

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 cross-modal reasoning, 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%, 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). CoCI demonstrates strong performance across three distinct supervision regimes: First, when using naturally occurring CIs in benchmarks with curated image pairs, we achieve improvements of up to 98.9% on NaturalBench, 69.5% on D3, and 37.6% on MMVP. Second, for scenarios with modest training data ($\sim 5k$ samples), we show that a lightweight neural classifier can effectively select CIs from similar images at inference time, improving NaturalBench performance by up to 36.8%. Third, for scenarios with no training data, we develop a caption-matching technique that selects CIs by comparing LVLM-generated descriptions of candidate images. Notably, on VQAv2, our method improves VQA performance even in pointwise evaluation settings without explicit contrast images. Our approach demonstrates the potential for enhancing 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 mul-

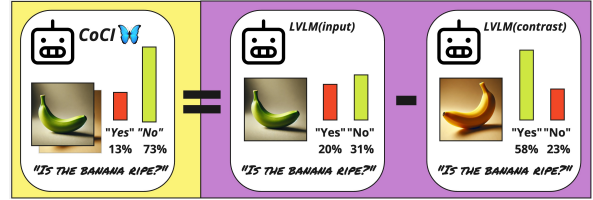


Figure 1: CoCI penalizes target image logits using those from a contrast image, weighted by hyperparameter α .

timodal search, manufacturing, and robotics. Recent benchmarks have exposed this limitation: on NaturalBench (Li et al., 2024a), which tests visual 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% respectively. 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 by conditioning next-token prediction on both text and images.¹ Generation proceeds by iteratively sampling tokens from the model’s predicted distributions until reaching an EOS token or length limit. The LVLM next-token prediction 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, I is the input image, and $\mathcal{V}$ is the model’s vocabulary.

2.2 Contrastive Decoding

Following Li et al. (2023), various Contrastive Decoding approaches have emerged (Sennrich et al.,

2024; Jin et al., 2024; Phan et al., 2024). We implement CoCI based on Sennrich et al. (2024)’s minimal variant:

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

CoCI penalizes token probabilities from the target image distribution $P(y|y_{<t}, I)$ with those from the contrast image distribution $P(y|y_{<t}, I')$. The parameter α controls 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 an MLP classifier to select them during inference. Given LVLM L and training triplets $\langle q, I, I' \rangle$ (binary question and image pairs with different answers), we: (a) Extract LVLM hidden states $h_{q,i} \in \mathbb{R}^{d_L}$ per image-question pair. (b) Concatenate states for image pairs: $h_{q,i,i'} \in \mathbb{R}^{2*d_L}$. (c) Create negative samples using the j least similar images from top- k similar images to I in dataset D^3 . (d) Train a three-layer MLP classifier.⁴ We train on NaturalBench (60% split) augmented with GPT-4-generated question paraphrases. At inference, we select the CI maximizing classifier score among k most similar images.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)

²We use $\alpha = 0.5$ for VQA and $\alpha = 0.8$ for open-ended generation.

³ $j = 5, k = 100$. Using flickr30k (Young et al., 2014) and open-clip (Ilharco et al., 2021; Cherti et al., 2023; Radford et al., 2021a; Schuhmann et al., 2022) with cosine similarity.

⁴See appendix A.1 and A.3 for implementation details.

⁵See table 2 for k value comparisons. Inference uses identical retrieval setup as training.

⁶We set $k = 5$ without tuning.

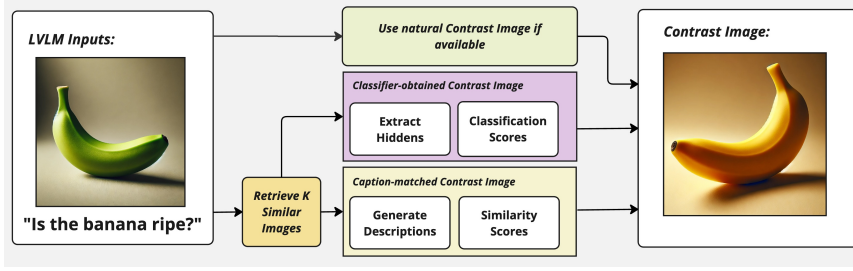


Figure 2: Illustration of the approaches we explore for obtaining a Contrast Image (CI).

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 using three leading LVLms⁷ on four benchmarks, three specifically targeting fine-grained visual discrimination:

NaturalBench (Li et al., 2024a) evaluates similar image discrimination through yes/no and multiple-choice questions, with different answers for paired images. The benchmark contains 1900 image pairs (two questions per pair), split into train (60%), dev (20%), and test (20%) sets. We measure image accuracy (both questions correct), question accuracy (per-question), and group accuracy (all four image-question combinations correct).

MMVP (Multimodal Visual Patterns) (Tong et al., 2024) evaluates visual difference detection through multiple-choice questions on 150 image pairs. Each pair differs in specific visual aspects (object state, position, or orientation). Success requires correct answers for both images in a pair.

D3 (Detect, Describe, Discriminate) (Gaur et al., 2024) assesses models’ ability to generate discriminative descriptions between similar images across 247 pairs. We adapt D3 for CoCI by treating it as a single-input task, generating separate descriptions per image. Evaluation follows the original self-retrieval protocol, measuring whether an

Model	Method	D3 (self-ret.)	MMVP (acc.)	NB (g-acc.)	VQAv2 (acc.)
Qwen2-VL	Baseline	30.8	46.0	30.8	72.66
	CoCI _{CAP}	34.8	48.7	31.3	74.33
	CoCI _{NAT}	52.2	63.3	46.6	-
LLaVA-OV	Baseline	25.1	52.7	28.2	61.66
	CoCI _{CAP}	31.6	57.3	31.6	73.66
	CoCI _{NAT}	38.1	66.7	56.1	-
Llama 3.2	Baseline	28.7	39.3	21.1	58
	CoCI _{CAP}	33.6	41.3	22.4	58
	CoCI _{NAT}	35.6	43.3	29.2	-

Table 1: CoCI performance comparison with provided CIs across benchmarks, with natural CIs (CoCI_{NAT}) and caption-matched CIs (CoCI_{CAP}).

image-text encoder correctly matches descriptions to their images.

VQAv2 (Goyal et al., 2017) serves as our general-purpose visual question answering benchmark. While not focused on fine-grained discrimination, we include it to demonstrate CoCI’s broader applicability. We evaluate on 300 validation set image-question pairs using exact match accuracy.

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, particularly on D3 where LLaVA improves from 25.1% to 31.6%, suggesting that contrasting against images with similar captions effectively guides visual discrimination. CoCI with caption matching improves performance on VQAv2 for two of the three tested models while maintaining baseline performance for Llama 3.2, demonstrating that CoCI enhances general-purpose VQA abilities beyond fine-grained visual discrimination tasks.

Throughout our experiments, Llama exhibits different behavior compared to other models - showing lower performance and reduced responsiveness

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

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.

to our methods. This pattern is evident in Table 2, where Qwen and LLaVA’s performance improves with larger candidate pools (k), peaking around $k=32$, while Llama performs best with small pools ($k=4$). This behavior could be attributed to two factors: First, while the hyperparameters worked well for Qwen and LLaVA, they may not be optimal for Llama without model-specific tuning. Second, Llama’s architectural differences, particularly its use of cross-attention, could lead to different behaviors in our contrastive decoding context. While exploring these architecture-specific considerations could be valuable, it is beyond the scope of this work.

In NaturalBench, G-Acc shows particularly strong improvement with natural CIs 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 modest gains. The substantial gap between natural CIs and other methods suggests that classifier-selected and caption-matched CIs, while beneficial, 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 adjustments (Huo et al., 2024), 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 fine-grained visual discrimination.

Alignment and grounding in LVLMs. Prior work has enhanced visual-textual alignment through object-level synthesis (Wang et al., 2024b), targeted fine-tuning (Lu et al., 2024), and 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 for modality alignment. Recent works leverage contrast pairs: (Le et al., 2023) and (Zhang et al., 2024) generate synthetic 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 requiring data generation or training, our method operates at inference time.⁹

6 Conclusion

We introduced Contrastive decoding with Contrast Images (CoCI), demonstrating its effectiveness in improving LVLMs’ fine-grained visual discrimination capabilities in both VQA and long-form generation tasks. While naturally occurring contrast pairs yielded the strongest gains, both classifier-based and caption-matching approaches provide meaningful improvements without requiring dataset curation or model training. We validated the generality of our method through experiments with caption-based contrast selection, showing that CoCI does not rely on pre-curated pairs but can leverage them when available. Notably, CoCI improves performance even on tasks that don’t explicitly measure fine-grained discrimination.

Our results show that contrastive decoding algorithms, when combined with strategic contrast image selection, improve LVLMs’ ability to make fine-grained distinctions and their overall VQA abilities, 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](#).
- Yash Goyal, Tejas Khot, Douglas Summers-Stay, Dhruv Batra, and Devi Parikh. 2017. [Making the v in vqa matter: Elevating the role of image understanding in visual question answering](#).
- 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 Alonsoius, 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, Ishan Misra, Ivan Evtimov, Jack Zhang, Jade Copet, Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park, Jay Mahadeokar, Jeet Shah, Jelmer van 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,

402	Joseph Rocca, Joshua Johnstun, Joshua Saxe, Jun-	Daniel Kreymer, Daniel Li, David Adkins, David	466
403	teng Jia, Kalyan Vasuden Alwala, Karthik Prasad,	Xu, Davide Testuggine, Delia David, Devi Parikh,	467
404	Kartikaya Upasani, Kate Plawiak, Ke Li, Kenneth	Diana Liskovich, Didem Foss, Dingkan Wang, Duc	468
405	Heafield, Kevin Stone, Khalid El-Arini, Krithika Iyer,	Le, Dustin Holland, Edward Dowling, Eissa Jamil,	469
406	Kshitiz Malik, Kuenley Chiu, Kunal Bhalla, Kushal	Elaine Montgomery, Eleonora Presani, Emily Hahn,	470
407	Lakhotia, Lauren Rantala-Yearly, Laurens van der	Emily Wood, Eric-Tuan Le, Erik Brinkman, Este-	471
408	Maaten, Lawrence Chen, Liang Tan, Liz Jenkins,	ban Arcaute, Evan Dunbar, Evan Smothers, Fei Sun,	472
409	Louis Martin, Lovish Madaan, Lubo Malo, Lukas	Felix Kreuk, Feng Tian, Filippas Kokkinos, Firat	473
410	Blecher, Lukas Landzaat, Luke de Oliveira, Madeline	Ozgenel, Francesco Caggioni, Frank Kanayet, Frank	474
411	Muzzi, Mahesh Pasupuleti, Mannat Singh, Manohar	Seide, Gabriela Medina Florez, Gabriella Schwarz,	475
412	Paluri, Marcin Kardas, Maria Tsimpoukelli, Mathew	Gada Badeer, Georgia Swee, Gil Halpern, Grant	476
413	Oldham, Mathieu Rita, Maya Pavlova, Melanie Kam-	Herman, Grigory Sizov, Guangyi, Zhang, Guna	477
414	badur, Mike Lewis, Min Si, Mitesh Kumar Singh,	Lakshminarayanan, Hakan Inan, Hamid Shojanaz-	478
415	Mona Hassan, Naman Goyal, Narjes Torabi, Niko-	eri, Han Zou, Hannah Wang, Hanwen Zha, Haroun	479
416	lay Bashlykov, Nikolay Bogoychev, Niladri Chatterji,	Habeeb, Harrison Rudolph, Helen Suk, Henry As-	480
417	Ning Zhang, Olivier Duchenne, Onur Çelebi, Patrick	pegren, Hunter Goldman, Hongyuan Zhan, Ibrahim	481
418	Alrassy, Pengchuan Zhang, Pengwei Li, Petar Va-	Damlaj, Igor Molybog, Igor Tufanov, Ilias Leontiadis,	482
419	sic, Peter Weng, Prajjwal Bhargava, Pratik Dubal,	Irina-Elena Veliche, Itai Gat, Jake Weissman, James	483
420	Praveen Krishnan, Punit Singh Koura, Puxin Xu,	Geboski, James Kohli, Janice Lam, Japhet Asher,	484
421	Qing He, Qingxiao Dong, Ragavan Srinivasan, Raj	Jean-Baptiste Gaya, Jeff Marcus, Jeff Tang, Jen-	485
422	Ganapathy, Ramon Calderer, Ricardo Silveira Cabral,	nifer Chan, Jenny Zhen, Jeremy Reizenstein, Jeremy	486
423	Robert Stojnic, Roberta Raileanu, Rohan Maheswari,	Teboul, Jessica Zhong, Jian Jin, Jingyi Yang, Joe	487
424	Rohit Girdhar, Rohit Patel, Romain Sauvestre, Ron-	Cummings, Jon Carvill, Jon Shepard, Jonathan Mc-	488
425	nie Polidoro, Roshan Sumbaly, Ross Taylor, Ruan	Phie, Jonathan Torres, Josh Ginsburg, Junjie Wang,	489
426	Silva, Rui Hou, Rui Wang, Saghar Hosseini, Sa-	Kai Wu, Kam Hou U, Karan Saxena, Kartikay Khan-	490
427	hana Chennabasappa, Sanjay Singh, Sean Bell, Seo-	delwal, Katayoun Zand, Kathy Matosich, Kaushik	491
428	hyun Sonia Kim, Sergey Edunov, Shaoliang Nie, Sha-	Veeraraghavan, Kelly Michelen, Keqian Li, Ki-	492
429	ran Narang, Sharath Raparthi, Sheng Shen, Shengye	ran Jagadeesh, Kun Huang, Kunal Chawla, Kyle	493
430	Wan, Shruti Bhosale, Shun Zhang, Simon Van-	Huang, Lailin Chen, Lakshya Garg, Lavender A,	494
431	denhende, Soumya Batra, Spencer Whitman, Sten	Leandro Silva, Lee Bell, Lei Zhang, Liangpeng	495
432	Sootla, Stephane Collot, Suchin Gururangan, Syd-	Guo, Licheng Yu, Liron Moshkovich, Luca Wehrst-	496
433	ney Borodinsky, Tamar Herman, Tara Fowler, Tarek	edt, Madian Khabsa, Manav Avalani, Manish Bhatt,	497
434	Sheasha, Thomas Georgiou, Thomas Scialom, Tobias	Martynas Mankus, Matan Hasson, Matthew Lennie,	498
435	Speckbacher, Todor Mihaylov, Tong Xiao, Ujjwal	Matthias Reso, Maxim Groshev, Maxim Naumov,	499
436	Karn, Vedanuj Goswami, Vibhor Gupta, Vignesh	Maya Lathi, Meghan Keneally, Miao Liu, Michael L.	500
437	Ramanathan, Viktor Kerkez, Vincent Gonguet, Vir-	Seltzer, Michal Valko, Michelle Restrepo, Mihir Pa-	501
438	ginie Do, Vish Vogeti, Vítor Albiero, Vladan Petro-	tel, Mik Vyatskov, Mikayel Samvelyan, Mike Clark,	502
439	vic, Weiwei Chu, Wenhan Xiong, Wenyin Fu, Whit-	Mike Macey, Mike Wang, Miquel Jubert Hermoso,	503
440	ney Meers, Xavier Martinet, Xiaodong Wang, Xi-	Mo Metanat, Mohammad Rastegari, Munish Bansal,	504
441	aofang Wang, Xiaoqing Ellen Tan, Xide Xia, Xin-	Nandhini Santhanam, Natascha Parks, Natasha	505
442	feng Xie, Xuchao Jia, Xuewei Wang, Yaelle Gold-	White, Navyata Bawa, Nayan Singhal, Nick Egebo,	506
443	schlag, Yashesh Gaur, Yasmine Babaei, Yi Wen,	Nicolas Usunier, Nikhil Mehta, Nikolay Pavlovich	507
444	Yiwen Song, Yuchen Zhang, Yue Li, Yuning Mao,	Laptev, Ning Dong, Norman Cheng, Oleg Chernoguz,	508
445	Zacharie Delpierre Coudert, Zheng Yan, Zhengxing	Olivia Hart, Omkar Salpekar, Ozlem Kalinli, Parkin	509
446	Chen, Zoe Papakipos, Aaditya Singh, Aayushi Sri-	Kent, Parth Parekh, Paul Saab, Pavan Balaji, Pe-	510
447	vastava, Abha Jain, Adam Kelsey, Adam Shajnfeld,	dro Rittner, Philip Bontrager, Pierre Roux, Piotr	511
448	Adithya Gangidi, Adolfo Victoria, Ahuva Goldstand,	Dollar, Polina Zvyagina, Prashant Ratanchandani,	512
449	Ajay Menon, Ajay Sharma, Alex Boesenberg, Alexei	Pritish Yuvraj, Qian Liang, Rachad Alao, Rachel	513
450	Baevski, Allie Feinstein, Amanda Kallet, Amit San-	Rodriguez, Rafi Ayub, Raghotham Murthy, Raghu	514
451	gani, Amos Teo, Anam Yunus, Andrei Lupu, An-	Nayani, Rahul Mitra, Rangaprabhu Parthasarathy,	515
452	dres Alvarado, Andrew Caples, Andrew Gu, Andrew	Raymond Li, Rebekkah Hogan, Robin Battey, Rocky	516
453	Ho, Andrew Poulton, Andrew Ryan, Ankit Ramchan-	Wang, Russ Howes, Ruty Rinott, Sachin Mehta,	517
454	dani, Annie Dong, Annie Franco, Anuj Goyal, Apar-	Sachin Siby, Sai Jayesh Bondu, Samyak Datta, Sara	518
455	jita Saraf, Arkabandhu Chowdhury, Ashley Gabriel,	Chugh, Sara Hunt, Sargun Dhillon, Sasha Sidorov,	519
456	Ashwin Bharambe, Assaf Eisenman, Azadeh Yaz-	Satadru Pan, Saurabh Mahajan, Saurabh Verma,	520
457	dan, Beau James, Ben Maurer, Benjamin Leonhardi,	Seiji Yamamoto, Sharadh Ramaswamy, Shaun Lind-	521
458	Bernie Huang, Beth Loyd, Beto De Paola, Bhargavi	say, Shaun Lindsay, Sheng Feng, Shenghao Lin,	522
459	Paranjape, Bing Liu, Bo Wu, Boyu Ni, Braden Han-	Shengxin Cindy Zha, Shishir Patil, Shiva Shankar,	523
460	cock, Bram Wasti, Brandon Spence, Brani Stojkovic,	Shuqiang Zhang, Shuqiang Zhang, Sinong Wang,	524
461	Brian Gamido, Britt Montalvo, Carl Parker, Carly	Sneha Agarwal, Soji Sajuyigbe, Soumith Chintala,	525
462	Burton, Catalina Mejia, Ce Liu, Changhan Wang,	Stephanie Max, Stephen Chen, Steve Kehoe, Steve	526
463	Changkyu Kim, Chao Zhou, Chester Hu, Ching-	Satterfield, Sudarshan Govindaprasad, Sumit Gupta,	527
464	Hsiang Chu, Chris Cai, Chris Tindal, Christoph Fe-	Summer Deng, Sungmin Cho, Sunny Virk, Suraj	528
465	ichtenhofer, Cynthia Gao, Damon Civin, Dana Beaty,	Subramanian, Sy Choudhury, Sydney Goldman, Tal	529

530	Remez, Tamar Glaser, Tamara Best, Thilo Koehler,	Xiang Lisa Li, Ari Holtzman, Daniel Fried, Percy Liang,	588
531	Thomas Robinson, Tianhe Li, Tianjun Zhang, Tim	Jason Eisner, Tatsunori Hashimoto, Luke Zettle-	589
532	Matthews, Timothy Chou, Tzook Shaked, Varun	moyer, and Mike Lewis. 2023. Contrastive decoding:	590
533	Vontimitta, Victoria Ajayi, Victoria Montanez, Vijai	Open-ended text generation as optimization.	591
534	Mohan, Vinay Satish Kumar, Vishal Mangla, Vlad		
535	Ionescu, Vlad Poenaru, Vlad Tiberiu Mihailescu,	Zhaowei Li, Qi Xu, Dong Zhang, Hang Song, Yiqing	592
536	Vladimir Ivanov, Wei Li, Wenchen Wang, Wen-	Cai, Qi Qi, Ran Zhou, Juntong Pan, Zefeng Li, Van Tu	593
537	wen Jiang, Wes Bouaziz, Will Constable, Xiaocheng	Vu, Zhida Huang, and Tao Wang. 2024c. Ground-	594
538	Tang, Xiaojian Wu, Xiaolan Wang, Xilun Wu, Xinbo	inggpt:language enhanced multi-modal grounding	595
539	Gao, Yaniv Kleinman, Yanjun Chen, Ye Hu, Ye Jia,	model.	596
540	Ye Qi, Yenda Li, Yilin Zhang, Ying Zhang, Yossi Adi,		
541	Youngjin Nam, Yu, Wang, Yu Zhao, Yuchen Hao,	Fuxiao Liu, Kevin Lin, Linjie Li, Jianfeng Wang, Yaser	597
542	Yundi Qian, Yunlu Li, Yuzi He, Zach Rait, Zachary	Yacoob, and Lijuan Wang. 2024a. Mitigating hal-	598
543	DeVito, Zef Rosnbrick, Zhaoduo Wen, Zhenyu Yang,	lucination in large multi-modal models via robust	599
544	Zhiwei Zhao, and Zhiyu Ma. 2024. The llama 3 herd	instruction tuning.	600
545	of models.		
546	Qidong Huang, Xiao wen Dong, Pan Zhang, Bin Wang,	Sheng Liu, Haotian Ye, Lei Xing, and James Zou. 2024b.	601
547	Conghui He, Jiaqi Wang, Dahua Lin, Weiming	Reducing hallucinations in vision-language models	602
548	Zhang, and Neng H. Yu. 2023. Opera: Alleviating	via latent space steering.	603
549	hallucination in multi-modal large language models		
550	via over-trust penalty and retrospection-allocation.	Ilya Loshchilov and Frank Hutter. 2019. Decoupled	604
551	2024 IEEE/CVF Conference on Computer Vision and	weight decay regularization.	605
552	Pattern Recognition (CVPR), pages 13418–13427.		
553	Fushuo Huo, Wenchao Xu, Zhong Zhang, Haozhao	Junyu Lu, Dixiang Zhang, Songxin Zhang, Zejian Xie,	606
554	Wang, Zhicheng Chen, and Peilin Zhao. 2024. Self-	Zhuoyang Song, Cong Lin, Jiaying Zhang, Bingyi	607
555	introspective decoding: Alleviating hallucinations for	Jing, and Pingjian Zhang. 2024. Lyrics: Boosting	608
556	large vision-language models.	fine-grained language-vision alignment and compre-	609
557	Gabriel Ilharco, Mitchell Wortsman, Ross Wightman,	hension via semantic-aware visual objects.	610
558	Cade Gordon, Nicholas Carlini, Rohan Taori, Achal		
559	Dave, Vaishaal Shankar, Hongseok Namkoong, John	Avshalom Manevich and Reut Tsarfaty. 2024. Mitigat-	611
560	Miller, Hannaneh Hajishirzi, Ali Farhadi, and Lud-	ing hallucinations in large vision-language models	612
561	wig Schmidt. 2021. Openclip . If you use this soft-	(LVLMs) via language-contrastive decoding (LCD).	613
562	ware, please cite it as below.	In Findings of the Association for Computational	614
563	Jing Jin, Houfeng Wang, Hao Zhang, Xiaoguang Li,	Linguistics: ACL 2024, pages 6008–6022, Bangkok,	615
564	and Zhijiang Guo. 2024. DVD: Dynamic con-	Thailand. Association for Computational Linguistics.	616
565	trastive decoding for knowledge amplification in		
566	multi-document question answering. In Proceed-	OpenAI, :, Aaron Hurst, Adam Lerer, Adam P. Goucher,	617
567	ings of the 2024 Conference on Empirical Methods	Adam Perelman, Aditya Ramesh, Aidan Clark,	618
568	in Natural Language Processing, pages 4624–4637,	AJ Ostrow, Akila Welihinda, Alan Hayes, Alec	619
569	Miami, Florida, USA. Association for Computational	Radford, Aleksander Mařdry, Alex Baker-Whitcomb,	620
570	Linguistics.	Alex Beutel, Alex Borzunov, Alex Carney, Alex	621
571	Tiep Le, Vasudev Lal, and Phillip Howard. 2023. Coco-	Chow, Alex Kirillov, Alex Nichol, Alex Paino, Alex	622
572	counterfactuals: Automatically constructed counter-	Renzin, Alex Tachard Passos, Alexander Kirillov,	623
573	factual examples for image-text pairs.	Alexi Christakis, Alexis Conneau, Ali Kamali, Allan	624
574	Sicong Leng, Hang Zhang, Guanzheng Chen, Xin	Jabri, Allison Moyer, Allison Tam, Amadou Crookes,	625
575	Li, Shijian Lu, Chunyan Miao, and Lidong Bing.	Amin Tootoochian, Amin Tootoonchian, Ananya	626
576	2023. Mitigating object hallucinations in large vision-	Kumar, Andrea Vallone, Andrej Karpathy, Andrew	627
577	language models through visual contrastive decoding.	Braunstein, Andrew Cann, Andrew Codisputi, An-	628
578	Baiqi Li, Zhiqiu Lin, Wenxuan Peng, Jean	drew Galu, Andrew Kondrich, Andrew Tulloch, An-	629
579	de Dieu Nyandwi, Daniel Jiang, Zixian Ma,	drey Mishchenko, Angela Baek, Angela Jiang, An-	630
580	Simran Khanuja, Ranjay Krishna, Graham Neubig,	toine Pelisse, Antonia Woodford, Anuj Gosalia, Arka	631
581	and Deva Ramanan. 2024a. Naturalbench: Evaluat-	Dhar, Ashley Pantuliano, Avi Nayak, Avital Oliver,	632
582	ing vision-language models on natural adversarial	Barret Zoph, Behrooz Ghorbani, Ben Leimberger,	633
583	samples.	Ben Rossen, Ben Sokolowsky, Ben Wang, Benjamin	634
584	Bo Li, Yuanhan Zhang, Dong Guo, Renrui Zhang, Feng	Zweig, Beth Hoover, Blake Samic, Bob McGrew,	635
585	Li, Hao Zhang, Kaichen Zhang, Peiyuan Zhang, Yan-	Bobby Spero, Bogo Gertler, Bowen Cheng, Brad	636
586	wei Li, Ziwei Liu, and Chunyuan Li. 2024b. Llava-	Lightcap, Brandon Walkin, Brendan Quinn, Brian	637
587	onevision: Easy visual task transfer.	Guarraci, Brian Hsu, Bright Kellogg, Brydon East-	638
		man, Camillo Lugaresi, Carroll Wainwright, Cary	639
		Bassin, Cary Hudson, Casey Chu, Chad Nelson,	640
		Chak Li, Chan Jun Shern, Channing Conger, Char-	641
		lotte Barette, Chelsea Voss, Chen Ding, Cheng Lu,	642
		Chong Zhang, Chris Beaumont, Chris Hallacy, Chris	643
		Koch, Christian Gibson, Christina Kim, Christine	644
		Choi, Christine McLeavey, Christopher Hesse, Clau-	645
		dia Fischer, Clemens Winter, Coley Czarnecki, Colin	646

647	Jarvis, Colin Wei, Constantin Koumouzelis, Dane	Chao, Paul McMillan, Pavel Belov, Peng Su, Pe-	711
648	Sherburn, Daniel Kappler, Daniel Levin, Daniel Levy,	ter Bak, Peter Bakkum, Peter Deng, Peter Dolan,	712
649	David Carr, David Farhi, David Mely, David Robin-	Peter Hoeschele, Peter Welinder, Phil Tillet, Philip	713
650	son, David Sasaki, Denny Jin, Dev Valladares, Dim-	Pronin, Philippe Tillet, Prafulla Dhariwal, Qiming	714
651	itris Tsipras, Doug Li, Duc Phong Nguyen, Duncan	Yuan, Rachel Dias, Rachel Lim, Rahul Arora, Ra-	715
652	Findlay, Edede Oiwoh, Edmund Wong, Ehsan As-	jan Troll, Randall Lin, Rapha Gontijo Lopes, Raul	716
653	dar, Elizabeth Proehl, Elizabeth Yang, Eric Antonow,	Puri, Reah Miyara, Reimar Leike, Renaud Gaubert,	717
654	Eric Kramer, Eric Peterson, Eric Sigler, Eric Wal-	Reza Zamani, Ricky Wang, Rob Donnelly, Rob	718
655	lace, Eugene Brevdo, Evan Mays, Farzad Khorasani,	Honsby, Rocky Smith, Rohan Sahai, Rohit Ramchan-	719
656	Felipe Petroski Such, Filippo Raso, Francis Zhang,	dani, Romain Huet, Rory Carmichael, Rowan Zellers,	720
657	Fred von Lohmann, Freddie Sulit, Gabriel Goh,	Roy Chen, Ruby Chen, Ruslan Nigmatullin, Ryan	721
658	Gene Oden, Geoff Salmon, Giulio Starace, Greg	Cheu, Saachi Jain, Sam Altman, Sam Schoenholz,	722
659	Brockman, Hadi Salman, Haiming Bao, Haitang	Sam Toizer, Samuel Miserendino, Sandhini Agar-	723
660	Hu, Hannah Wong, Haoyu Wang, Heather Schmidt,	wal, Sara Culver, Scott Ethersmith, Scott Gray, Sean	724
661	Heather Whitney, Heewoo Jun, Hendrik Kirchner,	Grove, Sean Metzger, Shamez Hermani, Shantanu	725
662	Henrique Ponde de Oliveira Pinto, Hongyu Ren,	Jain, Shengjia Zhao, Sherwin Wu, Shino Jomoto, Shi-	726
663	Huiwen Chang, Hyung Won Chung, Ian Kivlichan,	rong Wu, Shuaiqi, Xia, Sonia Phene, Spencer Papay,	727
664	Ian O’Connell, Ian O’Connell, Ian Osband, Ian Sil-	Srinivas Narayanan, Steve Coffey, Steve Lee, Stew-	728
665	ber, Ian Sohl, Ibrahim Okuyucu, Ikai Lan, Ilya	art Hall, Suchir Balaji, Tal Broda, Tal Stramer, Tao	729
666	Kostrikov, Ilya Sutskever, Ingmar Kanitscheider,	Xu, Tarun Gogineni, Taya Christianson, Ted Sanders,	730
667	Ishaan Gulrajani, Jacob Coxon, Jacob Menick, Jakub	Tejal Patwardhan, Thomas Cunninghamman, Thomas	731
668	Pachocki, James Aung, James Betker, James Crooks,	Degry, Thomas Dimson, Thomas Raoux, Thomas	732
669	James Lennon, Jamie Kiros, Jan Leike, Jane Park,	Shadwell, Tianhao Zheng, Todd Underwood, Todor	733
670	Jason Kwon, Jason Phang, Jason Teplitz, Jason	Markov, Toki Sherbakov, Tom Rubin, Tom Stasi,	734
671	Wei, Jason Wolfe, Jay Chen, Jeff Harris, Jenia Var-	Tomer Kaftan, Tristan Heywood, Troy Peterson, Tyce	735
672	avva, Jessica Gan Lee, Jessica Shieh, Ji Lin, Jiahui	Walters, Tyna Eloundou, Valerie Qi, Veit Moeller,	736
673	Yu, Jiayi Weng, Jie Tang, Jieqi Yu, Joanne Jang,	Vinnie Monaco, Vishal Kuo, Vlad Fomenko, Wayne	737
674	Joaquin Quinero Candela, Joe Beutler, Joe Lan-	Chang, Weiyi Zheng, Wenda Zhou, Wesam Manassra,	738
675	ders, Joel Parish, Johannes Heidecke, John Schul-	Will Sheu, Wojciech Zaremba, Yash Patil, Yilei Qian,	739
676	man, Jonathan Lachman, Jonathan McKay, Jonathan	Yongjik Kim, Youlong Cheng, Yu Zhang, Yuchen	740
677	Uesato, Jonathan Ward, Jong Wook Kim, Joost	He, Yuchen Zhang, Yujia Jin, Yunxing Dai, and Yuri	741
678	Huizinga, Jordan Sitkin, Jos Kraaijeveld, Josh Gross,	Malkov. 2024. Gpt-4o system card .	742
679	Josh Kaplan, Josh Snyder, Joshua Achiam, Joy Jiao,		
680	Joyce Lee, Juntang Zhuang, Justyn Harriman, Kai	Phuc Phan, Hieu Tran, and Long Phan. 2024. Distilla-	743
681	Fricke, Kai Hayashi, Karan Singhal, Katy Shi, Kavin	tion contrastive decoding: Improving llms reasoning	744
682	Karthik, Kayla Wood, Kendra Rimbach, Kenny Hsu,	with contrastive decoding and distillation .	745
683	Kenny Nguyen, Keren Gu-Lemberg, Kevin Button,		
684	Kevin Liu, Kiel Howe, Krithika Muthukumar, Kyle	Alec Radford, Jong Wook Kim, Chris Hallacy,	746
685	Luther, Lama Ahmad, Larry Kai, Lauren Itow, Lau-	A. Ramesh, Gabriel Goh, Sandhini Agarwal, Girish	747
686	ren Workman, Leher Pathak, Leo Chen, Li Jing, Lia	Sastry, Amanda Askell, Pamela Mishkin, Jack Clark,	748
687	Guy, Liam Fedus, Liang Zhou, Lien Mamitsuka, Lil-	Gretchen Krueger, and Ilya Sutskever. 2021a. Learn-	749
688	ian Weng, Lindsay McCallum, Lindsey Held, Long	ing transferable visual models from natural language	750
689	Ouyang, Louis Feuvrier, Lu Zhang, Lukas Kon-	supervision. In <i>ICML</i> .	751
690	draciuk, Lukasz Kaiser, Luke Hewitt, Luke Metz,		
691	Lyric Doshi, Mada Aflak, Maddie Simens, Madelaine	Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya	752
692	Boyd, Madeleine Thompson, Marat Dukhan, Mark	Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sas-	753
693	Chen, Mark Gray, Mark Hudnall, Marvin Zhang,	try, Amanda Askell, Pamela Mishkin, Jack Clark,	754
694	Marwan Aljubei, Mateusz Litwin, Matthew Zeng,	Gretchen Krueger, and Ilya Sutskever. 2021b. Learn-	755
695	Max Johnson, Maya Shetty, Mayank Gupta, Meghan	ing transferable visual models from natural language	756
696	Shah, Mehmet Yatbaz, Meng Jia Yang, Mengchao	supervision .	757
697	Zhong, Mia Glaese, Mianna Chen, Michael Jan-		
698	ner, Michael Lampe, Michael Petrov, Michael Wu,	Pritam Sarkar, Sayna Ebrahimi, Ali Etemad, Ahmad	758
699	Michele Wang, Michelle Fradin, Michelle Pokrass,	Beirami, Sercan Ö. Arik, and Tomas Pfister. 2024.	759
700	Miguel Castro, Miguel Oom Temudo de Castro,	Data-augmented phrase-level alignment for mitigat-	760
701	Mikhail Pavlov, Miles Brundage, Miles Wang, Mi-	ing object hallucination .	761
702	nal Khan, Mira Murati, Mo Bavarian, Molly Lin,		
703	Murat Yesildal, Nacho Soto, Natalia Gimelshein, Na-	Christoph Schuhmann, Romain Beaumont, Richard	762
704	talie Cone, Natalie Staudacher, Natalie Summers,	Vencu, Cade W Gordon, Ross Wightman, Mehdi	763
705	Natan LaFontaine, Neil Chowdhury, Nick Ryder,	Cherti, Theo Coombes, Aarush Katta, Clayton	764
706	Nick Stathas, Nick Turley, Nik Tezak, Niko Felix,	Mullis, Mitchell Wortsman, Patrick Schramowski,	765
707	Nithanth Kudige, Nitish Keskar, Noah Deutsch, Noel	Srivatsa R Kundurthy, Katherine Crowson, Lud-	766
708	Bundick, Nora Puckett, Ofir Nachum, Ola Okelola,	wig Schmidt, Robert Kaczmarczyk, and Jenia Jitsev.	767
709	Oleg Boiko, Oleg Murk, Oliver Jaffe, Olivia Watkins,	2022. LAION-5b: An open large-scale dataset for	768
710	Olivier Godement, Owen Campbell-Moore, Patrick	training next generation image-text models . In <i>Thirty-</i>	769

sixth Conference on Neural Information Processing
Systems Datasets and Benchmarks Track.

Rico Sennrich, Jannis Vamvas, and Alireza Mohammadshahi. 2024. [Mitigating hallucinations and off-target machine translation with source-contrastive and language-contrastive decoding](#).

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
```

We used *laion/CLIP-ViT-L-14-DataComp.XL-s13B-b90K* as the open-clip model for both image and text encoding throughout this work.

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.

A.4 Examples of Image Pairs and Their Descriptions from D3

The following pairs of images with their respective captions are sampled from D3 results, obtained using CoCI with Natural Contrast Images and Qwen 7B.



Baseline: The image shows a person holding a banana in front of their mouth. The person is wearing a green shirt and is standing against a plain wall. The banana is positioned in such a way that it appears to be covering the person's mouth, creating a humorous or playful effect. The background is simple and uncluttered, which helps to focus attention on the person and the banana.

CoCI: The image shows a person holding a banana close to their face. The person appears to be making an exaggerated facial expression, with their eyes wide open and eyebrows raised. The background is a plain, light-colored wall.



Baseline: The image shows a person holding a banana in front of their mouth, creating the illusion of a smile. The person is wearing a green shirt and is standing against a plain wall. The banana is positioned in such a way that it covers the lower part of the person's face, mimicking a smile. The lighting in the image is even, and the background is simple and uncluttered.

CoCI: The image depicts a man with a light complexion wearing a green shirt. He is leaning against a wall, which appears to be a combination of light and dark shades, possibly indicating a corner or a wall with different sections. The man is holding a banana in his right hand, and he has placed the banana horizontally across his mouth, creating a playful expression that resembles a mustache. The banana has some dark spots on it, suggesting it is ripe. The background is relatively plain, with no visible decorations or objects, focusing the viewer's attention on the man and the banana.

Figure 3: Comparison of baseline and CoCI captions. The baseline captions are very similar and provide no information that allows to discriminate between the images, while the left CoCI caption mentions the raised eyebrows.



Baseline: The image depicts the iconic Elizabeth Tower, commonly known as Big Ben, located in London, England. The tower is a prominent example of Gothic Revival architecture and is part of the Palace of Westminster. The photograph captures the upper portion of the tower, focusing on its clock face and the ornate details of its structure.

CoCI: The image showcases a close-up view of a prominent clock tower, which is characterized by its Gothic architectural style. **The tower is adorned with intricate details and features two large, circular clock faces, one on each visible side of the structure.** Both clocks are encased in ornate frames, which are embellished with decorative patterns and motifs.



Baseline: The image depicts the iconic Elizabeth Tower, commonly known as Big Ben, located in London, England. The tower is a prominent example of Gothic Revival architecture and is part of the Palace of Westminster. The clock face is visible in the center of the tower, displaying the time. The clock is encased in a decorative frame with intricate designs, and the hands of the clock are clearly visible. The tower is topped with a spire and a weathervane, which adds to its historical and architectural significance. The sky in the background is clear, providing a sharp contrast to the detailed stonework of the tower.

CoCI: The image depicts the iconic clock tower known as Big Ben, located in London, England. The tower is part of the Palace of Westminster and is one of the most recognizable landmarks in the world.

Figure 4: Comparison of baseline and CoCI captions. The left CoCI caption mentions the two visible clock faces.



Baseline: The image depicts a person engaged in windsurfing, a water sport that combines elements of surfing and sailing. The individual is standing on a windsurfing board, which is partially submerged in the water. ... The individual appears to be in a standing position, holding onto the sail with both hands, indicating that they are actively maneuvering the board.
CoCI: The image depicts a person engaging in the water sport of windsurfing near a beach. The individual is wearing a black wetsuit and appears to be holding onto the sail of the windsurfing board ... *The person appears to be standing in shallow water, with their upper body and arms visible above the waterline.*



Baseline: The image depicts a person engaged in windsurfing, an activity that combines elements of surfing and sailing. The windsurfer is standing on a surfboard, which is partially submerged in the water ... The individual's posture suggests they are actively maneuvering the sail, likely using their body weight and the wind to control the direction and speed of the board.
CoCI: The image depicts a windsurfer riding a wave in the ocean. The windsurfer is captured from the side, *showing their body leaning back as they navigate the wave* ... The windsurfer is wearing dark clothing and appears to be in motion, with the water splashing around their feet as they ride the wave. ...

Figure 5: Comparison of baseline and CoCI captions. The CoCI captions describe the state of the windsurfer: standing in the shallow water vs. riding the wave.