

Weak-to-Strong Extrapolation Expedites Alignment

Chujie Zheng^{1,2*} Ziqi Wang³ Heng Ji³ Minlie Huang^{1†} Nanyun Peng^{2†}

¹The CoAI Group, DCST, BNRist, Tsinghua University

²University of California, Los Angeles ³University of Illinois Urbana-Champaign

chujiezhengchn@gmail.com aihuang@tsinghua.edu.cn violetpeng@cs.ucla.edu

Abstract

Although the capabilities of large language models (LLMs) ideally scale up with increasing data and compute, they are inevitably constrained by limited resources in reality. Suppose we have a moderately trained LLM (e.g., trained to align with human preference) in hand, *can we further exploit its potential and cheaply acquire a stronger model?* In this paper, we propose a simple method called **EXPO** to boost LLMs' alignment with human preference. EXPO assumes that a medium-aligned model can be interpolated between a less-aligned (*weaker*) model, e.g., the initial SFT model, and a better-aligned (*stronger*) one, thereby directly obtaining this *stronger* model by *extrapolating* from the weights of the former two relatively *weaker* models. On the AlpacaEval 2.0 benchmark, we show that EXPO pushes models trained with less preference data (e.g., 10% or 20%) to reach and even surpass the fully-trained one, without any additional training. Furthermore, EXPO also significantly improves off-the-shelf DPO/RLHF models and exhibits decent scalability across model sizes from 7B to 70B. Our work demonstrates the efficacy of model extrapolation in exploiting LLMs' capabilities, suggesting a promising direction that deserves future exploration.

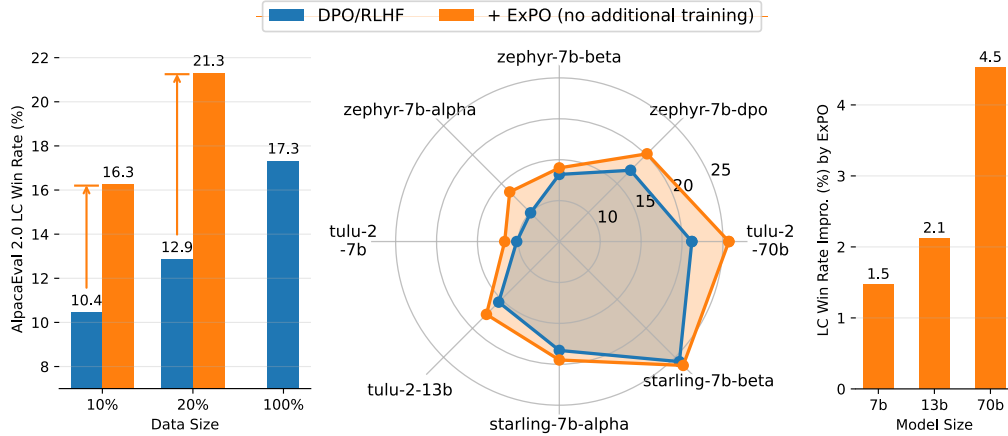


Figure 1: AlpacaEval 2.0 [22] evaluation results, where we report the length-controlled (LC) win rates [9] over the GPT-4 baseline. **Left:** With *no additional training*, EXPO boosts the models trained with less preference data to compete (10% data) and even outperform (20% data) the fully-trained model. **Middle:** EXPO also remarkably improves *off-the-shelf* DPO/RLHF models (by up to 6.8%). **Right:** The improvement by EXPO *scales up* with the increasing model size (from 7B to 70B), as verified on the Tulu-2 [18] model suite.

*Work done during Chujie's visit to UCLA. Project repository: <https://github.com/chujiezheng/LLM-Extrapolation>.

†Corresponding authors.

1 Introduction

Modern large language models (LLM) typically undergo additional fine-tuning to align with human expectations [29, 27, 28], including both supervised fine-tuning (SFT) on demonstration outputs [33, 40] and alignment training with human preference [5, 31]. Similar to the pre-training phase [15], the alignment of LLMs can also be continuously improved by increasing data and training steps [8, 5, 42]. However, in reality, alignment training is inevitably constrained by available resources and thus cannot grow indefinitely. *Suppose we have a moderately-trained LLM in hand, is it possible to further exploit its potential and cheaply acquire a stronger model?*

We draw inspiration from the literature of *model interpolation*, also known as model/weight averaging. It aims to integrate different models fine-tuned from the same base model into a unified one by interpolating between their weights [38, 19, 41], relying on the mode connectivity of neural networks [11, 10]. Previous work showed that with the basic uniform interpolation (i.e., using the same interpolation ratio for all the model modules), the obtained new model usually achieves trade-off performance between the original ones [26, 41, 23]. We similarly observe this phenomenon when we interpolate between an SFT model and a model further trained by direct preference optimization (DPO) [31] or reinforcement learning from human feedback (RLHF) [49], as shown in Figure 2.

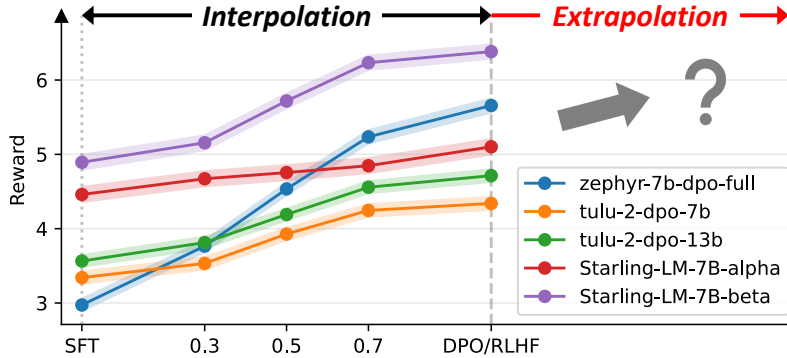


Figure 2: Calculating the reward scores (§3.1) on the UltraFeedback [7] development set, we observe that *model interpolation* usually gives trade-off performance between the two original models (e.g., an SFT model and the model further trained by DPO/RLHF), as similarly observed in previous literature [26, 41, 23]. This observation motivates our proposal of ExPO that cheaply obtains a *stronger* model from *weaker* models via *model extrapolation*.

We are intrigued by another question: *If we treat the DPO/RLHF model as an intermediate result of interpolation between the SFT model and some stronger model, can we obtain this stronger model by reversely **extrapolating** from the former two models’ weights?* If so, we can actually start with two relatively *weaker* models from the training process and straightforwardly obtain a *stronger* one. As indicated by the gray arrow in Figure 2, this could also improve off-the-shelf already-aligned models, such as the many open-sourced LLMs on HuggingFace.

Based on the above motivation, we propose a simple method called **ExPO** (*model extrapolation*) to boost LLMs’ alignment with human preference (§2). ExPO assumes a medium-aligned model $\mathcal{M}$ can be interpolated from a less-aligned (*weaker*) model $\mathcal{M}_w$ (e.g., the SFT model) and a better-aligned (*stronger*) one $\mathcal{M}_s$. Then, we can directly obtain this *stronger* model $\mathcal{M}_s$ by extrapolating from the weights of the two relatively *weaker* models $\mathcal{M}$ and $\mathcal{M}_w$, without any additional training on top of them.

Despite its simplicity, we demonstrate that ExPO is quite effective in improving the alignment of various LLMs, as summarized in Figure 1. Specifically, for standard DPO training, we show that ExPO pushes the models trained with less data (e.g., 10% or 20%) to reach and even surpass the fully-trained one, as evaluated on the AlpacaEval 2.0 benchmark [22, 9] (§3). Furthermore, ExPO also remarkably improves off-the-shelf DPO/RLHF models, by up to 6.8% on AlpacaEval 2.0 (§4), and manifests satisfactory scalability across model sizes from 7B to 70B. Our work demonstrates model extrapolation as a promising method for boosting LLMs’ alignment with human preference and better exploiting the capabilities of LLMs, which we believe deserves more future exploration.

2 Methodology

2.1 Overview

Inspired by the observation in Figure 2, we make the following assumption:

A model $\mathcal{M}$ can be interpolated between a weaker model $\mathcal{M}_w$ and a stronger model $\mathcal{M}_s$, which satisfy the relationship in terms of their alignment with human preference: $\mathcal{M}_w < \mathcal{M} < \mathcal{M}_s$.

Specifically, we suppose the medium-aligned model $\mathcal{M}$ (parameterized by θ) to be one that has been moderately trained for human preference alignment. We also suppose the less-aligned weaker model $\mathcal{M}_w$ (parameterized by θ_w) simply to be the SFT model used for initializing $\mathcal{M}$. The above assumption suggests that there exists a better-aligned stronger model $\mathcal{M}_s$ (parameterized by θ_s) and an interpolation coefficient $\gamma \in [0, 1]$ such that:

$$\theta = (1 - \gamma)\theta_w + \gamma\theta_s. \quad (1)$$

Here we consider the simplest form of uniform linear interpolation. With the substitution of $\alpha = 1/\gamma - 1 \in [0, +\infty)$, we can obtain the assumed *stronger* model $\mathcal{M}_s$ by *extrapolating* from the weights of the relatively *weaker* $\mathcal{M}_w$ and $\mathcal{M}$ (i.e., *weak-to-strong extrapolation*). Our proposed **EXPO** method is formulated as follows:

$$\theta_s = (1 + \alpha)\theta - \alpha\theta_w = \theta + \alpha(\theta - \theta_w) = \theta + \alpha\Delta\theta, \quad (2)$$

where the coefficient α serves as the hyperparameter that controls the length of extrapolation. In practice, α can be cheaply tuned as a decoding hyperparameter (like the sampling temperature) on a development set with no model training involved.

2.2 Insights on EXPO

We first use Figure 3 for an intuitive illustration of EXPO. Specifically, EXPO can be viewed as a “global gradient update”, based on the global weight change $\Delta\theta = \theta - \theta_w$ from the initial $\mathcal{M}_w$ to the final $\mathcal{M}$. The weight change $\Delta\theta$ indicates a direction in the parameter space, in which the model’s alignment with human preference is improved (measured by a reward score). Hence, EXPO essentially aims to amplify the learned reward signal through the extrapolation $\alpha\Delta\theta$.

Based on the above illustration, we identify two prerequisites for EXPO. **First**, the model $\mathcal{M}$ should have not yet been trained to its optimality. *This prerequisite is generally valid*, as evidenced by the most powerful LLMs such as GPT-4 and Claude that are undergoing constant optimization for better alignment. We will show in §4 that even the open-source models that have been extensively trained for human preference alignment still have significant room for further improvement.

Second, also more importantly, the weight change $\Delta\theta$ from $\mathcal{M}_w$ to $\mathcal{M}$ should be of “high quality”, meaning it should as accurately as possible indicate an extrapolation direction in which the alignment can get improved. In mainstream preference alignment algorithms such as DPO or RLHF, *this prerequisite can also be generally established*, as $\mathcal{M}$ is initialized from $\mathcal{M}_w$ (the SFT model) and is essentially trained to maximize the reward signal of human preference, either from preference data or reward models. Nonetheless, the “quality” of $\Delta\theta$ can vary depending on the training configuration of $\mathcal{M}$ and the capability of $\mathcal{M}_w$, as we will discuss in §3.3 and 4.2.

Combining the two prerequisites, when the model $\mathcal{M}$ initialized from its SFT checkpoint $\mathcal{M}_w$ has undergone moderate alignment training, it can potentially get better aligned by EXPO. We will experimentally verify this in §3 and 4. However, other model combinations for $\mathcal{M}_w$ and $\mathcal{M}$, such as a Base and an SFT models or two separately-trained RLHF models, usually cannot guarantee the second prerequisite. We will discuss this more in §4.3 in conjunction with empirical results.

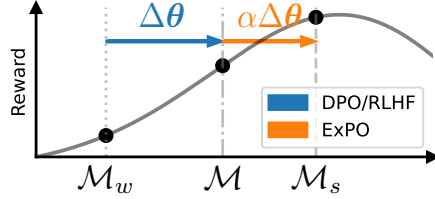


Figure 3: EXPO can be viewed as a “global gradient update” that moves the model weight along the direction of $\Delta\theta$ in which the model’s alignment with human preference is improved (measured by a reward score).

2.3 Highlights

We underline the following appealing properties of EXPO:

- **Simplicity:** EXPO is extremely simple and quick to implement. It merely involves performing extrapolation based on the weights of two checkpoints $\mathcal{M}_s$ and $\mathcal{M}$, which can be accomplished within just a few lines of code.
- **Efficiency:** EXPO needs no additional model training on top of $\mathcal{M}_s$ and $\mathcal{M}$. The only hyperparameter α is also cheap to tune as no training will be involved. Moreover, we believe more efficient means of hyperparameter search can be developed in future work, as evidenced by the advances in adaptive model interpolation [17, 23].
- **Scalability:** EXPO is in principle applicable to various LLMs, including those of large sizes or that have been extensively trained for human preference alignment. We will show in §4 that EXPO can improve off-the-shelf already-aligned models of varying sizes and capabilities.

3 Experiments

We first demonstrate the effectiveness of EXPO in a *controlled setting*, i.e., training the model $\mathcal{M}$ with less preference data, so we can know in advance that $\mathcal{M}$ *still has room for further improvement* (corresponding to the first prerequisite in §2.2). We show that EXPO endows the models trained using less data (e.g., 10% or 20%) with equivalent or superior performance to the fully-trained one.

3.1 Experimental Setup

Models To train models for human preference alignment, we refer to the alignment handbook³ [36], a widely-used code base released by HuggingFace for alignment training of LLMs. We follow their setup of training the Mistral-based [20] `zephyr-7b-sft-full` and `zephyr-7b-dpo-full` models [37]. Specifically, we use the same preference dataset but *varying data sizes* to train the models. We employ the same mainstream DPO [31] algorithm for alignment training, where the SFT model `zephyr-7b-sft-full` is used as the reference model in DPO and also used for initializing the policy models. We adopt the same hyperparameter configuration as `zephyr-7b-dpo-full` (see Appendix B) and train all the models on 4 A100 80GB GPUs. We use `zephyr-7b-dpo-full` as the fully-trained baseline (i.e., trained with 100% data).

Data We use the same preprocessed UltraFeedback⁴ [7] dataset for DPO training. UltraFeedback is a large-scale preference dataset, containing diverse instructions and response pairs with GPT-4-annotated preference labels. It has been popularly used by the open-source community for training aligned LLMs [18, 37, 48]. The preprocessed version provided by HuggingFace contains 61K and 1K preference data in the training and development set, respectively. Each data consists of an instruction and a pair of responses, with one labeled as preferred.

Evaluation We evaluate the models on AlpacaEval 2.0 [22], a leading and popular benchmark that assesses LLMs’ alignment with human preference. It contains 805 instructions representative of real user cases. For each instruction, the response of the evaluated model is compared head-to-head with that of the GPT-4 baseline. An evaluator based on GPT-4 (its version is `gpt-4-1106-preview` during our work) produces the probability of preferring the evaluated model, which provides an affordable and replicable alternative to human preference annotation. Then, the win rate over the GPT-4 baseline is computed as the expected preference probability on all the 805 instructions.

Recently, AlpacaEval 2.0 has introduced the new length-controlled (LC) win rate metric [9], which aims to alleviate the length bias of the GPT-4 evaluator (i.e., the prior preference toward longer responses) [30]. According to [9], *the LC win rate metric currently has the highest correlation (a Spearman correlation of 0.98) with real-world human evaluation* [47], which consolidates the reliability of AlpacaEval 2.0 evaluation.

³<https://github.com/huggingface/alignment-handbook>

⁴https://huggingface.co/datasets/HuggingFaceH4/ultrafeedback_binarized

For ExPO, we decide the optimal α based on the performance on the UltraFeedback *development set*, as evaluated by an open-source reward model⁵. It ranks among the top on RewardBench⁶ [21], a leaderboard that assesses the performance of reward models. More importantly, this reward model is not involved in either preference annotation or RLHF training of all the models we experiment with in this work, thus reducing the risk of reward hacking. In our experiments, we also report the average scores produced by the reward model on the 805 AlpacaEval 2.0 instructions as a reference.

3.2 Results

Table 1: AlpacaEval 2.0 evaluation results of models trained with less preference data. The DPO models are all initialized from the SFT model `zephyr-7b-sft-full`.

	Reward	Win Rate	LC Win Rate
SFT	3.42	4.7%	8.7%
DPO (full data)	6.16	14.7%	17.3%
+ ExPO, no training	6.52 (+0.36)	18.0% (+3.3%)	20.2% (+2.9%)
DPO (10% data)	3.97	5.9%	10.4%
+ ExPO, no training	6.57 (+2.60)	17.9% (+12.0%)	16.3% (+5.9%)
DPO (20% data)	4.70	8.6%	12.9%
+ ExPO, no training	6.95 (+2.25)	22.7% (+14.1%)	21.3% (+8.4%)

In Table 1, we show the performance of the models trained with less (10% and 20%) preference data as well as the results of further applying ExPO on top of them. As expected, training with less preference data results in lower-tier performance, as indicated by their LC win rates on AlpacaEval 2.0. For instance, compared to the 17.3% of using 100% data, using 10% and 20% only achieves the performance of 10.4% and 12.9%, respectively. However, after applying ExPO, using only 10% data can achieve competitive performance to the fully-trained model (16.3% vs. 17.3%), while using 20% of the data already achieves beyond that (21.3%), giving a remarkable advantage of 21.3% - 17.3% = 4%.

We also observe that the model trained with 20% data obtains a greater improvement from ExPO than that trained with 10% data (8.4% vs. 5.9%). It implies that the former gives a superior extrapolation direction $\Delta\theta$ to the latter, as illustrated in Figure 4. However, the “quality” of $\Delta\theta$ is not simply correlated with the amount of data, as shown in Table 1 where using 20% data slightly outperforms using full data when both applying ExPO (21.3% vs. 20.2%). This is because the increasing size can also amplify the biases in the preference data, which becomes more likely to be learned by the model $\mathcal{M}$ as shortcuts. We next analyze the impact of $\mathcal{M}$ ’s training configuration on $\Delta\theta$ in detail.

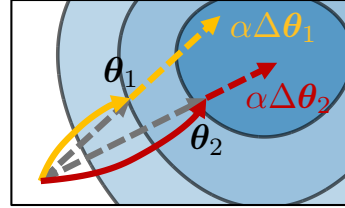


Figure 4: The “quality” of $\Delta\theta$ and the effectiveness of ExPO can vary depending on the training configurations of $\mathcal{M}$. Here, $\Delta\theta_2$ indicates a superior extrapolation direction to $\Delta\theta_1$.

3.3 Analysis

Comparison between Data Sizes Figure 5 presents the reward scores and output lengths on the UltraFeedback *development set* versus the extrapolation coefficient α values in ExPO. We have two main observations. **Firstly**, for the models $\mathcal{M}$ trained with different data sizes, the optimal α of ExPO varies and generally decreases as the data size increases, as indicated by the vertical dashed lines in the left part of Figure 5. This is because the larger data size usually leads to the more thorough convergence of training, even if only to a local optimum, which naturally narrows the viable range of the extrapolation coefficient α .

Secondly, the global optimal reward score (6.08) achieved by ExPO is obtained with a medium size (20%) of training data, rather than the smaller (5% or 10%) or larger (40%) ones. For the former (5%

⁵<https://huggingface.co/weqweasdas/RM-Mistral-7B>

⁶<https://huggingface.co/spaces/allenai/reward-bench>

and 10% data), although ExPO significantly improves the performance (from the reward score 3.13 to 4.79, and 3.59 to 5.82, respectively), *the limited data still cannot provide an accurate $\Delta\theta$* , thus capping the potential performance after model extrapolation. For the latter (40% data), we conjecture that the model may have learned the spurious features in preference data as shortcuts, especially the length bias⁷ [30] where the preferred responses are usually longer. As shown in the right part of Figure 5, for the model trained with 40% data, using a very small α results in a dramatic increase in the output length. In this case, $\Delta\theta$ *becomes more likely to contain the spurious features*, and in particular, the length bias can be amplified by model extrapolation. But this does not lead to sustained improvement of performance, as shown in the right part of Figure 5, where the optimal rewards typically correspond to moderate output lengths ranging between 500 and 600.

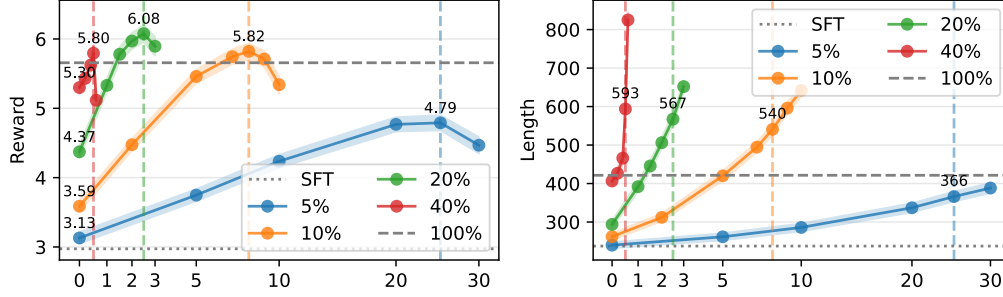


Figure 5: For the models $\mathcal{M}$ trained with varying data sizes, we plot the reward scores (**left**) and output lengths (**right**) on the UltraFeedback *development set* versus varying α values in ExPO.

Comparison with Hyperparameter Tuning As ExPO can be viewed as a “global gradient update” (§2.2), we also compare with simply tuning the training hyperparameters. Specifically, we use the same 20% training data but increase the learning rate or training epochs. From the left part of Figure 6, we observe that increasing the two hyperparameters indeed somewhat improves the original reward score. However, it is still inferior to the optimal reward score achieved by ExPO under the default configuration, and also noticeably impairs the gains from model extrapolation (the peak points are lower than that of the default configuration). This is probably because the model is *overfitted* to the training data and similarly learns the spurious features (like the length bias), thus failing to provide an accurate $\Delta\theta$. The overfitting issue can also be evidenced by the right part of Figure 6. The models trained with larger learning rates or for more epochs become prone to generating longer outputs with a small α , but do not obtain noticeable reward improvement (the left part of Figure 6), implying that $\Delta\theta$ is very likely to contain the spurious length feature rather than the true human preference.

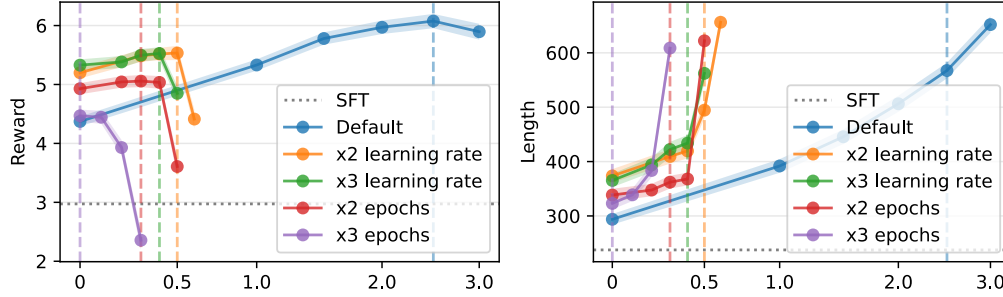


Figure 6: For the models trained using 20% data but with larger learning rates or for more epochs, we plot the reward scores (**left**) and output lengths (**right**) on the UltraFeedback *development set* versus varying α values in ExPO.

Based on the above empirical analysis, we emphasize the critical role of $\Delta\theta$ in ExPO. Particularly, we show that *the “quality” of $\Delta\theta$ requires the appropriate choice of the training configuration for $\mathcal{M}$* , including both the preference data and the training hyperparameters. In the subsequent §4.2, we will further discuss the impact of $\mathcal{M}_w$ ’s capability on the effectiveness of ExPO.

⁷The average lengths of the preferred and unpreferred responses in the UltraFeedback training set are 319 and 277 tokens, respectively.

4 Model Extrapolation Boosts Off-the-Shelf Models

We next demonstrate the impressive efficacy of EXPO in improving *off-the-shelf* already-aligned LLMs from HuggingFace, based on their SFT and DPO/RLHF checkpoints. We particularly underscore the *scalability* of EXPO across different model sizes and capabilities.

4.1 Experimental Setup

When selecting open-source LLMs for experiments, we found that many well-known aligned LLMs, such as LLaMA-2/3 [35, 1], Gemma [34], and Qwen [4], do not release the corresponding SFT checkpoints. Such an opacity hinders the feasibility of experimenting with these more representative models. To facilitate reproducible research, we select the following open-source DPO/RLHF models that (1) have also publicly accessible SFT checkpoints, (2) have disclosed the training data, and (3) are popularly downloaded on HuggingFace or have been evaluated on the AlpacaEval 2.0 leaderboard:

- `tulu-2-dpo-7/13/70b` [18], a LLaMA-2-based model suite. Since the three-sized models undergo *the same SFT and DPO training processes* (including both the data and configuration), they can serve as a reasonable testbed for the *scalability* of EXPO across different model sizes.
- `zephyr-7b-alpha/beta` and `zephyr-7b-dpo-full` [37], three Mistral-based models. They are trained with different hyperparameter configurations and on slightly different preference data.
- `Starling-LM-7B-alpha/beta` [48], two Mistral-based models. They are trained by the RLHF algorithm with different reward models.

Similar to §3.1, we select the optimal α from [0.1, 0.2, 0.3, 0.4, 0.5] based on the performance on the UltraFeedback *development set*, as evaluated by the aforementioned reward model.

4.2 Results

Table 2: AlpacaEval 2.0 evaluation results of off-the-shelf DPO/RLHF models. The **gray models’ scores** are copied from the official leaderboard for reference. For the models that have been officially evaluated, we report the higher one between our reproduced score[†] and that from the leaderboard[‡].

	Reward	Win Rate	LC Win Rate
Llama-2-70b-chat-hf	-	13.9%	17.4%
gpt-3.5-turbo-0613	-	14.1%	22.7%
Gemini Pro	-	18.2%	24.4%
claude-2.1	-	15.7%	25.3%
tulu-2-dpo-7b	5.09	8.5% [†]	10.2% [†]
+ EXPO	5.42 (+0.33)	11.5% (+3.0%)	11.7% (+1.5%)
tulu-2-dpo-13b	5.37	11.2% [†]	15.5% [†]
+ EXPO	5.89 (+0.52)	15.6% (+4.4%)	17.6% (+2.1%)
tulu-2-dpo-70b	5.84	16.0% [‡]	21.2% [‡]
+ EXPO	6.12 (+0.28)	23.0% (+7.0%)	25.7% (+4.5%)
zephyr-7b-alpha	4.68	8.4% [‡]	10.3% [‡]
+ EXPO	4.87 (+0.19)	10.6% (+2.2%)	13.6% (+3.3%)
zephyr-7b-beta	5.31	11.0% [‡]	13.2% [‡]
+ EXPO	5.40 (+0.09)	11.1% (+0.1%)	14.0% (+0.8%)
zephyr-7b-dpo-full	6.16	14.7%	17.3%
+ EXPO	6.52 (+0.36)	18.0% (+3.3%)	20.2% (+2.9%)
Starling-LM-7B-alpha	5.80	15.0% [†]	18.3% [†]
+ EXPO	5.98 (+0.18)	18.2% (+3.2%)	19.5% (+1.2%)
Starling-LM-7B-beta	7.12	26.6%	25.8%
+ EXPO	7.40 (+0.28)	29.6% (+3.0%)	26.4% (+0.6%)

The results in Table 2 demonstrate that ExPO enhances the performance of the already-aligned LLMs, by impressive increases of up to **6.8% LC win rate** and 10.5% basic win rate on AlpacaEval 2.0. The improvement is made across LLMs of various capabilities, from the weakest zephyr-7b-alpha and tulu-2-dpo-7b to the strongest Starling-LM-7B-beta and tulu-2-dpo-70b. It suggests that most open-source LLMs have not been aligned with human preference optimally, and ExPO enables the further exploitation of these models’ capabilities.

Specifically for the model suite Tulu-2, where the 7B/13B/70B models are trained using the same preference data and configuration, the enhancement by ExPO nicely *scales up* with the increasing model size. We conjecture that this is because the larger/stronger $\mathcal{M}_w$ enables the better learning of the reward signal in the preference data or reward models, leading to both a stronger $\mathcal{M}$ and a more accurate $\Delta\theta$, which together result in the greater improvement for $\mathcal{M}$ after model extrapolation. Therefore, with the same preference data and training configuration, *we optimistically expect the improvement by ExPO can also scale up as the capability of $\mathcal{M}_w$ increases.*

4.3 Discussion

Finally, we discuss the impact of model choices for $\mathcal{M}_w$ and $\mathcal{M}$ on the effectiveness of ExPO. In previous analyses and experiments, we choose $\mathcal{M}_w$ as an SFT model, and $\mathcal{M}$ as the model further trained for human preference alignment on top of $\mathcal{M}_w$. *Can other types of model combination $\mathcal{M}_w$ and $\mathcal{M}$, such as a Base and an SFT model, or two separately-trained RLHF models, be able to produce meaningful extrapolated models?* We experiment with the following types of combinations:

- **Base + SFT**: Mistral-7B-v0.1 [20] as $\mathcal{M}_w$ and Mistral-7B-Instruct-v0.1 as $\mathcal{M}$.
- **SFT 1 + SFT 2 (trained from different base models)**: Mistral-7B-Instruct-v0.1 as $\mathcal{M}_w$ and Mistral-7B-Instruct-v0.2 as $\mathcal{M}$.
- **SFT 1 + SFT 2 (same base)**: openchat_3.5 [39] as $\mathcal{M}_w$ and openchat-3.5-0106 as $\mathcal{M}$.
- **RLHF 1 + RLHF 2 (same base)**: gemma-7b-it [34] as $\mathcal{M}_w$ and gemma-1.1-7b-it as $\mathcal{M}$.
Note that it is not disclosed whether the two models are initialized from the same SFT model.

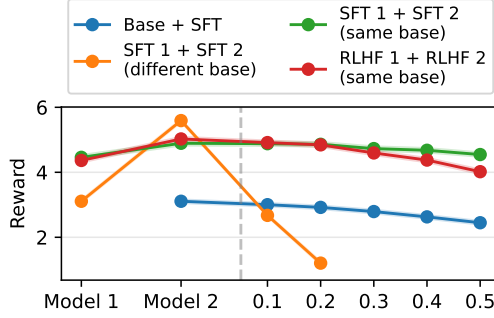


Figure 7: Reward scores of other types of model combinations on the UltraFeedback *development set*, with α varying from 0.1 to 0.5.

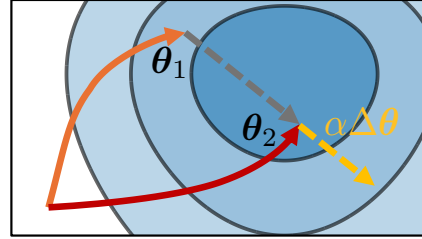


Figure 8: Extrapolation from two separately-trained models may not improve alignment, as their weight difference ($\Delta\theta$) usually cannot guarantee a direction along which the reward signal can get further amplified.

From the results shown in Figure 7, we find that extrapolating from two SFT models that are trained from different base models can easily lead to the model collapse, probably because they do not meet the requirement of mode connectivity [11, 10], namely, the same or close initialization. For the combination of Base and SFT, extrapolation degrades the performance. One cause is that the training from Base to SFT does not naturally reflect human preference, which is exactly why we need additional preference alignment training. Another cause is that compared to the Base model, the SFT one acquires the instruction-following ability and is also adapted to specified input/output formats [45]. ExPO can amplify both learned features (§2.2), but the latter does not aid in alignment and may instead similarly lead to model collapse. For the two separately-trained SFT or RLHF models, we find that they also cannot benefit from model extrapolation. We speculate that this is because $\mathcal{M}$ is not initialized from $\mathcal{M}_w$, so the path in the parameter space from θ_w to θ is not in the direction along which the reward signal can be amplified. As illustrated in Figure 8, even though $\mathcal{M}$ (θ_2) has not yet achieved optimality on its own optimization path, it still cannot be improved in another direction of

$\Delta\theta$. Overall, our method EXPO is currently applicable to the combination of an SFT model $\mathcal{M}_w$ and a model $\mathcal{M}$ further trained on top of the former, which is a very realistic combination choice, as modern LLMs that are trained to align with human preference are almost all initialized from their SFT checkpoints.

5 Related Work

LLM Alignment Modern LLMs are typically first pre-trained on massive textual corpora (resulting in a Base model) [6, 35, 1] and then trained to align with human expectations [27, 28, 35]. The alignment process generally contains two stages. In the first stage, an LLM is supervisedly fine-tuned (SFT) on *demonstration outputs* and learns to follow human instructions [40, 33]. In the second stage, the LLM is trained to learn *human preference* and assign higher probabilities to human-preferred outputs over the disfavored ones. This is usually implemented in the fashion of reinforcement learning (RL) [29, 5] or contrastive learning [44, 46, 31], as exemplified by the reinforcement learning from human feedback (RLHF) [49] and direct preference optimization (DPO) [31] algorithms, respectively. Similar to the scaling law in the pre-training phase [15], recent work also revealed that the capabilities of aligned models can also be constantly improved by scaling up the amount of alignment data [40, 33, 8] and increasing the training steps or iterations for human preference alignment [5, 42, 14]. However, the data and computation resources available in reality are always finite, which may prevent the full exploitation of models’ capabilities. Our work proposes the EXPO method to boost LLMs’ alignment with human preference in a simple, efficient, and scalable manner.

Model Merging and Interpolation Model merging is a recently focal technique for building powerful LLMs based on existing ones [2, 3]. It aims to integrate multiple models fine-tuned from the same base model into a unified one that retains the respective strengths [43, 12]. The simplest form of model merging is *model interpolation*, also known as model/weight averaging [26, 41, 23], which builds upon the mode connectivity of neural networks [11, 10]. In practice, the uniform interpolation usually results in trade-off performance between the two original models, as observed in previous literature [26, 41, 23] and our experiments in Figure 2. One approach to addressing this issue is to adaptively adjust the interpolation coefficient for different model modules (e.g., different model layers) [17, 23]. Our proposed EXPO method (§2) has a similar idea of blending model weights to improve the model capability, but works under a distinct premise and goal. Rather than integrating multiple strong models into a generalist, our method aims to use two relatively weaker models to produce a stronger model that can even surpass the limits of the fully-trained one (§3 and 4).

6 Conclusion

We present EXPO, a simple method to boost LLMs’ alignment with human preference. By extrapolating from the weights of an SFT model $\mathcal{M}_w$ and a further trained one $\mathcal{M}$, EXPO enables directly obtaining a better-aligned model without any additional training. We demonstrate the efficacy of EXPO across various LLMs, from those trained with limited preference data to the off-the-shelf ones from HuggingFace, where EXPO manifests decent scalability across varying model sizes and capabilities. Given its simplicity, efficiency, and scalability, we recommend EXPO as a promising approach for better exploiting LLMs’ capabilities, which deserves more future exploration.

Limitations & Future Work Our work is limited by the public accessibility to the checkpoints of the SFT and DPO/RLHF models. Thus unfortunately, we are unable to experiment with the more representative LLMs like LLaMA-2/3 [35, 1], Gemma [34], and Qwen [4]. We hope for more open-source efforts in increasing LLMs’ transparency and accessibility. Outside the scope of this study, there are several problems that may attract future research. First, since EXPO is based on the simplest uniform linear extrapolation (Equation 2, using the same α for all the model modules), future work may devise methods to adaptively search optimal α for different model modules. Second, while we currently rely on an external reward model for searching α , future work may get rid of such reliance by resorting to the capability of the models $\mathcal{M}$ and $\mathcal{M}_w$ themselves. Third, although our work provides intuitive illustrations for EXPO and empirically demonstrates its effectiveness, future work may establish theoretical explanations and analyses for its underlying mechanisms. Finally, it would also be interesting to apply EXPO to multi-modal LLMs like LLaVA [24] and other model architectures like Mamba [13].

Acknowledgements

We thank the open-source community, including the HuggingFace, AllenAI, and Nexusflow teams, for promoting the transparency of LLMs by releasing model checkpoints and disclosing training details. This work would not be possible without these efforts from the open-source community. We thank Wei Xiong for releasing the reward models and for the valuable discussion.

References

- [1] AI@Meta. Llama 3 model card, 2024.
- [2] Samuel Ainsworth, Jonathan Hayase, and Siddhartha Srinivasa. Git re-basin: Merging models modulo permutation symmetries. In *The Eleventh International Conference on Learning Representations*, 2023.
- [3] Takuya Akiba, Makoto Shing, Yujin Tang, Qi Sun, and David Ha. Evolutionary optimization of model merging recipes. *arXiv preprint arXiv:2403.13187*, 2024.
- [4] Jinze Bai, Shuai Bai, Yunfei Chu, Zeyu Cui, Kai Dang, Xiaodong Deng, Yang Fan, Wenbin Ge, Yu Han, Fei Huang, et al. Qwen technical report. *arXiv preprint arXiv:2309.16609*, 2023.
- [5] Yuntao Bai, Andy Jones, Kamal Ndousse, Amanda Askell, Anna Chen, Nova DasSarma, Dawn Drain, Stanislav Fort, Deep Ganguli, Tom Henighan, et al. Training a helpful and harmless assistant with reinforcement learning from human feedback. *arXiv preprint arXiv:2204.05862*, 2022.
- [6] Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. Language models are few-shot learners. In *Advances in neural information processing systems*, volume 33, pages 1877–1901, 2020.
- [7] Ganqu Cui, Lifan Yuan, Ning Ding, Guanming Yao, Wei Zhu, Yuan Ni, Guotong Xie, Zhiyuan Liu, and Maosong Sun. Ultrafeedback: Boosting language models with high-quality feedback. *arXiv preprint arXiv:2310.01377*, 2023.
- [8] Ning Ding, Yulin Chen, Bokai Xu, Yujia Qin, Shengding Hu, Zhiyuan Liu, Maosong Sun, and Bowen Zhou. Enhancing chat language models by scaling high-quality instructional conversations. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 3029–3051, Singapore, December 2023. Association for Computational Linguistics.
- [9] Yann Dubois, Balázs Galambosi, Percy Liang, and Tatsunori B Hashimoto. Length-controlled alpacaEval: A simple way to debias automatic evaluators. *arXiv preprint arXiv:2404.04475*, 2024.
- [10] Rahim Entezari, Hanie Sedghi, Olga Saukh, and Behnam Neyshabur. The role of permutation invariance in linear mode connectivity of neural networks. In *International Conference on Learning Representations*, 2022.
- [11] Timur Garipov, Pavel Izmailov, Dmitrii Podoprikin, Dmitry P Vetrov, and Andrew G Wilson. Loss surfaces, mode connectivity, and fast ensembling of dnns. In *Advances in neural information processing systems*, volume 31, 2018.
- [12] Charles Goddard, Shamane Siriwardhana, Malikeh Ehghaghi, Luke Meyers, Vlad Karpukhin, Brian Benedict, Mark McQuade, and Jacob Solawetz. Arcee’s mergekit: A toolkit for merging large language models. *arXiv preprint arXiv:2403.13257*, 2024.
- [13] Albert Gu and Tri Dao. Mamba: Linear-time sequence modeling with selective state spaces. *arXiv preprint arXiv:2312.00752*, 2023.
- [14] Shangmin Guo, Biao Zhang, Tianlin Liu, Tianqi Liu, Misha Khalman, Felipe Llinares, Alexandre Rame, Thomas Mesnard, Yao Zhao, Bilal Piot, et al. Direct language model alignment from online ai feedback. *arXiv preprint arXiv:2402.04792*, 2024.

- [15] Jordan Hoffmann, Sebastian Borgeaud, Arthur Mensch, Elena Buchatskaya, Trevor Cai, Eliza Rutherford, Diego de Las Casas, Lisa Anne Hendricks, Johannes Welbl, Aidan Clark, et al. Training compute-optimal large language models. *arXiv preprint arXiv:2203.15556*, 2022.
- [16] Ari Holtzman, Jan Buys, Li Du, Maxwell Forbes, and Yejin Choi. The curious case of neural text degeneration. In *International Conference on Learning Representations*, 2020.
- [17] Gabriel Ilharco, Marco Tulio Ribeiro, Mitchell Wortsman, Ludwig Schmidt, Hannaneh Hajishirzi, and Ali Farhadi. Editing models with task arithmetic. In *The Eleventh International Conference on Learning Representations*, 2023.
- [18] Hamish Ivison, Yizhong Wang, Valentina Pyatkin, Nathan Lambert, Matthew Peters, Pradeep Dasigi, Joel Jang, David Wadden, Noah A Smith, Iz Beltagy, et al. Camels in a changing climate: Enhancing lm adaptation with tulu 2. *arXiv preprint arXiv:2311.10702*, 2023.
- [19] P Izmailov, AG Wilson, D Podoprikin, D Vetrov, and T Garipov. Averaging weights leads to wider optima and better generalization. In *34th Conference on Uncertainty in Artificial Intelligence 2018, UAI 2018*, pages 876–885, 2018.
- [20] Albert Q Jiang, Alexandre Sablayrolles, Arthur Mensch, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Florian Bressand, Gianna Lengyel, Guillaume Lample, Lucile Saulnier, et al. Mistral 7b. *arXiv preprint arXiv:2310.06825*, 2023.
- [21] Nathan Lambert, Valentina Pyatkin, Jacob Morrison, LJ Miranda, Bill Yuchen Lin, Khyathi Chandu, Nouha Dziri, Sachin Kumar, Tom Zick, Yejin Choi, et al. Rewardbench: Evaluating reward models for language modeling. *arXiv preprint arXiv:2403.13787*, 2024.
- [22] Xuechen Li, Tianyi Zhang, Yann Dubois, Rohan Taori, Ishaan Gulrajani, Carlos Guestrin, Percy Liang, and Tatsunori B. Hashimoto. AlpacaEval: An automatic evaluator of instruction-following models. https://github.com/tatsu-lab/alpaca_eval, 2023.
- [23] Yong Lin, Hangyu Lin, Wei Xiong, Shizhe Diao, Jianmeng Liu, Jipeng Zhang, Rui Pan, Haoxiang Wang, Wenbin Hu, Hanning Zhang, Hanze Dong, Renjie Pi, Han Zhao, Nan Jiang, Heng Ji, Yuan Yao, and Tong Zhang. Mitigating the alignment tax of rlhf, 2023.
- [24] Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. Visual instruction tuning. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- [25] Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. In *International Conference on Learning Representations*, 2019.
- [26] Behnam Neyshabur, Hanie Sedghi, and Chiyuan Zhang. What is being transferred in transfer learning? In *Advances in neural information processing systems*, volume 33, pages 512–523, 2020.
- [27] OpenAI. <https://chat.openai.com/chat>, 2022.
- [28] OpenAI. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- [29] Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. Training language models to follow instructions with human feedback. *Advances in neural information processing systems*, 35:27730–27744, 2022.
- [30] Ryan Park, Rafael Rafailov, Stefano Ermon, and Chelsea Finn. Disentangling length from quality in direct preference optimization. *arXiv preprint arXiv:2403.19159*, 2024.
- [31] Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. Direct preference optimization: Your language model is secretly a reward model. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- [32] Samyam Rajbhandari, Jeff Rasley, Olatunji Ruwase, and Yuxiong He. Zero: memory optimizations toward training trillion parameter models. In *Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis*, pages 1–16, 2020.

- [33] Rohan Taori, Ishaan Gulrajani, Tianyi Zhang, Yann Dubois, Xuechen Li, Carlos Guestrin, Percy Liang, and Tatsunori B. Hashimoto. Stanford alpaca: An instruction-following llama model. https://github.com/tatsu-lab/stanford_alpaca, 2023.
- [34] Gemma Team, Thomas Mesnard, Cassidy Hardin, Robert Dadashi, Surya Bhupatiraju, Shreya Pathak, Laurent Sifre, Morgane Rivière, Mihir Sanjay Kale, Juliette Love, et al. Gemma: Open models based on gemini research and technology. *arXiv preprint arXiv:2403.08295*, 2024.
- [35] Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*, 2023.
- [36] Lewis Tunstall, Edward Beeching, Nathan Lambert, Nazneen Rajani, Shengyi Huang, Kashif Rasul, Alexander M. Rush, and Thomas Wolf. The alignment handbook. <https://github.com/huggingface/alignment-handbook>, 2023.
- [37] Lewis Tunstall, Edward Beeching, Nathan Lambert, Nazneen Rajani, Kashif Rasul, Younes Belkada, Shengyi Huang, Leandro von Werra, Clémentine Fourrier, Nathan Habib, et al. Zephyr: Direct distillation of lm alignment. *arXiv preprint arXiv:2310.16944*, 2023.
- [38] Joachim Utans. Weight averaging for neural networks and local resampling schemes. In *Proc. AAAI-96 Workshop on Integrating Multiple Learned Models*. AAAI Press, pages 133–138. Citeseer, 1996.
- [39] Guan Wang, Sijie Cheng, Xianyu Zhan, Xiangang Li, Sen Song, and Yang Liu. Openchat: Advancing open-source language models with mixed-quality data. In *The Twelfth International Conference on Learning Representations*, 2024.
- [40] Yizhong Wang, Yeganeh Kordi, Swaroop Mishra, Alisa Liu, Noah A. Smith, Daniel Khoshabi, and Hannaneh Hajishirzi. Self-instruct: Aligning language models with self-generated instructions. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 13484–13508, Toronto, Canada, July 2023. Association for Computational Linguistics.
- [41] Mitchell Wortsman, Gabriel Ilharco, Samir Ya Gadre, Rebecca Roelofs, Raphael Gontijo-Lopes, Ari S Morcos, Hongseok Namkoong, Ali Farhadi, Yair Carmon, Simon Kornblith, et al. Model soups: averaging weights of multiple fine-tuned models improves accuracy without increasing inference time. In *International conference on machine learning*, pages 23965–23998. PMLR, 2022.
- [42] Wei Xiong, Hanze Dong, Chenlu Ye, Ziqi Wang, Han Zhong, Heng Ji, Nan Jiang, and Tong 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*, 2024.
- [43] Le Yu, Bowen Yu, Haiyang Yu, Fei Huang, and Yongbin Li. Language models are super mario: Absorbing abilities from homologous models as a free lunch. *arXiv preprint arXiv:2311.03099*, 2023.
- [44] Yao Zhao, Mikhail Khalman, Rishabh Joshi, Shashi Narayan, Mohammad Saleh, and Peter J Liu. Calibrating sequence likelihood improves conditional language generation. In *The Eleventh International Conference on Learning Representations*, 2023.
- [45] Chujie Zheng. Chat templates for huggingface large language models. https://github.com/chujiezheng/chat_templates, 2024.
- [46] Chujie Zheng, Pei Ke, Zheng Zhang, and Minlie Huang. Click: Controllable text generation with sequence likelihood contrastive learning. In *Findings of the Association for Computational Linguistics: ACL 2023*, pages 1022–1040, Toronto, Canada, July 2023. Association for Computational Linguistics.

- [47] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, Hao Zhang, Joseph E. Gonzalez, and Ion Stoica. Judging LLM-as-a-judge with MT-bench and chatbot arena. In *Thirty-seventh Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2023.
- [48] Banghua Zhu, Evan Frick, Tianhao Wu, Hanlin Zhu, Karthik Ganesan, Wei-Lin Chiang, Jian Zhang, and Jiantao Jiao. Starling-7b: Improving llm helpfulness & harmlessness with rlaiif, November 2023.
- [49] Daniel M Ziegler, Nisan Stiennon, Jeffrey Wu, Tom B Brown, Alec Radford, Dario Amodei, Paul Christiano, and Geoffrey Irving. Fine-tuning language models from human preferences. *arXiv preprint arXiv:1909.08593*, 2019.

A Open-Source Models Used in This Work

Model	URL
RM-Mistral-7B	https://huggingface.co/weqweasdas/RM-Mistral-7B
tulu-2-7b	https://huggingface.co/allenai/tulu-2-7b
tulu-2-dpo-7b	https://huggingface.co/allenai/tulu-2-dpo-7b
tulu-2-13b	https://huggingface.co/allenai/tulu-2-13b
tulu-2-dpo-13b	https://huggingface.co/allenai/tulu-2-dpo-13b
tulu-2-70b	https://huggingface.co/allenai/tulu-2-70b
tulu-2-dpo-70b	https://huggingface.co/allenai/tulu-2-dpo-70b
mistral-7b-sft-alpha	https://huggingface.co/HuggingFaceH4/mistral-7b-sft-alpha
zephyr-7b-alpha	https://huggingface.co/HuggingFaceH4/zephyr-7b-alpha
mistral-7b-sft-beta	https://huggingface.co/HuggingFaceH4/mistral-7b-sft-beta
zephyr-7b-beta	https://huggingface.co/HuggingFaceH4/zephyr-7b-beta
zephyr-7b-sft-full	https://huggingface.co/alignment-handbook/zephyr-7b-sft-full
zephyr-7b-dpo-full	https://huggingface.co/alignment-handbook/zephyr-7b-dpo-full
openchat_3.5	https://huggingface.co/openchat/openchat_3.5
Starling-LM-7B-alpha	https://huggingface.co/berkeley-nest/Starling-LM-7B-alpha
openchat-3.5-0106	https://huggingface.co/openchat/openchat-3.5-0106
Starling-LM-7B-beta	https://huggingface.co/Nexusflow/Starling-LM-7B-beta
Mistral-7B-v0.1	https://huggingface.co/mistralai/Mistral-7B-v0.1
Mistral-7B-Instruct-v0.1	https://huggingface.co/mistralai/Mistral-7B-Instruct-v0.1
Mistral-7B-Instruct-v0.2	https://huggingface.co/mistralai/Mistral-7B-Instruct-v0.2
gemma-7b-it	https://huggingface.co/google/gemma-7b-it
gemma-1.1-7b-it	https://huggingface.co/google/gemma-1.1-7b-it

B More Implementation Details

For model training in §3, we adopt the global batch size 128 and gradient accumulation steps 4. We train the models on 4 A100 80GB GPUs, with ZeRO-3 offload [32] and gradient checkpointing for reducing GPU memory usage. We set the learning rate to $5e-7$, with the cosine scheduling and

warmup ratio of 0.1, and use the AdamW [25] optimizer to train the models for one epoch. For DPO, we follow `zephyr-7b-dpo-full` and set the coefficient β to 0.01.

For response generation, we employ the `vllm` library for high-throughput inference. We use top- k ($k = 40$) and nucleus sampling [16] ($p = 0.9$) with a temperature of 0.7. To avoid repetition in generated texts, we set both the factors of presence penalty and frequency penalty to 0.1.

For hyperparameter search, we manually try different values of α . We use the obtained model to generate responses on the UltraFeedback development set, score the responses with the reward model, and choose the optimal α corresponding to the highest average score. We list in Table 3 the search range and the optimal α in our experiments.

Table 3: Summary of hyperparameter search results.

SFT Model	DPO/RLHF Model	Search Range	Optimal α
<code>zephyr-7b-sft-full</code>	DPO (10% data)	[2, 5, 7, 8, 9, 10]	8
<code>zephyr-7b-sft-full</code>	DPO (20% data)	[1.0, 2.0, 2.5, 3.0]	2.5
<code>tulu-2-7b</code>	<code>tulu-2-dpo-7b</code>	[0.1, 0.3, 0.4, 0.5]	0.5
<code>tulu-2-13b</code>	<code>tulu-2-dpo-13b</code>	[0.1, 0.3, 0.4, 0.5]	0.5
<code>tulu-2-70b</code>	<code>tulu-2-dpo-70b</code>	[0.1, 0.3, 0.4, 0.5]	0.5
<code>mistral-7b-sft-alpha</code>	<code>zephyr-7b-alpha</code>	[0.1, 0.2, 0.3, 0.4, 0.5]	0.3
<code>mistral-7b-sft-beta</code>	<code>zephyr-7b-beta</code>	[0.1, 0.2, 0.3, 0.5]	0.1
<code>zephyr-7b-sft-full</code>	<code>zephyr-7b-dpo-full</code>	[0.1, 0.2, 0.3, 0.4, 0.5]	0.3
<code>openchat_3.5</code>	<code>Starling-LM-7B-alpha</code>	[0.1, 0.2, 0.3, 0.5]	0.2
<code>openchat-3.5-0106</code>	<code>Starling-LM-7B-beta</code>	[0.1, 0.3, 0.4, 0.5]	0.5