

MRFD: Multi-Region Fusion Decoding with Self-Consistency for Mitigating Hallucinations in LVLMs

Haonan Ge^λ Yiwei Wang^{λ*} Ming-Hsuan Yang^λ Yujun Cai[†]

^λ Department of Computer Science and Engineering, University of California at Merced

[†]The University of Queensland

gehaonan82@gmail.com wangyw.evan@gmail.com

Abstract

Large Vision-Language Models (LVLMs) have shown strong performance across multimodal tasks. However, they often produce hallucinations—text that is inconsistent with visual input, due to the limited ability to verify information in different regions of the image. To address this, we propose **Multi-Region Fusion Decoding (MRFD)**, a training-free decoding method that improves factual grounding by modeling inter-region consistency. MRFD identifies salient regions using cross-attention, generates initial responses for each, and computes reliability weights based on Jensen-Shannon Divergence (JSD) among the responses. These weights guide a consistency-aware fusion of per-region predictions, using region-aware prompts inspired by Chain-of-Thought reasoning. Experiments across multiple LVLMs and benchmarks show that MRFD significantly reduces hallucinations and improves response factuality without requiring model updates.

1 Introduction

Large Vision-Language Models (LVLMs) have emerged as powerful tools for multimodal understanding, achieving significant advances in image captioning, visual question answering, and visual reasoning (Alayrac et al., 2022; Liu et al., 2023b; Li et al., 2023a; Zhu et al., 2023; Xia et al., 2024; Li et al., 2025b,a). However, these models frequently generate hallucinated content, producing textual outputs inconsistent with the visual input. This phenomenon poses substantial challenges for reliability and practical deployment (Ji et al., 2023; Liu et al., 2024). Hallucinations typically manifest as misidentified objects, fabricated attributes, or omitted visual information, reducing the factual alignment between image and text.

* Corresponding author.

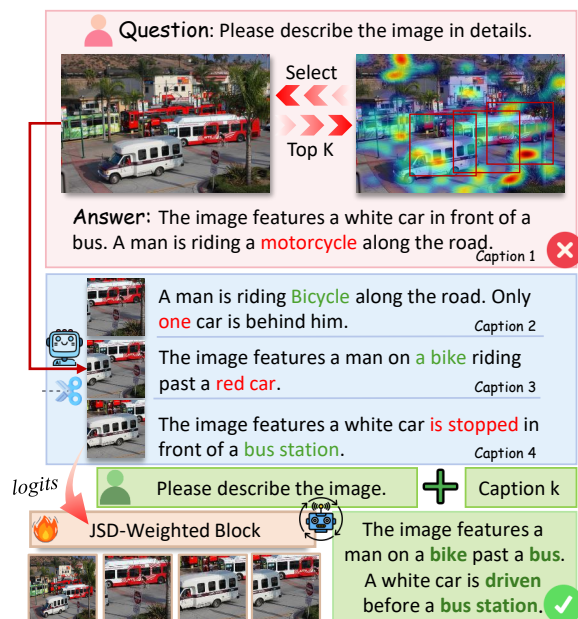


Figure 1: The MRFD process: leveraging multiple regional responses (Captions 1-4), a JSD-Weighted Block derives consistency weights to guide a prompted fusion decoding, yielding a more reliable output.

This issue is especially prominent in scenarios that require fine-grained understanding or interpretation of visually complex scenes. Current approaches include training-based solutions like data augmentation and fine-tuning, which improve factual grounding but require considerable resources and often lack generalization (Rohrbach et al., 2018; Gunjal et al., 2024). Training-free alternatives such as chain-of-thought prompting (Wei et al., 2022) and contrastive or corrective decoding methods (Li et al., 2023b; Leng et al., 2024; Favero et al., 2024; Woo et al., 2024; Huang et al., 2024; Zhang et al., 2025a) offer flexibility, but they often rely on carefully crafted prompts or fail to exploit intrinsic visual information fully.

A key limitation shared by many of these approaches is their tendency to process images holistically or to analyze regions in isolation, often with-

out dynamic mechanisms to assess the reliability of different visual cues or to reconcile potentially conflicting interpretations from multiple perspectives. Consequently, explicit consistency-based reasoning across various image segments is rarely integrated into current decoding procedures (Feng et al., 2024; Zhang et al., 2021).

Our analysis (Section 3) highlights that some image regions provide more trustworthy evidence than others. In particular, we observe that the consistency of region-level responses, which are quantified by Jensen-Shannon Divergence (JSD), correlates with their factual correctness (see Section 3.2). This strong correlation indicates that inter-region agreement is a key determinant of output quality: responses consistent with the consensus view from multiple regions are not only more factually accurate but also demonstrate higher reliability and lower hallucination rates.

To this end, we propose **Multi-Region Fusion Decoding (MRFD)**, a decoding strategy that aims to incorporate multi-perspective reasoning into the generation process. As illustrated in Figure 1, MRFD identifies several salient regions in the image based on attention mechanisms, then generates an initial response for each. By computing Jensen-Shannon Divergence (JSD) among these responses, the method derives a set of consistency-based reliability weights. These weights are used during decoding to combine predictions from all regions. Furthermore, inspired by CoT prompting, we construct region-aware prompts by combining the original question with the initial analysis of each region, preserving the localized context throughout the generation. The contributions of this work are:

- We propose a training-free decoding method that integrates multiple region-level perspectives, weighted by inter-region consistency, to reduce hallucinations in LVLMs.
- We introduce a JSD-based scoring scheme to quantify agreement among region-wise responses and guide reliability-aware fusion.
- We design a region-aware prompting strategy to enhance contextual grounding during generation without modifying model parameters.

2 Related Work

LVLm Hallucinations and Grounding Deficiencies. Large Vision-Language Models (LVLMs) frequently ‘hallucinate’—generating text unsupported

by visual input (Ji et al., 2023; Liu et al., 2024). This often stems from flawed visual grounding due to various factors such as data biases or poor connections between vision and language components (Han et al., 2024; Jiang et al., 2024; Shu et al., 2025; Wang et al., 2023a). Although standard attention mechanisms (Anderson et al., 2018) and region-based analyzers (Li et al., 2022; Kamath et al., 2021) aim to improve grounding, LVLMs still struggle to reliably combine information from multiple, potentially conflicting, image regions or to assess their mutual consistency.

Limitations in Advanced Decoding and Fusion Strategies. Advanced strategies to improve LVLm outputs also exhibit limitations. *Chain-of-Thought* (CoT) prompting (Wei et al., 2022; Zhang et al., 2023; Lyu et al., 2023), while enhancing reasoning, may not ensure its steps visually align with the image and can be sensitive to setup or resource-intensive. Many training-free corrective or contrastive decoding methods (e.g., (Li et al., 2023b; Leng et al., 2024; Favero et al., 2024; Woo et al., 2024; Huang et al., 2024; Zhang et al., 2025a; Wang et al., 2024)) struggle to capture finegrained local image features and risk missing valid information, resulting in poor grounding capabilities. Moreover, when attempting to fuse information from multiple sources—a concept for which consistency has proven beneficial in language modeling tasks (Wang et al., 2023b; Xiong et al., 2023)—LVLMs face distinct challenges with visual regions. Common simplistic aggregation techniques (like averaging) are often insufficient for visual data, where the trustworthiness of different regions can vary dramatically. A central challenge thus remains: developing principled methods to appropriately weigh and fuse evidence from diverse, contextually-understood visual regions based on their consistency.

3 Motivation

3.1 Global Decoding Misses Local Evidence

To explore whether visual grounding capability is related to the hallucinations in LVLMs, we analyze their attention patterns during question answering. Figure 2 shows a case where the model is asked: “Is there a laptop in the image?” When using the full image as input, the model outputs “NO,” even though a laptop is present. The attention map shows that the model distributes focus across unrelated regions.

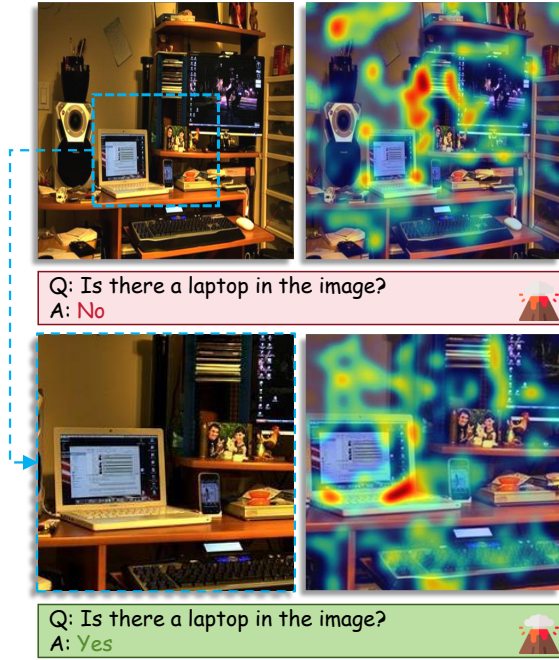


Figure 2: LVLm cross-attention patterns for "Is there a laptop in the image?". (upper) Full image input results in scattered attention and potential error. (lower) Cropped image input focused on the laptop yields concentrated attention and improved accuracy.

In contrast, when the image is cropped to a salient region based on high attention (the desk area), the model concentrates more narrowly on the relevant evidence and correctly answers "YES." The localized input guides the model to verify only what matters for the query, avoiding distractions from the rest of the image.

This comparison reveals a key shortcoming of global decoding: attention is easily diffused across many areas, some of which may mislead the model. In cluttered or ambiguous scenes, this often results in factual errors. Cropped regions help narrow the visual focus of the model, improving the quality of the answer by strengthening the localized foundation.

3.2 Region Consistency Reflects Reliability

However, not all regions—even attention-guided ones—are equally reliable. Some may be visually ambiguous or contextually misleading. To avoid overconfident on a single (possibly misleading) region, we adopt a self-consistency approach: comparing multiple region-level responses. When different views agree on an answer, that answer is more likely to be trustworthy.

To test this, we run experiments using LLaVA-1.5 on 3,000 MSCOCO validation samples, an-

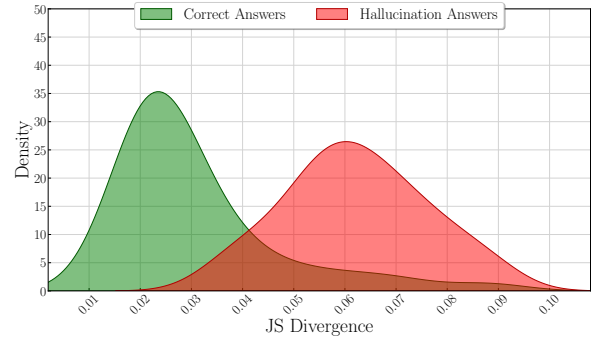


Figure 3: Density distribution of JS Divergence for correct versus hallucinated LVLm responses, indicating lower JSD correlates with higher factual accuracy.

notated with hallucination labels from the POPE benchmark. For each image-question pair, we generate several responses from attention-guided image patches, along with one from the full image. For the output of each cropped region, we calculate the Jensen-Shannon divergence (JSD) between its response distribution and the average response distribution across all regions.

Our results (illustrated in Figure 3) show a clear distinction: JS Divergence for hallucinated answers tends to concentrate around 0.06-0.07, whereas for correct answers, it centers around approximately 0.02. This indicates that the JSD between a specific regional response and the average across regions effectively reflects the level of hallucination of that regional response: a lower JSD correlates with fewer hallucinations. This finding resonates with the principle of self-consistency (Wang et al., 2023b), where agreement among multiple diverse outputs often signals higher quality and reliability. Motivated by this, we design a decoding strategy: A mechanism that aggregates evidence from multiple regions, and checks their agreement, can help improve the robustness and factual grounding of model predictions.

4 Method

4.1 Overview

Current LVLm decoding methods often suffer from hallucinations due to the lack of multi-perspective consistency checking, as discussed in Section 3.

To address this, we propose **Multi-Region Fusion Decoding (MRFD)**, a training-free decoding strategy that enhances answer reliability through: 1) Attention-Guided Region Selection: selecting salient regions based on attention maps; 2) Multi-Region Analysis with JSD-Based Weighting: as-

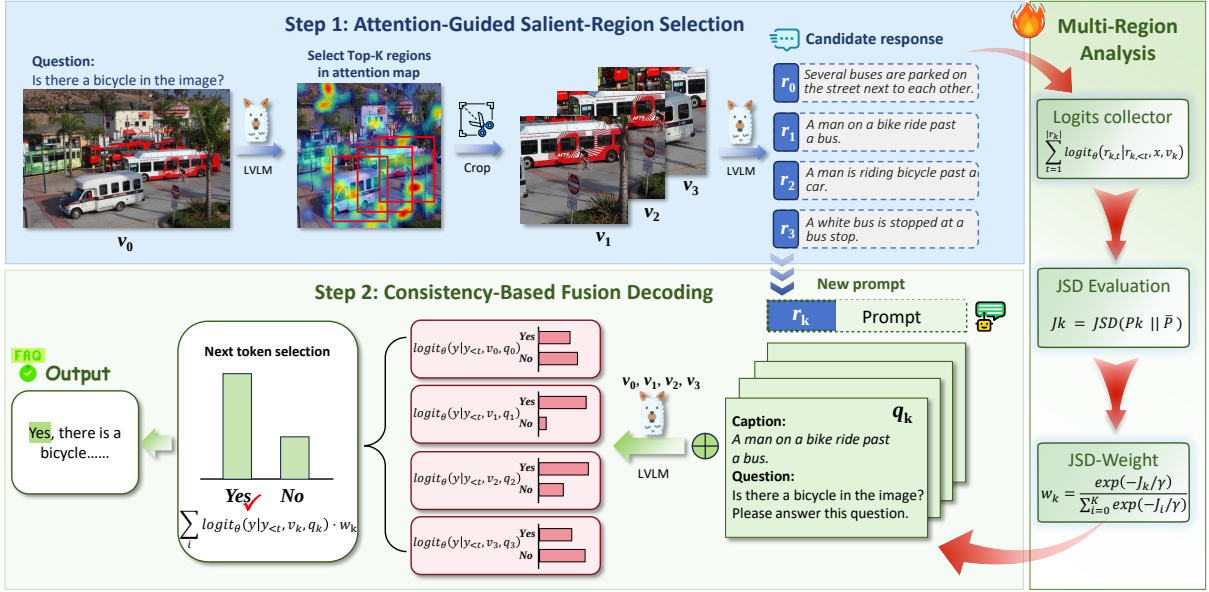


Figure 4: Overall framework of Multi-Region Fusion Decoding (MRFD): Step 1 uses attention to select and crop salient regions (v_k), generates candidate responses (r_k) per region, and computes JSD-based consistency weights (w_k) for each response. Step 2 forms new inputs per region with a candidate response and the original prompt. They are all processed in parallel, fusing per-region logits using the weights w_k during parallel decoding to select the output tokens.

sessing consistency across regions using Jensen-Shannon Divergence; 3) Consistency-Based Fusion Decoding: fusing predictions weighted by their mutual consistency. The full framework is illustrated in Figure 4.

4.2 Attention-Guided Region Selection

Building on the multi-perspective verification approach motivated in Section 3, the first crucial step in our MRFD’s framework is to identify and isolate multiple informative regions from the input image for focused and independent analysis.

Leveraging Cross-Attention for Region Identification. To identify image regions most relevant to the input query, we leverage the cross-attention mechanisms already present in LVLM architectures. Given an original image I_0 and a text query x , we compute the attention matrix A from the LVLM’s cross-attention layers. Let $H_T \in \mathbb{R}^{n \times d}$ represent the sequence of hidden states from the textual input, and $H_V \in \mathbb{R}^{m \times d}$ represent those from the visual input, where n and m denote sequence lengths and d represents the hidden dimensionality.

The cross-attention weight matrix $A \in \mathbb{R}^{n \times m}$, capturing alignment between textual and visual modalities, is computed using scaled dot-product

attention (Vaswani et al., 2017):

$$A = \text{softmax} \left(\frac{(H_T W^Q)(H_V W^K)^\top}{\sqrt{d_k}} \right) \quad (1)$$

where $W^Q \in \mathbb{R}^{d \times d_k}$ and $W^K \in \mathbb{R}^{d \times d_k}$ are projection matrices for queries and keys, and d_k is their dimensionality.

To obtain an overall map of visual focus, we aggregate these attention weights across the textual dimension:

$$a_j = \frac{1}{n} \sum_{i=1}^n A_{ij} \quad \text{for } j = 1, \dots, m \quad (2)$$

This produces an aggregated attention vector $a \in \mathbb{R}^{1 \times m}$. Assuming the visual hidden states correspond to a spatial grid of patches ($d' \times d' = m$), we reshape this vector into a 2D spatial attention map:

$$\hat{A} = \text{reshape}(a, (d', d')) \quad (3)$$

This spatial attention map $\hat{A}$ visually highlights the regions that the model finds most relevant to the input query. By utilizing the model’s own attention patterns, we can directly identify which parts of the image are being primarily considered when answering the query.

Selecting Salient Regions. Using the spatial attention map $\hat{A}$, we identify the top-K most salient

regions $R_1, R_2, \dots, R_K$ within the original image. To ensure diversity and comprehensive coverage, we employ an integral image approach to efficiently search for K non-overlapping or minimally overlapping regions that maximize summed attention scores. For practical implementation, we constrain these regions to be rectangular bounding boxes with a maximum overlap (IoU) threshold between any two regions to ensure diversity in the selected perspectives.

For each selected region R_k , we crop the original image to obtain a focused sub-image v_k . This process yields a set of $K+1$ visual inputs: the original full image $v_0 = I_0$ and K cropped region images $\{v_1, v_2, \dots, v_K\}$. These cropped regions represent the most informative parts of the image with respect to the input query, as determined by the model’s own attention mechanism.

4.3 Multi-Region Analysis with JSD-Based Weighting

Having selected salient regions, we now detail their independent analysis and our JSD-based weighting mechanism. This approach is crucial for assessing inter-region consistency (Section 3.2). For each selected region R_k , we process the corresponding cropped sub-image v_k through the LVLM to generate an initial analysis:

$$r_k = \text{LVLM}(v_k, x) \quad (4)$$

where r_k represents the LVLM’s response when considering only the visual information from region k . We also generate a response r_0 using the original full image v_0 . The resulting set $\{r_0, r_1, \dots, r_K\}$ provides multiple perspectives on the query, each grounded in different salient regions of the image.

To assess the reliability of information from each region, we measure the consistency between region-specific analyses using Jensen-Shannon Divergence (JSD). The JSD provides a symmetric measure of similarity between probability distributions and is defined as:

$$JSD(P||Q) = \frac{1}{2}D_{KL}(P||M) + \frac{1}{2}D_{KL}(Q||M) \quad (5)$$

where P and Q are two probability distributions, $M = \frac{1}{2}(P + Q)$, and D_{KL} is the Kullback-Leibler divergence.

We derive a representative vocabulary distribution P_k for each initial analysis sequence r_k by

averaging the next-token probability distributions computed during generation:

$$P_k = \text{softmax}\left(\frac{1}{|r_k|} \sum_{t=1}^{|r_k|} \text{logit}_\theta(y|r_{k,<t}, x, v_k)\right) \quad (6)$$

where $|r_k|$ is the length of r_k and logit_θ represents the model’s output logits. This distribution P_k captures the overall token prediction patterns when the model considers region k . We then compute the average distribution across all regions:

$$\bar{P} = \frac{1}{K+1} \sum_{i=0}^K P_i \quad (7)$$

For each region k , we calculate its divergence from this average distribution:

$$J_k = JSD(P_k||\bar{P}) \quad (8)$$

These divergence scores measure how much each region’s predictions deviate from the consensus. A lower J_k indicates that region k provides information more consistent with other regions, suggesting higher reliability.

Finally, we convert these divergence scores into weights with temperature parameter γ :

$$w_k = \frac{\exp(-J_k/\gamma)}{\sum_{i=0}^K \exp(-J_i/\gamma)} \quad (9)$$

The temperature parameter γ controls the sharpness of the weight distribution—a smaller γ creates more contrast between weights, while a larger γ leads to more uniform weighting.

4.4 JS-Weighted Integrative Fusion Decoding

With the consistency-based weights w_k computed for each visual input v_k ($k = 0 \dots K$), we now perform the final decoding step by fusing the next-token predictions in a manner that prioritizes consistent visual evidence.

Motivated by the effectiveness of Chain-of-Thought (CoT) prompting (Wei et al., 2022) in structuring reasoning, we adapt its core principle to enhance factual grounding within our multi-region fusion framework. For each region k , we construct a fixed region-aware prompt q_k by concatenating the original question x and the corresponding region’s initial analysis r_k :

$$q_k = \text{Concat}(x, r_k) \quad (10)$$

This prompt q_k encapsulates both the global query and the localized preliminary analysis derived from v_k . It serves as a static, enriched context for region k throughout the decoding process, providing region-specific grounding information without requiring dynamic prompt updates.

During autoregressive decoding, at each step t' , the LVLM produces next-token logits $\ell_k^{(t')}$ conditioned on the current partial output sequence $y_{<t'}$, the visual input v_k , and the fixed region-aware prompt q_k :

$$\ell_k^{(t')} = \text{logit}_\theta(y|y_{<t'}, v_k, q_k) \quad (11)$$

These logits represent the prediction from each region’s perspective given its pre-computed analysis. We aggregate them using the pre-computed consistency weights w_k :

$$\ell_{\text{fused}}^{(t')} = \sum_{k=0}^K w_k \cdot \ell_k^{(t')} \quad (12)$$

The final probability distribution for the next token is obtained via softmax:

$$P_{\text{fused}}^{(t')}(y) = \text{softmax}(\ell_{\text{fused}}^{(t')}) \quad (13)$$

The next token $\hat{y}_{t'}$ is then selected (e.g., sampling) from this fused distribution, and the process repeats autoregressively until completion.

5 Experiments

5.1 Experimental Settings

Evaluated LVLMs. We evaluate Multi-Region Fusion Decoding (MRFD) on two representative open-source LVLMs: LLaVA-1.5-7B (Liu et al., 2023a) and InstructBLIP-7B (Dai et al., 2023). These models adopt different vision-language interfaces—direct projection in LLaVA-1.5 and query-based encoding (Q-Former (Li et al., 2023a)) in InstructBLIP—enabling evaluation of MRFD’s generality. MRFD is applied as a training-free, decoding-time procedure on frozen models. Architectural details are provided in Appendix A.

Baselines. As a basic baseline, regular decoding samples tokens from the model’s post-softmax output probabilities. These include contrastive methods such as VCD (Leng et al., 2024) and M3ID (Favero et al., 2024), robustness-enhancing methods like RITUAL (Woo et al., 2024), and approaches employing correction or feedback, such as DeGF (Zhang et al., 2025b) and Woodpecker (Yin

et al., 2023). We also include comparisons with other relevant methods: HALC (Chen et al., 2024), and OPERA (Huang et al., 2024). Performance of these baselines is based on our re-implementations using publicly available code where possible. Detailed descriptions of each baseline’s methodology are provided in Appendix B.

Implementation Details. Across all experiments, our Multi-Region Fusion Decoding (MRFD) method selects $K = 3$ salient regions and employs a temperature of $\gamma = 0.02$ for the JSD-based weighting (Eq. 9). We utilize multinomial sampling for both stages of decoding. Other detailed settings are provided in Appendix C.

5.2 Datasets and Benchmarks

We evaluate MRFD on multiple benchmarks covering both hallucination detection and general vision-language understanding. Below we briefly describe the key datasets; full details and evaluation metrics are provided in Appendix D.

POPE (Li et al., 2023c): A Yes/No QA benchmark for object existence hallucination, built from MSCOCO, A-OKVQA, and GQA with various negative sampling strategies.

CHAIR (Rohrbach et al., 2018): Measures hallucinated object mentions in image captions via CHAIRi and CHAIRs scores on MSCOCO images.

MME (Fu et al., 2023): We focus on the MME-Hallucination subset, which tests object, count, position, and attribute hallucinations.

5.3 Results and Discussions

Results on POPE. In Table 1, we compare the performance of our MRFD method against other baselines on the POPE benchmark (Li et al., 2023c) under three different negative sampling settings (Random, Popular, Adversarial), across both LLaVA-1.5 and InstructBLIP. As shown, MRFD consistently outperforms other decoding methods on both LVLMs, achieving leading F1 scores across all six configurations, with improvements of up to 2.44% in accuracy, 6.52% in precision, and 3.05% in F1 score compared to the respective second-best approaches. This suggests that MRFD’s core strategy—identifying multiple salient regions via cross-attention, assessing their response consistency using JSD-weighting, and fusing their context-enriched predictions—enables LVLMs to better ground responses in relevant visual evidence, thereby effectively addressing object hallucinations.

Setting	Method	LLaVA-1.5 (Liu et al., 2023a)			InstructBLIP (Dai et al., 2023)		
		Acc. ↑	Prec. ↑	F1 ↑	Acc. ↑	Prec. ↑	F1 ↑
Random	Regular	82.42	78.30	83.67	79.85	80.33	83.45
	VCD (Leng et al., 2024)	84.69	80.30	85.85	84.47	83.07	84.81
	M3ID (Favero et al., 2024)	85.46	81.54	86.42	85.32	83.51	85.58
	RITUAL (Woo et al., 2024)	86.71	82.84	87.51	87.12	85.64	<u>87.23</u>
	DeGF (Zhang et al., 2025b)	<u>87.79</u>	<u>86.33</u>	88.08	<u>87.21</u>	89.01	86.70
	MRFD	88.15	88.91	88.23	88.03	<u>88.74</u>	88.01
Popular	Regular	76.57	71.23	79.56	75.17	70.90	77.54
	VCD (Leng et al., 2024)	77.30	71.61	80.57	78.12	73.77	80.10
	M3ID (Favero et al., 2024)	78.66	73.09	81.45	78.32	73.75	80.30
	RITUAL (Woo et al., 2024)	79.75	74.55	82.31	78.40	73.63	80.55
	DeGF (Zhang et al., 2025b)	81.94	78.33	83.31	80.37	78.84	81.24
	MRFD	81.99	78.81	83.29	82.17	79.52	83.15
Adversarial	Regular	71.09	65.77	75.93	71.02	66.58	74.70
	VCD (Leng et al., 2024)	71.13	65.28	76.37	73.07	68.50	76.36
	M3ID (Favero et al., 2024)	72.10	66.27	76.87	72.97	67.92	76.55
	RITUAL (Woo et al., 2024)	71.87	66.17	76.88	73.06	67.83	76.91
	DeGF (Zhang et al., 2025b)	<u>76.13</u>	<u>71.50</u>	<u>79.01</u>	<u>75.96</u>	<u>73.45</u>	<u>77.36</u>
	MRFD	77.99	76.16	79.22	77.72	74.26	79.72

Table 1: Results on POPE benchmark. Higher (↑) accuracy, precision, and F1 indicate better performance. The best results are bolded, and the second-best are underlined.

Method	LLaVA-1.5		InstructBLIP	
	Cs ↓	Ci ↓	Cs ↓	Ci ↓
Regular	26.2	9.4	31.2	11.1
VCD	24.4	7.9	30.0	10.1
M3ID	21.4	6.3	30.8	10.4
RITUAL	22.4	6.9	26.6	8.9
Woodpecker	24.9	7.5	31.2	10.8
HALC	21.7	7.1	24.5	8.0
DeGF	<u>18.4</u>	<u>6.1</u>	<u>24.0</u>	<u>7.7</u>
Ours (MRFD)	14.1	5.0	21.3	6.1

Table 2: Results on CHAIR benchmark for caption generation. We limited the maximum number of new tokens to 64. Lower (↓) CHAIRs (Cs) and CHAIRi (Ci) indicate less hallucination. Best results are bolded, second-best are underlined.

Moreover, while most decoding method tend to be overconfident in nonexistence, the consistency verification inherent in MRFD appears to promote more cautious and precise responses, which is evidenced by its strong precision, particularly in challenging adversarial settings (e.g., 76.16 on LLaVA-1.5 and 74.26 on InstructBLIP), highlighting its enhanced capability in filtering false positives and suppressing misinformation. Detailed results of POPE are attached in Appendix E.3.

Results on CHAIR. We evaluate MRFD’s effectiveness in mitigating object hallucination in open-

ended image captioning using the CHAIR benchmark (Rohrbach et al., 2018), reporting CHAIRs (Cs) and CHAIRi (Ci) scores (lower is better) for LLaVA-1.5 and InstructBLIP in Table 2. MRFD consistently achieves state-of-the-art performance on both LVLMS, significantly outperforming the strong DeGF baseline with relative CHAIR score reductions of up to 16.4% on LLaVA-1.5 and 20.8% on InstructBLIP. Given that the CHAIR task for image captioning necessitates comprehensive attention to diverse local details, MRFD’s superior performance underscores its advanced capability for robust multi-region analysis and effective integration of key information from various visual segments, leading to more factually grounded image captions. Detailed results of CHAIR are appended to Appendix E.1. Some qualitative examples are shown in Appendix H.

Results on MME-Hallucination. We evaluate MRFD on the MME hallucination subset (Fu et al., 2023), which assesses diverse hallucination types including object-level (Existence), attribute-level (Count, Color), and relation-level (Position) understanding in LVLMS. As illustrated in Figure 5 (presenting results for both LLaVA-1.5 and InstructBLIP), our MRFD method consistently demonstrates strong performance, achieving leading or state-of-the-art results across the majority of these hallucination categories for both evaluated LVLMS

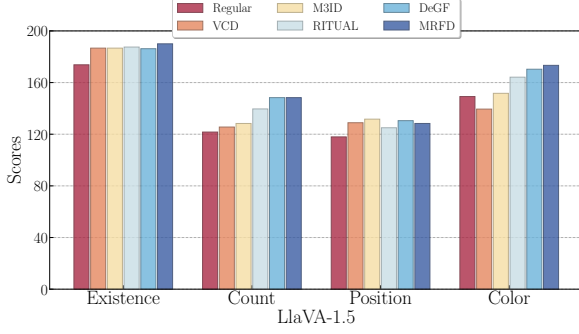


Figure 5: Experimental results of MME on a hallucination subset with different decoding strategies.

when compared to regular decoding and other advanced baselines. This robust performance across a spectrum of challenging hallucination types underscores that MRFD’s core mechanism—multi-region analysis coupled with consistency-driven fusion—effectively enhances fine-grained visual understanding and overall factual consistency. Detailed results of MME-Hallucination are appended to Appendix E.2

Efficiency Discussion. Despite employing multiple inference steps, MRFD maintains a competitive efficiency profile, particularly when compared to other advanced hallucination mitigation techniques that involve more complex iterative or feedback mechanisms. Details are in Appendix F.

5.4 Ablation Study

We conduct ablation studies to evaluate the impact of key components and hyperparameters in MRFD, using LLaVA-1.5 on the POPE-MSCOCO dataset.

First, we assess the contributions of MRFD’s core designs: JSD-based Consistency Weighting (CW), Fusion Prompt (FP), and attention-guided Region Selection (RS). We test our full MRFD against three main variants where these components are individually altered: (1) **MRFD w/o CW**, applying uniform fusion weights; (2) **MRFD w/o FP**, using only the original question for regional decoding; and (3) **MRFD w/o RS**, processing only the global image through the subsequent pipeline stages. As detailed in Table 3, full MRFD achieves an 86.21 F1 score. Removing CW degrades F1 performance by 2.94%, underscoring the importance of dynamic, consistency-based weighting. Omitting FP reduces F1 by 3.93%, highlighting the value of enriched regional context. Bypassing RS results in the largest F1 drop of 4.07%, emphasizing that robust multi-region analysis is funda-

Model Variants	Acc. $\uparrow$	Prec. $\uparrow$	F1 Score $\uparrow$
MRFD (Full)	86.50	88.11	86.21
w/o CW	83.76	84.27	83.74
w/o FP	82.87	83.95	82.70
w/o RS	82.77	83.94	82.58
Regular	80.57	78.84	81.59

Table 3: Ablation study with different model variants on POPE-COCO under the average of three settings .

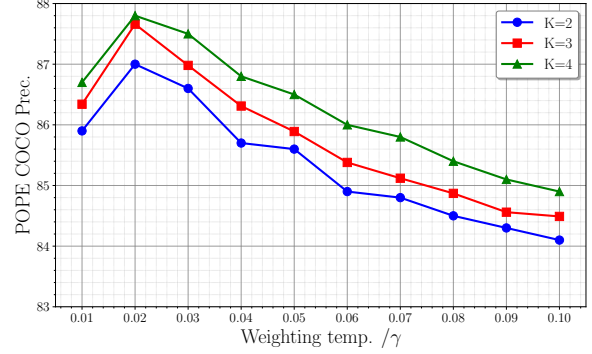


Figure 6: Sensitivity of POPE-COCO Precision (LLaVA-1.5) to JSD weighting temperature γ for $K = 2, 3, 4$ regions, with optimal performance around $\gamma = 0.02$.

mental. Despite these impacts, all three ablated MRFD variants still outperform Regular decoding (81.59 F1), while the complete MRFD configuration showcases the strongest synergistic benefits.

Second, we analyze the sensitivity to the JSD weighting temperature γ (Eq. 9). As depicted in Figure 6 (POPE COCO Precision for $K = 2, 3, 4$), performance peaks around $\gamma = 0.02$ ($K = 3$ at $\approx 87.8\%$ Prec.). Higher γ values (≥ 0.08) lead to more uniform weights and correspondingly reduced precision, approaching unweighted fusion. Conversely, very low γ values (≈ 0.01) cause over-reliance on a few regions due to extreme weighting, which can negatively impact robustness, particularly for larger K (e.g., performance degradation for $K = 4$). Thus, $\gamma = 0.02$ is adopted for providing an optimal balance. Further details are presented in Appendix G.

6 Conclusion

We presented Multi-Region Fusion Decoding (MRFD), a training-free approach that improves LLM reliability by mimicking self-consistency multi-view verification. MRFD identifies salient regions via attention, estimates their reliability using Jensen-Shannon Divergence (JSD), and

fuses region-level predictions using consistency-weighted, Chain-of-Thought-inspired prompts. Experiments show that MRFD effectively reduces hallucinations and improves factuality across various LVLMs and benchmarks, offering a simple yet robust decoding strategy.

Limitations

MRFD relies on the quality and interpretability of attention maps produced by the underlying LVLM, which may vary across architectures and input conditions. The assumption that inter-region consistency indicates factual correctness holds empirically but may not generalize to all reasoning tasks, especially those requiring abstract or commonsense inference. Additionally, while MRFD improves factuality in image-grounded tasks, it has not been evaluated in broader multimodal contexts such as video or dialogue-based grounding, which may involve more complex temporal or conversational dependencies.

Ethics Statement

Our work on Multi-Region Fusion Decoding (MRFD) aims to improve the reliability and factual grounding of LVLMs, thereby contributing to more trustworthy AI systems by reducing visual hallucinations. We believe this has positive implications for applications where factual accuracy is critical.

Acknowledgements

The work is partially supported by the NSF of the United States Grant CRII 2451683, an NVIDIA Academic Grants Program, University of California at Merced, and a UC Merced Faculty Research Award. The views and conclusions are those of the authors and should not reflect the official policy or position of the U.S. Government.

We also thank Prof. Kai-Wei Chang (Department of Computer Science, UCLA) for his insightful feedback on earlier drafts, including suggestions on clarifying terminology, adding appropriate cognitive-science references, discussing experimental settings, and presenting qualitative examples, as well as for pointing us to relevant related work. Any remaining errors are our own.

References

- Jean-Baptiste Alayrac, Jeff Donahue, Pauline Luc, Antoine Miech, Iain Barr, Yana Hasson, Karel Lenc, Arthur Mensch, Katherine Millican, Malcolm Reynolds, and 1 others. 2022. [Flamingo: a visual language model for few-shot learning](#). In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 35, pages 23716–23737.
- Peter Anderson, Xiaodong He, Chris Buehler, Damien Teney, Mark Johnson, Stephen Gould, and Lei Zhang. 2018. [Bottom-up and top-down attention for image captioning and visual question answering](#). In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 6077–6086.
- Zhaorun Chen, Zhuokai Zhao, Hongyin Luo, Huaxiu Yao, Bo Li, and Jiawei Zhou. 2024. [Halc: Object hallucination reduction via adaptive focal-contrast decoding](#). *Preprint*, arXiv:2403.00425.
- Wei-Lin Chiang, Zhuohan Li, Zi Lin, Ying Sheng, Zhanghao Wu, Hao Zhang, Lianmin Zheng, Siyuan Zhuang, Yonghao Zhuang, Joseph E. Gonzalez, Ion Stoica, and Eric P. Xing. 2023. [Vicuna: An open-source chatbot impressing gpt-4 with 90%* chatgpt quality](#). *Preprint*, arXiv:2303.13617.
- Wenliang Dai, Junnan Li, Dongxu Li, Anthony Meng Huat Tiong, Junqi Zhao, Weisheng Wang, Boyang Li, Pascale Fung, and Steven Hoi. 2023. [Instructblip: Towards general-purpose vision-language models with instruction tuning](#). *Preprint*, arXiv:2305.06500.
- Alessandro Favero, Luca Zancato, Matthew Trager, Siddharth Choudhary, Pramuditha Perera, Alessandro Achille, Ashwin Swaminathan, and Stefano Soatto. 2024. [Multi-modal hallucination control by visual information grounding](#). In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, Seattle, WA, USA. IEEE / CVF.
- Mingqian Feng, Yunlong Tang, Zeliang Zhang, and Chenliang Xu. 2024. [Do more details always introduce more hallucinations in lvlm-based image captioning?](#) *arXiv preprint arXiv:2406.12663*.
- Chaoyou Fu, Peixian Chen, Yunhang Shen, Yulei Qin, Mengdan Zhang, Xu Chen, Qing Predicted, Yu Qiao, and Ping Luo. 2023. [Mme: A comprehensive evaluation benchmark for multimodal large language models](#). *Preprint*, arXiv:2306.13394. MME-Hallucination evaluation is part of this benchmark.
- Sanyam Gunjal, Ben Malin, Florian Golemo, Kiriakos N. Kutulakos, and Aaron Courville. 2024. [Mitigating hallucinations in vision-language models via fine-grained vision-code data augmentation](#). In *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision (WACV)*, pages 3306–3316.

- Tianrui Han, Ruisheng Zhang, and Siyu Li. 2024. [Analyzing and mitigating object hallucination in large vision-language models from perspective of training data bias](#). *arXiv preprint arXiv:2402.16011*.
- Qidong Huang, Xiaoyi Dong, Pan Zhang, Bin Wang, Conghui He, Jiaqi Wang, Dahua Lin, Weiming Zhang, and Nenghai Yu. 2024. [OPERA: Alleviating hallucination in multi-modal large language models via over-trust penalty and retrospection-allocation](#). In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, Seattle, WA, USA. IEEE / CVF.
- Drew A Hudson and Christopher D Manning. 2019. [GQA: A new dataset for real-world visual reasoning and compositional question answering](#). In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 6700–6709, Long Beach, CA, USA. IEEE / CVF.
- Gabriel Ilharco, Mitchell Wortsman, Ross Wightman, Cade Gordon, Nicholas Carlini, Rohan Taori, Achal Dave, Vaishaal Shankar, Hongseok Namkoong, John Miller, Hannaneh Hajishirzi, Ali Farhadi, and Ludwig Schmidt. 2021. [Openclip](#). Version 2.20.0.
- Ziwei Ji, Nayeon Lee, Rita Frieske, Tiezheng Yu, Dan Su, Yan Xu, Etsuko Ishii, Yejin Jenny Bang, Andrea Madotto, and Pascale Fung. 2023. [Survey of hallucination in natural language generation](#). *ACM Computing Surveys*, 55(12):1–38.
- Yifan Jiang, Zhengyuan Zhu, Tianjun Zhao, and Dazhen She. 2024. [Revealing and mitigating hallucinations in large vision-language models](#). *arXiv preprint arXiv:2401.04071*.
- Aishwarya Kamath, Mannat Singh, Yann LeCun, Gabriel Synnaeve, Ishan Misra, and Nicolas Carion. 2021. [Mdetr – modulated detection for end-to-end multi-modal understanding](#). *Preprint*, arXiv:2104.12763.
- Seongyun Lee, Sue Hyun Park, Yongrae Jo, and Minjoon Seo. 2024. [VOLCANO: Mitigating multimodal hallucination through self-feedback guided revision](#). In *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL)*, Mexico City, Mexico. Association for Computational Linguistics.
- Yifan Leng, Yifan Li, Jiashi Zhang, Zidong Liu, Fang Wen, and Kaixiong Chen. 2024. [Mitigating object hallucinations in large vision-language models via visual contrastive decoding](#). In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, Seattle, WA, USA. IEEE / CVF.
- Junnan Li, Dongxu Li, Silvio Savarese, and Steven Hoi. 2023a. [BLIP-2: Bootstrapping language-image pre-training with frozen image encoders and large language models](#). In *Proceedings of the 40th International Conference on Machine Learning (ICML)*, volume 202 of *Proceedings of Machine Learning Research*, pages 19730–19748. PMLR.
- Liunian Harold Li, Pengchuan Li, Zhe Zhang, Kuaida Lee, Peng Li, Yelong Zhang, Kathy Lee, Steven Hoi, Jenq-Neng Wang, Linjie Wang, Luke Zettlemoyer, and Ching Yun Liu. 2022. [Grounded language-image pre-training](#). In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 10965–10975.
- Xiang Li, Lin Li, Zihang Geng, Percy Liu, and Jian Liu. 2023b. [Contrastive decoding: Open-ended text generation as optimization](#). In *International Conference on Learning Representations (ICLR)*.
- Yanshu Li, Yi Cao, Hongyang He, Qisen Cheng, Xiang Fu, Xi Xiao, Tianyang Wang, and Ruixiang Tang. 2025a. [M²iv: Towards efficient and fine-grained multimodal in-context learning via representation engineering](#). *Preprint*, arXiv:2504.04633.
- Yanshu Li, Tian Yun, Jianjiang Yang, Pinyuan Feng, Jinfa Huang, and Ruixiang Tang. 2025b. [Taco: Enhancing multimodal in-context learning via task mapping-guided sequence configuration](#). *arXiv preprint arXiv:2505.17098*.
- Yifan Li, Yifan Du, Kun Wang, Yilin Zhao, Jiashi Zhang, Zidong Liu, Kaixiong Chen, and Fang Wen. 2023c. [Evaluating object hallucination in large vision-language models](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 11221–11233, Singapore. Association for Computational Linguistics.
- Tsung-Yi Lin, Michael Maire, Serge Belongie, James Hays, Pietro Perona, Deva Ramanan, Piotr Dollár, and C Lawrence Zitnick. 2014. [Microsoft COCO: Common objects in context](#). In *European conference on computer vision (ECCV)*, pages 740–755. Springer.
- Hanchao Liu, Wenyuan Xue, Yifei Chen, Dapeng Chen, Xiutian Zhao, Ke Wang, Liping Hou, Rongjun Li, and Wei Peng. 2024. [A survey on hallucination in large vision-language models](#). *Preprint*, arXiv:2402.00253.
- Haotian Liu, Chunyuan Li, Yuheng Li, and Yong Jae Lee. 2023a. [Improved baselines with visual instruction tuning](#). *Preprint*, arXiv:2310.03744.
- Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. 2023b. [Visual instruction tuning](#). In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 36, pages 68178–68198.
- Yuanhan Liu, Haochen Wang, Jingwen Li, Yanwei Fu, Chunyuan Li, and Jianfeng Gao. 2023c. [MMBench: Is your multi-modal model an all-around player?](#) *Preprint*, arXiv:2307.06281.
- Qing Lyu, Shreya Chen, Johnny H GCL, Shichang Zhao, and Abhinav Shrivastava. 2023. [Faithful chain-of-thought reasoning](#). *arXiv preprint arXiv:2301.13379*.

- Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, Gretchen Krueger, and Ilya Sutskever. 2021. [Learning transferable visual models from natural language supervision](#). *Preprint*, arXiv:2103.00020.
- Anna Rohrbach, Marcus Rohrbach, Rong Hu, Trevor Darrell, and Bernt Schiele. 2018. [Object hallucination in image captioning](#). In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 4243–4248, Brussels, Belgium. Association for Computational Linguistics.
- Dustin Schwenk, M. R. Magnusson, Borun Pan, Marco Zompero, Ferda Ofli, Michael Feldman, Russel Gabard, Florian Metze, Adam Trischler, and Thomas Testole. 2022. [A-OKVQA: A benchmark for visual question answering using world knowledge](#). In *Proceedings of the European Conference on Computer Vision (ECCV)*, pages 541–558, Tel Aviv, Israel. Springer Nature Switzerland.
- Dong Shu, Haiyan Zhao, Jingyu Hu, Weiru Liu, Ali Payani, Lu Cheng, and Mengnan Du. 2025. [Large vision-language model alignment and misalignment: A survey through the lens of explainability](#). *Preprint*, arXiv:2501.01346.
- Quan Sun, Yuxin Fang, Ledell Wu, Xinlong Wang, and Yue Cao. 2023. [Eva-clip: Improved training techniques for clip at scale](#). *Preprint*, arXiv:2303.15389.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. [Attention is all you need](#). In *Advances in Neural Information Processing Systems (NeurIPS)*. Curran Associates, Inc.
- P. Viola and M. Jones. 2001. [Rapid object detection using a boosted cascade of simple features](#). In *Proceedings of the 2001 IEEE Computer Society Conference on Computer Vision and Pattern Recognition. CVPR 2001*, volume 1, pages I–I.
- Junyang Wang, Yiyang Zhou, Guohai Xu, Pengcheng Shi, Chenlin Zhao, Haiyang Xu, Qinghao Ye, Ming Yan, Ji Zhang, Jihua Zhu, Jitao Sang, and Haoyu Tang. 2023a. [Evaluation and analysis of hallucination in large vision-language models](#). *Preprint*, arXiv:2308.15126.
- Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc Le, Ed Chi, Sharan Narang, Aakanksha Chowdhery, and Denny Zhou. 2023b. [Self-consistency improves chain of thought reasoning in language models](#). In *International Conference on Learning Representations (ICLR)*.
- Zhaobin Wang, Xinyi Li, and Jing Chen. 2024. [Mitigating hallucination in lvlms via disturbance instruction](#). *arXiv preprint arXiv:2401.12876*.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, and Denny Zhou. 2022. [Chain-of-thought prompting elicits reasoning in large language models](#). In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 35, pages 24824–24837. Curran Associates, Inc.
- Sangmin Woo, Jaehyuk Jang, Donguk Kim, Yubin Choi, and Changick Kim. 2024. [Ritual: Random image transformations as a universal anti-hallucination lever in large vision language models](#). *Preprint*, arXiv:2405.17821.
- Haotian Xia, Zhengbang Yang, Yun Zhao, Yuqing Wang, Jingxi Li, Rhys Tracy, Zhuangdi Zhu, Yuanfang Wang, Hanjie Chen, and Weining Shen. 2024. [Language and multimodal models in sports: A survey of datasets and applications](#). *Preprint*, arXiv:2406.12252.
- Chen Xiong, Zhi Liu, Le Wang, and Lifan Liu. 2023. [Can llms evaluate their own accuracy? a study on math problem solving](#). *arXiv preprint arXiv:2305.07331*.
- Shukang Yin, Chaoyou Fu, Sirui Zhao, Tong Chen, Enhong Chen, Gaoang Wang, Qiao Yu, and Ping Luo. 2023. [Woodpecker: Hallucination correction for multimodal large language models](#). *Preprint*, arXiv:2310.16045.
- Ce Zhang, Zifu Wan, Zhehan Kan, Martin Q. Ma, Simon Stepputtis, Deva Ramanan, Russ Salakhutdinov, Louis-Philippe Morency, Katia Sycara, and Yaqi Xie. 2025a. [Self-correcting decoding with generative feedback for mitigating hallucinations in large vision-language models](#). *Preprint*, arXiv:2502.06130.
- Ce Zhang, Zifu Wan, Zhehan Kan, Martin Q. Ma, Simon Stepputtis, Deva Ramanan, Russ Salakhutdinov, Louis-Philippe Morency, Katia P. Sycara, and Yaqi Xie. 2025b. [Self-correcting decoding with generative feedback for mitigating hallucinations in large vision-language models](#). In *The Thirteenth International Conference on Learning Representations (ICLR)*.
- Wenqiao Zhang, Haochen Shi, Siliang Tang, Jun Xiao, Qiang Yu, and Yueting Zhuang. 2021. [Consensus graph representation learning for better grounded image captioning](#). In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing (EMNLP)*.
- Zhuosheng Zhang, Jing Liu, Zheyuan Luo, Shane Storks, Zhe Li, Binhang Wang, Chen Chen, Iryna Gurevych, Honglak Lee, Graham Liu, and Aston Zhang. 2023. [Multimodal chain-of-thought reasoning in language models](#). *arXiv preprint arXiv:2302.00923*.
- Deyao Zhu, Jun Chen, Xiaoqian Shen, Xiang Li, and Mohamed Elhoseiny. 2023. [MiniGPT-4: Enhancing vision-language understanding with advanced large language models](#). *arXiv preprint arXiv:2304.10592*.

A Evaluated LVLMs

We evaluate our proposed Multi-Region Fusion Decoding (MRFD) framework on three representative open-source Large Vision-Language Models (VLMs): **LLaVA-1.5-7B** (Liu et al., 2023a), and **InstructBLIP-7B** (Dai et al., 2023). For the visual encoder, LLaVA-1.5 uses ViT-L-336px pre-trained from CLIP-L/14-336px (Radford et al., 2021). InstructBLIP uses ViT-g/14 pre-trained from EVA-CLIP (Sun et al., 2023). Qwen-VL uses Openclip ViT-bigG (Ilharco et al., 2021). Regarding the language model module, LLaVA-1.5 and InstructBLIP both utilize Vicuna-7B (Chiang et al., 2023).

The vision-language interface varies across the evaluated models. LLaVA-1.5 directly projects visual features using MLP layers. InstructBLIP, conversely, adopts the Q-Former (Li et al., 2023a), which processes visual features via 32 learnable queries to generate a standardized input for the LLM. Our evaluation of MRFD encompasses these distinct strategies—direct MLP projection and query-based bottleneck (Q-Former)—to demonstrate its broad effectiveness in mitigating hallucinations regardless of the connection module employed.

B Baselines

We compare the performance of our MRFD method with several recent training-free decoding approaches designed to mitigate hallucination. These key baselines, along with their core methodologies, are described below:

- **VCD** (Leng et al., 2024): This method contrasts output distributions derived from an original visual input (v) and a distorted version (v'). Given a textual query and v , the model generates two distributions. The distorted input v' is created using pre-defined distortions (e.g., a Gaussian noise mask) to v . These two distributions are then contrasted to form the final output probability. For reproduction, we follow VCD’s default setting with $\alpha = 1$ (a parameter controlling the contrast strength) and use 500 noise steps to generate v' . VCD aims to enhance visual grounding by reducing reliance on language priors.
- **M3ID** (Favero et al., 2024): M3ID contrasts output distributions from original visual inputs against those from pure text inputs (lacking visual information). The final probability

distribution is a combination of the distribution conditioned on both vision and text, and a contrastive term derived from the difference between vision-conditioned and text-only conditioned distributions, balanced by a hyperparameter λ . We adhere to their recommended $\lambda = 0.02$. M3ID also seeks to improve visual grounding by emphasizing visual information.

- **RITUAL** (Woo et al., 2024): RITUAL applies common image transformations (e.g., crop, flip, color jitter) to the original visual input v , creating a transformed version $v^{(T)}$. It then generates the response by utilizing information from both the original (v) and transformed ($v^{(T)}$) images. The final probability distribution combines logits from both views, with the contribution of the transformed input adjusted by a balancing hyperparameter κ . We follow their official implementation, setting $\kappa = 3$. This approach aims to improve robustness through consistency across augmented views.
- **DeGF** (Zhang et al., 2025b): DeGF introduces a self-correction mechanism using feedback from text-to-image generative models. Specifically, it first generates an image based on the LVLM’s initial textual response. This generated image then acts as an auxiliary visual reference, providing self-feedback to the LVLM to verify and correct its initial response, often through complementary or contrastive decoding techniques.
- **Woodpecker** (Yin et al., 2023): Woodpecker is a post-hoc correction framework designed to mitigate hallucinations in the outputs of Multimodal Large Language Models (MLLMs). It operates by first prompting the MLLM itself to identify potential hallucinations (across several predefined types like object existence, attributes, etc.) in its initial response. If hallucinations are detected, Woodpecker then instructs the MLLM to revise and correct these identified errors.
- **HALC** (Chen et al., 2024): HALC (e.g., "Mitigating Object Hallucinations in Large Vision-Language Models via Cause Analysis and Post-hoc Correction") is a post-hoc method that first analyzes the potential causes of object hallucinations to identify objects in the re-

sponse that are likely to be hallucinated. Subsequently, it instructs the LVLM to verify the existence of these specific, suspect objects within the image and make corrections if they are indeed confirmed as hallucinations.

- **OPERA** (Huang et al., 2024): OPERA aims to alleviate visual relation hallucinations in LVLMs during decoding. It consists of two main components: an Over-trust Penalty (OP) term designed to penalize the model’s over-confidence on unreliable visual relations during token generation, and a Retrospection-Allocation (RA) mechanism that encourages the model to retrospect previously generated tokens and re-allocate attention to relevant visual regions for verification and potential correction.

We report the performance of these baselines based on our re-implementation using their released code bases where available.

C Implementation Details

In all experiments using MRFD, we set the number of regions $K = 3$ to obtain cropped images and the temperature $\gamma = 0.02$ for JSD-based weighting (Eq. 9). For the decoding process, we employ multinomial sampling in both stages. Specifically, during the first step to generate the initial analyses r_k , we use a sampling temperature of 0.7. In the second step to generate the final output sequence y , we use a lower sampling temperature of 0.1 after fusing the logits.

To obtain the aggregated spatial attention map $\hat{A}$ (Eq. 3), we first average the attention weights across all attention heads within the relevant layer(s). For **LLaVA-1.5**, we utilize the attention map from the final cross-attention layer. For **InstructBLIP**, which uses Q-Former, we identify the query token with the maximum aggregated attention score and use its corresponding attention map. The spatial dimension d' of the map $\hat{A}$ corresponds to the grid size of the visual patches ($m = d' \times d'$), which is 24×24 for LLaVA-1.5 (ViT-L-336px) and 16×16 for the ViT-g/14 used by InstructBLIP.

To efficiently identify the top-K salient regions $\{R_1, \dots, R_K\}$ based on $\hat{A}$, we employ an integral image approach (Viola and Jones, 2001) to quickly calculate the sum of attention scores within any rectangular bounding box. We search for the K non-overlapping or minimally overlapping rectan-

gular regions that maximize these summed attention scores. To ensure diversity in the selected regions, we enforce a maximum Intersection over Union (IoU) of 40% between any pair of selected bounding boxes R_i and R_j ($i \neq j$).

D Datasets and Benchmarks

We evaluate our MRFD framework on a diverse set of benchmarks targeting both hallucination detection and general vision-language capabilities.

- **POPE** (Li et al., 2023c): POPE (Polling-based Object Probing Evaluation) is a widely used benchmark for assessing object existence hallucination in LVLMs. It presents models with Yes/No questions concerning the presence of specific objects (e.g., "Is there a {object} in the image?"). The benchmark data is structured into three main subsets derived from MSCOCO (Lin et al., 2014), A-OKVQA (Schwenk et al., 2022), and GQA (Hudson and Manning, 2019). Each of these subsets is further divided based on three negative sampling strategies for non-existent objects: *random*, *popular*, and *adversarial*, which vary in difficulty. For evaluation, we report standard metrics including Accuracy, Precision, Recall, and F1 score.
- **MME** (Fu et al., 2023): MME serves as a comprehensive benchmark for evaluating overall LVLM perception and cognition. Our evaluation specifically utilizes the **MME-Hallucination** subset, which is designed to assess a range of common hallucination types. These are categorized into object-level assessments like *existence* and *count*, and attribute-level assessments such as object *position* and *color*. Questions in this subset are typically Yes/No queries. We report scores based on the official benchmark protocol, which often involves combined accuracy measures reflecting both question-level and image-level correctness.
- **CHAIR** (Rohrbach et al., 2018): The CHAIR (Caption Hallucination Assessment with Image Relevance) benchmark quantifies object hallucinations within the context of open-ended image captioning. LVLMs are prompted to generate descriptive captions for images, for which we, following prior work (Lee et al., 2024), use a random selection

of 500 images from the MSCOCO (Lin et al., 2014) validation set. The generated captions are then compared against ground-truth objects within the image to calculate the CHAIRi (instance-level) and CHAIRs (category-level) scores, where lower scores indicate fewer hallucinations.

$$\text{CHAIR}_S = \frac{\text{captions w/ hallucinated objects}}{\text{all captions}},$$

$$\text{CHAIR}_I = \frac{\text{hallucinated objects}}{\text{all mentioned objects}}.$$
(14)

- **MMBench (Liu et al., 2023c):** MMBench evaluates a broad spectrum of multimodal capabilities through carefully curated multiple-choice questions that span various cognitive dimensions and skills. For this benchmark, we adhere to the official evaluation protocol and report the overall accuracy score.

E Detailed Results of experiments

E.1 Detailed Results of CHAIR

Detailed results of CHAIR are shown in Table 4 and Table 5, reporting CHAIRs (Cs) and CHAIRi (Ci) scores (lower is better) for LLaVA-1.5 and InstructBLIP. MRFD consistently achieves state-of-the-art performance on both LVLMs.

Method	LLaVA-1.5		InstructBLIP	
	Cs ↓	Ci ↓	Cs ↓	Ci ↓
Regular	55.0	16.3	57.0	17.6
VCD	54.4	16.6	60.4	17.8
M3ID	56.6	15.7	62.2	18.1
RITUAL	49.6	14.8	<u>48.4</u>	14.5
Woodpecker	57.6	16.7	60.8	17.6
HALC	51.0	14.8	53.8	15.7
DeGF	<u>48.8</u>	<u>14.6</u>	49.2	<u>14.4</u>
Ours (MRFD)	37.1	9.2	38.2	10.9

Table 4: Results on CHAIR benchmark for caption generation. We limited the maximum number of new tokens to 128. Lower (↓) CHAIRs (Cs) and CHAIRi (Ci) indicate less hallucination. Best results are bolded, second-best are underlined.

E.2 Detailed Results of MME

In table 7 and Figure 7, we provide detailed results on the MME-Hallucination benchmark (Fu et al.,

Method	LLaVA-1.5		InstructBLIP	
	Cs ↓	Ci ↓	Cs ↓	Ci ↓
Regular	58.0	17.7	61.0	18.2
VCD	58.2	16.7	63.0	18.6
M3ID	56.8	16.1	65.8	19.9
RITUAL	51.0	15.1	50.4	15.3
DeGF	<u>49.8</u>	<u>14.7</u>	<u>49.8</u>	<u>15.1</u>
Ours (MRFD)	39.0	11.0	38.6	11.3

Table 5: Results on CHAIR benchmark for caption generation. We limited the maximum number of new tokens to 256. Lower (↓) CHAIRs (Cs) and CHAIRi (Ci) indicate less hallucination. Best results are bolded, second-best are underlined.

2023) for both LLaVA-1.5 and InstructBLIP. The table includes scores for object-level (existence, count) and attribute-level (position, color) tasks, and averages the score across three random seeds. The best results are bolded, and the second-best are underlined.

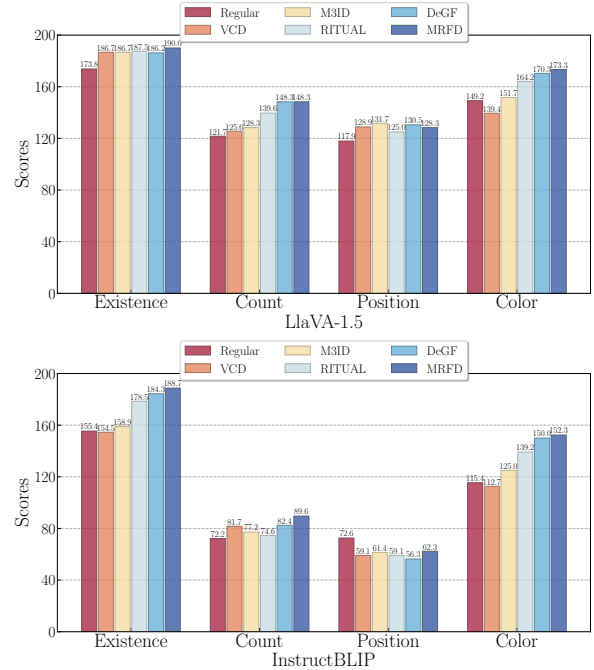


Figure 7: Experimental results of MME with different decoding strategies.

E.3 Detailed Results of POPE

In table 6, we provide detailed results on the POPE benchmark (Li et al., 2023c) for both LLaVA-1.5 and InstructBLIP across three different settings (Random, Popular, Adversarial). The table includes accuracy (Acc.), precision (Prec.), recall,

and F1 scores for each method. The best results are bolded, and the second-best are underlined.

F Efficiency Comparison

We report MRFD’s efficiency on the CHAIR benchmark with two backbones—InstructBLIP-7B and LLaVA-1.5-7B—under the same decoding setup (128 max tokens, RTX 3090 GPU). With multi-region ($K=3$) analysis and fusion, MRFD increases latency by about $\sim 2.84\times$ (from 3.59 s to 10.21 s) and peak GPU memory by $\sim 1.11\times$ (from 15,238 MB to 16,932 MB) relative to Regular decoding on InstructBLIP-7B (Table 8). On LLaVA-1.5-7B, MRFD shows a $\sim 3.30\times$ latency increase (3.44 s $\rightarrow$ 11.34 s) and a $\sim 1.11\times$ memory increase (15,778 MB $\rightarrow$ 17,458 MB) over Regular (Table 9). Despite this overhead—and compared to lighter baselines such as VCD or the post-hoc Woodpecker—MRFD remains markedly more efficient than iterative/feedback-heavy methods (DeGF, OPERA, HALC), while achieving the best CHAIRs (lower is better): 38.2 on InstructBLIP-7B and 37.1 on LLaVA-1.5-7B, representing $\sim 33\%$ and $\sim 32.5\%$ reductions from their respective Regular baselines. Overall, MRFD offers a compelling balance between computational cost and hallucination reduction.

Method	Avg. Lat. ↓	GPU Mem. ↓	CHAIRs ↓
Regular	3.59 s	15,238 MB	57.0
VCD	6.82 s	16,078 MB	60.4
OPERA	25.31 s	23,575 MB	62.2
Woodpecker	11.72 s	22,798 MB	60.8
HALC	21.84 s	22,689 MB	53.8
DeGF	14.69 s	18,726 MB	49.2
Ours (MRFD)	10.21 s	16,932 MB	38.2

Table 8: Efficiency comparison on **InstructBLIP-7B**. Lower is better for all metrics.

Method	Avg. Lat. ↓	GPU Mem. ↓	CHAIRs ↓
Regular	3.44 s	15,778 MB	55.0
VCD	6.91 s	16,634 MB	54.4
OPERA	24.70 s	22,706 MB	52.6
Woodpecker	10.68 s	22,199 MB	57.6
HALC	22.61 s	23,084 MB	51.0
DeGF	13.89 s	19,119 MB	48.8
Ours (MRFD)	11.34 s	17,458 MB	37.1

Table 9: Efficiency comparison on **LLaVA-1.5-7B**. Lower is better for all metrics.

G Detailed Ablation Study

G.1 Components Ablation

To assess the contributions of MRFD’s core components—attention-guided Region Selection (RS), JSD-based Consistency Weighting (CW), and Fusion Prompts (FP)—we conduct ablation studies. Results on POPE-COCO (average all settings, LLaVA-1.5) are in Table 10. The evaluated model variants are:

- (a) **MRFD (Full)** : Our complete method, integrating all three components.
- (b) **w/o CW** : Employs RS and FP, but uses uniform weights for fusing predictions from multiple regions, bypassing JSD-based consistency weighting.
- (c) **w/o FP** : Utilizes RS and CW, but omits the enriched Fusion Prompts, using only the original question for each regional decoding pass.
- (d) **w/o RS (Global Image + FP)** : Bypasses attention-guided region selection, operating solely on the global image. The Fusion Prompt is constructed based on the global image’s initial response. In this single-view context, the multi-region Consistency Weighting (CW) mechanism as defined is not applicable or becomes trivial (effectively a weight of 1 for the single view).
- (e) **RS only** : Leverages attention-guided regions, but with uniform fusion weights and only the original question as prompt, isolating the benefit of the multi-region perspective itself.
- (f) **FP only (Global Image)** : Applies the Fusion Prompt (derived from the global image’s initial response) directly to the global image decoding, without multi-region selection or any form of consistency weighting.
- (g) **Global Image (Token-Level Voting)** : Operates on the global image with the original prompt. The CW component is adapted to perform token-level voting by aggregating implicitly diversified predictions from the single global view, testing self-consistency benefits at the token level.
- (h) **Regular** : Standard greedy decoding using only the global image and the original question.

	Components			Acc. $\uparrow$	Prec. $\uparrow$	F1 $\uparrow$
	CW	FP	RS			
(a)	✓	✓	✓	86.50	88.11	86.21
(b)		✓	✓	83.76	84.27	83.74
(c)	✓		✓	82.87	83.95	82.70
(d)	✓	✓		82.77	83.94	82.58
(e)			✓	82.58	82.01	82.26
(f)		✓		81.96	81.36	82.14
(g)	✓			81.52	81.05	81.68
(h)				80.57	78.84	81.59

Table 10: Ablation study with different model variants on POPE-COCO under the average of three settings. CW: JSD-based Consistency Weighting, FP: Fusion Prompt, RS: Region Selection.

G.2 Decoding Strategies

To evaluate the impact of different decoding strategies on MRFD, we conduct an ablation study using the POPE-COCO benchmark under the adversarial setting with LLaVA-1.5. The results are summarized in Table 11. We compare three decoding strategies: low temperature sampling, high temperature sampling, and high temperature sampling with top-p filtering. The best results are bolded.

The low temperature strategy is the default setting in our experiments, which t is 0.1. The high temperature strategy is set to $t=0.5$, which increases the randomness of the sampling process. The high temperature + top-p strategy combines high temperature sampling with top-p filtering, where we set $p=0.7$ to retain the top 70% of the probability mass.

Decoding	Acc. $\uparrow$	Prec. $\uparrow$	F1 $\uparrow$
Low temp.	82.75	85.22	82.12
High Temp.	83.04	82.29	83.23
High Temp.+Top P	82.50	81.25	82.67

Table 11: Ablation study on different decoding strategies.

H Qualitative Examples

To provide a more intuitive understanding of MRFD’s impact on mitigating hallucinations and improving caption factuality, this section presents qualitative examples from our experiments on the CHAIR benchmark using LLaVA-1.5. We compare captions generated by our Multi-Region Fusion

Decoding (MRFD) method against those from the standard Regular baseline. These examples are selected to highlight differences in how the methods perceive and describe objects, attributes, and their relationships within the visual scene. Figure 8 and Figure 9 illustrate one such comparison.

Dataset	Setting	Method	LLaVA-1.5				InstructBLIP			
			Acc. $\uparrow$	Prec. $\uparrow$	Recall $\uparrow$	F1 $\uparrow$	Acc. $\uparrow$	Prec. $\uparrow$	Recall $\uparrow$	F1 $\uparrow$
MS-COCO	Random	Regular	83.13	81.94	85.00	83.44	83.07	83.02	83.26	83.08
		VCD	87.00	86.13	<u>88.18</u>	87.15	86.23	88.14	83.73	85.88
		M3ID	87.50	87.38	<u>87.67</u>	87.52	86.67	88.09	<u>84.79</u>	86.41
		RITUAL	88.87	89.23	88.39	<u>88.81</u>	88.83	90.48	86.80	88.60
		DeGF	89.03	91.20	86.41	88.74	88.83	93.73	82.42	87.71
		Ours	89.50	92.55	85.94	89.12	<u>88.52</u>	93.83	82.61	<u>87.86</u>
	Popular	Regular	81.17	78.28	86.26	82.08	77.00	73.82	83.68	78.44
		VCD	83.10	79.96	<u>88.34</u>	83.94	80.07	77.67	84.39	80.89
		M3ID	84.30	81.58	88.63	84.95	80.97	77.93	<u>86.19</u>	81.85
		RITUAL	85.83	84.17	88.26	86.17	81.97	78.90	87.26	<u>82.87</u>
		DeGF	<u>86.63</u>	87.75	84.86	86.28	<u>82.73</u>	84.02	80.27	<u>82.10</u>
		Ours	87.24	<u>86.56</u>	88.22	87.38	83.69	85.22	81.58	83.36
	Adversarial	Regular	77.43	73.31	86.26	79.26	74.60	71.26	82.46	76.45
		VCD	77.17	72.18	88.40	79.47	77.20	74.29	83.19	78.49
		M3ID	78.23	73.51	<u>88.28</u>	80.22	77.47	73.68	<u>85.48</u>	79.14
		RITUAL	78.80	74.43	87.74	80.54	78.73	74.57	87.21	<u>80.39</u>
		DeGF	81.63	80.59	83.33	81.94	80.30	80.90	79.33	80.11
		Ours	82.75	85.22	79.25	82.12	82.49	83.14	81.51	82.32
A-OKVQA	Random	Regular	81.90	76.63	91.78	83.53	80.63	76.82	87.75	81.92
		VCD	83.83	78.05	<u>94.13</u>	85.34	84.20	80.90	89.54	85.00
		M3ID	84.67	79.25	93.94	85.97	85.43	81.77	<u>91.20</u>	86.23
		RITUAL	85.17	79.79	94.21	86.40	87.13	83.92	91.87	<u>87.71</u>
		DeGF	86.93	84.28	90.80	87.42	87.40	87.67	86.86	87.26
		Ours	87.13	87.56	86.76	<u>87.16</u>	88.33	<u>86.81</u>	90.40	88.57
	Popular	Regular	75.07	68.58	92.53	78.77	75.17	70.15	87.60	77.91
		VCD	76.63	69.59	94.59	80.19	78.63	73.53	89.46	80.72
		M3ID	77.80	70.98	94.07	80.91	78.80	73.38	90.39	81.00
		RITUAL	78.83	71.99	<u>94.37</u>	81.68	78.73	72.83	91.68	81.17
		DeGF	80.90	75.68	91.05	82.66	81.47	78.61	86.47	82.35
		Ours	80.99	76.68	89.28	<u>82.51</u>	83.19	78.74	<u>91.00</u>	84.43
	Adversarial	Regular	67.23	61.56	91.81	73.70	69.87	64.54	88.20	74.54
		VCD	67.40	61.39	93.79	74.21	71.00	65.41	89.13	75.45
		M3ID	68.60	62.22	94.74	75.11	70.10	64.28	90.47	75.16
		RITUAL	68.57	62.26	<u>94.27</u>	74.99	70.27	64.15	91.89	75.55
		DeGF	72.70	66.70	<u>90.68</u>	76.86	73.93	69.36	85.70	76.67
		Ours	75.23	71.03	85.39	77.65	75.62	69.40	<u>91.78</u>	79.04
GQA	Random	Regular	82.23	76.32	93.47	84.03	79.67	76.05	86.62	80.99
		VCD	83.23	76.73	<u>95.38</u>	85.05	82.83	80.16	87.26	83.56
		M3ID	84.20	78.00	95.26	85.77	83.07	80.06	88.06	83.87
		RITUAL	86.10	80.30	95.66	87.31	84.87	82.52	<u>88.47</u>	<u>85.39</u>
		DeGF	87.40	83.51	93.20	88.09	85.40	85.64	84.61	85.12
		Ours	87.81	86.62	90.32	88.41	87.24	<u>85.57</u>	89.69	87.58
	Popular	Regular	73.47	66.83	93.20	77.84	73.33	68.72	85.67	76.26
		VCD	72.37	65.27	<u>95.60</u>	77.58	76.13	71.10	88.07	<u>78.68</u>
		M3ID	73.87	66.70	95.35	78.49	75.17	69.94	88.26	78.04
		RITUAL	74.80	67.50	95.66	79.15	74.50	69.17	<u>88.39</u>	77.61
		DeGF	78.10	71.56	93.25	80.98	76.90	73.89	83.20	78.27
		Ours	<u>77.72</u>	73.17	88.24	<u>80.00</u>	79.62	74.61	90.15	81.65
	Adversarial	Regular	68.60	62.43	93.41	74.84	68.60	63.94	85.31	73.10
		VCD	68.83	62.26	<u>95.67</u>	75.43	71.00	65.75	87.66	75.14
		M3ID	68.67	62.16	95.42	75.28	71.17	65.79	<u>88.19</u>	<u>75.36</u>
		RITUAL	68.23	61.75	95.81	75.10	70.17	64.76	88.48	74.78
		DeGF	74.07	67.42	93.14	78.22	73.63	70.08	80.92	75.11
		Ours	76.00	72.22	84.50	<u>77.88</u>	75.06	70.25	87.18	77.80

Table 6: Detailed results on POPE (Li et al., 2023d) benchmark. Higher ($\uparrow$) accuracy, precision, recall, and F1 indicate better performance. The best results are bolded, and the second-best are underlined.

Model	Method	Object-level		Attribute-level		Total Score ↑
		Existence ↑	Count ↑	Position ↑	Color ↑	
LLaVA-1.5	Regular	173.75	121.67	117.92	149.17	562.50
	VCD	186.67	125.56	128.89	139.45	580.56
	M3ID	186.67	128.33	<u>131.67</u>	151.67	598.11
	RITUAL	187.50	<u>139.58</u>	125.00	<u>164.17</u>	<u>616.25</u>
	DeGF	186.22	148.33	130.50	170.33	635.38
	Ours	190.00	148.33	128.33	173.33	640.00
InstructBLIP	Regular	155.42	72.17	72.58	115.43	415.60
	VCD	154.49	81.67	59.11	112.67	407.94
	M3ID	158.89	77.22	61.44	125.00	422.55
	RITUAL	178.50	74.58	59.08	139.17	451.33
	DeGF	<u>184.32</u>	<u>82.44</u>	56.33	<u>150.00</u>	<u>473.09</u>
	Ours	188.67	89.58	<u>62.33</u>	152.33	492.91

Table 7: Detailed performance on the MME-Hallucination benchmark. Scores are reported as mean. Higher scores ↑ indicate better performance. For each model group LLaVA-1.5, InstructBLIP, results for "Ours" are **bolded** if they are the best in that column. Other best results in a column are also **bolded**. Underlined values score typically represent the second-best performing method or a notable baseline. The "Ours" rows are highlighted.

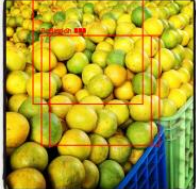
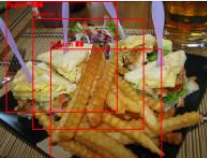
 Question: Provide a short description of the image. Answer: A bunch of green lemons.	Question: Please describe the image in details. Answer: The image features a large assortment of oranges, with a prominent display prominently displayed in an accessible location. Surrounding the crate of fresh oranges, there are various other boxes containing even more lemons of different sizes and colors.
 Question: Provide a short description of the image. Answer: A bunch of green and yellow lemons.	Question: Please describe the image in details. Answer: The image features a large pile of lemons, with many of them being green and yellow. The lemons are arranged in various positions, some stacked on top of each other and others sitting next to each other. The pile is quite large, occupying a significant portion of the image. The lemons are displayed in a way that showcases their vibrant colors and abundance.
 Question: Provide a short description of the image. Answer: A bunch of green and yellow lemons.	Question: Please describe the image in details. Answer: The image showcases two elephants splashing around and playing in a lake together. The first elephant, on the left side of the scene, is partially submerged, displaying its trunk out of the water as it interacts with the other elephant. The second elephant, located on the right side, is in full view and appears to be playfully prodding the first elephant.
 Question: Provide a short description of the image. Answer: A bunch of green and yellow lemons.	Question: Please describe the image in details. Answer: The image features two elephants in a body of water, possibly a lake or a river. One elephant is larger and appears to be an adult, while the other is smaller and likely a baby elephant. The baby elephant is resting its body on the larger elephant, creating a heartwarming scene.

Figure 8: Qualitative comparison of captions generated by Regular decoding (Red) and our MRFD (Green) on an image from the CHAIR evaluation set.

	Question: Provide a short description of the image. Answer: A piece of chicken in a sandwich.	Question: Please describe the image in details. Answer: The image shows a wooden dining table where a delicious combination of food and beverages is served. There are various sandwiches on a plate and a tray of chips, accompanied by cups containing drinks, possibly beers since two beer forks are present. In addition to the main dishes, the table is also set with three forks, two of which are positioned in similar directions and near the sandwiches, while three knives are placed around the plate and tray.
	Question: Provide a short description of the image. Answer: A side of french fries.	Question: Please describe the image in details. Answer: The image features a delicious meal consisting of four sandwich and a side of french fries. The sandwich is cut in half, revealing the meat and vegetables inside. The fries are scattered around the sandwich, with some pieces placed near the sandwich and others further away. In addition to the main dishes, the table is also set with four forks.
	Question: Provide a short description of the image. Answer: Two pieces of sandwich.	Question: Please describe the image in details. Answer: The scene is a public bathroom with a toilet mounted on the wall. A white urinal is attached below a podium, featuring an electronic device likely for display purposes. There is another toilet in the room, situated on the left side with the blanket displaying a picture on it, likely giving the toilet a unique appearance.
	Question: Provide a short description of the image. Answer: Sandwichs with two forks.	Question: Please describe the image in details. Answer: The image features a white toilet in a bathroom. The toilet is situated next to a wall, and there is a silver handle on the wall. The toilet is positioned in the lower left corner of the image. The bathroom appears to be clean and well-maintained.

	Question: Provide a short description of the image. Answer: A toilet in a bathroom.	Question: Please describe the image in details. Answer: The scene is a public bathroom with a toilet mounted on the wall. A white urinal is attached below a podium, featuring an electronic device likely for display purposes. There is another toilet in the room, situated on the left side with the blanket displaying a picture on it, likely giving the toilet a unique appearance.
	Question: Provide a short description of the image. Answer: The toilet seat is white.	Question: Please describe the image in details. Answer: The image features a white toilet in a bathroom. The toilet is situated next to a wall, and there is a silver handle on the wall. The toilet is positioned in the lower left corner of the image. The bathroom appears to be clean and well-maintained.
	Question: Provide a short description of the image. Answer: A silver handle on the wall.	Question: Please describe the image in details. Answer: The image features a white toilet in a bathroom. The toilet is situated next to a wall, and there is a silver handle on the wall. The toilet is positioned in the lower left corner of the image. The bathroom appears to be clean and well-maintained.
	Question: Provide a short description of the image. Answer: The toilet seat is white.	Question: Please describe the image in details. Answer: The image features a white toilet in a bathroom. The toilet is situated next to a wall, and there is a silver handle on the wall. The toilet is positioned in the lower left corner of the image. The bathroom appears to be clean and well-maintained.

Figure 9: Qualitative comparison of captions generated by Regular decoding (Red) and our MRFD (Green) on an image from the CHAIR evaluation set.