

Hierarchical Approaches for Domain-Specific Image Captioning: Classification, Distillation, and Optimization

Anonymous ACL submission

Abstract

This paper presents a novel hierarchical approach to improve image captioning, focusing on enhancing the accuracy, contextual appropriateness, and linguistic diversity of generated captions. We address key limitations of existing multimodal models, including inaccurate object relationships, missed details, and poor domain-specific understanding. Our method incorporates three main components: a classification-guided prompting system that utilizes domain-specific knowledge, a knowledge distillation framework that transfers captioning capabilities from GPT-4o to the LLaVA model, and an iterative Direct Preference Optimization (DPO) approach that refines caption quality. Extensive experiments demonstrate that our approach outperforms existing methods, achieving near-GPT-4o performance while maintaining computational efficiency. Additionally, we release a high-quality dataset of 9,840 image-caption pairs across 18 categories, providing valuable resources for future research in domain-specific image captioning.

1 Introduction

Image captioning, the task of automatically generating natural language descriptions for images, has emerged as a crucial indicator of models' ability to understand visual and language information effectively (Mokady et al., 2021; Li et al., 2023). Despite significant advances in this field, current state-of-the-art models still face substantial challenges in generating accurate, contextually appropriate, and comprehensive descriptions.

Recent multimodal models, while demonstrating impressive capabilities across various tasks (Li et al., 2023; Liu et al., 2024), exhibit notable limitations in image captioning. These limitations manifest in several aspects: inaccurate object relationship descriptions, missed subtle but important details, and inadequate ability to express domain-specific concepts. Moreover, traditional evaluation

metrics such as BLEU (Papineni et al., 2002) and METEOR (Banerjee and Lavie, 2005) often fail to capture these nuanced aspects of performance, potentially leading to misleading assessments of model capabilities.

To address these limitations, we propose a novel hierarchical approach that combines classification-guided prompting, knowledge distillation, and iterative preference optimization. Our method consists of three key components:

- A classification-first strategy that categorizes images into semantic classes, enabling the use of class-specific prompting templates for more precise and contextually relevant caption generation.
- A knowledge distillation framework that leverages GPT-4o's superior captioning abilities to create high-quality training pairs, which are then used to fine-tune a more efficient LLaVA model¹.
- An iterative Direct Preference Optimization (DPO) approach that continuously refines the model's output quality by learning from automatically generated preference pairs using a Pixtral (Agrawal et al., 2024) model for evaluation.

Through comprehensive experiments, we demonstrate that this hierarchical approach significantly improves captioning quality across multiple dimensions. The classification-guided prompting ensures domain-specific accuracy, while the knowledge distillation from GPT-4o provides rich linguistic diversity. Furthermore, the iterative DPO refinement helps align the model's outputs with quality standards, leading to more natural and contextually appropriate descriptions.

¹LLaVA serves as our base multimodal model due to its strong vision-language capabilities and computational efficiency.

Our extensive evaluation shows that the proposed approach achieves significant improvements over existing methods. Specifically, our model achieves a Pixtral score of 4.312 after DPO optimization, approaching the performance of GPT-4o (4.473) and GPT-4o-mini (4.293), while substantially outperforming the base LLaVA model (3.598). Human evaluation strongly correlates with our automatic metrics (Pearson correlation: 0.779), validating the effectiveness of our evaluation framework. These results demonstrate the success of combining semantic classification, knowledge distillation, and preference optimization in improving image captioning quality.

Our main contributions are summarized as follows:

- We propose a novel classification-guided prompting strategy that leverages domain-specific knowledge for more accurate and contextually appropriate caption generation.
- We present an effective knowledge distillation pipeline that efficiently transfers captioning capabilities from GPT-4o to a more practical LLaVA model.
- We develop an iterative DPO framework that uses automated evaluation for continuous model improvement, demonstrating consistent performance gains over multiple iterations.
- We release a comprehensive dataset of 9,840 high-quality image-caption pairs across 18 diverse categories, where each image is paired with a GPT-4o-generated caption. This curated dataset provides valuable training resources for future research in domain-specific image captioning.
- We introduce a fine-tuned DPO version of the LLaVA model, which demonstrates strong image captioning capabilities, further enhancing the practical application of the model.

2 Related Works

2.1 Image Captioning Models

Recent advances in image captioning have been largely driven by the development of large multimodal models. Traditional approaches relied on encoder-decoder architectures (Xu et al., 2016; Anderson et al., 2018), while more recent work has

focused on end-to-end multimodal training. Models like BLIP (Li et al., 2023) and LLaVA (Liu et al., 2024) have demonstrated impressive capabilities in understanding and describing visual content. However, these models often struggle with domain-specific details and contextual accuracy, motivating our classification-guided approach.

2.2 Caption Evaluation Metrics

The evaluation of image captioning systems has evolved through several generations:

Reference-based Metrics Traditional metrics such as BLEU (Papineni et al., 2002), METEOR (Banerjee and Lavie, 2005), and CIDEr (Vedantam et al., 2015) evaluate captions by comparing them with human-written references. While widely used, these metrics often fail to capture semantic accuracy and contextual appropriateness (Anderson et al., 2016).

Learning-based Metrics More recent approaches like BERTScore (Zhang* et al., 2020) and CLIP-Score (Hessel et al., 2022) leverage pre-trained models to assess caption quality. These metrics show better correlation with human judgments but still lack explainability and domain-specific understanding.

LLM-based Evaluation The emergence of large language models has introduced new possibilities for caption evaluation. Recent work has shown that LLMs can provide more nuanced and context-aware evaluations (Liu et al., 2023). Our work builds on this direction by using the Pixtral model for automated evaluation and preference learning.

2.3 Knowledge Distillation and Model Alignment

Knowledge distillation has proven effective in transferring capabilities from larger to smaller models (Hinton et al., 2015). In the multimodal domain, recent work has shown success in distilling vision-language capabilities (Li et al., 2023). Our approach extends this to image captioning by distilling knowledge from GPT-4o to a more efficient LLaVA model.

Direct Preference Optimization DPO has emerged as a powerful technique for aligning language models with human preferences (Rafailov et al., 2024). Recent work has demonstrated its effectiveness in improving text generation quality (Liu et al., 2023). We adapt this approach to image

captioning by using automated evaluations to guide the optimization process.

2.4 Category-Specific Generation

The idea of leveraging category information for improved generation has been explored in various contexts. Previous work has shown that domain-specific prompting can improve text generation quality (Liu et al., 2024). Our work extends this concept to image captioning by introducing a classification-guided prompting strategy that adapts to different visual domains.

3 Method

Our approach consists of three main components: (1) a classification-guided prompting system, (2) a knowledge distillation pipeline from GPT-4o, and (3) an iterative DPO optimization process. Figure 1 illustrates the complete architecture of our proposed method. In this section, we describe each component in detail.

3.1 Classification-Guided Prompting

The effectiveness of image captioning heavily depends on the model’s ability to recognize and describe domain-specific content accurately. To address this challenge, we propose a classification-first approach that leverages a fine-tuned large multimodal model (LMM) for image classification, followed by category-specific prompting for caption generation.

Classification Model Training We fine-tune a LLaVA model to serve as our classifier. The training process consists of several key steps:

1) **Category Definition:** We define 18 comprehensive categories $\mathcal{C}$ covering most common image scenarios:

- Scene types: Animal, Architecture, Art, City, Landscape
- Daily activities: Daily_Life, Food, Love, Medical, Sports
- Technical: Autonomous_Driving, Smart_Cities, Transportation
- Visual content: Cartoon, Chart, Person, Plant
- Others: Other

2) **Training Data Collection:** For each category $c \in \mathcal{C}$, we collect images from multiple sources,

including Flickr, JAAD (Rasouli et al., 2017, 2018), ScienceQA (Lu et al., 2022), ChartQA (Masry et al., 2022), etc.:

$$D_{\text{train}} = (I_i, c_i)_{i=1}^N \quad (1)$$

where I_i represents an image and c_i its corresponding category label.

3) **Model Fine-tuning:** We fine-tune the LLaVA model using our collected dataset:

$$\mathcal{L}_{\text{cls}} = - \sum_{i=1}^N \log p_{\theta}(c_i | I_i) \quad (2)$$

where θ represents the model parameters.

Classification Process Given a new image I , the classification is performed as:

$$c = f_{\text{LLaVA}}(I) \in \mathcal{C} \quad (3)$$

where f_{LLaVA} is our fine-tuned LLaVA classifier.

Structured Prompt Design For each category c , we design a specialized prompting template P_c to guide the caption generation process. Representative examples of our prompting templates are provided in Appendix A, and the complete set is available in our repository.

Caption Generation The final caption generation process involves:

$$C_{\text{final}} = g_{\text{LLaVA}}(I, P_c) \quad (4)$$

where g_{LLaVA} is another LLaVA model fine-tuned on GPT-4o distilled data (detailed in Section 3.2).

3.2 Knowledge Distillation from GPT-4o

As mentioned in Section 3.1, our caption generation model g_{LLaVA} is fine-tuned through knowledge distillation from GPT-4o. We leverage the same dataset used for classifier training to create high-quality caption pairs.

High-Quality Caption Generation For each image I in our training dataset, we use GPT-4o to generate captions:

$$C_{\text{GPT-4o}} = g_{\text{GPT-4o}}(I, P_c) \quad (5)$$

where $g_{\text{GPT-4o}}$ represents the GPT-4o model’s caption generation function.

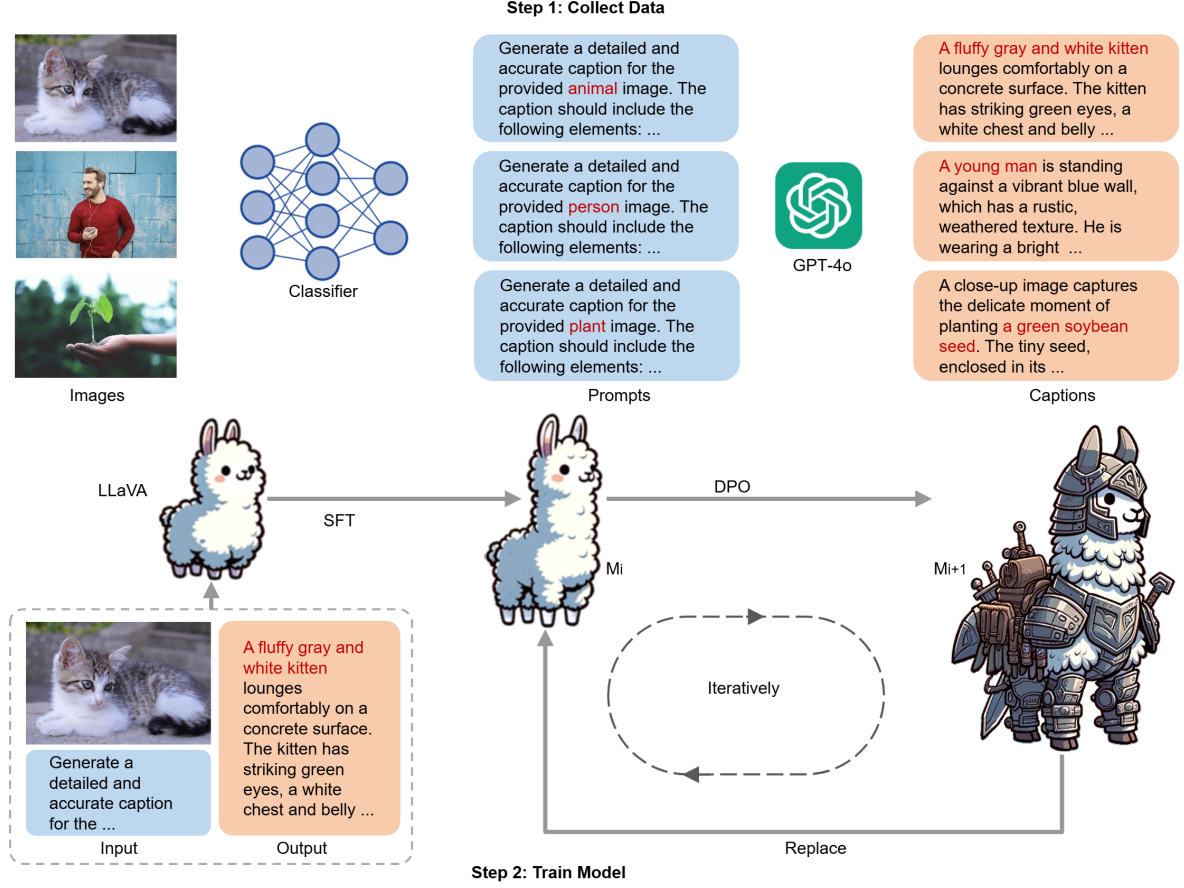


Figure 1: method

Training Data Creation We create training pairs by combining:

$$D_{\text{distill}} = \{(I_i, P_{c_i}, C_{\text{GPT-4o}, i})\}_{i=1}^N \quad (6)$$

where N is the total number of images across all categories.

LLaVA Fine-tuning We fine-tune the LLaVA model using these high-quality training pairs with both full parameter fine-tuning and LoRA approaches. For both versions, the supervised learning objective is:

$$\mathcal{L}_{\text{SFT}} = - \sum_{i=1}^N \log p_{\theta}(C_{\text{GPT-4o}, i} | I_i, P_{c_i}) \quad (7)$$

where θ represents the parameters of the LLaVA model. The detailed training configurations for both approaches are provided in Appendix B.

3.3 Iterative DPO Optimization

To further improve the quality of generated captions, we implement an iterative Direct Preference Optimization (DPO) process using a Pixtral model

for preference scoring. We leverage both full-parameter and LoRA fine-tuned models to generate diverse candidate captions.

Caption Generation and Scoring For each image, we use both the full-parameter and LoRA fine-tuned models to generate five captions each, resulting in ten candidate captions per image. These captions are then evaluated using the Pixtral model, which assigns a score between 1 and 5 to each caption following a structured evaluation prompt (detailed in Appendix C).

Preference Pair Selection From the ten candidates, we identify the highest and lowest scoring captions. If their score difference exceeds 2 points (on the 1-5 scale), we use this pair for DPO training. This threshold ensures we only learn from pairs with significant quality differences.

DPO Training The DPO objective is defined as:

$$\mathcal{L}_{\text{DPO}} = -\mathbb{E}_{(I, C_w, C_l)} [\log \sigma(\beta(r_{\theta}(C_w) - r_{\theta}(C_l)))] \quad (8)$$

where $r_{\theta}(C)$ is the reward score for caption C , β is temperature parameter, and σ is sigmoid function.

Iterative Refinement The DPO process consists of:

1. Generate ten captions per image
2. Score captions using Pixtral
3. Select pairs with significant score differences
4. Update model using DPO objective
5. Repeat until convergence

The detailed training configurations and implementation details are provided in Appendix B.

4 Experiments

4.1 Experimental Setup

Dataset We collect and curate a new dataset from multiple sources including Flickr, JAAD, ScienceQA, ChartQA, etc., covering 18 diverse categories including animal, architecture, art, autonomous driving, cartoon, chart, city, daily life, food, landscape, love, medical, other, person, plant, smart cities, sports, and transportation. The dataset consists of 9,840 training images and 796 validation images, with approximately balanced distribution across categories. All images are manually verified to ensure quality and correct categorization. The validation set maintains a similar category distribution as the training set to ensure fair evaluation.

For each image in the training set, we generate a high-quality caption using GPT-4o following our category-specific prompting strategy. These captions are carefully crafted to capture domain-specific details and contextual information relevant to each category. This results in a rich dataset where each image is paired with a detailed, contextually appropriate caption that can serve as high-quality training data for future research.

The dataset is designed with several key features:

- **Category Balance:** Training and validation sets maintain approximately balanced distribution across all categories, ensuring comprehensive coverage across domains.
- **Quality Control:** All images undergo manual verification for visual quality and category correctness.
- **Domain Diversity:** The 18 categories cover a wide range of scenarios from daily life to specialized domains like medical imaging and autonomous driving.

- **High-Quality Captions:** Each training image is paired with a GPT-4o generated caption that follows category-specific guidelines.

We plan to release this dataset along with our code and models to facilitate future research in domain-specific image captioning.

Evaluation Metrics We employ two types of evaluation:

- **Automated Evaluation:** The Pixtral model serves as our primary metric, assigning scores from 1 to 5 based on caption quality. We evaluate all generated captions on the entire validation set of 796 images.
- **Human Evaluation:** Three expert annotators evaluate 54 randomly sampled image-caption pairs using the same 1-5 scale. The sample size was chosen to ensure thorough evaluation while maintaining feasible annotation effort.

Baselines We compare our method with several strong baselines:

- **GPT-4o:** A large-scale language model developed by OpenAI, known for its advanced natural language understanding and generation capabilities. It serves as a benchmark for state-of-the-art performance in various language tasks.
- **GPT-4o-mini:** A more compact version of GPT-4o, designed to offer competitive performance with reduced computational requirements. This model maintains a balance between efficiency and effectiveness, making it suitable for applications with limited resources.
- **Claude-3-Haiku:** A large language model developed by Anthropic, optimized for speed and affordability. It is designed to process 21,000 tokens per second for prompts under 32,000 tokens, making it suitable for enterprise workloads that require quick analysis of large datasets.
- **LLaVA-v1.5-7B:** An open-source multimodal model fine-tuned from LLaMA and Vicuna on GPT-generated multimodal instruction-following data. It is an auto-regressive language model based on the transformer architecture, trained to understand and generate

both text and images. The model was trained in September 2023 and is licensed under the LLAMA 2 Community License.

- **LLaVA-FT**: A version of LLaVA model that has undergone full-parameter fine-tuning on our curated dataset of 9,840 image-caption pairs, trained without Direct Preference Optimization (DPO). This model represents the effectiveness of traditional fine-tuning approaches on our domain-specific dataset.
- **LLaVA-LoRA**: A variant of the LLaVA model fine-tuned on our dataset using Low-Rank Adaptation (LoRA) techniques, without DPO. This approach demonstrates the efficacy of parameter-efficient fine-tuning methods on our carefully curated training data while maintaining model performance.
- **Qwen-VL-Chat** (Bai et al., 2023): A multimodal chatbot model that integrates vision and language understanding, enabling it to process and generate responses based on both textual and visual inputs. This model is designed to handle a wide range of multimodal tasks, including image captioning and visual question answering.
- **Qwen2-VL-7B-Instruct** (Wang et al., 2024): An advanced multimodal model that builds upon the Qwen-VL-Chat architecture, incorporating instruction-following capabilities to improve its performance on tasks requiring specific guidance. The 7B refers to the model’s parameter size, indicating a balance between computational efficiency and performance.

4.2 Main Results

The experimental results reveal several comprehensive findings:

Base Model Comparisons: Among the models, we observe a clear performance hierarchy. While the closed-source GPT-4o maintains the highest performance (4.473), setting a strong upper bound, other closed-source models like GPT-4o-mini (4.293) and Claude-3-Haiku (4.268) also show strong performance. For open-source models, Qwen2-VL-7B-Instruct achieves 4.224 and Qwen-VL-Chat reaches 3.833. LLaVA-v1.5-7B serves as our baseline with a score of 3.598.

Fine-tuning Effectiveness: Our fine-tuning approaches show significant improvements over the

Model	Pixtral Score	Δ	#Params
GPT-4o	4.473	-	-
GPT-4o-mini	4.293	-	-
Claude-3-Haiku	4.268	-	-
LLaVA-v1.5-7B (Base)	3.598	-	7B
Qwen-VL-Chat	3.833	+0.235	7B
Qwen2-VL-7B-Instruct	4.224	+0.626	7B
LLaVA-FT (Full)	4.078	+0.480	7B
LLaVA-LoRA	4.034	+0.436	7B
DPO (1 iteration)	4.113	+0.515	7B
DPO (2 iterations)	4.168	+0.570	7B
DPO (3 iterations)	4.290	+0.692	7B
DPO (4 iterations)	4.312	+0.714	7B

Table 1: Performance comparison on the validation set (796 images). Scores range from 1 to 5. Δ shows absolute improvement over the base model. #Params shows the total number of parameters.

base model. Full parameter fine-tuning achieves a score of 4.078 (+0.480 over base), demonstrating the effectiveness of comprehensive model adaptation. LoRA fine-tuning reaches 4.034 (+0.436 over base), nearly matching full fine-tuning while being parameter-efficient. Both approaches outperform Qwen-VL-Chat but still trail behind more advanced models like Claude-3-Haiku and Qwen2-VL-7B-Instruct, suggesting room for further improvement.

DPO Improvement: The iterative DPO process shows consistent and meaningful improvements. First iteration achieves +0.515 over base model, surpassing both fine-tuning approaches. Subsequent iterations show continuous improvement, with the second iteration adding +0.055 (reaching 4.168), the third iteration showing the largest incremental gain of +0.122 (achieving 4.290), and the fourth iteration providing a modest +0.022 improvement (reaching 4.312). The diminishing returns after the third iteration suggest convergence of the optimization process. Notably, our final DPO model outperforms both Qwen2-VL-7B-Instruct by +0.088 points and Claude-3-Haiku by +0.044 points, demonstrating the effectiveness of our approach.

Gap Analysis with Leading Models: Our final model (DPO 4 iterations) achieves a score of 4.312, which is particularly noteworthy as it surpasses both closed-source models (GPT-4o-mini and Claude-3-Haiku) and open-source models (Qwen2-VL-7B-Instruct and Qwen-VL-Chat). While there remains a small gap of 0.161 points from GPT-4o, our model outperforms GPT-4o-mini by +0.019 points, Claude-3-Haiku by +0.044

points, and the best open-source model Qwen2-VL-7B-Instruct by +0.088 points. This achievement is especially significant considering we maintain the computational efficiency of a 7B parameter model, demonstrating that our approach effectively narrows the performance gap between open-source and closed-source models in the field of image captioning.

Computational Efficiency: All our models maintain the original 7B parameter count, offering practical deployment capabilities compared to larger models like GPT-4o and Claude-3-Haiku. This enables efficient inference time while achieving competitive performance, providing a scalable solution for real-world applications.

4.3 Human Validation Study

To validate the reliability of Pixtral as an automated evaluation metric, we conducted a comprehensive human evaluation study. The study included 54 image-caption pairs randomly sampled from the validation set, evaluated by three experts with ML/CV backgrounds and experience in image captioning. Each annotator independently scored all pairs following detailed evaluation criteria, using a 1-5 scale with 1 point increments, which aligns with the Pixtral scoring scale.

Correlation Analysis The comparison between human and Pixtral scores shows strong alignment with a Pearson correlation of 0.762 (p-value: 0.000), Kendall’s Tau of 0.545, and Spearman correlation of 0.595. These strong correlations across different statistical measures validate Pixtral’s effectiveness as an automated evaluation metric.

Error Analysis The difference between Pixtral and human scores demonstrates reasonable agreement, with a Mean Absolute Error (MAE) of 0.414 and Root Mean Square Error (RMSE) of 0.572. These values indicate that Pixtral’s scores typically deviate from human consensus by less than half a point on the 5-point scale, suggesting strong practical reliability.

Inter-annotator Agreement The Cohen’s Kappa coefficients between annotators are 0.507, 0.617, and 0.466 for pairs 1-2, 1-3, and 2-3, respectively. These values reflect moderate to good inter-annotator agreement, highlighting that while there are some subjective differences in caption evaluation, the overall consistency is still acceptable. These results further underscore the

advantage of using a consistent automated metric like Pixtral, and the importance of averaging multiple human judgments to mitigate individual biases.

4.4 Qualitative Analysis

We demonstrate our model’s capability through several test cases on images not present in our training set. As shown in Fig. 1, our classification-guided prompting effectively handles various domain-specific characteristics. For animal images, our model accurately captures physical attributes (e.g., "fluffy gray and white kitten") and behavioral aspects (e.g., "lounges comfortably"). For person descriptions, it successfully combines human attributes with environmental context (e.g., "standing against a vibrant blue wall"). In plant-related cases, it shows capability in describing fine-grained details (e.g., "green soybean seed"). These examples demonstrate our model’s robust performance on novel images. However, challenges remain in several aspects. The model sometimes struggles with complex scenes involving multiple objects or actions, may have difficulty with rare or unusual visual concepts, and occasionally misses subtle contextual details. These limitations suggest directions for future improvements in our approach.

5 Discussion and Conclusion

5.1 Key Findings

Our experimental results support several key conclusions:

- 1) The effectiveness of our approach in improving caption quality, evidenced by both automated and human evaluation. Through classification-guided prompting, knowledge distillation, and DPO refinement, our method demonstrates consistent improvements over the base model.
- 2) The validity of Pixtral as an automated evaluation metric, supported by strong correlation with human judgments. This validation enables efficient automatic evaluation of large-scale image captioning systems.
- 3) The efficiency of our pipeline, with our final model (DPO 4 iterations) achieving strong performance (4.312) while maintaining practical deployment requirements. This demonstrates the effectiveness of our knowledge distillation approach in transferring capabilities from larger models.
- 4) The complementary benefits of different components: classification-guided prompting for lan-

guage quality improvement, knowledge distillation for model capability transfer, and DPO for continuous optimization. Each component contributes to the overall performance gain, leading to consistent improvements over the base model.

5) We introduce a fine-tuned DPO version of the LLaVA model, which demonstrates strong image captioning capabilities, further enhancing the practical application of the model.

5.2 Dataset Contribution

A significant outcome of this work is our curated dataset, which includes 9,840 training images and 796 validation images, spanning 18 diverse categories. This dataset offers several unique advantages:

- **High-quality GPT-4o-generated captions:** Each image is paired with a detailed caption generated by the GPT-4o model, which accurately captures domain-specific details. We have carefully designed category-specific prompting strategies to ensure that the descriptions reflect the unique context and information of each category.
- **Balanced category distribution:** We ensured that both the training and validation sets have a balanced distribution of categories, ensuring comprehensive coverage across all categories. This balance helps the model perform well across different types of images, avoiding bias towards any specific category and improving generalizability.
- **Manual verification ensuring data quality and category accuracy:** All images have undergone manual verification to ensure their quality and the accuracy of their category labels. Expert reviewers manually checked each image to ensure it is correctly categorized and meets the required standards, eliminating the potential impact of incorrect labels.
- **Diverse domain coverage:** The dataset spans a wide range of domains, from common scenarios to specialized fields, including animal, architecture, art, daily life, food, medical, smart cities, and more. These categories include everyday scenes as well as specialized fields such as smart transportation, healthcare, and autonomous driving.

- **Rich Contextual Information and Details:** The captions provided with each image extend beyond mere visual descriptions, incorporating comprehensive contextual information. These captions transcend simple visual representation, integrating broader background and nuanced details to provide profound insights into the scene, environment, and wider context.

We believe that this dataset will serve as a valuable resource for future research in domain-specific image captioning and multimodal understanding. It provides researchers with a high-quality, manually verified data foundation that can significantly advance research in image understanding, caption generation, and cross-modal learning. Additionally, the dataset’s diversity and high-quality captions offer a rich source of experimental data for training and evaluating domain-specific generation models, contributing to the further development of more refined models for image captioning.

5.3 Conclusion

We have presented a hierarchical approach to image captioning that combines classification-guided prompting, knowledge distillation, and iterative preference optimization. Our method demonstrates significant improvements over baseline approaches, achieving performance comparable to much larger models while maintaining practical deployment requirements. We also introduce a fine-tuned DPO version of the LLaVA model, which further enhances the practical application of the model for image captioning. Together with our released dataset, this work contributes both methodological advances and valuable resources to the field of image captioning. We believe these contributions will facilitate future research in domain-specific visual understanding and caption generation.

5.4 Limitations and Future Work

Although our approach performs well, it has some limitations. The current classification system with 18 categories may not capture all image scenarios, and the model faces challenges with complex scenes containing multiple objects. Additionally, the DPO method shows slow convergence in later iterations. Future work should focus on improving classification systems and enhancing the model’s ability to handle complex scenes.

References

- Pravesh Agrawal, Szymon Antoniak, Emma Bou Hanna, Baptiste Bout, Devendra Chaplot, Jessica Chudnovsky, Diogo Costa, Baudouin De Monicault, Saurabh Garg, Theophile Gervet, Soham Ghosh, Amélie Héliou, Paul Jacob, Albert Q. Jiang, Kartik Khandelwal, Timothée Lacroix, Guillaume Lample, Diego Las Casas, Thibaut Lavril, Teven Le Scao, Andy Lo, William Marshall, Louis Martin, Arthur Mensch, Pavankumar Muddireddy, Valera Nemychnikova, Marie Pellat, Patrick Von Platen, Nikhil Raghuraman, Baptiste Rozière, Alexandre Sablayrolles, Lucile Saulnier, Romain Sauvestre, Wendy Shang, Roman Soletskyi, Lawrence Stewart, Pierre Stock, Joachim Studnia, Sandeep Subramanian, Sagar Vaze, Thomas Wang, and Sophia Yang. 2024. *Pixtral 12b*. *Preprint*, arXiv:2410.07073.
- Peter Anderson, Basura Fernando, Mark Johnson, and Stephen Gould. 2016. *Spice: Semantic propositional image caption evaluation*. *Preprint*, arXiv:1607.08822.
- Peter Anderson, Xiaodong He, Chris Buehler, Damien Teney, Mark Johnson, Stephen Gould, and Lei Zhang. 2018. *Bottom-up and top-down attention for image captioning and visual question answering*. *Preprint*, arXiv:1707.07998.
- Jinze Bai, Shuai Bai, Shusheng Yang, Shijie Wang, Sinan Tan, Peng Wang, Junyang Lin, Chang Zhou, and Jingren Zhou. 2023. Qwen-vl: A frontier large vision-language model with versatile abilities. *arXiv preprint arXiv:2308.12966*.
- Satanjeev Banerjee and Alon Lavie. 2005. *METEOR: An automatic metric for MT evaluation with improved correlation with human judgments*. In *Proceedings of the ACL Workshop on Intrinsic and Extrinsic Evaluation Measures for Machine Translation and/or Summarization*, pages 65–72, Ann Arbor, Michigan. Association for Computational Linguistics.
- Jack Hessel, Ari Holtzman, Maxwell Forbes, Ronan Le Bras, and Yejin Choi. 2022. *Clipscore: A reference-free evaluation metric for image captioning*. *Preprint*, arXiv:2104.08718.
- Geoffrey Hinton, Oriol Vinyals, and Jeff Dean. 2015. *Distilling the knowledge in a neural network*. *Preprint*, arXiv:1503.02531.
- Junnan Li, Dongxu Li, Silvio Savarese, and Steven Hoi. 2023. *BLIP-2: Bootstrapping Language-Image Pre-training With Frozen Image Encoders and Large Language Models*. *arXiv preprint*.
- Haotian Liu, Chunyuan Li, Yuheng Li, and Yong Jae Lee. 2024. *Improved baselines with visual instruction tuning*. *Preprint*, arXiv:2310.03744.
- Yang Liu, Dan Iter, Yichong Xu, Shuohang Wang, Ruochen Xu, and Chenguang Zhu. 2023. *G-eval: NLG evaluation using gpt-4 with better human alignment*. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 2511–2522, Singapore. Association for Computational Linguistics.
- Pan Lu, Swaroop Mishra, Tony Xia, Liang Qiu, Kai-Wei Chang, Song-Chun Zhu, Oyvind Taffjord, Peter Clark, and Ashwin Kalyan. 2022. Learn to explain: Multimodal reasoning via thought chains for science question answering. In *The 36th Conference on Neural Information Processing Systems (NeurIPS)*.
- Ahmed Masry, Do Long, Jia Qing Tan, Shafiq Joty, and Enamul Hoque. 2022. *ChartQA: A benchmark for question answering about charts with visual and logical reasoning*. In *Findings of the Association for Computational Linguistics: ACL 2022*, pages 2263–2279, Dublin, Ireland. Association for Computational Linguistics.
- Ron Mokady, Amir Hertz, and Amit H Bermano. 2021. Clipcap: Clip prefix for image captioning. *arXiv preprint arXiv:2111.09734*.
- Kishore Papineni, Salim Roukos, Todd Ward, and Wei-Jing Zhu. 2002. *Bleu: a method for automatic evaluation of machine translation*. In *Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics*, pages 311–318, Philadelphia, Pennsylvania, USA. Association for Computational Linguistics.
- Rafael Rafailov, Archit Sharma, Eric Mitchell, Stefano Ermon, Christopher D. Manning, and Chelsea Finn. 2024. *Direct preference optimization: Your language model is secretly a reward model*. *Preprint*, arXiv:2305.18290.
- Amir Rasouli, Iuliia Kotseruba, and John K Tsotsos. 2017. Are they going to cross? a benchmark dataset and baseline for pedestrian crosswalk behavior. In *ICCVW*, pages 206–213.
- Amir Rasouli, Iuliia Kotseruba, and John K Tsotsos. 2018. It’s not all about size: On the role of data properties in pedestrian detection. In *ECCVW*.
- Ramakrishna Vedantam, C. Lawrence Zitnick, and Devi Parikh. 2015. *Cider: Consensus-based image description evaluation*. *Preprint*, arXiv:1411.5726.
- Peng Wang, Shuai Bai, Sinan Tan, Shijie Wang, Zhihao Fan, Jinze Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, Yang Fan, Kai Dang, Mengfei Du, Xuancheng Ren, Rui Men, Dayiheng Liu, Chang Zhou, Jingren Zhou, and Junyang Lin. 2024. Qwen2-vl: Enhancing vision-language model’s perception of the world at any resolution. *arXiv preprint arXiv:2409.12191*.
- Kelvin Xu, Jimmy Ba, Ryan Kiros, Kyunghyun Cho, Aaron Courville, Ruslan Salakhutdinov, Richard Zemel, and Yoshua Bengio. 2016. *Show, attend and tell: Neural image caption generation with visual attention*. *Preprint*, arXiv:1502.03044.

Tianyi Zhang*, Varsha Kishore*, Felix Wu*, Kilian Q.
Weinberger, and Yoav Artzi. 2020. [Bertscore: Evaluating text generation with bert](#). In *International Conference on Learning Representations*.

A Category-Specific Prompts

This appendix provides the complete set of category-specific prompts used in our research. Each prompt follows the structured design described in Section 3.1. Below we present the prompts for all 18 categories, illustrating our comprehensive approach to domain-specific image captioning.

A.1 Animal Category: Biological Entity Focus

System Prompt:

Generate a detailed and accurate caption for the provided animal image. The caption should include the following elements:

1. Accuracy:

- Ensure the description is precise and free of errors.

2. Animal Species:

- Mention the species of the animal if identifiable

3. Physical Characteristics:

- Describe the physical features of the animal (e.g., color, size, distinguishing marks)

4. Behavior:

- Describe the behavior or activity of the animal in the image

5. Habitat:

- Mention the environment or setting where the animal is located (e.g., forest, desert, ocean)

6. Contextual Details:

- Provide any additional relevant information (e.g., weather conditions, time of day)

User Text: "You are given an image of an animal. Please generate a detailed and accurate caption for the image without any additional information."

A.2 Architecture Category: Built Environment Focus

System Prompt:

Generate a detailed and accurate caption for the provided architecture image. The

caption should include the following elements:

1. Accuracy:

- Ensure the description is precise and free of errors.

2. Building Type:

- Mention the type of building (e.g., residential, commercial, historical)

3. Architectural Style:

- Describe the architectural style (e.g., modern, Gothic, Baroque)

4. Physical Characteristics:

- Describe the physical features of the building (e.g., height, materials, distinctive elements)

5. Surroundings:

- Mention the building's surroundings (e.g., urban area, countryside)

6. Contextual Details:

- Provide any additional relevant information (e.g., weather conditions, time of day)

User Text: "You are given an image of an architectural structure. Please generate a detailed and accurate caption for the image without any additional information."

A.3 Art Category: Creative Work Focus

System Prompt:

Generate a detailed and accurate caption for the provided art image. The caption should include the following elements:

1. Accuracy:

- Ensure the description is precise and free of errors.

2. Art Form:

- Mention the form of art (e.g., painting, sculpture, digital art)

3. Style:

- Describe the style of the artwork (e.g., abstract, realism, impressionism)

4. Elements and Composition:

861	- Describe the elements and composition		
862	of the artwork (e.g., colors, shapes, sub-		
863	jects)		
864	5. Context:		
865	- Mention any relevant context (e.g.,		
866	artist, period, cultural significance)		
867	6. Emotional and Interpretative Ele-		
868	ments:		
869	- Provide interpretation or emotional ele-		
870	ments conveyed by the artwork		
871	User Text: "You are given an image of		
872	an artwork. Please generate a detailed		
873	and accurate caption for the image with-		
874	out any additional information."		
875	A.4 Autonomous Driving Category: Vehicle		
876	Environment Focus		
877	System Prompt:		
878	Generate a detailed and accurate caption		
879	for the provided autonomous driving im-		
880	age. The caption should include the fol-		
881	lowing elements:		
882	1. Accuracy:		
883	- Ensure the description is precise and		
884	free of errors.		
885	2. Scene Overview:		
886	- Briefly describe the overall scene (e.g.,		
887	urban street, highway, suburban area)		
888	- Mention weather conditions and time of		
889	day (e.g., daytime, night, rainy)		
890	3. Road Environment:		
891	- Describe road type (e.g., two-way four-		
892	lane, one-way street)		
893	- Mention road surface conditions (e.g.,		
894	dry, wet, snowy)		
895	- Describe road markings and traffic signs		
896	4. Traffic Conditions:		
897	- Describe positions and behaviors of sur-		
898	rounding vehicles		
899	- Mention pedestrians, cyclists, and other		
900	road users		
901	- Describe traffic light status (if visible)		
902	5. Ego Vehicle Status:		
903	- Describe the autonomous vehicle's po-		
904	sition and driving status		
	- Mention speed and relative position to		905
	other vehicles		906
	6. Potential Hazards:		907
	- Point out any potential dangers or situa-		908
	tions requiring special attention		909
	7. Autonomous Driving Relevant Ele-		910
	ments:		911
	- Describe elements particularly impor-		912
	tant for the autonomous driving system		913
	- Mention situations that may require spe-		914
	cial handling		915
	User Text: "You are given an image re-		916
	lated to autonomous driving. Please gen-		917
	erate a detailed and accurate caption for		918
	the image without any additional infor-		919
	mation."		920
	A.5 Cartoon Category: Animated Content		921
	Focus		922
	System Prompt:		923
	Generate a detailed and accurate caption		924
	for the provided cartoon image. The		925
	caption should include the following ele-		926
	ments:		927
	1. Accuracy:		928
	- Ensure the description is precise and		929
	free of errors.		930
	2. Character Description:		931
	- Mention the main characters and their		932
	characteristics		933
	3. Content Description:		934
	- Describe the content of the image		935
	4. Scene Description:		936
	- Describe the scene or setting (e.g., loca-		937
	tion, time of day)		938
	5. Actions and Interactions:		939
	- Describe the actions and interactions of		940
	the characters		941
	6. Visual Style:		942
	- Mention the visual style of the cartoon		943
	(e.g., color palette, drawing style)		944
	7. Contextual Details:		945
	- Provide any additional relevant informa-		946
	tion (e.g., mood, tone)		947

948	User Text: "You are given an image of a	- Mention notable landmarks or buildings	992
949	cartoon. Please generate a detailed and		
950	accurate caption for the image without	4. Activity and Movement:	993
951	any additional information."	- Describe any visible activity or move-	994
		ment (e.g., traffic, pedestrians)	995
952	A.6 Chart Category: Data Visualization	5. Weather and Time of Day:	996
953	Focus	- Mention weather conditions and time of	997
954	System Prompt:	day	998
955	Generate a detailed and accurate cap-	6. Contextual Details:	999
956	tion for the provided chart image. The	- Provide any additional relevant infor-	1000
957	caption should include the following ele-	mation (e.g., cultural or historical signifi-	1001
958	ments:	cance)	1002
959	1. Accuracy:	User Text: "You are given an image of	1003
960	- Ensure the description is precise and	a city. Please generate a detailed and	1004
961	free of errors.	accurate caption for the image without	1005
962	2. Details:	any additional information."	1006
963	- Include relevant specifics such as the	A.8 Daily Life Category: Everyday Scene	1007
964	type of chart, data points, trends, and key	Focus	1008
965	observations.	System Prompt:	1009
966	3. Consistency:	Generate a detailed and accurate caption	1010
967	- Maintain a consistent style and format.	for the provided daily life image. The	1011
968	4. Readability:	caption should include the following ele-	1012
969	- Write in clear, concise language, avoid-	ments:	1013
970	ing unnecessary complexity.	1. Accuracy:	1014
971	5. Insights or Recommendations:	- Ensure the description is precise and	1015
972	- If appropriate, provide simple insights	free of errors.	1016
973	or recommendations based on the chart	2. Scene Overview:	1017
974	data.	- Briefly describe the overall scene (e.g.,	1018
975	User Text: "You are given an image of a	home, office, street)	1019
976	chart image. Please generate a detailed	3. Activities:	1020
977	and accurate caption for the image with-	- Describe the activities of people in the	1021
978	out any additional information."	scene	1022
979	A.7 City Category: Urban Environment	4. Objects and Environment:	1023
980	Focus	- Mention notable objects and environ-	1024
981	System Prompt:	mental details	1025
982	Generate a detailed and accurate caption	5. Emotional and Social Context:	1026
983	for the provided city image. The caption	- Describe the emotional and social con-	1027
984	should include the following elements:	text of the scene	1028
985	1. Accuracy:	6. Contextual Details:	1029
986	- Ensure the description is precise and	- Provide any additional relevant infor-	1030
987	free of errors.	mation (e.g., time of day, weather)	1031
988	2. Cityscape Overview:	User Text: "You are given an image de-	1032
989	- Briefly describe the overall cityscape	picturing daily life. Please generate a de-	1033
990	(e.g., skyline, street view)	tailed and accurate caption for the image	1034
991	3. Landmarks and Buildings:	without any additional information."	1035

1122	1. Accuracy:	User Text: "You are given an image of	1166
1123	- Ensure the description is precise and	a person. Please generate a detailed and	1167
1124	free of errors.	accurate caption for the image without	1168
1125	2. Scene Overview:	any additional information."	1169
1126	- Briefly describe the overall scene	A.14 Plant Category: Botanical Focus	1170
1127	3. Medical Procedures:	System Prompt:	1171
1128	- Describe any visible medical proce-	Generate a detailed and accurate cap-	1172
1129	dures or activities	tion for the provided plant image. The	1173
1130	4. Medical Equipment:	caption should include the following ele-	1174
1131	- Mention notable medical equipment and	ments:	1175
1132	instruments	1. Accuracy:	1176
1133	5. Contextual Details:	- Ensure the description is precise and	1177
1134	- Provide any additional relevant informa-	free of errors.	1178
1135	tion (e.g., condition being treated, emo-	2. Plant Species:	1179
1136	tional atmosphere)	- Mention the species of the plant if iden-	1180
1137	User Text: "You are given an image de-	tifiable	1181
1138	picting a medical scene. Please generate	3. Physical Characteristics:	1182
1139	a detailed and accurate caption for the	- Describe the physical features of the	1183
1140	image without any additional informa-	plant (e.g., color, size, distinctive ele-	1184
1141	tion."	ments)	1185
1142	A.13 Person Category: Human Subject Focus	4. Environment:	1186
1143	System Prompt:	- Mention the environment or setting	1187
1144	Generate a detailed and accurate cap-	where the plant is located (e.g., garden,	1188
1145	tion for the provided person image. The	forest)	1189
1146	caption should include the following ele-	5. Contextual Details:	1190
1147	ments:	- Provide any additional relevant informa-	1191
1148	1. Accuracy:	tion (e.g., season, time of day)	1192
1149	- Ensure the description is precise and	User Text: "You are given an image of	1193
1150	free of errors.	a plant. Please generate a detailed and	1194
1151	2. Physical Appearance:	accurate caption for the image without	1195
1152	- Describe the person's physical appear-	any additional information."	1196
1153	ance (e.g., age, gender, clothing)	A.15 Smart Cities Category: Urban	1197
1154	3. Actions and Activities:	Technology Focus	1198
1155	- Mention any actions or activities the	System Prompt: Generate a detailed and	1199
1156	person is engaged in	accurate caption for the provided smart	1200
1157	4. Emotional State:	city image. The caption should include	1201
1158	- Describe the person's emotional state or	the following elements: 1. Accuracy:	1202
1159	expression	- Ensure the description is precise and	1203
1160	5. Context and Setting:	free of errors. 2. Cityscape Overview:	1204
1161	- Provide context about the setting or	- Briefly describe the overall cityscape	1205
1162	background	(e.g., skyline, street view) 3. Smart In-	1206
1163	6. Contextual Details:	frastructure: - Mention notable smart	1207
1164	- Provide any additional relevant informa-	infrastructure elements (e.g., smart build-	1208
1165	tion (e.g., weather, time of day)	ings, IoT devices) 4. Technological Ele-	1209
		ments: - Describe visible technological	1210

1211	elements (e.g., sensors, automated systems) 5. Traffic and Transportation:	User Text: "You are given an image of a transportation scene. Please generate a detailed and accurate caption for the image without any additional information."	1259
1212	- Mention traffic and transportation systems (e.g., autonomous vehicles, smart traffic lights) 6. Contextual Details:		1260
1213	- Provide any additional relevant information (e.g., time of day, weather)		1261
1214			1262
1215			
1216			
1217			
1218	User Text: "You are given an image of a smart city. Please generate a detailed and accurate caption for the image without any additional information."		
1219			
1220			
1221			
1222	A.16 Sports Category: Athletic Activity Focus	A.18 Other Category: General Content Focus	1263
1223	System Prompt: Generate a detailed and accurate caption for the provided sports image. The caption should include the following elements: 1. Accuracy: - Ensure the description is precise and free of errors. 2. Sport Type: - Mention the type of sport being played 3. Players and Actions: - Describe the players and their actions 4. Equipment: - Mention any visible sports equipment 5. Setting: - Describe the setting (e.g., stadium, field) 6. Contextual Details: - Provide any additional relevant information (e.g., score, spectators, weather)	System Prompt: Generate a detailed and accurate caption for the provided image. The caption should include the following elements: 1. Accuracy: - Ensure the description is precise and free of errors. 2. Content Description: - Provide a clear and relevant description of the content in the image. 3. In-depth Understanding: - Go beyond surface-level details by interpreting the underlying meaning or significance of the scene, including any implied relationships, emotions, or actions. 4. Contextual Details: - Include any necessary contextual details (e.g., cultural, historical, social significance) to enhance the understanding of the image. 5. Visual and Environmental Elements: - Describe notable visual and environmental elements present in the image, paying attention to subtle features that may influence the overall interpretation.	1264
1224			1265
1225			1266
1226			1267
1227			1268
1228			1269
1229			1270
1230			1271
1231			1272
1232			1273
1233			1274
1234			1275
1235			1276
1236			1277
1237	User Text: "You are given an image of a sports scene. Please generate a detailed and accurate caption for the image without any additional information."		1278
1238			1279
1239			1280
1240			1281
1241	A.17 Transportation Category: Mobility Focus	User Text: "You are given an image. Please generate a detailed and accurate caption for the image without any additional information, focusing on both surface details and deeper interpretations."	1286
1242			1287
1243	System Prompt: Generate a detailed and accurate caption for the provided transportation image. The caption should include the following elements: 1. Accuracy: - Ensure the description is precise and free of errors. 2. Vehicle Type: - Mention the type of vehicle(s) shown 3. Scene Overview: - Briefly describe the overall scene (e.g., road, railway, airport) 4. Activity and Movement: - Describe any visible activity or movement of the vehicles 5. Environment: - Mention the environment or setting (e.g., urban, rural) 6. Contextual Details: - Provide any additional relevant information (e.g., weather conditions, time of day)		1288
1244			1289
1245			1290
1246			
1247			
1248			
1249			
1250			
1251			
1252			
1253			
1254			
1255			
1256			
1257			
1258			
		We plan to make these prompts and associated implementations available to facilitate future research in domain-specific image captioning.	1291
			1292
			1293
		B Training Details	1294
		B.1 Model Configurations	1295
		Classification Model	1296
		• Base model: LLaVA-v1.5-7B	1297
		• Training strategy: Full parameter fine-tuning	1298
		• Learning rate: 2e-5	1299
		• Batch size: 16 per device	1300
		• Number of epochs: 3	1301

1302	Caption Model (Full Parameter)	Based on the criteria provided, does the caption meet the requirements?	1338
1303	• Base model: LLaVA-v1.5-7B	Please rate the caption on a scale of 1 to 5, where:	1339
1304	• Training data: GPT-4o generated captions		1340
1305	• Learning rate: 2e-5	• 1 - Poor: Does not meet the criteria at all.	1341
1306	• Batch size: 16 per device	• 2 - Fair: Meets a few criteria but has significant issues.	1342
1307	• Training epochs: 3	• 3 - Good: Meets most criteria with minor issues.	1343
1308	Caption Model (LoRA)	• 4 - Very Good: Meets all criteria with very few issues.	1344
1309	• LoRA rank: 128	• 5 - Excellent: Perfectly meets all criteria.	1345
1310	• LoRA alpha: 256		1346
1311	• LoRA dropout: 0	Be as objective as possible. After providing your explanation, you must rate the response on a scale of 1 to 5 by strictly following this format: "Rating: [[rating]]", for example: "Rating: [[5]]".	1347
1312	• Learning rate: 2e-4		1348
1313	• Target modules: all	Here is the image and the corresponding caption.	1349
1314	DPO Training	Caption: {caption to be evaluated}	1350
1315	• LoRA rank: 8		1351
1316	• LoRA alpha: 16		1352
1317	• Learning rate: 2e-6		1353
1318	• Batch size: 2 per device		1354
1319	• Gradient accumulation steps: 8		1355
1320	B.2 Computing Infrastructure		1356
1321	• Hardware: 8 NVIDIA A100 GPUs (80GB)		1357
1322	• Framework: LlamaFactory with DeepSpeed		1358
1323	• Mixed precision: BF16		1359
1324	For detailed training scripts and configurations, we will make these publicly available in our upcoming repository to facilitate future research in this area.		1360
1325			1361
1326			1362
1327			1363
1328	C Pixtral Evaluation Implementation		
1329	The Pixtral evaluation system uses a standardized prompt structure for evaluating captions:		
1330			
1331	You are a helpful and precise assistant for evaluating the quality of captions. A caption will be provided for a particular image. Below are the criteria for evaluation:		
1332			
1333			
1334			
1335			
1336	<Criteria Begin> {category-specific criteria} <Criteria End>		
1337			