

Debiasing Multimodal Large Language Models via Noise-Aware Preference Optimization

Zefeng Zhang^{1,2} Hengzhu Tang³ Jiawei Sheng^{1,2} Zhenyu Zhang^{3*}
 Yiming Ren^{1,2} Zhenyang Li³ Dawei Yin³ Duohe Ma^{1,2†} Tingwen Liu^{1,2†}

¹Institute of Information Engineering, Chinese Academy of Sciences

²School of Cyber Security, University of Chinese Academy of Sciences, ³Baidu Inc.

{zhangzefeng, shengjiawei, renyiming, maduohe, liutingwen}@iie.ac.cn

{tanghengzhu, zhangzhenyu07, zhenyounglee, yindawei}@baidu.com

Abstract

Multimodal Large Language Models (MLLMs) excel in various tasks, yet often struggle with modality bias, where the model tends to rely heavily on a single modality and overlook critical information in other modalities, which leads to incorrect focus and generating irrelevant responses. In this paper, we propose using the paradigm of preference optimization to solve the modality bias problem, including RLAIF-V-Bias, a debiased preference optimization dataset, and a Noise-Aware Preference Optimization (NaPO) algorithm. Specifically, we first construct the dataset by introducing perturbations to reduce the informational content of certain modalities, compelling the model to rely on a specific modality when generating negative responses. To address the inevitable noise in automatically constructed data, we combine the noise-robust Mean Absolute Error (MAE) with the Binary Cross-Entropy (BCE) in Direct Preference Optimization (DPO) by a negative Box-Cox transformation, and dynamically adjust the algorithm’s noise robustness based on the evaluated noise levels in the data. Extensive experiments validate our approach, demonstrating not only its effectiveness in mitigating modality bias but also its significant role in minimizing hallucinations. The code and data is available at <https://github.com/zhangzef/NaPO>.

1. Introduction

Multimodal Large Language Models (MLLMs) have achieved powerful multimodal understanding capabilities [46, 61, 64] through pretraining on large-scale image-text data, significantly advancing AI research [2, 4, 31, 33, 34, 56, 69]. These MLLMs fuse large-scale pre-trained vision models into the representation space of the Large Lan-

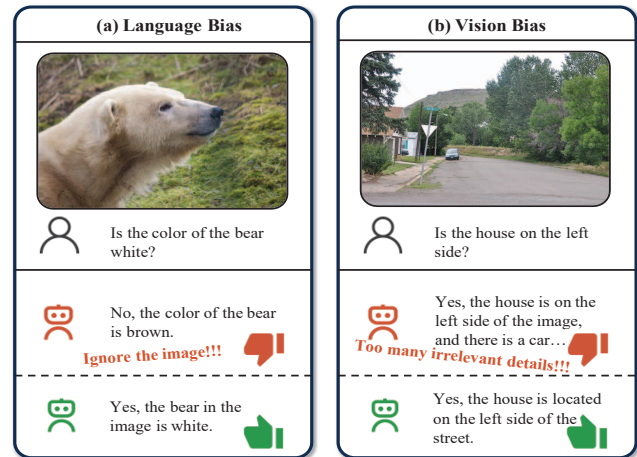


Figure 1. Examples of different types of modality-biased responses and their preferred counterparts. **Left:** The model relies excessively on prior knowledge, assuming a bear is brown while overlooking the image, which shows a polar bear. **Right:** Although the model answers the question correctly, it provides unnecessary image details that are irrelevant to the question.

guage Models (LLMs), allowing the LLMs access to the visual representations. However, MLLMs continue to struggle with the *modality bias* [11, 28, 62], where the model tends to rely heavily on one of the involved modalities and overlook critical information from other modalities, leading to incorrect focuses and irrelevant responses. Specifically, in models with both text and image inputs, modality biases are mainly manifested as the *language bias* and *vision bias*. Figure 1 (a) shows an example of the *language bias*: for the question “Is the color of the bear white?”, the model overly relies on the language priors [70] that most bears are brown, and overlooks the critical information of “polar bear” with white color in the input image. This bias inevitably leads to incorrect responses and would be even worse when larger

*Project lead.

†Corresponding author.

LLMs dominate MLLMs. Besides, Figure 1 (b) shows an example of the *vision bias*: when asking “Is the house on the left side?”, the model focuses too much on image details, resulting in a lack of accurate understanding of the textual question. Such modality bias would add trivial irrelevant information in model responses, exacerbating the challenge of instruct following [67].

An ideal MLLM ought to be modality-unbiased, effectively integrating information from all modalities to provide accurate and complete answers [62]. Existing studies have attempted balanced dataset distribution for training [19, 27], or devise strategies to identifying and mitigating bias during training or inference [24, 38, 49]. However, these studies mostly require additional large-scale supervised fine-tuning, which risks losing valuable existing knowledge in MLLMs [16, 51]. In contrast, we notice that debiasing MLLMs can be seen as a *preference optimization* problem [37] in LLMs, which can enhance alignment with human preferences by increasing the probability gap between preferred (unbiased) and non-preferred (biased) generated responses. In other words, by increasing the generation probability of unbiased responses over biased ones, the model is expected to incorporate critical information from all modalities, thus alleviating the bias on a single modality. In this way, this debiasing design adjusts the model preference on response generation, yet remains most existing knowledge in MLLMs with original capabilities. Additionally, it is still challenging to derive high-quality datasets for preference optimization. To our knowledge, there are few preference optimization datasets specifically for MLLM debiasing, and automatically constructed datasets often contain significant noise since biased responses, though present, are not always of low quality. This makes standard preference optimization algorithms struggle when facing potentially noisy (incorrect) unbiased and biased preference optimization data.

To this end, starting from the queries in RLAIF-V [58], we design a data construction method to generate biased data by perturbing other modalities to prompt the model to rely excessively on a single modality and produce biased responses. Specifically, we generate language-biased and vision-biased responses by selectively masking visual and textual information in the input, attach such biased responses to RLAIF-V as negative samples, and finally achieve a new preference optimization dataset with modality bias, termed RLAIF-V-Bias. Besides, we propose introducing Noise-Aware Preference Optimization (NaPO) to dynamically identify noisy data and reduce optimization weights for these samples. Specifically, NaPO builds on Generalized Cross Entropy [63] by applying a negative Box-Cox transformation [6], combining noise-robust Mean Absolute Error (MAE)[18] with Binary Cross-Entropy (BCE) from Direct Preference Optimization (DPO)¹. NaPO’s noise robustness coefficient is dynamically

adjusted by assessing the noise level of training samples.

To evaluate effectiveness, we take LLaVA-v1.5-7b [33] as the base model to generate negative samples and perform preference optimizations, following the setup of RLAIF-V [58]. Next, we evaluate the proposed RLAIF-V-Bias dataset and the NaPO algorithm on VLind-Bench [28] (a benchmark for language priors and commonsense biases in MLLMs) as well as on common hallucination benchmarks: Object HalBench [40], MMHalBench [44], and AMBER [50]. Compared to the original training set and DPO, our approach showed approximately 19.5% and 18.6% improvements in reducing bias and language priors, with a notable reduction in hallucinations relative to the baseline. We also tested the effectiveness of our method and dataset on models with varying parameter sizes.

Our contributions are threefold: First, we develop a data construction method based on the causes of modality bias and practice it to create a debiasing-oriented preference optimization dataset RLAIF-V-Bias. Second, we propose NaPO, which applies the negative Box-Cox transformation to DPO, enabling it to adjust the loss function’s noise robustness, and we further design a dynamic noise assessment method that allows NaPO to adapt its noise robustness dynamically during training based on data analysis. Lastly, we demonstrate through experiments on language-prior and hallucination benchmarks that our method effectively mitigates the modality bias problem in MLLMs.

2. Preliminaries

2.1. Preference Optimization

Preference optimization aims to align LLMs with human preferences, thereby enhancing their responsiveness to human needs. The RL-based preference optimization method first trains a Supervised Fine-Tuning (SFT) model using human-labeled preference data to obtain a reward model. Then, it simulates the environment through the reward model and uses algorithms such as Proximal Policy Optimization (PPO) to maximize the LMs’ reward, achieving alignment with human preferences [13]. Due to the complexity and high resource consumption of RL-based methods, many studies focus on designing simplified loss functions that enable LLMs to align with human preferences directly using this loss function and human-labeled preference data. DPO [39] is one of the promising approaches. Given an input x and a response y , DPO defines its reward $r(x, y)$ as:

$$r(x, y) = \beta \log \frac{\pi_{\theta}(y|x)}{\pi_{ref}(y|x)} + \beta \log Z(x), \quad (1)$$

where $Z(x)$ is a partition function, β is a hyperparameter that controls the deviation from the reference model, π_{θ} and π_{ref} denote the policy model and the reference model, respectively. Given a preference optimization sample (x, y_w, y_l) ,

¹ See Section 2.2 for a detailed comparison of the two loss functions.

where y_w denotes the preferred response and y_l the rejected response, DPO aligns the LLM with human values by maximizing the reward margin between y_w and y_l based on the Bradley-Terry model [7]:

$$\begin{cases} \psi_\Sigma(x, y_w, y_l) = \beta \log \frac{\pi_\theta(y_w|x)}{\pi_{ref}(y_w|x)} - \beta \log \frac{\pi_\theta(y_l|x)}{\pi_{ref}(y_l|x)}, \\ \psi_\mu(x, y_w, y_l) = \frac{\beta}{|y_w|} \log \frac{\pi_\theta(y_w|x)}{\pi_{ref}(y_w|x)} - \frac{\beta}{|y_l|} \log \frac{\pi_\theta(y_l|x)}{\pi_{ref}(y_l|x)}, \end{cases} \quad (2)$$

where ψ_Σ denotes calculating the reward margin using the sum of log probabilities (logP), while ψ_μ denotes calculating the reward margin using the average logP. $|y|$ denotes the length of the token sequence. With the default setting of ψ_Σ , DPO uses a BCE loss to enhance the reward difference between the y_w and y_l for LLMs:

$$\mathcal{L}_{DPO} = -\log \sigma \left(\beta \log \frac{\pi_\theta(y_w|x)}{\pi_{ref}(y_w|x)} - \beta \log \frac{\pi_\theta(y_l|x)}{\pi_{ref}(y_l|x)} \right). \quad (3)$$

2.2. Noise Robustness Analysis

A loss function is considered noise-robust if it minimizes risk similarly with both noisy and noise-free labels [18]. This means the loss function can suppress or reduce the impact of noisy data on the optimization process. As a commonly used classification loss, BCE converges quickly but is prone to overfitting on noisy data, whereas MAE is noise-robust but converges slowly, which can result in undertrained models. Here we compare these two loss functions from the perspectives of symmetric loss and gradients. Now, let's consider a simple binary classification scenario. For a training sample $\{x, y\}$, $y \in \{0, 1\}$, MAE and BCE can be formalized as:

$$\begin{cases} \mathcal{L}_{MAE} = |y - f(x)|, \\ \mathcal{L}_{BCE} = y \log(f(x)) + (1 - y) \log(1 - f(x)). \end{cases} \quad (4)$$

For the first perspective, Ghosh et al. [18] have proven that symmetric loss functions exhibit superior noise robustness. A loss function $\mathcal{L}$ is a symmetric loss if and only if it satisfies the following equation for any x and classifier f [18]:

$$\mathcal{L}(f(x), y) + \mathcal{L}(f(x), 1 - y) = C, \quad (5)$$

where C is a constant. When the loss function $\mathcal{L}$ is MAE, $C = 1$. However, when the loss function $\mathcal{L}$ is BCE, the result of the Equation (5) is $-(\log(f(x)) + \log(1 - f(x)))$, which is not a constant. Therefore, MAE has better noise robustness compared to BCE. From the perspective of gradients, the gradients of the two are as follows:

$$\frac{\partial \mathcal{L}(f_\theta(x), y)}{\partial \theta} = \begin{cases} -\frac{1}{f_\theta(x)} \nabla_\theta f_\theta(x) & \text{for BCE,} \\ -\nabla_\theta f_\theta(x) & \text{for MAE.} \end{cases} \quad (6)$$

So the smaller $f_\theta(x)$ or larger $\frac{1}{f_\theta(x)}$, are implicitly weighed more than samples with predictions that agree more with

provided labels in the gradient update. This implies that in BCE training, greater emphasis is placed on harder samples, helping the model to converge quickly but potentially causing it to overfit on noisy data. In summary, while BCE *converges quickly*, it tends to overfit on noisy data, whereas MAE is *noise-robust* but converges more slowly, which may lead to suboptimal training performance [63].

3. Method

This section details modality bias mitigation in MLLMs within the Bradley-Terry preference optimization framework, achieved by reducing biased response probabilities and increasing ground truth likelihoods. The approach consists of two main components: biased response generation and noise-aware preference optimization. In Section 3.1, we describe the biased response generation process, while Section 3.2 covers noise-aware preference optimization and includes an analysis of data for noise distribution to guide bias reduction.

3.1. Modality Biase Response Generation

A biased response is one where the model disproportionately relies on a single modality. To encourage this, we introduce controlled disturbances to other modalities, reducing their informational weight and increasing reliance on the target modality. We generate biased responses using the RLAI-F-V dataset [58], a multimodal preference optimization dataset built through iterative feedback from open-source models.

Language-biased response generation. Language bias refers to the model's excessive reliance on prior knowledge or textual information when generating answers. For a multimodal input $x = (v, t)$, where v is the visual information and t is the instruction and textual context, we reduce the model's reliance on the visual modality to encourage language-biased responses. Specifically, by masking out all visual information, we guide the model to produce responses that depend primarily on textual information and prior knowledge:

$$y_{lb} = MLLM([MASK]; t), \quad (7)$$

where y_{lb} is the language-biased response, $[MASK]$ denotes the mask tokens used to mask all visual information.

Vision-biased response generation. Visual bias occurs when the model focuses too heavily on visual information, producing irrelevant image details in its responses. For a multimodal input (v, t) , we reduce the model's reliance on textual input to encourage vision-biased responses. Similar to the language-biased response generation, by masking out all textual information, we guide the model to generate responses that depend mainly on visual information:

$$y_{vb} = MLLM(v; [MASK]), \quad (8)$$

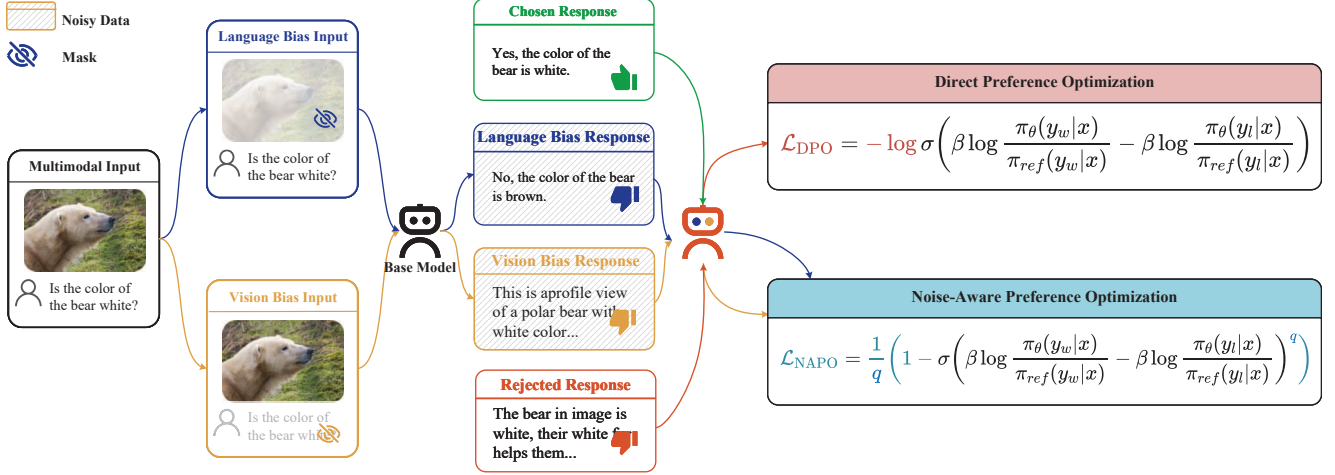


Figure 2. **Method details.** First, biased responses are constructed by using masking to guide the model toward over-relying on prompts and generating responses based on the base model. Next, NaPO is applied for noise-robust preference optimization to counteract noise in automatically constructed data, dynamically assessing data noise levels to calculate NaPO’s noise robustness coefficient q (see Equation (12)). Here we assumed that the original data is of high quality, so DPO is used to train on it directly. Additional experiments were conducted with NaPO on the original data, and the results can be found in Appendix A.

where y_{vb} refers to vision-biased response, and $[MASK]$ denotes the mask tokens used to mask all textual information.

3.2. Noise-Aware Preference Optimization

Automatically constructed data inevitably contain noise, allowing some unbiased responses to appear among the generated biased ones, which adversely affects model training. Standard preference optimization methods, however, lack robustness to noise and incorporating noise robustness often compromises model convergence [63]. In this section, we propose a NaPO algorithm. We first start by combining the BCE from DPO with the noise-robust MAE using a negative Box-Cox transformation. Then, based on data analysis, we introduce an adaptive noise-aware method to dynamically adjust the noise robustness coefficient in NaPO.

For a given preference optimization training sample (x, y_w, y_l) , DPO uses BCE loss to fit the reward margin activated by the sigmoid function to 1, enhancing the reward difference between the preferred sample y_w and the non-preferred sample y_l . As shown in Section 2.2, BCE converges quickly but can be prone to overfitting, whereas MAE is highly noise-robust but converges more slowly. To combine these advantages, we use the negative Box-Cox transformation as the loss to achieve both fast convergence (from BCE) and noise robustness (from MAE):

$$\mathcal{L}_{\text{NaPO}}(x, y_w, y_l) = \frac{1}{q} \left(1 - \sigma \left(\beta \log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \beta \log \frac{\pi_{\theta}(y_l|x)}{\pi_{\text{ref}}(y_l|x)} \right) \right)^q, \quad (9)$$

where $q \in (0, 1]$ is the noise robustness coefficient. By L’Hôpital’s rule, as q approaches 0, the function’s limit converges to the BCE:

$$\lim_{q \rightarrow 0} \frac{1}{q} (1 - x^q) = -\log(x). \quad (10)$$

Therefore, when $q \in (0, 1]$, the range of values for $\mathcal{L}_{\text{NaPO}}$ is:

$$\mathcal{L}_{\text{MAE}} \leq \mathcal{L}_{\text{NaPO}} < \mathcal{L}_{\text{BCE}}. \quad (11)$$

We can observe the differences between the various loss functions in Figure 3. As the value of q increases, the NaPO becomes more similar to MAE, resulting in stronger noise robustness but slower convergence. Conversely, as q decreases, NaPO approaches BCE, which reduces noise robustness but accelerates convergence. This also indicates that the upper bound of NaPO is DPO.

Adaptive noise robustness coefficient (q). The long-tail issue in LLMs’ knowledge [26] causes varying degrees of biases in responses to common and rare questions, making a fixed noise robustness coefficient insufficient and less flexible. For instance, the model may answer correctly about a brown bear’s color without visual input (coincidentally unbiased response), and produce an incorrect response when asked about a polar bear (biased response). In this section, we would like to design a dynamic noise robustness coefficient for biased samples based on data observations.

For a given input x , the sum of log probabilities of the answer y can reflect the LLM’s confidence, where a smaller log probability margin indicates greater similarity to the ground

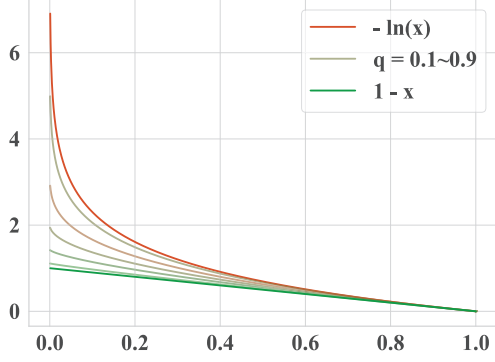


Figure 3. **Comparison of different loss functions.** We plotted the function $(1 - x^q)q^{-1}$ for values of q in the range (0.1, 0.3, 0.5, 0.7, 0.9), and compared it with both MAE $(1 - x)$ and BCE $-\ln(x)$. By adjusting the value of q , we can balance the noise robustness and the rapid convergence ability of NaPO.

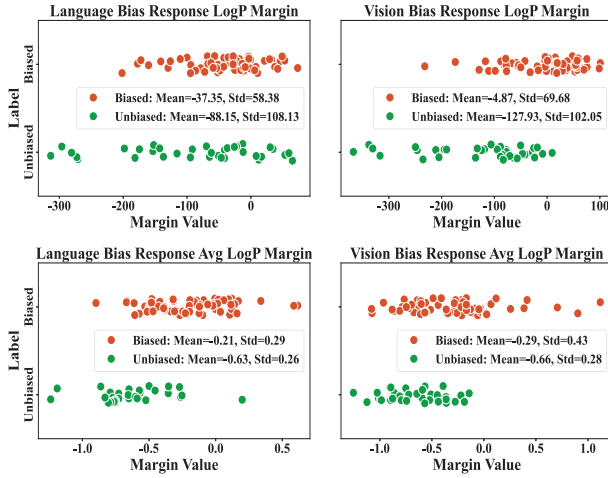


Figure 4. **Analysis of noise and margin distribution in automatically constructed data.** The first row shows LogP margins between each biased response type and the ground truth, while the second row shows avg LogP margins. In language-biased responses, biased (noise-free) data have a higher avg LogP margin than unbiased (noise) data. Similarly, in vision-biased responses, biased (noise-free) data show a higher LogP margin than unbiased (noise) data.

truth. We randomly sample 100 data points and manually annotate each automatically generated response as either biased (noise-free) or unbiased (noise) for ease of observation. Figure 4 shows that, for language-biased responses, the avg LogP margin of noise data is generally smaller than that of noise-free data. Similarly, in vision-biased responses, the LogP margin for noise data is smaller than that for noise-free data². Since a larger coefficient q indicates stronger

noise robustness, we observe an inverse trend, i.e., q and the margin value ought to be generally inversely proportional. Based on the above observations, we dynamically derive the coefficient q as follows:

$$q = 1 - \sigma(\alpha\psi(x, y_w, y_l)), \quad (12)$$

where $\psi(x, y_w, y_l)$ is the reward margin formula ψ in Equation (2), and α is a scaling factor that normalizes the reward margin within the high-gradient range of the sigmoid function for effective reward capture. We use the ψ_μ to calculate q for language-biased responses and the ψ_Σ for vision-biased responses based on the above observation. To ensure training stability, we use the batch-average reward margin to calculate ψ and finally derive q .

Final optimization objective. From Figure 4, we observe that the margin value is positively correlated with data quality: noisy samples typically have smaller margins, while noise-free samples generally have larger margins. Therefore, we introduce a margin-based dynamic weight in the final optimization objective to balance the loss functions:

$$\gamma_i = \frac{\psi_\Sigma(x, y_w, y_i)}{\sum_i \psi_\Sigma(x, y_w, y_i)}, \quad i \in [y_l; y_{lb}; y_{vb}], \quad (13)$$

where $\psi_\Sigma(x, y_w, y_i)$ is the calculation of the reward margin using the sum of LogP in Equation (2). Based on, we obtain the final optimization objective as follows:

$$\begin{aligned} \mathcal{L}_\gamma = & \gamma_{y_l} \mathcal{L}_{\text{DPO}}(x, y_w, y_l) \\ & + \gamma_{y_{lb}} \mathcal{L}_{\text{NaPO}}(x, y_w, y_{lb}) \\ & + \gamma_{y_{vb}} \mathcal{L}_{\text{NaPO}}(x, y_w, y_{vb}), \end{aligned} \quad (14)$$

where $\mathcal{L}_\gamma$ means the weighted sum of the $\mathcal{L}_{\text{NaPO}}$ and $\mathcal{L}_{\text{DPO}}$ with γ . In this way, we adaptively adjust the weights of the loss function to balance the optimization of the general preference, language bias, and vision bias.

4. Experiments

4.1. Experimental Setup

In this section, we briefly introduce the implementation details, baselines, and evaluation settings.

Implementation details. Following RLAIF-V [58], we use LLaVA-v1.5-7B as the backbone model to construct our training dataset RLAIF-V-Bias, which is an extension of RLAIF-V, containing both the original data from RLAIF-V and additional bias data. Consistent with RLAIF-V settings, we set β to 0.1, used a learning rate of 5e-7, trained for 4 epochs, and set the batch size to 4. On an 8xA100 80GB

Table 4.

²A detailed analysis of this phenomenon is provided in the Appendix B.2, and experimental evidence supporting this conclusion can be found in

machine, data construction took 24 hours, and training took 7 hours. We also tested our method on LLaVA-v1.5-13B, using the same parameters as for LLaVA-v1.5-7B but training for only one epoch. For NaPO, we calculate the noise robustness coefficient q for language-biased data using the reward margin based on avg logp with α set to 0.5. For vision-biased data, we calculate q using the reward margin based on logp, with α set to 0.01. For dynamic loss weights, we derive the values of a , b , and c based on the logp reward margin during training. In practice, both q and the dynamic loss weight a , b , c are truncated to the range $[0.01, 1]$.

Baseline approaches. We primarily compare our method with standard DPO [39]. Although MDPO [48] is designed for hallucination issues, its approach to addressing language bias is similar to ours; thus, we reproduced MDPO using default settings for comparison. We also include results from other multimodal LLMs—GPT-4V [2], LLaVA v1.5-13B [33], POVID [68], OPERA [25], VCD [29], EOS [54], HA-DPO [66], HALVA [41], RLHF-V [57], HSA-DPO [55], Silkie [32], and RLAIIF-V [58]—for reference. However, these results are not directly comparable due to differences in base models, preference data, and alignment methods. Notably, RLAIIF-V’s training involves iterative feedback from multiple open-source models, making it challenging to reproduce exactly; thus, we present DPO-based results under default settings for reference.

Evaluation benchmarks. We tested our method on a benchmark for evaluating language and commonsense bias in MLLMs, as well as on two hallucination-specific evaluation sets for MLLMs. VLind-Bench [28] measures language priors and commonsense bias in MLLMs. It has two of the main metrics, Language Prior (LP) and Commonsense Bias (CB), which are used to evaluate the linguistic and visual bias of the model, respectively. Object HalBench [40] is a standard benchmark for object hallucination. Following Yu et al. [58], we augment it with eight diverse prompts for 300 instances and report CHAIR scores [40] for hallucination rate at the response (CHAIRs) and object (CHAIRi) levels. AMBER [50] is a multimodal LLM hallucination benchmark with detailed object annotations, focusing on generative tasks with 1K images. Using the official evaluation tool, we report CHAIR score variants, object coverage, hallucinated response rate, and hallucination overlap with human cognition. We additionally tested on the GPT-4-based hallucination evaluation dataset MMHalBench [44], a practical question-answering benchmark with eight question categories and 12 object topics, using GPT-4 [2] to assess response quality (scored from zero to six) and hallucination rate.

4.2. Main Results

Table 1 presents our main experimental results, from which we can draw the following conclusions: (1) *Our method effectively mitigates modality bias in MLLMs.* Compared to the strongest baseline, our approach achieves an average improvement of 19% on the modality bias benchmark for MLLMs. (2) *The optimization objective is crucial in addressing modality bias in MLLMs.* Although MDPO is designed for hallucination issues, its approach is similar to ours in addressing language bias. However, the results on the bias benchmark indicate that MDPO does not effectively mitigate modality bias in MLLMs. (3) *There is a connection between hallucination and modality bias issues in MLLMs.* In addition to DPO, our enhancements—bias preference optimization data and NaPO—show that our approach reduces hallucination in MLLMs while mitigating modality bias. These observations underline the effectiveness of our method in addressing both modality bias and hallucination.

4.3. Variant Analysis

To better understand the contribution of each module in our method, we selected VLind-Bench, a bias benchmark, and Object HalBench, a hallucination benchmark, as our main evaluation sets. We conducted detailed variant experiments, as shown in Table 2. The experiments are divided into three groups: full data variant analysis (A), language-bias variant analysis (LB), and vision-bias variant analysis (VB). Here, *w/o* denotes *without*, and *repl.* indicates *replace with*.

Through observation and analysis of the variant experiment results, we can draw the following conclusions: (1) *Dynamic weighting effectively balances the optimization magnitude of different losses.* By comparing A1 and A2 in Table 2, we observe a noticeable decline in overall model performance when dynamic weighting is removed. (2) *DPO performs poorly when handling noisy data.* By comparing A3, LB2, and VB2 in Table 2 with their default settings, we observe that DPO not only increases hallucination but also degrades the model’s performance on bias-related issues. (3) *Language-bias data is more effective than vision-bias data in addressing the model’s language prior issues, while vision-bias data better alleviates commonsense bias.* Comparing LB1 and VB1, we observe that LB1 is more effective for language prior issues, whereas VB1 performs better on commonsense bias issues. (4) *The two types of data exhibit a synergistic effect.* Comparing the full dataset (A2) with LB1 and VB1 individually, we observe that mixed training with both data types enhances the model’s performance on both bias issues while also reducing hallucinations. (5) *Vision-bias data is more effective in suppressing hallucinations in the model.* This can be understood in two ways: first, vision-bias data reduces the model’s focus on irrelevant details, preventing it from outputting unrelated elements. Secondly, comparing VB1, LB1, and A2, we observe that vision-bias

Model	VLindBench		Object HalBench		AMBER				MMHalBench	
	CB $\uparrow$	LP $\uparrow$	CHAIRs $\downarrow$	CHAIRi $\downarrow$	CHAIRs $\downarrow$	Cover. $\uparrow$	HalRate $\downarrow$	Cog. $\downarrow$	Score $\uparrow$	HalRate $\downarrow$
GPT-4V	91.1	75.6	13.6	7.3	4.6	67.1	30.7	2.6	3.49	0.28
LLaVA-v1.5-7B	0.0	0.0	53.6	25.2	7.8	51.0	36.4	4.2	2.11	0.54
+ HACL	-	-	-	-	-	-	-	-	2.13	0.50
+ OPERA	-	-	45.1	22.3	-	-	-	-	2.15	0.54
+ POVID	-	-	48.1	24.4	-	-	-	-	2.08	0.56
+ VCD	-	-	48.8	24.3	-	-	-	-	2.12	0.54
+ EOS	-	-	40.3	17.8	5.1	49.1	22.7	2.0	2.03	0.59
+ HA-DPO	-	-	39.9	19.9	6.7	49.8	30.9	3.3	1.97	0.60
+ HALVA	-	-	-	-	6.6	53.0	32.2	3.4	2.25	0.54
+ MDPO-10K	-	-	35.7	9.8	4.4	52.4	24.5	2.4	2.39	0.54
+ RLAIF-V-Iterative	54.3	35.3	8.5	4.3	-	-	-	-	3.06	0.29
LLaVA-v1.5-13B	31.5	20.9	53.3	14.5	8.5	50.9	37.6	4.2	3.03	0.47
+ RLHF-V	-	-	12.2	7.5	6.3	46.1	25.1	2.1	2.81	0.49
+ HSA-DPO	-	-	5.2	3.2	2.1	47.3	13.4	1.2	2.61	0.48
+ HALVA	-	-	-	-	6.4	52.6	30.4	3.2	2.58	0.45
LLaVA-v1.5-7B	0.0	0.0	53.6	25.2	7.8	51.0	36.4	4.2	2.11	0.54
+ $\mathcal{L}_{\text{DPO}}$ with RLAIF-V	39.4	25.4	32.0	8.5	4.9	52.0	23.4	1.6	3.23	0.38
+ $\mathcal{L}_{\text{MDPO}}$ with RLAIF-V	0.3	0.4	35.3	10.5	<u>4.2</u>	<u>53.1</u>	<u>22.4</u>	2.2	<u>3.28</u>	0.42
+ $\mathcal{L}_{\gamma}$ with RLAIF-V-Bias	58.9	44.0	25.7	6.2	4.0	54.1	20.7	1.4	3.31	0.35
LLaVA-v1.5-13B	31.5	20.9	53.3	14.5	8.5	50.9	37.6	4.2	3.03	0.47
+ $\mathcal{L}_{\text{DPO}}$ with RLAIF-V	37.1	<u>21.2</u>	<u>25.8</u>	<u>6.3</u>	3.3	50.5	19.9	<u>1.3</u>	3.39	0.35
+ $\mathcal{L}_{\text{MDPO}}$ with RLAIF-V	32.8	16.9	30.3	9.1	3.4	<u>52.6</u>	18.1	1.4	3.72	0.30
+ $\mathcal{L}_{\gamma}$ with RLAIF-V-Bias	42.1	25.1	23.7	5.9	3.5	55.7	<u>19.0</u>	1.2	<u>3.55</u>	<u>0.33</u>

Table 1. **Main experimental results.** We evaluated our method based on LLaVA-v1.5-7B on bias and hallucination benchmarks, using DPO and MDPO as primary comparisons. Due to differences in training data, model scale, and training strategies, we included additional results for reference. Our method showed an average improvement of 19% on the bias benchmark and a notable reduction in hallucinations across benchmarks.

No.	Variant	VLindBench		Object HalBench	
		CB $\uparrow$	LP $\uparrow$	CHAIRs $\downarrow$	CHAIRi $\downarrow$
A1	-	58.9	44.0	25.7	6.2
A2	w/o γ_i	50.0	38.2	27.7	8.0
A3	repl. DPO	43.4	32.2	29.0	8.3
LB1	w/o VB	40.4	36.4	28.0	6.4
LB2	repl. DPO	44.7	29.9	28.3	6.8
VB1	w/o LB	62.3	31.4	26.3	7.6
VB2	repl. DPO	55.3	28.5	27.7	7.4

Table 2. **Variant analysis.** We divided the variant experiments into three groups based on the data scale. Using controlled variable ablations, we conducted a detailed analysis of the contributions of dynamic weighting, NaPO, language-bias data, and vision-bias data to the final results.

data leads to significantly lower hallucination rates.

4.4. Further Analysis

Quantitative analysis of training data. In this part, we aim to construct a training dataset of equal size to RLAIF-V, while effectively leveraging the advantages of different data types within RLAIF-V-Bias for comparison. All experimental settings are consistent with those in the main

Data	VLindBench		Object HalBench	
	CB $\uparrow$	LP $\uparrow$	CHAIRs $\downarrow$	CHAIRi $\downarrow$
RLAIF-V-Bias	58.9	44.0	25.7	6.2
RLAIF-V	39.4	25.4	32.0	8.5
RLAIF-V-Bias (Random)	51.3	40.9	28.3	7.2

Table 3. **Quantitative analysis of training data.** We use the same training settings as in the main experiment but vary the amount of training data for comparison. “Random” indicates that, for each batch, we randomly select one type of data from the original, language-bias, or vision-bias data, ensuring the total training data volume matches the original RLAIF-V dataset. The results show that *while reducing data volume leads to performance degradation, our data is more effective than the original data at the same volume.*

experiment. Specifically, for each batch, we randomly select one type of training data—original, language-bias, or vision-bias—as the current batch’s training data. This approach can be viewed as randomly selecting one of the three losses in the final optimization objective (Equation (14)) for each batch, setting its weight to 1 while setting the others to 0, to enable a quantitative analysis of the training data. Observation and analysis of the experimental results in Table 3 indicate that *with an equivalent amount of training data, RLAIF-V-Bias*

Margin	VLindBench		Object HalBench	
	CB $\uparrow$	LP $\uparrow$	CHAIRs $\downarrow$	CHAIRi $\downarrow$
LB with ψ_Σ	34.8	22.2	28.7	7.1
LB with ψ_μ (default)	40.4	36.4	28.0	6.4
VB with ψ_Σ (default)	62.3	41.4	26.3	7.6
VB with ψ_μ	52.3	29.5	28.7	7.8

Table 4. **Relationship between q calculation method and data type.** We used the same hyperparameter settings as in the main experiment and tested various methods for calculating q across different datasets, where the calculation of the ψ_μ and ψ_Σ is shown in Equation (2). The results indicate that *using inappropriate noise estimation methods to calculate the noise robustness coefficient leads to performance degradation.*

(Random) still outperforms RLAIIF-V in mitigating hallucination and bias issues, which demonstrate the effective of our motivation. However, the performance of the RLAIIF-V-Bias (Random) still falls short compared to the RLAIIF-V-Bias.

Noise evaluation metric analysis. In this chapter, we analyze the impact of different noise evaluation methods on the final results by applying distinct noise assessment methods to various data types. All experimental settings and hyperparameters remain consistent with those in the main experiment. Specifically, for the scale factor α in Equation (12), we use the default settings: $\alpha = 0.5$ for reward margins ψ_μ calculated using average logp, and $\alpha = 0.01$ for reward margins ψ_Σ calculated using logp, to ensure comparability of results. Observing the experimental results in Table 4, we see that *using inappropriate noise evaluation methods leads to a sharp decline in model performance.* This may be because incorrect noise assessment increases the gradient for noisy data in the loss function while decreasing it for non-noisy data. As a result, the model may overfit to noisy data and overlook high-quality data during optimization. This outcome also supports the pattern we summarized in Section 3.2 based on our analysis of Figure 4.

5. Related Work

5.1. Modality Bias in MLLMs

Modality bias occurs when a model overly relies on one modality or prior knowledge, neglecting other relevant modalities [11, 21]. This leads to a focus on incorrect information and weakens generalization, manifesting as language bias [1, 3, 8, 70] or vision bias [22, 43]. Related to the hallucination problem [62], modality bias has been extensively studied in VQA tasks [3, 8, 11, 19, 21, 36, 70]. However, most prior work has focused on balanced datasets and complex training strategies, which do not generalize well to MLLMs. Recent studies addressing modality bias in MLLMs include benchmarks for evaluating bias [28],

datasets for assessing modality bias [11], and methods like contrastive decoding to reduce language priors [62]. While MDPO [48] takes an approach similar to ours, it targets hallucination by lowering the probability of preferred samples when images are absent, differing from our method that directly reduces the probability of biased responses without images. Additionally, works such as [14, 42, 52, 60] address social and group biases in vision-language models like CLIP. Although similar in theme, these studies primarily focus on social biases, which diverges from our focus on modality bias in MLLMs.

5.2. Preference Optimization

Preference optimization algorithms enhance LLMs by aligning them with human values [37]. While RLHF[5, 10, 12, 13, 37, 47] is effective, its complexity has led to exploration of simpler alternatives. RAFT [15] selects optimal samples with existing reward models, RRHF [59] uses a simpler ranking loss, and DPO [39] employs a preference-based loss for improved stability. SLiC-HF [65], KTO [17], RSO [35], and ORPO [23] focus on preference calibration and efficient modeling. β -DPO [53] dynamically adjusts β based on data distribution. With the increasing scale of models and training datasets, obtaining human-annotated gold data has become increasingly challenging [9]. Recently, LLM-synthesized data [45] has garnered significant attention. However, synthesized data inevitably introduces noise. Current research primarily focuses on leveraging automatically constructed data [45] and AI-based feedback [20, 30] to address existing challenges, while relatively little attention has been paid to enhancing the noise robustness of training methods. In contrast, NaPO integrates noise-robust MAE into DPO and dynamically adjusts noise robustness based on data noise, enhancing stability and resilience in noisy environments.

6. Conclusion

We introduce RLAIIF-V-Bias, a dataset designed to optimize preferences in MLLMs and reduce modality bias by including both language and visual bias data. To handle noise in automatically constructed data, we propose NaPO, which combines BCE loss from DPO with noise-robust MAE loss using a negative Box-Cox transformation, allowing dynamic noise detection and robustness adjustment. Experimental results demonstrate that our method effectively mitigates bias and significantly reduces hallucinations in the model outputs. Compared to DPO, NaPO exhibits stronger resilience to noise in automatically constructed data. However, two clouds still linger behind our research: *whether NaPO can robustly handle noise originating from LLM-synthesized data in broader contexts, and whether bias is always harmful.*

Acknowledgement

We would like to thank the anonymous reviewers for their comments. This work is supported by the National Natural Science Foundation of China (No.62406319), the Youth Innovation Promotion Association of CAS (No.2021153), and the Postdoctoral Fellowship Program of CPSF (No.GZC20232968).

References

- [1] Ehsan Abbasnejad, Damien Teney, Amin Parvaneh, Javen Shi, and Anton van den Hengel. Counterfactual vision and language learning. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 10044–10054, 2020. 8
- [2] Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altmerschmidt, Sam Altman, Shyamal Anadkat, et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023. 1, 6
- [3] Aishwarya Agrawal, Dhruv Batra, Devi Parikh, and Anirudha Kembhavi. Don’t just assume; look and answer: Overcoming priors for visual question answering. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 4971–4980, 2018. 8
- [4] Jinze Bai, Shuai Bai, Shusheng Yang, Shijie Wang, Sinan Tan, Peng Wang, Junyang Lin, Chang Zhou, and Jingren Zhou. Qwen-vl: A frontier large vision-language model with versatile abilities. *arXiv preprint arXiv:2308.12966*, 2023. 1
- [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. 8
- [6] George EP Box and David R Cox. An analysis of transformations. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 26(2):211–243, 1964. 2
- [7] Ralph Allan Bradley and Milton E Terry. Rank analysis of incomplete block designs: I. the method of paired comparisons. *Biometrika*, 39(3/4):324–345, 1952. 3
- [8] Remi Cadene, Corentin Dancette, Matthieu Cord, Devi Parikh, et al. Rubi: Reducing unimodal biases for visual question answering. *Advances in neural information processing systems*, 32, 2019. 8
- [9] Boxi Cao, Keming Lu, Xinyu Lu, Jiawei Chen, Mengjie Ren, Hao Xiang, Peilin Liu, Yaojie Lu, Ben He, Xianpei Han, et al. Towards scalable automated alignment of llms: A survey. *arXiv preprint arXiv:2406.01252*, 2024. 8
- [10] Huayu Chen, Guande He, Lifan Yuan, Ganqu Cui, Hang Su, and Jun Zhu. Noise contrastive alignment of language models with explicit rewards. *arXiv preprint arXiv:2402.05369*, 2024. 8
- [11] Meiqi Chen, Yixin Cao, Yan Zhang, and Chaochao Lu. Quantifying and mitigating unimodal biases in multimodal large language models: A causal perspective. *arXiv preprint arXiv:2403.18346*, 2024. 1, 8
- [12] Sayak Ray Chowdhury, Anush Kini, and Nagarajan Natarajan. Provably robust dpo: Aligning language models with noisy feedback. *arXiv preprint arXiv:2403.00409*, 2024. 8
- [13] Paul F Christiano, Jan Leike, Tom Brown, Miljan Martic, Shane Legg, and Dario Amodei. Deep reinforcement learning from human preferences. *Advances in neural information processing systems*, 30, 2017. 2, 8
- [14] Ching-Yao Chuang, Varun Jampani, Yuanzhen Li, Antonio Torralba, and Stefanie Jegelka. Debiasing vision-language models via biased prompts. *arXiv preprint arXiv:2302.00070*, 2023. 8
- [15] Hanze Dong, Wei Xiong, Deepanshu Goyal, Yihan Zhang, Winnie Chow, Rui Pan, Shizhe Diao, Jipeng Zhang, Kashun Shum, and Tong Zhang. Raft: Reward ranked finetuning for generative foundation model alignment. *arXiv preprint arXiv:2304.06767*, 2023. 8
- [16] Qingxiu Dong, Damai Dai, Yifan Song, Jingjing Xu, Zhifang Sui, and Lei Li. Calibrating factual knowledge in pretrained language models. In *Findings of the Association for Computational Linguistics: EMNLP 2022, Abu Dhabi, United Arab Emirates, December 7-11, 2022*, pages 5937–5947. Association for Computational Linguistics, 2022. 2
- [17] Kawin Ethayarajh, Winnie Xu, Niklas Muennighoff, Dan Jurafsky, and Douwe Kiela. Kto: Model alignment as prospect theoretic optimization. *arXiv preprint arXiv:2402.01306*, 2024. 8
- [18] Aritra Ghosh, Naresh Manwani, and PS Sastry. Making risk minimization tolerant to label noise. *Neurocomputing*, 160: 93–107, 2015. 2, 3
- [19] Tejas Gokhale, Pratyay Banerjee, Chitta Baral, and Yezhou Yang. Mutant: A training paradigm for out-of-distribution generalization in visual question answering. *arXiv preprint arXiv:2009.08566*, 2020. 2, 8
- [20] Jiawei Gu, Xuhui Jiang, Zhichao Shi, Hexiang Tan, Xuehao Zhai, Chengjin Xu, Wei Li, Yinghan Shen, Shengjie Ma, Honghao Liu, et al. A survey on llm-as-a-judge. *arXiv preprint arXiv:2411.15594*, 2024. 8
- [21] Yangyang Guo, Liqiang Nie, Harry Cheng, Zhiyong Cheng, Mohan Kankanhalli, and Alberto Del Bimbo. On modality bias recognition and reduction. *ACM Transactions on Multimedia Computing, Communications and Applications*, 19(3): 1–22, 2023. 8
- [22] Vipul Gupta, Zhuowan Li, Adam Kortylewski, Chenyu Zhang, Yingwei Li, and Alan Yuille. Swapmix: Diagnosing and regularizing the over-reliance on visual context in visual question answering. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 5078–5088, 2022. 8
- [23] Jiwoo Hong, Noah Lee, and James Thorne. Orpo: Monolithic preference optimization without reference model. *arXiv preprint arXiv:2403.07691*, 2(4):5, 2024. 8
- [24] Jianqiang Huang, Yu Qin, Jiabin Qi, Qianru Sun, and Hanwang Zhang. Deconfounded visual grounding. In *Thirty-Sixth AAAI Conference on Artificial Intelligence, AAAI 2022, Thirty-Fourth Conference on Innovative Applications of Artificial Intelligence, IAAI 2022, The Twelveth Symposium on Educational Advances in Artificial Intelligence, EAAI 2022 Virtual*

- Event, February 22 - March 1, 2022, pages 998–1006. AAAI Press, 2022. 2
- [25] Qidong Huang, Xiaoyi Dong, Pan Zhang, Bin Wang, Conghui He, Jiaqi Wang, Dahua Lin, Weiming Zhang, and Nenghai Yu. Opera: Alleviating hallucination in multi-modal large language models via over-trust penalty and retrospection-allocation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 13418–13427, 2024. 6
- [26] Nikhil Kandpal, Haikang Deng, Adam Roberts, Eric Wallace, and Colin Raffel. Large language models struggle to learn long-tail knowledge. In *International Conference on Machine Learning*, pages 15696–15707. PMLR, 2023. 4
- [27] Camila Kolling, Martin D. Móre, Nathan Gavenski, Eduardo H. P. Pooch, Otávio Parraga, and Rodrigo C. Barros. Efficient counterfactual debiasing for visual question answering. In *IEEE/CVF Winter Conference on Applications of Computer Vision, WACV 2022, Waikoloa, HI, USA, January 3-8, 2022*, pages 2572–2581. IEEE, 2022. 2
- [28] Kang-il Lee, Minbeom Kim, Seunghyun Yoon, Minsung Kim, Dongryeol Lee, Hyukhun Koh, and Kyomin Jung. Vlind-bench: Measuring language priors in large vision-language models. *arXiv preprint arXiv:2406.08702*, 2024. 1, 2, 6, 8
- [29] Sicong Leng, Hang Zhang, Guanzheng Chen, Xin Li, Shijian Lu, Chunyan Miao, and Lidong Bing. Mitigating object hallucinations in large vision-language models through visual contrastive decoding. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 13872–13882, 2024. 6
- [30] Dawei Li, Bohan Jiang, Liangjie Huang, Alimohammad Beigi, Chengshuai Zhao, Zhen Tan, Amrita Bhattacharjee, Yuxuan Jiang, Canyu Chen, Tianhao Wu, et al. From generation to judgment: Opportunities and challenges of llm-as-a-judge. *arXiv preprint arXiv:2411.16594*, 2024. 8
- [31] Junnan Li, Dongxu Li, Silvio Savarese, and Steven Hoi. Blip-2: Bootstrapping language-image pre-training with frozen image encoders and large language models. In *International conference on machine learning*, pages 19730–19742. PMLR, 2023. 1
- [32] Lei Li, Zhihui Xie, Mukai Li, Shunian Chen, Peiyi Wang, Liang Chen, Yazheng Yang, Benyou Wang, and Lingpeng Kong. Silk: Preference distillation for large visual language models. *arXiv preprint arXiv:2312.10665*, 2023. 6
- [33] Haotian Liu, Chunyuan Li, Yuheng Li, and Yong Jae Lee. Improved baselines with visual instruction tuning. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 26296–26306, 2024. 1, 2, 6
- [34] Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. Visual instruction tuning. *Advances in neural information processing systems*, 36, 2024. 1
- [35] Tianqi Liu, Yao Zhao, Rishabh Joshi, Misha Khalman, Mohammad Saleh, Peter J Liu, and Jialu Liu. Statistical rejection sampling improves preference optimization. *arXiv preprint arXiv:2309.06657*, 2023. 8
- [36] Yulei Niu, Kaihua Tang, Hanwang Zhang, Zhiwu Lu, Xian-Sheng Hua, and Ji-Rong Wen. Counterfactual vqa: A cause-effect look at language bias. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 12700–12710, 2021. 8
- [37] 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. 2, 8
- [38] Vaidehi Patil, Adyasha Maharana, and Mohit Bansal. Debiasing multimodal models via causal information minimization. In *Findings of the Association for Computational Linguistics: EMNLP 2023, Singapore, December 6-10, 2023*, pages 4108–4123. Association for Computational Linguistics, 2023. 2
- [39] Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. Direct preference optimization: Your language model is secretly a reward model. *Advances in Neural Information Processing Systems*, 36, 2024. 2, 6, 8
- [40] Anna Rohrbach, Lisa Anne Hendricks, Kaylee Burns, Trevor Darrell, and Kate Saenko. Object hallucination in image captioning. *arXiv preprint arXiv:1809.02156*, 2018. 2, 6
- [41] Pritam Sarkar, Sayna Ebrahimi, Ali Etemad, Ahmad Beirami, Sercan Ö Arık, and Tomas Pfister. Mitigating object hallucination via data augmented contrastive tuning. *arXiv preprint arXiv:2405.18654*, 2024. 6
- [42] Ashish Seth, Mayur Hemani, and Chirag Agarwal. Dear: Debiasing vision-language models with additive residuals. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 6820–6829, 2023. 8
- [43] Qingyi Si, Fandong Meng, Mingyu Zheng, Zheng Lin, Yuanxin Liu, Peng Fu, Yanan Cao, Weiping Wang, and Jie Zhou. Language prior is not the only shortcut: A benchmark for shortcut learning in vqa. *arXiv preprint arXiv:2210.04692*, 2022. 8
- [44] Zhiqing Sun, Sheng Shen, Shengcao Cao, Haotian Liu, Chunyuan Li, Yikang Shen, Chuang Gan, Liang-Yan Gui, Yu-Xiong Wang, Yiming Yang, et al. Aligning large multimodal models with factually augmented rlhf. *arXiv preprint arXiv:2309.14525*, 2023. 2, 6
- [45] Zhen Tan, Dawei Li, Song Wang, Alimohammad Beigi, Bohan Jiang, Amrita Bhattacharjee, Mansoor Karami, Jundong Li, Lu Cheng, and Huan Liu. Large language models for data annotation and synthesis: A survey. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 930–957, 2024. 8
- [46] Hengzhu Tang, Zefeng Zhang, Zhiping Li, Zhenyu Zhang, Xing Wu, Li Gao, Suqi Cheng, and Dawei Yin. Multi-branch collaborative learning network for video quality assessment in industrial video search. *arXiv preprint arXiv:2502.05924*, 2025. 1
- [47] 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. 8
- [48] Fei Wang, Wenxuan Zhou, James Y Huang, Nan Xu, Sheng Zhang, Hoifung Poon, and Muhao Chen. mdpo: Conditional

- preference optimization for multimodal large language models. *arXiv preprint arXiv:2406.11839*, 2024. 6, 8
- [49] Jialu Wang, Yang Liu, and Xin Eric Wang. Are gender-neutral queries really gender-neutral? mitigating gender bias in image search. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing, EMNLP 2021, Virtual Event / Punta Cana, Dominican Republic, 7-11 November, 2021*, pages 1995–2008. Association for Computational Linguistics, 2021. 2
- [50] Junyang Wang, Yuhang Wang, Guohai Xu, Jing Zhang, Yukai Gu, Haitao Jia, Ming Yan, Ji Zhang, and Jitao Sang. An llm-free multi-dimensional benchmark for mllms hallucination evaluation. *arXiv preprint arXiv:2311.07397*, 2023. 2, 6
- [51] Zecheng Wang, Xinye Li, Zhanyue Qin, Chunshan Li, Zhiying Tu, Dianhui Chu, and Dianbo Sui. Can we debias multimodal large language models via model editing? In *Proceedings of the 32nd ACM International Conference on Multimedia, MM 2024, Melbourne, VIC, Australia, 28 October 2024 - 1 November 2024*, pages 3219–3228. ACM, 2024. 2
- [52] Zecheng Wang, Xinye Li, Zhanyue Qin, Chunshan Li, Zhiying Tu, Dianhui Chu, and Dianbo Sui. Can we debias multimodal large language models via model editing? In *Proceedings of the 32nd ACM International Conference on Multimedia, pages 3219–3228*, 2024. 8
- [53] Junkang Wu, Yuexiang Xie, Zhengyi Yang, Jiancan Wu, Jinyang Gao, Bolin Ding, Xiang Wang, and Xiangnan He. β -dpo: Direct preference optimization with dynamic β . *arXiv preprint arXiv:2407.08639*, 2024. 8
- [54] Yue Wu, Zhiqing Sun, Huizhuo Yuan, Kaixuan Ji, Yiming Yang, and Quanquan Gu. Self-play preference optimization for language model alignment. *arXiv preprint arXiv:2405.00675*, 2024. 6
- [55] Wenyi Xiao, Ziwei Huang, Leilei Gan, Wangui He, Haoyuan Li, Zhelun Yu, Hao Jiang, Fei Wu, and Linchao Zhu. Detecting and mitigating hallucination in large vision language models via fine-grained ai feedback. *arXiv preprint arXiv:2404.14233*, 2024. 6
- [56] Qinghao Ye, Haiyang Xu, Guohai Xu, Jiabo Ye, Ming Yan, Yiyang Zhou, Junyang Wang, Anwen Hu, Pengcheng Shi, Yaya Shi, et al. mplug-owl: Modularization empowers large language models with multimodality. *arXiv preprint arXiv:2304.14178*, 2023. 1
- [57] Tianyu Yu, Yuan Yao, Haoye Zhang, Taiwen He, Yifeng Han, Ganqu Cui, Jinyi Hu, Zhiyuan Liu, Hai-Tao Zheng, Maosong Sun, et al. Rlhv: Towards trustworthy mllms via behavior alignment from fine-grained correctional human feedback. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 13807–13816, 2024. 6
- [58] Tianyu Yu, Haoye Zhang, Yuan Yao, Yunkai Dang, Da Chen, Xiaoman Lu, Ganqu Cui, Taiwen He, Zhiyuan Liu, Tat-Seng Chua, et al. Rlaif-v: Aligning mllms through open-source ai feedback for super gpt-4v trustworthiness. *arXiv preprint arXiv:2405.17220*, 2024. 2, 3, 5, 6
- [59] Zheng Yuan, Hongyi Yuan, Chuanqi Tan, Wei Wang, Songfang Huang, and Fei Huang. Rrhf: Rank responses to align language models with human feedback without tears. *arXiv preprint arXiv:2304.05302*, 2023. 8
- [60] Jie Zhang, Sibow Wang, Xiangkui Cao, Zheng Yuan, Shiguang Shan, Xilin Chen, and Wen Gao. Vlbiasbench: A comprehensive benchmark for evaluating bias in large vision-language model. *arXiv preprint arXiv:2406.14194*, 2024. 8
- [61] Wenyuan Zhang, Jiawei Sheng, Shuaiyi Nie, Zefeng Zhang, Xinghua Zhang, Yongquan He, and Tingwen Liu. Revealing the challenge of detecting character knowledge errors in llm role-playing. *arXiv preprint arXiv:2409.11726*, 2024. 1
- [62] Yi-Fan Zhang, Weichen Yu, Qingsong Wen, Xue Wang, Zhang Zhang, Liang Wang, Rong Jin, and Tieniu Tan. Debiasing large visual language models. *arXiv preprint arXiv:2403.05262*, 2024. 1, 2, 8
- [63] Zhilu Zhang and Mert Sabuncu. Generalized cross entropy loss for training deep neural networks with noisy labels. *Advances in neural information processing systems*, 31, 2018. 2, 3, 4
- [64] Zefeng Zhang, Jiawei Sheng, Chuang Zhang, Yunzhi Liang, Wenyuan Zhang, Siqi Wang, and Tingwen Liu. Optimal transport guided correlation assignment for multimodal entity linking. *arXiv preprint arXiv:2406.01934*, 2024. 1
- [65] Yao Zhao, Rishabh Joshi, Tianqi Liu, Misha Khalman, Mohammad Saleh, and Peter J Liu. Slic-hf: Sequence likelihood calibration with human feedback. *arXiv preprint arXiv:2305.10425*, 2023. 8
- [66] Zhiyuan Zhao, Bin Wang, Linke Ouyang, Xiaoyi Dong, Jiaqi Wang, and Conghui He. Beyond hallucinations: Enhancing llms through hallucination-aware direct preference optimization. *arXiv preprint arXiv:2311.16839*, 2023. 6
- [67] Jeffrey Zhou, Tianjian Lu, Swaroop Mishra, Siddhartha Brahma, Sujoy Basu, Yi Luan, Denny Zhou, and Le Hou. Instruction-following evaluation for large language models. *CoRR*, abs/2311.07911, 2023. 2
- [68] Yiyang Zhou, Chenhang Cui, Rafael Rafailov, Chelsea Finn, and Huaxiu Yao. Aligning modalities in vision large language models via preference fine-tuning. *arXiv preprint arXiv:2402.11411*, 2024. 6
- [69] Deyao Zhu, Jun Chen, Xiaoqian Shen, Xiang Li, and Mohamed Elhoseiny. Minigpt-4: Enhancing vision-language understanding with advanced large language models. *arXiv preprint arXiv:2304.10592*, 2023. 1
- [70] Xi Zhu, Zhendong Mao, Chunxiao Liu, Peng Zhang, Bin Wang, and Yongdong Zhang. Overcoming language priors with self-supervised learning for visual question answering. *arXiv preprint arXiv:2012.11528*, 2020. 1, 8