
Reinforcing Spatial Reasoning in Vision-Language Models with Interwoven Thinking and Visual Drawing

Junfei Wu^{123*}, Jian Guan^{3*}, Kaituo Feng⁴, Qiang Liu¹², Shu Wu^{12†}, Liang Wang¹²,
Wei Wu^{3†}, Tieniu Tan¹²⁵

¹NLPR, MAIS, Institute of Automation, Chinese Academy of Sciences.

²University of Chinese Academy of Sciences. ³Ant Group.

⁴CUHK MMLab. ⁵Nanjing University.

junfei.wu@cripac.ia.ac.cn, shu.wu@nlpr.ia.ac.cn,
{jianguanthu, wuwei19850318}@gmail.com

 Code: <https://github.com/AntResearchNLP/ViLaSR>

 Model: <https://huggingface.co/AntResearchNLP/ViLaSR>

Abstract

As textual reasoning with large language models (LLMs) has advanced significantly, there has been growing interest in enhancing the multimodal reasoning capabilities of large vision-language models (LVLMs). However, existing methods primarily approach multimodal reasoning in a straightforward, text-centric manner, where both reasoning and answer derivation are conducted purely through text, with the only difference being the presence of multimodal input. As a result, these methods often encounter fundamental limitations in spatial reasoning tasks that demand precise geometric understanding and continuous spatial tracking—capabilities that humans achieve through mental visualization and manipulation. To address the limitations, we propose drawing to reason in space, a novel paradigm that enables LVLMs to reason through elementary drawing operations in the visual space. By equipping models with basic drawing operations, including annotating bounding boxes and drawing auxiliary lines, we empower them to express and analyze spatial relationships through direct visual manipulation, meanwhile avoiding the performance ceiling imposed by specialized perception tools in previous tool-integrated reasoning approaches. To cultivate this capability, we develop a three-stage training framework: cold-start training with synthetic data to establish basic drawing abilities, reflective rejection sampling to enhance self-reflection behaviors, and reinforcement learning to directly optimize for target rewards. Extensive experiments demonstrate that our model, named **VILASR**, consistently outperforms existing methods across diverse spatial reasoning benchmarks, involving maze navigation, static spatial reasoning, video-based reasoning, and multi-view-based reasoning tasks, with an average improvement of 18.4%. Ablation studies reveal the critical role of each training stage, where reflective rejection sampling strengthens the model’s self-correction capabilities, and reinforcement learning effectively unlocks its reasoning potential.

1 Introduction

Large language models (LLMs) have exhibited remarkable reasoning capabilities in complex tasks such as mathematical problem-solving [66, 35] and code generation [6], particularly through the

*Equal Contribution. †Corresponding Author.

“slow thinking” paradigm exemplified by OpenAI o1 [26] and DeepSeek R1 [10], which enables extended reasoning with in-depth self-reflection. Encouraged by the success, a growing body of research is now aiming to adapt similar techniques to large vision-language models (LVLMs) to enhance their capabilities in image and video reasoning [72, 11, 50, 13].

While current LVLMs excel at basic visual perception tasks such as object detection [12] and narrative understanding [40, 9], they often struggle with spatial reasoning [64, 70], which requires understanding spatial relationships among objects and tracking their dynamic evolution [60]—capabilities that are crucial for real-world applications such as robotics [49] and augmented reality [17]. Indeed, even when such deep visual understanding is required, LVLMs still rely solely on text-based reasoning [72], assuming that visual information can be perfectly translated into textual semantic space [23]. Unfortunately, such translation is inherently challenging [33]: spatial details are inevitably lost when converting visual information to text, and describing dynamic changes of object positions becomes prohibitively complex in textual space (see examples of GPT-4o in Figure 1). Taking inspiration from human cognition, where spatial reasoning relies on mental visualization and dynamic manipulation [3], we advocate a vision-centric reasoning paradigm where models actively edit and re-encode visual information at each reasoning step, dynamically supplementing spatial details and relationships. This “thinking with images” approach, while validated by OpenAI o3 [47], remains underexplored in open-source research.

In parallel with our work, recent studies have begun to integrate visual tools to enable vision-centric reasoning [8, 51, 21, 31]. However, these approaches exhibit limitations along two key dimensions. First, the reasoning capabilities of LVLMs are constrained by black-box perception tools (e.g., grounding and OCR systems), resulting in not only fixed perception capabilities but also fragmented reasoning composed of disconnected tool invocations that undermine coherence and holistic planning. Second, these methods heavily draw upon reasoning data curated based on human priors, which often exhibit oversimplified logic compared to the complexity of spatial reasoning tasks, with problem decomposition and tool invocation interleaved in a simplistic and linear fashion [8, 51]. Such adherence to prescribed reasoning patterns prevents LVLMs from developing the capacity for critical reflection on tool outputs—a capability that has proven crucial for advanced reasoning [45]. These fundamental limitations call for a more flexible and intrinsic approach to vision-centric reasoning.

In light of these challenges, we propose “drawing to reason in space,” a novel, versatile reasoning paradigm that empowers LVLMs to reason through elementary drawing operations, as exemplified in Figure 1. Through simple yet powerful operations, including bounding boxes for object localization and auxiliary lines for relationship analysis, this paradigm enables direct visual interaction for spatial problem-solving, mirroring human behavior while avoiding the pitfalls of dependence on external perception tools. Based on this paradigm, we develop **VILASR**, a **V**ision-**L**anguage model that achieves sophisticated **S**patial **R**easoning through interwoven thinking and visual drawing. To realize this vision, we develop a three-stage training framework (Figure 2). First, we introduce cold-start training with synthetic data to establish basic visual drawing abilities. Second, we design a reflective rejection sampling mechanism that selectively reinforces reasoning paths demonstrating both correct answers and self-correction behaviors, enabling models to revise their visual operations based on intermediate results. Finally, we employ reinforcement learning (RL) with carefully designed rewards that balance answer correctness and reasoning format, incentivizing both accurate spatial understanding and coherent visual thinking processes.

Extensive experiments across five diverse spatial reasoning benchmarks demonstrate the effectiveness of VILASR, which achieves substantial improvements over strong baselines by 18.4% on average, involving challenging scenarios that require sequential spatial planning (e.g., maze navigation), temporal relationship tracking (e.g., video-based reasoning), and information integration from multiple perspectives. Ablation studies reveal the critical role of each training stage. Particularly, reflective rejection sampling significantly enhances the model’s self-correction capabilities, with the frequency of reflection behaviors doubling compared to models trained without this intermediate stage. Further inference-time scaling experiments reveal that each training stage progressively enhances the model’s reasoning potential, with the final RL optimization effectively consolidating multi-attempt capabilities while maintaining strong single-attempt performance, significantly narrowing the performance gap compared to earlier training stages.

In summary, our contributions are threefold:

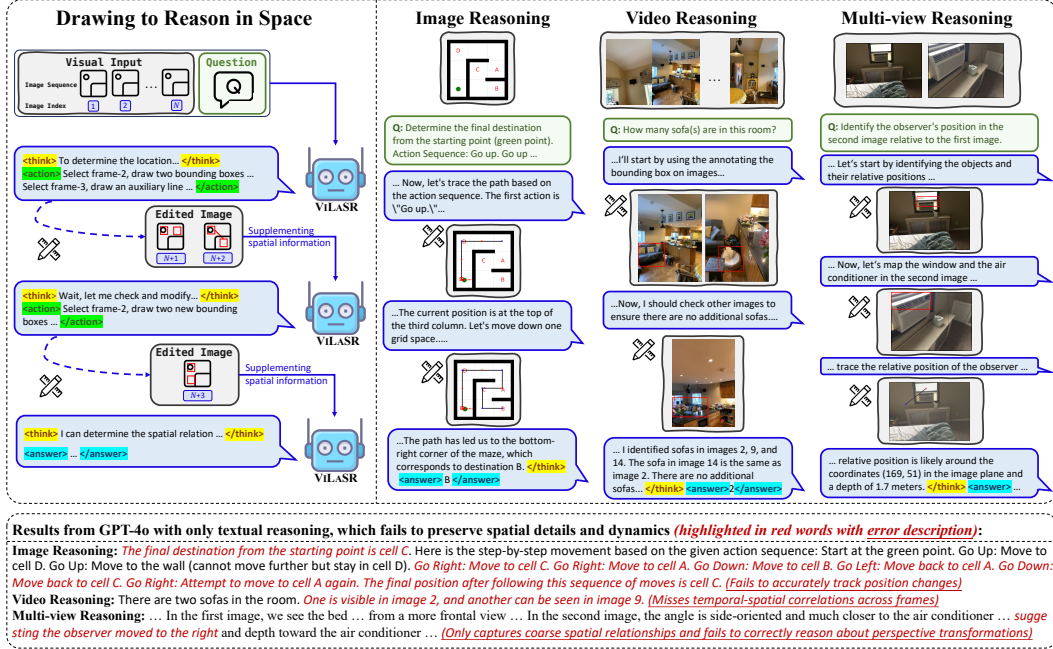


Figure 1: **Top left:** Overview of the “drawing to reason in space” paradigm, which enables visual reasoning through iterative thinking and drawing operations. **Top right:** Examples across three spatial reasoning tasks, demonstrating how ViLASR decomposes complex problems into interpretable visual reasoning steps. **Bottom:** Corresponding results from GPT-4o with only textual reasoning.

- I. We propose drawing to reason in space, a novel reasoning paradigm that enables LVLMs to perform spatial reasoning through interpretable visual operations;
- II. We develop a principled training framework that effectively cultivates models’ visual reasoning capabilities through progressive stages;
- III. We demonstrate state-of-the-art performance across multiple spatial reasoning benchmarks and provide insights into the development of visual reasoning abilities in LVLMs.

2 Related works

2.1 Reasoning in LLMs and LVLMs

Instilling advanced reasoning capabilities within LLMs and LVLMs remains a significant challenge. Current approaches can be broadly categorized into three directions: **(1) Prompt engineering:** crafting prompts to elicit latent reasoning capabilities in pre-trained models [66, 81] and enable tool usage to complement model capabilities with external knowledge [1, 18] or specialized functionalities [57, 19, 61, 73]. These prompt-based approaches heavily depend on models’ instruction-following capabilities and often suffer from prompt sensitivity. **(2) Supervised fine-tuning:** developing tailored datasets to enhance specific reasoning capabilities [75, 67, 82] and training models to effectively utilize external tools [54, 53, 18, 8, 55, 51]. Such approaches are inherently bounded by the quality and scale of training data, limiting their potential for advanced reasoning. **(3) Reinforcement learning:** engineering reward mechanisms to incentivize desired reasoning patterns [29, 56, 10] and optimize tool usage strategies [34]. While promising in textual domains, RL approaches for tool-augmented reasoning have yet to be fully investigated in multimodal scenarios. Current LVLMs either confine their reasoning to textual form [72], or rely on specialized perception tools, being constrained by tool capabilities [8, 51] and lacking reflective reasoning [55]. In contrast, we pioneer a staged training recipe that enables models to both express and reflect upon their spatial understanding through iterative drawing operations.

2.2 Visual spatial reasoning

Visual spatial reasoning, as a crucial component of multimodal intelligence, extends beyond general visual perception to encompass two particularly challenging fundamental capabilities [15, 77]: relational reasoning, which involves understanding distances, directions, and spatial common sense between objects [38, 60]; and perspective transformation, which requires holding and manipulating spatial relationships [42, 70]. Despite the remarkable progress of current LVLMs in basic visual tasks [24, 36, 39, 9], numerous benchmarks have revealed significant challenges in spatial reasoning [43, 27, 69, 32, 70]. This limitation is partly attributed to the nature of current training datasets [22], which primarily focus on visual perception rather than spatial understanding. Recent works have attempted to address this challenge through various approaches. For image-based spatial reasoning, several studies have proposed synthetic datasets [38, 5, 7, 4]. For video-based spatial reasoning, approaches have emerged either by incorporating 3D representations as bridging knowledge [83] or by tracking objects across frames [37]. However, systematic approaches to enhance LVLMs’ spatial reasoning capabilities that can generalize across both images and videos remain largely unexplored. Our work fills this gap by developing a principled reasoning framework in LVLMs.

3 Methodology

In this section, we present our approach to advancing spatial reasoning capabilities in LVLMs. Given a spatial reasoning question Q and a visual input $I = \{I_n\}_{n=1}^N$, where each I_n represents a single image, $N = 1$ corresponds to a single image input, and $N > 1$ denotes a video sequence or multiple image inputs, our goal is to enable LVLMs to derive the answer A through iterative visual drawing and thinking. We first introduce our reasoning paradigm (3.1), and then detail the training framework (3.2) that cultivates this reasoning capability.

3.1 Drawing to reason in space

We propose “drawing to reason in space,” a reasoning paradigm that empowers LVLMs to decompose complex spatial reasoning tasks into a sequence of interpretable visual drawing and thinking steps, as illustrated in Figure 1. Formally, given the visual input I and question Q , the LVLM $\mathcal{M}$ generates a multi-step reasoning path R to derive the final answer. The reasoning path R can be represented as a T -step chain: $R = \{(r_t, e_t, o_t)\}_{t=1}^T$, where each step interleaves natural language reasoning r_t , drawing operations e_t , and observed results o_t from executing the drawing operations [74]. This iterative process continues until $\mathcal{M}$ reaches a conclusive answer in the final reasoning step r_T . Next, we elaborate on the core components of the paradigm.

Drawing operations for spatial reasoning. We equip LVLMs with two essential drawing operations $\mathcal{T} = \{\mathcal{T}_{\text{box}}, \mathcal{T}_{\text{line}}\}$ for bounding box annotation and auxiliary line drawing, respectively. These operations are fundamental to spatial reasoning as they enable explicit representation of object locations and their spatial transitions—bounding boxes anchor object positions while auxiliary lines visualize spatial trajectories and relationships². Each operation $\tau \in \mathcal{T}$ accepts three parameters:

- k : the index of the target image from either the original input I or previous drawing outputs $\{o_i\}$;
- p : single or multiple spatial coordinates for annotating bounding boxes or drawing auxiliary lines;
- l : semantic labels describing the annotated content.

The execution output of each drawing operation is an annotated version of the target image, preserving the original content while overlaying the specified visual elements. While this minimal operation set proves effective for spatial reasoning, it can be readily extended with additional visual manipulation tools in future work.

Per-step spatial reasoning. At the t -th reasoning step, $\mathcal{M}$ generates both natural language reasoning r_t and drawing operations e_t based on the entire interaction history:

$$(r_t, e_t = \{e_t^j = (k_t^j, p_t^j, l_t^j)\}_{j=1}^{m_t}) = \mathcal{M}(I, Q, R_{<t}), \quad (1)$$

²Since this work focuses on inter-object spatial relationships rather than object-specific details, we adopt only basic drawing operations rather than common visual manipulations like zooming or cropping [8, 47].

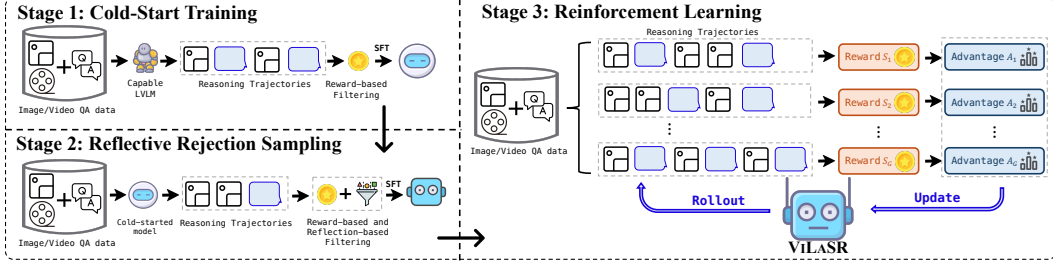


Figure 2: Overview of the three-stage training framework of ViLASR.

where e_t specifies a set of m_t operations to be executed, e_t^j denote each single operation with three necessary parameters, and $R_{<t} = \{(r_i, e_i, o_i)\}_{i<t}$ means the reasoning history before step t . The operations in e_t can target different images through distinct indices k_t^j , which is crucial for both image and video reasoning as key information may come from multiple previous drawing outputs or different video frames. After execution, these operations output a set of annotated images $o_t = \{o_t^j\}_{j=1}^{m_t}$, which are sequentially indexed to maintain order. To enable index-based image retrieval for drawing, $\mathcal{M}$ has access to the indices of all available images, including both the original visual input and operation-generated outputs. The reasoning process terminates under two conditions: (i) when r_t reaches a final answer, in which case $e_t = \emptyset$ and $o_t = \emptyset$; or (ii) when the number of reasoning steps reaches a predefined maximum T_{max} .

3.2 Training framework

Our training framework consists of three stages: cold-start training with synthetic offline data to establish basic visual interaction capabilities, reflective rejection sampling to cultivate reflection behaviors, and reinforcement learning for incentivizing reasoning potentials, as illustrated in Figure 2. During training, we focus on multiple-choice questions (e.g., “Which object is to the left of the chair? A. table ...”) and numerical questions (e.g., “How many tables appear in this video?”) due to their amenability to automated correctness evaluation, strategically excluding free-form answer questions (e.g., describing motion trajectories). In what follows, we first present our reward function design that guides all three stages, followed by detailed descriptions of the stage-wise training procedures.

Reward function. For reward design, we propose a rule-based function that combines answer correctness and reasoning format adherence:

$$S = \mathbb{1}(S_{\text{correct}} > \beta) \cdot (S_{\text{correct}}(A, \hat{A}) + S_{\text{format}}(R)), \quad (2)$$

where β is a threshold that ensures format rewards are only granted when a minimum correctness threshold is met, preventing reward hacking where the model might optimize for format adherence at the expense of task accuracy. $\hat{A}$ is the predicted answer that we extract from the model’s final reasoning step r_T using rule-based parsing. Specifically, we compute S_{correct} as follows:

$$S_{\text{correct}}(A, \hat{A}) = \begin{cases} \mathbb{1}(A = \hat{A}), & \text{for multiple-choice questions,} \\ \frac{1}{|\mathcal{C}|} \sum_{\theta \in \mathcal{C}} \mathbb{1}\left(\frac{|A - \hat{A}|}{A} < 1 - \theta\right), & \text{for numerical questions.} \end{cases} \quad (3)$$

For multiple-choice questions, we simply check the exact match between the predicted and ground-truth answers. For numerical questions, we adopt Mean Relative Accuracy (MRA) [70] for reward computation, which provides a more robust evaluation than conventional metrics like absolute error or fixed thresholding. Instead, MRA examines the prediction accuracy across multiple confidence levels $\mathcal{C} = \{0.50, 0.55, \dots, 0.95\}$. For each threshold $\theta \in \mathcal{C}$, it checks if the relative error between predicted value A and ground truth $\hat{A}$ falls below $(1 - \theta)$. The final score averages these binary outcomes, effectively measuring the model’s precision across different stringency levels.

For S_{format} , we evaluate the quality of reasoning format based on the structural validity of R , assigning a score of 1 if all operations in the reasoning path R are executable, and 0 otherwise.

Cold-start training. Prompting alone often fails to elicit effective visual manipulation abilities in LLMs for reasoning [19, 8]. Therefore, we initialize models’ ability to reason in space with

drawing operations through supervised learning on a synthetic dataset $\mathcal{D}_{\text{cold}}$. The training objective is to minimize the average negative log-likelihood over all reasoning and operation tokens:

$$\mathcal{L}_{\text{cold}} = \mathbb{E}_{(I, Q, A, R = \{(r_t, e_t, o_t)\}_{t=1}^T) \sim \mathcal{D}_{\text{cold}}} \left(-\frac{1}{N} \sum_{t=1}^T \log p(r_t, e_t | I, Q, R_{<t}) \right) \quad (4)$$

where $N = \sum_{t=1}^T (|r_t| + |e_t|)$ denotes the total number of tokens in reasoning steps and drawing operations. To construct $\mathcal{D}_{\text{cold}}$, we first collect a diverse set of image and video question-answering pairs from publicly available datasets, comprising visual inputs I , questions Q , and ground-truth answers A . We then leverage Qwen2.5-72B-VL [2] to generate reasoning paths following our reasoning paradigm described in §3.1. The generated paths are subsequently filtered based on rule-based correctness and format checking in Eq. 2 to ensure high-quality demonstrations of spatial reasoning. Appendix B shows more details.

Reflective rejection sampling. The success of RL often relies on the model’s initial capability to exhibit reflective behaviors [14]. In our case, the ability to reflect on and revise drawing operations based on observed execution output is crucial. Accordingly, we define reflective behavior as the recurrence of identical labels across different time steps within the same reasoning process. Formally, for a reasoning path R , reflection occurs when:

$$\exists(t_1, t_2, u, v) : (l_{t_1}^u = l_{t_2}^v) \wedge (e_{t_1}^u \neq e_{t_2}^v) \text{ in } R = \{(r_t, e_t = \{e_t^j = (k_t^j, p_t^j, l_t^j)\}_{j=1}^{m_t}, o_t)\}_{t=1}^T, \quad (5)$$

where $l_{t_1}^u$ and $l_{t_2}^v$ represent semantic labels assigned to drawing operations at different reasoning steps, and e_t^j denote the parameters of one drawing operation. However, after cold-start training, we observe that the resulting model, denoted as $\mathcal{M}_{\text{cold}}$, infrequently exhibits such reflective behavior (see §4.3 for more details), potentially limiting the effectiveness of subsequent RL optimization. To address this limitation, we introduce a novel reflective rejection sampling mechanism. Given a batch of spatial reasoning examples $\mathcal{D}_{\text{reflect}}$, we first use $\mathcal{M}_{\text{cold}}$ to sample corresponding reasoning paths, and then continue fine-tune $\mathcal{M}_{\text{cold}}$ with the following objective:

$$\mathcal{L}_{\text{reflect}} = \mathbb{E}_{(I, Q, A) \sim \mathcal{D}_{\text{reflect}}, R \sim \mathcal{M}_{\text{cold}}(\cdot | I, Q)} \left(-\frac{1}{N} \sum_{t=1}^T \log p(r_t, e_t | I, Q, R_{<t}) \right) \phi, \quad (6)$$

where ϕ acts as a binary filter that equals 1 if and only if the reasoning path R not only yields the correct answer A and meets format criteria, but also satisfies the reflection condition in Eq. 5. This selective training strategy encourages the model to learn from high-quality reasoning paths that demonstrate both reflective thinking and correct reasoning, facilitating the development of self-correction capabilities crucial for subsequent RL optimization.

Reinforcement learning. We further optimize the model using RL with carefully designed rollout policies. During policy rollout, we monitor the reasoning process and apply early termination to avoid inefficient or circular reasoning patterns. Specifically, the rollout is terminated and the reward is set to zero when any of the following conditions are met: (1) the model fails to generate any drawing operations, i.e., $e_t = \emptyset$ while r_t has not reached a final answer; (2) the number of accumulated images exceeds a predefined threshold α ; or (3) a drawing operation duplicates a previously executed one, i.e., $\exists t_1, t_2, u, v : e_{t_1}^u = e_{t_2}^v$.

With the rollout policies defined above and reward functions in Eq. 2, we optimize the policy using GRPO [10] without the KL penalty term [20]:

$$\mathcal{L}_{\text{RL}} = \mathbb{E}_{\substack{(I, Q, A) \sim \mathcal{D}_{\text{rl}} \\ \{R_i\}_{i=1}^G \sim p_{\text{old}}(\cdot | I, Q)}} \left(-\frac{1}{G} \sum_{i=1}^G \frac{1}{N_i} \sum_{t=1}^{T_i} \min(\rho_{i,t} A_i, \text{clip}(\rho_{i,t}, 1 - \epsilon, 1 + \epsilon) A_i) \right), \quad (7)$$

$$\rho_{i,t} = \frac{p(r_{i,t}, e_{i,t} | I, Q, R_{i,<t})}{p_{\text{old}}(r_{i,t}, e_{i,t} | I, Q, R_{i,<t})}, A_i = \frac{S_i - \text{mean}(\{S_j\}_{j=1}^G)}{\text{std}(\{S_j\}_{j=1}^G)}, \quad (8)$$

where G is the number of rollout reasoning paths, $R_i = \{(r_{i,t}, e_{i,t}, o_{i,t})\}_{t=1}^{T_i}$ is the i -th reasoning path, N_i denotes the total length of the R_i except the tool outputs, S_i is the reward of R_i , and p and p_{old} represent the current and old policy distributions, respectively. The normalized score $A_{i,t}$ reflects the relative quality of each reasoning path within the rollout group, enabling the model to distinguish between learnable and poor reasoning trajectories.

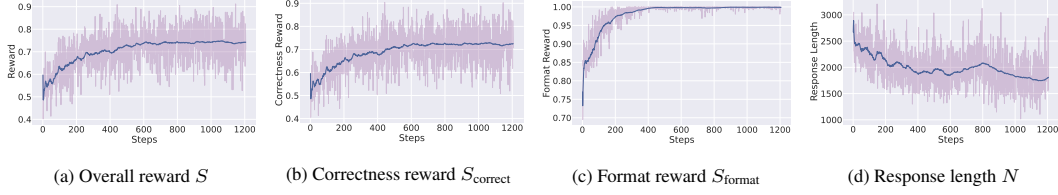


Figure 3: RL training curves of ViLASR.

4 Experiment

4.1 Experimental setups

Evaluation benchmark and metrics. To evaluate the effectiveness and generalization of ViLASR, we conduct experiments on five benchmarks covering three categories:

- **Image spatial reasoning:** focusing on static relationships and sequential planning, including (1) Maze [25], specifically designed for navigation assessment; (2) SpatialEval-Real [65], demanding real-world spatial relation understanding.
- **Video spatial reasoning:** requiring temporal relationship tracking, including VSI-Bench [70], which tests visual spatial understanding over temporal sequences.
- **Multi-view spatial reasoning:** challenging models to integrate information from multiple perspectives, including SPAR-Bench [78], MMSI-Bench [71].

These benchmarks consist exclusively of multiple-choice and numerical questions. We evaluate the model performance using accuracy for multiple-choice questions and Mean Relative Accuracy (MRA) [70] for numerical questions. Notably, these metrics follow the same formulation as our training objective (Eq. 3). Appendix C.2 shows the benchmark statistics.

Implementation details. We implement ViLASR based on Qwen2.5-VL-7B [2]. During the training phase, we process all visual inputs at a maximum resolution of $256 \times 28 \times 28$, with video clips uniformly sampled at 16 frames. In the evaluation stage, we maintain the 16-frame count and the $256 \times 28 \times 28$ frame resolution for VSI-Bench, while increase the resolution to $448 \times 28 \times 28$ for other benchmarks. The training is conducted on a cluster of 16 NVIDIA A100 (80G) GPUs. For both cold-start training stage and reflective fine-tuning stage, we optimize the model with a learning rate of 1×10^{-5} for three epochs, requiring approximately 24 and 3 hours, respectively. The subsequent RL optimization is implemented using the VERL framework [58], where we set the training batch size to 32 and generate 8 candidate reasoning paths per question. We set the maximum cumulative image number α to 42. And the reward threshold β in Eq. 2 is set to 0.0.

The training dynamics are shown in Figure 3. As illustrated in (a-c), both overall reward S and its components (S_{correct} and S_{format}) demonstrate steady improvements, validating the effectiveness of our RL optimization. Notably, while achieving better task performance, we observe a decrease in response length N from 2,500 to 1,800 tokens (d). This can be attributed to two factors: First, during cold-start training, we set a minimum of three reasoning steps in synthetic data to encourage in-depth reasoning, while RL optimization subsequently drives the model toward more efficient use of drawing operations. Second, the limited training data may not fully expose the model to complex scenarios requiring lengthy reasoning chains. We expect that increasing training data diversity could potentially lead to longer but necessary reasoning steps for more challenging spatial tasks [68]. We also present an analysis of the runtime performance in Appendix E.2.

Baselines. We compare ViLASR with various representative models and methods as follows:

- **Proprietary LVLMS:** GPT-4o [44], Gemini-1.5-Pro [62], Gemini-2.0-Flash [16], OpenAI o3 and o4-mini [46];
- **Open-source LVLMS:** These models range from small-sized models including Qwen2.5-VL-7B [2], LLaVA-NeXT-Video-7B [80] and LLaVA-OneVision-7B [30] to large-sized models including Kimi-VL-A3B-Instruct (16B) [63], Qwen2.5-VL-72B [2], LLaVA-NeXT-Video-72B [80], LLaVA-OneVision-72B [30];

Table 1: Performance comparison across spatial reasoning benchmarks. Gray-shaded rows represent large-sized models (>7B parameters). For non-shaded rows, **bold** and underlined numbers indicate the best and second-best results. *Italic* numbers in gray-shaded rows indicate performance below non-shaded best results. *Improvement* refers to the absolute improvement of ViLASR compared with Qwen2.5-VL-7B w/o reasoning. [†] and [‡] indicate results from VSI-Bench and VSI-Bench (tiny) set [70] respectively. * Results from [63]. ** Results from [71]. N/A: not support multiple-image input. OpenAI o4-mini is evaluated on a small subset of benchmarks due to its high cost.

Method	Tool	Reasoning	Image		Video	Mutli-view	
			MAZE	SpatialEval-Real	VSI-Bench	SPAR-Bench	MMSI-Bench
Proprietary LVLMS							
GPT-4o	✗	✗	48.8	60.7	34.0 [†]	33.6	30.3**
GPT-4o	✗	☑	63.1	65.1	-	38.1	-
Gemini-1.5-Pro	✗	✗	-	-	45.4 [†]	-	-
Gemini-2.0-Flash	✗	✗	40.2	-	45.4 [‡]	33.4	-
Gemini-2.0-Flash	✗	☑	61.7	-	-	28.0	-
OpenAI o3	✗	☑	-	-	-	-	41.0**
OpenAI o4-mini	✗	☑	79.0	-	-	46.2	-
Open-source LVLMS							
Qwen2.5-VL-7B	✗	✗	33.7	58.5	32.7	31.7	26.9
Qwen2.5-VL-7B	✗	☑	36.5	54.1	26.2	31.6	27.1
LLaVA-NeXT-Video-7B	✗	✗	34.7	68.1	35.6 [†]	31.3	26.8
LLaVA-OneVision-7B	✗	✗	30.8	62.9	32.4 [†]	30.6	24.5**
Kimi-VL-A3B-Instruct (16B)	✗	✗	45.0	68.9	37.4*	33.0	-
Qwen2.5-VL-72B	✗	✗	50.5	68.1	36.0	37.4	30.7**
LLaVA-NeXT-Video-72B	✗	✗	43.2	72.5	40.9 [†]	35.6	28.3
LLaVA-OneVision-72B	✗	✗	46.3	68.8	40.2 [†]	34.4	28.4**
Representative methods for multimodal reasoning							
CogCoM-17B [51]	☑	☑	29.8	49.3	N/A	N/A	N/A
VisCoT-7B [55]	☑	☑	26.0	43.2	N/A	N/A	N/A
SpaceR-7B [48]	✗	☑	<u>38.6</u>	62.7	<u>43.5</u>	<u>37.1</u>	<u>28.8</u>
VisProg [19]	☑	☑	22.6	26.8	N/A	N/A	N/A
ViperGPT [61]	☑	☑	27.3	45.9	10.6	23.7	25.4
Ours							
ViLASR	☑	☑	98.2	<u>63.9</u>	45.4	37.6	30.2
Improvement	N/A	N/A	+64.5	+5.4	+12.7	+5.9	+3.3

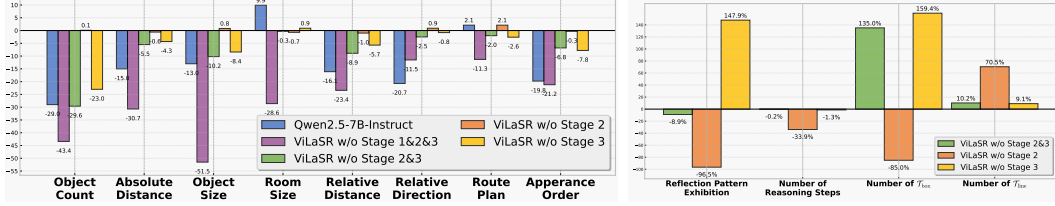
- **Representative models focused on multimodal reasoning:** CoGCoM [51], which enables step-by-step visual reasoning through image manipulations with specialized perception tools (e.g., grounding) but is limited to single-image inputs; VisCoT [55], which first identifies key regions through bounding box annotations followed by textual reasoning chains on single images; SpaceR [48], which is a contemporary work that extends spatial reasoning to video understanding through task-specific optimization on automatically curated video QA data; VisProg [19], which pioneers visual programming by prompting an LLM to generate pseudocode orchestrating modular vision APIs for compositional tasks; and ViperGPT [61], which extends this to generate executable Python with richer control flow (e.g., loops, conditionals) for vision–language interaction.

Furthermore, we perform ablation studies to evaluate each individual training stage to assess their contributions to the final performance. Implementation details of the baselines are show in Appendix C.1.

4.2 Main results

As shown in Table 1, ViLASR shows strong performance across various spatial reasoning benchmarks, from maze navigation to complex video understanding. Our analysis reveals several key findings:

(1) **Limited visual reasoning capabilities in open-source LLMs.** Our experiments reveal a significant disparity between proprietary and open-source LLMs in their ability to leverage reasoning processes. Proprietary models consistently outperform their non-reasoning counterparts by a substantial margin across different benchmarks (e.g., GPT-4o, Gemini-2.0-Flash). However, open-source models like Qwen2.5-VL-7B show minimal or even negative improvements with reasoning enabled, suggesting a significant gap in their fundamental visual reasoning capabilities. This observation highlights a critical direction for future research: enhancing the multimodal reasoning abilities of



(a) Absolute change in scores compared with ViLaSR on eight subsets of VSI-Bench. (b) Relative change (%) in four key behavioral metrics compared with ViLaSR.

Figure 4: Ablation study results on VSI-Bench regarding three training stages: cold-start training (Stage 1), reflective rejection sampling (Stage 2), and reinforcement learning (Stage 3).

open-source models, potentially through improved architectural designs and more sophisticated training strategies that better integrate visual and linguistic information.

(2) Strong performance of ViLaSR on image-based spatial tasks. ViLaSR demonstrates exceptional capabilities in both sequential spatial planning (MAZE) and static spatial understanding (SpatialEval). The remarkable performance on maze navigation, in particular, highlights the advantage of our “drawing to reason in space” paradigm: through iterative drawing operations, the model effectively breaks down multi-step navigation sequences into interpretable visual steps, thereby capturing and tracking spatial state transitions. In contrast, existing methods show limited performance due to various constraints: COGCOM is restricted by the capabilities of its perception tools, and VisCoT lacks the ability to reflect and revise its visual operations since it only performs one-time grounding. VisProg and ViperGPT exhibit inferior performance because they often generate malformed or non-executable programs on complex questions. Moreover, despite leveraging modular APIs, these methods frequently struggle with basic spatial relationships, consistent with recent findings [41].

(3) Competitive results of ViLaSR on video and multi-view reasoning. Video and multi-view reasoning present unique challenges beyond static image understanding, as it requires tracking spatial relationships across multiple viewpoints and temporal transitions. While most existing models struggle with this increased complexity, ViLaSR achieves state-of-the-art performance on all benchmarks, surpassing significantly larger open-source models such as LLaVA-OneVision-72B. This success can be attributed to our flexible visual operations that effectively handle dynamic spatial relationships. In contrast, existing approaches face fundamental limitations: COGCOM and VisCoT lack video processing capabilities entirely, while SpaceR’s text-centric reasoning fails to fully exploit visual temporal information.

4.3 Ablation study

To comprehensively evaluate our training framework, we conduct ablation studies on VSI-Bench’s eight spatial reasoning subtasks and analyze four key behavioral metrics. The behavioral metrics include: (1) *Reflection Pattern Exhibition*, measuring the ratio of examples exhibiting reflection behavior as defined in Eq. 5; (2) *Number of Reasoning Steps*, indicating the average number of reasoning steps (T) required to answer each question; and (3,4) *Numbers of T_{box} and T_{line}* , reflecting the average number of drawing operations utilized per question. The results are shown in Figure 4(a) for task performance and Figure 4(b) for behavioral patterns. We draw the following insights:

(1) Cold-start training: Essential for spatial reasoning foundation. Comparing “ViLaSR w/o Stage 1&2&3” with “ViLaSR w/o Stage 2&3”, we observe substantial performance improvements across all subtasks. Notably, when equipped with only our “drawing to reason in space” paradigm without training (i.e., “ViLaSR w/o Stage 1&2&3”), the model performs even worse than the Qwen2.5-VL-7B backbone. This degradation highlights that sophisticated spatial reasoning capabilities cannot emerge solely with prompting but must be learned through dedicated training.

(2) Reflection sampling: Key to self-correction ability. The comparison between ViLaSR and “ViLaSR w/o Stage 2” reveals the crucial role of reflective rejection sampling. Removing this stage leads to a 96.5% decrease in reflection pattern exhibition, fundamentally altering the model’s behavior and performance. Behaviorally, we observe a 33.9% reduction in reasoning steps and dramatically different patterns in drawing operation usage: a 85.0% decrease in T_{box} but a 70.5%

increase in $\mathcal{T}_{\text{line}}$. Such changes particularly affect tasks requiring precise object localization and measurement: “Absolute Distance” (-0.6%), “Room Size” (-0.7%), and “Relative Distance” (-1.0%). In contrast, tasks that primarily rely on directional judgment or categorical reasoning show minimal impact or even slight improvements (e.g., “Object Count”: +0.1%, “Relative Direction”: +0.9%), as they can be solved with simpler spatial relationships indicated by lines. This pattern reveals that reflection capability fundamentally shapes how the model builds spatial understanding. Without it, the model shows a tendency to make quick spatial judgments through auxiliary lines without sufficient self-verification. Based on this observation, we hypothesize specific mechanisms through which reflection enhances spatial reasoning: re-examining object locations through bbox annotations and carefully evaluating different spatial relationships before making final judgments (see examples in Appendix §D). This verification-driven approach, rather than making multiple “guessing” attempts through auxiliary lines, appears to be key to accurate spatial reasoning.

(3) RL optimization: Dense rewards enable fine-grained learning. Comparing ViLaSR with “ViLaSR w/o Stage 3”, we observe consistent performance decreases across most subtasks. The substantial increases in both $\mathcal{T}_{\text{box}}$ (+159.4%) and $\mathcal{T}_{\text{line}}$ (+9.1%) usage without RL suggest that RL optimization helps the model learn to use drawing operations more selectively. Moreover, tasks requiring precise numerical answers (“Object Count,” “Absolute Distance,” “Object Size” and “Room Size”) show greater average performance gaps without RL compared to multiple-choice questions (-9.21% vs. -4.07%). This disparity highlights a key advantage of RL optimization: while supervised learning merely maximizes the probability of correct answers, RL provides dense reward signals based on numerical proximity to ground truth, enabling more effective learning for precise spatial measurements.

4.4 Inference-time scaling

We further evaluate the effectiveness of our training framework through the lens of inference-time scaling analysis such as pass@ k evaluation (sampling k outputs in parallel and selecting the best). Prior research has demonstrated: (1) pass@ k serves as an empirical upper bound of model capability with sufficiently large k [59], while pass@1 reflects its single-attempt performance; (2) RL has been shown to effectively narrow this gap by consolidating pass@ k capabilities into pass@1 performance [76]. This motivates our analysis of how each training stage expands model capability.

Results in Figure 5 reveal the progressive enhancement of model capability across training stages: (1) Cold-start training substantially improves both the empirical upper bound (pass@8) and single-attempt performance (pass@1) compared to Qwen2.5-VL-7B; (2) Reflective rejection sampling further elevates both metrics, demonstrating expanded reasoning potential; (3) RL optimization not only pushes the empirical upper bound (pass@8) but also significantly narrows the gap with single-attempt performance (pass@1), indicating effective consolidation of multi-attempt capabilities that aligns with previous studies about RL’s optimization effects [76].

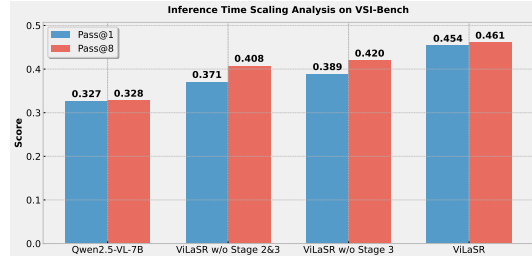


Figure 5: Analysis of inference-time scaling behavior using pass@1 and pass@8 performance.

5 Conclusion

This work presents a principled approach to enhancing spatial reasoning capabilities in LVLMS through visual drawing operations. By enabling direct manipulation in the visual space through the “drawing to reason in space” paradigm, we bridge the gap between text-centric reasoning and human-like spatial cognition. Through careful empirical validation, we demonstrate that our three-stage training framework successfully cultivates sophisticated reasoning pattern. Extensive experiments across diverse spatial reasoning benchmarks validate the effectiveness of our approach, showing particular strengths in maze navigation and temporal-spatial understanding tasks. While our results are promising, several important directions remain for future exploration, such as extending our drawing operations to handle more complex 3D spatial relationships, and investigating more efficient training strategies.

Acknowledgements

We thank the anonymous reviewers for their insightful comments. This work was supported in part by the National Natural Science Foundation of China (NSFC) under Grant Nos. 62141608, 62236010, 62372454, and 62576339, and in part by the Ant Group Research Intern Program.

References

- [1] Akari Asai, Zeqiu Wu, Yizhong Wang, Avirup Sil, and Hannaneh Hajishirzi. Self-RAG: Learning to retrieve, generate, and critique through self-reflection. In The Twelfth International Conference on Learning Representations, 2024.
- [2] Shuai Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, Sibao Song, Kai Dang, Peng Wang, Shijie Wang, Jun Tang, et al. Qwen2. 5-vl technical report. arXiv preprint arXiv:2502.13923, 2025.
- [3] Neil Burgess. Spatial cognition and the brain. Annals of the New York Academy of Sciences, 1124(1):77–97, 2008.
- [4] Wenxiao Cai, Iaroslav Ponomarenko, Jianhao Yuan, Xiaoqi Li, Wankou Yang, Hao Dong, and Bo Zhao. Spatialbot: Precise spatial understanding with vision language models, 2025.
- [5] Boyuan Chen, Zhuo Xu, Sean Kirmani, Brain Ichter, Dorsa Sadigh, Leonidas Guibas, and Fei Xia. SpatialVLM: Endowing vision-language models with spatial reasoning capabilities. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), pages 14455–14465, June 2024.
- [6] Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Ponde de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, Alex Ray, Raul Puri, Gretchen Krueger, Michael Petrov, Heidy Khlaaf, Girish Sastry, Pamela Mishkin, Brooke Chan, Scott Gray, Nick Ryder, Mikhail Pavlov, Alethea Power, Lukasz Kaiser, Mohammad Bavarian, Clemens Winter, Philippe Tillet, Felipe Petroski Such, Dave Cummings, Matthias Plappert, Fotios Chantzis, Elizabeth Barnes, Ariel Herbert-Voss, William Hebgen Guss, Alex Nichol, Alex Paino, Nikolas Tezak, Jie Tang, Igor Babuschkin, Suchir Balaji, Shantanu Jain, William Saunders, Christopher Hesse, Andrew N. Carr, Jan Leike, Josh Achiam, Vedant Misra, Evan Morikawa, Alec Radford, Matthew Knight, Miles Brundage, Mira Murati, Katie Mayer, Peter Welinder, Bob McGrew, Dario Amodei, Sam McCandlish, Ilya Sutskever, and Wojciech Zaremba. Evaluating large language models trained on code, 2021.
- [7] An-Chieh Cheng, Hongxu Yin, Yang Fu, Qiushan Guo, Ruihan Yang, Jan Kautz, Xiaolong Wang, and Sifei Liu. SpatialRGPT: Grounded spatial reasoning in vision-language models. In The Thirty-eighth Annual Conference on Neural Information Processing Systems, 2024.
- [8] Chuanqi Cheng, Jian Guan, Wei Wu, and Rui Yan. From the least to the most: Building a plug-and-play visual reasoner via data synthesis. In Yaser Al-Onaizan, Mohit Bansal, and Yun-Nung Chen, editors, Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing, pages 4941–4957, Miami, Florida, USA, November 2024. Association for Computational Linguistics.
- [9] Chuanqi Cheng, Jian Guan, Wei Wu, and Rui Yan. Scaling video-language models to 10k frames via hierarchical differential distillation, 2025.
- [10] DeepSeek-AI. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning, 2025.
- [11] Yihe Deng, Hritik Bansal, Fan Yin, Nanyun Peng, Wei Wang, and Kai-Wei Chang. Open-vlthinker: An early exploration to complex vision-language reasoning via iterative self-improvement. arXiv preprint arXiv:2503.17352, 2025.
- [12] Yu Du, Fangyun Wei, Ziheng Zhang, Miaoqing Shi, Yue Gao, and Guoqi Li. Learning to prompt for open-vocabulary object detection with vision-language model. In Proceedings of the IEEE/CVF conference on computer vision and pattern recognition, pages 14084–14093, 2022.

- [13] Kaituo Feng, Kaixiong Gong, Bohao Li, Zonghao Guo, Yibing Wang, Tianshuo Peng, Benyou Wang, and Xiangyu Yue. Video-r1: Reinforcing video reasoning in mllms. arXiv preprint arXiv:2503.21776, 2025.
- [14] Kanishk Gandhi, Ayush Chakravarthy, Anikait Singh, Nathan Lile, and Noah D. Goodman. Cognitive behaviors that enable self-improving reasoners, or, four habits of highly effective stars, 2025.
- [15] Howard E Gardner. Frames of mind: The theory of multiple intelligences. Basic books, 2011.
- [16] Google. Introducing gemini 2.0: our new ai model for the agentic era, 2024.
- [17] Kristen Grauman, Andrew Westbury, Eugene Byrne, Zachary Chavis, Antonino Furnari, Rohit Girdhar, Jackson Hamburger, Hao Jiang, Miao Liu, Xingyu Liu, Miguel Martin, Tushar Nagarajan, Ilija Radosavovic, Santhosh K. Ramakrishnan, Fiona Ryan, Jayant Sharma, Michael Wray, Mengmeng Xu, Eric Z. Xu, Chen Zhao, Siddhant Bansal, Dhruv Batra, Vincent Cartillier, Sean Crane, Tien Do, Morrie Doulaty, Akshay Erapalli, Christoph Feichtenhofer, Adriano Fragomeni, Qichen Fu, Christian Fuegen, Abrahm Kahsay Gebreselasie, Cristina González, James M. Hillis, Xuhua Huang, Yifei Huang, Wenqi Jia, Weslie Khoo, Jáchym Kolár, Satwik Kottur, Anurag Kumar, Federico Landini, Chao Li, Yanghao Li, Zhenqiang Li, Karttikeya Mangalam, Raghava Modhugu, Jonathan Munro, Tullie Murrell, Takumi Nishiyasu, Will Price, Paola Ruiz Puentes, Merey Ramazanova, Leda Sari, Kiran K. Somasundaram, Audrey Southerland, Yusuke Sugano, Ruijie Tao, Minh Vo, Yuchen Wang, Xindi Wu, Takuma Yagi, Yunyi Zhu, Pablo Arbeláez, David J. Crandall, Dima Damen, Giovanni Maria Farinella, Bernard Ghanem, Vamsi Krishna Ithapu, C. V. Jawahar, Hanbyul Joo, Kris Kitani, Haizhou Li, Richard A. Newcombe, Aude Oliva, Hyun Soo Park, James M. Rehg, Yoichi Sato, Jianbo Shi, Mike Zheng Shou, Antonio Torralba, Lorenzo Torresani, Mingfei Yan, and Jitendra Malik. Ego4d: Around the world in 3,000 hours of egocentric video. 2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), pages 18973–18990, 2021.
- [18] Jian Guan, Wei Wu, Peng Xu, Hongning Wang, Minlie Huang, et al. Amor: A recipe for building adaptable modular knowledge agents through process feedback. Advances in Neural Information Processing Systems, 37:126118–126148, 2024.
- [19] Tanmay Gupta and Aniruddha Kembhavi. Visual programming: Compositional visual reasoning without training. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, pages 14953–14962, 2023.
- [20] Jingcheng Hu, Yinmin Zhang, Qi Han, Daxin Jiang, Xiangyu Zhang, and Heung-Yeung Shum. Open-reasoner-zero: An open source approach to scaling up reinforcement learning on the base model, 2025.
- [21] Yushi Hu, Weijia Shi, Xingyu Fu, Dan Roth, Mari Ostendorf, Luke Zettlemoyer, Noah A Smith, and Ranjay Krishna. Visual sketchpad: Sketching as a visual chain of thought for multimodal language models. In The Thirty-eighth Annual Conference on Neural Information Processing Systems, 2024.
- [22] Drew A Hudson and Christopher D Manning. Gqa: A new dataset for real-world visual reasoning and compositional question answering. In Proceedings of the IEEE/CVF conference on computer vision and pattern recognition, pages 6700–6709, 2019.
- [23] Minyoung Huh, Brian Cheung, Tongzhou Wang, and Phillip Isola. Position: The platonic representation hypothesis. In Forty-first International Conference on Machine Learning, 2024.
- [24] Aaron Hurst, Adam Lerer, Adam P Goucher, Adam Perelman, Aditya Ramesh, Aidan Clark, AJ Ostrow, Akila Welihinda, Alan Hayes, Alec Radford, et al. Gpt-4o system card. arXiv preprint arXiv:2410.21276, 2024.
- [25] Michael Igorevich Ivanitskiy, Rusheb Shah, Alex F. Spies, Tilman Räuher, Dan Valentine, Can Rager, Lucia Quirke, Chris Mathwin, Guillaume Corlouer, Cecilia Diniz Behn, and Samy Wu Fung. A configurable library for generating and manipulating maze datasets, 2023.

- [26] Aaron Jaech, Adam Kalai, Adam Lerer, Adam Richardson, Ahmed El-Kishky, Aiden Low, Alec Helyar, Aleksander Madry, Alex Beutel, Alex Carney, et al. Openai o1 system card. arXiv preprint arXiv:2412.16720, 2024.
- [27] Amita Kamath, Jack Hessel, and Kai-Wei Chang. What’s "up" with vision-language models? investigating their struggle with spatial reasoning. arXiv preprint arXiv:2310.19785, 2023.
- [28] Alina Kuznetsova, Hassan Rom, Neil Alldrin, Jasper Uijlings, Ivan Krasin, Jordi Pont-Tuset, Shahab Kamali, Stefan Popov, Matteo Mallocci, Alexander Kolesnikov, et al. The open images dataset v4: Unified image classification, object detection, and visual relationship detection at scale. International journal of computer vision, 128(7):1956–1981, 2020.
- [29] Hung Le, Yue Wang, Akhilesh Deepak Gotmare, Silvio Savarese, and Steven Chu Hong Hoi. Coderl: Mastering code generation through pretrained models and deep reinforcement learning. Advances in Neural Information Processing Systems, 35:21314–21328, 2022.
- [30] Bo Li, Yuanhan Zhang, Dong Guo, Renrui Zhang, Feng Li, Hao Zhang, Kaichen Zhang, Peiyuan Zhang, Yanwei Li, Ziwei Liu, and Chunyuan Li. Llava-onevision: Easy visual task transfer, 2024.
- [31] Chengzu Li, Wenshan Wu, Huanyu Zhang, Yan Xia, Shaoguang Mao, Li Dong, Ivan Vulić, and Furu Wei. Imagine while reasoning in space: Multimodal visualization-of-thought. arXiv preprint arXiv:2501.07542, 2025.
- [32] Chengzu Li, Caiqi Zhang, Han Zhou, Nigel Collier, Anna Korhonen, and Ivan Vulić. Topviewr: Vision-language models as top-view spatial reasoners, 2024.
- [33] Songtao Li and Hao Tang. Multimodal alignment and fusion: A survey. arXiv preprint arXiv:2411.17040, 2024.
- [34] Xuefeng Li, Haoyang Zou, and Pengfei Liu. Torl: Scaling tool-integrated rl, 2025.
- [35] Hunter Lightman, Vineet Kosaraju, Yura Burda, Harri Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. Let’s verify step by step. arXiv preprint arXiv:2305.20050, 2023.
- [36] Ji Lin, Hongxu Yin, Wei Ping, Pavlo Molchanov, Mohammad Shoeybi, and Song Han. Vila: On pre-training for visual language models. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, pages 26689–26699, 2024.
- [37] Benlin Liu, Yuhao Dong, Yiqin Wang, Zixian Ma, Yansong Tang, Luming Tang, Yongming Rao, Wei-Chiu Ma, and Ranjay Krishna. Coarse correspondences boost spatial-temporal reasoning in multimodal language model, 2024.
- [38] Fangyu Liu, Guy Emerson, and Nigel Collier. Visual spatial reasoning. Transactions of the Association for Computational Linguistics, 11:635–651, 2023.
- [39] Haotian Liu, Chunyuan Li, Yuheng Li, and Yong Jae Lee. Improved baselines with visual instruction tuning. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, pages 26296–26306, 2024.
- [40] Muhammad Maaz, Hanoona Rasheed, Salman Khan, and Fahad Khan. Video-chatgpt: Towards detailed video understanding via large vision and language models. In Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers), pages 12585–12602, 2024.
- [41] Damiano Marsili, Rohun Agrawal, Yisong Yue, and Georgia Gkioxari. Visual agentic ai for spatial reasoning with a dynamic api. In Proceedings of the Computer Vision and Pattern Recognition Conference, pages 19446–19455, 2025.
- [42] Julia McAfoose and BT Baune. Exploring visual-spatial working memory: A critical review of concepts and models. Neuropsychology review, 19:130–142, 2009.

- [43] Roshanak Mirzaee, Hossein Rajaby Faghihi, Qiang Ning, and Parisa Kordjamshidi. SPARTQA: A textual question answering benchmark for spatial reasoning. In Kristina Toutanova, Anna Rumshisky, Luke Zettlemoyer, Dilek Hakkani-Tur, Iz Beltagy, Steven Bethard, Ryan Cotterell, Tanmoy Chakraborty, and Yichao Zhou, editors, Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, pages 4582–4598, Online, June 2021. Association for Computational Linguistics.
- [44] OpenAI. Hello gpt-4o. In OpenAI Blog, 2024.
- [45] OpenAI. Introducing openai o1-preview. <https://openai.com/index/introducing-openai-o1-preview/>, 2024.
- [46] OpenAI. Introducing openai o3 and o4-mini, 2025.
- [47] OpenAI. Thinking with images, 2025.
- [48] Kun Ouyang, Yuanxin Liu, Haoning Wu, Yi Liu, Hao Zhou, Jie Zhou, Fandong Meng, and Xu Sun. Spacer: Reinforcing mllms in video spatial reasoning, 2025.
- [49] Abby O’Neill, Abdul Rehman, Abhiram Maddukuri, Abhishek Gupta, Abhishek Padalkar, Abraham Lee, Acorn Pooley, Agrim Gupta, Ajay Mandlekar, Ajinkya Jain, Albert Tung, Alex Bewley, Alex Herzog, Alex Irpan, Alexander Khazatsky, Anant Rai, Anchit Gupta, Andrew Wang, Anikait Singh, Animesh Garg, Aniruddha Kembhavi, Annie Xie, Anthony Brohan, Antonin Raffin, Archit Sharma, Arefeh Yavary, Arhan Jain, Ashwin Balakrishna, Ayzaan Wahid, Ben Burgess-Limerick, Beomjoon Kim, Bernhard Schölkopf, Blake Wulfe, Brian Ichter, Cewu Lu, Charles Xu, Charlotte Le, Chelsea Finn, Chen Wang, Chenfeng Xu, Cheng Chi, Chenguang Huang, Christine Chan, Christopher Agia, Chuer Pan, Chuyuan Fu, Coline Devin, Danfei Xu, Daniel Morton, Danny Driess, Daphne Chen, Deepak Pathak, Dhruv Shah, Dieter Buehler, Dinesh Jayaraman, Dmitry Kalashnikov, Dorsa Sadigh, Edward Johns, Ethan Foster, Fangchen Liu, Federico Ceola, Fei Xia, Feiyu Zhao, Freek Stulp, Gaoyue Zhou, Gaurav S. Sukhatme, Gautam Salhotra, Ge Yan, Gilbert Feng, Giulio Schiavi, Glen Berseth, Gregory Kahn, Guanzhi Wang, Hao Su, Hao-Shu Fang, Haochen Shi, Henghui Bao, Heni Ben Amor, Henrik I Christensen, Hiroki Furuta, Homer Walke, Hongjie Fang, Huy Ha, Igor Mordatch, Ilija Radosavovic, Isabel Leal, Jacky Liang, Jad Abou-Chakra, Jaehyung Kim, Jaimyn Drake, Jan Peters, Jan Schneider, Jasmine Hsu, Jeannette Bohg, Jeffrey Bingham, Jeffrey Wu, Jensen Gao, Jiaheng Hu, Jiajun Wu, Jialin Wu, Jiankai Sun, Jianlan Luo, Jiayuan Gu, Jie Tan, Jihoon Oh, Jimmy Wu, Jingpei Lu, Jingyun Yang, Jitendra Malik, João Silvério, Joey Hejna, Jonathan Booyer, Jonathan Tompson, Jonathan Yang, Jordi Salvador, Joseph J. Lim, Junhyek Han, Kaiyuan Wang, Kanishka Rao, Karl Pertsch, Karol Hausman, Keegan Go, Keerthana Gopalakrishnan, Ken Goldberg, Kendra Byrne, Kenneth Oslund, Kento Kawaharazuka, Kevin Black, Kevin Lin, Kevin Zhang, Kiana Ehsani, Kiran Lekkala, Kirsty Ellis, Krishan Rana, Krishnan Srinivasan, Kuan Fang, Kunal Pratap Singh, Kuo-Hao Zeng, Kyle Hatch, Kyle Hsu, Laurent Itti, Lawrence Yunliang Chen, Lerrel Pinto, Li Fei-Fei, Liam Tan, Linxi Jim Fan, Lionel Ott, Lisa Lee, Luca Weihs, Magnum Chen, Marion Lepert, Marius Memmel, Masayoshi Tomizuka, Masha Itkina, Mateo Guaman Castro, Max Spero, Maximilian Du, Michael Ahn, Michael C. Yip, Mingtong Zhang, Mingyu Ding, Minh Heo, Mohan Kumar Srirama, Mohit Sharma, Moo Jin Kim, Naoaki Kanazawa, Nicklas Hansen, Nicolas Heess, Nikhil J Joshi, Niko Suenderhauf, Ning Liu, Norman Di Palo, Nur Muhammad Mahi Shafiullah, Oier Mees, Oliver Kroemer, Osbert Bastani, Pannag R Sanketi, Patrick Tree Miller, Patrick Yin, Paul Wohlhart, Peng Xu, Peter David Fagan, Peter Mitrano, Pierre Sermanet, Pieter Abbeel, Priya Sundaresan, Qiuyu Chen, Quan Vuong, Rafael Rafailov, Ran Tian, Ria Doshi, Roberto Martín-Martín, Rohan Bajjal, Rosario Scalise, Rose Hendrix, Roy Lin, Runjia Qian, Ruohan Zhang, Russell Mendonca, Rutav Shah, Ryan Hoque, Ryan Julian, Samuel Bustamante, Sean Kirmani, Sergey Levine, Shan Lin, Sherry Moore, Shikhar Bahl, Shivin Dass, Shubham Sonawani, Shuran Song, Sichun Xu, Siddhant Halder, Siddharth Karamcheti, Simeon Adebola, Simon Guist, Soroush Nasiriany, Stefan Schaal, Stefan Welker, Stephen Tian, Subramanian Ramamoorthy, Sudeep Dasari, Suneel Belkhale, Sungjae Park, Suraj Nair, Suvir Mirchandani, Takayuki Osa, Tanmay Gupta, Tatsuya Harada, Tatsuya Matsushima, Ted Xiao, Thomas Kollar, Tianhe Yu, Tianli Ding, Todor Davchev, Tony Z. Zhao, Travis Armstrong, Trevor Darrell, Trinity Chung, Vidhi Jain, Vincent Vanhoucke, Wei Zhan, Wenxuan Zhou, Wolfram Burgard, Xi Chen, Xiaolong Wang, Xinghao Zhu, Xinyang Geng, Xiyuan Liu, Xu Liangwei, Xuanlin Li, Yao Lu, Yecheng Jason

- Ma, Yejin Kim, Yevgen Chebotar, Yifan Zhou, Yifeng Zhu, Yilin Wu, Ying Xu, Yixuan Wang, Yonatan Bisk, Yoonyoung Cho, Youngwoon Lee, Yuchen Cui, Yue Cao, Yueh-Hua Wu, Yujin Tang, Yuke Zhu, Yunchu Zhang, Yunfan Jiang, Yunshuang Li, Yunzhu Li, Yusuke Iwasawa, Yutaka Matsuo, Zehan Ma, Zhuo Xu, Zichen Jeff Cui, Zichen Zhang, and Zipeng Lin. Open x-embodiment: Robotic learning datasets and rt-x models : Open x-embodiment collaboration0. In *2024 IEEE International Conference on Robotics and Automation (ICRA)*, pages 6892–6903, 2024.
- [50] Yi Peng, Xiaokun Wang, Yichen Wei, Jiangbo Pei, Weijie Qiu, Ai Jian, Yunzhuo Hao, Jiachun Pan, Tianyidan Xie, Li Ge, et al. Skywork r1v: pioneering multimodal reasoning with chain-of-thought. *arXiv preprint arXiv:2504.05599*, 2025.
- [51] Ji Qi, Ming Ding, Weihang Wang, Yushi Bai, Qingsong Lv, Wenyi Hong, Bin Xu, Lei Hou, Juanzi Li, Yuxiao Dong, and Jie Tang. Cogcom: A visual language model with chain-of-manipulations reasoning. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [52] Zhangyang Qi, Zhixiong Zhang, Ye Fang, Jiaqi Wang, and Hengshuang Zhao. Gpt4scene: Understand 3d scenes from videos with vision-language models. *arXiv preprint arXiv:2501.01428*, 2025.
- [53] Yujia Qin, Shihao Liang, Yining Ye, Kunlun Zhu, Lan Yan, Yaxi Lu, Yankai Lin, Xin Cong, Xiangru Tang, Bill Qian, Sihan Zhao, Lauren Hong, Runchu Tian, Ruobing Xie, Jie Zhou, Mark Gerstein, dahai li, Zhiyuan Liu, and Maosong Sun. ToolLLM: Facilitating large language models to master 16000+ real-world APIs. In *The Twelfth International Conference on Learning Representations*, 2024.
- [54] Timo Schick, Jane Dwivedi-Yu, Roberto Dessì, Roberta Raileanu, Maria Lomeli, Eric Hambro, Luke Zettlemoyer, Nicola Cancedda, and Thomas Scialom. Toolformer: Language models can teach themselves to use tools. *Advances in Neural Information Processing Systems*, 36, 2024.
- [55] Hao Shao, Shengju Qian, Han Xiao, Guanglu Song, Zhuofan Zong, Letian Wang, Yu Liu, and Hongsheng Li. Visual cot: Advancing multi-modal language models with a comprehensive dataset and benchmark for chain-of-thought reasoning. *Advances in Neural Information Processing Systems*, 37:8612–8642, 2024.
- [56] Bo Shen, Jiaxin Zhang, Taihong Chen, Daoguang Zan, Bing Geng, An Fu, Muhan Zeng, Ailun Yu, Jichuan Ji, Jingyang Zhao, et al. Pangu-coder2: Boosting large language models for code with ranking feedback. *arXiv preprint arXiv:2307.14936*, 2023.
- [57] Yongliang Shen, Kaitao Song, Xu Tan, Dongsheng Li, Weiming Lu, and Yueting Zhuang. Hugginggpt: Solving ai tasks with chatgpt and its friends in huggingface. In *Advances in Neural Information Processing Systems*, 2023.
- [58] Guangming Sheng, Chi Zhang, Zilingfeng Ye, Xibin Wu, Wang Zhang, Ru Zhang, Yanghua Peng, Haibin Lin, and Chuan Wu. Hybridflow: A flexible and efficient rlhf framework. *arXiv preprint arXiv: 2409.19256*, 2024.
- [59] Charlie Victor Snell, Jaehoon Lee, Kelvin Xu, and Aviral Kumar. Scaling LLM test-time compute optimally can be more effective than scaling parameters for reasoning. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [60] Chan Hee Song, Valts Blukis, Jonathan Tremblay, Stephen Tyree, Yu Su, and Stan Birchfield. RoboSpatial: Teaching spatial understanding to 2D and 3D vision-language models for robotics. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2025. Oral Presentation.
- [61] Dídac Surís, Sachit Menon, and Carl Vondrick. Vipergpt: Visual inference via python execution for reasoning. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 11888–11898, 2023.
- [62] Gemini Team, Petko Georgiev, Ving Ian Lei, Ryan Burnell, Libin Bai, Anmol Gulati, Garrett Tanzer, Damien Vincent, Zhufeng Pan, Shibo Wang, et al. Gemini 1.5: Unlocking multimodal understanding across millions of tokens of context. *arXiv preprint arXiv:2403.05530*, 2024.

- [63] Kimi Team, Angang Du, Bohong Yin, Bowei Xing, Bowen Qu, Bowen Wang, Cheng Chen, Chenlin Zhang, Chenzhuang Du, Chu Wei, Congcong Wang, Dehao Zhang, Dikang Du, Dongliang Wang, Enming Yuan, Enzhe Lu, Fang Li, Flood Sung, Guangda Wei, Guokun Lai, Han Zhu, Hao Ding, Hao Hu, Hao Yang, Hao Zhang, Haoning Wu, Haotian Yao, Haoyu Lu, Heng Wang, Hongcheng Gao, Huabin Zheng, Jiaming Li, Jianlin Su, Jianzhou Wang, Jiaqi Deng, Jiezhong Qiu, Jin Xie, Jinhong Wang, Jingyuan Liu, Junjie Yan, Kun Ouyang, Liang Chen, Lin Sui, Longhui Yu, Mengfan Dong, Mengnan Dong, Nuo Xu, Pengyu Cheng, Qizheng Gu, Runjie Zhou, Shaowei Liu, Sihan Cao, Tao Yu, Tianhui Song, Tongtong Bai, Wei Song, Weiran He, Weixiao Huang, Weixin Xu, Xiaokun Yuan, Xingcheng Yao, Xingzhe Wu, Xinxing Zu, Xinyu Zhou, Xinyuan Wang, Y. Charles, Yan Zhong, Yang Li, Yangyang Hu, Yanru Chen, Yejie Wang, Yibo Liu, Yibo Miao, Yidao Qin, Yimin Chen, Yiping Bao, Yiqin Wang, Yongsheng Kang, Yuanxin Liu, Yulun Du, Yuxin Wu, Yuzhi Wang, Yuzi Yan, Zaida Zhou, Zhaowei Li, Zhejun Jiang, Zheng Zhang, Zhilin Yang, Zhiqi Huang, Zihao Huang, Zijia Zhao, Ziwei Chen, and Zongyu Lin. Kimi-vl technical report, 2025.
- [64] Jiayu Wang, Yifei Ming, Zhenmei Shi, Vibhav Vineet, Xin Wang, Yixuan Li, and Neel Joshi. Is a picture worth a thousand words? delving into spatial reasoning for vision language models. In A. Globerson, L. Mackey, D. Belgrave, A. Fan, U. Paquet, J. Tomczak, and C. Zhang, editors, Advances in Neural Information Processing Systems, volume 37, pages 75392–75421. Curran Associates, Inc., 2024.
- [65] Jiayu Wang, Yifei Ming, Zhenmei Shi, Vibhav Vineet, Xin Wang, Yixuan Li, and Neel Joshi. Is a picture worth a thousand words? delving into spatial reasoning for vision language models. In The Thirty-Eighth Annual Conference on Neural Information Processing Systems, 2024.
- [66] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, et al. Chain-of-thought prompting elicits reasoning in large language models. Advances in neural information processing systems, 35:24824–24837, 2022.
- [67] Jiaxin Wen, Jian Guan, Hongning Wang, Wei Wu, and Minlie Huang. Codeplan: Unlocking reasoning potential in large language models by scaling code-form planning. In The Thirteenth International Conference on Learning Representations, 2025.
- [68] Tian Xie, Zitian Gao, Qingnan Ren, Haoming Luo, Yuqian Hong, Bryan Dai, Joey Zhou, Kai Qiu, Zhirong Wu, and Chong Luo. Logic-rl: Unleashing llm reasoning with rule-based reinforcement learning, 2025.
- [69] Yutaro Yamada, Yihan Bao, Andrew K Lampinen, Jungo Kasai, and Ilker Yildirim. Evaluating spatial understanding of large language models. arXiv preprint arXiv:2310.14540, 2023.
- [70] Jihan Yang, Shusheng Yang, Anjali Gupta, Rilyn Han, Li Fei-Fei, and Saining Xie. Thinking in Space: How Multimodal Large Language Models See, Remember and Recall Spaces. arXiv preprint arXiv:2412.14171, 2024.
- [71] Sihan Yang, Runsen Xu, Yiman Xie, Sizhe Yang, Mo Li, Jingli Lin, Chenming Zhu, Xiaochen Chen, Haodong Duan, Xiangyu Yue, Dahua Lin, Tai Wang, and Jiangmiao Pang. Mmsi-bench: A benchmark for multi-image spatial intelligence, 2025.
- [72] Yi Yang, Xiaoxuan He, Hongkun Pan, Xiyan Jiang, Yan Deng, Xingtao Yang, Haoyu Lu, Dacheng Yin, Fengyun Rao, Minfeng Zhu, et al. R1-onevision: Advancing generalized multimodal reasoning through cross-modal formalization. arXiv preprint arXiv:2503.10615, 2025.
- [73] Zhengyuan Yang, Linjie Li, Jianfeng Wang, Kevin Lin, Ehsan Azarnasab, Faisal Ahmed, Zicheng Liu, Ce Liu, Michael Zeng, and Lijuan Wang. Mm-react: Prompting chatgpt for multimodal reasoning and action. arXiv preprint arXiv:2303.11381, 2023.
- [74] Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik R Narasimhan, and Yuan Cao. React: Synergizing reasoning and acting in language models. In The Eleventh International Conference on Learning Representations, 2022.

- [75] Longhui Yu, Weisen Jiang, Han Shi, YU Jincheng, Zhengying Liu, Yu Zhang, James Kwok, Zhenguo Li, Adrian Weller, and Weiyang Liu. Metamath: Bootstrap your own mathematical questions for large language models. In The Twelfth International Conference on Learning Representations, 2023.
- [76] Yang Yue, Zhiqi Chen, Rui Lu, Andrew Zhao, Zhaokai Wang, Yang Yue, Shiji Song, and Gao Huang. Does reinforcement learning really incentivize reasoning capacity in llms beyond the base model?, 2025.
- [77] Huanyu Zhang, Chengzu Li, Wenshan Wu, Shaoguang Mao, Yifan Zhang, Haochen Tian, Ivan Vulić, Zhang Zhang, Liang Wang, Tieniu Tan, and Furu Wei. Scaling and beyond: Advancing spatial reasoning in mllms requires new recipes, 2025.
- [78] Jiahui Zhang, Yurui Chen, Yanpeng Zhou, Yueming Xu, Ze Huang, Jilin Mei, Junhui Chen, Yu-Jie Yuan, Xinyue Cai, Guowei Huang, Xingyue Quan, Hang Xu, and Li Zhang. From flatland to space: Teaching vision-language models to perceive and reason in 3d, 2025.
- [79] Peng Zhang, Can Li, Liang Qiao, Zhanzhan Cheng, Shiliang Pu, Yi Niu, and Fei Wu. Vsr: a unified framework for document layout analysis combining vision, semantics and relations. In Document Analysis and Recognition–ICDAR 2021: 16th International Conference, Lausanne, Switzerland, September 5–10, 2021, Proceedings, Part I 16, pages 115–130. Springer, 2021.
- [80] Yuanhan Zhang, Jinming Wu, Wei Li, Bo Li, Zejun Ma, Ziwei Liu, and Chunyuan Li. Video instruction tuning with synthetic data. arXiv preprint arXiv:2410.02713, 2024.
- [81] Zhuosheng Zhang, Aston Zhang, Mu Li, hai zhao, George Karypis, and Alex Smola. Multimodal chain-of-thought reasoning in language models. Transactions on Machine Learning Research, 2024.
- [82] Xueliang Zhao, Wei Wu, Jian Guan, and Lingpeng Kong. Promptcot: Synthesizing olympiad-level problems for mathematical reasoning in large language models, 2025.
- [83] Chenming Zhu, Tai Wang, Wenwei Zhang, Jiangmiao Pang, and Xihui Liu. Llava-3d: A simple yet effective pathway to empowering lmms with 3d-awareness, 2025.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: We clearly state its key contributions in both abstract and introduction: proposing a novel “drawing to reason in space” paradigm, developing a three-stage training framework, and demonstrating state-of-the-art performance across multiple spatial reasoning benchmarks. These claims are thoroughly supported by detailed methodology descriptions and comprehensive experimental results.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We create a separate “Limitations” section, i.e., §F.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [NA]

Justification: Our paper does not introduce theoretical results.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: The paper provides comprehensive reproducibility information across all key components: (1) detailed description of the three-stage training framework with specific hyperparameters, (2) explicit formulation of reward functions and filtering criteria, (3) clear specification of dataset construction including maze generation parameters and data filtering processes, and (4) complete model architecture details based on the open-source Qwen2.5-VL backbone.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).

- (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: We open-source the data, code, and model, with links provided in this paper.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: The paper provides comprehensive experimental details, including: (1) explicit dataset splits and construction process, (2) complete hyperparameters for all three training stages, (3) detailed reward function formulations and thresholds, and (4) specific model configurations based on Qwen2.5-VL. Key implementation choices and their rationales are discussed in the main text, with full technical details provided in the appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: Please see §C.4.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: The paper specifies compute requirements for each training stage in §4.1.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: We have conformed the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: We create a separate “Broader impact” section, i.e., §G.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [\[Yes\]](#)

Justification: We discuss our safeguards in §G.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [\[Yes\]](#)

Justification: The paper properly credits all existing assets. All datasets are used within their license terms, with proper citations to original papers and links to source repositories provided in the appendix.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.

- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [\[Yes\]](#)

Justification: The paper thoroughly documents new assets through detailed specifications of the three-stage training framework and model architecture and comprehensive documentation of our dataset construction process. We include a README document with our code in the supplementary material.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [\[NA\]](#)

Justification: We do not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [\[NA\]](#)

Justification: Our research focuses on algorithmic development and model training using existing public datasets, without involving any human subjects or crowdsourcing efforts. Therefore, IRB approval was not required for this work.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [Yes]

Justification: Yes, our work involves LLMs in two significant aspects: (1) we leverage Qwen2.5-VL-72B to generate initial reasoning trajectories for constructing the ViLASR-ColdStart-33k dataset, and (2) our method builds upon and fine-tunes the Qwen2.5-VL model architecture. Both uses are explicitly documented in our methodology section, with detailed specifications of how these LLMs contribute to our training pipeline.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

A Appendix outline

In these supplementary materials, we provide:

- Dataset construction (Appendix B);
- Experimental setup and full evaluation results (Appendix C);
- Visualization results (Appendix D).
- Complexity comparison between ViLASR and baselines (Appendix E);
- Discussion on limitations (Appendix F) and broader impact (Appendix G) of ViLASR.

B Dataset construction

We created diverse datasets spanning three categories of spatial reasoning tasks: maze navigation that tests path planning abilities, static image understanding that focuses on spatial relationships, and video comprehension that captures temporal evolution and multi-view spatial reasoning. Building upon basic spatial understanding in static images through bounding box annotation and auxiliary line drawing, our video data further extends these capabilities by incorporating temporal dynamics and multi-view perspectives, leading to more comprehensive spatial reasoning abilities.

Our data collection pipeline leverages multiple public datasets, with careful balancing across different task types to ensure comprehensive coverage. This effort resulted in three primary datasets: **ViLASR-ColdStart-33k**, a supervised dataset containing curated reasoning paths, **ViLASR-RRS-8k**, specifically designed for the reflective rejection sampling stage, and **ViLASR-RL-40k**, optimized for reinforcement learning. For ViLASR-ColdStart-33k construction, we leveraged Qwen2.5-VL-72B [2] to generate initial reasoning trajectories following our “drawing to reason in space” paradigm detailed in §3.1. A rigorous filtering process based on answer correctness and format validity helped identify the most reliable samples. We applied two format criteria: (1) ensuring all drawing operations are executable and the final answer can be correctly parsed, and (2) requiring at least three rounds of thinking to ensure sufficient spatial reasoning depth. After cold-start training, we constructed ViLASR-RRS-8k by selecting instances that are generated by the cold-started model and exhibit clear self-correction patterns and correct final answers, providing ideal training examples for cultivating reflective behaviors.

Table 2: Dataset distribution across different training stages.

Subset (Modality)	ViLASR-ColdStart-33k	ViLASR-RRS-8k	ViLASR-RL-40k
Maze Navigation (Image)	7,000	1,000	10,000
VQA (Image)	10,000	3,000	10,000
GPT4Scene [52] (Video)	6,000	2,000	10,000
SR-91k [48] (Video)	10,000	1,800	10,000
Total	33,000	7,800	40,000

Table 2 presents the detailed distribution of these datasets. Our source data encompasses a broad range of spatial reasoning scenarios, drawing from GQA [22], VSR [79], OpenImages [28], OpenSpaces [5] and SpaceLLaVA [5]. For maze navigation tasks, we procedurally generated mazes with varying grid sizes (from 3×3 to 6×6) using a depth-first search algorithm³, maintaining approximately equal proportions across different grid sizes. Each maze consists of a clearly marked starting point and four candidate destinations (labeled as A, B, C, and D). The questions are formulated as multiple-choice problems where the model must determine which destination would be reached by following a given action sequence (e.g., “Determine the final destination from the starting point (green point). Action Sequence: Go up. Go up ...”). This automated generation process ensures diverse maze layouts and action sequences while maintaining problem solvability, creating a controlled environment for evaluating sequential spatial reasoning capabilities.

³<https://github.com/understanding-search/maze-dataset>

Table 3: APIs used for proprietary models evaluation.

Model	API Name	Provider
GPT-4o	gpt-4o-2024-08-06	OpenAI
Gemini-1.5-Flash	gemini-1.5-flash	Google
Gemini-1.5-Pro	gemini-1.5-pro	Google
Gemini-2.0-Flash	gemini-2.0-flash	Google
OpenAI o4-mini	o4-mini-2025-04-16	OpenAI

Table 4: Number of examples in our evaluation benchmark.

	Image	Video	Multi-view	
MAZE	SpatialEval-Real	VSI-Bench	SPAR-Bench	MMSI-Bench
2,000	135	5,130	7,211	1,000

C Experimental setup and full evaluation results

C.1 Baseline implementation

We evaluate ViLASR against various state-of-the-art models and methods.

For closed-source LVLMs, we directly query their official APIs as shown in Table 3. We evaluate these models using zero-shot prompting by directly providing the input images and questions. Under the setting without reasoning, we explicitly prompt these models to output answers directly without intermediate reasoning steps, i.e., “Answer with the option’s letter from the given choices directly” for multiple-choice questions and “Please answer with a single numerical value (e.g., 42 or 3.14)” for numerical responses.

For open-source LVLMs ranging from 7B parameters (Qwen2.5-VL-7B [2], LLaVA-NeXT-Video-7B [80], LLaVA-OneVision-7B [30]) to 72B parameters (Qwen2.5-VL-72B [2], LLaVA-NeXT-Video-72B [80], LLaVA-OneVision-72B [30]). Kimi-VL-A3B-Instruct(16B) [63] leverages Mixture-of-Experts (MoE) architecture with 16.0B total parameters, while dynamically activating 2.8B parameters during inference. We conduct zero-shot evaluation using their standard prompting formats.

For specialized reasoning models, we evaluate using their officially released checkpoints and prompts: COGCOM [51] with its built-in perception tools, VisCoT [55] with its bounding box annotation pipeline, and SpaceR [48] with its video understanding capabilities.

C.2 Benchmark statistics

Table 4 shows the benchmark statistics.

C.3 Detailed results on VSI-Bench

Table 5 provides detailed results on VSI-Bench. ViLASR achieves state-of-the-art performance with an average accuracy of 45.4%, outperforming all baseline methods by a significant margin (+12.7%). The improvements are particularly pronounced in tasks requiring precise spatial measurements and object localization: “Object Count” (+29.0%), “Absolute Distance” (+15.0%), and “Object Size” (+13.0%). This aligns with our case study observations (see §D) where ViLASR demonstrates superior capabilities in systematic measurement and spatial reasoning through drawing operations.

Furthermore, the strong performance in “Relative Distance” (+16.1%) and “Relative Direction” (+20.7%) demonstrates ViLASR’s effectiveness in comparative spatial reasoning. By explicitly drawing auxiliary lines to connect and measure between objects, our model can more accurately assess relative positions and orientations. The significant improvement in “Appearance Order” (+19.8%) further highlights ViLASR’s capability in temporal-spatial reasoning, where systematic annotation of objects across multiple frames helps track and verify their sequential relationships.

Table 5: Results on VSI-Bench. Gray-shaded rows represent large-sized models ($>7B$ parameters). [†] Results from [70]. [‡] Results from VSI-Bench (tiny) set [70]. * Results from [63].

Method	Sub-tasks								Average
	Numerical questions				Multiple-choice questions				
	Object Count	Absolute Distance	Object Size	Room Size	Relative Distance	Relative Direction	Route Plan	Apperance Order	
Proprietary LVLMS									
GPT-4o	46.2	5.3	43.8	38.2	37.0	41.3	31.5	28.5	34.0 [†]
Gemini-1.5-Flash	49.8	30.8	53.5	54.4	37.7	41.0	31.5	37.8	42.1 [†]
Gemini-1.5-Pro	56.2	30.9	64.1	43.6	51.3	46.3	36.0	34.6	45.4 [†]
Gemini-2.0-Flash	52.4	30.6	66.7	31.8	56.0	46.3	24.5	55.1	45.4 [‡]
Open-source LVLMS									
Qwen2.5-VL-7B	34.5	19.4	47.6	40.8	32.8	24.5	<u>32.5</u>	29.4	32.7
LLaVA-NeXT-Video-7B	48.5	14.0	47.8	24.2	<u>43.5</u>	42.4	34.0	30.6	35.6 [†]
LLaVA-OneVision-7B	47.7	20.2	47.4	12.3	42.5	35.2	29.4	24.4	32.4 [†]
Kimi-VL-A3B-Instruct-16B	-	-	-	-	-	-	-	-	37.4 [*]
Qwen2.5-VL-72B	33.9	27.2	59.3	28.5	47.2	35.3	22.2	34.5	36.0
LLaVA-NeXT-Video-72B	48.9	22.8	57.4	35.3	42.4	36.7	35.0	48.6	40.9 [†]
LLaVA-OneVision-72B	43.5	23.9	57.6	37.5	42.5	39.9	32.5	44.6	40.2 [†]
Representative methods for multimodal reasoning									
SpaceR-7B	<u>61.9</u>	<u>28.6</u>	60.9	<u>35.2</u>	38.2	46.0	31.4	<u>45.6</u>	<u>43.5</u>
Ours									
ViLASR	63.5	34.4	<u>60.6</u>	30.9	48.9	<u>45.2</u>	30.4	49.2	45.4
Improvement	+29.0	+15.0	+13.0	-9.9	+16.1	+20.7	-2.1	+19.8	+12.7

Interestingly, while showing strong performance in most categories, ViLASR exhibits slight decreases in “Room Size” (-9.9%) and “Route Plan” (-2.1%). This suggests a limitation in tasks requiring holistic reasoning: “Room Size” estimation demands understanding the complete room layout from partial views and reasoning about unseen spaces, while “Route Plan” needs global path planning across multiple viewpoints. Unlike localized spatial measurements that can be solved through explicit drawing operations, these tasks require more sophisticated global inference capabilities to integrate information across different views and time points. This reveals potential directions for extending our drawing-based framework: we could introduce specialized drawing tools for layout reconstruction and space completion. Such extensions would be natural additions to our current drawing operation set, enabling more comprehensive spatial reasoning capabilities.

C.4 Statistical significance analysis

We conduct rigorous statistical significance tests to validate our experimental results. Using paired t -tests, we find that ViLASR significantly outperforms the base model (Qwen2.5-VL-7B) on MAZE, SpatialEval-Real, VSI-Bench, SPAR-Bench and MMSI-Bench ($p < 0.05$). These results demonstrate the robust and consistent advantages of our approach across different spatial reasoning scenarios.

C.5 Prompt template

In this section, we present the comprehensive prompt templates utilized in ViLASR, including both the system prompt and user prompt, as illustrated in Figure 6 and Figure 7, respectively. We use these prompt templates for both generating reasoning paths in cold start data and inference of ViLASR.

C.6 System prompt of ViLASR

The system prompt utilized in the ViLASR reasoning framework is presented in Figure 6.

C.7 Query prompt template of ViLASR

The query prompt utilized in the ViLASR reasoning framework is presented in Figure 7. Our framework utilizes three distinct types of prompts throughout the reasoning process:

System prompt

Guidance:

You are a spatial reasoning assistant with access to two powerful visualization tools.

Your task is to break down complex spatial problems and iteratively refine your solution through visualization feedback.

Available tools:

You can use the following two tools to visualize. After each tool usage, you must wait for and analyze the visualization feedback before proceeding.

1. **Object Mapper**

- Purpose: Identifies and maps key items in the space

- Input format: JSON

```
```json
[{{
 "index": i, # Image index
 "bbox_2d": [x1, y1, x2, y2],
 "label": "object name/description"
}}]
```
```

- Output: Generates bounding boxes for visual inspection of the i-th image

2. **Path Tracer**

- Purpose: Plots movement or connections between points

- Input format: JSON

```
```json
[{{
 "index": i, # Image index
 "start_point_2d": [x1, y1],
 "end_point_2d": [x2, y2],
 "label": "trace_description"
}}]
```
```

- Output: Generates visual paths for verification of the i-th image

Required Output Format:

For each reasoning step, you must structure your response as follows:

<think> [Your detailed reasoning process] </think> Action: [Object Mapper/Path Tracer]

```
```json
[JSON format coordinates]
```
```

After your reasoning and iteratively refine your solution through visualization feedback, you should arrive at a final answer and structure your response as follows:

<think> [Your detailed reasoning process] </think> Action: Answer

<answer> [Your final answer] </answer>

Please NOTE the following reasoning techniques:

1. Initial Analysis

- Break down the spatial problem

- Plan your approach

2. Iterative Reasoning for Each Step

- Choose appropriate tool

- Provide absolute coordinates in JSON format (The top-left corner of the image is (0, 0) and the bottom-right corner is (**{width}**, **{height}**))

- Observe the visualization output

- Reflect on the visualization:

* Is the placement/path accurate?

* Does it align with your reasoning?

* What adjustments are needed?

- Backtrack and Adjust:

* If errors found, backtrack to previous step to modify actions or decisions as needed

Figure 6: System prompt used in ViLASR.

- **Initial Query Prompt** (Image/Video Query Prompt): Initiates the reasoning process by establishing the initial visual context and query.
- **Follow-up Query Prompt**: Guides subsequent reasoning steps by incorporating intermediate results and accumulated context.
- **Final Answer Prompt**: Terminates the reasoning process and enforces answer generation when either the maximum number of processed images or reasoning steps is reached.



Figure 7: Query prompt and output template used in ViLASR.

D Visualization results

To illustrate the effectiveness of our approach, we present three representative examples from the MAZE, SPAR-Bench, and VSI-Bench benchmarks in Figures 8, 9, and 10, respectively. Note that the “Path Tracer” tool refers to drawing auxiliary lines, and the “Object Mapper” tool refers to annotating bounding boxes.

In the maze navigation task (Figure 8), while Qwen2.5-VL-7B provides incorrect answers directly and GPT-4o attempts textual reasoning but fails to accurately track spatial transitions, ViLASR successfully decomposes the task into interpretable steps. By using the “Path Tracer” tool to visualize and verify each movement through auxiliary lines, our model ensures accurate navigation through the maze, leading to the correct destination.

For multi-view reasoning (Figure 9), ViLASR demonstrates sophisticated spatial understanding by first mapping key objects (“window,” “air conditioner”) in both images using “Object Mapper”, then systematically analyzing their relative positions. Through careful verification using “Path Tracer,” it correctly determines the observer’s position change, while both baseline models struggle with perspective transformations and make incorrect judgments about movement direction.

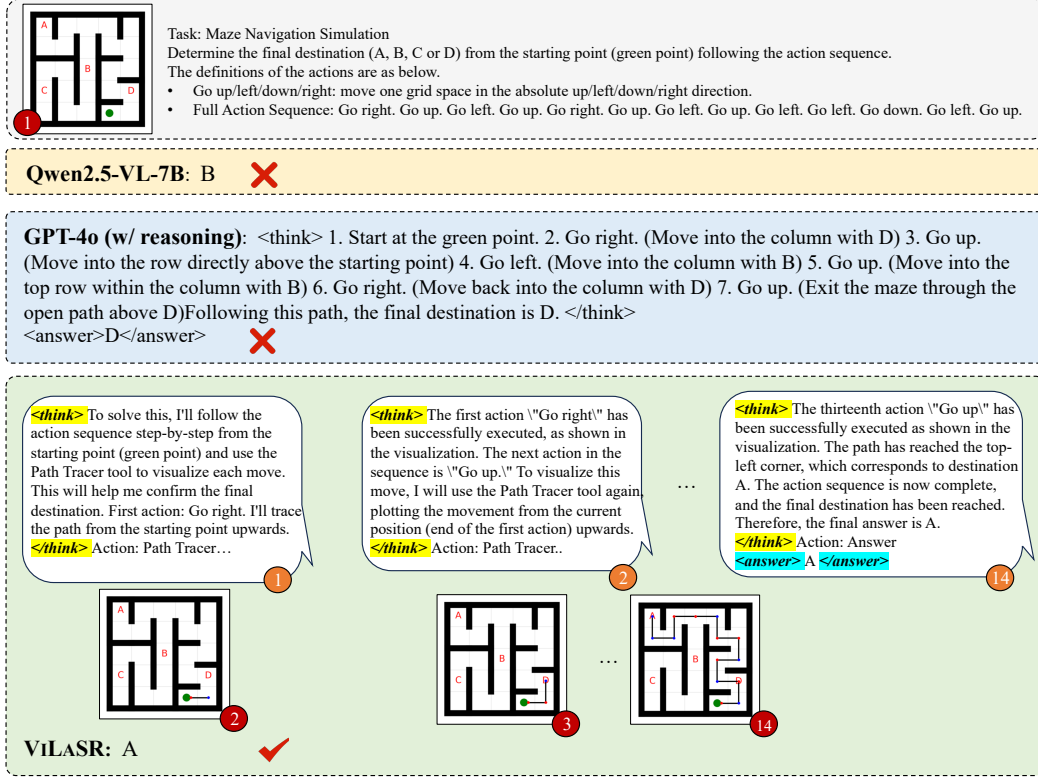


Figure 8: A visualization example of spatial reasoning approaches on the Maze benchmark, including Qwen-2.5-VL-7B, GPT-4o, and ViLASR.

In the video spatial reasoning task (Figure 10), ViLASR exhibits strong reflection capability and systematic problem-solving. When initial attempts to locate and measure the telephone in images 5 and 12 fail, it self-corrects and identifies the correct object in image 13. Furthermore, it shows a sophisticated measurement strategy by using reference objects (“monitor,” “headphones”) to establish scale and convert pixel measurements to centimeters. In contrast, baseline models either fail to locate the target object (Qwen2.5-VL-7B) or make rough estimations without proper measurement (GPT-4o).

These cases illustrate how drawing operations enable more reliable spatial reasoning by grounding abstract relationships in concrete visual representations.

E Complexity analysis

To thoroughly assess the computational efficiency of ViLASR, we present a comprehensive analysis comparing our approach with the base model Qwen2.5-VL-7B model. This analysis encompasses both theoretical complexity bounds and empirical resource utilization, providing insights into the framework’s scalability and computational characteristics.

E.1 Theoretical time complexity

We analyze the computational complexity across three model variants: the base model (performing direct answer generation), standard reasoning models (generating reasoning chains followed by answer derivation), and our ViLASR model (employing iterative visual drawing and thinking). Let M denote the input length (including both image and text tokens), N the answer length, L the per-step reasoning path length, and S the number of reasoning steps, where $S = 1$ for both base model and standard reasoning approaches. Empirically, ViLASR achieves shorter per-step reasoning

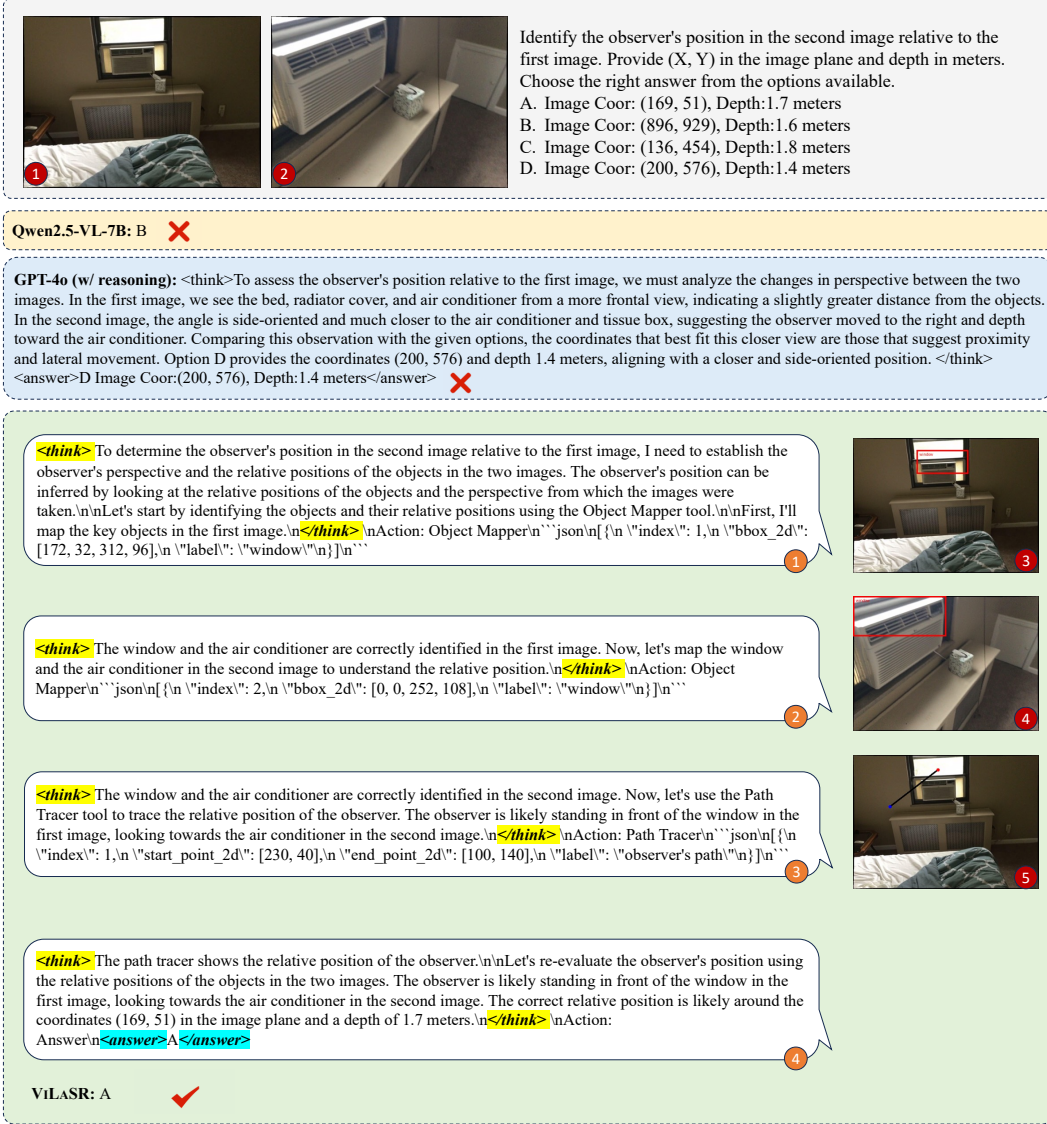


Figure 9: A visualization example of spatial reasoning approaches on the SPAR-Bench benchmark, including Qwen-2.5-VL-7B, GPT-4o, and ViLASR.

Table 6: Complexity analysis of different reasoning approaches.

| Model Type | Input Length | Output Length | Time Complexity |
|--------------------|------------------------|-----------------|---|
| Base Model | M | N | $\mathcal{O}(M \cdot N)$ |
| Standard Reasoning | M | $L + N$ | $\mathcal{O}(M \cdot (L + N))$ |
| ViLASR (Ours) | $M + \sum_{i=1}^S M_i$ | $S \cdot L + N$ | $\mathcal{O}(S^2 \cdot L \cdot (L + M) + S \cdot (L + M)(L + N))$ |

paths (smaller L) compared to single-step approaches due to its decomposition of reasoning into interpretable visual operations. Table 6 presents a comparative analysis of their time complexities.

The total computational complexity is:

$$T = \sum_{k=1}^{S-1} [(M + \sum_{i=1}^{k-1} (L_i + M_i)) \cdot L_k] + (M + \sum_{i=1}^{S-1} (L_i + M_i)) \cdot (L_S + N), \quad (9)$$

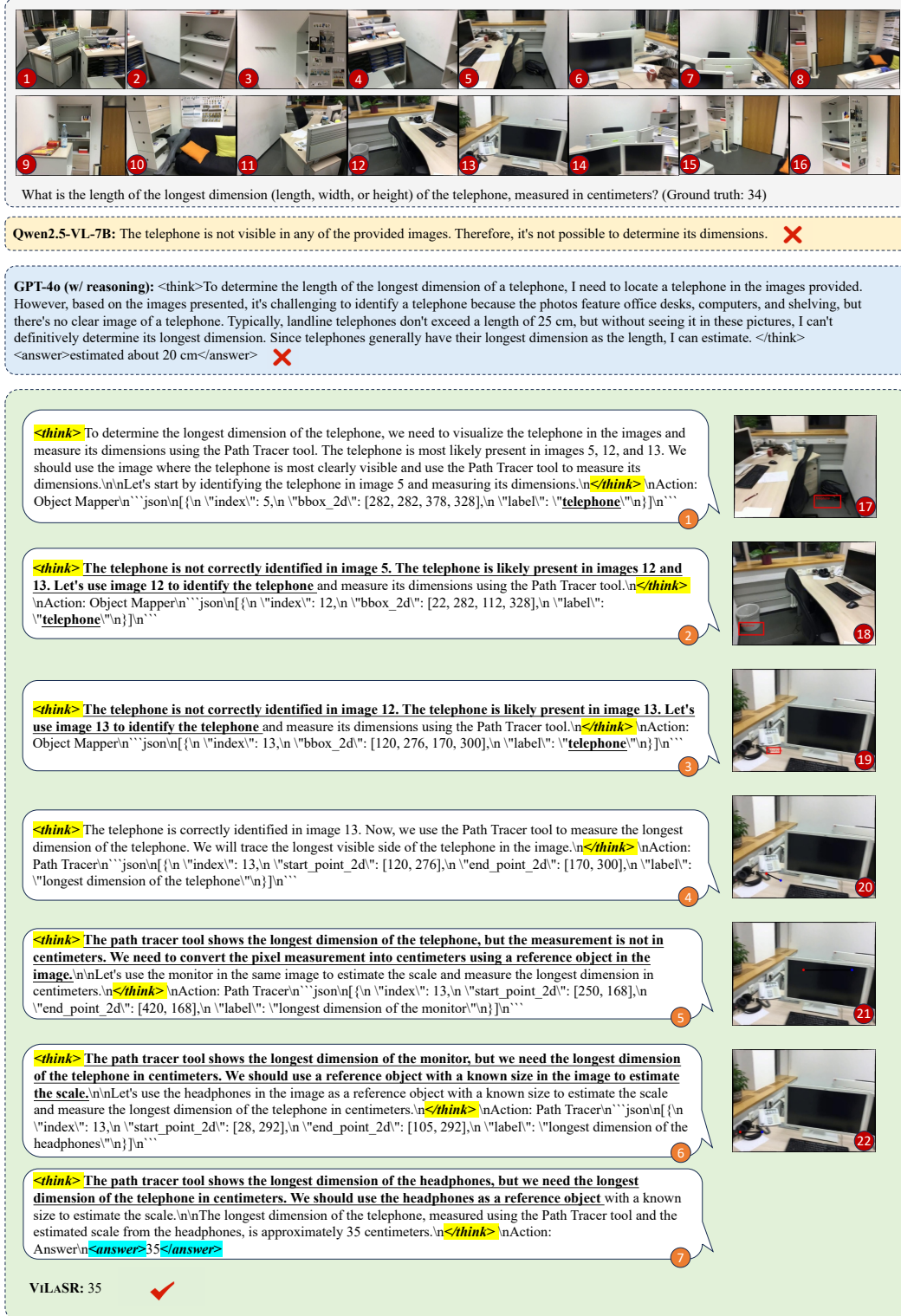


Figure 10: A visualization example of spatial reasoning approaches on the VSI-Bench benchmark, including Qwen-2.5-VL-7B, GPT-4o, and ViLASR. We highlight the generated words of ViLASR that exhibit reflection behavior in **bold**.

Table 7: Average runtime comparison in seconds per sample (ViLASR vs. backbone) across different benchmarks.

| Method | Average Runtime (s/sample) | | |
|---------------------|----------------------------|------------|-----------|
| | Maze | EmbSpatial | VSI-Bench |
| Qwen2.5-VL-7B (COT) | 1.0 | 0.9 | 1.9 |
| ViLASR | 4.4 | 1.2 | 7.8 |

where L_k is the length of the reasoning path generated at step k , and M_k is the additional multimodal context incorporated at step k .

Assuming each reasoning step has approximately equal length, i.e., $L_i \approx L$, $M_i \approx M$ for all i , we can approximate:

$$T = \mathcal{O}(S^2 \cdot L \cdot (L + M) + S \cdot (L + M)(L + N)). \quad (10)$$

E.2 Practical runtime

To provide empirical evidence for our complexity analysis, we evaluate the runtime performance of ViLASR and compare it against the backbone model, Qwen2.5-VL-7B (CoT), across various benchmarks. Table 7 illustrates the average runtime time per sample.

F Limitations

While our work demonstrates promising results in spatial reasoning tasks, several limitations warrant discussion:

First, our training framework primarily focuses on multiple-choice and numerical questions due to their amenability to automated evaluation, excluding more complex spatial reasoning scenarios that require free-form textual descriptions (e.g., detailed motion trajectory analysis). This limitation constrains the model’s ability to handle open-ended spatial reasoning tasks.

Second, a potential risk in our reinforcement learning stage is the coupling between the reward function and the final evaluation metrics. Since our reward is directly tied to achieving the correct final answer on specific benchmarks, the model might learn to exploit dataset-specific shortcuts to maximize its reward, rather than developing a truly generalizable spatial reasoning ability. This is a common challenge in reward-based optimization that warrants further investigation into more process-oriented reward schemes.

Third, the effectiveness of our drawing operations is inherently limited by the 2D nature of the visual interface. Complex 3D spatial relationships and viewpoint changes, which are common in real-world scenarios, may not be adequately captured by our current drawing primitives. This limitation particularly affects the model’s performance on tasks involving 3D object relationships or dynamic camera movements.

Finally, the computational cost of our three-stage training pipeline, especially during the reinforcement learning stage, may limit its accessibility to researchers with limited computational resources.

G Broader impact

Our work on enhancing spatial reasoning capabilities in vision-language models has both positive and potential negative societal implications that warrant careful consideration.

On the positive side, improved spatial reasoning capabilities could significantly benefit various applications in robotics, autonomous navigation, and assistive technologies. For instance, more accurate spatial understanding could help robots better navigate complex environments and assist people with visual impairments in daily tasks. In educational settings, these models could provide

interactive tools for teaching spatial concepts and geometric reasoning. Additionally, the interpretable nature of our drawing-based reasoning approach enhances model transparency, potentially increasing trust and adoption in critical applications.

However, several potential negative impacts need to be addressed. First, the enhanced spatial reasoning capabilities could be misused for surveillance purposes, enabling more sophisticated tracking and monitoring systems that could infringe on privacy rights. Second, there might be accessibility issues as the drawing-based reasoning approach assumes users have access to and can interact with visual interfaces, potentially excluding certain user groups.

To mitigate these concerns, we recommend: (1) implementing strict usage guidelines and access controls when deploying these models in sensitive applications, (2) exploring alternative interaction modalities to make the technology more accessible. We also encourage future research to focus on developing privacy-preserving spatial reasoning techniques.

To ensure responsible deployment of our technology, we have implemented several safeguards in our release. First, our model access will be gated through an API that requires user agreement to usage guidelines, specifically prohibiting applications in surveillance or privacy-invasive systems. Second, we provide detailed documentation about the model’s capabilities and limitations, along with best practices for responsible implementation in different application scenarios.