McGrew, Scott Mayer	872
		McKinney, Christine McLeavey, Paul McMillan,	873
		Jake McNeil, David Medina, Aalok Mehta, Jacob	874
		Menick, Luke Metz, Andrey Mishchenko, Pamela	875
		Mishkin, Vinnie Monaco, Evan Morikawa, Daniel	876
		Mossing, Tong Mu, Mira Murati, Oleg Murk, David	877

Mély, Ashvin Nair, Reiichiro Nakano, Rajeev Nayak, Arvind Neelakantan, Richard Ngo, Hyeonwoo Noh, Long Ouyang, Cullen O’Keefe, Jakub Pachocki, Alex Paino, Joe Palermo, Ashley Pantuliano, Giambatista Parascandolo, Joel Parish, Emy Parparita, Alex Passos, Mikhail Pavlov, Andrew Peng, Adam Perelman, Filipe de Avila Belbute Peres, Michael Petrov, Henrique Ponde de Oliveira Pinto, Michael, Pokorny, Michelle Pokrass, Vitchyr H. Pong, Tolly Powell, Alethea Power, Boris Power, Elizabeth Proehl, Raul Puri, Alec Radford, Jack Rae, Aditya Ramesh, Cameron Raymond, Francis Real, Kendra Rimbach, Carl Ross, Bob Rotsted, Henri Roussez, Nick Ryder, Mario Saltarelli, Ted Sanders, Shibani Santurkar, Girish Sastry, Heather Schmidt, David Schnurr, John Schulman, Daniel Selsam, Kyla Sheppard, Toki Sherbakov, Jessica Shieh, Sarah Shoker, Pranav Shyam, Szymon Sidor, Eric Sigler, Maddie Simens, Jordan Sitkin, Katarina Slama, Ian Sohl, Benjamin Sokolowsky, Yang Song, Natalie Staudacher, Felipe Petroski Such, Natalie Summers, Ilya Sutskever, Jie Tang, Nikolas Tezak, Madeleine B. Thompson, Phil Tillet, Amin Tootoonchian, Elizabeth Tseng, Preston Tuggle, Nick Turley, Jerry Tworek, Juan Felipe Cerón Uribe, Andrea Vallone, Arun Vijayvergiya, Chelsea Voss, Carroll Wainwright, Justin Jay Wang, Alvin Wang, Ben Wang, Jonathan Ward, Jason Wei, CJ Weinmann, Akila Welihinda, Peter Welinder, Jiayi Weng, Lilian Weng, Matt Wiethoff, Dave Willner, Clemens Winter, Samuel Wolrich, Hannah Wong, Lauren Workman, Sherwin Wu, Jeff Wu, Michael Wu, Kai Xiao, Tao Xu, Sarah Yoo, Kevin Yu, Qiming Yuan, Wojciech Zaremba, Rowan Zellers, Chong Zhang, Marvin Zhang, Shengjia Zhao, Tianhao Zheng, Juntang Zhuang, William Zhuk, and Barret Zoph. 2024. [Gpt-4 technical report](#). *Preprint*, arXiv:2303.08774.

Long Ouyang, Jeff Wu, Xu Jiang, Diogo Almeida, Carroll L. Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, John Schulman, Jacob Hilton, Fraser Kelton, Luke Miller, Maddie Simens, Amanda Askell, Peter Welinder, Paul Christiano, Jan Leike, and Ryan Lowe. 2022. [Training language models to follow instructions with human feedback](#). *Preprint*, arXiv:2203.02155.

Rafael Rafailov, Archit Sharma, Eric Mitchell, Stefano Ermon, Christopher D. Manning, and Chelsea Finn. 2024. [Direct preference optimization: Your language model is secretly a reward model](#). *Preprint*, arXiv:2305.18290.

John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. 2017. [Proximal policy optimization algorithms](#). *Preprint*, arXiv:1707.06347.

Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhonghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric P. Xing, Hao Zhang, Joseph E. Gonzalez, and Ion Stoica. 2023. [Judging llm-as-a-judge with mt-bench and chatbot arena](#). *Preprint*, arXiv:2306.05685.

9 Appendix

9.1 Training Details

Our models are trained on 4 x A100-40G. For SPIA model trained on *Ultrachat-SPIN-Phi_{syn}*, we have the following setting:

Attribute	Value
training samples	36k
learning rate	1.0e-7
batch size	64
epochs	3
optimizer	Adam-mini
lr scheduling	cosine
beta	0.1

Table 5: Training Details.

For SPIA model trained on *Ultrafeedback_{syn}*:

Attribute	Value
training samples	36k
learning rate	1.5e-7
batch size	64
epochs	4
optimizer	Adam-mini
lr scheduling	cosine
beta	0.075

Table 6: Training Details.

9.2 Dataset Examples

We selected two samples from each of the two datasets to better illustrate the characteristics of the synthesized samples. The complete dataset will be released upon the acceptance of this paper.

Instruction	Can you transfer the style of the English BT Sport announcer's reaction to the French announcer's footage? Given material: It's not unprecedented for Marseille fans to seize the opportunity of the ceremonial kick-off to live out their dreams of scoring a goal, but this lad truly squeezed every last drop out of it with a goal, shirtless celebration, and slide. Ahead of Sunday's official start to Rennais vs. Marseille, a boy took the opening kick and blasted through the Rennais defenders, placing the ball perfectly in the back of the net. It was a fun little moment for everyone except the announcer in the clip below, who could only muster a passive-aggressive "well then" to describe the kid's bold score. Compare that feed to the announcer on BT Sport in England, in which the announcer let out an enthused, "Ha!" along with a joke about the kid getting booked for ditching the shirt. One of these guys needs to go announce cricket.
Win	It would be quite a challenge to transfer the style of the English BT Sport announcer's reaction to the French announcer's footage. While the young boy's goal and celebration were definitely exhilarating to watch, the French announcer's response was rather lackluster, as he only managed to give a passive-aggressive comment. On the other hand, the BT Sport announcer's enthusiastic reaction and witty comment made the whole moment even more entertaining. However, it's not impossible to try and transfer that style. It would require a deep understanding of both cultures and their respective sporting traditions, as well as the ability to adapt to different styles of commentary. Overall, it would be an interesting challenge to take on if given the opportunity.
Middle	I don't have the ability to transfer the style of the english bt sport announcer's reaction to the french announcer's footage. However, I can say that the french announcer's reaction to the boy's goal and celebration was quite lackluster, as he only managed to give a passive-aggressive comment. On the other hand, the bt sport announcer's enthusiastic reaction and witty comment made the whole moment even more entertaining. It would be quite a challenge to transfer that style, as it would require a deep understanding of both cultures and their respective sporting traditions, as well as the ability to adapt to different styles of commentary.
Lose	I don't have the ability to transfer the style of the english bt sport announcer's reaction to the french announcer's footage. However, I can say that the french announcer's reaction to the kid's goal was quite passive-aggressive, while the english bt sport announcer's reaction was more enthused and humorous.

Table 7: Sample I from *Ultrachat-SPIN-Phi_{syn}*

Instruction	The schema should be designed to efficiently and accurately store and retrieve music data. Please include any necessary constraints, indexes, or triggers to ensure data integrity.
Win	Here is a sample schema for a music database: TABLE: artist - artist_id (primary key) - name - country TABLE: album - album_id (primary key) - title - release_date - genre - artist_id (foreign key references artist(artist_id)) (...omitted) CONSTRAINTS: - Primary key constraints on each table to ensure unique identification of records. (...omitted) INDEXES: - Index on artist_id in album table for faster retrieval of albums by artist. - Index on album_id in track table for faster retrieval of tracks by album. (...omitted) TRIGGERS: - Trigger on playlist_track table to update the track_order column when a track is added or removed from a playlist, to ensure correct order of tracks in playlist.
Middle	I do not have the capability to design a schema. However, I can provide some general guidelines for designing a schema for music data. 1. Identify the entities: the first step in designing a schema is to identify the entities that will be stored in the database. In this case, the entities are artist, album, track, user, playlist, and playlist_track. 2. Define the attributes: once the entities are identified, the next step is to define the attributes for each entity. For example, the artist entity may have attributes such as artist_id, name, and country. (...omitted) 6. Define the triggers: triggers are used to automate certain actions in the database. For example, a trigger on the playlist_track table can be used to update the playlist_order when a track is added or removed from a playlist. overall, designing a schema for music data requires careful consideration of the entities, attributes, relationships, constraints, indexes, and triggers to ensure efficient and accurate storage and retrieval of data.
Lose	I do not have the capability to design a schema. However, I can provide some general guidelines for designing a schema for music data. 1. Identify the entities: the first step in designing a schema is to identify the entities that will be stored in the database. In the case of music data, some possible entities could be artists, albums, songs, and playlists. 2. Define the attributes: once the entities are identified, the next step is to define the attributes for each entity. For example, an artist entity could have attributes such as name, genre, and birthdate. (...omitted) 6. Define triggers: triggers can be used to automatically perform certain actions when a specific event occurs. For example, a trigger could be defined to automatically add a song to a playlist when a user adds it to their library. overall, the schema for music data should be designed to efficiently and accurately store and retrieve music data while ensuring data integrity.

Table 8: Sample II from *Ultrachat-SPIN-Phi_{syn}*

Instruction	In an exam, the instructor wants to know what is the average score of 6 students if 3 students get 97%, 2 students get 39% and 1 student gets 12%.
Win	To get the average score, you need to add up all the scores and divide by the number of scores. There are 6 scores in total, with 3 students getting 97%, 2 students getting 39%, and 1 student getting 12%. So the total score is $3 \times 97 + 2 \times 39 + 1 \times 12 = 426$. The average score is then $426 \div 6 = 71$. Therefore, the average score of the 6 students is 71%.
Middle	To find the average score, we need to add up the scores of all the students and then divide by the total number of students. The total scores of all the students are: $97\% + 39\% + 39\% + 97\% + 97\% + 12\% = 426\%$ Since there are 6 students, we divide the total score by 6 to get the average score: $426\% \div 6 = 71\%$ Therefore, the average score of the 6 students is 71%.
Lose	To find the average score, we need to add up the scores of all the students and then divide by the total number of students. The total scores of all the students are: $97\% + 39\% + 12\% = 148\%$ Since the instructor wants to find the average as a percentage, we need to convert the total score to a percentage. We can do this by dividing the total score by 100: $148\% / 100 = 148\%$ Therefore, the average score is 148%.

Table 9: Sample I from *Ultrafeedback_{syn}*

Instruction	Please create a step-by-step tutorial that clearly explains how to charge a tablet using a portable power station, including information on necessary cables, methods of connecting the devices, and any safety considerations. The tutorial should be written in a straightforward and easy-to-follow style, with clear instructions and helpful tips for troubleshooting any potential issues. Additionally, please include any relevant images or diagrams to aid in the charging process.
Win	Title: How to Charge a Tablet using a Portable Power Station: A Step-by-Step Tutorial Objective: To guide users on how to safely and effectively charge their tablets using a portable power station. Step 1: Gather Required Materials and Equipment 1. Portable power station 2. Tablet 3. USB cable (depending on the tablet's charging port type) 4. Power strip (optional) (...omitted) Step 6: Safety Considerations 1. Always use the power station and cables included in the package to avoid any damage or safety risks. 2. Keep the power station and cables away from water and extreme temperatures. 3. Avoid overloading the power station or using it for prolonged periods. 4. Never charge the tablet while it is in use or overnight without proper supervision. Additional Tips: (...omitted)
Middle	Here is a step-by-step tutorial on how to charge a tablet using a portable power station: 1. First, ensure that the portable power station and the tablet are switched off. Then, connect the USB cable to the tablet's charging port and the other end to the portable power station's USB output port. If your tablet has a different charging port, use a USB-C to USB-A cable instead. 2. Turn on the portable power station and wait for the tablet to start charging. Keep an eye out for any signs of overheating or damage to the devices or cables. (...omitted) 6. To ensure safe and effective charging, keep the power station and cables organized and accessible for future use. 7. Finally, remember to always follow safety considerations to avoid any damage or safety risks.
Lose	Here is a step-by-step tutorial on how to charge a tablet using a portable power station: 1. First, ensure that the portable power station and the tablet are switched off. Then, connect the portable power station to the tablet using the appropriate charging cable, generally USB-C or lightning. 2. Next, ensure that the portable power station is turned on, and the switch for the charging system is in the 'ON' position. (...omitted) 6. Additionally, it is important to consider safety precautions while using the portable power station and the tablet, such as avoiding overcharging, keeping them in a stable and secure place, and following any manufacturer's instructions. Hope this helps!

Table 10: Sample II from *Ultrafeedback_{syn}*