

Making LVLMs Look Twice: Contrastive Decoding with Contrast Images

Anonymous ACL submission

Abstract

Large Vision-Language Models (LVLMs) are becoming increasingly popular for text-vision tasks requiring reasoning over both modalities, but often struggle with fine-grained visual discrimination. This limitation is evident in recent benchmarks like NaturalBench and D3, where closed models such as GPT-4o achieve only 39.6% accuracy, and open-source models perform below random chance (25%). We introduce Contrastive decoding with Contrast Images (CoCI), which adjusts LVLM outputs by contrasting them against outputs for similar images (Contrast Images - CIs). We first evaluate CoCI using naturally occurring CIs in benchmarks with curated image pairs, achieving improvements of up to 98.9% on NaturalBench, 69.5% on D3, and 37.6% on MMVP. For real-world applications where natural CIs are unavailable, we show that given sufficient training data, a lightweight neural classifier can effectively select CIs from similar images at inference time, improving NaturalBench performance by up to 36.8%. For scenarios lacking training data, we develop a caption-matching technique that selects CIs by comparing LVLM-generated descriptions of candidate images. Our method demonstrates the potential for improving LVLMs at inference time through different CI selection approaches, each suited to different data availability scenarios.

1 introduction

Large Vision-Language Models (LVLMs) are becoming increasingly popular for text-vision tasks that require reasoning over both modalities. However, they often struggle with fine-grained visual discrimination — that is, the ability to tell two similar yet distinct images apart — a crucial capability for real-world applications such as multimodal search, manufacturing, and robotics. Recent benchmarks have exposed this limitation: on NaturalBench (Li et al., 2024a), which tests vi-

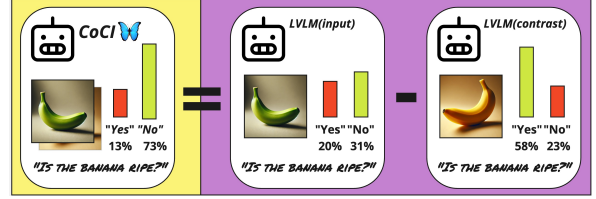


Figure 1: CoCI penalizes target image logits using those from a contrast image, weighted by hyperparameter α .

sual question answering over closely related images, state-of-the-art closed models like GPT-4o (OpenAI et al., 2024) achieve only 39.6% accuracy. Similarly, on the D3 benchmark (Gaur et al., 2024), which requires describing differences between paired images, open-source models perform below random chance (25%).

Efforts to address fine-grained visual discrimination in LVLMs are still under-explored. Current strategies addressing other LVLM shortcomings often rely on fine-tuning with specialized datasets (Wang et al., 2023; Chen et al., 2023; Liu et al., 2024a; Sarkar et al., 2024), multi-step correction pipelines (Yin et al., 2023; Zhou et al., 2023), or inference-time methods (Leng et al., 2023; Manevich and Tsarfaty, 2024; Liu et al., 2024b; Huang et al., 2023). Inference-time methods are particularly appealing as they do not require expensive model training and are less prone to compounding errors that can affect multi-step systems.

Building on the advantages of inference-time methods, we propose Contrastive decoding with Contrast Images (CoCI), an approach specifically designed to improve fine-grained visual discrimination in LVLMs. CoCI penalizes LVLM next-token probabilities with those obtained by feeding a different, contrasting image input (See Figure 1).

We evaluate CoCI across three different supervision regimes. First, using naturally occurring Contrast Images in curated benchmarks like NaturalBench, D3 and MMVP, we demonstrate improvements up to 98.9%, 69.5%, 37.6% respec-

tively. This establishes a performance ceiling for CoCI when ideal CIs are available. For applications where natural CIs are unavailable but training data exists, we show that a lightweight classifier can effectively select CIs from visually similar images at inference time, improving NaturalBench performance by up to 36.5%. In settings without training data, we propose a caption-matching technique that selects CIs at inference time by comparing LVLM-generated descriptions of candidate images.

Experiments with leading LVLMs — Qwen2-VL, LLaVA-OneVision, and Llama 3.2 (Wang et al., 2024a; Li et al., 2024b; Grattafiori et al., 2024) — establish the potential of contrastive decoding strategies with contrastive images for improved multimodal reasoning in real-world tasks.

2 Contrastive Decoding with Contrast Images (CoCI)

We present CoCI, a method to improve LVLM outputs by penalizing token probabilities that are likely under a contrast image. The choice of contrast image is crucial: e.g., when querying about fruit ripeness with an input image of an unripe banana, contrasting against an image of a ripe banana provides strong contrastive signal, while an image of a ripe pear offers weaker contrast and an image of a bus provides no useful signal and may degrade performance. This intuition guides our CI selection strategies across different scenarios. Before formalizing this intuition, we first review key concepts in LVLM text generation.

2.1 Preliminaries: Text Generation in LVLMs

LVLMs extend LLMs’ next-token prediction capability by conditioning on both text and images.¹ Generation (also called decoding) proceeds by sampling from next-token distributions, concatenating selected tokens with the context, and feeding this back into the model until an EOS token or length limit is reached. The LVLM next-token prediction formula is:

$$\text{LVLM}_t(y_{<t}, I) = P(y|y_{<t}, I) \quad \forall y \in \mathcal{V} \quad (1)$$

where $y_{<t}$ is the token prefix including both input and generated text up to position t , I is the input image, and $\mathcal{V}$ is the model’s vocabulary.

¹In this work, we focus on the single image input case.

2.2 Contrastive Decoding

Since Li et al. (2023) introduced Contrastive Decoding, multiple variants have been explored (Sennrich et al., 2024; Jin et al., 2024; Phan et al., 2024), varying in weighting mechanisms, truncation strategies, and probability space operations. We implement CoCI based on Sennrich et al. (2024)’s minimal-hyperparameter variant as follows:

$$\begin{aligned} \text{CoCI}_t(y_{<t}, I, I') = \\ \log \left(P(y|y_{<t}, I) - \alpha P(y|y_{<t}, I') \right) \quad \forall y \in \mathcal{V} \end{aligned} \quad (2)$$

At each timestep t , CoCI penalizes token probabilities from distribution $P(y|y_{<t}, I)$ with those from $P(y|y_{<t}, I')$, where I is the target image and I' is the contrast image. The hyperparameter α controls the contrastive penalty strength.²

2.3 Obtaining Contrast Images

We propose three approaches for obtaining CIs:

Naturally occurring CIs. Many tasks naturally provide pairs of images that can serve as contrast images (CIs). For instance, a home assistant robot searching for “the blue ceramic mug with a chip on the handle” needs to distinguish between similar cups to find the exact match. We evaluate this scenario using LVLM benchmarks with curated image pairs designed to test fine-grained discrimination capabilities. These paired images serve as natural CIs in our experiments.

Classifier-obtained CIs. For cases without natural CIs, we train a lightweight classifier to select them at inference time. Given an LVLM L and training triplets $\langle q, I, I' \rangle$ where q is a binary question and I, I' are images with different ground-truth answers, we (a) Extract LVLM hidden states $h_{q,i} \in R^{d_L}$ for each image-question pair. (b) Concatenate states for image pairs to get $h_{q,i,i'} \in R^{2*d_L}$. (c) Create negative samples using the j least similar images from the top- k similar images to I in dataset D ³. (d) Train a three-layer MLP as a CI classifier⁴

²Throughout this work, we set $\alpha = 0.5$ for VQA tasks, and $\alpha = 0.8$ for open-ended generation tasks, without tuning.

³We set $j = 5$, $k = 100$ without tuning. For D , we use flickr30k (Young et al., 2014). We use open-clip with checkpoint “laion/CLIP-ViT-L-14-DataComp.XL-s13B-b90K” as the image encoder (Ilharco et al., 2021; Cherti et al., 2023; Radford et al., 2021a; Schuhmann et al., 2022). We use cosine-similarity as the similarity score throughout this work.

⁴See appendix A.1 and A.3 for PyTorch code, training setup and ablations.

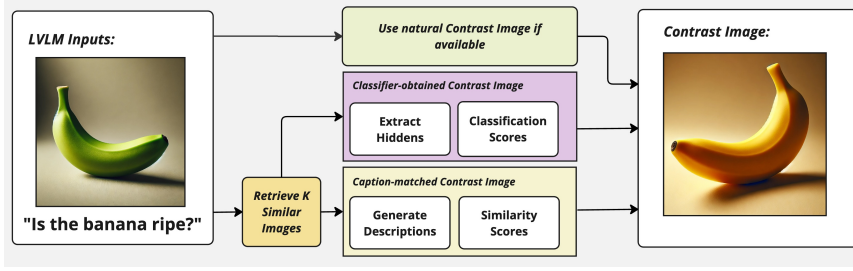


Figure 2: Illustration of the approaches we explore for obtaining a Contrast Image (CI).

For training, we use 60% of NaturalBench data, augmented with GPT-4o-generated question paraphrases. At inference, we select the CI that maximizes the classifier score among the k most similar images to the input.⁵

Caption-matched CIs. For scenarios without training data, we select CIs by comparing LVLM-generated image descriptions. Given an input image, we (a) Retrieve k similar images⁶. (b) Generate LVLM descriptions for all $k + 1$ images. (c) Embed descriptions using a text encoder. (d) Select the image whose description is most similar to the input image’s description

2.4 Research Hypothesis

We test whether: (a) Contrastive decoding with CIs improves LVLM fine-grained reasoning, (b) A lightweight classifier trained on LVLM hidden states can effectively select CIs, and (c) Images with similar LVLM descriptions can serve as CIs.

3 Experiments

We evaluate CoCI with three leading LVLMs⁷, on three benchmarks testing fine-grained visual discrimination:

NaturalBench (Li et al., 2024a) tests discrimination between similar images through yes/no and multiple-choice questions, where answers differ between paired images. The benchmark contains 1900 image pairs with two questions per pair. We split the data into train (60%), dev (20%), and test (20%) sets. Metrics include image accuracy (both questions correct per image), question accuracy (per-question correctness), and group accuracy (G-Acc, requiring correct answers for all four image-question combinations).

⁵See table 2 for a comparison of using different k values at inference time. Our inference-time retrieval setup is identical to the one used in training (i.e. dataset, embedding model).

⁶We set $k = 5$ without tuning.

⁷See appendix A.2 for details on the checkpoints we used.

Model	Method	D3 (self-ret.)	MMVP (acc.)	NB (g-acc.)
Qwen2-VL	Baseline	30.8	46.0	30.8
	CoCI _{CAP}	34.8	48.7	31.3
	CoCI _{NAT}	52.2	63.3	46.6
LLaVA-OV	Baseline	25.1	52.7	28.2
	CoCI _{CAP}	31.6	57.3	31.6
	CoCI _{NAT}	38.1	66.7	56.1
Llama 3.2	Baseline	28.7	39.3	21.1
	CoCI _{CAP}	33.6	41.3	22.4
	CoCI _{NAT}	35.6	43.3	29.2

Table 1: CoCI performance comparison with provided CIs across three benchmarks, with natural CIs (CoCI_{NAT}) and caption-matched CIs (CoCI_{CAP}). Random baseline score for all metrics is 25.

MMVP (Multimodal Visual Patterns) (Tong et al., 2024) evaluates visual difference detection through multiple-choice questions on image pairs. Each pair differs in at least one visual aspect (e.g., object state, position, or orientation), with questions targeting these differences. The dataset comprises 150 image pairs with one question per pair. A model is considered correct only if it answers correctly on both images in a pair.

D3 (Detect, Describe, Discriminate) (Gaur et al., 2024) evaluates LVLMs’ ability to generate discriminative descriptions that uniquely distinguish between similar images and consists of 247 image pairs. We adapt D3 to evaluate CoCI by reformulating it as a single-input task, where the model describes each image separately. We use the original self-retrieval evaluation protocol, which tests whether an image-text encoder can match the descriptions back to their respective images.

4 Results and Discussion

In Table 1 we can see that using natural CIs yields substantial improvements: up to 21.4 points on D3 (Qwen), 17.3 points on MMVP (LLaVA), and 27.9 points on NaturalBench (LLaVA). Caption-matched CIs show moderate but consistent gains,

Model	Method	Q-acc	I-acc	Acc	G-acc
Qwen2-VL	Baseline	55.3	59.3	76.8	30.8
	$Cls_{k=4}$	55.5	58.8	76.4	32.1
	$Cls_{k=8}$	56.3	58.9	76.7	32.4
	$Cls_{k=16}$	57.4	60.1	77.2	33.7
	$Cls_{k=32}$	57.8	60.1	77.4	34.2
	$Cls_{k=64}$	58.2	60.8	77.9	33.9
LLaVA-OV	Baseline	53.8	56.1	74.6	28.2
	$Cls_{k=4}$	59.2	59.6	77.6	35.3
	$Cls_{k=8}$	57.8	60.1	77.5	34.5
	$Cls_{k=16}$	57.6	58.7	77.0	33.4
	$Cls_{k=32}$	60.3	62.1	78.5	38.4
	$Cls_{k=64}$	59.7	62.1	78.2	37.6
Llama 3.2	Baseline	46.3	50.5	71.8	21.1
	$Cls_{k=4}$	49.2	52.8	73.2	23.2
	$Cls_{k=8}$	49.1	52.2	73.1	21.8
	$Cls_{k=16}$	48.8	52.4	73.1	22.4
	$Cls_{k=32}$	49.9	52.5	73.7	22.1
	$Cls_{k=64}$	49.7	52.5	73.6	22.1

Table 2: CoCI accuracy metrics on the NaturalBench test set with CIs chosen using a lightweight classifier. $k = j$ denotes the classifier ran on the j most similar images to the input image.

particularly on D3 where LLaVA improves from 25.1% to 31.6%, suggesting that contrasting against images with similar captions effectively guides visual discrimination.

Table 2 shows that for Qwen and LLaVA, performance improves with larger candidate pools (k), peaking around $k=32$. Llama performs best with small pools ($k=4$), perhaps due to differences in its architecture affecting classifier effectiveness on its hidden states.

In NaturalBench, G-Acc shows particularly strong improvement with natural CIs (e.g., from 28.2% to 56.1% for LLaVA-OV) as it requires consistency across all image-question combinations. This pattern persists with classifier-selected CIs, where G-Acc improves by up to 10.2 points while other metrics show more modest gains.

The substantial performance gap between natural CIs and other methods indicates that while classifier-selected and caption-matched CIs provide improvements, they don’t yet capture all aspects that make natural pairs effective.⁸

5 Related Work

Inference-time methods for enhancing multimodal reasoning. Recent work has focused on hallucination reduction through confidence-based probability adjustments (Huo et al., 2024), compar-

ison with semantic references (Yang et al., 2024), and contrastive decoding with perturbed inputs (Leng et al., 2023; Manevich and Tsarfaty, 2024). Our work extends these approaches to address fine-grained visual discrimination.

Alignment and grounding in LVLMS. Prior work has enhanced visual-textual alignment through object-level information synthesis (Wang et al., 2024b), targeted fine-tuning (Lu et al., 2024), and specialized dataset construction (Li et al., 2024c). While these methods improve foundational capabilities, they don’t directly address fine-grained discrimination.

Contrastive examples in multimodal models. CLIP (Radford et al., 2021b) established contrastive learning between images and text for modality alignment. Recent works focus on leveraging contrast pairs: (Le et al., 2023) and (Zhang et al., 2024) generate synthetic contrast datasets using text-to-image models, while (Abbasnejad et al., 2020) and (Zhou et al., 2024) use contrastive examples to address dataset biases. Unlike these approaches that require data generation or model training, our method operates at inference time.⁹

6 Conclusion

We introduced Contrastive decoding with Contrast Images (CoCI), showing its effectiveness in improving LVLMS’ fine-grained visual discrimination capabilities on multiple benchmarks. While naturally occurring contrast pairs yielded the strongest performance gains, both classifier-based and caption-matching approaches can provide meaningful improvements without requiring dataset curation, model training or extensive computational resources.

The effectiveness of our approach across different supervision regimes establishes CoCI as a practical solution for improving LVLMS performance at inference time. Our results suggest that contrastive decoding strategies, when combined with appropriate contrast image selection, can enhance LVLMS’ ability to make fine-grained visual distinctions, opening new avenues for improving multimodal reasoning through inference-time techniques.

⁸See appendix A.3 for ablation tests with different CI selection strategies.

⁹Classifier-selected CIs require minimal preprocessing compared to model finetuning or dataset curation.

7 Limitations

CoCI has several limitations worth noting. While we demonstrate its effectiveness with classifier-based and caption-matching approaches, the substantial performance gap between natural and automatically selected CIs indicates significant headroom for finding more effective contrast images. We tested simple selection methods to establish the viability of the approach, leaving the exploration of more sophisticated CI selection strategies to future work. Additionally, our evaluation focuses primarily on VQA and self-retrieval protocols; exploring additional evaluation methods could reveal other aspects of how CoCI affects LVLM generations.

The method introduces additional computation at inference time, running the LVLM twice per generation step and requiring CI selection overhead. While this aligns with the growing trend of leveraging test-time compute for improved performance, the current implementation could be optimized. Future work could explore more efficient implementations of contrastive decoding and investigate fusing operations like hidden state extraction with the generation procedure to reduce computational overhead.

Our implementation uses Flickr30k as the image database for CI selection - using larger, more diverse image collections could improve performance. Alternative image retrieval models and similarity scoring methods could also enhance CI selection. Additionally, our approach does not address cases where multiple contrasts might be informative - we only use a single contrast image, while some scenarios might benefit from multiple contrasting viewpoints.

The experiments use a fixed contrastive weight (α) across tasks within each category (VQA/generation). A more nuanced approach to setting this parameter, dynamically per sample or per token, based on the specific input or task, could yield better results.

While CoCI improves visual discrimination, it could potentially amplify biases present in contrast image databases or introduce new failure modes when inappropriate contrast images are selected. These risks should be carefully evaluated before deployment in sensitive applications.

Finally, our experiments focus exclusively on English-language benchmarks. Extending CoCI to multilingual settings and investigating how contrastive decoding approaches perform across differ-

ent languages represents an important direction for future research.

References

- Ehsan Abbasnejad, Damien Teney, Amin Parvaneh, Javen Shi, and Anton van den Hengel. 2020. [Counterfactual vision and language learning](#). In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 10041–10051.
- Zhiyang Chen, Yousong Zhu, Yufei Zhan, Zhaowen Li, Chaoyang Zhao, Jinqiao Wang, and Ming Tang. 2023. [Mitigating hallucination in visual language models with visual supervision](#).
- Mehdi Cherti, Romain Beaumont, Ross Wightman, Mitchell Wortsman, Gabriel Ilharco, Cade Gordon, Christoph Schuhmann, Ludwig Schmidt, and Jenia Jitsev. 2023. Reproducible scaling laws for contrastive language-image learning. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 2818–2829.
- Manu Gaur, Darshan Singh S, and Makarand Tapaswi. 2024. [Detect, describe, discriminate: Moving beyond vqa for mllm evaluation](#).
- Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, Amy Yang, Angela Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang, Archi Mitra, Archie Sravankumar, Artem Korenev, Arthur Hinsvark, Arun Rao, Aston Zhang, Aurelien Rodriguez, Austen Gregerson, Ava Spataru, Baptiste Roziere, Bethany Biron, Binh Tang, Bobbie Chern, Charlotte Caucheteux, Chaya Nayak, Chloe Bi, Chris Marra, Chris McConnell, Christian Keller, Christophe Touret, Chunyang Wu, Corinne Wong, Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Allonsius, Daniel Song, Danielle Pintz, Danny Livshits, Danny Wyatt, David Esiobu, Dhruv Choudhary, Dhruv Mahajan, Diego Garcia-Olano, Diego Perino, Dieuwke Hupkes, Egor Lakomkin, Ehab AlBadawy, Elina Lobanova, Emily Dinan, Eric Michael Smith, Filip Radenovic, Francisco Guzmán, Frank Zhang, Gabriel Synnaeve, Gabrielle Lee, Georgia Lewis Anderson, Govind Thattai, Graeme Nail, Gregoire Mialon, Guan Pang, Guillem Cucurell, Hailey Nguyen, Hannah Korevaar, Hu Xu, Hugo Touvron, Iliyan Zarov, Imanol Arrieta Ibarra, Isabel Kloumann, Is-han Misra, Ivan Evtimov, Jack Zhang, Jade Copet, Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park, Jay Mahadeokar, Jeet Shah, Jelmervan der Linde, Jennifer Billock, Jenny Hong, Jenya Lee, Jeremy Fu, Jianfeng Chi, Jianyu Huang, Jiawen Liu, Jie Wang, Jiecao Yu, Joanna Bitton, Joe Spisak, Jongsoo Park, Joseph Rocca, Joshua Johnstun, Joshua Saxe, Junteng Jia, Kalyan Vasuden Alwala, Karthik Prasad, Kartikeya Upasani, Kate Plawiak, Ke Li, Kenneth Heafield, Kevin Stone, Khalid El-Arini, Krithika Iyer, Kshitiz Malik, Kuenley Chiu, Kunal Bhalla, Kushal

394	Lakhotia, Lauren Rantala-Yearly, Laurens van der Maaten, Lawrence Chen, Liang Tan, Liz Jenkins,	Emily Wood, Eric-Tuan Le, Erik Brinkman, Este-	458
395	Louis Martin, Lovish Madaan, Lubo Malo, Lukas	ban Arcaute, Evan Dunbar, Evan Smothers, Fei Sun,	459
396	Blecher, Lukas Landzaat, Luke de Oliveira, Madeline	Felix Kreuk, Feng Tian, Filippas Kokkinos, Firat	460
397	Muzzi, Mahesh Pasupuleti, Mannat Singh, Manohar	Ozgenel, Francesco Caggioni, Frank Kanayet, Frank	461
398	Paluri, Marcin Kardas, Maria Tsimpoukelli, Mathew	Seide, Gabriela Medina Florez, Gabriella Schwarz,	462
399	Oldham, Mathieu Rita, Maya Pavlova, Melanie Kam-	Gada Badeer, Georgia Swee, Gil Halpern, Grant	463
400	badur, Mike Lewis, Min Si, Mitesh Kumar Singh,	Herman, Grigory Sizov, Guangyi, Zhang, Guna	464
401	Mona Hassan, Naman Goyal, Narjes Torabi, Niko-	Lakshminarayanan, Hakan Inan, Hamid Shojanaz-	465
402	lay Bashlykov, Nikolay Bogoychev, Niladri Chatterji,	eri, Han Zou, Hannah Wang, Hanwen Zha, Haroun	466
403	Ning Zhang, Olivier Duchenne, Onur Çelebi, Patrick	Habeeb, Harrison Rudolph, Helen Suk, Henry As-	467
404	Alrassy, Pengchuan Zhang, Pengwei Li, Petar Vas-	pegren, Hunter Goldman, Hongyuan Zhan, Ibrahim	468
405	sic, Peter Weng, Prajjwal Bhargava, Pratik Dubal,	Damlaj, Igor Molybog, Igor Tufanov, Ilias Leontiadis,	469
406	Praveen Krishnan, Punit Singh Koura, Puxin Xu,	Irina-Elena Veliche, Itai Gat, Jake Weissman, James	470
407	Qing He, Qingxiao Dong, Ragavan Srinivasan, Raj	Geboski, James Kohli, Janice Lam, Japhet Asher,	471
408	Ganapathy, Ramon Calderer, Ricardo Silveira Cabral,	Jean-Baptiste Gaya, Jeff Marcus, Jeff Tang, Jen-	472
409	Robert Stojnic, Roberta Raileanu, Rohan Maheswari,	nifer Chan, Jenny Zhen, Jeremy Reizenstein, Jeremy	473
410	Rohit Girdhar, Rohit Patel, Romain Sauvestre, Ron-	Teboul, Jessica Zhong, Jian Jin, Jingyi Yang, Joe	474
411	nie Polidoro, Roshan Sumbaly, Ross Taylor, Ruan	Cummings, Jon Carvill, Jon Shepard, Jonathan Mc-	475
412	Silva, Rui Hou, Rui Wang, Saghar Hosseini, Sa-	Phie, Jonathan Torres, Josh Ginsburg, Junjie Wang,	476
413	hana Chennabasappa, Sanjay Singh, Sean Bell, Seo-	Kai Wu, Kam Hou U, Karan Saxena, Kartikay Khan-	477
414	hyun Sonia Kim, Sergey Edunov, Shaoliang Nie, Sha-	delwal, Katayoun Zand, Kathy Matosich, Kaushik	478
415	ran Narang, Sharath Rapparth, Sheng Shen, Shengye	Veeraraghavan, Kelly Michelena, Keqian Li, Ki-	479
416	Wan, Shruti Bhosale, Shun Zhang, Simon Van-	ran Jagadeesh, Kun Huang, Kunal Chawla, Kyle	480
417	denhende, Soumya Batra, Spencer Whitman, Sten	Huang, Lailin Chen, Lakshya Garg, Lavender A,	481
418	Sootla, Stephane Collot, Suchin Gururangan, Syd-	Leandro Silva, Lee Bell, Lei Zhang, Liangpeng	482
419	ney Borodinsky, Tamar Herman, Tara Fowler, Tarek	Guo, Licheng Yu, Liron Moshkovich, Luca Wehrst-	483
420	Sheasha, Thomas Georgiou, Thomas Scialom, Tobias	edt, Madian Khabsa, Manav Avalani, Manish Bhatt,	484
421	Speckbacher, Todor Mihaylov, Tong Xiao, Ujjwal	Martynas Mankus, Matan Hasson, Matthew Lennie,	485
422	Karn, Vedanuj Goswami, Vibhor Gupta, Vignesh	Matthias Reso, Maxim Groshev, Maxim Naumov,	486
423	Ramanathan, Viktor Kerkez, Vincent Gonguet, Vir-	Maya Lathi, Meghan Keneally, Miao Liu, Michael L.	487
424	ginie Do, Vish Vogeti, Vitor Albiero, Vladan Petro-	Seltzer, Michal Valko, Michelle Restrepo, Mihir Pat-	488
425	vic, Weiwei Chu, Wenhan Xiong, Wenyin Fu, Whit-	tel, Mik Vyatskov, Mikayel Samvelyan, Mike Clark,	489
426	ney Meers, Xavier Martinet, Xiaodong Wang, Xi-	Mike Macey, Mike Wang, Miquel Jubert Hermoso,	490
427	aofang Wang, Xiaoqing Ellen Tan, Xide Xia, Xin-	Mo Metanat, Mohammad Rastegari, Munish Bansal,	491
428	feng Xie, Xuchao Jia, Xuwei Wang, Yaelle Gold-	Nandhini Santhanam, Natascha Parks, Natasha	492
429	schlag, Yashesh Gaur, Yasmine Babaei, Yi Wen,	White, Navyata Bawa, Nayan Singhal, Nick Egebo,	493
430	Yiwen Song, Yuchen Zhang, Yue Li, Yuning Mao,	Nicolas Usunier, Nikhil Mehta, Nikolay Pavlovich	494
431	Zacharie Delpierre Coudert, Zheng Yan, Zhengxing	Laptev, Ning Dong, Norman Cheng, Oleg Chernoguz,	495
432	Chen, Zoe Papakipos, Aaditya Singh, Aayushi Sri-	Olivia Hart, Omkar Salpekar, Ozlem Kalinli, Parkin	496
433	vastava, Abha Jain, Adam Kelsey, Adam Shajnfeld,	Kent, Parth Parekh, Paul Saab, Pavan Balaji, Pe-	497
434	Adithya Gangidi, Adolfo Victoria, Ahuva Goldstand,	dro Rittner, Philip Bontrager, Pierre Roux, Piotr	498
435	Ajay Menon, Ajay Sharma, Alex Boesenberg, Alexei	Dollar, Polina Zvyagina, Prashant Ratanchandani,	499
436	Baevski, Allie Feinstein, Amanda Kallet, Amit San-	Pritish Yuvraj, Qian Liang, Rachad Alao, Rachel	500
437	gani, Amos Teo, Anam Yunus, Andrei Lupu, An-	Rodriguez, Rafi Ayub, Raghotham Murthy, Raghu	501
438	dres Alvarado, Andrew Caples, Andrew Gu, Andrew	Nayani, Rahul Mitra, Rangaprabhu Parthasarathy,	502
439	Ho, Andrew Poulton, Andrew Ryan, Ankit Ramchan-	Raymond Li, Rebekkah Hogan, Robin Battey, Rocky	503
440	dani, Annie Dong, Annie Franco, Anuj Goyal, Apar-	Wang, Russ Howes, Ruty Rinott, Sachin Mehta,	504
441	jita Saraf, Arkabandhu Chowdhury, Ashley Gabriel,	Sachin Siby, Sai Jayesh Bondu, Samyak Datta, Sara	505
442	Ashwin Bharambe, Assaf Eisenman, Azadeh Yaz-	Chugh, Sara Hunt, Sargun Dhillon, Sasha Sidorov,	506
443	dan, Beau James, Ben Maurer, Benjamin Leonhardi,	Satadru Pan, Saurabh Mahajan, Saurabh Verma,	507
444	Bernie Huang, Beth Loyd, Beto De Paola, Bhargavi	Seiji Yamamoto, Sharadh Ramaswamy, Shaun Lind-	508
445	Paranjape, Bing Liu, Bo Wu, Boyu Ni, Braden Hancock,	say, Shaun Lindsay, Sheng Feng, Shenghao Lin,	509
446	Bram Wasti, Brandon Spence, Brani Stojkovic,	Shengxin Cindy Zha, Shishir Patil, Shiva Shankar,	510
447	Brian Gamido, Britt Montalvo, Carl Parker, Carly	Shuqiang Zhang, Shuqiang Zhang, Sinong Wang,	511
448	Burton, Catalina Mejia, Ce Liu, Changan Wang,	Sneha Agarwal, Soji Sajuyigbe, Soumith Chintala,	512
449	Changkyu Kim, Chao Zhou, Chester Hu, Ching-	Stephanie Max, Stephen Chen, Steve Kehoe, Steve	513
450	Hsiang Chu, Chris Cai, Chris Tindal, Christoph Fe-	Satterfield, Sudarshan Govindaprasad, Sumit Gupta,	514
451	ichtenhofer, Cynthia Gao, Damon Civin, Dana Beaty,	Summer Deng, Sungmin Cho, Sunny Virk, Suraj	515
452	Daniel Kreymer, Daniel Li, David Adkins, David	Subramanian, Sy Choudhury, Sydney Goldman, Tal	516
453	Xu, Davide Testuggine, Delia David, Devi Parikh,	Remez, Tamar Glaser, Tamara Best, Thilo Koehler,	517
454	Diana Liskovich, Didem Foss, Dingkan Wang, Duc	Thomas Robinson, Tianhe Li, Tianjun Zhang, Tim	518
455	Le, Dustin Holland, Edward Dowling, Eissa Jamil,	Matthews, Timothy Chou, Tzook Shaked, Varun	519
456	Elaine Montgomery, Eleonora Presani, Emily Hahn,	Vontimitta, Victoria Ajayi, Victoria Montanez, Vijai	520
457		Mohan, Vinay Satish Kumar, Vishal Mangla, Vlad	521

522	Ionescu, Vlad Poenaru, Vlad Tiberiu Mihailescu,	Zhaowei Li, Qi Xu, Dong Zhang, Hang Song, Yiqing	579
523	Vladimir Ivanov, Wei Li, Wenchen Wang, Wen-	Cai, Qi Qi, Ran Zhou, Junting Pan, Zefeng Li, Van Tu	580
524	wen Jiang, Wes Bouaziz, Will Constable, Xiaocheng	Vu, Zhida Huang, and Tao Wang. 2024c. Ground-	581
525	Tang, Xiaojian Wu, Xiaolan Wang, Xilun Wu, Xinbo	inggpt:language enhanced multi-modal grounding	582
526	Gao, Yaniv Kleinman, Yanjun Chen, Ye Hu, Ye Jia,	model .	583
527	Ye Qi, Yenda Li, Yilin Zhang, Ying Zhang, Yossi Adi,		
528	Youngjin Nam, Yu, Wang, Yu Zhao, Yuchen Hao,	Fuxiao Liu, Kevin Lin, Linjie Li, Jianfeng Wang, Yaser	584
529	Yundi Qian, Yunlu Li, Yuzi He, Zach Rait, Zachary	Yacoob, and Lijuan Wang. 2024a. Mitigating hal-	585
530	DeVito, Zef Rosnbrick, Zhaoduo Wen, Zhenyu Yang,	lucination in large multi-modal models via robust	586
531	Zhiwei Zhao, and Zhiyu Ma. 2024. The llama 3 herd	instruction tuning .	587
532	of models .		
533	Qidong Huang, Xiao wen Dong, Pan Zhang, Bin Wang,	Sheng Liu, Haotian Ye, Lei Xing, and James Zou. 2024b.	588
534	Conghui He, Jiaqi Wang, Dahua Lin, Weiming	Reducing hallucinations in vision-language models	589
535	Zhang, and Neng H. Yu. 2023. Opera: Alleviating	via latent space steering .	590
536	hallucination in multi-modal large language models		
537	via over-trust penalty and retrospection-allocation .	Ilya Loshchilov and Frank Hutter. 2019. Decoupled	591
538	2024 <i>IEEE/CVF Conference on Computer Vision and</i>	weight decay regularization .	592
539	<i>Pattern Recognition (CVPR)</i> , pages 13418–13427.		
540	Fushuo Huo, Wenchao Xu, Zhong Zhang, Haozhao	Junyu Lu, Dixiang Zhang, Songxin Zhang, Zejian Xie,	593
541	Wang, Zhicheng Chen, and Peilin Zhao. 2024. Self-	Zhuoyang Song, Cong Lin, Jiaying Zhang, Bingyi	594
542	introspective decoding: Alleviating hallucinations for	Jing, and Pingjian Zhang. 2024. Lyrics: Boosting	595
543	large vision-language models .	fine-grained language-vision alignment and compre-	596
544		hension via semantic-aware visual objects .	597
545	Gabriel Ilharco, Mitchell Wortsman, Ross Wightman,		
546	Cade Gordon, Nicholas Carlini, Rohan Taori, Achal	Avshalom Manevich and Reut Tsarfaty. 2024. Mitigat-	598
547	Dave, Vaishaal Shankar, Hongseok Namkoong, John	ing hallucinations in large vision-language models	599
548	Miller, Hannaneh Hajishirzi, Ali Farhadi, and Lud-	(LVLMs) via language-contrastive decoding (LCD) .	600
549	wig Schmidt. 2021. Openclip . If you use this soft-	In <i>Findings of the Association for Computational</i>	601
550	ware, please cite it as below.	<i>Linguistics: ACL 2024</i> , pages 6008–6022, Bangkok,	602
551	Jing Jin, Houfeng Wang, Hao Zhang, Xiaoguang Li,	Thailand. Association for Computational Linguistics.	603
552	and Zhijiang Guo. 2024. DVD: Dynamic con-		
553	trastive decoding for knowledge amplification in	OpenAI, :, Aaron Hurst, Adam Lerer, Adam P. Goucher,	604
554	multi-document question answering . In <i>Proceed-</i>	Adam Perelman, Aditya Ramesh, Aidan Clark,	605
555	<i>ings of the 2024 Conference on Empirical Methods</i>	AJ Ostrow, Akila Welihinda, Alan Hayes, Alec	606
556	<i>in Natural Language Processing</i> , pages 4624–4637,	Radford, Aleksander Mądry, Alex Baker-Whitcomb,	607
557	Miami, Florida, USA. Association for Computational	Alex Beutel, Alex Borzunov, Alex Carney, Alex	608
558	Linguistics.	Chow, Alex Kirillov, Alex Nichol, Alex Paino, Alex	609
559	Tiep Le, Vasudev Lal, and Phillip Howard. 2023. Coco-	Renzin, Alex Tachard Passos, Alexander Kirillov,	610
560	counterfactuals: Automatically constructed counter-	Alexi Christakis, Alexis Conneau, Ali Kamali, Allan	611
561	factual examples for image-text pairs .	Jabri, Allison Moyer, Allison Tam, Amadou Crookes,	612
562	Sicong Leng, Hang Zhang, Guanzheng Chen, Xin	Amin Tootoochian, Amin Tootoonchian, Ananya	613
563	Li, Shijian Lu, Chunyan Miao, and Lidong Bing.	Kumar, Andrea Vallone, Andrej Karpathy, Andrew	614
564	2023. Mitigating object hallucinations in large vision-	Braunstein, Andrew Cann, Andrew Codisputi, An-	615
565	language models through visual contrastive decoding .	drew Galu, Andrew Kondrich, Andrew Tulloch, An-	616
566	Baiqi Li, Zhiqiu Lin, Wenxuan Peng, Jean	drey Mishchenko, Angela Baek, Angela Jiang, An-	617
567	de Dieu Nyandwi, Daniel Jiang, Zixian Ma,	toine Pelisse, Antonia Woodford, Anuj Gosalia, Arka	618
568	Simran Khanuja, Ranjay Krishna, Graham Neubig,	Dhar, Ashley Pantuliano, Avi Nayak, Avital Oliver,	619
569	and Deva Ramanan. 2024a. Naturalbench: Evaluat-	Barret Zoph, Behrooz Ghorbani, Ben Leimberger,	620
570	ing vision-language models on natural adversarial	Ben Rossen, Ben Sokolowsky, Ben Wang, Benjamin	621
571	samples .	Zweig, Beth Hoover, Blake Samic, Bob McGrew,	622
572	Bo Li, Yuanhan Zhang, Dong Guo, Renrui Zhang, Feng	Bobby Spero, Bogo Gertler, Bowen Cheng, Brad	623
573	Li, Hao Zhang, Kaichen Zhang, Peiyuan Zhang, Yan-	Lightcap, Brandon Walkin, Brendan Quinn, Brian	624
574	wei Li, Ziwei Liu, and Chunyuan Li. 2024b. Llava-	Guarraci, Brian Hsu, Bright Kellogg, Brydon East-	625
575	onevision: Easy visual task transfer .	man, Camillo Lugaresi, Carroll Wainwright, Cary	626
576	Xiang Lisa Li, Ari Holtzman, Daniel Fried, Percy Liang,	Bassin, Cary Hudson, Casey Chu, Chad Nelson,	627
577	Jason Eisner, Tatsunori Hashimoto, Luke Zettle-	Chak Li, Chan Jun Shern, Channing Conger, Char-	628
578	moyer, and Mike Lewis. 2023. Contrastive decoding:	lotte Barette, Chelsea Voss, Chen Ding, Cheng Lu,	629
	Open-ended text generation as optimization .	Chong Zhang, Chris Beaumont, Chris Hallacy, Chris	630
		Koch, Christian Gibson, Christina Kim, Christine	631
		Choi, Christine McLeavey, Christopher Hesse, Clau-	632
		dia Fischer, Clemens Winter, Coley Czarnecki, Colin	633
		Jarvis, Colin Wei, Constantin Koumouzelis, Dane	634
		Sherburn, Daniel Kappler, Daniel Levin, Daniel Levy,	635
		David Carr, David Farhi, David Mely, David Robin-	636
		son, David Sasaki, Denny Jin, Dev Valladares, Dim-	637
		itris Tsipras, Doug Li, Duc Phong Nguyen, Duncan	638

639	Findlay, Edede Oiwoh, Edmund Wong, Ehsan As-	jan Troll, Randall Lin, Rapha Gontijo Lopes, Raul	703
640	dar, Elizabeth Proehl, Elizabeth Yang, Eric Antonow,	Puri, Reah Miyara, Reimar Leike, Renaud Gaubert,	704
641	Eric Kramer, Eric Peterson, Eric Sigler, Eric Wal-	Reza Zamani, Ricky Wang, Rob Donnelly, Rob	705
642	lace, Eugene Brevdo, Evan Mays, Farzad Khorasani,	Honsby, Rocky Smith, Rohan Sahai, Rohit Ramchan-	706
643	Felipe Petroski Such, Filippo Raso, Francis Zhang,	dani, Romain Huet, Rory Carmichael, Rowan Zellers,	707
644	Fred von Lohmann, Freddie Sulit, Gabriel Goh,	Roy Chen, Ruby Chen, Ruslan Nigmatullin, Ryan	708
645	Gene Oden, Geoff Salmon, Giulio Starace, Greg	Cheu, Saachi Jain, Sam Altman, Sam Schoenholz,	709
646	Brockman, Hadi Salman, Haiming Bao, Haitang	Sam Toizer, Samuel Miserendino, Sandhini Agar-	710
647	Hu, Hannah Wong, Haoyu Wang, Heather Schmidt,	wal, Sara Culver, Scott Ethersmith, Scott Gray, Sean	711
648	Heather Whitney, Heewoo Jun, Hendrik Kirchner,	Grove, Sean Metzger, Shamez Hermani, Shantanu	712
649	Henrique Ponde de Oliveira Pinto, Hongyu Ren,	Jain, Shengjia Zhao, Sherwin Wu, Shino Jomoto, Shi-	713
650	Huiwen Chang, Hyung Won Chung, Ian Kivlichan,	rong Wu, Shuaiqi, Xia, Sonia Phene, Spencer Papay,	714
651	Ian O’Connell, Ian O’Connell, Ian Osband, Ian Sil-	Srinivas Narayanan, Steve Coffey, Steve Lee, Stew-	715
652	ber, Ian Sohl, Ibrahim Okuyucu, Ikai Lan, Ilya	art Hall, Suchir Balaji, Tal Broda, Tal Stramer, Tao	716
653	Kostrikov, Ilya Sutskever, Ingmar Kanitscheider,	Xu, Tarun Gogineni, Taya Christianson, Ted Sanders,	717
654	Ishaan Gulrajani, Jacob Coxon, Jacob Menick, Jakub	Tejal Patwardhan, Thomas Cunningham, Thomas	718
655	Pachocki, James Aung, James Betker, James Crooks,	Degry, Thomas Dimson, Thomas Raoux, Thomas	719
656	James Lennon, Jamie Kiros, Jan Leike, Jane Park,	Shadwell, Tianhao Zheng, Todd Underwood, Todor	720
657	Jason Kwon, Jason Phang, Jason Teplitz, Jason	Markov, Toki Sherbakov, Tom Rubin, Tom Stasi,	721
658	Wei, Jason Wolfe, Jay Chen, Jeff Harris, Jenia Var-	Tomer Kaftan, Tristan Heywood, Troy Peterson, Tyce	722
659	avva, Jessica Gan Lee, Jessica Shieh, Ji Lin, Jiahui	Walters, Tyna Eloundou, Valerie Qi, Veit Moeller,	723
660	Yu, Jiayi Weng, Jie Tang, Jieqi Yu, Joanne Jang,	Vinnie Monaco, Vishal Kuo, Vlad Fomenko, Wayne	724
661	Joaquin Quinero Candela, Joe Beutler, Joe Lan-	Chang, Weiyi Zheng, Wenda Zhou, Wesam Manassra,	725
662	ders, Joel Parish, Johannes Heidecke, John Schul-	Will Sheu, Wojciech Zaremba, Yash Patil, Yilei Qian,	726
663	man, Jonathan Lachman, Jonathan McKay, Jonathan	Yongjik Kim, Youlong Cheng, Yu Zhang, Yuchen	727
664	Uesato, Jonathan Ward, Jong Wook Kim, Joost	He, Yuchen Zhang, Yujia Jin, Yunxing Dai, and Yury	728
665	Huizinga, Jordan Sitkin, Jos Kraaijeveld, Josh Gross,	Malkov. 2024. Gpt-4o system card .	729
666	Josh Kaplan, Josh Snyder, Joshua Achiam, Joy Jiao,		
667	Joyce Lee, Juntang Zhuang, Justyn Harriman, Kai	Phuc Phan, Hieu Tran, and Long Phan. 2024. Distilla-	730
668	Fricke, Kai Hayashi, Karan Singhal, Katy Shi, Kevin	tion contrastive decoding: Improving llms reasoning	731
669	Karthik, Kayla Wood, Kendra Rimbach, Kenny Hsu,	with contrastive decoding and distillation .	732
670	Kenny Nguyen, Keren Gu-Lemberg, Kevin Button,		
671	Kevin Liu, Kiel Howe, Krithika Muthukumar, Kyle	Alec Radford, Jong Wook Kim, Chris Hallacy,	733
672	Luther, Lama Ahmad, Larry Kai, Lauren Itow, Lau-	A. Ramesh, Gabriel Goh, Sandhini Agarwal, Girish	734
673	ren Workman, Leher Pathak, Leo Chen, Li Jing, Lia	Sastry, Amanda Askell, Pamela Mishkin, Jack Clark,	735
674	Guy, Liam Fedus, Liang Zhou, Lien Mamitsuka, Lil-	Gretchen Krueger, and Ilya Sutskever. 2021a. Learn-	736
675	ian Weng, Lindsay McCallum, Lindsey Held, Long	ing transferable visual models from natural language	737
676	Ouyang, Louis Feuvrier, Lu Zhang, Lukas Kon-	supervision . In <i>ICML</i> .	738
677	draciuk, Lukasz Kaiser, Luke Hewitt, Luke Metz,		
678	Lyric Doshi, Mada Aflak, Maddie Simens, Madelaine	Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya	739
679	Boyd, Madeleine Thompson, Marat Dukhan, Mark	Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sas-	740
680	Chen, Mark Gray, Mark Hudnall, Marvin Zhang,	try, Amanda Askell, Pamela Mishkin, Jack Clark,	741
681	Marwan Aljubeih, Mateusz Litwin, Matthew Zeng,	Gretchen Krueger, and Ilya Sutskever. 2021b. Learn-	742
682	Max Johnson, Maya Shetty, Mayank Gupta, Meghan	ing transferable visual models from natural language	743
683	Shah, Mehmet Yatbaz, Meng Jia Yang, Mengchao	supervision .	744
684	Zhong, Mia Glaese, Mianna Chen, Michael Jan-		
685	ner, Michael Lampe, Michael Petrov, Michael Wu,	Pritam Sarkar, Sayna Ebrahimi, Ali Etemad, Ahmad	745
686	Michele Wang, Michelle Fradin, Michelle Pokrass,	Beirami, Sercan Ö. Arik, and Tomas Pfister. 2024.	746
687	Miguel Castro, Miguel Oom Temudo de Castro,	Data-augmented phrase-level alignment for mitigat-	747
688	Mikhail Pavlov, Miles Brundage, Miles Wang, Mil-	ing object hallucination .	748
689	nal Khan, Mira Murati, Mo Bavarian, Molly Lin,		
690	Murat Yesildal, Nacho Soto, Natalia Gimelshein, Na-	Christoph Schuhmann, Romain Beaumont, Richard	749
691	talie Cone, Natalie Staudacher, Natalie Summers,	Vencu, Cade W Gordon, Ross Wightman, Mehdi	750
692	Natan LaFontaine, Neil Chowdhury, Nick Ryder,	Cherti, Theo Coombes, Aarush Katta, Clayton	751
693	Nick Stathas, Nick Turley, Nik Tezak, Niko Felix,	Mullis, Mitchell Wortsman, Patrick Schramowski,	752
694	Nithanth Kudige, Nitish Keskar, Noah Deutsch, Noel	Srivatsa R Kundurthy, Katherine Crowson, Lud-	753
695	Bundick, Nora Puckett, Ofir Nachum, Ola Okelola,	wig Schmidt, Robert Kaczmarczyk, and Jenia Jitsev.	754
696	Oleg Boiko, Oleg Murk, Oliver Jaffe, Olivia Watkins,	2022. LAION-5b: An open large-scale dataset for	755
697	Olivier Godement, Owen Campbell-Moore, Patrick	training next generation image-text models . In <i>Thirty-</i>	756
698	Chao, Paul McMillan, Pavel Belov, Peng Su, Pe-	<i>sixth Conference on Neural Information Processing</i>	757
699	ter Bak, Peter Bakkum, Peter Deng, Peter Dolan,	<i>Systems Datasets and Benchmarks Track</i> .	758
700	Peter Hoeschele, Peter Welinder, Phil Tillet, Philip		
701	Pronin, Philippe Tillet, Prafulla Dhariwal, Qiming	Rico Sennrich, Jannis Vamvas, and Alireza Moham-	759
702	Yuan, Rachel Dias, Rachel Lim, Rahul Arora, Ra-	madshahi. 2024. Mitigating hallucinations and off-	760
		target machine translation with source-contrastive	761
		and language-contrastive decoding .	762

- Shengbang Tong, Zhuang Liu, Yuexiang Zhai, Yi Ma, Yann LeCun, and Saining Xie. 2024. [Eyes wide shut? exploring the visual shortcomings of multi-modal llms.](#)
- Lei Wang, Jiabang He, Shenshen Li, Ning Liu, and Ee-Peng Lim. 2023. [Mitigating fine-grained hallucination by fine-tuning large vision-language models with caption rewrites.](#)
- 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. 2024a. [Qwen2-vl: Enhancing vision-language model’s perception of the world at any resolution.](#)
- Wei Wang, Zhaowei Li, Qi Xu, Linfeng Li, YiQing Cai, Botian Jiang, Hang Song, Xingcan Hu, Pengyu Wang, and Li Xiao. 2024b. [Advancing fine-grained visual understanding with multi-scale alignment in multi-modal models.](#)
- Dingchen Yang, Bowen Cao, Guang Chen, and Changjun Jiang. 2024. [Pensieve: Retrospect-then-compare mitigates visual hallucination.](#)
- Shukang Yin, Chaoyou Fu, Sirui Zhao, Tong Xu, Hao Wang, Dianbo Sui, Yunhang Shen, Ke Li, Xing Sun, and Enhong Chen. 2023. [Woodpecker: Hallucination correction for multimodal large language models.](#)
- Peter Young, Alice Lai, Micah Hodosh, and Julia Hockenmaier. 2014. [From image descriptions to visual denotations: New similarity metrics for semantic inference over event descriptions.](#) *Transactions of the Association for Computational Linguistics*, 2:67–78.
- Jianrui Zhang, Mu Cai, Tengyang Xie, and Yong Jae Lee. 2024. [Countercurate: Enhancing physical and semantic visio-linguistic compositional reasoning via counterfactual examples.](#)
- Baohang Zhou, Ying Zhang, Kehui Song, Hongru Wang, Yu Zhao, Xuhui Sui, and Xiaojie Yuan. 2024. [MCIL: Multimodal counterfactual instance learning for low-resource entity-based multimodal information extraction.](#) In *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, pages 11101–11110, Torino, Italia. ELRA and ICCL.
- Yiyang Zhou, Chenhang Cui, Jaehong Yoon, Linjun Zhang, Zhun Deng, Chelsea Finn, Mohit Bansal, and Huaxiu Yao. 2023. [Analyzing and mitigating object hallucination in large vision-language models.](#) *ArXiv*, abs/2310.00754.

A Appendix

A.1 Lightweight Classifier Implementation Details

Below is the PyTorch code of the lightweight classifier.

```
class Classifier(torch.nn.Module):
    def __init__(self, input_dim: int):
        super(Classifier, self).__init__()
        # factor of 2 due to concatentaion of target and candidate features
        self.linear1 = torch.nn.Linear(input_dim * 2, input_dim)
        self.linear2 = torch.nn.Linear(input_dim, input_dim)
        self.linear3 = torch.nn.Linear(input_dim, 1)
        self.dropout = torch.nn.Dropout(p=0.3)

    def forward(self, x) -> torch.Tensor:
        x = self.dropout(self.linear1(x))
        x = F.relu(x)
        x = self.dropout(self.linear2(x))
        x = F.relu(x)
        x = self.linear3(x)
        return x
```

We trained a classifier per tested LVLM, all with the following parameters, using the AdamW ([Loshchilov and Hutter, 2019](#)) optimizer.

```
batch_size=256
num_epochs=13
learning_rate=3e-4
weight_decay=1e-6
```

A.2 LVLM Checkpoints Tested

The following are the LVLM checkpoints we tested CoCI with:

```
Qwen/Qwen2-VL-7B-Instruct
llava-hf/llava-onevision-qwen2-7b-ov-hf
meta-llama/Llama-3.2-11B-Vision-Instruct
```

A.3 Effect of Choosing a Contrast Image on NaturalBench Performance

842

Method	Setting	Q-acc	I-acc	Acc	G-acc
CoCI ablations	Baseline	51.6	55.4	75.1	25.6
	CI $\leftarrow$ Random (out of top-5 most similar to input)	49.6	52.1	73.8	23.2
	CI $\leftarrow$ Natural	71.8	70.8	84.3	51.6
	CI $\leftarrow$ Most similar to input	49.7	52.5	73.6	23.9
	CI $\leftarrow$ Most similar to Natural	60.3	60.7	78.9	35.0
	CI $\leftarrow$ Least similar to Natural	46.7	48.9	72.6	21.8
Classifier	$k = 4$	51.7	54.3	74.5	26.6
	$k = 8$	53.0	55.4	75.3	26.6
	$k = 16$	54.3	56.8	76.1	29.2
	$k = 32$	52.2	54.6	75.1	25.8
	$k = 64$	51.8	53.9	74.7	26.3
	$k = 100$	52.1	54.1	74.8	25.5
Classifier+augmentations	$k = 4$	52.0	54.3	74.6	27.1
	$k = 8$	52.8	55.9	75.0	27.9
	$k = 16$	54.5	57.8	76.1	29.2
	$k = 32$	54.9	58.2	75.9	30.0
	$k = 64$	54.7	57.9	76.1	30.3
	$k = 100$	54.7	58.0	76.1	30.0

Table 3: CoCI performance on the NaturalBench dev set with different CI selection methods, using Qwen2-VL. Classifier+augmentations indicates training data augmentation with GPT-4o paraphrased questions and standard image augmentations. Using natural CIs provides the strongest performance gains, with a 26-point improvement in group accuracy over baseline (51.6% vs 25.6%). Selecting CIs by similarity to natural CIs improves performance significantly (35.0% G-acc), while using the least similar images performs worse than baseline (21.8%), validating the importance of CI selection strategy. Random CI selection hurts performance (23.2% G-acc) even when restricted to similar images, highlighting that similarity alone is insufficient. Training with augmented data provides modest but consistent improvements across all metrics, with G-acc increasing by about 4 points compared to the non-augmented classifier. The augmented classifier also demonstrates more robust performance, maintaining consistent scores across different k values compared to the higher variance seen in the non-augmented version.