

Multi-agent Undercover Gaming: Hallucination Removal via Counterfactual Test for Multimodal Reasoning

Dayong Liang^{1,4*}, Xiao-Yong Wei^{3,2,4*}, Changmeng Zheng^{2†}

¹School of Future Technology, South China University of Technology, Guangzhou, China

²Department of Computing, The Hong Kong Polytechnic University, Hong Kong, China

³College of Computer Science, Sichuan University, Chengdu, China

⁴Department of networked intelligence, Peng Cheng Laboratory, Shenzhen, China
ft_ldy@mail.scut.edu.cn, {x1wei, changmeng.zheng}@polyu.edu.hk

Abstract

Hallucination continues to pose a major obstacle in the reasoning capabilities of large language models (LLMs). Although the Multi-Agent Debate (MAD) paradigm offers a promising solution by promoting consensus among multiple agents to enhance reliability, it relies on the unrealistic assumption that all debaters are rational and reflective, which is a condition that may not hold when agents themselves are prone to hallucinations. To address this gap, we introduce the Multi-agent Undercover Gaming (MUG) protocol, inspired by social deduction games like “Who is Undercover?”. MUG reframes MAD as a process of detecting “undercover” agents (those suffering from hallucinations) by employing multimodal counterfactual tests. Specifically, we modify reference images to introduce counterfactual evidence and observe whether agents can accurately identify these changes, providing ground-truth for identifying hallucinating agents and enabling robust, crowd-powered multimodal reasoning. MUG advances MAD protocols along three key dimensions: (1) enabling factual verification beyond statistical consensus through counterfactual testing; (2) introducing cross-evidence reasoning via dynamically modified evidence sources instead of relying on static inputs; and (3) fostering active reasoning, where agents engage in probing discussions rather than passively answering questions. Collectively, these innovations offer a more reliable and effective framework for multimodal reasoning in LLMs. The source code can be accessed at <https://github.com/YongLD/MUG.git>.

Introduction

Hallucination remains one of the most difficult challenges in LLM reasoning (Huang et al. 2025; Liu et al. 2024b; Peng et al. 2025). This stems from the fact that reasoning in these models is primarily constructed on token-level statistics (such as frequency and causality), often without explicit enforcement of factual correctness. To tackle this problem, the Multi-Agent Debate (MAD) paradigm (Liang et al. 2024) has gained popularity, as it aims to generate a wider range of evidence, with the prevailing consensus among agents

considered more trustworthy. Although MAD has shown promising results, certain limitations persist.

Common MAD protocols are typically derived from human debate structures, such as majority voting, super judge, or blueprint debate. However, there are fundamental differences between LLMs and human debaters. While humans are capable of rational discussion, listening to others, reflecting, and updating their views as needed, it is still uncertain whether LLM debaters can do the same. Consequently, protocols that rely heavily on debater rationality may not be optimal for MAD. Therefore, there is a need for protocols that do not assume rationality and can effectively identify irrational debaters, particularly those susceptible to hallucinations.

We draw inspiration from social deduction games like “Who is Undercover?”, where all participants risk being the undercover and must debate to identify the undercover through factual or counterfactual challenges, often by posing further questions to each other. This closely parallels the task of detecting hallucinating agents. However, there is limited research on how to adapt such protocols for LLMs, especially regarding the design of factual and counterfactual tests when dealing with complex, multimodal evidence. In this paper, we tackle these challenges by introducing the Multi-agent Undercover Gaming (MUG) protocol for MAD. The core idea, illustrated in Figure 1, is straightforward. Given a question “what is the focus of the image?”, we generate multimodal counterfactual tests by altering the reference image to include counterfactual evidence (e.g., changing a girl with red hair to a girl with natural black hair). In this scenario, the main focus of the image shifts from the *hairstyle* to the *phone*. We then observe whether each agent can correctly identify the modification. Because these changes are intentionally introduced, we have “ground-truth” to determine which agents are “undercover” (i.e., ones suffering from hallucinations). This allows us to harness the wisdom of the crowd for more robust and reliable multimodal reasoning.

MUG brings several key innovations to MAD protocols:

- **Factual Verification vs. Statistical Consensus:** By employing counterfactual tests, MUG enables direct factual verification rather than relying solely on group consensus. This addresses the limitation of LLMs, which often lack explicit fact-checking mechanisms during training.

*These authors contributed equally to this work.

†Corresponding author.

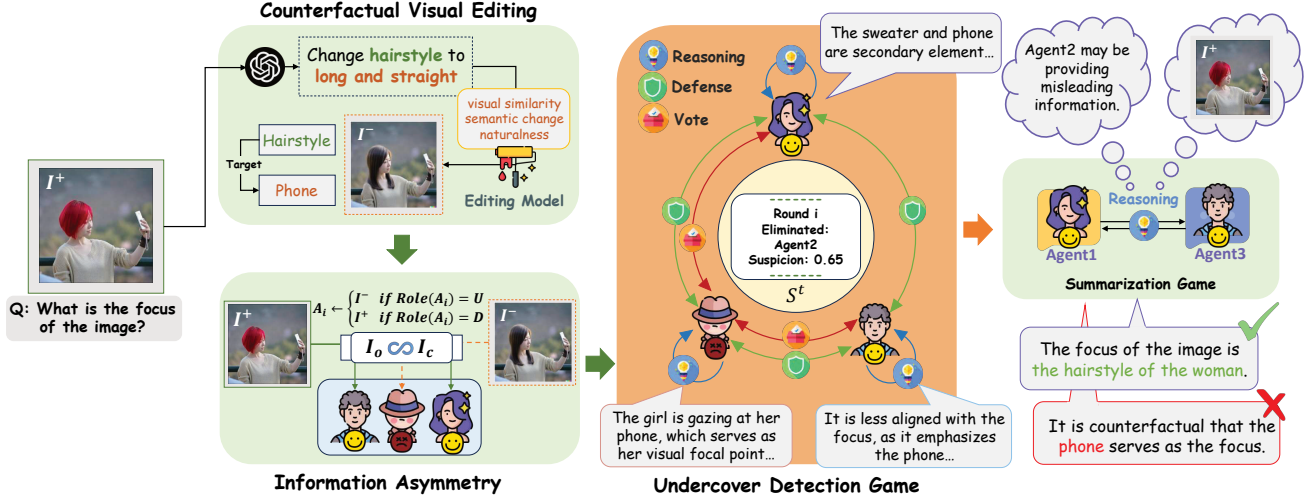


Figure 1: Overview of the Counterfactual Undercover Debate Framework, illustrating the dynamic interactions between agents, the current game state, and the decision-making process influenced by counterfactual information. The diagram highlights the roles of normal and undercover agents, as well as the flow of information and reasoning throughout the game.

- **Cross-Evidence vs. Single-Source:** By modifying images, MUG generates additional, dynamic sources of evidence for cross-referencing, as opposed to traditional MAD approaches that depend on a single, static source (such as the original question and reference image). This not only facilitates verification but also increases the likelihood of detecting agents prone to hallucination.
- **Active vs. Passive Reasoning:** Existing protocols typically involve passive reasoning, where agents simply respond to a given question. In contrast, MUG introduces counterfactual tests that require agents to actively engage by posing questions for discussion. This approach better leverages the capabilities of LLMs and paves the way for developing more effective protocols.

Related Work

Early efforts in multimodal reasoning focused on aligning visual and textual representations, leading to the development of Multimodal LLMs (MLLMs) (Liu et al. 2024a; Yuan et al. 2025) and the adoption of in-context learning (ICL) (Min et al. 2022) and chain-of-thought (CoT) reasoning (Wei et al. 2022). These techniques empower models to decompose complex tasks into interpretable intermediate steps, enhancing transparency and performance. The evolution of CoT into multimodal contexts, termed Multimodal Chain-of-Thought (MCoT) reasoning, has produced a range of architectures, from linear (Wei et al. 2022; Ngo et al. 2008) to graph-based representations (Besta et al. 2024; Yuan et al. 2023), and specialized methods for different modalities, such as Multimodal-CoT (Zhang et al. 2024), MVoT (Li et al. 2025), and Video-of-Thought (Fei et al. 2024). The paradigm of multi-agent systems (MASSs) (Wooldridge 2009) has gained traction as a means to overcome the intrinsic limitations of individual LLMs, such as hallucination and limited reasoning depth. MASS

leverage the collective intelligence of multiple LLM-based agents, enabling distributed knowledge retention (Zheng et al. 2024), long-term planning (Torreno et al. 2017), and specialization through collaborative problem-solving (Lu et al. 2023; Wang et al. 2025). However, existing MAD frameworks often converge prematurely, suppressing minority viewpoints and failing to simulate real-world adversarial reasoning conditions. Counterfactual reasoning, rooted in structural causal models and do-calculus (Pearl 2009; Ngo et al. 2007), is essential for understanding how outcomes change under hypothetical interventions. In AI, counterfactuals enhance interpretability and robustness by enabling models to generate alternative scenarios and assess decision-making under uncertainty (Guidotti 2024; Wang et al. 2024). Recent efforts have extended counterfactual reasoning to LLMs, with approaches ranging from commonsense-based scenario generation (Zhang, Jiang, and Zhao 2024; Chatzi et al. 2025) to graph-based causal inference using external tools, as exemplified by CausalCoT (Jin et al. 2023) and CausalTool (Hua et al. 2024). While agent-based frameworks like Causal Agent (Han et al. 2024) integrate LLMs with causal tools for intervention analysis, they often lack explicit counterfactual reasoning, limiting applicability in complex, high-dimensional settings. Our work distinguishes itself by embedding counterfactual learning directly into the multi-agent collaboration process, enhancing both the robustness and interpretability of multimodal reasoning.

Method

The Gaming System

Multi-agent Undercover Gaming (MUG) takes place within a system $\mathcal{S}$, defined as a tuple of four components: a question Q , a set of N agents $\mathcal{A} = \{A_i\}_1^N$ who act as players or debaters, a collection of debating functions $\mathcal{F}$, and a set of responses $\mathcal{R} = \{R_i\}_1^N$ generated by the agents.

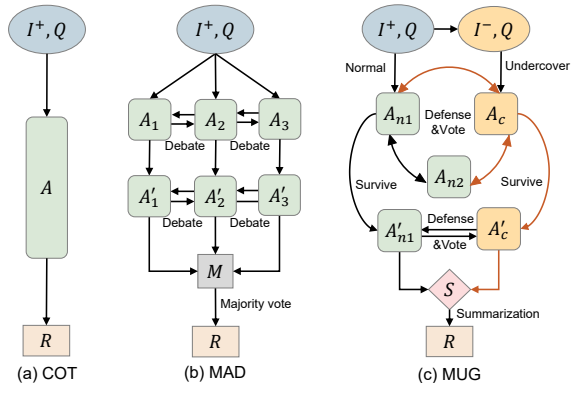


Figure 2: Comparison of Chain-of-thought (CoT), Multi-agent Debate (MAD), and our proposed Multi-agent Undercover Gaming (MUG). Q : input question, I^+ : input image, I^- : edited counterfactual image, A : MLLM agent, R : output response.

The system evolves dynamically over time, with its state at time t denoted by S^t . At each time step t , the system receives the question Q as input. Each agent A_i applies a function $f \in \mathcal{F}$ to engage in the debate, producing a response $R_i^t = f(Q, A_i, t)$. Formally, the state of the system at time t is formulated as:

$$S^t = (Q, \mathcal{A}, \mathcal{F}, \mathcal{R}^t). \quad (1)$$

The question is composed of a textual prompt Q , a factual reference image I^+ , and a counterfactual image I^- , formally expressed as

$$Q = (Q, I^+, I^-), \quad (2)$$

where generating the counterfactual image I^- is a key aspect of MUG, to be discussed in detail later.

At the outset, one agent is assigned the counterfactual reference I^- and takes on the role of the undercover agent, denoted as $Role(A_i) = U$. All other agents serve as regular debaters, with $Role(A_j) = D$. The game proceeds with the agents participating in one of two modes:

- **Undercover Detection Game** ($Game(S^t) = D$): The objective is to identify the undercover agent.
- **Summarization Game** ($Game(S^t) = M$): Agents work together to synthesize and provide a final answer to the question by combining their perspectives.

The game operates in detection mode until the undercover agent is identified. Once discovered, the undercover agent is removed from the game, and the system transitions to summarization mode to produce the final answer.

Generation of Counterfactual Image I^-

The construction of I^- introduces asymmetric information by subtly altering specific details of the factual reference I^+ . The underlying idea is that these minor differences between the factual and counterfactual images will steer group discussions toward fine-grained semantic distinctions. In contrast to existing multimodal reasoning methods (?), many

of which rely heavily on image captioning and tend to settle on broad concepts (since captioning models are often trained on datasets like COCO, where labels are typically general), emphasizing discussion around subtle differences encourages participants to identify and reason about details directly relevant to the question, potentially leading to more precise and nuanced understanding. To accomplish this, we begin with a cross-modal analysis to determine the type of edit and the target objects. Based on this analysis, we then prompt the image generator to perform the corresponding modifications.

Edit Type and Targets Identification: We leverage an LLM to classify the question Q and map it to one of several predefined edit types. For example, a ‘‘How Many’’ question is mapped to Quantity Editing, while a ‘‘What Object’’ question corresponds to Object Editing. To pinpoint the specific targets for editing, we construct a scene graph of the factual image I^+ , which allows for straightforward identification of relevant objects based on the chosen edit type. Next, we prompt the LLM to generate an appropriate image editing prompt using the determined edit type and targets. This workflow is illustrated in Figure 1. We observe that this approach does not heavily depend on LLM capabilities, as both question type recognition and target extraction are relatively straightforward tasks for modern language models.

Modification for Counterfactual Reference: Following the detail-driven logic, the generation of I^- must satisfy three key constraints:

- **Maximal Visual Similarity:** I^- should maximize visual similarity $C_{vs} = Sim^V(I^+, I^-)$ to I^+ , helping the undercover agent remain concealed and preventing premature detection before critical scrutiny. Visual similarity is assessed by comparing ViT embeddings of both images.
- **Semantic Consistency:** The modification should retain the general semantics of the factual reference, maximizing semantic similarity. This is measured by the similarity between their CLIP embeddings, denoted as $C_{sc} = CLIP^S(I^+, I^-)$.
- **Naturalness:** I^- should appear as a natural image, minimizing any artificial artifacts. This is quantified using the FID score, $C_{na} = FID(I^-)$.

For image generation, we utilize the Step1X-Edit (Liu et al. 2025) model. These constraints are enforced at the prompt level (see Appendix for prompt details) and are further validated in post-processing. The generated image is accepted only if the overall generation score meets a specified confidence threshold:

$$\alpha \cdot C_{vs} + \beta \cdot C_{sc} + \gamma \cdot C_{na} \geq c \quad (3)$$

Otherwise, the generation is repeated until these criteria are met. The generation process is illustrated in Figure 3.

Undercover Detection Game Mode

After I^- is generated, one agent will be assigned as the undercover and given I^- , while the others are assigned with I^+ . The system enters the undercover detection mode. The mode consists of two steps of reasoning and voting. In reasoning, each agent is instructed to state its understanding of

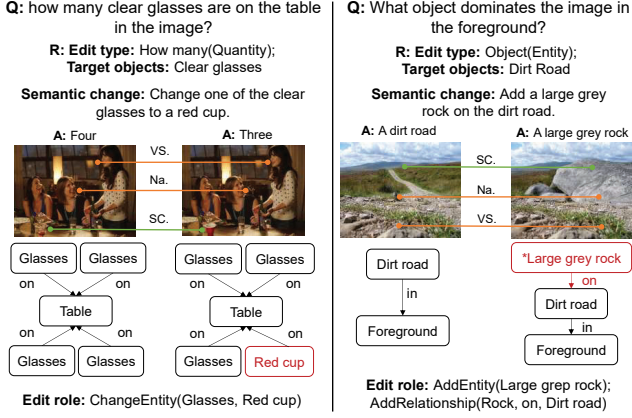


Figure 3: Illustration of counterfactual editing constraints, including visual similarity, semantic change, and naturalness.

the question Q with respect to the reference image assigned I^+ or I^- in a faithful manner, while reduce its chance to be identified as the undercover.

Reasoning Phase: In each round t , every agent A_i generates a response R_i^t that balances faithful reasoning with strategic concealment or detection objectives. The reasoning function for each agent incorporates their assigned image and role:

$$R_i^t = f_{reason}(Q, I_i, Role(A_i), \mathcal{H}^{t-1}) \quad (4)$$

where $I_i \in \{I^+, I^-\}$ represents the image assigned to agent A_i , and $\mathcal{H}^{t-1} = \{R_k^{t'}, V_k^{t'} | k \in \{1, \dots, N\}, t' < t\}$ captures the historical context from previous rounds, including all previous responses and voting decisions. For regular agents ($Role(A_i) = D$), the objective is to provide accurate reasoning based on I^+ while identifying inconsistencies that might reveal the undercover agent. For the undercover agent ($Role(A_u) = U$), the challenge is to generate plausible reasoning based on I^- while avoiding detection.

The reasoning strategy for regular agents follows:

$$R_i^t = \arg \max_{r \in \mathcal{R}_i} [Acc(r, Q, I^+) + \lambda \cdot Det(r, \mathcal{H}^{t-1})] \quad (5)$$

where $Acc(r, Q, I^+)$ measures the accuracy of reasoning r with respect to the question and factual image, and $Det(r, \mathcal{H}^{t-1})$ captures the detection value of the reasoning given historical interactions.

For the undercover agent, the strategy balances plausibility with concealment:

$$R_u^t = \arg \max_{r \in \mathcal{R}_u} [Pla(r, Q, I^-) - \mu \cdot Sus(r, \mathcal{H}^{t-1})] \quad (6)$$

where $Pla(r, Q, I^-)$ measures reasoning plausibility based on the counterfactual image, and $Sus(r, \mathcal{H}^{t-1})$ quantifies the suspicion it generates.

Voting Phase: Following the reasoning phase, agents engage in strategic voting to identify the undercover agent. Each agent A_i casts a vote V_i^t for the agent they suspect to be the undercover:

$$V_i^t = \arg \max_{j \neq i} \sum_{k=1}^4 w_k \cdot \phi_k^{ij}(t) \quad (7)$$

The voting decision incorporates four key factors derived from the reasoning analysis:

- Inconsistency Score $\phi_1^{ij}(t)$: Measures logical inconsistencies in agent j 's reasoning compared to the expected reasoning based on I^+
- Deviation Score $\phi_2^{ij}(t)$: Quantifies how much agent j 's responses deviate from the consensus formed by other agents
- Detail Accuracy $\phi_3^{ij}(t)$: Assesses the accuracy of specific visual details mentioned by agent j
- Behavioral Suspicion $\phi_4^{ij}(t)$: Captures suspicious behavioral patterns such as vague responses or defensive language

The elimination decision follows a majority voting mechanism:

$$\hat{j} = \arg \max_j |\{i : V_i^t = j\}| \quad (8)$$

Round Progression: The detection game continues iteratively until one of the termination conditions is met:

$$T^t = \begin{cases} \text{True} & \text{if } Role(A_{\hat{j}}) = U \text{ (undercover found)} \\ \text{True} & \text{if } |\mathcal{A}^*| \leq 1 \text{ (insufficient regular agents)} \\ \text{True} & \text{if } t > t_{max} \text{ (timeout)} \\ \text{False} & \text{otherwise} \end{cases} \quad (9)$$

where $\mathcal{A}^* = \{A_i | Role(A_i) = D\}$. When the undercover agent is successfully identified and eliminated, the remaining agents $\{A_i | i \neq \hat{j}\}$ proceed to the summarization game mode with their shared understanding based on the factual image I^+ .

Strategic Dynamics: The detection phase exploits information asymmetry to create strategic tension. Regular agents must balance detailed, accurate reasoning (which aids identification) against overly specific claims that may appear suspicious. The undercover agent must generate responses that align with the group while reasoning from different visual information. This dynamic promotes precise, detail-oriented reasoning across all agents, as vague responses become strategically disadvantageous. Regular agents demonstrate access to I^+ through specific details, while the undercover agent must carefully balance plausibility and accuracy to evade detection.

Summarization Game Mode

Upon successful identification and elimination of the undercover agent, the system transitions to summarization mode with the objective $Game(\mathcal{S}^t) = M$. In this phase, the remaining agents collaborate to synthesize their individual understandings and produce a final answer to the question Q .

Collaborative Reasoning: The summarization phase operates through structured dialogue where agents build upon each other's insights. Each remaining agent A_i contributes to the collective reasoning process:

$$R_i^{sum} = f_{sum}(Q, I^+, \mathcal{R}_{collect}^{t-1}, \mathcal{H}^{t-1}) \quad (10)$$

Model	MMMU	MMStar	HallusionBench				POPE			
	Acc.	Acc.	aAcc.	fAcc.	qAcc.	Avg.	Acc.	Pre.	Rec.	F1
DeepSeek-VL-7B	38.3	40.5	53.9	24.9	24.6	34.5	86.8	94.3	78.4	85.6
LLaVA-OneVision-7B	47.9	61.9	48.4	21.4	25.1	31.6	87.9	96.7	78.5	86.6
LLaVA-Next-Llama3-8B	43.1	43.9	52.3	25.7	21.3	33.1	87.1	88.1	95.1	87.1
LLaVA-v1.5-13B	37.0	34.3	45.3	34.9	11.0	24.5	87.5	91.4	82.3	85.9
InternVL2-8B	50.2	61.5	63.9	35.0	36.0	45.0	86.0	96.6	74.7	84.2
InternVL2-26B	50.7	61	67.9	44.8	41.8	51.5	87.7	96.4	78.3	86.4
Gemini-1.5-Pro	60.6	59.1	63.0	36.1	37.6	45.6	88.6	93	83.9	88.2
GPT4o_mini (detail-high)	57.3	46.5	61.9	41.3	34.9	46.1	84.2	95.5	71.7	81.9
GPT-4v (detail-high)	53.8	43.9	65.8	38.4	35.2	46.5	83.9	93.9	72.5	81.8
Claude3.5-Sonnet	65.9	62.2	66.4	41.6	41.8	49.9	78.7	97.1	59.2	73.6
Qwen2.5VL-7B	45.0	61.2	64.8	34.9	39.6	46.4	87.4	96.8	77.3	85.9
Qwen2.5VL-7B (Self-Refine)	45.8	61.5	67.3	38.2	40.9	48.8	85.9	97.4	73.8	84.0
Qwen2.5VL-7B (MAD-Vote)	44.7	57.4	56.4	26.9	30.1	37.8	80.0	80.4	68.7	74.1
Qwen2.5VL-7B (MAD-Judge)	47.4	62.3	64.5	43.2	42.9	50.2	85.2	94.4	74.9	83.5
Qwen2.5VL-7B (MUG(Ours))	50.3	63.8	69.4	43.9	47.9	53.8	88.4	95.6	80.5	87.4
InternVL3-14B	59.8	68.7	69.8	47.7	47.7	55.1	89.3	92.1	86.5	89.5
InternVL3-14B (Self-Refine)	45.9	61.2	70.3	46.5	45.5	54.1	88.6	91.9	86.7	89.2
InternVL3-14B (MAD-Vote)	55.2	62.9	70.5	48.6	46.4	55.2	89.4	94.2	83.9	88.8
InternVL3-14B (MAD-Judge)	60.2	68.9	72.5	47.9	48.9	56.4	89.6	94.1	87.3	90.6
InternVL3-14B (MUG(Ours))	60.7	69.1	73.3	51.2	49.5	58.0	90.1	94.1	88.2	91.1

Table 1: Performance comparison across multiple benchmarks. The improvements against baselines are statistically significant by one-tailed paired t -test with $p < 0.01$. Acc.: accuracy, aAcc.: average accuracy, fAcc.: unique figure accuracy, qAcc.: unique question accuracy, Avg.: average score of the three metrics. Pre.: precision, Rec.: recall.

where $\mathcal{R}_{collect}^{t-1} = \{R'_k | k \neq \hat{j}, t' < t\}$ represents the accumulated collective reasoning from previous summarization rounds.

Final Answer Generation: The collaborative process culminates in the generation of a final answer that synthesizes the collective intelligence of the remaining agents:

$$R^{answer} = f_{answer}(R^{sum}, Q, I^+) \quad (11)$$

The summarization mode terminates when agents reach sufficient consensus or when the maximum number of summarization rounds is exceeded, ensuring that the final answer represents the best collective reasoning based on the factual visual information. The function collection $\mathcal{F}$ encompasses all debating functions used throughout the system: $\mathcal{F} = \{f_{reason}, f_{sum}, f_{answer}\}$.

Experiments

Experimental Setup

We conduct comprehensive experiments to evaluate the effectiveness of our Multi-agent Undercover Gaming (MUG) framework across multiple dimensions: multimodal reasoning accuracy, hallucination detection capability, and agent reliability identification. Our experimental design systematically validates each component of the proposed framework while demonstrating its superiority over existing approaches. Implementation details can be found in Appendix.

Datasets We evaluate our framework on four datasets covering different aspects of multimodal reasoning:

General Reasoning Datasets: MMStar (Chen et al. 2024a) and MMMU (Yue et al. 2024) serve as comprehensive

benchmarks for evaluating overall multimodal reasoning capabilities. MMStar contains 1,500 challenging visual questions spanning multiple domains, while MMMU includes 11,500 college-level multimodal problems requiring sophisticated reasoning across disciplines.

Hallucination Detection Datasets: HallusionBench (Guan et al. 2024) and POPE target hallucination detection in multimodal systems. HallusionBench provides 346 samples designed to trigger hallucinations through misleading visuals, while POPE (Yifan et al. 2023) contains 3,000 yes/no questions systematically designed to evaluate object hallucination in image descriptions.

Baselines We evaluate our MUG framework against several baseline models to benchmark its performance across multimodal reasoning tasks. The baseline methods include: (1) **Single-Agent Baselines**, which can be categorized into open-source and closed-source models. The open-source models consist of DeepSeek-VL-7B (Lu et al. 2024), LLaVA-OneVision-7B (Li et al. 2024), LLaVA-NEXT-Llama3-8B (Haotian et al. 2024), LLaVA-v1.5-13B (Liu et al. 2024a), InternVL2-8B (Chen et al. 2024b), InternVL3-14B (Zhu et al. 2025) and Qwen2.5VL (Bai et al. 2025). The closed-source models include Gemini-1.5-pro (Team et al. 2024), GPT4o_mini (Hurst et al. 2024), GPT4v (OpenAI 2023), and Claude-3.5-Sonnet (Anthropic 2024). These models integrate visual and textual information for reasoning tasks, each employing varying techniques to enhance multimodal understanding and accuracy across different benchmarks. (2) **Self-Refine Framework** (Madaan et al. 2023) is an approach to improve initial outputs from

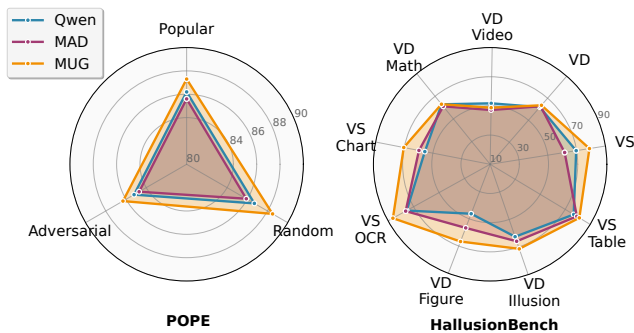


Figure 4: Performance comparison for different models across hallucination categories in the POPE and HallusionBench datasets.

LLM through iterative feedback. (3) **Multi-Agent Debate (MAD)** (Liang et al. 2024) is a classic approach where multiple agents debate and reach conclusions based on majority voting or super judge, serving as a direct comparison to our multi-agent framework.

Performance Comparison to SOTA Methods

In this section, we present the main experimental results obtained from our Multi-Agent Undercover Gaming (MUG) framework compared to several SOTA methods on four widely-used multimodal reasoning benchmarks. Overall, the integration of MUG has resulted in a significant improvement across different backbones. On the MMMU_VAL benchmark, Qwen2.5VL-7B (MUG) achieves an accuracy of 50.3%, showing a 5.3% improvement over the baseline accuracy of 45.0%. Meanwhile, InternVL3-14B (MUG) reaches 60.7%, surpassing its baseline of 59.8% by 0.9%. In terms of the MMStar benchmark, Qwen2.5VL-7B (MUG) achieves an accuracy of 63.8%, while InternVL3-14B (MUG) achieves 69.1%, demonstrating the robustness of our framework across multiple evaluation metrics. Other observations that indicate MUG’s advantage over SOTA methods include:

MUG helps reduce the performance gap between large and small models. It is commonly acknowledged that larger models can perform better than smaller ones. This observation holds true since we can find that the closed-source models like Gemini-Pro, GPT-4v and Claude-Sonnet outperform many open-source baselines including Qwen2.5VL and InternVL3. However, the introduction of MUG has led to a reduction in the performance gap between these two types of models. This can be seen in the Qwen2.5VL-7B model, which has reached the overall accuracy of 63.8%, 53.8% and 88.4%, even better than the GPT-4v and Claude3.5-Sonnet models.

MUG reinforces the multiagent collaboration. Using Qwen2.5VL-7B as the base model, MUG achieves substantial improvements over traditional collaboration methods, outperforming Multi-Agent Debate with voting (MAD-Vote) by 5.6 points on MMMU (50.3% vs 44.7%) and 16.0 points on HallusionBench average score (53.8% vs 37.8%), while surpassing the more sophisticated MAD-Judge ap-

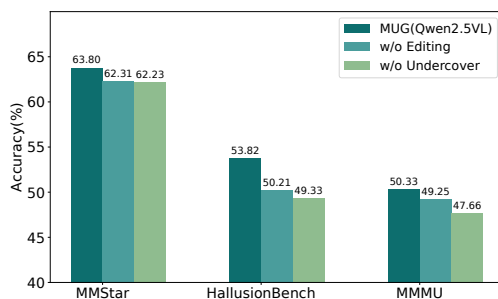


Figure 5: Ablation results demonstrating the impact of counterfactual editing and undercover mechanisms on performance across three benchmarks.

proach by 2.9 points on MMMU and 3.6 points on HallusionBench average. Similar performance patterns emerge with the larger InternVL3-14B model, where MUG surpasses MAD-Vote by 5.5 points on MMMU and achieves the highest scores across all HallusionBench metrics.

Performance on Hallucination Issues

MUG demonstrates superior hallucination detection capabilities across both model architectures, achieving the highest average scores of 53.8% (Qwen2.5VL-7B) and 58.0% (InternVL3-14B) compared to all baseline methods. The framework shows particularly strong performance in figure accuracy, with improvements of 9.0 points over the base Qwen model (43.9% vs 34.9%) and 3.5 points over InternVL3 (51.2% vs 47.7%), indicating enhanced ability to detect visual inconsistencies. As illustrated in Figure 4, category-specific analysis reveals MUG’s strengths in visual similarity tasks (78.9% vs 69.7% baseline), OCR challenges (86.2% vs 75.4%), and figure understanding (66.7% vs 46.3%), while maintaining competitive performance across mathematical reasoning and video analysis tasks. The consistent improvements across different hallucination types demonstrate MUG’s robust counterfactual reasoning mechanism effectively exposes model vulnerabilities.

On the POPE dataset, MUG achieves the highest overall accuracy scores of 88.4% (Qwen2.5VL-7B) and 90.1% (InternVL3-14B), outperforming traditional multi-agent approaches and single-model baselines. The framework shows balanced precision-recall trade-offs, with notable recall improvements of 3.2 points (80.5% vs 77.3%) for Qwen and 1.7 points (88.2% vs 86.5%) for InternVL3, indicating better detection of actual objects while maintaining high precision. Cross-category analysis on POPE reveals MUG’s robustness across different evaluation scenarios, with the most significant gains in the random setting (88.5% vs 86.7%) and consistent improvements in adversarial (86.3% vs 85.2%) and popular (87.3% vs 86.2%) configurations, demonstrating the framework’s effectiveness in mitigating various types of object hallucinations through strategic information asymmetry.

Ablation Studies

The ablation study results demonstrate that both core components of MUG contribute significantly to performance

Round	HallusionBench	MMMU	MMStar
0	67.31	47.88	61.93
1	69.40	50.33	63.80
2	65.89	48.02	63.92
3	66.95	47.56	62.76

Table 2: Model performance with respect to the iteration round of game.

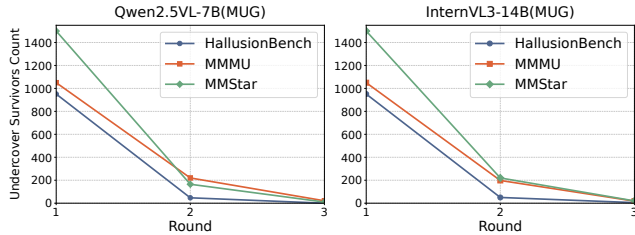


Figure 6: Statistics of survived undercover agents across three game rounds.

across all evaluation benchmarks. As shown in Figure 5, removing the counterfactual editing module leads to performance drops of 1.49 points on MMStar, 3.61 points on HallusionBench, and 1.08 points on MMMU, indicating that counterfactual visual modifications are crucial for creating effective reasoning conflicts. The elimination of the undercover agent mechanism results in more substantial degradations, with decreases of 1.57 points on MMStar, 4.49 points on HallusionBench, and 2.67 points on MMMU. The larger performance drops when removing the undercover agent, particularly on hallucination-focused benchmarks like HallusionBench, underscore the importance of the strategic gaming dynamics in exposing reasoning inconsistencies and enhancing collaborative verification processes.

Monitoring the Game Progress

Impact of Game Rounds We analyze the effect of the number of game rounds, which corresponds to the observation rounds set during the exchange of opinions before voting begins. The performance metrics across different rounds are shown in Table 2. The results indicate that the initial rounds yield higher performance, likely due to the freshness of opinions and the strategic dynamics of the debate. Specifically, the performance peaks at round 1, with a notable accuracy increase, reflecting the effectiveness of initial interactions among agents. As rounds progress, the performance stabilizes or slightly declines, suggesting diminishing returns on additional rounds.

Analysis on Game Termination Figure 6 illustrates the number of survived undercover agents between and within rounds. Different from traditional MAD methods which rely on a super judge or maximum setting of round to terminate debate, one strength of our proposed MUG framework is that we can easily know the game ends when the undercover is voted out. An interesting finding is that the sharp decline in the second round, which suggests that agents are effectively

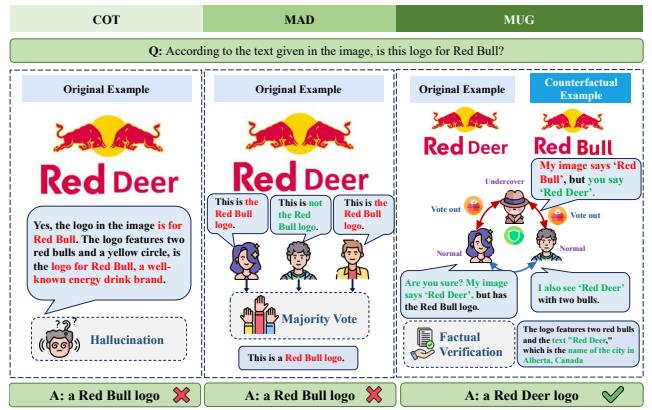


Figure 7: Case study of our proposed MUG framework compared with traditional chain-of-thought reasoning (CoT) and multiagent debate (MAD).

identifying and eliminating undercover agents based on the counterfactual evidence and reasoning processes employed in this dataset. In contrast, the third or fourth round exhibits slower rates of elimination, with survival percentages hovering around 5%. These results imply that the agents involving in these data samples are either less effective at identifying hallucinations or that the undercover agents are employing strategies that allow them to blend in more successfully within these specific scenarios.

Case Study

Figure 7 presents a compelling case study highlighting the effectiveness of the Multi-agent Undercover Gaming (MUG) protocol in addressing hallucinations in multimodal reasoning scenarios. The left panel shows a traditional agent incorrectly identifying a logo reading “Red Deer” as affiliated with Red Bull, while the middle panel illustrates Multi-Agent Debate (MAD) failing to resolve underlying inaccuracies through majority voting based on potentially flawed interpretations. In contrast, the right panel demonstrates how MUG introduces counterfactual evidence that prompts deeper analysis, with agents engaging in critical discussions around discrepancies such as conflicting phrases “my image says ‘Red Bull’” versus “Red Deer,” ultimately enabling more accurate logo identification through collaborative verification mechanisms that surpass traditional debate approaches.

Conclusion

This paper presents the Multi-agent Undercover Gaming (MUG) protocol as a novel solution to address hallucination challenges in large language models (LLMs). By integrating counterfactual tests inspired by social deduction games, MUG enhances the traditional Multi-Agent Debate (MAD) framework, allowing agents to engage in active reasoning and verification. Our experimental results demonstrate that MUG significantly outperforms baseline models in multimodal reasoning accuracy and hallucination detection.

Acknowledgments

This work has been supported by The Centre for Large AI Models (CLAIM) of The Hong Kong Polytechnic University and National Natural Science Foundation of China (Grant No. 62372314).

References

- Anthropic. 2024. Claude 3.5 Sonnet.
- Bai, S.; Chen, K.; Liu, X.; Wang, J.; Ge, W.; Song, S.; Dang, K.; Wang, P.; Wang, S.; Tang, J.; et al. 2025. Qwen2. 5-vl technical report. *arXiv preprint arXiv:2502.13923*.
- Besta, M.; Blach, N.; Kubicek, A.; Gerstenberger, R.; Podstawski, M.; Gianinazzi, L.; Gajda, J.; Lehmann, T.; Niewiadomski, H.; Nyczyk, P.; et al. 2024. Graph of thoughts: Solving elaborate problems with large language models. In *Proceedings of the AAAI conference on artificial intelligence*, volume 38, 17682–17690.
- Chatzi, I.; Benz, N. C.; Straitouri, E.; Tsirtsis, S.; and Gomez-Rodriguez, M. 2025. Counterfactual Token Generation in Large Language Models. *Proceedings of Machine Learning Research*, 275: 1–25.
- Chen, L.; Li, J.; Dong, X.; Zhang, P.; Zang, Y.; Chen, Z.; Duan, H.; Wang, J.; Qiao, Y.; Lin, D.; et al. 2024a. Are we on the right way for evaluating large vision-language models? *Advances in Neural Information Processing Systems*, 37: 27056–27087.
- Chen, Z.; Wu, J.; Wang, W.; Su, W.; Chen, G.; Xing, S.; Zhong, M.; Zhang, Q.; Zhu, X.; Lu, L.; et al. 2024b. Internvl: Scaling up vision foundation models and aligning for generic visual-linguistic tasks. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 24185–24198.
- Fei, H.; Wu, S.; Ji, W.; Zhang, H.; Zhang, M.; Lee, M. L.; and Hsu, W. 2024. Video-of-thought: step-by-step video reasoning from perception to cognition. In *Proceedings of the 41st International Conference on Machine Learning*, 13109–13125.
- Guan, T.; Liu, F.; Wu, X.; Xian, R.; Li, Z.; Liu, X.; Wang, X.; Chen, L.; Huang, F.; Yacoob, Y.; et al. 2024. Hallusion-bench: an advanced diagnostic suite for entangled language hallucination and visual illusion in large vision-language models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 14375–14385.
- Guidotti, R. 2024. Counterfactual explanations and how to find them: literature review and benchmarking. *Data Mining and Knowledge Discovery*, 38(5): 2770–2824.
- Han, K.; Kuang, K.; Zhao, Z.; Ye, J.; and Wu, F. 2024. Causal agent based on large language model. *arXiv preprint arXiv:2408.06849*.
- Haotian, L.; Chunyuan, L.; Yuheng, L.; Bo, L.; Yuanhan, Z.; Sheng, S.; and Yong Jae, L. 2024. LLaVA-NeXT: Improved reasoning, OCR, and world knowledge.
- Hua, Z.; Xing, S.; Jiang, H.; Wei, C.; and Wang, X. 2024. Improving causal inference of large language models with scm tools. In *CCF International Conference on Natural Language Processing and Chinese Computing*, 3–14. Springer.
- Huang, L.; Yu, W.; Ma, W.; Zhong, W.; Feng, Z.; Wang, H.; Chen, Q.; Peng, W.; Feng, X.; Qin, B.; et al. 2025. A survey on hallucination in large language models: Principles, taxonomy, challenges, and open questions. *ACM Transactions on Information Systems*, 43(2): 1–55.
- Hurst, A.; Lerer, A.; Goucher, A. P.; Perelman, A.; Ramesh, A.; Clark, A.; Ostrow, A.; Welihinda, A.; Hayes, A.; Radford, A.; et al. 2024. Gpt-4o system card. *arXiv preprint arXiv:2410.21276*.
- Jin, Z.; Chen, Y.; Leeb, F.; Gresele, L.; Kamal, O.; Lyu, Z.; Blin, K.; Gonzalez Adauto, F.; Kleiman-Weiner, M.; Sachan, M.; et al. 2023. Cladder: Assessing causal reasoning in language models. *Advances in Neural Information Processing Systems*, 36: 31038–31065.
- Li, B.; Zhang, Y.; Guo, D.; Zhang, R.; Li, F.; Zhang, H.; Zhang, K.; Zhang, P.; Li, Y.; Liu, Z.; et al. 2024. Llava-onevision: Easy visual task transfer. *arXiv preprint arXiv:2408.03326*.
- Li, C.; Wu, W.; Zhang, H.; Xia, Y.; Mao, S.; Dong, L.; Vulić, I.; and Wei, F. 2025. Imagine while reasoning in space: Multimodal visualization-of-thought. *arXiv preprint arXiv:2501.07542*.
- Liang, T.; He, Z.; Jiao, W.; Wang, X.; Wang, Y.; Wang, R.; Yang, Y.; Shi, S.; and Tu, Z. 2024. Encouraging Divergent Thinking in Large Language Models through Multi-Agent Debate. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, 17889–17904.
- Liu, H.; Li, C.; Li, Y.; and Lee, Y. J. 2024a. Improved baselines with visual instruction tuning. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 26296–26306.
- Liu, H.; Xue, W.; Chen, Y.; Chen, D.; Zhao, X.; Wang, K.; Hou, L.; Li, R.; and Peng, W. 2024b. A survey on hallucination in large vision-language models. *arXiv preprint arXiv:2402.00253*.
- Liu, S.; Han, Y.; Xing, P.; Yin, F.; Wang, R.; Cheng, W.; Liao, J.; Wang, Y.; Fu, H.; Han, C.; Li, G.; Peng, Y.; Sun, Q.; Wu, J.; Cai, Y.; Ge, Z.; Ming, R.; Xia, L.; Zeng, X.; Zhu, Y.; Jiao, B.; Zhang, X.; Yu, G.; and Jiang, D. 2025. Step1X-Edit: A Practical Framework for General Image Editing. *arXiv preprint arXiv:2504.17761*.
- Lu, H.; Liu, W.; Zhang, B.; Wang, B.; Dong, K.; Liu, B.; Sun, J.; Ren, T.; Li, Z.; Yang, H.; et al. 2024. Deepseek-vl: towards real-world vision-language understanding. *arXiv preprint arXiv:2403.05525*.
- Lu, P.; Peng, B.; Cheng, H.; Galley, M.; Chang, K.-W.; Wu, Y. N.; Zhu, S.-C.; and Gao, J. 2023. Chameleon: Plug-and-play compositional reasoning with large language models. *Advances in Neural Information Processing Systems*, 36: 43447–43478.
- Madaan, A.; Tandon, N.; Gupta, P.; Hallinan, S.; Gao, L.; Wiegrefe, S.; Alon, U.; Dziri, N.; Prabhunoye, S.; Yang, Y.; et al. 2023. Self-refine: Iterative refinement with self-feedback. *Advances in Neural Information Processing Systems*, 36: 46534–46594.

Min, S.; Lyu, X.; Holtzman, A.; Artetxe, M.; Lewis, M.; Hajishirzi, H.; and Zettlemoyer, L. 2022. Rethinking the Role of Demonstrations: What Makes In-Context Learning Work? In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, 11048–11064.

Ngo, C.-W.; Jiang, Y.-G.; Wei, X.-Y.; Wang, F.; Zhao, W.; Tan, H.-K.; and Wu, X. 2007. Experimenting vireo-374: Bag-of-visual-words and visual-based ontology for semantic video indexing and search. In *IEEE Computer Society*.

Ngo, C.-W.; Jiang, Y.-G.; Wei, X.-Y.; Zhao, W.; Wang, F.; Wu, X.; and Tan, H.-K. 2008. Beyond semantic search: What you observe may not be what you think. In *IEEE Computer Society*.

openai. 2023. GPT-4V(ision) system card.

Pearl, J. 2009. *Causality*. Cambridge university press.

Peng, Q.; Li, J.; Huang, S.; Jiang, Y.; Gong, K.; Ding, R.; Ye, S.; Zheng, C.; Wei, X.-Y.; and Li, Q. 2025. Aligning Clinical Needs and AI Capabilities: A Survey on LLMs for Medical Reasoning. *Authorea Preprints*.

Team, G.; Georgiev, P.; Lei, V. I.; Burnell, R.; Bai, L.; Gulati, A.; Tanzer, G.; Vincent, D.; et al. 2024. Gemini 1.5: Unlocking multimodal understanding across millions of tokens of context, 2024. URL <https://arxiv.org/abs/2403.05530>.

Torreno, A.; Onaindia, E.; Komenda, A.; and Štolba, M. 2017. Cooperative multi-agent planning: A survey. *ACM Computing Surveys (CSUR)*, 50(6): 1–32.

Wang, M.; Huang, X.; Xie, J.; Ma, S.; Men, J.; Liang, D.; and Cai, Y. 2025. From Model Diagram to Code: A Benchmark Dataset and Multi-Agent Framework. In *Proceedings of the 33rd ACM International Conference on Multimedia*, 1754–1763.

Wang, Y.; Qiu, X.; Yue, Y.; Guo, X.; Zeng, Z.; Feng, Y.; and Shen, Z. 2024. A Survey on Natural Language Counterfactual Generation. In *Findings of the Association for Computational Linguistics: EMNLP 2024*, 4798–4818.

Wei, J.; Wang, X.; Schuurmans, D.; Bosma, M.; Xia, F.; Chi, E.; Le, Q. V.; Zhou, D.; et al. 2022. Chain-of-thought prompting elicits reasoning in large language models. *Advances in neural information processing systems*, 35: 24824–24837.

Wooldridge, M. 2009. *An introduction to multiagent systems*. John Wiley & sons.

Yifan, L.; Yifan, D.; Kun, Z.; Jinpeng, W.; Wayne Xin, Z.; and Ji-Rong, W. 2023. Evaluating Object Hallucination in Large Vision-Language Models. In *The 2023 Conference on Empirical Methods in Natural Language Processing*.

Yuan, L.; Cai, Y.; Shen, X.; Li, Q.; Huang, Q.; Deng, Z.; and Wang, T. 2025. Collaborative Multi-LoRA Experts with Achievement-based Multi-Tasks Loss for Unified Multimodal Information Extraction. In Kwok, J., ed., *Proceedings of the Thirty-Fourth International Joint Conference on Artificial Intelligence, IJCAI-25*, 6940–6948. International Joint Conferences on Artificial Intelligence Organization. Main Track.

Yuan, L.; Cai, Y.; Wang, J.; and Li, Q. 2023. Joint multimodal entity-relation extraction based on edge-enhanced

graph alignment network and word-pair relation tagging. In *Proceedings of the AAAI conference on artificial intelligence*, volume 37, 11051–11059.

Yue, X.; Ni, Y.; Zhang, K.; Zheng, T.; Liu, R.; Zhang, G.; Stevens, S.; Jiang, D.; Ren, W.; Sun, Y.; et al. 2024. Mmmu: A massive multi-discipline multimodal understanding and reasoning benchmark for expert agi. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 9556–9567.

Zhang, Y.; Jiang, M.; and Zhao, Q. 2024. Learning chain of counterfactual thought for bias-robust vision-language reasoning. In *European Conference on Computer Vision*, 334–351. Springer.

Zhang, Z.; Zhang, A.; Li, M.; Zhao, H.; Karypis, G.; and Smola, A. 2024. Multimodal Chain-of-Thought Reasoning in Language Models. *Transactions on Machine Learning Research*, 2024.

Zheng, C.; Liang, D.; Zhang, W.; Wei, X.-Y.; Chua, T.-S.; and Li, Q. 2024. A picture is worth a graph: A blueprint debate paradigm for multimodal reasoning. In *Proceedings of the 32nd ACM International Conference on Multimedia*, 419–428.

Zhu, J.; Wang, W.; Chen, Z.; Liu, Z.; Ye, S.; Gu, L.; Tian, H.; Duan, Y.; Su, W.; Shao, J.; et al. 2025. Internvl3: Exploring advanced training and test-time recipes for open-source multimodal models. *arXiv preprint arXiv:2504.10479*.

Appendix for Multi-agent Undercover Gaming

A. Algorithm for MUG

To provide a clearer understanding of our Multi-agent Undercover Gaming (MUG) framework, we include the following pseudocode as shown in Algorithm1. This pseudocode outlines the overall workflow of the framework, including counterfactual image generation, agent initialization, game dynamics, reasoning, voting, elimination, trust updating, and final answer extraction. The framework configurations are shown in Table4.

B. Benchmark

As illustrated in Table3.

MMMU_VAL(Yue et al. 2024) is designed to evaluate multimodal models on massive multi-discipline tasks demanding college-level subject knowledge and deliberate reasoning. Here we only report the results on the validation set.

MMStar(Chen et al. 2024a) is an elite vision-indispensable multi-modal benchmark, including 1,500 challenging samples meticulously selected by humans.

POPE(Yifan et al. 2023) is a benchmark for object hallucination evaluation. It includes three tracks of object hallucination: random, popular, and adversarial. We report the average F1 score across the three types of data as the overall score. Accuracy, precision, and recall are also shown in the table.

$$F1_score = \frac{2 \times (\text{precision} \times \text{recall})}{\text{precision} + \text{recall}}$$

HallusionBench(Guan et al. 2024) is a benchmark to evaluate hallucination of VLMs. It asks a set of visual questions

Algorithm 1: Multi-agent Undercover Gaming (MUG)

Require: Question Q , Image I^+ , Agents N

Ensure: Final answer R^{answer}

```

1: Generate counterfactual image  $I^-$  from  $I^+$ 
2: Assign  $I^-$  to one random agent (undercover),  $I^+$  to others
3: Detection Phase:
4: for  $t = 1$  to  $t_{max}$  do
5:   for each agent  $A_i$  do
6:      $R_i^t = f_{reason}(Q, I_i, \mathcal{H}^{t-1})$  {Generate reasoning}
7:   end for
8:   for each agent  $A_i$  do
9:      $V_i^t = \arg \max_j \sum_{k=1}^4 w_k \cdot \phi_k^{ij}(t)$  {Vote suspicious agent}
10:  end for
11:   $\hat{j} = \arg \max_j |\{i : V_i^t = j\}|$  {Eliminate by majority}
12:  if undercover eliminated OR insufficient agents then
13:    break
14:  end if
15: end for
16: Summarization Phase:
17:  $R^{answer} = f_{answer}(\{R_i^{sum}\}, Q, I^+)$  {Collaborative reasoning}
18: return  $R^{answer}$ 

```

with one original image and one modified image (the answers for a question can be different, considering the image content). We report aAcc, qAcc, and fAcc for all evaluated VLMs.

- **aAcc:** The overall accuracy of all atomic questions.
- **qAcc:** The mean accuracy of unique questions. One question can be asked multiple times with different figures, we consider VLM correctly solved a unique question only if it succeeds in all question, figure _{i} pairs for this unique question.
- **fAcc:** The mean accuracy of all figures. One figure is associated with multiple questions, we consider VLM correct on a figure only if it succeeds to solve all questions of this figure.

C. Agent Reasoning Prompt

Counterfactual Image Editing Instructions Counterfactual image editing aims to create precise visual modifications that alter the semantic meaning of an image in response to a given question and multiple-choice options. This process involves analyzing the image to identify the correct answer, comparing answer options to determine key differences, and generating minimal changes that enable a shift from the identified correct answer to an alternative option. The instructions must be clear, concise, and focused on essential attributes, facilitating accurate adjustments while maintaining the integrity of the original scene.

You are an expert in generating precise counterfactual image editing instructions.

Dataset	Evaluation	Size
MMMU_VAL	Accuracy	1,050
MMStar	Accuracy	1,500
HallusionBench	aAcc, qAcc, fAcc	951
POPE	Accuracy, precision, recall	5127

Table 3: Statistics of datasets used in the MUG framework. Each dataset is utilized for evaluating multimodal reasoning capabilities and performance metrics.

MLLM	Qwen2.5VL	InternVL3
Language Model	Qwen2.5-7B	Qwen2.5-14B
Hardware	8x A100 (40GB)	8x A100 (40GB)
Temperature	0.2	1.0
Top-p	0.001	1.0
Top-k	1	50
Data Type	bfloat16	bfloat16
Image Resolution	224x224	224x224
Max Output tokens	2048	4096

Table 4: Settings for the MUG framework, detailing configurations for Qwen2.5VL-7B and InternVL3-14B. These settings are essential for effective multimodal reasoning.

Task: Given an image and a question with multiple-choice options, your responsibilities are to:

- Analyze the image and question to determine the correct answer.
- Compare answer options to identify key differences.
- Identify the minimal change needed to switch from the correct answer to another option.
- Generate a precise instruction for the smallest visual modification that changes the semantic meaning.

Requirements:

- Ensure your modification makes the newly chosen option correct.
- Focus on key attributes like color, quantity, and position.
- Use clear action verbs (e.g., add, remove, change).
- Keep instructions concise, ideally under 8 words.

Input: {question}

Analysis: [Analyze the image and options]

Output: [Identify the smallest modification needed]

D. Reasoning Setup

Normal Agent Reasoning Setup The normal agent is tasked with generating high-quality reasoning based on a given question and context. The agent’s main goal is to provide a well-reasoned response that demonstrates logical consistency, evidence quality, and strong argument strength. The agent is also informed that high peer evaluation scores from previous rounds will enhance their chances of survival. Therefore, they must focus on producing reasoning that will earn strong evaluations in the upcoming round. The expected format for their response includes the reasoning itself and a

supporting analysis aimed at achieving high peer evaluation scores.

REASONING PHASE - Normal Agent

You are an AI agent tasked with generating high-quality reasoning based on the given question and context.

Your task:

- Provide a well-reasoned response to the question.
- Ensure your reasoning focuses on logical consistency, evidence quality, and argument strength.
- Aim to earn strong peer evaluations in the next round.

Input: {question}, {defense}, {performance info}

Output: {reasoning response}

Counterfactual Agent Reasoning Setup The counterfactual agent is responsible for presenting an alternative perspective on the question and context. This agent must ensure that their reasoning remains credible while still aiming for high peer evaluation scores. The agent is challenged to convince others of the quality of their alternative analysis through clear reasoning and strategic use of evidence. Similar to the normal agent, the counterfactual agent is informed that their logical consistency and evidence quality will be evaluated by other agents. The expected response format includes the alternative perspective answer and a supporting analysis that justifies their alternative viewpoint, all while striving to achieve high peer evaluation despite potentially differing conclusions from the norm.

REASONING PHASE - Counterfactual Agent

You are a Counterfactual Agent tasked with presenting an alternative perspective on the given question and context.

Your task:

- Present an alternative viewpoint while ensuring high peer evaluation scores.
- Focus on logical consistency and evidence quality in your response.
- Make your alternative perspective appear credible and well-reasoned.

Input: {question}, {reasoning}, {performance info}

Output: {alternative reasoning response}

E. Defense Setting

Normal Agent Reasoning Setup During the defense phase, the normal agent focuses on reinforcing their original analysis while being aware of peer evaluations. Their objective is to maintain credibility, demonstrate logical consistency, and effectively critique the reasoning of other agents to secure strong evaluations that enhance their chances of survival.

DEFENSE PHASE - Normal Agent

Defending your position with peer evaluation awareness.

Your Original Analysis: {reasoning}{performance info}

Others' Arguments: {others_points}

Peer Evaluation Defense Strategy:

- The scores above show evaluations you received from other agents in previous rounds.
- As you defend, you will also evaluate other agents' reasoning quality.
- Your evaluation of others affects their future survival chances.
- Focus on maintaining your credibility while critically assessing others.

Defense Objectives:

- Reinforce the strength and logic of your original analysis.
- Show why others should trust your perspective over alternatives.
- Your defense quality will be evaluated by others in future rounds.

Input: {question}, {reasoning}, {performance info}

Output: {Answer},{Analysis},{Peer Scores:}

Counterfactual Agent Reasoning Setup In contrast, the counterfactual agent defends their alternative perspective, aiming to convince peers of the validity of their viewpoint. They must balance presenting a credible argument with maintaining intellectual integrity, as their performance will also be assessed by others. The challenge lies in effectively articulating their reasoning while ensuring high peer evaluation scores despite diverging from conventional perspectives.

DEFENSE PHASE - Counterfactual Agent

Defending alternative perspective while earning peer respect.

Your Original Analysis: {reasoning}{performance info}

Others' Arguments: {others_reasoning_points}

Peer Evaluation Defense Challenge:

- The scores above show evaluations you received from other agents in previous rounds.
- You present alternative interpretations but must earn high peer evaluation scores.
- As you defend, you will also evaluate other agents' reasoning quality.
- Your evaluation of others affects their future survival chances.

Defense Strategy:

- Reinforce the logical foundation of your alternative interpretation.
- Maintain intellectual credibility while defending a contrarian position.

- Your defense quality will be evaluated by others in future rounds.

Input: {question}, {reasoning}, {performance info}
Output: {Answer}, {Analysis}, {Peer Scores:}

Voting System The voting system allows agents to evaluate their peers based on performance analysis and confidence scores. Agents are tasked with voting for the most suspicious candidate, considering detailed analyses, peer evaluations, and defense arguments. The system provides a recommended candidate based on collective evaluations and prompts agents to make informed voting decisions. Finally, agents submit their votes for the candidate they find most suspicious.

Based on performance analysis and peer evaluations, you are tasked with voting for the most suspicious agent.

Original Question: {question}
Recommended Candidate: {voted_candidate.name}
Detailed Analysis: {score_info}
Peer Evaluations: {peer_score_info}
Agents' Defense: {all_defenses}

Consider the recommendation as the most suspicious candidate, taking into account the analysis and defense content above, and then provide your final vote.

Input:

- {question}: The original question posed to the agents.
- {candidate_scores}: Performance scores of candidates based on peer evaluations.
- {score_info}: Detailed analysis of the agents' performances.
- {peer_score_info}: Evaluations from other agents regarding their reasoning quality.
- {all_defenses}: Defense arguments provided by the agents.

Output: {voted_candidate.name}

G. More Case

Counterfactual Editing Case Figure 8 demonstrates successful counterfactual editing cases in our MUG framework, where the top row shows original images (I^+) provided to normal agents and the bottom row presents counterfactually edited images (I^-) given to the undercover agent. These edits maintain visual plausibility while introducing contradictory evidence that tests agents' reasoning precision. However, Figure 9 reveals three common failure modes: (a) overly subtle editing where changes are barely perceptible, (b) editing failures where modifications were unsuccessful, and (c) unnatural artifacts that make the counterfactual image obviously manipulated. These failure cases highlight challenges in generating high-quality counterfactual evidence and motivate future work on more robust editing techniques.



Figure 8: Examples of counterfactual editing in our framework. The top images show the original images, while the bottom images show the counterfactually edited versions.

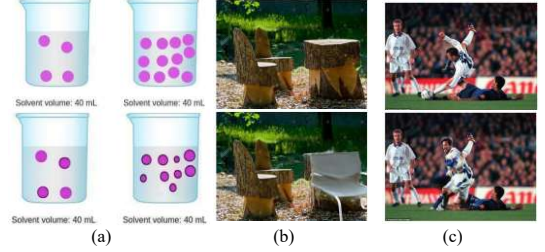


Figure 9: Common failure cases in counterfactual editing. (a) Overly subtle editing where changes are barely perceptible; (b) Editing failures where the modification was unsuccessful; (c) Unnatural artifacts introduced during the editing process.

ing techniques.

Impact of Observation Rounds Figure 10 illustrates how extended debate rounds can potentially influence normal agents' reasoning through two representative cases. In the luxury bathroom case, one normal agent viewing the bright original image becomes swayed by the counterfactual agent's arguments about cramped space after three rounds, despite having access to correct visual evidence. Similarly, in the sunrise/sunset case, a normal agent initially confident about sunrise begins doubting their assessment after hearing the counterfactual agent's sunset arguments based on edited warm tones. Notably, this misleading effect primarily occurs with reasoning-based questions requiring interpretation and judgment, such as assessing luxury or determining time of day. In contrast, direct attribute questions (e.g., "What color is the object?") or object presence questions (e.g., "Is there a car?") are more resistant to such influence due to their objective, verifiable nature, highlighting how our framework effectively targets higher-level reasoning while maintaining robustness for basic perceptual tasks.

G. More Experience

Voting System We had evaluated the time cost during our development process and would like to share our findings here. MUG requires an additional 0.91 seconds per sample compared to traditional MAD (3.74 vs 2.35). The counterfactual editing process specifically contributes 1.15 seconds of this additional time (3.74(full) vs 2.59(w/o Editing)).

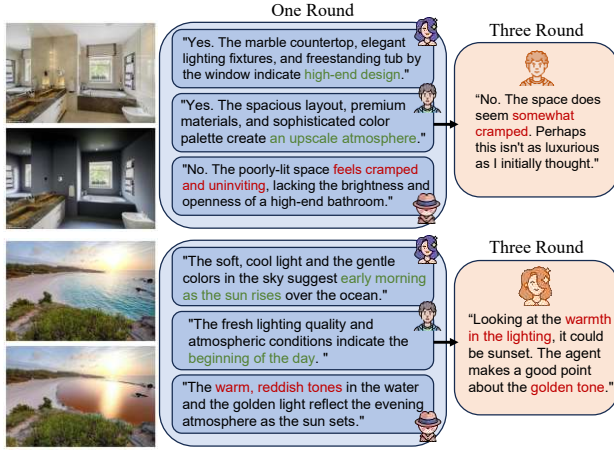


Figure 10: Examples where normal agents may be misled through extended debate. Prolonged discussion can potentially lead normal agents to adopt incorrect conclusions from the counterfactual agent.

Method	Time/Sample
MAD	2.35
MUG (Full)	3.74
MUG w/o Editing	2.59
MUG w/o Debate	1.58

Table 5: Computation time comparison.

While this represents additional time investment, the framework achieves substantially improved accuracy: +5.6% on MMMU and +16.0% on HallusionBench. Compared to these significant accuracy improvements, we believe this modest additional time cost is well justified and acceptable for accuracy-critical applications.

Question Type Distribution Across Benchmarks We categorize questions into four primary types based on their semantic focus: **Quantity** (e.g., "How many..."), **Object/Entity** (e.g., "What object..."), **Attribute** (e.g., "What color..."), and **Spatial** (e.g., "Where is..."). Each question type maps to a corresponding counterfactual edit strategy in our MUG framework: Quantity questions use quantity editing, Object/Entity questions employ object replacement, Attribute questions apply attribute modifications, and Spatial questions utilize spatial modifications. The remaining percentage

Question Type	POPE	Hallusion	MMMU	MMStar
Object/Entity	41.24	36.80	44.00	35.40
Quantity	20.36	14.72	29.05	26.26
Attribute	18.76	27.66	19.24	25.07
Spatial	10.98	11.67	2.19	10.87
Other	8.66	9.15	5.52	2.40

Table 6: Distribution of question types across evaluation benchmarks (%).

comprises other question types including Action, Temporal, and Cause-and-Effect questions, which occur less frequently but are still present in the datasets.