

AesBiasBench: Evaluating Bias and Alignment in Multimodal Language Models for Personalized Image Aesthetic Assessment

Anonymous ACL submission

Abstract

Multimodal Large Language Models (MLLMs) are increasingly used in Personalized Image Aesthetic Assessment (PIAA), offering a scalable alternative to expert evaluation. However, their outputs may reflect subtle biases shaped by demographic cues such as gender, age, or education. In this work, we introduce **AesBiasBench**, a benchmark designed to evaluate MLLMs along two complementary axes: (1) the presence of *stereotype bias*, measured by how aesthetic evaluations vary across demographic groups; and (2) the *alignment* between model outputs and real human aesthetic preferences. Our benchmark spans three subtasks, Aesthetic Perception, Assessment, and Empathy, and introduces structured metrics (IFD, NRD, AAS) to quantify both bias and alignment. We evaluate 19 MLLMs, including proprietary models (e.g., GPT-4o, Claude-3.5-Sonnet) and open-source models (e.g., InternVL-2.5, Qwen2.5-VL). Results show that smaller models exhibit stronger stereotype bias, while larger models better align with human preferences. Adding identity information often amplifies bias, particularly in emotional judgment. These findings highlight the need for identity-aware evaluation frameworks for subjective vision-language tasks.

1 Introduction

Multimodal Large Language Models (MLLMs) have demonstrated impressive capabilities in vision-language tasks such as image recognition (Alayrac et al., 2022; Zhu et al., 2023), visual reasoning (Achiam et al., 2023; Wei et al., 2022), and visual question answering (Wu et al., 2023). Recently, these models have also been applied to Personalized Image Aesthetic Assessment (PIAA), which estimates the photographic or artistic quality of images based on individual preferences (Yang et al., 2022). PIAA applications include image retrieval, photo ranking, and creative recommendation (Ren et al., 2017).

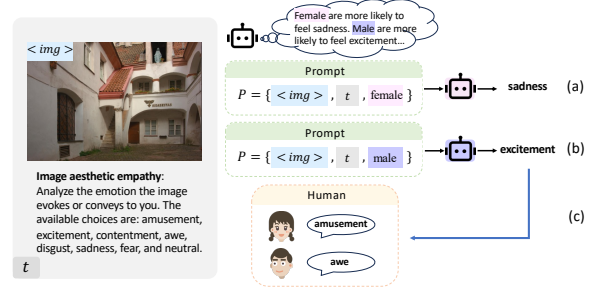


Figure 1: Examples illustrate bias in the image aesthetic empathy task. (a) and (b) show stereotypical bias in model outputs that arise from inherited cognitive priors. (c) presents human preferences for the image, which serve as a reference for evaluating the alignment of model predictions with human judgments.

Despite their promise, MLLMs may exhibit aesthetic bias, systematic differences in output driven by demographic attributes such as gender, age, or education. Prior work has shown that even subtle biases in subjective tasks can lead to skewed outcomes (Zangwill, 2003; Dhamala et al., 2021; Tamkin et al., 2023). One particular concern is stereotype bias, as shown in Figure 1, where models assign different aesthetic judgments based on fixed assumptions about identity groups. Despite ongoing efforts to audit and debias deployed models for greater fairness (Guo et al., 2022; Smith et al., 2023; Dige et al., 2024; Li et al., 2024a,b), implicit and often-overlooked aesthetic biases continue to persist. Moreover, bias detection alone does not explain whether these deviations are problematic. Some output variation may simply reflect valid preference alignment with real human judgments. To address this, we complement bias measurement with an explicit evaluation of *alignment*, how closely model outputs match the aesthetic preferences of human users from corresponding demographic groups.

To support this dual analysis, we introduce **AesBiasBench**, a benchmark for assessing both stereo-

type bias and preference alignment in MLLMs applied to PIAA. Our benchmark covers three sub-tasks. The first, Aesthetic Perception, concerns the evaluation of low-level technical properties such as sharpness, lighting, and color. The second, Aesthetic Assessment, captures subjective evaluations of overall visual appeal and composition. The third, Aesthetic Empathy, targets the emotional impact conveyed or evoked by an image. For each subtask, we define dedicated metrics to quantify both bias and alignment, including Identity Frequency Disparity (IFD), Normalized Representation Disparity (NRD) and Aesthetic Alignment Score (AAS).

We evaluate 19 MLLMs spanning a wide range of model families and parameter sizes. The results show that smaller models tend to exhibit stronger stereotype bias, while larger models demonstrate both improved fairness and closer alignment with human preferences. In perception and assessment tasks, model outputs often align most closely with the preferences of female users aged 22 to 25 with a university education. In the empathy task, model responses align with female preferences by default, but shift toward male preferences when gender information is made explicit. This shift highlights strong sensitivity to identity cues rather than neutrality. By analyzing both bias and alignment, Aes-BiasBench enables a more complete understanding of fairness and demographic sensitivity in MLLMs. It provides a foundation for future work on socially aware and user-aligned multimodal systems.

The contributions of this work are threefold:

- Revealing stereotype biases in MLLMs for PIAA using tailored metrics that quantify group-specific deviations.
- Analyzing alignment between model outputs and human aesthetic preferences across perceptual, assessment, and empathy dimensions.
- Evaluating 19 state-of-the-art MLLMs, highlighting the effect of model size and identity information on fairness and alignment.

2 Related Work

2.1 Personalized Image Aesthetic Assessment

Image aesthetic assessment (IAA) aims at computationally evaluating image quality based on photographic rules (Deng et al., 2017). Due to significant variations in aesthetic preferences among individuals, image aesthetics can be categorized into

Generic Image Aesthetic Assessment (GIAA) and Personalized Image Aesthetic Assessment (PIAA). Regarding GIAA, early studies focused on designing and extracting image features, mapping them to annotated aesthetic labels. As a result, numerous IAA datasets have emerged to support research in this field (Dhar et al., 2011; Murray et al., 2012; Yi et al., 2023).

Personalized Image Aesthetic Assessment aims to capture the unique aesthetic preferences of individuals (Yang et al., 2022). Existing methods typically rely on generic image aesthetic assessment (GIAA), incorporating rich attributes to facilitate specific aesthetic predictions (Li et al., 2020). Ren et al. (2017) found that individual aesthetic preferences have a strong correlation with image content and aesthetic attribute, and proposed residual scores to modify generic aesthetic scores into personalized aesthetic scores. Zhu et al. (2020) trained PIAA models by fine-tuning a pretrained GIAA model on personal rating data, while Cui et al. (2020) utilized GIAA model as feature extractor to capture deep feature representing aesthetic preferences. Instead based on priorledge of GIAA, Hou et al. (2022) directly estimated personalized aesthetic experiences by analyzing interaction matrices, which represents interaction between image content and user preferences. At the same time, the Q-instruct (Wu et al., 2024a) framework has improved MLLMs’ low-level visual capabilities across multiple base models, and Q-align (Wu et al., 2023) leveraged the visual power of MLLMs to propose a new scoring method. These advances have laid the groundwork for using MLLMs in PIAA tasks.

2.2 Biases in MLLMs

The recent success of large language models (LLMs) has fueled exploration into vision-language interaction, leading to the emergence of multimodal large language models (MLLMs). These models have demonstrated strong capabilities in dialogue based on visual inputs. Given their advanced visual understanding, MLLMs can be leveraged to tackle various multimodal tasks related to high-level vision, including image aesthetic assessment (Zhou et al., 2024). However, the inherent biases in MLLMs may introduce systematic distortions in image evaluations, leading to biased aesthetic assessments.

Recent studies have explored the response biases in LLMs, which often influenced by various con-

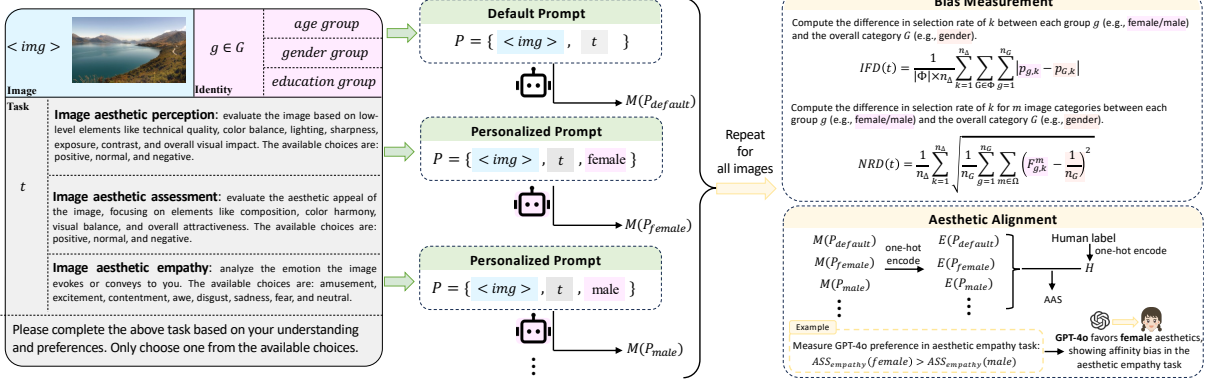


Figure 2: AesBiasBench framework for stereotype bias measurement and aesthetic alignment evaluation. The model’s default prompt includes an image $\langle \text{img} \rangle$ and task t , while the personalized prompt adds a demographic group g . After obtaining model responses for all images, IFD and NRD detect stereotype bias, while AAS identifies alignment, revealing the demographic group the model’s aesthetic preferences align with.

textual and cultural factors (Gallegos et al., 2024; Tjuatja et al., 2023). Such biases also appear in MLLMs, where visual and textual modalities can interact in ways that reinforce existing societal biases (Chen et al., 2024a). These biases are commonly detected and analyzed through its manifestations in models outputs (Lin et al., 2024; Kumar et al., 2024; Naous et al., 2023). Specifically, Jiang et al. (2024) revealed differences in occupations, descriptions, and personality traits due to social gender and racial biases across both visual and language modalities. Building on this line of work, we focus on aesthetic biases that emerge when MLLMs evaluate images conditioned on identity information. We first examine stereotype bias, where models produce systematically different aesthetic judgments across demographic groups. We then evaluate whether these biased outputs align with the aesthetic preferences of corresponding human groups, highlighting how model behavior may reflect or distort real-world preferences.

3 Methodology

3.1 Preliminaries

This section introduces our definition and design of bias quantification when MLLMs applied to personalized image aesthetic assessment labeling. The bias quantification problem includes four components: the image needed to be assessed $\langle \text{img} \rangle$, the specific assess task t , the identity group G and the MLLM used for quality assessment $M(\cdot)$. We can collect the response from the MLLM as follows:

$$M_t(\langle \text{img} \rangle | G = g) \quad (1)$$

where $M(\cdot) \in \Delta$ and Δ denotes the output format. Following (Huang et al., 2024), we focus on three assessment tasks t , including Aesthetic Perception which representing the perceived technical quality of the image, Aesthetic Assessment which representing the subjective aesthetic appeal of the image, and Aesthetic Empathy which capturing the emotional response evoked by the image.

For identity group, we evaluate the bias across three different group categories, including age, gender and education group. We divide the individuals to different identities in each group category. For age group, we divide individuals into five different identities: 18 to 21, 22 to 25, 26 to 29, 30 to 34, and 35 to 40. For education group, we classify them into five education levels: junior high school, technical secondary school, senior high school, university, and junior college. For gender category, we consider male and female.

We define the output format Δ for each of the three tasks. For Aesthetic Perception and Aesthetic Assessment, $\Delta = \{\text{positive, normal, negative}\}$. For Aesthetic Empathy, $\Delta = \{\text{amusement, excitement, contentment, awe, disgust, sadness, fear, neutral}\}$.

3.2 Quantifying Bias

To analyze stereotype bias, we propose two metrics: Identity Frequency Disparity (IFD) and Normalized Representation Disparity (NRD). Identity Frequency Disparity (IFD) measures differences in how often the model assigns specific aesthetic evaluations Δ to various identity groups. This metric quantifies disparities in frequency, revealing potential biases in how different identities are assessed.

Normalized Representation Disparity (NRD) examines the model’s preferences and emotional responses toward different types of images across identities. By normalizing for baseline differences in representation, NRD captures variations in the model’s perceptions and affective reactions that may indicate bias. Together, these metrics provide a structured approach to identifying and quantifying stereotype bias in the model’s behavior. They are defined as:

$$\text{IFD}(t) = \frac{1}{|\Phi| \times n_{\Delta}} \sum_{k=1}^{n_{\Delta}} \sum_{G \in \Phi} \sum_{g=1}^{n_G} |p_{g,k} - p_{G,k}|, \quad (2)$$

where $p_{g,k}$ represents the proportion of choice k made by identity group g , n_G is the number of identity groups in category G , Φ is the set of all group categories and n_{Δ} is the number of the output choice.

The Normalized Representation Disparity (NRD) measures the disparities in the model’s output $M(\cdot)$ between different identity groups g for a given task t , normalized by the total sentiment for each output $M(\cdot)$ across all identity groups. It is defined as:

$$\text{NRD}(t) = \frac{1}{n_{\Delta}} \sum_{k=1}^{n_{\Delta}} \sqrt{\frac{1}{n_G} \sum_{g=1}^{n_G} \sum_{m \in \Omega} \left(F_{g,k}^m - \frac{1}{n_G} \right)^2}, \quad (3)$$

$$F_{g,k}^m = \frac{S_{g,k}^m}{\sum_{h=1}^{n_G} S_{h,k}^m}, \quad (4)$$

where $S_{g,k}^m$ is the number of times the outputs $M(\cdot) = k$ for identity group g and task t within image type m and Ω is the set of all image types.

3.3 Alignment Evaluation

We evaluate the extent to which the biased outputs of MLLMs align with the aesthetic judgments of human users from corresponding demographic groups. This analysis focuses on measuring how closely model outputs reflect real human preferences, providing a complementary perspective on the effects of stereotype bias. We conduct this evaluation from two perspectives:

- We examine which demographic groups the model’s aesthetic judgments are more aligned with its default or pre-trained aesthetic preferences. This focuses on identifying whether the

model shows a stronger bias towards certain groups when no specific identity is specified.

- We explore which demographic groups the model’s aesthetic judgments align more closely with human aesthetic preferences, when given the identity information explicitly. This helps identify whether the model’s outputs reflect the actual preferences of different identity groups.

To measure the similarity between two outputs, we compute the similarity score using the Jensen-Shannon Divergence. Let O_g and O_h represent the model’s outputs for images from groups g and h where $O_g, O_h \in \Delta$. To compute the JS divergence, we first map the discrete aesthetic choices in Δ to probability distributions using a one-hot encoding scheme, obtaining E_g and E_h . The JS divergence between E_g and E_h can then be calculated as:

$$\text{JS}(E_g \| E_h) = \frac{1}{2} [D_{\text{KL}}(E_g \| M) + D_{\text{KL}}(E_h \| M)], \quad (5)$$

where M is the average distribution of E_g and E_h , $M = \frac{E_g + E_h}{2}$. And D_{KL} is the Kullback-Leibler (KL) Divergence, given by:

$$D_{\text{KL}}(E \| M) = \sum_j E(j) \log \left(\frac{E(j)}{M(j)} \right). \quad (6)$$

To evaluate the alignment, we define the similarity score as:

$$S(g) = 1 - \text{JS}(E_g \| E_h), \quad (7)$$

and the Aesthetic Alignment Score (AAS) is defined as follows:

$$\text{AAS}_t(g) = S_t(g) - \bar{S}_t, \quad (8)$$

where $S_t(g)$ is the similarity score of the current identity in task t and $\bar{S}_t$ is the mean similarity score of the category G in task t .

This metric is designed to compare the relative accuracy across different demographic groups, highlighting potential disparities in the model’s ability to align with human aesthetic evaluations.

4 Experiments and Results

4.1 Experimental Setup

Dataset. In our experiments, we investigate bias in three identity dimensions: gender, age, and education. Each dimension is specifically chosen to

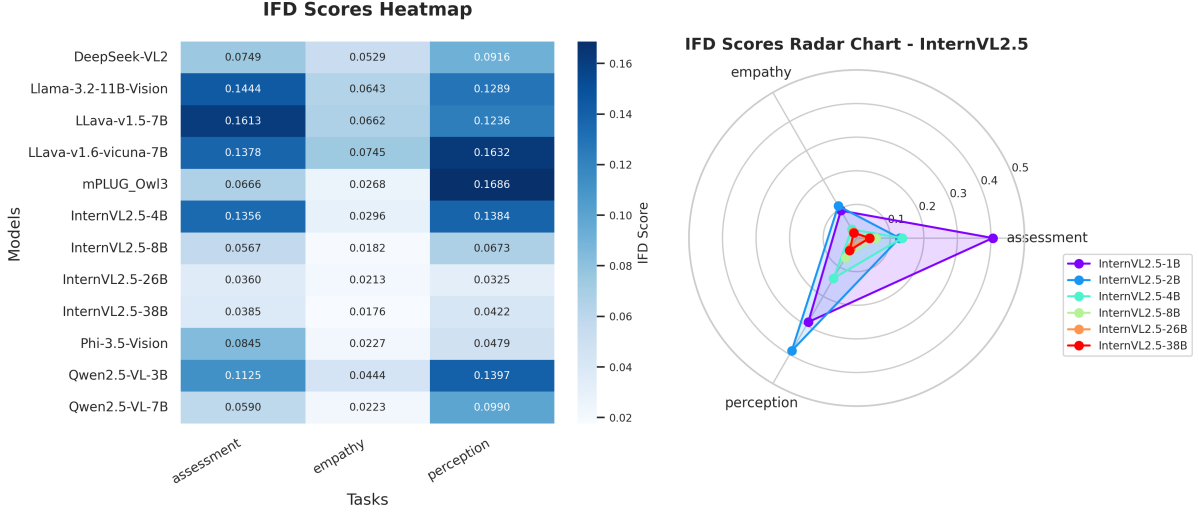


Figure 3: Left: IFD scores heatmap across diverse set of models. Right: Radar chart of IFD scores for InternVL-2.5 series models, showing variations by model size. A higher IFD indicates a greater degree of stereotype bias.

investigate societal biases in aesthetic perceptions toward the respective groups. We perform extensive testing on a well-established dataset for personalized image aesthetic assessment (Yang et al., 2022), the Personalized Image Aesthetics Database with Rich Attributes (PARA). PARA comprises 31,220 images annotated by 438 human raters with rich feature annotations. Built upon it, we generate three types of task evaluations for the 31,220 images: aesthetic perception, aesthetic assessment, and emotional perception. The three tasks are evaluated by IFD, NRD, AAS, and similarity score to examine both stereotype bias and aesthetic alignment bias.

To align with human raters, who typically judge based on discrete text-defined levels, we convert continuous scores in the PARA dataset into discrete rating levels. In AesBiasBench (ABB), this ensures consistency with the output formats Δ and facilitates fair comparisons between model outputs and human judgments. We adopt equidistant intervals to convert scores into rating levels as Wu et al. (2023), which is to uniformly divide the range between the highest score (M) and lowest score (m) into three distinct intervals.

$$L(s) = l_i \quad (9)$$

where $m + \frac{i-1}{3} \times (M - m) < s \leq m + \frac{i}{3} \times (M - m)$.

Models. In this work, we investigate a diverse set of models, including **InternVL2.5** (1B, 2B, 4B, 8B, 26B, 38B) (Chen et al., 2024b,c,d), **Qwen2.5-VL** (3B, 7B) (Yang et al., 2024), **LLaVA** (v1.5-7B,

v1.6-vicuna-7B) (Liu et al., 2023a,b), **LLaMA-3.2** (11B-vision-instruct) (Grattafiori et al., 2024), **mPLUG-Owl3** (7B) (Ye et al., 2024), **Mono-InternVL** (2B) (Luo et al., 2024), **Phi-3.6-Vision** (Abdin et al., 2024), **GLM-4V** (9B) (GLM et al., 2024), and **DeepSeek** (VL2) (Wu et al., 2024b). We also include closed-source models such as **Claude-3.5-Sonnet**, **Gemini-2.0-flash** and **GPT-4o** in our analysis. This comprehensive selection enables a systematic evaluation of biases across a wide range of architectures and scales. With this setup, we can compare bias variations within the same model series across different sizes, as well as across models of similar sizes. Such comparisons provide deeper insights into how model architecture, scale, and training paradigms influence bias.

4.2 Stereotype Bias Analysis

4.2.1 Existence of Bias in MLLMs

We quantify stereotype bias in MLLMs performing PIAA using two metrics: Identity Fairness Deviation (IFD) and Normalized Response Deviation (NRD). The heatmap in Figure 3 shows the IFD scores across multiple models, indicating substantial identity-related biases, where higher IFD values reflect stronger bias. Among these, the InternVL2.5 model series consistently shows lower IFD values, suggesting better fairness across demographic identities.

Additionally, Figure 4 (left) illustrates NRD scores, confirming strong biases, particularly evident in empathy-driven aesthetic tasks. Gender is

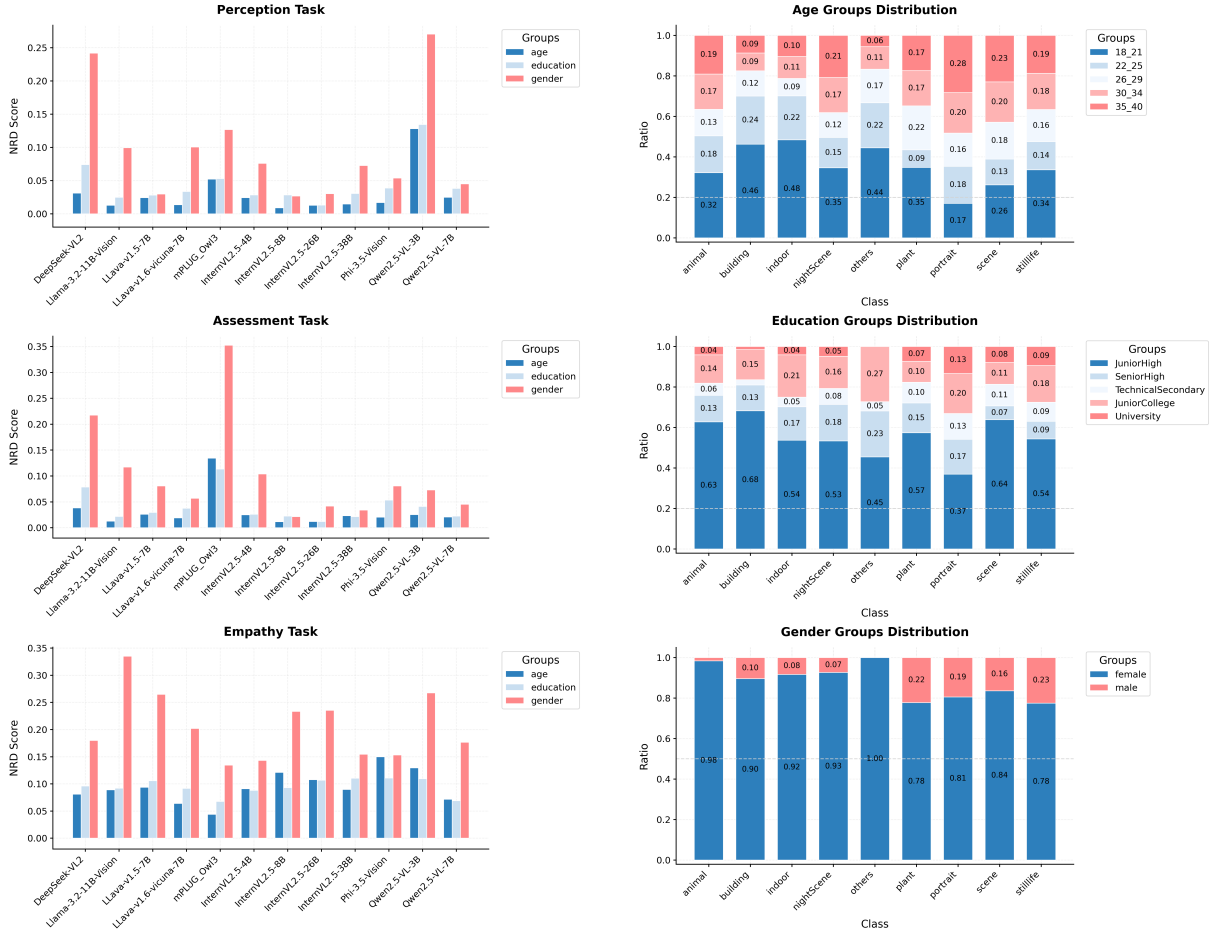


Figure 4: Left: NRD scores for age, gender, and education across three tasks. Right: $F_{g,k}^m$ scores of fear emotion from different groups for aesthetic empathy task in Claude-3.5-Sonnet, illustrating stereotype bias.

consistently identified as a major influencing factor, with notably higher NRD scores across all evaluated models. This emphasizes significant differences in the emotional perception of images among different demographic groups.

To further illustrate this, Figure 4 right provides a detailed example using Claude-3.5-Sonnet in the empathy task. The model predicts that younger individuals, those with lower education levels, and females are more likely to exhibit fear responses. These results suggest that advanced models encode systematic differences in emotional aesthetic judgment across demographic groups in emotional aesthetic judgment, reinforcing the presence of subtle but persistent stereotypical biases in MLLMs.

4.2.2 Impact of Model Size on Bias

The radar chart in Figure 3 right shows the IFD scores across the InternVL2.5 series. The results reveal a clear inverse relationship between model size and stereotype bias: as the model size increases from 1B to 38B, the IFD scores consistently de-

crease. InternVL2.5-1B shows the highest level of bias, followed by 2B, 4B, and 8B, with each larger model displaying progressively lower bias. The largest models, 26B and 38B, yield the most stable and fair outputs. This trend indicates that identity-related bias decreases consistently with increasing model size.

This pattern is not limited to the InternVL2.5 series. Similar trends are observed in other model families, where smaller variants consistently exhibit higher IFD scores than their larger counterparts, indicating stronger stereotype bias. While this may appear to reflect the effect of model capacity alone, it is likely influenced by differences in training data scale and diversity as well. Larger models are often trained on broader and more balanced datasets, which may provide better coverage of identity-related variations and contribute to more equitable outputs.

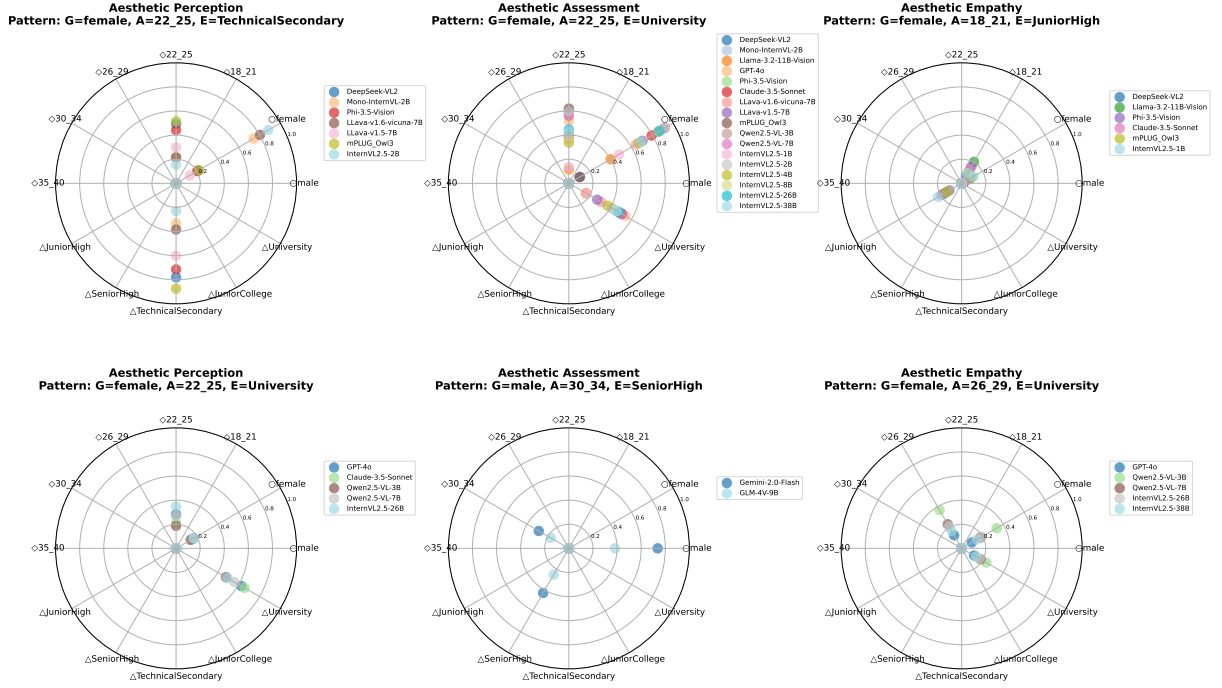


Figure 5: AAS of the model on three tasks without identity information, showing the two most common identity patterns for each task. $\circ$, $\diamond$, and $\triangle$ represent groups by gender, age, and education respectively.

4.3 Aesthetic Alignment Analysis

4.3.1 Default Aesthetic Preferences of Models

We begin by analyzing the default aesthetic alignment of MLLMs when no identity information is provided in the prompt. Using the Aesthetic Alignment Score (AAS), we measure the similarity between model outputs and the aesthetic preferences of different demographic groups across the three tasks.

The heatmap and radar plots in Figure 5 and summary statistics in Table 1 reveal clear and consistent demographic biases across tasks. All three tasks show a strong alignment with **female** aesthetic preferences, with 17 out of 19 models exhibiting this pattern. In terms of age, the **22–25** group dominates in Perception and Assessment, while Empathy shows a shift toward the younger 18–21 group. Educational alignment is more task-specific. The most consistent pattern appears in the Assessment task, where nearly all models align with the same group: **female**, aged **22–25**, with a **university** education.

Task-specific patterns also emerge. As shown in Figure 5, the points in the radar plot for the Empathy task are more tightly clustered, indicating that the AAS values are generally lower compared to the other tasks. This aligns with the observation

	Perception	Assessment	Empathy
Gender	female (17)	female (17)	female (17)
Age	22_25 (12)	22_25 (17)	18_21 (8)
Education	Tech (7)	University (17)	Junior (7)

Table 1: The number of models exhibiting the highest AAS with different demographic groups across three tasks. The table summarizes results from 19 models.

in Figure 3, where the Empathy task also exhibits lower IFD values. Together, these results show that the default models are more fair in the Empathy task and exhibit weaker alignment with human aesthetic preferences.

4.3.2 Sensitivity to Identity in Aesthetic Preferences

To further examine how identity information influences aesthetic alignment, we analyze the consistency of identity patterns across tasks after explicitly including demographic attributes in the prompts.

As shown in Figure 6 and summary statistics in Table 2, adding explicit identity information reduces the number of models that share the same dominant aesthetic pattern. This shift reflects that model outputs are sensitive to demographic descriptors, indicating the absence of neutral or identity-invariant behavior. It indicates that aesthetic out-

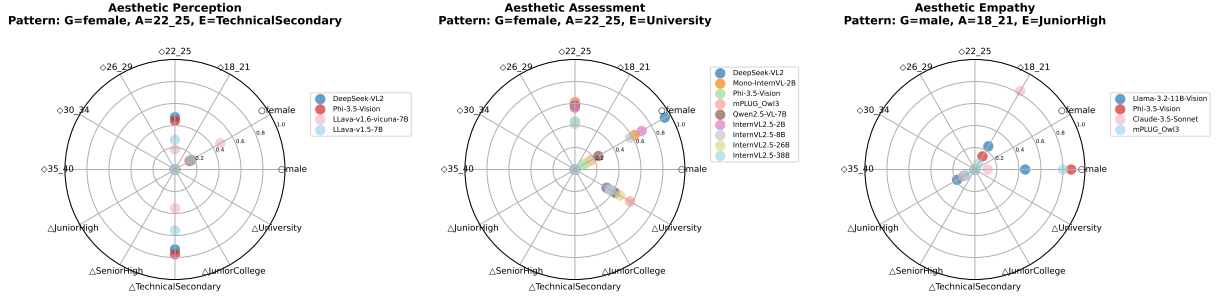


Figure 6: AAS of the model on three tasks with identity information, showing the two most common identity patterns for each task. $\circ$, $\diamond$, and $\triangle$ represent groups by gender, age, and education respectively.

	Perception	Assessment	Empathy
Gender	female (15)	female (14)	male (17)
Age	22_25 (10)	22_25 (15)	30_34 (8)
Education	Junior (7)	University (10)	University (6)

Table 2: The number of models exhibiting the highest AAS with different demographic groups across three tasks when explicit identity attributes are provided. The table summarizes results from 19 models.

Model	$\Delta S_E(M)$	$\Delta S_E(F)$	Δ
DeepSeek-VL2	-0.0535	-0.0749	0.0214
Llama-3.2-11B-Vision	0.0074	-0.0021	0.0095
GPT-4o	0.0395	-0.0748	0.1143
Phi-3.5-Vision	-0.0113	-0.0293	0.0180
Claude-3.5-Sonnet	0.1180	-0.1166	0.2346
Gemini-2.0-Flash	0.0274	-0.3780	0.4054
GLM-4V-9B	-0.0743	-0.1015	0.0272
mPLUG_Owl3	-0.0047	-0.0226	0.0179
Qwen2.5-VL-3B	0.0330	0.0196	0.0134
Qwen2.5-VL-7B	0.0287	0.0198	0.0089
InternVL2.5-1B	0.0220	-0.0244	0.0464
InternVL2.5-2B	0.0187	-0.1971	0.2158
InternVL2.5-4B	0.0085	-0.0022	0.0107
InternVL2.5-8B	-0.0007	-0.0126	0.0119
InternVL2.5-26B	-0.0160	-0.0324	0.0164

Table 3: $\Delta S_E(M)$ and $\Delta S_E(F)$ denote the changes in similarity scores to male and female aesthetic preferences, respectively, after adding gender identity to the prompt in the empathy task. Δ represents the incremental gain of male over female, computed as $\Delta S_E(M) - \Delta S_E(F)$. The top 3 highest and lowest Δ values are highlighted using soft red and blue gradients.

puts are systematically influenced by identity descriptors, revealing latent social biases in the models.

In particular, Table 2 shows a striking shift in the Empathy task: 17 models align with **male** identities, which is a complete reversal from the identity-agnostic setting, where 17 models had aligned with **female**. Table 3 illustrates this bias sensitivity, showing increased alignment with male preferences when gender is added.

As shown in Table 3, most models show a greater increase in similarity to male preferences after gender is specified, indicating higher sensitivity to male identity. Instead of exposing more balanced behavior, the inclusion of gender information reveals stronger model bias, with responses becoming more aligned to **male**-associated aesthetic patterns—a deviation possibly reflecting differences in training data composition or architectural design.

5 Conclusion

This paper introduced **AesBiasBench**, a benchmark for evaluating biases in MLLMs on PIAA tasks. To quantify stereotype bias, we proposed two metrics: IFD and NRD. In addition, we used the AAS to measure how model outputs correspond to human aesthetic preferences across demographic groups. Key findings include: (1) Stereotype bias is prevalent across models, with smaller models showing more pronounced deviations and larger

models exhibiting lower IFD and NRD scores, indicating increased fairness with scale. (2) Model outputs align disproportionately with certain demographic groups, notably, female individuals aged 22–25 with university education—even when identity information is not provided. (3) Adding identity descriptors amplifies existing biases, as shown in the empathy task where alignment shifts more strongly toward male preferences, revealing heightened sensitivity to demographic cues rather than neutrality. These results highlight the importance of identity-aware evaluation and point to the need for fairness-oriented design in future MLLMs used for subjective and socially-influenced tasks.

limitations

Our AesBiasBench evaluates existing MLLMs along two complementary axes: (1) stereotype bias and (2) human preference alignment. To make the results more reliable, we identify two possible limitations:

First, the analysis is restricted to three identity attributes: age, gender, and education. While these dimensions capture important aspects of demographic variation, other factors, such as culture, race, and religion, may also influence aesthetic preferences and model behavior. Incorporating a broader range of identity dimensions could enable a more comprehensive understanding of demographic bias in MLLMs.

Second, we evaluate 19 MLLMs, including proprietary models (e.g., GPT-4o, Claude-3.5-Sonnet) and open-source models (e.g., InternVL2.5 and Qwen2.5-VL series). While this selection spans a range of model families and sizes, future work could explore a broader set of architectures, training strategies, and deployment contexts, which may reveal additional forms of bias or alternative alignment.

Ethics

In this study, we constructed AesBiasBench using the publicly accessible Personalized Image Aesthetics Database with Rich Attributes (PARA). No original data collection was conducted; all analyses relied solely on pre-existing dataset resources. To the best of our knowledge, the PARA dataset was developed in strict adherence to academic and scientific data collection protocols, ensuring compliance with ethical standards for research involving human subjects.

Our research does not involve any personally identifiable information (PII) or process private/sensitive user data. The demographic attributes utilized (e.g., age groups, gender, education levels) are provided in the PARA dataset as anonymized and aggregated metadata, with no individual-level data accessible. This design ensures that no participant can be re-identified through the study’s analyses.

The core objective of this research is to systematically uncover and characterize biased behaviors of multimodal large language models (MLLMs) in personalized aesthetic judgment tasks. By quantifying demographic disparities in model outputs, we aim to foster greater awareness within the research

community and contribute to the development of more equitable, transparent, and socially accountable AI systems. Our work aligns with the broader ethical imperative to promote fairness in machine learning, particularly in applications impacting human values and societal norms.

References

- Marah Abdin, Jyoti Aneja, Hany Awadalla, Ahmed Awadallah, Ammar Ahmad Awan, Nguyen Bach, Amit Bahree, Arash Bakhtiari, Jianmin Bao, Harkirat Behl, and 1 others. 2024. Phi-3 technical report: A highly capable language model locally on your phone. *arXiv preprint arXiv:2404.14219*.
- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altenschmidt, Sam Altman, Shyamal Anadkat, and 1 others. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- Jean-Baptiste Alayrac, Jeff Donahue, Pauline Luc, Antoine Miech, Iain Barr, Yana Hasson, Karel Lenc, Arthur Mensch, Katherine Millican, Malcolm Reynolds, and 1 others. 2022. Flamingo: a visual language model for few-shot learning. *Advances in neural information processing systems*, 35:23716–23736.
- Dongping Chen, Ruoxi Chen, Shilin Zhang, Yinyu Liu, Yaochen Wang, Huichi Zhou, Qihui Zhang, Yao Wan, Pan Zhou, and Lichao Sun. 2024a. Mllm-as-a-judge: Assessing multimodal llm-as-a-judge with vision-language benchmark. *arXiv preprint arXiv:2402.04788*.
- Zhe Chen, Weiyun Wang, Yue Cao, Yangzhou Liu, Zhangwei Gao, Erfei Cui, Jinguo Zhu, Shenglong Ye, Hao Tian, Zhaoyang Liu, and 1 others. 2024b. Expanding performance boundaries of open-source multimodal models with model, data, and test-time scaling. *arXiv preprint arXiv:2412.05271*.
- Zhe Chen, Weiyun Wang, Hao Tian, Shenglong Ye, Zhangwei Gao, Erfei Cui, Wenwen Tong, Kongzhi Hu, Jiapeng Luo, Zheng Ma, and 1 others. 2024c. How far are we to gpt-4v? closing the gap to commercial multimodal models with open-source suites. *arXiv preprint arXiv:2404.16821*.
- Zhe Chen, Jiannan Wu, Wenhui Wang, Weijie Su, Guo Chen, Sen Xing, Muyan Zhong, Qinglong Zhang, Xizhou Zhu, Lewei Lu, and 1 others. 2024d. Internvl: Scaling up vision foundation models and aligning for generic visual-linguistic tasks. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 24185–24198.
- Chaoran Cui, Wenya Yang, Cheng Shi, Meng Wang, Xiushan Nie, and Yilong Yin. 2020. Personalized image quality assessment with social-sensed aesthetic preference. *Information Sciences*, 512:780–794.
- Yubin Deng, Chen Change Loy, and Xiaoou Tang. 2017. Image aesthetic assessment: An experimental survey. *IEEE Signal Processing Magazine*, 34(4):80–106.
- J Dhamala, T Sun, V Kumar, S Krishna, Y Puk-sachatkun, K-W Chang, and R Gupta. 2021. Bold: Dataset and metrics for measuring biases in open-ended language generation. In *Proceedings of the 2021 ACM conference on Fairness, Accountability, and Transparency*, pages 862–872.
- Sagnik Dhar, Vicente Ordonez, and Tamara L Berg. 2011. High level describable attributes for predicting aesthetics and interestingness. In *CVPR 2011*, pages 1657–1664. IEEE.
- Omkar Dige, Diljot Singh, Tsz Fung Yau, Qixuan Zhang, Borna Bolandraftar, Xiaodan Zhu, and Faiza Khan Khattak. 2024. Mitigating social biases in language models through unlearning. *arXiv preprint arXiv:2406.13551*.
- Isabel O Gallegos, Ryan A Rossi, Joe Barrow, Md Mehrab Tanjim, Sungchul Kim, Franck Dernoncourt, Tong Yu, Ruiyi Zhang, and Nesreen K Ahmed. 2024. Bias and fairness in large language models: A survey. *Computational Linguistics*, pages 1–79.
- Team GLM, Aohan Zeng, Bin Xu, Bowen Wang, Chenhui Zhang, Da Yin, Diego Rojas, Guanyu Feng, Hanlin Zhao, Hanyu Lai, Hao Yu, Hongning Wang, Jiadai Sun, Jiajie Zhang, Jiale Cheng, Jiayi Gui, Jie Tang, Jing Zhang, Juanzi Li, and 37 others. 2024. *Chatglm: A family of large language models from glm-130b to glm-4 all tools*. *Preprint*, arXiv:2406.12793.
- Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, and 1 others. 2024. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*.
- Yue Guo, Yi Yang, and Ahmed Abbasi. 2022. Auto-debias: Debiasing masked language models with automated biased prompts. In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1012–1023.
- Jingwen Hou, Weisi Lin, Guanghui Yue, Weide Liu, and Baoquan Zhao. 2022. Interaction-matrix based personalized image aesthetics assessment. *IEEE Transactions on Multimedia*, 25:5263–5278.
- Yipo Huang, Quan Yuan, Xiangfei Sheng, Zhichao Yang, Haoning Wu, Pengfei Chen, Yuzhe Yang, Leida Li, and Weisi Lin. 2024. Aesbench: An expert benchmark for multimodal large language models on image aesthetics perception. *arXiv preprint arXiv:2401.08276*.
- Yukun Jiang, Zheng Li, Xinyue Shen, Yugeng Liu, Michael Backes, and Yang Zhang. 2024. Modscan: Measuring stereotypical bias in large vision-language models from vision and language modalities. *arXiv preprint arXiv:2410.06967*.
- Abhishek Kumar, Sarfaroze Yunusov, and Ali Emami. 2024. Subtle biases need subtler measures: Dual metrics for evaluating representative and affinity bias in large language models. *arXiv preprint arXiv:2405.14555*.
- Cheng Li, Mengzhuo Chen, Jindong Wang, Sunayana Sitaram, and Xing Xie. 2024a. Culturellm: Incorporating cultural differences into large language models. *arXiv preprint arXiv:2402.10946*.

678	Cheng Li, Damien Teney, Linyi Yang, Jindong Wang,	1 others. 2022. Emergent abilities of large language	732
679	Qingsong Wen, Xing Xie, and Jindong Wang. 2024b.	models. <i>arXiv preprint arXiv:2206.07682</i> .	733
680	Culturepark: Boosting cross - cultural understand-		
681	ing in large language models. <i>arXiv preprint</i>	Haoning Wu, Zicheng Zhang, Erli Zhang, Chaofeng	734
682	<i>arXiv:2405.15145</i> .	Chen, Liang Liao, Annan Wang, Kaixin Xu, Chunyi	735
683	Leida Li, Hancheng Zhu, Sicheng Zhao, Guiguang Ding,	Li, Jingwen Hou, Guangtao Zhai, and 1 others. 2024a.	736
684	and Weisi Lin. 2020. Personality-assisted multi-task	Q-instruct: Improving low-level visual abilities for	737
685	learning for generic and personalized image aesthet-	multi-modality foundation models. In <i>Proceedings</i>	738
686	ics assessment. <i>IEEE Transactions on Image Pro-</i>	<i>of the IEEE/CVF Conference on Computer Vision</i>	739
687	<i>cessing</i> , 29:3898–3910.	<i>and Pattern Recognition</i> , pages 25490–25500.	740
688	Luyang Lin, Lingzhi Wang, Jinsong Guo, and Kam-	Haoning Wu, Zicheng Zhang, Weixia Zhang, Chaofeng	741
689	Fai Wong. 2024. Investigating bias in llm-based	Chen, Liang Liao, Chunyi Li, Yixuan Gao, An-	742
690	bias detection: Disparities between llms and human	nan Wang, Erli Zhang, Wenxiu Sun, and 1 others.	743
691	perception. <i>arXiv preprint arXiv:2403.14896</i> .	2023. Q-align: Teaching llms for visual scor-	744
692	Haotian Liu, Chunyuan Li, Yuheng Li, and Yong Jae	ing via discrete text-defined levels. <i>arXiv preprint</i>	745
693	Lee. 2023a. Improved baselines with visual instruc-	<i>arXiv:2312.17090</i> .	746
694	tion tuning.	Zhiyu Wu, Xiaokang Chen, Zizheng Pan, Xingchao	747
695	Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae	Liu, Wen Liu, Damai Dai, Huazuo Gao, Yiyang	748
696	Lee. 2023b. Visual instruction tuning. In <i>NeurIPS</i> .	Ma, Chengyue Wu, Bingxuan Wang, and 1 others.	749
697	Gen Luo, Xue Yang, Wenhan Dou, Zhaokai Wang, Ji-	2024b. Deepseek-vl2: Mixture-of-experts vision-	750
698	awen Liu, Jifeng Dai, Yu Qiao, and Xizhou Zhu.	language models for advanced multimodal under-	751
699	2024. Mono-internvl: Pushing the boundaries of	standing. <i>arXiv preprint arXiv:2412.10302</i> .	752
700	monolithic multimodal large language models with		
701	endogenous visual pre-training. <i>arXiv preprint</i>	An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui,	753
702	<i>arXiv:2410.08202</i> .	Bo Zheng, Bowen Yu, Chengyuan Li, Dayiheng Liu,	754
703	Naila Murray, Luca Marchesotti, and Florent Perronnin.	Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jian-	755
704	2012. Ava: A large-scale database for aesthetic vi-	hong Tu, Jianwei Zhang, Jianxin Yang, Jiaxi Yang,	756
705	sual analysis. In <i>2012 IEEE conference on computer</i>	Jingren Zhou, Junyang Lin, Kai Dang, and 22 oth-	757
706	<i>vision and pattern recognition</i> , pages 2408–2415.	ers. 2024. Qwen2.5 technical report. <i>arXiv preprint</i>	758
707	IEEE.	<i>arXiv:2412.15115</i> .	759
708	Tarek Naous, Michael J Ryan, Alan Ritter, and Wei	Yuzhe Yang, Liwu Xu, Leida Li, Nan Qie, Yaqian Li,	760
709	Xu. 2023. Having beer after prayer? measuring	Peng Zhang, and Yandong Guo. 2022. Personalized	761
710	cultural bias in large language models. <i>arXiv preprint</i>	image aesthetics assessment with rich attributes. In	762
711	<i>arXiv:2305.14456</i> .	<i>Proceedings of the IEEE/CVF Conference on Com-</i>	763
712	Jian Ren, Xiaohui Shen, Zhe Lin, Radomir Mech, and	<i>puter Vision and Pattern Recognition</i> , pages 19861–	764
713	David J Foran. 2017. Personalized image aesthetics.	19869.	765
714	In <i>Proceedings of the IEEE international conference</i>	Jiabo Ye, Haiyang Xu, Haowei Liu, Anwen Hu, Ming	766
715	<i>on computer vision</i> , pages 638–647.	Yan, Qi Qian, Ji Zhang, Fei Huang, and Jingren Zhou.	767
716	Brandon Smith, Miguel Farinha, Siobhan Macken-	2024. mplug-owl3: Towards long image-sequence	768
717	zie Hall, Hannah Rose Kirk, Aleksandar Shtedritski,	understanding in multi-modal large language models .	769
718	and Max Bain. 2023. Balancing the picture: Debias-	<i>Preprint</i> , arXiv:2408.04840.	770
719	ing vision-language datasets with synthetic contrast	Ran Yi, Haoyuan Tian, Zhihao Gu, Yu-Kun Lai, and	771
720	sets. <i>arXiv preprint arXiv:2305.15407</i> .	Paul L Rosin. 2023. Towards artistic image aesthetics	772
721	A Tamkin, A Askill, L Lovitt, E Durmus, N Joseph,	assessment: a large-scale dataset and a new method.	773
722	S Kravec, K Nguyen, J Kaplan, and D Ganguli. 2023.	In <i>Proceedings of the IEEE/CVF Conference on Com-</i>	774
723	Evaluating and mitigating discrimination in language	<i>puter Vision and Pattern Recognition</i> , pages 22388–	775
724	model decisions. <i>arXiv preprint arXiv:2312.03689</i> .	22397.	776
725	Lindia Tjautja, Valerie Chen, Sherry Tongshuang Wu,	Nick Zangwill. 2003. Aesthetic judgment.	777
726	Ameet Talwalkar, and Graham Neubig. 2023. Do	Zhaokun Zhou, Qiulin Wang, Bin Lin, Yiwei Su, Rui	778
727	llms exhibit human-like response biases? a case study	Chen, Xin Tao, Amin Zheng, Li Yuan, Pengfei Wan,	779
728	in survey design. <i>arXiv preprint arXiv:2311.04076</i> .	and Di Zhang. 2024. Uniaa: A unified multi-modal	780
729	Jason Wei, Yi Tay, Rishi Bommasani, Colin Raffel,	image aesthetic assessment baseline and benchmark.	781
730	Barret Zoph, Sebastian Borgeaud, Dani Yogatama,	<i>arXiv preprint arXiv:2404.09619</i> .	782
731	Maarten Bosma, Denny Zhou, Donald Metzler, and	Deyao Zhu, Jun Chen, Xiaoqian Shen, Xiang Li, and	783
		Mohamed Elhoseiny. 2023. Minigpt-4: Enhancing	784
		vision-language understanding with advanced large	785
		language models. <i>arXiv preprint arXiv:2304.10592</i> .	786

Hancheng Zhu, Leida Li, Jinjian Wu, Sicheng Zhao,
Guiguang Ding, and Guangming Shi. 2020. Personal-
alized image aesthetics assessment via meta-learning
with bilevel gradient optimization. *IEEE Transac-
tions on Cybernetics*, 52(3):1798–1811.