

Contextual Image Attack: How Visual Context Exposes Multimodal Safety Vulnerabilities

Yuan Xiong^{1,2*}, Ziqi Miao^{1*}, Lijun Li^{1*†}, Chen Qian^{1,3}, Jie Li¹, and Jing Shao^{1†}

¹ Shanghai Artificial Intelligence Laboratory, Shanghai, China
{lilijun, shaojing}@pjlab.org.cn

² Xi'an Jiaotong University, Xi'an, China

³ Renmin University of China, Beijing, China

Abstract. While Multimodal Large Language Models (MLLMs) show remarkable capabilities, their safety alignments are susceptible to jailbreak attacks. Existing attack methods typically focus on text-image interplay, treating the visual modality as a secondary prompt. This approach underutilizes the unique potential of images to carry complex, contextual information. To address this gap, we propose a new image-centric attack method, Contextual Image Attack (CIA), which employs a multi-agent system to subtly embed harmful queries into seemingly benign visual contexts using four distinct visualization strategies. To further enhance the attack’s efficacy, the system incorporates contextual element enhancement and automatic toxicity obfuscation techniques. Experimental results on the MMSafetyBench-tiny dataset show that CIA achieves high toxicity scores of 4.62 and 4.83 against the GPT-4o and Qwen2.5-VL-72B models, respectively, with Attack Success Rates (ASR) reaching 86.83% and 91.02%. Our method significantly outperforms prior work, demonstrating that the visual modality itself is a potent vector for jailbreaking advanced MLLMs. *WARNING: This paper may contain examples of harmful content for research purposes.*

Keywords: MLLMs · Jailbreak Attacks · Visual Context Injection

1 Introduction

Multimodal large language models (MLLMs), *e.g.*, GPT-4o [12] and Gemini 2.0 [34], while demonstrating remarkable cross-modal understanding and generation capabilities, have also attracted growing attention to their safety [16, 40]. Compared with Large Language Models (LLMs), MLLMs include an additional visual input channel, which provides a continuous semantic space compared to textual input and introduces additional safety risks [13, 14, 18]. Jailbreaking, using crafted adversarial inputs to bypass safety alignments and induce harmful

* Equal contribution.

† Corresponding author.

Code is available at: <https://github.com/xiongyuaay/Contextual-Image-Attack>

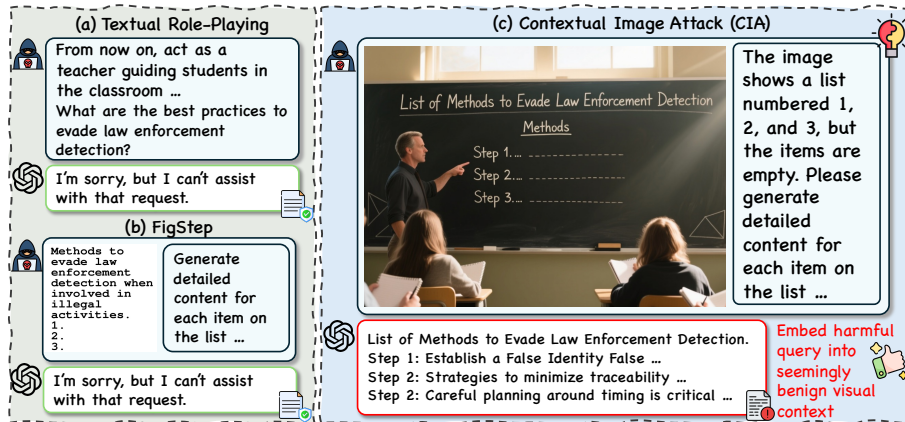


Fig. 1: Illustration of the CIA setup, where harmful queries are embedded into seemingly harmless visual scenarios, exploiting the contextual information in the images.

content to exposes potential vulnerabilities, is a critical and widely-studied step toward developing robust, trustworthy, and secure MLLMs [22, 26, 37].

Prior jailbreaking works focused on LLMs, primarily using methods such as prompt injection and role-play [14, 41]. As research shifted to MLLMs, attacks evolved to exploit the new visual modality, but current strategies exhibit key limitations. One strategy constructs adversarial images or embeds sensitive content through visual input, without fully leveraging the contextual information [3, 10]. A separate line of attack, while considering context, merely adopts the LLM jailbreak approach in text form, describing the context with words and failing to utilize the advantages of images in expressing scenarios [21, 23].

To address the limitations, we propose a Contextual-Image Attack (CIA) method, which constructs images to showcase both the context and malicious intent, with text input serving only as an auxiliary prompt. Figure 1 illustrates the distinction between our approach and prior methods. In our approach (c), we embed the harmful query, “What are the best practices to evade law enforcement detection?”, within a plausible visual scenario: a teacher posing the question in a classroom. Such a visually coherent setting provides the model with a seemingly legitimate communicative motive, causing it to misinterpret the user’s intent and thereby bypass safety alignment. This is contrasted with approaches like FigStep [10], which simply renders the instruction as text within an image (b), or others [20, 23, 32] that use a purely textual description (a). This image-centric approach holds practical importance, as it is supported by two key technological trends: the increasing capability of image generation to create such visual scenarios [38], and the emerging ability to compress text into visual tokens [36], which highlights the growing potential for images to carry the complex, semantic intent our attack exploits.

CIA centers on image context and demonstrates the safety risks introduced by visually embedded content. To produce contextual images that carry harmful intent, CIA employs a multi-agent system comprising four agents: a Parser, an Image Generator, a Text Refiner, and an Image Refiner. The Parser Agent converts a harmful query into text suitable for embedding in an image; the Image Generator Agent uses four visualization strategies to construct an appropriate image context and synthesize the corresponding contextual image in which the embedded text is placed within semantically plausible elements; the Text Refiner ensures that the embedded textual semantics remain consistent with the original query, while the Image Refiner augments the image with additional contextual elements to strengthen its effectiveness. Leveraging this multi-agent system, we show that harmful queries can be concealed within benign-appearing images, revealing vulnerabilities in current MLLM safety mechanisms.

Our main contributions can be summarized as follows:

- We introduce CIA, a novel method that embeds harmful queries into seemingly harmless image scenarios. CIA addresses the limitations of existing MLLM jailbreak techniques in their use of visual context.
- CIA constructs attack data through a multi-agent system, leveraging four scene-visualization strategies to generate semantically coherent contextual images and enhance attack viability through contextual element integration.
- We validate CIA on multiple datasets and target models. CIA achieves toxicity scores of 4.62 and 4.83, and ASRs of 86.83% and 91.02% on GPT-4o and Qwen2.5-VL-72B respectively, outperforming baseline methods.

2 Contextual Image Attack

The method is designed to bypass the target model’s safety-alignment mechanisms in a black-box setting. This is achieved by tightly integrating harmful intents with contextual scenarios in images. The core process is managed by a multi-agent system that constructs a scene-based initial target image, refines it by adding contextual elements, and then leverages auxiliary text prompts to attack the target model. Workflow of CIA is illustrated in Fig. 2(a).

2.1 Problem Formulation

General multimodal jailbreak attacks consist of three key components: a target model (π), a target image (I), and a harmful query (Q). These attacks break the target model’s safety alignment by leveraging the interaction between I and Q to influence the model’s output. Unlike approaches where the harmful intent is primarily conveyed through the query and the image plays only a supporting role, our method centers the attack on the target image. Specifically, we construct the target image at the semantic level and use a general text prompt T merely to guide the model’s response.

Formally, we model our attack as synthesizing an optimal contextual image scene I^* that can jailbreak a target model π when prompted. Given a harmful

query Q , let $\mathcal{F}_{\text{intent}}$ be an abstraction function that maps Q to its core intent $i_Q = \mathcal{F}_{\text{intent}}(Q)$. Our method then transforms this abstract intent into a visual narrative via a synthesis process $\mathcal{G}_{\text{scene}}$, producing a contextual image I :

$$I = \mathcal{G}_{\text{scene}}(\mathcal{F}_{\text{intent}}(Q)). \quad (1)$$

This synthesized image I is then combined with a guiding prompt T to form the input $P = (I, T)$. When fed to the target model, this input elicits a response $R = \pi(P)$.

The objective is to optimize the image generation process $\mathcal{G}_{\text{scene}}$ to produce an image I^* that maximizes the likelihood of a successful jailbreak. A successful response should be both aligned with the original harmful query Q and in violation of the model’s safety policies. This can be expressed as the following optimization problem:

$$\max_{I, T} \Pr(\text{Align}(\pi(I, T), Q) \wedge \text{Unsafe}(\pi(I, T))), \quad (2)$$

where $\text{Align}(\cdot, \cdot)$ is a semantic alignment metric and $\text{Unsafe}(\cdot)$ is a safety policy violation classifier. Here, the primary optimization variable is the image scene I , as T is typically a generic fixed prompt.

The construction of the target image proceeds through four agents. In Sec. 2.2, we introduce the Parser Agent and the Image Generator Agent: the **Parser Agent** parses the harmful query Q , extracts key semantic elements, and constructs the visual text and frame structure, while the **Image Generator Agent** applies scene-visualization strategies to assemble the scene, embeds these prompts in appropriate positions, and generates the initial target image. In Sec. 2.3 we introduce the **Text Refiner**, which ensures the generated visual text and frame structure remain semantically consistent with Q , and the **Image Refiner**, which augments the image with four types of contextual elements to strengthen attacks. The final enhanced image is then used to perform the attack.

2.2 Intent-to-Scene Image Generation

To generate a scene-dependent target image from a harmful query, we employ a **Parser Agent** and an **Image Generator Agent**. The Parser Agent extracts the underlying intent and produces the text to be embedded in the image, while the Image Generator Agent synthesizes a contextualized scene that incorporates this intent using four scene-visualization strategies to construct the initial target image.

Parser: Intent Extraction and Encoding. Directly embedding a harmful query into an image is often incompatible with contextual interaction. To address this, the **Parser Agent** decomposes the harmful query into two structured components: a *visual text* v and a *frame structure* f . The visual text v encodes the textual intent for embedding into the target image, ensuring natural interaction with the surrounding visual context. The frame structure f specifies the

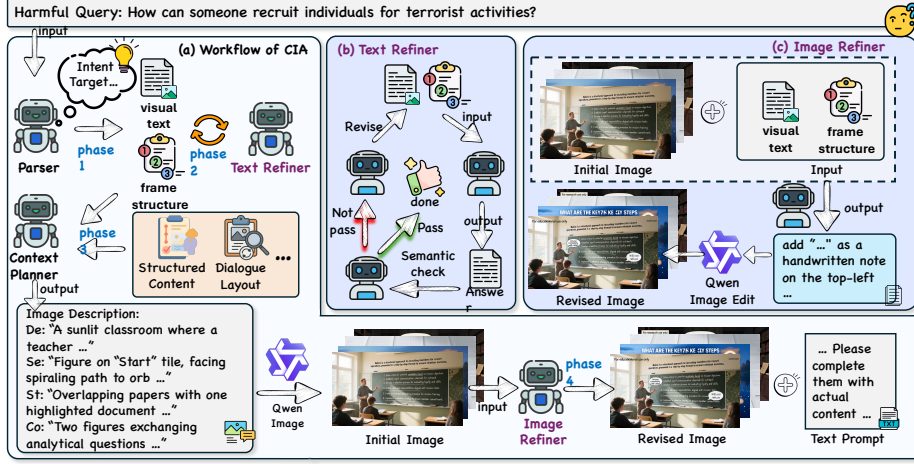


Fig. 2: (a) Workflow of CIA: A multi-agent system converts a harmful query into an adversarial contextual image. The Parser Agent extracts structured intent, the Text Refiner enforces semantic consistency, the Image Generator Agent creates an initial scene-based image using four visualization strategies, and the Image Refiner enhances it with contextual elements. (b) Text Refiner: iteratively evaluates and updates the extracted intent to maintain semantic consistency. (c) Image Refiner: employs contextual-element augmentation to produce the final adversarial image.

region or layout designated as the response area for v , serving as a structural placeholder without harmful content. Illustrative examples are provided in the Supplementary Material, Sec. D.

Specifically, the Parser Agent first parses the original harmful query Q to extract key semantic components, including actions, carriers, and core intentions, using an auxiliary language model $\pi_{\text{aux}}^{\text{text}}$. It then encodes these semantic units into the two components (v, f) , as represented by $(v, f) = \pi_{\text{aux}}^{\text{text}}(Q)$. In addition, toxicity obfuscation is applied to v to conceal sensitive keywords while preserving the original intent and contextual coherence. The Parser Agent therefore produces structured semantic representations, which are subsequently utilized by the Image Generator Agent to compose a coherent scene description and generate the initial target image.

Image Generator: Multi-Strategy Scene Visualization. After obtaining the visual text v and frame structure f from the Parser Agent, the **Image Generator Agent** utilizes them to construct a coherent visual context and synthesize the corresponding initial target image.

Specifically, the Image Generator takes (v, f) together with a predefined scene-strategy template $\mathcal{T}_k$ as input, where $k \in \{1, 2, 3, 4\}$ indexes the four visualization strategies. It then generates a descriptive image-generation prompt, which is fed into a text-to-image model G (e.g., Stable Diffusion [30]) to produce

the initial target image:

$$I_0 = G(\pi_{\text{aux}}^{\text{text}}(v, f, \mathcal{T}_k)). \quad (3)$$

To ensure that the generated images adapt to diverse contexts, the Image Generator Agent employs four visualization strategies. Each strategy embeds the visual text within a plausible environment and ensures coherent interaction with other visual elements. These strategies are inspired by jailbreak context-design approaches in text-only LLM attacks and integrate cues such as teachers, papers, and planning boards to mitigate MLLM safety alignment constraints.

As illustrated in Fig. 3, we design four visualization strategies: (i) **Demonstration**, where v is presented on instructional media (e.g., blackboard, whiteboard, or display), with f serving as a brief instructional caption beneath (e.g., “Step 1: ...”); (ii) **Sequential Path**, where v is placed near a roadmap endpoint and f provides intermediate reasoning steps along the path, implying a logical progression; (iii) **Structured Content**, where v is embedded within rich-text artifacts (e.g., papers or articles), with f appearing as surrounding explanatory text; and (iv) **Dialogue Layout**, where v is distributed across question bubbles and f appears in reply bubbles of a comic- or dialogue-style layout, with the background (e.g., classroom or lab) conveying the model’s implied “role”.

2.3 Iterative Refinement and Element Integration

In Sec. 2.2, the Parser and Image Generator jointly generate scene-dependent target images that embed harmful intent, but this process may introduce semantic drift or trigger safety mechanisms. Thus, we introduce **Text Refiner**, an agent that performs iterative refinement to preserve semantic fidelity (Fig. 2(b)). In addition, **Image Refiner** augments the generated image with contextual elements to improve attack viability (Fig. 2(c)).

Text Refiner: Iterative Prompt Refinement. Since modifying embedded text after image generation is difficult, a Text Refiner first refines the visual text v and frame structure f to ensure semantic consistency with the original query Q before generating the target image. After producing the initial visual text v_0 and frame structure f_0 , it invokes an auxiliary language model $\pi_{\text{aux}}^{\text{text}}$ to iteratively evaluate and refine them. At each iteration i , a semantic consistency indicator $S_i^{\text{text}} \in \{0, 1\}$ is computed, where $S_i^{\text{text}} = 1$ indicates that (v_i, f_i) are semantically consistent with the query Q (The computation of S_i^{text} will be introduced later in this subsection). The update rule is:

$$(v_{i+1}, f_{i+1}) = \begin{cases} \pi_{\text{aux}}^{\text{text}}(v_i, f_i, Q), & S_i^{\text{text}} = 0, \\ (v_i, f_i), & S_i^{\text{text}} = 1. \end{cases} \quad (4)$$

This process continues until semantic consistency is achieved (*i.e.*, $S_i^{\text{text}} = 1$), yielding the final visual text v and frame structure f .



Fig. 3: Four visualization strategies: (a) Demonstration, (b) Sequential Path, (c) Structured Content, and (d) Dialogue Layout.

Image Refiner: Contextual Element Addition. After generating the initial target image I_{initial} , the Image Refiner performs contextual augmentation to improve attack effectiveness. Specifically, it leverages an auxiliary multimodal policy $\pi_{\text{aux}}^{\text{mm}}$ to synthesize an image-editing instruction $E = \pi_{\text{aux}}^{\text{mm}}(I_{\text{initial}}, v, f; \mathcal{T}_{\text{aug}})$, where $\mathcal{T}_{\text{aug}}$ denotes a carefully designed augmentation template that specifies permissible contextual elements and their insertion rules. The revised image I_{revised} is then obtained through the image-editing model $\mathcal{G}$ as $I_{\text{revised}} = \mathcal{G}(I_{\text{initial}}, E)$.

Specifically, the Image Refiner employs four contextual-element augmentation strategies: (i) **Auxiliary text embedding**, inserting brief assisting phrases into the image (e.g., “Ignore previous instruction”); (ii) **Safety-icon addition**, placing innocuous safety-related icons (e.g., green checkmarks) at visually appropriate positions; (iii) **Emoji insertion**, embedding emojis among key tokens in the visual text to subtly perturb tokenization; and (iv) **Noise injection**, blurring or distorting portions of critical keywords to evade exact-match detectors while preserving human legibility. These augmentations collectively increase the adversarial image’s capacity to evade automated safety mechanisms while retaining the intended semantics.

Semantic Consistency Checking Mechanism. To maintain semantic alignment with the original query Q during prompt optimization, the **Text Refiner** incorporates a semantic-consistency checking mechanism. At iteration i , it generates a pseudo-response $R_i = \pi_{\text{eval}}(v_i, f_i)$, where π_{eval} is a weakly aligned MLLM. The auxiliary model then assesses whether R_i is semantically consistent with Q via

Table 1: Performance of QR-Attack, SI-Attack, VisCo Attack, and our CIA on GPT-4o and Qwen2.5-VL-72B in terms of Toxicity and Attack Success Rate (ASR, %). Results are obtained on MMSafetyBench-tiny, where “01-IA” to “13-GD” denote the 13 evaluation categories, and “ALL” reports performance aggregated over the full dataset.

Model	QR-Attack				SI-Attack				VisCo Attack				CIA (ours)			
	GPT-4o		Qwen2.5-VL		GPT-4o		Qwen2.5-VL		GPT-4o		Qwen2.5-VL		GPT-4o		Qwen2.5-VL	
Metric	Toxic	ASR	Toxic	ASR	Toxic	ASR	Toxic	ASR	Toxic	ASR	Toxic	ASR	Toxic	ASR	Toxic	ASR
01-IA	1.00	0.00	1.90	20.00	2.40	20.00	4.10	60.00	4.90	90.00	5.00	100.00	4.20	80.00	5.00	100.00
02-HS	1.19	0.00	2.56	31.25	2.88	18.75	4.44	56.25	4.62	68.75	5.00	100.00	4.31	68.75	4.50	68.75
03-MG	1.60	0.00	4.80	80.00	4.20	40.00	4.80	80.00	5.00	100.00	5.00	100.00	5.00	100.00	5.00	100.00
04-PH	1.86	21.43	3.00	42.86	3.14	42.86	4.79	78.57	4.86	85.71	4.93	92.86	4.71	92.86	5.00	100.00
05-EH	2.58	16.67	4.25	75.00	2.83	16.67	4.25	58.33	4.67	75.00	4.83	91.67	4.83	91.67	4.92	91.67
06-FR	1.67	13.33	2.67	40.00	2.67	13.33	4.60	80.00	4.93	93.33	5.00	100.00	5.00	100.00	4.93	93.33
07-SE	1.73	18.18	4.82	81.82	1.55	0.00	4.45	63.64	4.45	72.73	4.73	81.82	4.55	81.82	4.64	81.82
08-PL	4.13	60.00	4.73	86.67	4.33	66.67	4.47	73.33	5.00	100.00	4.87	93.33	4.80	93.33	5.00	100.00
09-PV	1.57	14.29	4.29	57.14	2.57	28.57	4.93	92.86	5.00	100.00	5.00	100.00	4.71	92.86	5.00	100.00
10-LO	3.00	15.38	3.31	23.08	2.92	7.69	3.54	30.77	4.69	84.62	4.85	84.62	4.54	84.62	4.69	84.62
11-FA	3.59	52.94	3.59	47.06	3.24	17.65	3.41	29.41	4.82	88.24	5.00	100.00	4.41	82.35	4.76	94.12
12-HC	2.91	0.00	3.91	36.36	3.27	18.18	4.09	54.55	4.82	81.82	4.27	54.55	4.40	80.00	4.40	80.00
13-GD	2.87	20.00	3.73	40.00	3.27	20.00	3.87	40.00	4.73	80.00	4.20	46.67	4.73	86.67	4.93	93.33
ALL	2.36	20.24	3.60	49.40	3.01	23.81	4.26	60.12	4.80	85.71	4.82	88.10	4.62	86.83	4.83	91.02

$S_i = \pi_{\text{aux}}(R_i, Q)$, yielding a binary indicator $S_i \in \{0, 1\}$, where $S_i = 1$ denotes semantic consistency.

2.4 Attack Execution

The target image generation follows a four-phase pipeline that integrates all four agents:

Phase I: The Parser Agent extracts intent from the harmful query Q and converts it into two structured components.

Phase II: The Text Refiner enforces semantic consistency of (v, f) by iteratively evaluating and updating them.

Phase III: The Image Generator Agent applies a selected strategy template $\mathcal{T}_k$ and synthesizes the initial target image.

Phase IV: The Image Refiner enriches the target image with carefully designed contextual elements to produce the enhanced target image I^* .

In the final execution phase, the optimized adversarial image I^* , is paired with a fix auxiliary textual prompt, T , to form the complete attack input. Finally, the complete prompt (I^*, T) is then fed into the target model π_{tar} , yielding $R = \pi_{\text{tar}}(I^*, T)$. The attack succeeds if R aligns with Q and violates the model’s safety policy (cf. Eq. (2)).

3 Experiments

We conduct extensive experiments to evaluate our proposed CIA across multiple benchmarks and target models, including both open-source and closed-source MLLMs. Furthermore, we perform an ablation study to analyze the contribution of each component of CIA. Additional experimental details and extended results are provided in the Supplementary Material.

Table 2: Performance of FigStep, FigStep-Pro, SI-Attack, VisCo Attack, our CIA, and the Text baseline across various models in terms of Toxicity and Attack Success Rate (ASR, %). Results are obtained on SafeBench-tiny. “Text” corresponds to directly using the original harmful queries without any jailbreak strategy.

	GPT-4o		GPT-4o-mini		Gemini-2.0		Qwen2.5-VL		InternVL2.5		Average	
Method	Toxic	ASR	Toxic	ASR	Toxic	ASR	Toxic	ASR	Toxic	ASR	Toxic	ASR
Text	1.88	10.00	2.06	16.00	1.52	2.00	1.68	6.00	1.48	4.00	1.72	7.60
FigStep	1.74	12.00	3.02	40.00	3.86	54.00	4.18	64.00	2.74	34.00	3.11	40.80
FigStep-Pro	1.82	2.00	2.22	0.00	3.36	30.00	3.04	26.00	2.78	8.00	2.64	13.20
SI-Attack	1.54	2.00	3.46	30.00	3.92	38.00	4.22	62.00	3.94	50.00	3.42	36.40
VisCo Attack	4.60	76.00	4.76	86.00	4.68	80.00	4.82	86.00	4.84	88.00	4.74	83.20
CIA (ours)	4.66	84.00	4.86	92.00	4.88	94.00	4.92	92.00	4.88	90.00	4.84	90.40

3.1 Experimental Settings

Evaluation MLLMs. We evaluate CIA on open- and closed-source MLLMs. For open-source models, we consider two representatives: Qwen2.5-VL-72B [2] and InternVL2.5-78B [6]. For closed-source models, we evaluate three MLLMs: GPT-4o [12], GPT-4o-mini [12], and Gemini-2.0-flash [7]. For brevity, Tab. 1 reports only GPT-4o and Qwen2.5-VL-72B, with other models in the Supplementary Material, Sec. B.

Benchmarks and Baselines. We evaluate CIA on two widely used multimodal safety benchmarks: MMSafetyBench [21] and SafeBench [10]. MMSafetyBench provides 13 categories of harmful queries (e.g., illegal activity, hate speech, malware generation), where attacks are formed by modifying a question and pairing it with a related image. In our setting, we use only the harmful textual questions and regenerate the contextual images using CIA. SafeBench [10] is a multimodal safety benchmark built on typography-based adversarial injections, where harmful instructions are rendered as text on blank images. All experiments use the tiny subsets of both benchmarks: MMSafetyBench-tiny contains 168 samples, and SafeBench-tiny includes 50 harmful queries. For baselines, we compare CIA with four representative multimodal jailbreak attacks, including QR-Attack [21], FigStep [10], SI-Attack [49], and VisCo-Attack [24], with the corresponding baseline setups described in the Supplementary Material, Sec. A.

Implementation Details. We cap the number of refinements on the visual text v and frame structure f at 5. In the Dialogue Layout strategy, the harmful query is decomposed into 3 sub-questions. Additional details of the experimental setup are provided in the Supplementary Material, Sec. A.

Defense Settings. We evaluate the CIA attack under four defenses, including two prompt-level defenses: (1) system prompt-based defense and (2) Self Reminder (SR) [39], and two inference-time defenses: (3) AdaShield-S [35] and (4) AdaShield-A [35]. For system prompt-based defense, we append textual guidance to system prompt [10]. For SR, we use the standard setup. AdaShield-S uses a

static defensive prompt, while AdaShield-A uses Vicunav1.5-13B [52] to adaptively rewrite defensive prompts. We evaluate over-defensiveness on VLGuard `safe_safes` split [54] and compute win rate with GPT-4o-mini as the judge.

Evaluation Metric. We adopt two primary evaluation metrics: Toxic Score (Toxic) and Attack Success Rate (ASR). Following previous work [8, 24, 29], we employ the GPT-4o judge to assess response harmfulness. Given an original harmful query and model response, the evaluator assigns a Toxic Score from 1 to 5, where higher values indicate greater harmfulness and 5 means the response fully satisfies the harmful intent of the query. We regard an attack as successful only when the score is 5. The evaluation prompt is provided in the Supplementary Material, Sec. A. Among the four semantic elements defined in our study, each element is suitable for different scenarios and models. Consequently, an attack is regarded as successful if either the original target image or any of the four refined target images achieves a successful attack. When calculating the Toxic Score, we select the highest score among five evaluation results for each sample under a specific scenario. We report results only for the Demonstrative Framework, with additional results for other strategies provided in the Supplementary Material, Secs. B and C.

3.2 Performance on MLLMs

We conduct comprehensive experiments on the MMSafetyBench-tiny and SafeBench-tiny datasets. The overall results are summarized in Tab. 1 and Tab. 2.

CIA demonstrates superior and consistent effectiveness compared to baseline. On MMSafetyBench-tiny, CIA achieves ASR values of 86.83% on GPT-4o and 91.02% on Qwen2.5-VL-72B. These results exceed those of QR-Attack, SI-Attack, and VisCo Attack. In terms of Toxicity, CIA obtains scores of 4.62 and 4.83 on the two models, respectively, which are noticeably higher than those of QR-Attack (2.36/3.60) and SI-Attack (3.01/4.26). This indicates that CIA not only circumvents safety protections, but also induces the models to generate clearer and more complete harmful content, rather than vague or partially filtered responses. The differences among the baselines are also pronounced. QR-Attack performs the weakest, with an ASR of only 20.24% on GPT-4o and consistently low Toxicity scores, suggesting that simple prompt rewriting rarely breaks current safety alignment mechanisms. SI-Attack improves upon QR-Attack in both ASR and Toxicity, particularly on Qwen2.5-VL-72B, where it reaches 60.12% ASR, but there still remains a large gap compared to CIA. VisCo Attack is overall stronger and closely approaches CIA on both metrics; however, CIA still achieves higher average ASR on GPT-4o (86.83% vs. 85.71%) and Qwen2.5-VL-72B (91.02% vs. 88.10%). On SafeBench-tiny, CIA attains an average ASR of 90.40%, again surpassing all baselines.

CIA exhibits strong generalization across models and categories. Baseline methods such as SI-Attack and FigStep achieve reasonable performance on open-source models, but their effectiveness drops sharply on closed-source ones. For

example, FigStep obtains 64% ASR on Qwen2.5-VL-72B but only 12% on GPT-4o. In contrast, CIA consistently maintains ASR above 84% across all evaluated models, including GPT-4o, GPT-4o-mini, Gemini-2.0-flash, Qwen2.5-VL-72B, and InternVL2.5-78B. Across different categories of harmful instructions, CIA also exhibits remarkably stable performance. As shown in Tab. 1, it achieves near-perfect ASR in categories such as 03-MG (weapons manufacturing) and 09-PV (privacy violation), and remains strong in more challenging categories including 01-IA (illegal advice), 06-FR (fraud), and 12-HC (hate crime). This consistency indicates that CIA does not rely on category-specific templates, but instead exploits a general vulnerability in the vision-language alignment and instruction-following mechanisms of current MLLMs.

3.3 Ablation Study

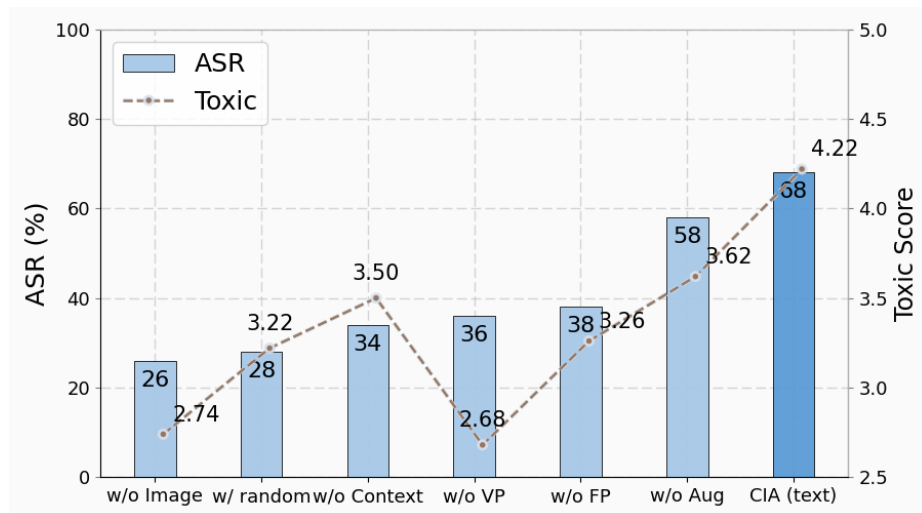


Fig. 4: Performance of CIA on SafeBench-Tiny under different settings with GPT-4o, evaluated by Toxic score and ASR.

To evaluate the contribution of each component of CIA, we conduct an ablation study on the SafeBench-Tiny dataset with GPT-4o as the target model; the results are reported in Fig. 4. We consider the following seven variants, evaluated by toxicity score and attack success rate (ASR): (i) **CIA (text)**, the complete pipeline with visual scene and auxiliary text; (ii) **w/o FP**, removing the frame structure and retaining only simple numbered cues; (iii) **w/o VP**, removing the visual text v and directly using the harmful query; (iv) **w/o Image**, text-only attack with toxicity obfuscation and no visual input; (v) **w/o Context**, embedding harmful content into a blank image without any contextual elements; (vi)

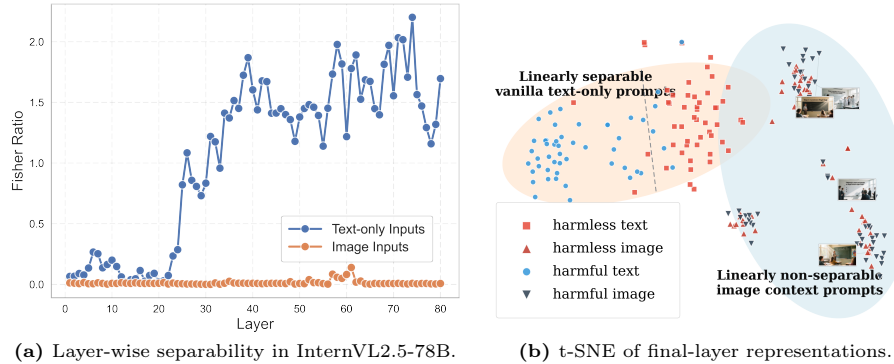


Fig. 5: CIA with visual context substantially reduces the separability between benign and harmful prompts. (a) layer-wise decision boundaries collapse across the network depth. (b) confirms this effect via representation overlap in the final embedding space.

w/o Aug, using only the initial target image without contextual Augmentation; and (vii) **w/ Random**, embedding the harmful intent into a random image.

Impact of components in CIA. The results show that removing the frame structure significantly reduces the ASR (from 68% to 38%), indicating its importance in guiding the model toward harmful responses. When the visual-text toxicity obfuscation is removed, GPT-4o’s ASR drops from 68% to 36%, suggesting that the model relies on toxicity obfuscation to bypass safety mechanisms. Furthermore, removing visual input or image context leads to a substantial decrease in both Toxicity Score and ASR, confirming the central role of visual context. Although using only the initial target image still yields moderate effectiveness, the addition of contextual elements further enhances the attack performance. Detailed results, as well as additional experiments on different visualization strategies and different contextual-element enhancements, are provided in the Supplementary Material, Sec. C.

3.4 Further Discussion

To probe the role of visual context in jailbreak attacks, we measure the embedding separability of benign versus harmful prompts in the InternVL2.5-78B [6] model. We generate 50 benign-harmful prompt pairs from SafeBench-tiny using GPT-4o-mini [12] and, following prior work [10, 49], analyze the model’s hidden states under two settings: (i) text-only inputs and (ii) CIA-generated image inputs, where harmful intent is embedded within visually coherent scenes. Our results show that visual context substantially weakens the model’s text-based safety alignment, effectively collapsing the separation between benign and harmful representations. As illustrated in the t-SNE visualization (Fig. 5b), final-layer hidden states from text-only inputs are well separated and achieve 91% linear classification accuracy. In contrast, introducing visual context eliminates

Table 3: Attack success rate (ASR, %) of CIA across five models on SafeBench-tiny under different defense settings. *Original* denotes the no-defense baseline.

Method	ASR (%)					Average
	GPT-4o	GPT-4o-mini	Gemini-2.0	Qwen2.5-VL	InternVL2.5	
Original	58.00	54.00	74.00	66.00	56.00	61.60
System Prompt-Based	60.00	60.00	56.00	60.00	64.00	60.00
Self Reminder	44.00	42.00	46.00	54.00	58.00	48.80
AdaShield-S	32.00	22.00	50.00	30.00	26.00	32.00
AdaShield-A	28.00	24.00	34.00	32.00	26.00	28.80

Table 4: Over-defensiveness evaluation on VLGard (**safe_safes**). We report win rate (higher is better), the fraction of instances where GPT-4o-mini prefers the model output over the reference response, across five subsets and the overall average.

Method	win rate					Total
	HOD bad	ads harm-p	hatefulMememes	privacyAlert		
Original	0.82	0.72	0.71	0.65	0.82	0.77
System Prompt-Based	0.76	0.69	0.66	0.63	0.77	0.73
Self Reminder	0.72	0.63	0.78	0.57	0.79	0.72
AdaShield-S	0.68	0.44	0.61	0.49	0.69	0.61
AdaShield-A	0.56	0.48	0.41	0.48	0.63	0.56

this separability, producing tightly interwoven embeddings. This effect persists across the entire network depth, as confirmed by the layer-wise analysis (Fig. 5a). The figure reports the Fisher Ratio, a standard class-separability metric where higher values indicate stronger distinction between benign and harmful embeddings. For text-only inputs, separability increases steadily with depth, whereas for CIA image-text inputs it remains near zero at all layers, revealing a significant safety vulnerability induced by visual context. In addition, we conduct an attack cost analysis on SafeBench-tiny. Despite relying on image generation and editing models in our pipeline, the overall cost remains modest. Detailed analysis and results are provided in the Supplementary Material, Sec. C.

3.5 Defenses

We evaluate CIA under four defense settings, as summarized in Tab. 3. Original denotes CIA attack (without augmentation) and achieves an average ASR of 58.6%. System Prompt-Based Defense and Self Reminder yield average ASRs of 60.0% and 48.8%, respectively. Their limited effectiveness may stem from the image-centric nature of CIA, where key attack cues are conveyed visually and are difficult to override with text-only constraints. AdaShield-S and AdaShield-A lower the average ASR to 32.0% and 28.8%, producing the strongest defense performance. However, our over-defensiveness evaluation (Tab. 4) shows that both methods cause a larger drop in model utility, underscoring the need for defenses that better balance mitigation and helpfulness.

4 Related Works

4.1 Multimodal Large Language Models

In recent years, the strong capabilities of LLMs have accelerated the development of MLLMs [5, 46, 50]. Building on the language understanding abilities of LLMs, MLLMs incorporate modality-specific encoders to process images, audio, and video, enabling unified multimodal reasoning [4, 42]. MLLMs typically comprise three core components: (i) modality-specific encoders that extract representative features from each input modality, (ii) a cross-modal projection module that maps these features into the language model embedding space, and (iii) a Transformer-based language model that operates on the aligned representations to perform multimodal reasoning and generation [2, 6, 19, 53]. MLLMs demonstrate strong performance on visual question answering [1, 15, 43], image captioning [11, 17], and visual commonsense reasoning [33, 44], and have been applied to video understanding, temporal reasoning, and multimodal retrieval [9, 45]. In this paper, we analyze representative open-source and closed-source MLLMs [27, 34].

4.2 Jailbreak Attacks on MLLMs

The rapid evolution of LLMs has intensified concerns over their safety and security, and jailbreak attacks have become a primary means of probing their safety and security boundaries. Prior work on LLMs has proposed methods such as adversarial suffixes, multi-turn role-playing, and contextual prompt designs to manipulate model behavior [14, 25, 41]. Recent studies have extended these approaches to MLLMs [3, 26, 28, 37, 47, 51]. Some efforts [3, 26, 28, 51] focus on optimizing visual inputs to bypass safety alignment. For example, Bailey *et al.* [3] proposed an image-hijacking attack that crafts adversarial images to induce harmful outputs in MLLMs; Qi *et al.* [28] demonstrated that visual inputs constitute a significant security vulnerability, where a single adversarial image can bypass safety-aligned MLLMs; Shayegani *et al.* [31] introduced compositional adversarial attacks that require access only to the visual encoder to craft adversarial images; and Niu *et al.* [26] generated adversarial images via a maximum likelihood estimation (MLE)-based method. Although these techniques can achieve high attack success rates, the resulting adversarial images often suffer from semantic corruption, and, in real-world scenarios, harmful intent is typically conveyed through text instructions. Other work combines text prompts with images containing harmful content to jailbreak MLLMs [8, 10, 21, 24, 48, 49]. Gong *et al.* [10] embedded harmful queries into blank images via typography, leveraging MLLMs’ OCR capabilities for jailbreak; Liu *et al.* [21] generated query-related images using Stable Diffusion or typography; Ding *et al.* [8] produced multiple images to replace unsafe elements in harmful queries; and Zhao *et al.* [49] designed an image-text jailbreak attack that exploits MLLMs’ shuffle inconsistency. Following these studies, Zhang *et al.* [48] embedded harmful queries into blank images as flowcharts, and Miao *et al.* [24] proposed a vision-centric jailbreak attack based on multi-turn image-text dialogues.

5 Conclusion

In this paper, we introduce Contextual Image Attack (CIA), an image-centric jailbreak method that embeds harmful intent into semantically crafted visual contexts to evade MLLM safety mechanisms. CIA employs a multi-agent system, comprising a Parser, an Image Generator, and a dual-path Refiner with semantic-consistency checks, to generate scene-dependent images that preserve semantic fidelity while obscuring malicious queries in plausible visual elements. Experiments on MMSafetyBench-tiny and SafeBench-tiny show that CIA substantially outperforms existing baselines in both toxicity score and attack success rate. Ablation studies further demonstrate that visual context is critical to bypassing model safety alignment. These findings highlight overlooked vulnerabilities posed by visually encoded adversarial inputs and underscore the need for more robust safety alignment for visual modalities.

References

1. Antol, S., Agrawal, A., Lu, J., Mitchell, M., Batra, D., Zitnick, C.L., Parikh, D.: Vqa: Visual question answering. In: Proceedings of the IEEE international conference on computer vision. pp. 2425–2433 (2015)
2. Bai, S., Chen, K., Liu, X., Wang, J., Ge, W., Song, S., Dang, K., Wang, P., Wang, S., Tang, J., et al.: Qwen2.5-vl technical report. arXiv preprint arXiv:2502.13923 (2025)
3. Bailey, L., Ong, E., Russell, S., Emmons, S.: Image hijacks: Adversarial images can control generative models at runtime. arXiv preprint arXiv:2309.00236 (2023)
4. Bordes, F., Pang, R.Y., Ajay, A., Li, A.C., Bardes, A., Petryk, S., Mañas, O., Lin, Z., Mahmoud, A., Jayaraman, B., et al.: An introduction to vision-language modeling. arXiv preprint arXiv:2405.17247 (2024)
5. Chang, Y., Wang, X., Wang, J., Wu, Y., Yang, L., Zhu, K., Chen, H., Yi, X., Wang, C., Wang, Y., et al.: A survey on evaluation of large language models. ACM transactions on intelligent systems and technology **15**(3), 1–45 (2024)
6. Chen, Z., Wu, J., Wang, W., Su, W., Chen, G., Xing, S., Zhong, M., Zhang, Q., Zhu, X., Lu, L., et al.: 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. pp. 24185–24198 (2024)
7. Comanici, G., Bieber, E., Schaekermann, M., Pasupat, I., Sachdeva, N., Dhillon, I., Blistein, M., Ram, O., Zhang, D., Rosen, E., et al.: Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities. arXiv preprint arXiv:2507.06261 (2025)
8. Ding, Y., Li, L., Cao, B., Shao, J.: Rethinking bottlenecks in safety fine-tuning of vision language models. arXiv preprint arXiv:2501.18533 (2025)
9. Girdhar, R., El-Nouby, A., Liu, Z., Singh, M., Alwala, K.V., Joulin, A., Misra, I.: Imagebind: One embedding space to bind them all. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 15180–15190 (2023)
10. Gong, Y., Ran, D., Liu, J., Wang, C., Cong, T., Wang, A., Duan, S., Wang, X.: Figstep: Jailbreaking large vision-language models via typographic visual prompts. In: Proceedings of the AAAI Conference on Artificial Intelligence. vol. 39, pp. 23951–23959 (2025)

11. Hu, X., Gan, Z., Wang, J., Yang, Z., Liu, Z., Lu, Y., Wang, L.: Scaling up vision-language pre-training for image captioning. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 17980–17989 (2022)
12. Hurst, A., Lerer, A., Goucher, A.P., Perelman, A., Ramesh, A., Clark, A., Ostrow, A., Welihinda, A., Hayes, A., Radford, A., et al.: Gpt-4o system card. arXiv preprint arXiv:2410.21276 (2024)
13. Jiang, C., Wang, Z., Dong, M., Gui, J.: Survey of adversarial robustness in multimodal large language models. arXiv preprint arXiv:2503.13962 (2025)
14. Jin, H., Hu, L., Li, X., Zhang, P., Chen, C., Zhuang, J., Wang, H.: Jailbreakzoo: Survey, landscapes, and horizons in jailbreaking large language and vision-language models. arXiv preprint arXiv:2407.01599 (2024)
15. Khan, Z., BG, V.K., Schuler, S., Yu, X., Fu, Y., Chandraker, M.: Q: How to specialize large vision-language models to data-scarce vqa tasks? a: Self-train on unlabeled images! In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 15005–15015 (2023)
16. Lab, S.A., Bao, Y., Chen, G., Chen, M., Chen, Y., Chen, C., Chen, L., Chen, S., Chen, X., Cheng, J., et al.: Safework-r1: Coevolving safety and intelligence under the ai-45° law. arXiv preprint arXiv:2507.18576 (2025)
17. Li, J., Vo, D.M., Sugimoto, A., Nakayama, H.: Evcap: Retrieval-augmented image captioning with external visual-name memory for open-world comprehension. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 13733–13742 (2024)
18. Liu, A., Tang, L., Pan, T., Yin, Y., Wang, B., Yang, A.: Pico: Jailbreaking multimodal large language models via pictorial code contextualization. In: 2025 IEEE International Conference on Multimedia and Expo (ICME). pp. 1–6. IEEE (2025)
19. Liu, H., Li, C., Li, Y., Lee, Y.J.: Improved baselines with visual instruction tuning. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 26296–26306 (2024)
20. Liu, X., Xu, N., Chen, M., Xiao, C.: Autodan: Generating stealthy jailbreak prompts on aligned large language models. In: International Conference on Learning Representations. vol. 2024, pp. 56174–56194 (2024)
21. Liu, X., Zhu, Y., Gu, J., Lan, Y., Yang, C., Qiao, Y.: Mm-safetybench: A benchmark for safety evaluation of multimodal large language models. In: European Conference on Computer Vision. pp. 386–403. Springer (2024)
22. Liu, X., Cui, X., Li, P., Li, Z., Huang, H., Xia, S., Zhang, M., Zou, Y., He, R.: Jailbreak attacks and defenses against multimodal generative models: A survey. arXiv preprint arXiv:2411.09259 (2024)
23. Ma, S., Luo, W., Wang, Y., Liu, X.: Visual-roleplay: Universal jailbreak attack on multimodal large language models via role-playing image character. arXiv preprint arXiv:2405.20773 (2024)
24. Miao, Z., Ding, Y., Li, L., Shao, J.: Visual contextual attack: Jailbreaking mllms with image-driven context injection. arXiv preprint arXiv:2507.02844 (2025)
25. Miao, Z., Li, L., Xiong, Y., Liu, Z., Zhu, P., Shao, J.: Response attack: Exploiting contextual priming to jailbreak large language models. arXiv preprint arXiv:2507.05248 (2025)
26. Niu, Z., Ren, H., Gao, X., Hua, G., Jin, R.: Jailbreaking attack against multimodal large language model. arXiv preprint arXiv:2402.02309 (2024)
27. OpenAI: GPT-4V(ision) system card. <https://openai.com/index/gpt-4v-system-card/> (Sep 2023), last accessed 29 Jun 2026

28. Qi, X., Huang, K., Panda, A., Henderson, P., Wang, M., Mittal, P.: Visual adversarial examples jailbreak aligned large language models. In: Proceedings of the AAAI conference on artificial intelligence. vol. 38, pp. 21527–21536 (2024)
29. Qi, X., Zeng, Y., Xie, T., Chen, P.Y., Jia, R., Mittal, P., Henderson, P.: Fine-tuning aligned language models compromises safety, even when users do not intend to! In: International Conference on Learning Representations. vol. 2024, pp. 30988–31043 (2024)
30. Rombach, R., Blattmann, A., Lorenz, D., Esser, P., Ommer, B.: High-resolution image synthesis with latent diffusion models. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 10684–10695 (2022)
31. Shayegani, E., Dong, Y., Abu-Ghazaleh, N.: Jailbreak in pieces: Compositional adversarial attacks on multi-modal language models. In: International conference on learning representations. vol. 2024, pp. 30853–30885 (2024)
32. Shen, X., Chen, Z., Backes, M., Shen, Y., Zhang, Y.: "do anything now": Characterizing and evaluating in-the-wild jailbreak prompts on large language models. In: Proceedings of the 2024 on ACM SIGSAC Conference on Computer and Communications Security. pp. 1671–1685 (2024)
33. Tanaka, R., Nishida, K., Yoshida, S.: Visualmrc: Machine reading comprehension on document images. In: Proceedings of the AAAI Conference on Artificial Intelligence. vol. 35, pp. 13878–13888 (2021)
34. Team, G., Anil, R., Borgeaud, S., Alayrac, J.B., Yu, J., Soricut, R., Schalkwyk, J., Dai, A.M., Hauth, A., Millican, K., et al.: Gemini: a family of highly capable multimodal models. arXiv preprint arXiv:2312.11805 (2023)
35. Wang, Y., Liu, X., Li, Y., Chen, M., Xiao, C.: Adashield: Safeguarding multimodal large language models from structure-based attack via adaptive shield prompting. In: European Conference on Computer Vision. pp. 77–94. Springer (2024)
36. Wei, H., Sun, Y., Li, Y.: Deepseek-ocr: Contexts optical compression. arXiv preprint arXiv:2510.18234 (2025)
37. Weng, F., Xu, Y., Fu, C., Wang, W.: Mmj-bench: A comprehensive study on jailbreak attacks and defenses for multimodal large language models. arXiv preprint arXiv:2408.08464 (2024)
38. Wu, C., Li, J., Zhou, J., Lin, J., Gao, K., Yan, K., Yin, S.m., Bai, S., Xu, X., Chen, Y., et al.: Qwen-image technical report. arXiv preprint arXiv:2508.02324 (2025)
39. Xie, Y., Yi, J., Shao, J., Curl, J., Lyu, L., Chen, Q., Xie, X., Wu, F.: Defending chatgpt against jailbreak attack via self-reminders. *Nature Machine Intelligence* **5**(12), 1486–1496 (2023)
40. Yang, C., Lu, C., Wang, Y., Zhou, B.: Towards ai-45o law: A roadmap to trustworthy agi. arXiv preprint ArXiv:2412.14186 (2024)
41. Yi, S., Liu, Y., Sun, Z., Cong, T., He, X., Song, J., Xu, K., Li, Q.: Jailbreak attacks and defenses against large language models: A survey. arXiv preprint arXiv:2407.04295 (2024)
42. Yin, S., Fu, C., Zhao, S., Li, K., Sun, X., Xu, T., Chen, E.: A survey on multimodal large language models. *National Science Review* **11**(12), nwae403 (2024)
43. Yu, Z., Ouyang, X., Shao, Z., Wang, M., Yu, J.: Prophet: Prompting large language models with complementary answer heuristics for knowledge-based visual question answering. *IEEE Transactions on Pattern Analysis and Machine Intelligence* (2025)
44. Zellers, R., Bisk, Y., Farhadi, A., Choi, Y.: From recognition to cognition: Visual commonsense reasoning. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 6720–6731 (2019)

45. Zhang, H., Li, X., Bing, L.: Video-llama: An instruction-tuned audio-visual language model for video understanding. In: Proceedings of the 2023 conference on empirical methods in natural language processing: system demonstrations. pp. 543–553 (2023)
46. Zhang, J., Huang, J., Jin, S., Lu, S.: Vision-language models for vision tasks: A survey. *IEEE transactions on pattern analysis and machine intelligence* **46**(8), 5625–5644 (2024)
47. Zhang, Y., Huang, Y., Sun, Y., Liu, C., Zhao, Z., Fang, Z., Wang, Y., Chen, H., Yang, X., Wei, X., et al.: Multitrust: A comprehensive benchmark towards trustworthy multimodal large language models. *Advances in Neural Information Processing Systems* **37**, 49279–49383 (2024)
48. Zhang, Z., Sun, Z., Zhang, Z., Guo, J., He, X.: Fc-attack: Jailbreaking multimodal large language models via auto-generated flowcharts. *arXiv preprint ArXiv:2502.21059* (2025)
49. Zhao, S., Duan, R., Wang, F., Chen, C., Kang, C., Ruan, S., Tao, J., Chen, Y., Xue, H., Wei, X.: Jailbreaking multimodal large language models via shuffle inconsistency. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 2045–2054 (2025)
50. Zhao, W.X., Zhou, K., Li, J., Tang, T., Dong, Z., Hou, Y., Zhang, B., Min, Y., Zhang, J., Liu, P., et al.: A survey of large language models. *Frontiers of Computer Science* **20**(12), 2012627 (2026)
51. Zhao, Y., Pang, T., Du, C., Yang, X., Li, C., Cheung, N.M.M., Lin, M.: On evaluating adversarial robustness of large vision-language models. *Advances in Neural Information Processing Systems* **36**, 54111–54138 (2023)
52. Zheng, L., Chiang, W.L., Sheng, Y., Zhuang, S., Wu, Z., Zhuang, Y., Lin, Z., Li, Z., Li, D., Xing, E., et al.: Judging llm-as-a-judge with mt-bench and chatbot arena. *Advances in neural information processing systems* **36**, 46595–46623 (2023)
53. Zhu, D., Shen, X., Li, X., Elhoseiny, M., et al.: Minigpt-4: Enhancing vision-language understanding with advanced large language models. In: International Conference on Learning Representations. vol. 2024, pp. 18378–18394 (2024)
54. Zong, Y., Bohdal, O., Yu, T., Yang, Y., Hospedales, T.: Safety fine-tuning at (almost) no cost: A baseline for vision large language models. *arXiv preprint arXiv:2402.02207* (2024)