

BEYOND SEEING: EVALUATING MULTIMODAL LLMs ON TOOL-ENABLED IMAGE PERCEPTION, TRANSFORMATION, AND REASONING

Xingang Guo^{1,2}, Utkarsh Tyagi¹, Advait Gosai¹, Paula Vergara¹,
Ernesto Gabriel Hernández Montoya¹, Chen Bo Calvin Zhang¹, Bin Hu²,
Yunzhong He¹, Bing Liu¹, Rakshith Sharma Srinivasa¹

¹ScaleAI, ²University of Illinois at Urbana-Champaign

Leaderboard page: <https://scale.com/leaderboard/vtb>

ABSTRACT

Multimodal Large Language Models (MLLMs) are increasingly applied in real-world scenarios where user-provided images are often imperfect, requiring active image manipulations such as cropping, editing, or enhancement to uncover salient visual cues. Beyond static visual perception, MLLMs must also *think with images*: dynamically transforming visual content and integrating it with other tools to solve complex tasks. However, this shift from treating vision as passive context to a manipulable cognitive workspace remains underexplored. Most existing benchmarks still follow a *think about images* paradigm, where images are regarded as static inputs. To address this gap, we introduce VISUALTOOLBENCH, a visual tool-use reasoning benchmark that rigorously evaluates MLLMs’ ability to perceive, transform, and reason across complex visual-textual tasks under the *think-with-images* paradigm. VISUALTOOLBENCH comprises 1,204 challenging, open-ended vision tasks (603 single-turn, 601 multi-turn) spanning across five diverse domains, each paired with detailed rubrics to enable systematic evaluation. Our evaluation shows that current MLLMs struggle with tasks requiring effective integration of vision and general-purpose tools. Even the strongest model (GPT-5-think) reaches only 18.68% pass rate. We further observe divergent tool-use behaviors, with OpenAI models benefiting from diverse image manipulations while Gemini-2.5-pro shows no improvement. By introducing the first benchmark centered on *think with images*, VISUALTOOLBENCH offers critical insights for advancing visual intelligence in MLLMs.

1 INTRODUCTION

Multimodal Large Language Models (MLLMs), which integrate visual and textual understanding, have advanced rapidly in recent years and achieve impressive performance on a wide range of vision-language tasks, including image grounding (Rasheed et al., 2024; Zhang et al., 2024), image-based science problems (Zou et al., 2024; Lu et al., 2023; Yan et al., 2025b), visual question answering (Kuang et al., 2025; Liu et al., 2023a), optical character recognition (OCR) (Chen et al., 2025; Huang et al., 2025a), and spatial reasoning (Yang et al., 2025; Wu et al., 2025a; Tang et al., 2025). Current frontier MLLMs can interpret, describe, and reason about complex visual scenes in natural language, narrowing the gap between human and machine perception (Yin et al., 2024).

However, real-world use cases often need sophisticated processing of visual input and MLLMs may need to dynamically interact with them. For example, users often submit photos that are rotated, underexposed, cluttered, or poorly framed. In such cases, MLLMs need to work through multiple reasoning steps, including image manipulation, information extraction via tool-usage to arrive at the final answer. Standard models without tool support typically struggle under such degradations, revealing the need for active visual manipulation to solve such harder tasks. This distinction has recently been framed as *thinking about images* versus *thinking with images* (Su et al., 2025; OpenAI, 2025). Current multimodal benchmarks mainly adopt the former paradigm of *thinking about images* and focus on perception and reasoning over fixed, static images. The latter, by contrast, emphasizes

Table 1: Comparison of VISUALTOOLBENCH with representative multi-modal benchmarks.

Benchmark	Dynamic Vision Tool	Rubrics	Expert-Curated	Reasoning	Multi-turn
ScienceQA (Lu et al., 2022)	✗	✗	✓	✓	✗
MathVista (Lu et al., 2023)	✗	✗	✓	✓	✗
MMMU (Yue et al., 2024)	✗	✗	✓	✓	✗
V* (Wu & Xie, 2024)	✗	✗	✗	✗	✗
GTA (Wang et al., 2024a)	✗	✗	✗	✗	✗
ChartQA (Wang et al., 2024b)	✗	✗	✓	✓	✗
MMDU (Liu et al., 2024a)	✗	✗	✗	✓	✓
m & m's (Ma et al., 2024a)	✗	✗	✗	✗	✗
VISTA (Scale AI, 2025)	✗	✓	✓	✓	✗
VISUALTOOLBENCH (Ours)	✓	✓	✓	✓	✓

interactive, tool-augmented reasoning, where models autonomously manipulate visual inputs (e.g., cropping, editing, or enhancing) to extract fine-grained information for problem solving. Equipping MLLMs with such vision-specific tools during evaluation is therefore essential for robust and generalizable reasoning. These capabilities transform visual inputs from passive perception into a dynamic cognitive workspace, enabling MLLMs to tackle tasks that would otherwise be infeasible. Existing benchmarks remain inadequate for capturing this dimension (Su et al., 2025c).

To bridge this gap, we introduce VISUALTOOLBENCH, a challenging benchmark for vision and general-purpose tool-use that systematically evaluates how well MLLMs can perceive, transform, and reason about images under the *think with image* paradigm¹. Our key design principles are as follows:

- **Non-trivial visual perception.** Critical visual content is not easily accessible, models must apply appropriate image transformations (e.g., cropping, editing, or enhancement) to extract key visual details for better reasoning.
- **Realistic task settings.** Both prompts and images are designed to reflect practical, real-world scenarios rather than synthetic or overly simplified cases, ensuring that the benchmark closely mirrors real-world user needs.
- **Implicit tool-use requirements.** Tasks do not explicitly instruct the model which tool to use; instead, models must infer when and how to invoke tools based on contextual cues, making evaluation more faithful to realistic usage.
- **Multi-step, compositional reasoning.** Tasks are designed such that require combining visual transformations with multi-step reasoning (e.g., applying a sequence of tools, integrating extracted information, and synthesizing results), testing model’s ability to plan and execute complex workflows.

To reflect real-world applications, we design five complementary task categories that capture diverse aspects of model performance. Two cases are tailored to single-turn tasks and three targeting multi-turn tasks (Sec.2.1). Each task is authored by qualified contributors with proper training and undergoes multiple review stages to ensure high-quality data samples (Sec.2.2). To further capture complex and realistic scenarios, tasks are presented in an open-ended format. Further, each task is accompanied by a set of rubrics that span multiple dimensions. These criteria are used for systematic model evaluations (Sec.2.3). Finally, VISUALTOOLBENCH supports both a dynamic vision tool, which expose a flexible Python API for generating image manipulation code and re-ingesting processed images into the reasoning process, and general-purpose tools: web search, Python interpreter, calculator, and historical weather lookup, aiding in retrieval and computation for more advanced tasks (Sec.2.4).

By releasing VISUALTOOLBENCH and its accompanying evaluation toolkit, we aim to catalyze the development of MLLMs that seamlessly integrate image perception, tool use, and reasoning into a unified competency stack. Our contributions are four-fold:

¹See Table 1 for an overall comparison between VISUALTOOLBENCH and existing multimodal benchmarks.

Statistic	Number
Total questions	1,204
- STEM	238 (19.7%)
- Medical	238 (19.7%)
- Finance	243 (20.2%)
- Sports	240 (20.0%)
- Generalist	245 (20.4%)
Single-turn	603 (50.1%)
- Region Switch Q&A	281 (46.6%)
- Hybrid Tool-use	322 (53.4%)
Multi-turn	601 (49.9%)
- Follow-up Test	198 (32.9%)
- Temporal Reasoning	205 (34.2%)
- Progressive Reasoning	198 (32.9%)
Total number of rubrics	7,777
Total number of images	2,893
Average prompt length	48.41
Average answer length	128.93

Table 2: Statistics of VISUALTOOLBENCH.

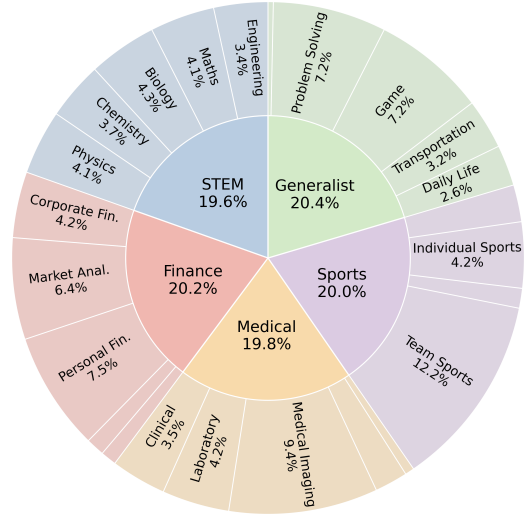


Figure 1: Topic distribution.

1. **The first *think with image*-oriented multimodal benchmark.** VISUALTOOLBENCH is the first benchmark to systematically evaluate MLLMs on tasks that require active visual manipulation to solve complex reasoning problems.
2. **Rubric-based, multi-dimensional evaluation.** Moving beyond binary correctness or exact string matching, we design detailed rubrics that capture partial credit across multiple categories. This richer scoring framework provides nuanced diagnostic insights into both the strengths and limitations of MLLMs.
3. **Large-scale and systematic evaluations.** We evaluate 16 representative MLLMs with function-calling capabilities, covering both reasoning and non-reasoning, as well as open- and closed-source models, under consistent settings. To support this evaluation, we developed a dedicated toolkit for vision tool use that allows models to access transformed images during reasoning and preserves complete tool-use trajectories. Our results reveal substantial performance gaps, with all models achieving pass rates below 20%.
4. **Comprehensive error and tool-use analysis.** We provide a detailed failure analysis and an in-depth examination of tool-use behaviors. Most failures stem from visual perception errors, highlighting the inefficiency of current models in using vision tools to extract key content. Furthermore, our study reveals divergent tool-use behaviors: the top-performing model, GPT-5, leverages diverse image manipulations to achieve clear gains over its no-tool baseline, whereas Gemini-2.5-pro does not gain improvement from tool access. These findings underscore the critical role of effective tool use in advancing MLLMs' performance.

2 VISUALTOOLBENCH

In this section, we present VISUALTOOLBENCH, a challenging visual reasoning benchmark that evaluates MLLMs' ability to perceive, transform, and reason on real-world tasks. The benchmark includes both single-turn and multi-turn interactions and incorporates five complementary task categories to probe different aspects of MLLM capabilities. Tasks are open-ended to reflect realistic scenarios, and each task is accompanied by detailed rubrics to support systematic evaluation. Table 2 summarizes the key statistics of VISUALTOOLBENCH, while Figure 1 illustrates the topic distribution of tasks.

2.1 TASK CATEGORY

We design five complementary task categories, each targeting a critical aspect of real-world use case. Together, they assess not only visual perception but also the efficiency of tool use and the depth of

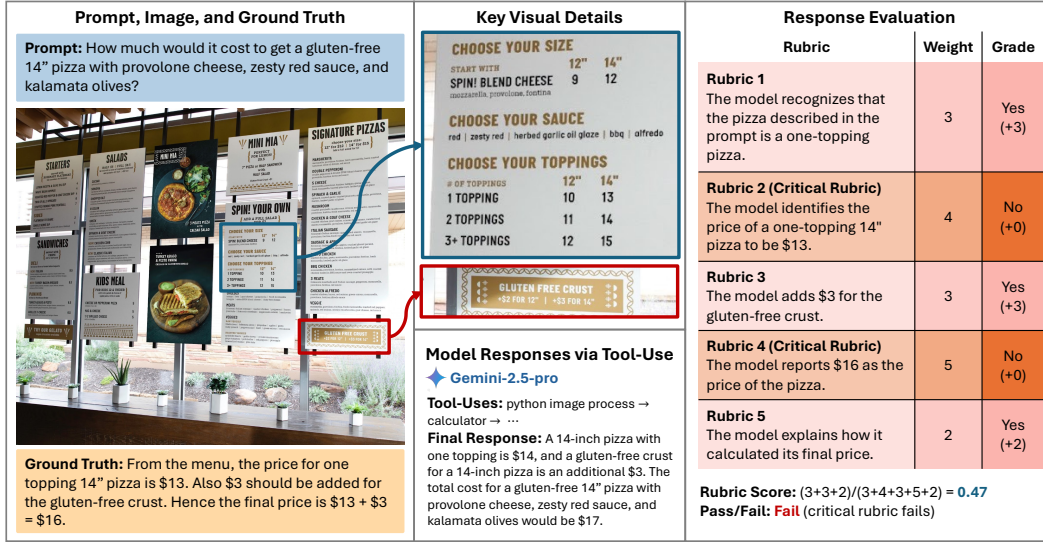


Figure 2: Demonstration example from VISUALTOOLBENCH (single-turn, generalist domain, region switch Q&A). The key visual content needed to solve the task is distributed across different regions of the image, requiring the model to crop multiple regions of interest (RoIs) for accurate perception and reasoning. Each task is paired with a detailed set of rubrics to evaluate model’s responses. From these rubrics, we derive both a weighted rubric score between 0 and 1 and a binary pass/fail outcome, depending on whether critical rubrics are satisfied.

multimodal reasoning. These categories are designed to mirror practical user scenarios, requiring models to think with images rather than relying solely on static perception.

Region-Switch Q&A (Single-Turn). The model answers a reasoning task that draws on information from multiple, spatially distinct regions of interest (RoIs) within a single image. Critical details may be small or dispersed, requiring the model to correctly identify, crop, and focus on relevant RoIs while disregarding irrelevant content. Success in this category reflects spatial selectivity, accurate region localization, and effective tool use for RoI extraction. Figure 2 illustrates a single-turn benchmark task for this evaluation focus in the generalist domain.

Hybrid Tool Reasoning (Single-Turn). The model must combine both vision-specific tools (e.g., python image processing) with other general-purpose tools (e.g., calculator, Python interpreter, or web search) to solve complex, multi-step tasks. This category evaluates the model’s ability to accurately call multiple tools and orchestrate heterogeneous tool outputs into a coherent reasoning chain towards solving complex visual reasoning tasks.

Follow-up Test (Multi-Turn). In this scenario, the first-round user query is intentionally underspecified or ambiguous. The model must engage in clarifying dialogue and ask follow-up questions before producing an answer. This tests conversational proactivity, uncertainty management, and the ability to self-correct, all essential skills for real-world deployments where users may provide incomplete or noisy instructions.

Temporal Visual Reasoning (Multi-Turn). Here the model reasons over a sequence of images across multiple turns, requiring it to detect temporal changes, track motion, or infer causal relationships among multiple images. Tasks may involve following the progression of an event, monitoring evolving states, or interpreting multi-step visual instructions.

Progressive Visual Reasoning (Multi-Turn). The model solves a series of interdependent questions about the same image, where later queries could build upon earlier answers. This requires the model to maintain internal consistency, remember prior outputs, and construct a layered understanding of the scene. Success in this category demonstrates long-horizon reasoning, contextual memory, and the ability to sustain a coherent reasoning trajectory.

Additional examples of VISUALTOOLBENCH are provided in Appendix E.

2.2 DATA COLLECTION

All VISUALTOOLBENCH tasks are authored by human contributors with diverse domain expertise. To ensure benchmark quality and realism, we adopt a rigorous multi-stage data collection pipeline:

1. **Contributor Training.** Contributors are first introduced to the project scope, task categories, and submission requirements. They are instructed to provide a task prompt, input image, reference answer, reference tool-use chain, and a set of evaluation rubrics.
2. **Initial Task Design.** Drawing on their domain expertise, contributors design tasks by selecting the appropriate domain and aligning with the specified task category. Each submission includes a text prompt and associated image(s), a golden answer, a reference tool-use chain that demonstrates a valid solution path, and well-defined evaluation rubrics.
3. **Initial Model Response Grading.** Contributors are presented with responses from three representative models (o3, Gemini-2.5-pro, and o4-mini) to their designed tasks. They then grade these responses against the rubrics². A task is selected only if at least two of the three models fail, thereby ensuring that the benchmark captures genuinely challenging cases.
4. **First-Round Review.** A reviewer evaluates each task for realism, necessity of dynamic image-based reasoning, correctness of the reference answer, and appropriateness of the rubrics. Tasks with minor issues may be revised, while those that are fundamentally unsound (e.g., not requiring genuine visual-tool use) are discarded.
5. **Second-Round Review.** A second independent reviewer validates the first-round decision, ensuring consistency and reliability across the benchmark.
6. **Final Integration.** Tasks that pass both review stages are incorporated into the benchmark, ensuring high quality, broad domain coverage, and diverse reasoning requirements.

This layered pipeline ensures that every task is original, realistic, and rigorously validated, resulting in a benchmark that robustly evaluates genuine visual intelligence.

2.3 RUBRIC-BASED EVALUATIONS

We adopt rubric-based evaluation to capture nuanced aspects of model performance beyond correctness alone (Arora et al., 2025; Starace et al., 2025; Scale AI, 2025; Lin et al., 2024; Sirdeshmukh et al., 2025; Fast et al., 2024; Gunjal et al., 2025; Guo et al., 2025). For each task, contributors provide a comprehensive set of rubric criteria to assess model responses. A rubric item may range from specific factual requirements (e.g., providing the correct final short answer) to broader aspects of desirable behavior (e.g., presenting key intermediate steps). Specifically, rubric items are organized into five main categories for VISUALTOOLBENCH:

1. **Visual Understanding:** Correct identification, extraction, and explanation of relevant visual elements such as text, objects, or spatial relationships.
2. **Truthfulness:** Accuracy of all factual statements and correctness of the final answer.
3. **Instruction Following:** Precise adherence to the task requirements specified in the prompt.
4. **Reasoning:** Use of clear, step-by-step logic with justified inferences and calculations.
5. **Presentation:** Clarity, coherence, structure, and appropriate formatting of the response.

Each rubric criterion is assigned a weight $w \in \{1, 2, 3, 4, 5\}$ by the task contributor, where higher weights indicate greater importance. To evaluate a model’s response, an auto-grader examines each rubric criterion independently and determines whether the response satisfies it. If the criterion is met, full points are awarded; otherwise, no points are given. The weighted rubric score for a task is then calculated as the sum of satisfied items, normalized by the total rubric weights. Rubric items with $w \geq 4$ are designated as **critical rubrics**. These rubrics typically correspond to essential aspects such as truthfulness and key visual understanding, and satisfying them indicates that the model has solved the task in a substantive way. Failure to meet any critical rubric results in an overall failure for

²These annotations serve as golden human labels for rubric-based evaluation and enable subsequent analysis of LLM-as-judge versus human-judge alignment (Sec. B.7), facilitating scalable evaluation.

that task, and this binary outcome is used to compute a task-level accuracy in VISUALTOOLBENCH. All rubrics together contribute to the weighted rubrics score, enabling more fine-grained analysis. The right panel of Figure 2 illustrates rubric-based grading.

2.4 TOOL SET

To enable broad, tool-augmented reasoning under the *think with images* paradigm, VISUALTOOLBENCH exposes a standardized API with six carefully selected tools: `python_image_processing`, `python_interpreter`, `web_search`, `calculator`, `browser_get_page_text`, and `historical_weather`. This compact yet diverse toolset spans core capabilities for image manipulation, computation, retrieval, and domain-specific lookups. Among them, the vision tool `python_image_processing` is particularly central: it supports arbitrary manipulations such as cropping, editing, and brightness/contrast adjustments, enabling models to iteratively refine visual inputs and use images as an interactive scratchpad. This versatility makes it the cornerstone of our benchmark’s *think-with-images* setup. Detailed tool descriptions are provided in Appendix D.

3 EXPERIMENT RESULTS

Evaluation Setup. We conduct large-scale evaluations using LiteLLM’s function-calling interface (LiteLLM). Models are given the supported tools and invoke them by emitting the corresponding call arguments. For conventional (non-vision) tools, outputs are textual and are appended to the dialogue as a `tool` message. Vision tools, in contrast, return transformed images. We observe that placing encoded images directly in a `tool` message does not make them perceptible to the model. To ensure effective re-ingestion of visual results, after a vision tool executes and saves its outputs, we insert an additional `user` message containing the encoded image(s). This preserves models’ ability to *think with images*, wherein newly produced images inform later reasoning. We set a cap of 20 tool calls per task, while human reference trajectories usually need less than 5 tool calls.

Baseline Models. We benchmark 16 representative MLLMs with function-calling capabilities, covering both reasoning and non-reasoning as well as open- and closed-source models. A complete list of models, along with their endpoints and configuration details, is provided in Appendix C.1.

Evaluation Metrics. We report two main metrics derived from rubric-based judgments: *Average Pass Rate (APR)* and *Average Rubric Score (ARS)*. (i) APR. Each task specifies a set of *critical* rubrics. A model’s response *passes* only if all it passes all the critical rubrics; otherwise it *fails*. APR is then the percentage of tasks that pass across the dataset. (ii) ARS. Each rubric is assigned an integer weight from 1 to 5 by the contributor to indicate its importance. For a model’s response on a task, we compute a weighted proportion of satisfied rubrics: the total weight of satisfied rubrics divided by the total weight of all rubrics for that task, as a weighted rubric score. The dataset-level ARS is then the average of these per-task rubric scores. Formal definitions and implementation details are provided in Appendix B.

3.1 MAIN RESULTS

We present the overall APR results in Table 3. Figure 3 represents the APR across task categories for top five performing models on VISUALTOOLBENCH. We make the following observations.

VISUALTOOLBENCH is highly challenging. From Table 3, it is clear that VISUALTOOLBENCH poses a challenging vision tool-use reasoning benchmark. Specifically, even the best-performing model, GPT-5-think, achieves only an 18.68% overall pass rate, and 11 out of 16 MLLMs obtain APRs below 10%. This highlights the limitations of current MLLMs and underscores the substantial room for improvement on visual-reasoning tasks where critical content is not directly accessible and must be extracted through vision tools.

OpenAI models outperform others. Models from OpenAI, including GPT-5, GPT-5-think, and o3, lead performance with APRs above 13%, showing a clear margin over competing models. This may be attributed to their specific training for solving *think with images* tasks (OpenAI, 2025). On the other hand, Gemini-2.5-pro also demonstrates relatively strong performance (11.75%), due to its advanced visual perception capabilities (Comanici et al., 2025).

Table 3: APR (%) results of the evaluated models across domains (averaged across two independent trails). Domain abbreviations: **Med** (Medicine), **Fin** (Finance), **Spt** (Sports), and **Gen** (Generalist). The best results in each column are highlighted with a red background, and the second-best results are highlighted in blue.

Model	Overall	Single-Turn					Multi-Turn				
		STEM	Med	Fin	Sprt	Gen	STEM	Med	Fin	Sprt	Gen
Open-Source Models											
Llama4-Maverick	1.41	3.88	2.07	1.22	2.52	0.81	0.41	1.71	0.00	0.83	0.83
Llama4-Scout	1.58	2.16	3.31	1.63	2.52	2.82	0.00	2.35	0.43	0.00	0.95
Closed-Source Models											
GPT-4.1	5.52	5.19	11.57	2.03	8.40	3.63	6.56	6.41	1.67	5.79	4.13
o3	13.74	23.83	22.73	12.20	18.07	17.34	11.07	7.69	11.25	7.02	6.61
o4-mini	11.12	15.16	19.01	11.38	14.71	18.95	5.33	8.55	5.83	5.79	6.20
GPT-5	16.96	29.31	24.79	24.80	16.81	26.21	14.29	4.79	11.44	7.94	9.15
GPT-5-think	18.68	28.13	26.03	24.39	22.69	29.03	15.57	7.69	12.92	10.74	9.50
Gemini-2.5-pro	11.75	19.83	16.53	17.07	14.71	18.95	7.37	7.42	2.74	7.77	5.12
Gemini-2.5-flash	4.69	6.90	6.20	3.66	9.66	8.06	3.75	3.96	0.00	3.19	1.39
Claude-sonnet-4	4.48	4.33	7.85	4.07	7.14	6.85	3.28	5.98	1.67	2.07	1.65
Claude-opus-4.1	4.71	6.03	9.92	4.07	7.14	7.26	3.95	5.13	0.00	2.48	1.28
Claude-sonnet-4.5	5.60	6.51	11.28	5.10	8.84	7.98	3.69	4.73	1.67	4.20	1.67
Claude-sonnet-4-think	4.44	2.16	8.26	4.07	7.56	8.87	3.28	4.27	1.67	1.65	2.48
Claude-opus-4.1-think	5.16	5.17	9.50	7.32	11.34	7.26	1.64	3.42	0.85	4.13	0.88
Claude-sonnet-4.5-think	6.20	9.48	13.91	6.91	6.89	8.72	3.34	3.04	1.69	4.71	3.03
Nova-Premier	2.00	3.02	5.79	0.81	2.52	2.42	2.09	1.77	0.00	1.66	0.00

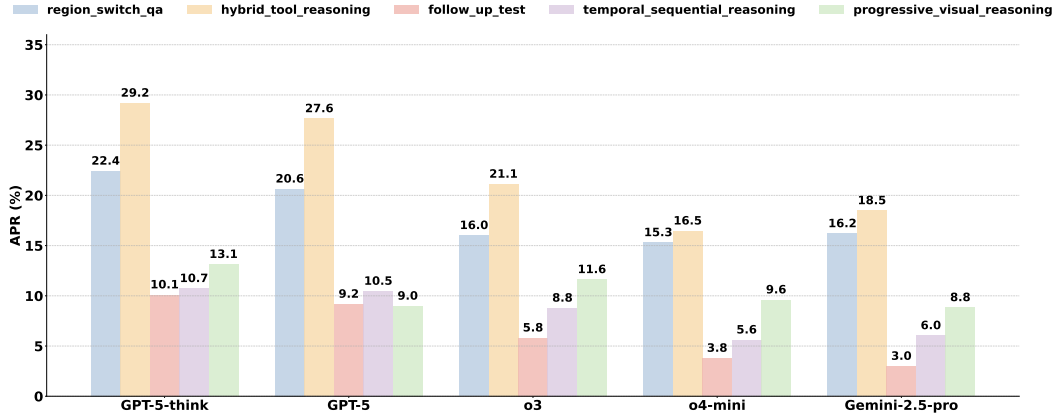


Figure 3: APR across task categories for the top five models on VISUALTOOLBENCH. Bars report model performance on each category, with exact APR values labeled above.

Multi-turn tasks are more difficult than single-turn tasks. From Table 3 and Figure 3, we see that single-turn tasks (region-switch Q&A, hybrid-tool reasoning) achieve higher pass rates compared to multi-turn tasks (follow-up test, temporal sequential reasoning, and progressive visual reasoning). This is expected, as multi-turn tasks involve 2 to 5 conversational turns, introducing more opportunities for errors and compounding reasoning challenges.

3.2 TOOL-USE ANALYSIS

In this section, we perform tool-call analysis on the evaluated models. First, we quantify tool-use behavior with three descriptive metrics computed from execution traces: **proactivity**, **tool-call success rate**, and **tool-call volume**. Proactivity is the fraction of tasks in which at least one tool is invoked, capturing a model’s tendency to integrate tools into its reasoning. Success rate measures the proportion of invocations that return a schema-valid, non-empty result, reflecting model’s adherence to tool specifications. Tool-call volume is the average number of tool calls per task, indicating how heavily a model relies on tools to solve tasks. Precise definitions are provided in Appendix B.4.

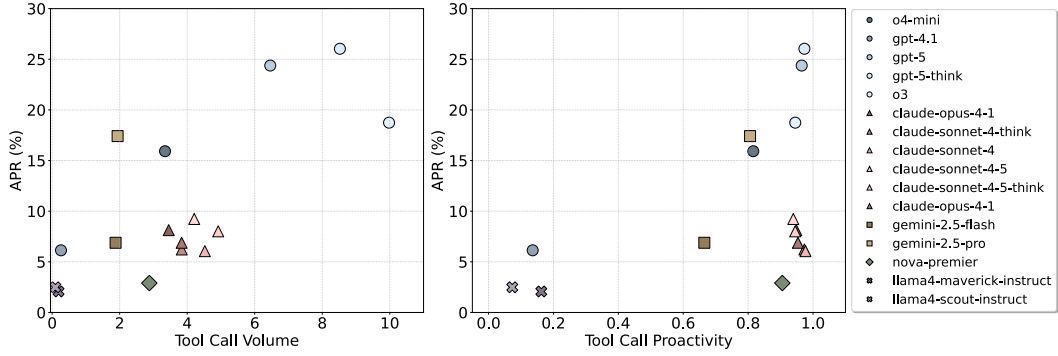


Figure 4: Relationship between APR and tool-use behaviors across models. Left: APR versus tool-use proactivity. Right: APR versus average tool call volume per task.

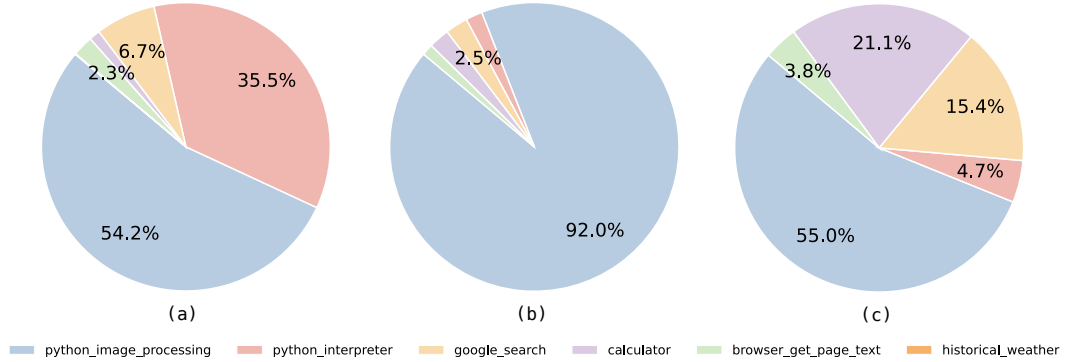


Figure 5: Tool-use distribution of top three APR models. Each subplot shows the percentage share of tool calls. Total number of tool calls: (a) o3: 6004; (b) GPT-5: 3805; (c) Gemini-2.5-pro: 1082.

Models use tools ineffectively. Figure 4 plots model APR against tool-use proactivity (left) and average tool-call volume per task (right). We observe a clear positive correlation: models that use tools more proactively and with higher call volume tend to achieve better performance. OpenAI models cluster in the upper right, showing both higher APR and stronger tool-use behaviors, while other families (Claude, Gemini, Llama, Nova, PixaI) remain concentrated in the lower ranges with limited tool use and weaker performance. These findings indicate that most current MLLMs underexploit tool capabilities, and that proactive, consistent tool use is a key driver of success. More detailed tool-call analyses are provided in Appendix B.5.

Vision tool is the most called tool. Figure 5 presents the tool-use distributions of the top three APR models (o3, GPT-5-think, and Gemini-2.5-Pro). Across all three models, more than 50% of tool calls are to the `python_image_processing` (92% for GPT-5), underscoring that image manipulation is the primary operation required to solve tasks in our benchmark. Other tools such as `python_interpreter`, `web_search`, and `calculator` are invoked less frequent, reinforcing the inherently *think with images* nature of VISUALTOOLBENCH.

Models exhibit diverse patterns of image manipulation. We do not constrain models to a predefined set of image operations; instead, they can invoke arbitrary manipulations via Python through the vision tool. To better understand tool usage in practice, we group the issued calls into eight representative categories: cropping, resizing, rotation, flipping, brightness adjustment, contrast adjustment, editing (e.g., annotation, inpainting, drawing), and others. Figure 6 reports the distribution of these operations across the top-performing models. GPT-5-think and GPT-5 stand out with both higher volumes and broader diversity of manipulations, reflecting more active exploration of tool

Model	Error Type	Percentage
GPT-5	visual perception	71.73%
	reasoning	11.56%
	calculation	2.79%
	others	13.92%
Gemini-2.5-pro	visual perception	78.01%
	reasoning	12.24%
	calculation	5.74%
	others	4.01%
Claude-opus-4.1	visual perception	82.11%
	reasoning	9.36%
	calculation	1.84%
	others	6.69%

Table 4: Error type statistics.

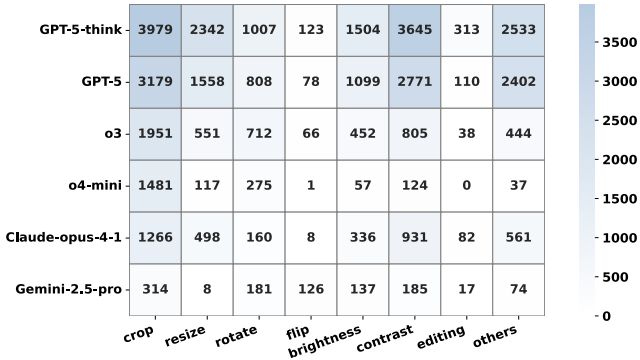


Figure 6: Image manipulation operations counts.

capabilities³. o3 and o4-mini also perform frequent manipulations but with narrower operation profiles, while Claude-opus-4.1 and Gemini-2.5-pro show comparatively limited usage. Illustrative cases and qualitative edge examples are provided in Appendix D.2.

3.3 ERROR ANALYSIS

In this section, we analyze errors for three representative models (GPT-5, Gemini-2.5-pro, and Claude-opus-4.1) using the following four general categories: visual-perception error, reasoning, calculation, and others. Table 4 summarizes the distribution of error types. Across all models, visual-perception errors are the most common failure mode. By contrast, calculation mistakes are rare, and reasoning errors occur only occasionally. More discussions are provided in Appendix B.6.

3.4 ABLATION STUDY

Figure 7 compares APR across four different settings: tool-use with strong and weak system prompts, no vision tool, and no tools. For GPT-5, removing tools or weakening the prompt causes significant performance drops ($\approx 11 - 14\%$), showing clear gains from tool-augmented reasoning. Surprisingly, Gemini-2.5-pro performs better **without tools** (+2.7%), implying the tool-use have counter-effect on its performance. In contrast, GPT-5 benefits from tools, boosting its performance further when enabled. These divergent trends likely reflect training differences: GPT-5 appears reinforced on tool-centric workflows, relying on iterative edits to offset weaker perception, whereas Gemini-2.5-pro, with stronger native vision, was likely exposed to fewer tool-use demonstrations and thus degrades when applying unnecessary tool calls. Despite these differences, all evaluated models achieve only modest APR on VISUALTOOLBENCH, indicating substantial headroom for improving when and how models use tools.

4 RELATED WORK

Most existing multimodal benchmarks remain limited to *think about images*: they focus on passive visual Q&A without active interactions (Rahmanzadehgervi et al., 2024; Yue et al., 2024; Wang et al., 2024b; Lu et al., 2023; Zou et al., 2024; Wu & Xie, 2024; Scale AI, 2025), restrict tool use to basic operations like cropping (Wang et al., 2024a; Ma et al., 2024a), or evaluate tool-agnostic multi-turn dialogues without essential integration (Liu et al., 2024a; Yan et al., 2025a); see also (Li et al., 2024) for a broader overview.

Recent efforts toward teaching MLLMs to *think with images* include prompting-based methods that leverage language mediation, visual input manipulation, or expert integration (Zeng et al., 2022;

³Although GPT-5 and GPT-5-think issue fewer tool calls than o3 (see the first column of Table 7), they perform more image manipulations overall as shown in Figure 6. A deeper inspection shows that GPT-5 and GPT-5-think often execute multiple operations within a single vision-tool call, indicating greater tool-use efficiency compared to o3 and contributing to their superior performance. See a demonstration in Appendix D.3.

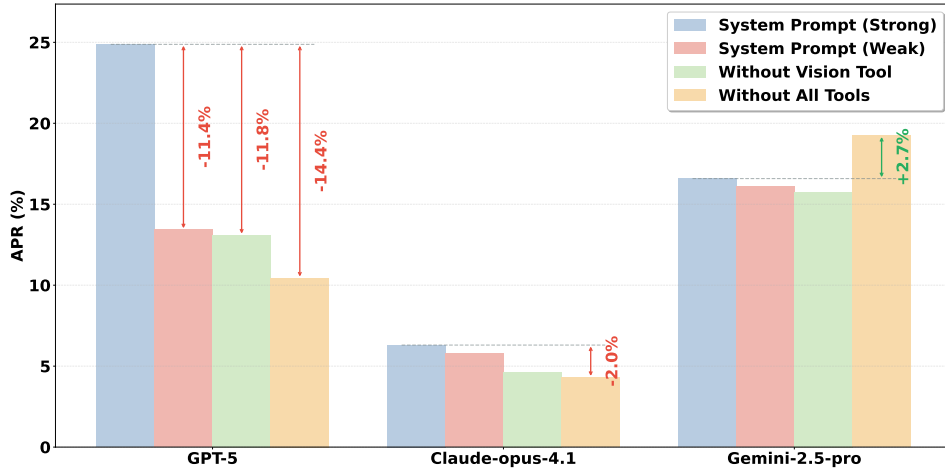


Figure 7: Impact of system prompts and tool availability on model performance (APR). This figure compares the average pass rate (APR) of three representative MLLMs: GPT-5, Claude-opus-4.1, and Gemini-2.5-pro, under four evaluation settings: strong system prompt, weak system prompt, without vision tool, and without all tools.

Yang et al., 2023b; Wu et al., 2024a; Liu et al., 2025a), supervised fine-tuning to enable tool invocation and intrinsic manipulations such as cropping, grounding, and dynamic attention (Liu et al., 2023b; Shao et al., 2024), and reinforcement learning for adaptive exploration and tool orchestration (Wang et al., 2025c; Fan et al., 2025; Huang et al., 2025b; Zheng et al., 2025). Despite these advances, evaluations remain constrained to visually passive benchmarks. VISUALTOOLBENCH addresses this gap by providing a rigorous testbed for genuine visual intelligence under the *think with images* paradigm. See broader discussion in (Su et al., 2025c) and more related work in Appendix A.

5 CONCLUSION

In this work, we introduced VISUALTOOLBENCH, a benchmark designed to evaluate MLLMs under the *think with image* paradigm. Unlike prior efforts that treat images as static inputs, our benchmark emphasizes active visual manipulation through tool use across diverse domains and task types. Our experiments show that current MLLMs continue to struggle on tasks requiring dynamic visual manipulations. We hope VISUALTOOLBENCH will drive progress toward models that can more effectively think with images and tackle challenging real-world scenarios.

REFERENCES

- Rahul K Arora, Jason Wei, Rebecca Soskin Hicks, Preston Bowman, Joaquin Quiñonero-Candela, Foivos Tsimpourlas, Michael Sharman, Meghan Shah, Andrea Vallone, Alex Beutel, et al. Health-bench: Evaluating large language models towards improved human health. *arXiv preprint arXiv:2505.08775*, 2025.
- Sule Bai, Mingxing Li, Yong Liu, Jing Tang, Haoji Zhang, Lei Sun, Xiangxiang Chu, and Yansong Tang. Univg-r1: Reasoning guided universal visual grounding with reinforcement learning. *arXiv preprint arXiv:2505.14231*, 2025a.
- Tianyi Bai, Zengjie Hu, Fupeng Sun, Jiantao Qiu, Yizhen Jiang, Guangxin He, Bohan Zeng, Conghui He, Binhang Yuan, and Wentao Zhang. Multi-step visual reasoning with visual tokens scaling and verification. *arXiv preprint arXiv:2506.07235*, 2025b.
- Jing Bi, Junjia Guo, Susan Liang, Guangyu Sun, Luchuan Song, Yunlong Tang, Jinxi He, Jiarui Wu, Ali Vosoughi, Chen Chen, et al. Verify: A benchmark of visual explanation and reasoning for investigating multimodal reasoning fidelity. *arXiv preprint arXiv:2503.11557*, 2025.
- Song Chen, Xinyu Guo, Yadong Li, Tao Zhang, Mingan Lin, Dongdong Kuang, Youwei Zhang, Lingfeng Ming, Fengyu Zhang, Yuran Wang, et al. Ocean-ocr: Towards general ocr application via a vision-language model. *arXiv preprint arXiv:2501.15558*, 2025.
- Zihui Cheng, Qiguang Chen, Xiao Xu, Jiaqi Wang, Weiyun Wang, Hao Fei, Yidong Wang, Alex Jinpeng Wang, Zhi Chen, Wanxiang Che, et al. Visual thoughts: A unified perspective of understanding multimodal chain-of-thought. *arXiv preprint arXiv:2505.15510*, 2025.
- Jiwan Chung, Junhyeok Kim, Siyeol Kim, Jaeyoung Lee, Min Soo Kim, and Youngjae Yu. Don’t look only once: Towards multimodal interactive reasoning with selective visual revisitation. *arXiv preprint arXiv:2505.18842*, 2025.
- Gheorghe Comanici, Eric Bieber, Mike Schaekermann, Ice Pasupat, Noveen Sachdeva, Inderjit Dhillon, Marcel Blistein, Ori Ram, Dan Zhang, Evan Rosen, et al. Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities. *arXiv preprint arXiv:2507.06261*, 2025.
- Yue Fan, Xuehai He, Diji Yang, Kaizhi Zheng, Ching-Chen Kuo, Yuting Zheng, Sravana Jyothi Narayanaraju, Xinze Guan, and Xin Eric Wang. Grit: Teaching mllms to think with images. *arXiv preprint arXiv:2505.15879*, 2025.
- Dennis Fast, Lisa C Adams, Felix Busch, Conor Fallon, Marc Huppertz, Robert Siepmann, Philipp Prucker, Nadine Bayerl, Daniel Truhn, Marcus Makowski, et al. Autonomous medical evaluation for guideline adherence of large language models. *NPJ Digital Medicine*, 7(1):358, 2024.
- Anisha Gunjal, Anthony Wang, Elaine Lau, Vaskar Nath, Bing Liu, and Sean Hendryx. Rubrics as rewards: Reinforcement learning beyond verifiable domains. *arXiv preprint arXiv:2507.17746*, 2025.
- Xingang Guo, Yaxin Li, Xiangyi Kong, Yilan Jiang, Xiayu Zhao, Zhihua Gong, Yufan Zhang, Daixuan Li, Tianle Sang, Beixiao Zhu, Gregory Jun, Yingbing Huang, Yiqi Liu, Yuqi Xue, Rahul Dev Kundu, Qi Jian Lim, Yizhou Zhao, Luke Alexander Granger, Mohamed Badr Younis, Darioush Keivan, Nippun Sabharwal, Shreyanka Sinha, Prakhar Agarwal, Kojo Vandyck, Hanlin Mai, Zichen Wang, Aditya Venkatesh, Ayush Barik, Jiankun Yang, Chongying Yue, Jingjie He, Libin Wang, Licheng Xu, Hao Chen, Jinwen Wang, Liujun Xu, Rushabh Shetty, Ziheng Guo, Dahui Song, Manvi Jha, Weijie Liang, Weiman Yan, Bryan Zhang, Sahil Bhandary Karnoor, Jialiang Zhang, Rutva Pandya, Xinyi Gong, Mithesh Ballae Ganesh, Feize Shi, Ruiling Xu, Yifan Zhang, Yanfeng Ouyang, Lianhui Qin, Elyse Rosenbaum, Corey Snyder, Peter Seiler, Geir Dullerud, Xiaojia Shelly Zhang, Zuofu Cheng, Pavan Kumar Hanumolu, Jian Huang, Mayank Kulkarni, Mahdi Namazifar, Huan Zhang, and Bin Hu. Toward engineering AGI: Benchmarking the engineering design capabilities of llms. In *arXiv preprint arXiv:2509.16204*, 2025.
- Yushi Hu, Hang Hua, Zhengyuan Yang, Weijia Shi, Noah A Smith, and Jiebo Luo. Promptcap: Prompt-guided task-aware image captioning. *arXiv preprint arXiv:2211.09699*, 2022.

- Mingxin Huang, Yongxin Shi, Dezhi Peng, Songxuan Lai, Zecheng Xie, and Lianwen Jin. Ocr-reasoning benchmark: Unveiling the true capabilities of mllms in complex text-rich image reasoning. *arXiv preprint arXiv:2505.17163*, 2025a.
- Zeyi Huang, Yuyang Ji, Anirudh Sundara Rajan, Zefan Cai, Wen Xiao, Haohan Wang, Junjie Hu, and Yong Jae Lee. Visualtoolagent (vista): A reinforcement learning framework for visual tool selection. *arXiv preprint arXiv:2505.20289*, 2025b.
- Jiayi Kuang, Ying Shen, Jingyou Xie, Haohao Luo, Zhe Xu, Ronghao Li, Yinghui Li, Xianfeng Cheng, Xika Lin, and Yu Han. Natural language understanding and inference with mllm in visual question answering: A survey. *ACM Computing Surveys*, 57(8):1–36, 2025.
- Geng Li, Jinglin Xu, Yunzhen Zhao, and Yuxin Peng. Dyfo: A training-free dynamic focus visual search for enhancing lmms in fine-grained visual understanding. In *Proceedings of the Computer Vision and Pattern Recognition Conference*, pp. 9098–9108, 2025.
- Lin Li, Guikun Chen, Hanrong Shi, Jun Xiao, and Long Chen. A survey on multimodal benchmarks: In the era of large ai models. *arXiv preprint arXiv:2409.18142*, 2024.
- Bill Yuchen Lin, Yuntian Deng, Khyathi Chandu, Faeze Brahman, Abhilasha Ravichander, Valentina Pyatkin, Nouha Dziri, Ronan Le Bras, and Yejin Choi. Wildbench: Benchmarking llms with challenging tasks from real users in the wild. *arXiv preprint arXiv:2406.04770*, 2024.
- LiteLLM. Litellm documentation: Function call. https://docs.litellm.ai/docs/completion/function_call. Accessed: 2025-09-17.
- Dairu Liu, Ziyue Wang, Minyuan Ruan, Fuwen Luo, Chi Chen, Peng Li, and Yang Liu. Visual abstract thinking empowers multimodal reasoning. *arXiv preprint arXiv:2505.20164*, 2025a.
- Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. Visual instruction tuning. *Advances in neural information processing systems*, 36:34892–34916, 2023a.
- Shilong Liu, Hao Cheng, Haotian Liu, Hao Zhang, Feng Li, Tianhe Ren, Xueyan Zou, Jianwei Yang, Hang Su, Jun Zhu, Lei Zhang, Jianfeng Gao, and Chunyuan Li. Llava-plus: Learning to use tools for creating multimodal agents. *arXiv:2311.05437*, 2023b.
- Yuqi Liu, Bohao Peng, Zhisheng Zhong, Zihao Yue, Fanbin Lu, Bei Yu, and Jiaya Jia. Seg-zero: Reasoning-chain guided segmentation via cognitive reinforcement. *arXiv preprint arXiv:2503.06520*, 2025b.
- Yuqi Liu, Tianyuan Qu, Zhisheng Zhong, Bohao Peng, Shu Liu, Bei Yu, and Jiaya Jia. Vision-reasoner: Unified visual perception and reasoning via reinforcement learning. *arXiv preprint arXiv:2505.12081*, 2025c.
- Ziyu Liu, Tao Chu, Yuhang Zang, Xilin Wei, Xiaoyi Dong, Pan Zhang, Zijian Liang, Yuanjun Xiong, Yu Qiao, Dahua Lin, et al. Mmdu: A multi-turn multi-image dialog understanding benchmark and instruction-tuning dataset for lvlms. *arXiv preprint arXiv:2406.11833*, 2024a.
- Zuyan Liu, Yuhao Dong, Yongming Rao, Jie Zhou, and Jiwen Lu. Chain-of-spot: Interactive reasoning improves large vision-language models. *arXiv preprint arXiv:2403.12966*, 2024b.
- Pan Lu, Swaroop Mishra, Tanglin Xia, Liang Qiu, Kai-Wei Chang, Song-Chun Zhu, Oyvind Tafjord, Peter Clark, and Ashwin Kalyan. Learn to explain: Multimodal reasoning via thought chains for science question answering. *Advances in Neural Information Processing Systems*, 35:2507–2521, 2022.
- Pan Lu, Hritik Bansal, Tony Xia, Jiacheng Liu, Chunyuan Li, Hannaneh Hajishirzi, Hao Cheng, Kai-Wei Chang, Michel Galley, and Jianfeng Gao. Mathvista: Evaluating mathematical reasoning of foundation models in visual contexts. *arXiv preprint arXiv:2310.02255*, 2023.
- Yan Ma, Linge Du, Xuyang Shen, Shaoxiang Chen, Pengfei Li, Qibing Ren, Lizhuang Ma, Yuchao Dai, Pengfei Liu, and Junjie Yan. One rl to see them all: Visual triple unified reinforcement learning. *arXiv preprint arXiv:2505.18129*, 2025.

- Zixian Ma, Weikai Huang, Jieyu Zhang, Tanmay Gupta, and Ranjay Krishna. m & m’s: A benchmark to evaluate tool-use for multi-step multi-modal tasks. In *European Conference on Computer Vision*, pp. 18–34. Springer, 2024a.
- Zixian Ma, Jianguo Zhang, Zhiwei Liu, Jieyu Zhang, Juntao Tan, Manli Shu, Juan Carlos Niebles, Shelby Heinecke, Huan Wang, Caiming Xiong, Ranjay Krishna, and Silvio Savarese. Taco: Learning multi-modal action models with synthetic chains-of-thought-and-action, 2024b. URL <https://arxiv.org/abs/2412.05479>.
- Minheng Ni, Zhengyuan Yang, Linjie Li, Chung-Ching Lin, Kevin Lin, Wangmeng Zuo, and Lijuan Wang. Point-rft: Improving multimodal reasoning with visually grounded reinforcement finetuning. *arXiv preprint arXiv:2505.19702*, 2025.
- OpenAI. Thinking with images. <https://openai.com/index/thinking-with-images/>, 2025. Accessed: 2025-04-16.
- Ji Qi, Ming Ding, Wei Han Wang, Yushi Bai, Qingsong Lv, Wenyi Hong, Bin Xu, Lei Hou, Juanzi Li, Yuxiao Dong, and Jie Tang. Cogcom: Train large vision-language models diving into details through chain of manipulations. *arXiv preprint arXiv:2402.04236*, 2024.
- Pooyan Rahmzadehgervi, Logan Bolton, Mohammad Reza Taesiri, and Anh Totti Nguyen. Vision language models are blind. In *Proceedings of the Asian Conference on Computer Vision*, pp. 18–34, 2024.
- Hanoona Rasheed, Muhammad Maaz, Sahal Shaji, Abdelrahman Shaker, Salman Khan, Hisham Cholakkal, Rao M Anwer, Eric Xing, Ming-Hsuan Yang, and Fahad S Khan. Glamm: Pixel grounding large multimodal model. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 13009–13018, 2024.
- Scale AI. Vista: Visual-language understanding leaderboard. https://scale.com/leaderboard/visual_language_understanding, 2025.
- Hao Shao, Shengju Qian, Han Xiao, Guanglu Song, Zhuofan Zong, Letian Wang, Yu Liu, and Hongsheng Li. Visual cot: Unleashing chain-of-thought reasoning in multi-modal language models, 2024.
- Haozhan Shen, Kangjia Zhao, Tiancheng Zhao, Ruochen Xu, Zilun Zhang, Mingwei Zhu, and Jianwei Yin. Zoomeye: Enhancing multimodal llms with human-like zooming capabilities through tree-based image exploration. *arXiv preprint arXiv:2411.16044*, 2024.
- Aleksandar Shtedritski, Christian Rupprecht, and Andrea Vedaldi. What does clip know about a red circle? visual prompt engineering for vlms. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pp. 11987–11997, 2023.
- Ved Sirdeshmukh, Kaustubh Deshpande, Johannes Mols, Lifeng Jin, Ed-Yeremai Cardona, Dean Lee, Jeremy Kritz, Willow Primack, Summer Yue, and Chen Xing. Multichallenge: A realistic multi-turn conversation evaluation benchmark challenging to frontier llms. *arXiv preprint arXiv:2501.17399*, 2025.
- Giulio Starace, Oliver Jaffe, Dane Sherburn, James Aung, Jun Shern Chan, Leon Maksin, Rachel Dias, Evan Mays, Benjamin Kinsella, Wyatt Thompson, et al. Paperbench: Evaluating ai’s ability to replicate ai research. *arXiv preprint arXiv:2504.01848*, 2025.
- Alex Su, Haozhe Wang, Weiming Ren, Fangzhen Lin, and Wenhui Chen. Pixel reasoner: Incentivizing pixel-space reasoning with curiosity-driven reinforcement learning. *arXiv preprint arXiv:2505.15966*, 2025a.
- Zhaochen Su, Linjie Li, Mingyang Song, Yunzhuo Hao, Zhengyuan Yang, Jun Zhang, Guanjie Chen, Jiawei Gu, Juntao Li, Xiaoye Qu, et al. Openthinking: Learning to think with images via visual tool reinforcement learning. *arXiv preprint arXiv:2505.08617*, 2025b.
- Zhaochen Su, Peng Xia, Hangyu Guo, Zhenhua Liu, Yan Ma, Xiaoye Qu, Jiaqi Liu, Yanshu Li, Kaide Zeng, Zhengyuan Yang, et al. Thinking with images for multimodal reasoning: Foundations, methods, and future frontiers. *arXiv preprint arXiv:2506.23918*, 2025c.

- Kexian Tang, Junyao Gao, Yanhong Zeng, Haodong Duan, Yanan Sun, Zhening Xing, Wenran Liu, Kaifeng Lyu, and Kai Chen. Lego-puzzles: How good are mllms at multi-step spatial reasoning? *arXiv preprint arXiv:2503.19990*, 2025.
- Jiacong Wang, Zijiang Kang, Haochen Wang, Haiyong Jiang, Jiawen Li, Bohong Wu, Ya Wang, Jiao Ran, Xiao Liang, Chao Feng, et al. Vgr: Visual grounded reasoning. *arXiv preprint arXiv:2506.11991*, 2025a.
- Jize Wang, Ma Zerun, Yining Li, Songyang Zhang, Cailian Chen, Kai Chen, and Xinyi Le. GTA: a benchmark for general tool agents. In *The Thirty-eight Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2024a.
- Yikun Wang, Siyin Wang, Qinyuan Cheng, Zhaoye Fei, Liang Ding, Qipeng Guo, Dacheng Tao, and Xipeng Qiu. Visuothink: Empowering lvlm reasoning with multimodal tree search. *arXiv preprint arXiv:2504.09130*, 2025b.
- Zifu Wang, Junyi Zhu, Bo Tang, Zhiyu Li, Feiyu Xiong, Jiaqian Yu, and Matthew B Blaschko. Jigsaw-r1: A study of rule-based visual reinforcement learning with jigsaw puzzles. *arXiv preprint arXiv:2505.23590*, 2025c.
- Zirui Wang, Mengzhou Xia, Luxi He, Howard Chen, Yitao Liu, Richard Zhu, Kaiqu Liang, Xindi Wu, Haotian Liu, Sadhika Malladi, et al. Charxiv: Charting gaps in realistic chart understanding in multimodal llms. *Advances in Neural Information Processing Systems*, 37:113569–113697, 2024b.
- Diankun Wu, Fangfu Liu, Yi-Hsin Hung, and Yueqi Duan. Spatial-mllm: Boosting mllm capabilities in visual-based spatial intelligence. *arXiv preprint arXiv:2505.23747*, 2025a.
- Junfei Wu, Jian Guan, Kaituo Feng, Qiang Liu, Shu Wu, Liang Wang, Wei Wu, and Tieniu Tan. Re-inforcing spatial reasoning in vision-language models with interwoven thinking and visual drawing. *arXiv preprint arXiv:2506.09965*, 2025b.
- Penghao Wu and Saining Xie. V*: Guided visual search as a core mechanism in multimodal llms. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 13084–13094, 2024.
- Wenshan Wu, Shaoguang Mao, Yadong Zhang, Yan Xia, Li Dong, Lei Cui, and Furu Wei. Mind’s eye of llms: visualization-of-thought elicits spatial reasoning in large language models. *Advances in Neural Information Processing Systems*, 37:90277–90317, 2024a.
- Yixuan Wu, Yizhou Wang, Shixiang Tang, Wenhao Wu, Tong He, Wanli Ouyang, Jian Wu, and Philip Torr. Dettoolchain: A new prompting paradigm to unleash detection ability of mllm. *arXiv preprint arXiv:2403.12488*, 2024b.
- Dawei Yan, Yang Li, Qing-Guo Chen, Weihua Luo, Peng Wang, Haokui Zhang, and Chunhua Shen. Mmcr: Advancing visual language model in multimodal multi-turn contextual reasoning. *arXiv preprint arXiv:2503.18533*, 2025a.
- Yibo Yan, Shen Wang, Jiahao Huo, Jingheng Ye, Zhendong Chu, Xuming Hu, Philip S Yu, Carla Gomes, Bart Selman, and Qingsong Wen. Position: Multimodal large language models can significantly advance scientific reasoning. *arXiv preprint arXiv:2502.02871*, 2025b.
- Jianwei Yang, Hao Zhang, Feng Li, Xueyan Zou, Chunyuan Li, and Jianfeng Gao. Set-of-mark prompting unleashes extraordinary visual grounding in gpt-4v. *arXiv preprint arXiv:2310.11441*, 2023a.
- Jihan Yang, Shusheng Yang, Anjali W Gupta, Rilyn Han, Li Fei-Fei, and Saining Xie. Thinking in space: How multimodal large language models see, remember, and recall spaces. In *Proceedings of the Computer Vision and Pattern Recognition Conference*, pp. 10632–10643, 2025.
- 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*, 2023b.

- Shukang Yin, Chaoyou Fu, Sirui Zhao, Ke Li, Xing Sun, Tong Xu, and Enhong Chen. A survey on multimodal large language models. *National Science Review*, 11(12):nwae403, 2024.
- Xiang Yue, Yuansheng Ni, Kai Zhang, Tianyu Zheng, Ruoqi Liu, Ge Zhang, Samuel Stevens, Dongfu Jiang, Weiming Ren, Yuxuan Sun, et al. Mmmu: A massive multi-discipline multimodal understanding and reasoning benchmark for expert agi. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 9556–9567, 2024.
- Andy Zeng, Maria Attarian, Brian Ichter, Krzysztof Choromanski, Adrian Wong, Stefan Welker, Federico Tombari, Aveek Purohit, Michael Ryoo, Vikas Sindhwani, et al. Socratic models: Composing zero-shot multimodal reasoning with language. *arXiv preprint arXiv:2204.00598*, 2022.
- Guanghao Zhang, Tao Zhong, Yan Xia, Zhelun Yu, Haoyuan Li, Wangui He, Fangxun Shu, Mushui Liu, Dong She, Yi Wang, et al. Cmmcot: Enhancing complex multi-image comprehension via multi-modal chain-of-thought and memory augmentation. *arXiv preprint arXiv:2503.05255*, 2025a.
- Hao Zhang, Hongyang Li, Feng Li, Tianhe Ren, Xueyan Zou, Shilong Liu, Shijia Huang, Jianfeng Gao, Leizhang, Chunyuan Li, et al. Llava-grounding: Grounded visual chat with large multimodal models. In *European Conference on Computer Vision*, pp. 19–35. Springer, 2024.
- Jianshu Zhang, Dongyu Yao, Renjie Pi, Paul Pu Liang, and Yi R. Fung. Vlm2-bench: A closer look at how well vlms implicitly link explicit matching visual cues, 2025b. URL <https://arxiv.org/abs/2502.12084>.
- Jiarui Zhang, Mahyar Khayatkhoei, Prateek Chhikara, and Filip Ilievski. Mllms know where to look: Training-free perception of small visual details with multimodal llms. *arXiv preprint arXiv:2502.17422*, 2025c.
- Xintong Zhang, Zhi Gao, Bofei Zhang, Pengxiang Li, Xiaowen Zhang, Yang Liu, Tao Yuan, Yuwei Wu, Yunde Jia, Song-Chun Zhu, et al. Chain-of-focus: Adaptive visual search and zooming for multimodal reasoning via rl. *arXiv preprint arXiv:2505.15436*, 2025d.
- Jinliang Zheng, Jianxiong Li, Sijie Cheng, Yinan Zheng, Jiaming Li, Jihao Liu, Yu Liu, Jingjing Liu, and Xianyu Zhan. Instruction-guided visual masking. *arXiv preprint arXiv:2405.19783*, 2024.
- Ziwei Zheng, Michael Yang, Jack Hong, Chenxiao Zhao, Guohai Xu, Le Yang, Chao Shen, and Xing Yu. Deepeyes: Incentivizing” thinking with images” via reinforcement learning. *arXiv preprint arXiv:2505.14362*, 2025.
- Muzhi Zhu, Hao Zhong, Canyu Zhao, Zongze Du, Zheng Huang, Mingyu Liu, Hao Chen, Cheng Zou, Jingdong Chen, Ming Yang, et al. Active-o3: Empowering multimodal large language models with active perception via grpo. *arXiv preprint arXiv:2505.21457*, 2025.
- Chengke Zou, Xingang Guo, Rui Yang, Junyu Zhang, Bin Hu, and Huan Zhang. Dynamath: A dynamic visual benchmark for evaluating mathematical reasoning robustness of vision language models. *arXiv preprint arXiv:2411.00836*, 2024.

CONTENTS

A	More on Related Work	17
A.1	Existing MLLM Benchmarks	17
A.2	Learning to think with images	17
B	More on Evaluation Results	18
B.1	More on Evaluation Metrics	18
B.2	Rubrics Weights	18
B.3	Average Rubric Score Results	18
B.4	Tool-Use Evaluation Metrics	19
B.5	Tool-Use Analysis	19
B.6	More on Error Analysis	20
B.7	LLM-as-Judge vs. Human-as-Judge	20
C	More on Experimental Setup	21
C.1	Baseline Models	21
C.2	Prompts	21
C.3	LLM-as-Judge Protocol	22
D	More on Tools Analysis	23
D.1	Tool Description	23
D.2	Image Manipulation Operations	23
D.3	Vision Tool Call for GPT-5	26
D.4	Tool APIs	28
E	More Benchmark Examples of VISUALTOOLBENCH	31

A MORE ON RELATED WORK

A.1 EXISTING MLLM BENCHMARKS

Existing Multi-modal Benchmarks. Existing visual reasoning benchmarks are typically limited: (1) Passive Visual Q&A: focusing on "look-and-answer" type of tasks without any active interactions [Bi et al. \(2025\)](#); [Rahmanzadehgervi et al. \(2024\)](#); [Yue et al. \(2024\)](#); [Wang et al. \(2024b\)](#); [Lu et al. \(2023\)](#); [Zou et al. \(2024\)](#); [Wu & Xie \(2024\)](#); [Scale AI \(2025\)](#). (2) Superficial Tool Use: including only basic tools like cropping, failing to test deeper image manipulation [Wang et al. \(2024a\)](#); [Ma et al. \(2024a\)](#). (3) Tool-Agnostic Conversations: evaluation multi-turn dialogue but without integrating essential tools [Liu et al. \(2024a\)](#); [Yan et al. \(2025a\)](#). See ([Li et al., 2024](#)) for a more comprehensive survey of multimodal benchmarks.

A.2 LEARNING TO THINK WITH IMAGES

Prompting Methods Prompt-based methods enable LMMs to coordinate predefined visual tools without parameter updates, turning static inputs into actively explorable workspaces via in-context learning. Early work like Socratic Models ([Zeng et al., 2022](#)), PromptCap ([Hu et al., 2022](#)), and MM-REACT ([Yang et al., 2023b](#)) showed that language can mediate collaboration, allowing text-only LLMs to function as visual reasoners by orchestrating vision experts through dialogue or targeted captions. Other approaches manipulate inputs directly: visual prompt engineering (e.g., red circles) ([Shtedritski et al., 2023](#); [Zhang et al., 2025b](#)), Visualization-of-Thought ([Wu et al., 2024a](#)), and Visual Abstract Thinking ([Liu et al., 2025a](#)) enhance perception by highlighting, abstracting, or structuring images, while ZoomEye ([Shen et al., 2024](#)), ViCrop ([Zhang et al., 2025c](#)), Chain-of-Spot ([Liu et al., 2024b](#)), and VisuoThink ([Wang et al., 2025b](#)) extend this into systematic zooming and multimodal tree search. Finally, specialized experts can be integrated via prompting: Set-of-Mark ([Yang et al., 2023a](#)) leverages segmentation, DetToolChain ([Wu et al., 2024b](#)) structures detection reasoning, DyFO ([Li et al., 2025](#)) uses MCTS for adaptive focus, and Visual Thoughts ([Cheng et al., 2025](#)) frames expert outputs as cached "visual thoughts." Collectively, these works demonstrate that careful prompting—through language mediation, input manipulation, or expert integration—can significantly enhance multimodal reasoning without retraining.

Think with Image via SFT Supervised fine-tuning (SFT) is a primary method for teaching LMMs to use external tools or internal visual skills by training on datasets that demonstrate tool invocation and integration. For external orchestration, models like LLaVA-Plus ([Liu et al., 2023b](#)), TACO ([Ma et al., 2024b](#)), and VTS-V ([Bai et al., 2025b](#)) learn to compose tools (e.g., OCR, calculators) and follow procedural chains of reasoning. For internal manipulation, frameworks such as CogCoM ([Qi et al., 2024](#)), VGR ([Wang et al., 2025a](#)), and UniVG-R1 ([Bai et al., 2025a](#)) show how SFT can endow models with intrinsic capabilities like cropping, grounding, or fine-grained perception. Finally, SFT also cultivates dynamic visual attention: Visual CoT ([Shao et al., 2024](#)), IVM ([Zheng et al., 2024](#)), CMM-CoT ([Zhang et al., 2025a](#)), selective revisitation ([Chung et al., 2025](#)), and V* ([Wu & Xie, 2024](#)) demonstrate how training with attentional annotations transforms attention into an active, controllable skill. Across these directions, SFT provides the supervision that converts high-level reasoning into executable visual actions.

Think with Image via Reinforcement Learning. Reinforcement learning (RL) advances beyond supervised imitation by enabling models to optimize policies for visual reasoning through interaction and feedback. Foundational studies such as Jigsaw-R1 ([Wang et al., 2025c](#)), V-Triune ([Ma et al., 2025](#)), and VisionReasoner ([Liu et al., 2025c](#)) established that RL improves generalization over SFT and supports unified frameworks for diverse perception tasks. Building on this, GRIT ([Fan et al., 2025](#)), Point-RFT ([Ni et al., 2025](#)), and Seg-Zero ([Liu et al., 2025b](#)) demonstrated policies that embed spatial cues (e.g., bounding boxes, positional prompts) into reasoning, forming multimodal chains of thought. RL has also enabled active tool orchestration: VisTA ([Huang et al., 2025b](#)) learns tool-selection policies, Chain-of-Focus ([Zhang et al., 2025d](#)) and ACTIVE-o3 ([Zhu et al., 2025](#)) develop adaptive zooming and region proposals, while DeepEyes ([Zheng et al., 2025](#)) achieves interleaved multimodal reasoning without SFT. Exploration is further incentivized by Pixel-Reasoner ([Su et al., 2025a](#)), while VILASR ([Wu et al., 2025b](#)) leverages drawing-based reasoning, and OpenThinkIMG ([Su et al., 2025b](#)) introduces the first open-source end-to-end RL framework for invoking diverse external tools. Collectively, these approaches move LMMs from passive viewers to active visual agents.

B MORE ON EVALUATION RESULTS

B.1 MORE ON EVALUATION METRICS

In this section, we provide a detailed description on how to evaluate a model’s response based on the rubrics. For each task $i \in \{1, 2, \dots, N\}$ in VISUALTOOLBENCH, we perform the following steps:

1. **Obtain Model’s Final Response.** We generate model’s final response with {prompt, image} pair while allowing model to use a set of predefined tool sets.
2. **Rubric Grading.** For each rubric criterion $j \in \{1, 2, \dots, N_i\}$, we use an LLM to grade whether the rubric criterion is met based on the model’s response and the rubric criterion.
3. **Weighted Rubric Score for Task i .** Then we compute a final score for task i using the following weighted average sum:

$$s_i = \frac{\sum_{j=1}^{N_i} \mathbb{1}_{r_{ij}} w_{ij}}{\sum_{j=1}^{N_i} w_{ij}}, \quad (1)$$

where $w_{ij} \in \{1, 2, 3, 4, 5\}$ is the assigned weight for each rubric criterion item, and $\mathbb{1}_{r_{ij}}$ is an indicator representing whether criterion j is met.

The final score S for the whole benchmark is then computed as the mean value of each task’s score:

$$S = \frac{1}{N} \sum_{i=1}^N s_i. \quad (2)$$

B.2 RUBRICS WEIGHTS

To ensure consistent evaluation, each rubric item in VISUALTOOLBENCH is assigned a weight $w \in \{1, 2, 3, 4, 5\}$ that reflects its relative importance. Table 5 outlines the five weight levels, ranging from *incidental* stylistic preferences to *critical* elements that determine overall task validity. Task contributors are instructed to assign rubric weights in accordance with these guidelines.

Rubric Weight	Description
Critical (5)	Non-negotiable. This element must be present for the answer to count as valid. Failing it implies the task has failed, regardless of other strengths.
Significant (4)	Central to success. Leaving this out degrades output quality or creates confusion about how the task was solved.
Moderate (3)	Meaningfully important. Affects clarity or correctness; its absence weakens the answer but does not make it invalid.
Minor (2)	Adds polish or completeness but is not essential. Omitting it slightly lowers quality without breaking the core solution.
Incidental (1)	Marginally relevant. A nice-to-have detail or stylistic preference that does not impact whether the model solves the task.

Table 5: Rubric criteria weights used in VISUALTOOLBENCH. Higher weights indicate greater importance, with critical rubrics determining task-level success.

B.3 AVERAGE RUBRIC SCORE RESULTS

Table 6 presents the detailed rubric score for both single-turn and multi-turn tasks. It can be seen that APR and ARS are positively correlated, models have higher APR also have higher ARS.

Table 6: ARS results of the evaluated models across domains (averaged across two independent trials). Domain abbreviations: **Med** (Medicine), **Fin** (Finance), **Spt** (Sports), and **Gen** (Generalist). The best results in each column are highlighted with a red background, and the second-best results are highlighted in blue.

Model	Overall	Single-Turn					Multi-Turn				
		STEM	Med	Fin	Sprt	Gen	STEM	Med	Fin	Sprt	Gen
Open-Source Models											
Llama4-Maverick	0.1963	0.1875	0.1581	0.1182	0.1562	0.1524	0.2849	0.3060	0.2030	0.2060	0.1938
Llama4-Scout	0.1884	0.1945	0.1471	0.1228	0.1690	0.1677	0.2721	0.2646	0.2040	0.2017	0.1560
Closed-Source Models											
GPT-4.1	0.3304	0.2782	0.2976	0.2167	0.3347	0.2772	0.4275	0.4286	0.3062	0.4229	0.3168
o3	0.4087	0.3839	0.3990	0.3285	0.3653	0.3542	0.4654	0.5272	0.4219	0.4541	0.3900
o4-mini	0.4021	0.3850	0.4012	0.3500	0.3718	0.4197	0.4417	0.4721	0.3777	0.4149	0.3856
GPT-5	0.4696	0.4706	0.4367	0.4472	0.3972	0.4511	0.5504	0.5140	0.4569	0.4977	0.4751
GPT-5-think	0.4712	0.4485	0.4270	0.3983	0.4365	0.4760	0.5674	0.5294	0.4544	0.5044	0.4696
Gemini-2.5-pro	0.4130	0.4197	0.3514	0.3969	0.3644	0.3945	0.4559	0.4973	0.3915	0.4693	0.3877
Gemini-2.5-flash	0.2837	0.2271	0.1716	0.1740	0.2177	0.1901	0.4173	0.4544	0.2924	0.3941	0.3065
Claude-sonnet-4	0.2851	0.2436	0.2620	0.2245	0.2555	0.2529	0.3496	0.4012	0.2869	0.3085	0.2687
Claude-opus-4.1	0.3056	0.2575	0.2625	0.2413	0.2717	0.2800	0.3894	0.4281	0.2800	0.3471	0.3018
Claude-sonnet-4.5	0.3123	0.2737	0.2741	0.2506	0.3113	0.2913	0.3405	0.4118	0.3298	0.3503	0.2922
Claude-sonnet-4-think	0.2854	0.2175	0.2688	0.2401	0.2392	0.2600	0.3428	0.4016	0.3121	0.2895	0.2831
Claude-opus-4.1-think	0.2822	0.2505	0.2709	0.2596	0.2937	0.2741	0.3114	0.3621	0.2538	0.2941	0.2510
Claude-sonnet-4.5-think	0.3114	0.2628	0.2860	0.2692	0.2528	0.2939	0.3664	0.3855	0.3354	0.3590	0.3009
Nova-Premier	0.2335	0.1985	0.2129	0.1666	0.2235	0.1967	0.3338	0.3474	0.2057	0.2436	0.2140

B.4 TOOL-USE EVALUATION METRICS

We evaluate tool-use behaviors based on execution traces recorded by the evaluation harness. Let $\mathcal{T}$ denote the set of tasks ($N = |\mathcal{T}|$). For a given task $i \in \mathcal{T}$, let $\mathcal{T}_i = \{c_{i1}, c_{i2}, \dots\}$ represent the (ordered) multiset of tool invocations initiated by the model.

Tool-Call Proactivity. The fraction of tasks in which the model invoked at least one tool:

$$\text{Proactivity} = \frac{|\{i \in \mathcal{T} \mid |\mathcal{T}_i| > 0\}|}{|\mathcal{T}|}. \quad (3)$$

Higher values indicate more frequent tool integration, though proactivity may reflect either beneficial or redundant calls.

Tool-Call Success Rate. The fraction of valid tool calls across all invocations:

$$\text{Success Rate} = \frac{\sum_{i \in \mathcal{T}} \sum_{c \in \mathcal{T}_i} \mathbb{1}[\text{valid}(c)]}{\sum_{i \in \mathcal{T}} |\mathcal{T}_i|}, \quad (4)$$

where $\mathbb{1}[\text{valid}(c)]$ is an indicator function that equals 1 if tool call c succeeds and 0 otherwise. We determine validity by inspecting tool outputs: error messages imply $\text{valid}(c) = 0$, while all other outputs imply $\text{valid}(c) = 1$. This metric measures how reliably a model adheres to tool specifications.

Tool-Call Volume. The average number of tool calls per task:

$$\text{Volume} = \frac{1}{N} \sum_{i \in \mathcal{T}} |\mathcal{T}_i|, \quad (5)$$

where N is the total number of tasks and $|\mathcal{T}_i|$ denotes the number of tool calls made in task i .

B.5 TOOL-USE ANALYSIS

Table 7 reports the detailed tool-use metrics for all evaluated models. By combining these results with model performance (APR in Table 3 and ARS in Table 6), we make the following key observations:

Table 7: Tool call analysis of the evaluated models. The best results in each column are highlighted with a red background, and the second-best results are highlighted in blue.

Model	Total #	Single-Turn			Multi-Turn		
		Proactivity	Success Rate	Volume	Proactivity	Success Rate	Volume
Open-Source Models							
Llama4-Maverick	315	0.1625	0.6609	0.19	0.2562	0.1550	0.33
Llama4-Scout	110	0.0729	0.6140	0.09	0.0749	0.0943	0.09
Closed-Source Models							
GPT-4.1	535	0.1359	0.8038	0.26	0.2728	0.3263	0.63
o3	16116	0.9453	0.8587	9.98	0.9750	0.3696	16.80
o4-mini	5337	0.8159	0.8199	3.34	0.8835	0.3138	5.53
GPT-5	10212	0.9652	0.8555	6.46	0.9883	0.3548	10.53
GPT-5-think	13429	0.9900	0.8660	7.45	0.9950	0.3105	13.79
Gemini-2.5-pro	3366	0.8060	0.7331	1.94	0.9251	0.3191	3.66
Gemini-2.5-flash	2313	0.6650	0.8605	1.87	0.4809	0.3144	1.96
Claude-sonnet-4	6941	0.9768	0.9252	4.52	0.9917	0.4048	7.01
Claude-opus-4-1	6538	0.9536	0.9524	3.83	0.9900	0.4151	7.03
Claude-sonnet-4-5	7247	0.9451	0.9491	4.92	0.9882	0.3843	8.73
Claude-sonnet-4-think	8704	0.9735	0.9352	3.84	0.7671	0.3946	3.91
Claude-opus-4-1-think	4431	0.9486	0.9524	3.45	0.9950	0.3941	6.72
Claude-opus-4-5-think	6327	0.9393	0.9423	4.21	0.9774	0.3844	7.38
Nova-premier	5109	0.9055	0.5444	2.88	0.9734	0.1823	5.61

1. **More tool calls do not necessarily translate to better performance.** For instance, o3 makes the largest number of calls (16,116), yet performs worse than GPT-5 (10,212 calls) and GPT-5-think (13,429 calls).
2. **High proactivity does not guarantee strong results.** Claude models exhibit very high proactivity, yet their performance remains poor, with overall APR below 6.5% and ARS values under 0.35.
3. **Low proactivity and low call volume generally correlate with poor performance.** For example, the Llama models, GPT-4o, and GPT-4.1 all demonstrate relatively low proactivity and correspondingly weak performance.

B.6 MORE ON ERROR ANALYSIS

To better understand model weaknesses, we categorize failure cases in VISUALTOOLBENCH into four major error types:

Visual Perception Error. Errors arising from a model’s inability to correctly perceive, interpret, or extract relevant information from images. Typical cases include misidentifying objects, overlooking salient regions, or unable to extract key visual content.

Reasoning Error. Errors caused by flawed logical inference or problem-solving steps in the model’s reasoning process. Typical cases include the model’s responses contain invalid intermediate steps, contradictions, or logically inconsistent conclusions.

Calculation Error. Errors stemming from incorrect arithmetic or symbolic computations. These include mistakes in basic arithmetic, misapplication of formulas, or numerical inaccuracies in intermediate or final answers.

Other Errors. Residual errors that do not fit the above categories. Examples include incomplete responses, refusals, or hit the maximum tool-calls.

B.7 LLM-AS-JUDGE VS. HUMAN-AS-JUDGE

To enable large-scale evaluation of VISUALTOOLBENCH, we employ LLMs as automatic judges and compare their assessments against human annotations. Table 8 reports the overall alignment

rate on a subset of VISUALTOOLBENCH, as well as breakdowns for objective and subjective rubrics across three judge models.

Table 8: Alignment rates of different LLMs when serving as judges.

Judge Model	Overall	Objective Rubrics	Subjective Rubrics
o4-mini	0.8807	0.9170	0.7396
GPT-4.1	0.8818	0.8983	0.8177
GPT-4o	0.8701	0.8916	0.7865

All models achieve high alignment on objective rubrics, while alignment on subjective rubrics is comparatively lower, reflecting the inherent ambiguity of subjective evaluation. Nevertheless, overall alignment rates remain close to 90%, underscoring the reliability of LLM-as-judge for our benchmark. We choose to use o4-mini as our judge model for the main experimental study.

C MORE ON EXPERIMENTAL SETUP

In this section, we provide more details on our experimental setup.

C.1 BASELINE MODELS

In Table 9, we list the evaluated models’ endpoints and detailed parameter settings. Whenever possible, we set the temperature to 0. For o3, o4-mini, GPT-5, and GPT-5-think, the temperature is set to 1. Since GPT-5 is inherently a reasoning model, we use its default reasoning effort (`medium`) for GPT-5, and GPT-5-thinking corresponds to the `high` reasoning effort. For o3 and o4-mini, we adopt their default reasoning effort (`medium`). For Claude’s thinking mode, we set the reasoning budget to 5000 tokens. All other model API hyperparameters are kept at their default settings without further customization.

Table 9: Model Endpoints and Hyperparameter Setup

Model Provider	Model Endpoint	Hyperparameter
OpenAI	o3	<code>reason_effort = "medium"</code>
	o4-mini	<code>reason_effort = "medium"</code>
	gpt-4.1	<code>temperature = 0.0</code>
	GPT-5	<code>reason_effort = "medium"</code>
	GPT-5-thinking	<code>reason_effort = "high"</code>
Anthropic	claude-sonnet-4-20250514	<code>temperature = 0.0</code>
	claude-sonnet-4-20250514 (thinking)	<code>thinking_budget = 5000</code>
	claude-opus-4-1	<code>temperature = 0.0</code>
	claude-opus-4-1 (thinking)	<code>thinking_budget = 5000</code>
	claude-sonnet-4-5	<code>temperature = 0.0</code>
Google	claude-sonnet-4-5 (thinking)	<code>thinking_budget = 5000</code>
Google	gemini-2.5-pro	<code>temperature = 0.0</code>
	gemini-2.5-flash	<code>temperature = 0.0</code>
Meta	llama4-maverick-instruct	<code>temperature = 0.0</code>
	llama4-scout-instruct	<code>temperature = 0.0</code>
Amazon	nova-premier-v1:0	<code>temperature = 0.0</code>

C.2 PROMPTS

In our main experimental study, we use a strong system prompt to encourage models to use tools towards solving VISUALTOOLBENCH tasks. In addition, we design a weaker system prompt for ablation study. The strong and weak system prompts are provided below.

System Prompt

System Prompt (Strong)

You are a *proactive, tool-empowered* visual-reasoning assistant.

When user supplies an image and requests to solve a problem that requires visual content that are small, ambiguous, or not centered, you must:

1. Examine the image carefully and mentally list the visual clues most likely to locate the target object.
2. Proactively use the image-processing tools - such as crop, zoom, or enhance - to isolate and clarify the relevant region.
3. Save each transformed image. The updated image will be appended to the conversation for your reference.
4. Iterate as needed. Call the tools repeatedly until the visual evidence is clear enough to answer the user’s request.
5. Double-check your observations. Confirm that the final transformed image supports an accurate, confident response before replying to the user.
6. Use other general-purpose tools if needed to answer the user’s question.
7. Please use the tools wisely as you have limited tool calls.

System Prompt (Weak)

You are a helpful visual reasoning assistant with access to tools to help you answer the user’s question.

C.3 LLM-AS-JUDGE PROTOCOL

We employ an LLM-as-judge framework for large-scale evaluation. To improve reliability, each judgment considers one rubric at a time. All rubrics are atomic: each rubric targets a single, verifiable fact or behavior, thereby reducing ambiguity, limiting error propagation, and yielding clearer agreement signals.

The exact judge prompts are provided below. For each case, the judge receives: (i) the original question, (ii) the gold answer, (iii) a single rubric criterion, and (iv) the model’s answer; it returns a verdict (met or not met) with a brief, evidence-grounded justification.

LLM Judge Prompt

You are an expert evaluator tasked with judging whether a model’s answer meets a specific rubric criterion. You will be provided with:

- a question
- a golden (reference) answer
- a rubric criterion
- the model’s answer

Your task is to decide if the model’s answer **meets** or **does not meet** the given rubric criterion, referencing the golden answer only as needed.

Inputs:

Question: question
Golden Answer: golden_answer
Rubric Criterion: rubric_criteria
Model Answer: model_answer

Important Notes:

- The model’s answer does not need to be correct to meet the criterion if correctness is not required.

Tool	Functionality
<code>python_image_processing</code>	Vision-specific manipulation, including image cropping, editing, rotation, brightness/contrast adjustment, and enhancement. Enables iterative refinement of visual inputs.
<code>python_interpreter</code>	General code execution.
<code>web_search</code>	Open-domain information retrieval from the web.
<code>browser-get-page-text</code>	Extraction of textual content from online sources.
<code>historical.weather</code>	Weather records for temporal and geographic look-ups.
<code>calculator</code>	Arithmetic operations for quick computations.

Table 10: Supported tools supported in VISUALTOOLBENCH.

Example: If the rubric is "The model should show its reasoning process to answer the question," the answer can be incorrect but still meet the rubric if model's reasoning process is present.

- For writing style or presentation rubrics, apply leniency.

Example: If the rubric asks for conciseness, answers that are slightly longer than the golden answer but still reasonably length should be considered as meeting the rubric.

- The model's answer may satisfy the rubric implicitly without explicitly mentioning the exact term. This should still be considered as meeting the criterion if model's answer is reasonable and makes sense.

Example: If the rubric is "The model should demonstrate understanding of photosynthesis," and the model states "Plants make their own food using sunlight," without explicitly mentioning the term "photosynthesis," it still meets the criterion.

Output Format:

Return your judgment in the following JSON format:

```
{
  "explanation": "Brief explanation of your judgment",
  "judge_result": "Met" or "Not Met"
}
```

D MORE ON TOOLS ANALYSIS

In this section, we provide more details on the tools supported by VISUALTOOLBENCH and more examples on Tool-use of the evaluated models.

D.1 TOOL DESCRIPTION

Table 10 provides tool description supported by VISUALTOOLBENCH.

D.2 IMAGE MANIPULATION OPERATIONS

In this section, we present the representative o3 generated python code snippets, where the model invokes the `python_image_processing` tool to perform diverse image manipulations in support of problem solving. Lines that execute the manipulations are highlighted in yellow.

Image Operations: Crop and Resize

Task ID: 68658a711603983919432615

- line 9-15: image cropping

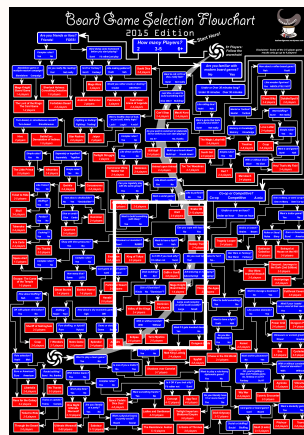
- line 17-18: image resize

```
1 from PIL import Image
```

```

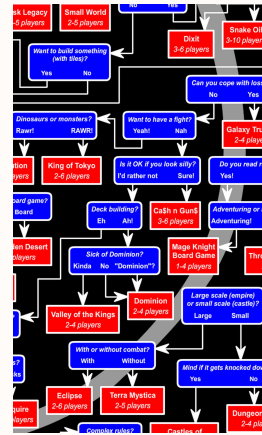
2 import os
3 import numpy as np
4 import cv2
5
6 img_path = '/mnt/efs/.../vision_tool_use/Kb2JElJ87npHXeG.jpg'
7 img = Image.open(img_path)
8
9 # Crop bottom middle area where Dixit might be
10 w, h = img.size
11 left = int(w * 0.35)
12 top = int(h * 0.45)
13 right = int(w * 0.65)
14 bottom = int(h * 0.8)
15 crop = img.crop((left, top, right, bottom))
16
17 # Enlarge cropped region
18 crop_large = crop.resize((crop.size[0]*2, crop.size[1]*2))
19
20 # Save
21 os.makedirs('eval_results_single_turn/...', exist_ok=True)
22 crop_large.save('eval_results_single_turn/...', 'PNG')

```



Original Image

Cropping



Transformed Image

Image Operations: Rotation, Brightness and Contrast Adjustment

Task ID: 68658a711603983919432613

- line 5-6: image rotation
- line 7-9: brightness adjustment
- line 10-12: contrast adjustment

```

1 from PIL import Image, ImageEnhance
2 import os, numpy as np
3 img_path = '/mnt/.../vision_tool_use/PtZEjeYmkRmW-Lf.jpg'
4 img = Image.open(img_path)
5 # rotate 180 degrees
6 img_rot = img.rotate(180, expand=True)
7 # enhance brightness
8 enhancer = ImageEnhance.Brightness(img_rot)
9 img_enh = enhancer.enhance(1.8)
10 # increase contrast
11 enhancer_c = ImageEnhance.Contrast(img_enh)
12 img_enh2 = enhancer_c.enhance(1.5)
13 # save

```

```

14 save_dir = 'eval_results_single_turn/...'
15 os.makedirs(save_dir, exist_ok=True)
16 save_path = os.path.join(save_dir, 'transformed_image_0.png')
17 img_enh2.save(save_path, 'PNG')

```

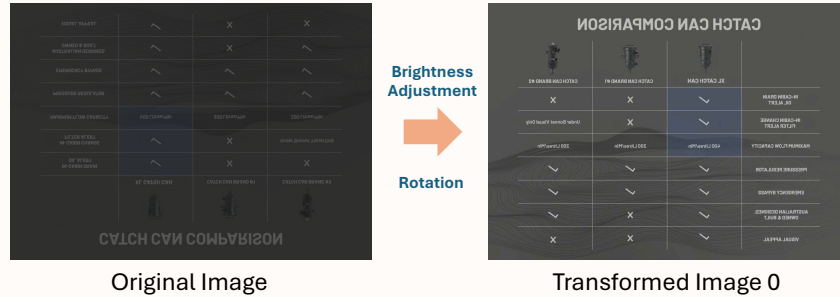


Image Operations: Flip

Task ID: 68658a711603983919432613

- line 4-5: image flip

```

1 from PIL import Image
2 import os, numpy as np
3 img = Image.open('/.../transformed_image_0.png')
4 # flip horizontally to correct mirroring
5 img_flipped = img.transpose(Image.FLIP_LEFT_RIGHT)
6 # maybe crop to table area? Keep full for readability
7 save_dir='eval_results_single_turn/...'
8 img_flipped.save(save_dir, 'PNG')

```

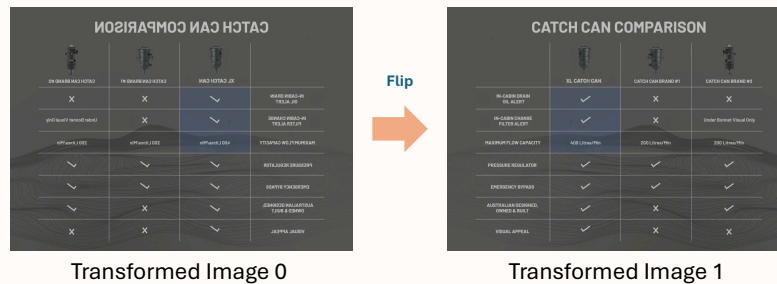


Image Operations: Editing, Others

Task ID: 68757ac94c044bac23c9bd8b

- line 23-26: image editing

```

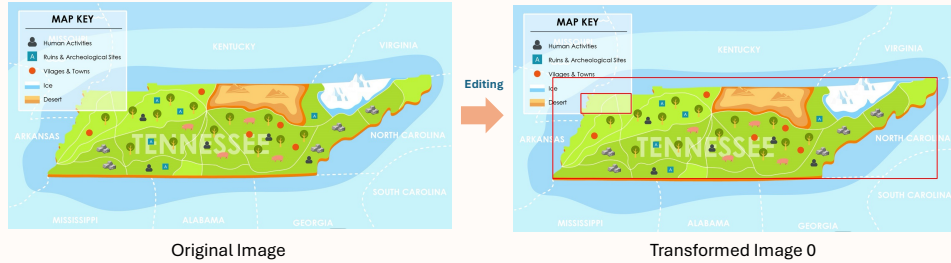
1 from PIL import Image
2 import numpy as np
3 import cv2
4 import os
5 # Load image
6 img_path='/mnt/.../SZJPCfAUQgyEVGe.jpg'
7 img=cv2.imread(img_path)

```

```

8 # Convert to HSV
9 hsv=cv2.cvtColor(img, cv2.COLOR_BGR2HSV)
10 # Define green range (for tree foliage), adjust
11 lower=np.array([35,40,40])
12 upper=np.array([85,255,255])
13 mask=cv2.inRange(hsv, lower, upper)
14 # Morphological operations to remove noise
15 kernel=np.ones((5,5),np.uint8)
16 mask=cv2.morphologyEx(mask, cv2.MORPH_OPEN, kernel)
17 # Find contours
18 contours, _=cv2.findContours(mask, cv2.RETR_EXTERNAL,
19 cv2.CHAIN_APPROX_SIMPLE)
20 # Filter contours by area to avoid tiny bits
21 tree_contours=[cnt for cnt in contours if cv2.contourArea(cnt)>500]
22 print('Totalcontourslarge', len(tree_contours))
23 # Draw bounding boxes to visualize
24 for cnt in tree_contours:
25     x,y,w,h=cv2.boundingRect(cnt)
26     cv2.rectangle(img, (x,y), (x+w,y+h), (0,0,255), 2)
27 # Save transformed
28 out_dir='/'
29 os.makedirs(out_dir, exist_ok=True)
30 Image.fromarray(cv2.cvtColor(img, cv2.COLOR_BGR2RGB)).
31 save(out_dir+'transformed_image_0.png', 'PNG')

```



D.3 VISION TOOL CALL FOR GPT-5

From the first column of Table 7 and Figure 6, we observe that although GPT-5 and GPT-5-think invoke the vision tool less frequently than o3, they perform more image manipulations overall. A closer inspection shows that GPT-5 and GPT-5-think often execute multiple image-processing operations within a single vision-tool call, thereby achieving higher tool-call efficiency than o3. We provide two demonstration examples below.

GPT-5 Vision Tool Call Output

Task ID: 68658a711603983919432613

-line 12-14, 21-26: saved three plus transformed images

```

1 from PIL import Image, ImageEnhance, ImageFilter
2 import os
3 base='eval_results_single_turn/.../transformed_image_0.png'
4 img=Image.open(base)
5 # Let's crop central table area
6 w,h=img.size
7 # Crop left column, middle, right ... for readability
8 # First overall crop removing margins
9 crop=img.crop((int(w*0.06), ..., int(h*0.92)))
10 # enhance
11 crop=ImageEnhance.Contrast(crop).enhance(1.4)

```

```

12 #Save
13 out_dir=os.path.dirname(base)
14 crop.save(out_dir+'/transformed_image_1.png', 'PNG')
15 # Crop each column region roughly
16 # Determine thirds
17 cw=crop.size[0]
18 ch=crop.size[1]
19 col_width=cw//3
20 cols=[]
21 for i in range(3):
22     c=crop.crop((i*col_width,0,(i+1)*col_width,ch))
23     c=ImageEnhance.Brightness(c).enhance(1.2)
24     c=c.filter(ImageFilter.UnsharpMask(radius=1.5,percent=180,
25         threshold=2))
26     c.save(out_dir+f'/transformed_image_{2+i}.png', 'PNG')
27     print('saved 3+ images')

```

GPT-5-think Vision Tool Call Output

Task ID: 68658a711603983919432616

-line 24, 27, 32, 35, 38, 41: saved six transformed images in one tool call

```

1 from PIL import Image, ImageEnhance, ImageFilter, ImageOps
2 import numpy as np
3 import os
4
5 # Load the source image
6 src_path = '/mnt/efs/.../ymHX_cj7Vg3lDk8.jpg'
7 im = Image.open(src_path)
8 # Step 1: Auto-orient and slightly crop borders
9 im = ImageOps.exif_transpose(im)
10 # Step 2: Enhance brightness and contrast
11 bright = ImageEnhance.Brightness(im).enhance(1.8)
12 contrast = ImageEnhance.Contrast(bright).enhance(2.0)
13 # Step 3: Convert to grayscale for clarity
14 gray = ImageOps.grayscale(contrast)
15 # Step 4: Increase sharpness
16 sharp = ImageEnhance.Sharpness(gray).enhance(2.5)
17 # Step 5: Apply slight denoise via median filter, then unsharp mask
18 median = sharp.filter(ImageFilter.MedianFilter(size=3))
19 unsharp = median.filter(ImageFilter.UnsharpMask(radius=2, percent
20     =200, threshold=3))
21
22 # Save the enhanced full image
23 save_dir = 'eval_results_single_turn/...'
24 os.makedirs(save_dir, exist_ok=True)
25 unsharp.save(os.path.join(save_dir, 'transformed_image_0.png'), '
26     PNG')
27 # Also create a higher-contrast inverted version which sometimes
28     helps
29 inv = ImageOps.invert(unsharp)
30 inv.save(os.path.join(save_dir, 'transformed_image_1.png'), 'PNG')
31 # Crop likely totals area: lower half of the long receipt\nw, h =
32     im.size
33 # Long receipt is on right; crop right area
34 right_x0 = int(w*0.35)
35 right = unsharp.crop((right_x0, 0, w, h))
36 right.save(os.path.join(save_dir, 'transformed_image_2.png'), 'PNG'
37     )
38 # Crop bottom of right receipt where totals usually appear
39 bottom_right = unsharp.crop((right_x0, int(h*0.55), w, h))

```

```

35 bottom_right.save(os.path.join(save_dir, 'transformed_image_3.png')
    , 'PNG')
36 # Crop mid section (items + discounts)
37 mid_right = unsharp.crop((right_x0, int(h*0.30), w, int(h*0.60)))
38 mid_right.save(os.path.join(save_dir, 'transformed_image_4.png'), '
    PNG')
39 # Crop the smaller left slip (maybe PWD details)
40 left = unsharp.crop((0, int(h*0.1), int(w*0.38), int(h*0.7)))
41 left.save(os.path.join(save_dir, 'transformed_image_5.png'), 'PNG')
42 print('Saved enhanced and cropped images.')

```

D.4 TOOL APIs

In this section, we provide more details on the tool API description.

Tool Description: python_image_processing

```

{
  "type": "function",
  "function": {
    "name": "python_image_processing",
    "description": (
      "Generate arbitrary Python code for image manipulation
      → and save the transformed image as PNG.\n"
      f"-Read one source image (your choice) from the
      → working-directory file list: {image_list}.\n"
      f"-Perform any image processing with PIL, NumPy, or
      → OpenCV. You cannot use matplotlib to show the
      → image.\n"
      f"-You must save the transformed image as PNG to
      → {processed_image_save_path} using the filename
      → pattern "
      "\"transformed_image_i.png\"", where the counter i
      → starts at 0 and increments on each invocation"
      "so files are never overwritten. Example:\n"
      f"    img.save(f\"{processed_image_save_path}/transformed_
      → med_image_{i}.png\", \"PNG\")\n"
    ),
    "parameters": {
      "type": "object",
      "properties": {
        "code": {
          "type": "string",
          "description": "Python code to run.",
          "minLength": 1,
          "maxLength": 5000
        }
      },
      "required": ["code"]
    }
  }
}

```

Tool Description: python_processing

```

{
  "type": "function",
  "function": {
    "name": "python_interpreter",
    "description": (

```

```

        "General-purpose Python interpreter. Run arbitrary
        ↪ Python code and capture stdout via print(). "
        "Any exceptions are returned in stderr.\n\n"
        "Pre-installed packages:\n"
        "  • numpy\n"
        "  • pandas\n"
        "  • requests\n"
        "  • scipy\n"
        "  • scikit-learn\n"
        "  • simpy\n"
        "  • tabulate\n"
        "  • beautifulsoup4\n"
        "  • yfinance"
    ),
    "parameters": {
        "type": "object",
        "properties": {
            "code": {
                "type": "string",
                "description": "Python code to run.",
                "minLength": 1,
                "maxLength": 5000
            }
        },
        "required": ["code"]
    }
}

```

Tool Description: web_search

```

{
    "type": "function",
    "function": {
        "name": "web_search",
        "description": (
            "Perform a Google search and return relevant results. "
            "Useful for finding current information, news, or facts
            ↪ about topics."
        ),
        "parameters": {
            "type": "object",
            "properties": {
                "query": {
                    "type": "string",
                    "description": "The search query to look up"
                },
                "num_results": {
                    "type": "integer",
                    "description": "Number of results to return
                    ↪ (1-10)",
                    "default": 5
                }
            },
            "required": ["query"]
        }
    }
}

```

Tool Description: browser_get_page_text

```
{
  "type": "function",
  "function": {
    "name": "browser_get_page_text",
    "description": (
      "Fetch a web page and extract its text content. "
      "Useful for reading articles, documentation, or any web"
      "↪ page content."
    ),
    "parameters": {
      "type": "object",
      "properties": {
        "url": {
          "type": "string",
          "description": "The URL of the web page to"
            "↪ fetch"
        }
      },
      "required": ["url"]
    }
  }
}
```

Tool Description: historical_weather

```
{
  "type": "function",
  "function": {
    "name": "historical_weather",
    "description": (
      "Get historical weather data for a specific location"
      "↪ and date. "
      "Useful for analyzing past weather patterns or events."
    ),
    "parameters": {
      "type": "object",
      "properties": {
        "location": {
          "type": "string",
          "description": "City name or coordinates (e.g.,"
            "↪ 'New York, NY' or '40.7128,-74.0060')"
        },
        "date": {
          "type": "string",
          "description": "Date in YYYY-MM-DD format"
        }
      },
      "required": ["location", "date"]
    }
  }
}
```

Tool Description: calculator

```
{
  "type": "function",
  "function": {
    "name": "calculator",
```

```

    "description": "A calculator tool that can perform basic
    ↪ arithmetic operations including +, -, *, /, %, ^, sqrt,
    ↪ sin, cos, tan, log, exp, and parentheses.",
    "parameters": {
      "type": "object",
      "properties": {
        "expression": {"type": "string", "description":
          ↪ "The expression to evaluate, e.g. \"2 * 3.14 *
          ↪ 5\"."}
      },
      "required": ["expression"]
    }
  }
}

```

E MORE BENCHMARK EXAMPLES OF VISUALTOOLBENCH

In this section, we provide more demonstration examples of VISUALTOOLBENCH. Figure 8 and Figure 9 provide two more examples with model responses and grading process.

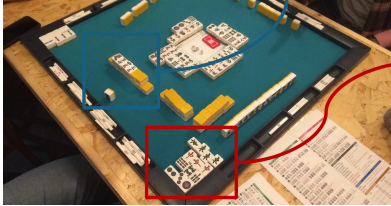
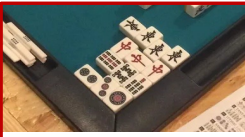
Prompt, Image, and Ground Truth		Key Visual Details		Response Evaluation		
Prompt: I am playing Japanese mahjong and sitting opposite to the current dealer now. This is the east wind. I am tenpai now and just waiting for Papin (8 pin). Please list all the yaku I have right now. If I win this hand by Tsumo, what is the han and fu of my hand and how many points can I get from each player? 		 				
Ground Truth: According to the image, the user has the following yaku currently: Red Dragon (a triple of red dragon) ... If the user wins this hand by Tsumo, this hand is a Mangan with 4 Han 50 Fu, here are the breakdowns ... the user can get 4000 points from the dealer, and 2000 points each from the rest two players. Hence in total 8000 points.		Model Responses via Tool-Use  o3 Tool-Uses: python image process Final Response: From what can be seen in the picture, your concealed hand (menzen) is 1p 1p 3p 4p 5p(red) 6p 7p 6s 7s 8s 3m 4m 5m + drawn-tile 8p ... So you receive 1000 + 500 + 500 = 2000 points in total (ignoring any riichi sticks or honba that might be on the table)				
		Rubric	Weight	Grade		
		Rubric 1 (Critical Rubric) The model gets the correct yaku list for the user's hand.	5	No (+0)		
		Rubric 2 The model gets the dora is papin based on the image.	3	No (+0)		
		Rubric 3 (Critical Rubric) The model gets the correct Han and Fu if the user win by Tsumo, and gets the correct points to collect from each player (in total 8000 points).	5	No (+0)		
		Rubric 4 The model's response in a well-structured, logically ordered format.	2	Yes (+2)		
		Rubric Score: $2/(5+3+5+2) = 0.13$ Pass/Fail: Fail (critical rubric fails)				

Figure 8: Demonstration example from VISUALTOOLBENCH (single-turn, generalist domain, region switch Q&A).

Below, we present additional benchmark examples covering both single-turn and multi-turn settings across all five task categories. Each example includes the task prompt, image input, golden answer, reference tool trajectory, and evaluation rubrics.

Single-Turn, Generalist, Hybrid-Reasoning

Prompt: Check out the menu from this new coffee shop I'm meeting a date at. I want to have a hot coffee ready for him when he gets here. Give me a suggestion for someone who doesn't like steamed milk or sweet-tasting drinks. I don't want it to be the cheapest type of coffee on the menu.

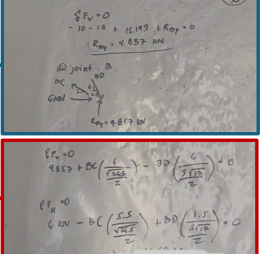
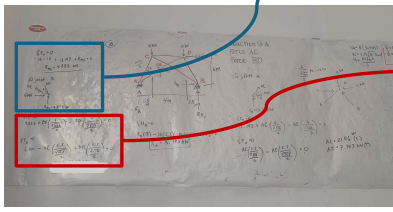
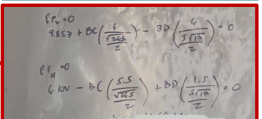
Prompt, Image, and Ground Truth	Key Visual Details	Response Evaluation		
Prompt: What are the equations referring to $\sum(F_v)=0$ and $\sum(F_H)=0$ at the left part in the image? If I only consider the equation for $\sum(F_v)=0$ in the left part of the image and let BD be the weight of 1 m³ hydrated loose lime according to Table 204-1 of NSCP 2015, what will be the value of BC? Do not include any units in the final answer and don't state if it's compression or tension."		Rubric 1 (Critical Rubric) The response determines the equations referring to ... at the left part in the image as ... respectively.	5	Yes (+5)
		Rubric 2 The response mentions that the weight of 1 m³ hydrated loose lime according to Table 204-1 of NSCP 2015 is 5 kN.	3	Yes (+3)
		Rubric 3 The response states the equation as ... by using BD=5.	3	No (+0)
		Rubric 4 (Critical Rubric) The response concludes that the ... approximately -0.0088.	5	Yes (+5)
		Rubric 5 The response must report the final answer with no unit ...	3	No (+0)
		Rubric 6 The response should be concise and ...	2	No (+0)
Ground Truth: Equilibrium Equations: Vertical forces ($\sum F_v = 0$): $4.857 + BC \times (6 / (\sqrt{265 / 2})) - BD \times (6 / (3\sqrt{17 / 2})) = 0$. Horizontal forces ($\sum F_H = 0$): $6 - BC \times (5.5 / (\sqrt{265 / 2})) + BD \times (1.5 / (3\sqrt{17 / 2})) = 0$. According to Table 204-1 of NSCP 2015, the weight of 1 m³ of hydrated loose lime is 5 kN, so: $BD = 5$. Substituting into the vertical equation: $4.857 + BC \times (6 / (\sqrt{265 / 2})) - 5 \times (6 / (3\sqrt{17 / 2})) = 0$. Solving for BC: $BC \approx -0.0088$.	Model Responses via Tool-Use ◆ Gemini-2.5-pro Tool-Uses: python image process Final Response: The equations referring to $\sum(F_v)=0$ and $\sum(F_H)=0$ are the following, respectively ... Based on Table 204-1 of NSCP 2015, the weight of 1 m³ ... $BC=0.072$.	Rubric Score: $(5+3+5)/(5+3+5+3+2) = 0.62$ Pass/Fail: Pass (all critical rubric are met)		

Figure 9: Demonstration example from VISUALTOOLBENCH (single-turn, STEM domain, region switch Q&A).



Golden Answer: That's a great way to start your date! Looking at the menu, an Americano would meet your requirements. It doesn't contain steamed milk, and it's not sweet. At \$4, it's not the cheapest item on the menu, either. Let me know if you need any other suggestions!

Reference Tool Trajectory:

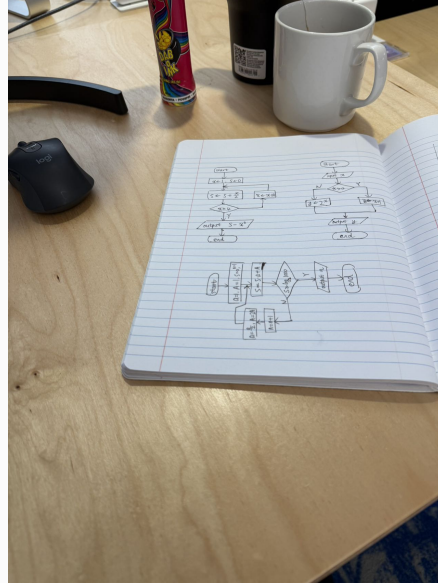
1. `python_image_processing`: Load image to enhance it by flipping horizontally.
2. `python_image_processing`: Crop the image for items that qualify as coffee and are not the cheapest on the menu.
3. `web_search`: Perform a web search to identify which types of coffee contain steamed milk and are sweet-tasting.

Rubrics:

Description	Weight	Type	Category	Critical
Suggests Americano coffee as the beverage.	5	Objective	Reasoning, Instruction following, Truthfulness, Visual Understanding	Yes
States that the Americano is not sweet and does not contain milk.	5	Objective	Instruction following, Truthfulness	Yes
States that the Americano is not the cheapest coffee on the menu	4	Objective	Visual Understanding	No

Single-Turn, STEM, Region Switch Q&A

Prompt: For the top-left flowchart, what operation updates 'S' inside the loop, and what condition terminates it? Also, for the top-right flowchart, what are the two possible calculations for 'y', and what condition decides which one is performed?



Golden Answer: The flowchart on the top left defines a loop. The process inside this loop updates the variable 'S' with the specific mathematical operation $S \leftarrow S + x/2$. The loop is controlled by the diamond shape, which sets the exact termination condition to $x \geq 4$. The flowchart on the top-right describes a conditional process based on the input 'x'. The diamond shape, $x > 0$, is the condition that determines which path is taken. If the condition is true, the calculation performed is $y \leftarrow x + 1$; if it is false, the alternate calculation performed is $y \leftarrow 2^x$.

Reference Tool Trajectory:

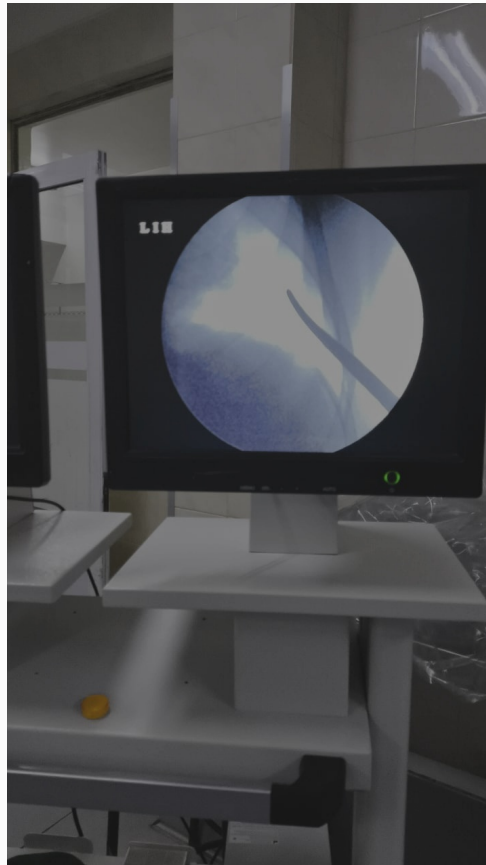
1. python_image_processing: Crop the top-left flowchart and read the text inside the loop to identify the iterative operation on the 'S' variable ($S \leftarrow S + x/2$) and the text in the diamond shape to identify the loop's termination condition ($x \geq 4$).
2. python_image_processing: Crop the the top-right flowchart and read the text in the diamond to find the deciding condition ($x > 0$) and trace its "Yes" and "No" paths to find the two possible calculations for the output 'y' ($y \leftarrow x + 1$ and $y \leftarrow 2^x$).

Description	Weight	Type	Category	Critical
The model identifies the main road the route follows as Kimball Avenue.	5	Objective	Reasoning, Instruction following, Truthfulness, Visual Understanding	Yes
The response states that the user should be picked up at 7:45 AM to catch the bus arriving at Crossroads Mall entrance by 8:30 AM.	5	Objective	Reasoning, Instruction following, Truthfulness, Visual Understanding	Yes

Multi-Turn, Medical, Temporal Sequential Reasoning

Turn 1

Prompt: The medical team is treating a young patient admitted to the emergency room with a diagnosis of pneumothorax after a MVA. During surgery, a chest drain is placed under fluoroscopic guidance. Based on this image, describe the correct trajectory and placement of the instruments. Also, describe the criteria for confirming the safety of the procedure, as well as the possible complications of improper drain placement.



Golden Answer: The image shows the introduction of scissors into the thoracic cavity of a patient with pneumothorax. The instrument's path is carefully directed toward the pleural area, taking great care to avoid direct contact with central structures such as the heart. The

orientation is appropriate, as it follows a lateral and slightly superior path, allowing safe access to the pleural space. Correct placement and safety criteria:

- The instrument is directed toward the pleural cavity and not the mediastinum.
- It is kept away from the cardiac silhouette and large vessels.
- The progressive opening of the tract facilitates subsequent placement of the drainage tube.

Possible complications of incorrect placement:

1. Cardiac or large vessel injury (cardiac tamponade, massive hemothorax).
2. Lung injury with subcutaneous emphysema or hemothorax.
3. Ineffective extrapleural placement.

In conclusion, the image shows a technically correct procedure in the scissors dissection phase, maintaining safety by avoiding vital structures before definitive placement of the chest drain.

Reference Tool Trajectory:

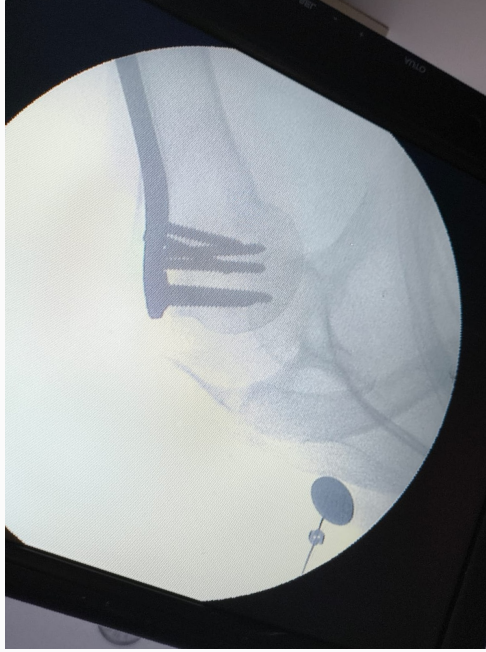
1. `python_image_processing`: Apply a brightness and contrast adjustment to clearly display the image.
2. `python_image_processing`: Apply a crop operation to extract only the area of the monitor containing the image.
3. `web_image_processing`: Search for information about the pneumothorax intervention procedure and its potential complications from improper drain placement.

Rubrics:

Description	Weight	Type	Category	Critical
Identify that the chest drain orientation follows a lateral and slightly superior path to safely access the pleural space.	5	Objective	Instruction following, Truthfulness, Visual Understanding	Yes
Identify safety criteria: instrument trajectory into the pleural cavity while avoiding the cardiac silhouette and large vessels.	5	Objective	Reasoning, Instruction following, Truthfulness	Yes
Identify possible complications of improper drain placement, including cardiac or vessel injury, lung injury, and ineffective extrapleural placement.	4	Objective	Reasoning, Instruction following, Truthfulness	Yes

Turn 2

Prompt: The patient also has a fracture due to the MVA, which was treated along with the pneumothorax. The patient returns for routine follow-up after surgical fixation of a fracture. The radiograph shows the surgical site with the internal fixation hardware in place. Identify the position of the fracture and the hardware holding it. Also, identify the progress of bone healing and if there are any potential complications, such as infection, loosening, hardware migration, or malunion.



Golden Answer: The follow-up radiograph of the proximal humerus shows internal fixation with a plate and multiple screws appropriately placed in the fracture. The hardware is intact, with no signs of loosening, migration, or fracture. The fracture is still slightly appreciable but shows early union with callus formation, consistent with a normal healing process at this stage. There is no evidence of infection, osteolysis, or malunion. Overall, the postoperative outcome of the proximal humerus fracture is favourable, with correct alignment and satisfactory progression.

Reference Tool Trajectory:

1. `python_image_processing`: Apply a flip operation to flip it vertically and extract it in the correct orientation.
2. `python_image_processing`: Apply a rotation operation using the estimated coordinates to extract only the laptop screen area with the imaging findings.

Rubrics:

Description	Weight	Type	Category	Critical
Identifies absence of infection and other complications.	5	Objective	reasoning, truthfulness, visual understanding	Yes
Identify that the proximal humerus fracture is treated with internal fixation hardware (a plate and multiple screws) appropriately placed.	5	Objective	visual understanding	Yes
Identifies that callus is forming at the fracture site.	4	Objective	Reasoning, Visual Understanding	Yes
Identify that this is a shoulder radiograph.	2	Objective	Visual Understanding	No
Identifies that the radiograph image quality is sufficient for evaluation	1	Objective	Visual Understanding	No

Multi-Turn, Generalist, Follow-up Test

Turn 1

Prompt: Shown is a map of multiple train lines throughout a city, with each station/stop marked by a white dot (with intersecting lines indicating where one can transfer lines). Suppose I start at the northmost station of Line Seven; then, without using external sources, what is the least number of required stops to reach the southmost station of Line One? Give all the reasoning necessary to solve the question.



Golden Answer: I am sorry, but it appears the numbers corresponding to each train line are not listed in the image given. To determine both the northmost station of Line Seven, as well as the southmost station of Line One, the line number corresponding to each color (or at least the colors associated with Lines One and Seven) must be given. Can you please provide this information?

Reference Tool Trajectory:

1. `python_image_processing`: Flip the image horizontally, and increase brightness and contrast as needed.
2. `python_image_processing`: Locate the train line map on the screen of the computer in the upper/middle portion of the image. Identify that the line numbers are not listed/are missing from the picture.

Rubrics:

Description	Weight	Type	Category	Critical
The model identifies that the line numbers corresponding to each color are not provided in the image and are necessary to solve the problem.	5	Objective	Visual Understanding	Yes
The model asks a follow-up clarification question requesting the mapping between line numbers and their colors for Lines One and Seven.	4	Objective	Instruction following	Yes

=====

Turn 2

Prompt: That's my bad, I forgot to provide you with the line numbers on the map. Can you now please solve the same question? Shown is a map of multiple train lines throughout a city, with each station/stop marked by a white dot (with intersecting lines indicating where one can transfer lines). Suppose I start at the northmost station of Line Seven; then, without using external sources, what is

the least number of required stops to reach the southmost station of Line One? Give all the reasoning necessary to solve the question.



Golden Answer: Using the compass at bottom right (north is up), identify Line 1 (yellow) and Line 7 (light green). The northernmost Line 7 station is Malvern Centre and the southernmost Line 1 station is Union. The minimum-stop route is: Malvern Centre → Neilson → west via Washburn → Sheppard McCowan → Kennedy → south to Warden → Pape → Garrard → Queen → Union, avoiding longer branches via Brenyon or Ionview. Counting stops on this path (excluding the starting station) gives 25.

Reference Tool Trajectory:

1. `python_image_processing`: Flip the image horizontally, and increase brightness and contrast as needed.
2. `python_image_processing`: Locate the train line map on the computer screen in the middle/upper portion of the image. Further, zoom into the upper portion of the screen to identify the train line color pairing as: light green (Line Seven) and yellow (Line One). Next, locate the compass symbol near the bottom right portion of the screen, indicating that the top of the image is the northmost portion. Finally, identify all possible station stops (marked as white dots) between the northmost station of Line Seven (Malvern Centre) and the southmost station of Line One (Union).

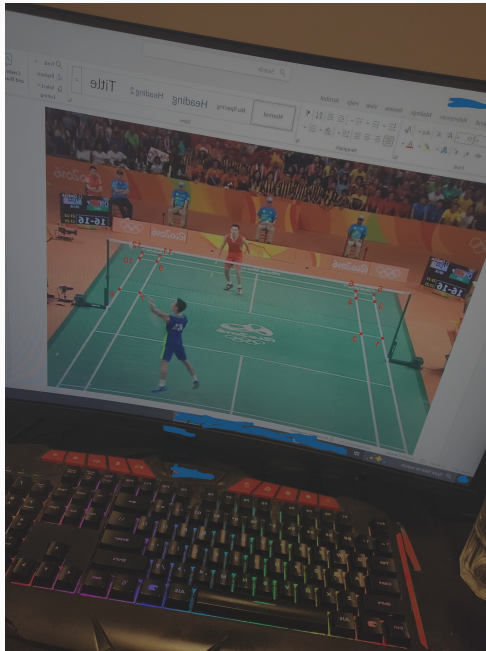
Rubrics:

Description	Weight	Type	Category	Critical
Model correctly identifies the northmost station of Line Seven as Malvern Centre.	5	Objective	Reasoning, Truthfulness, Visual Understanding	Yes
Model correctly identifies the southmost station of Line One as Union Station.	5	Objective	Reasoning, Instruction following, Truthfulness, Visual Understanding	Yes
The model states that the minimum number of stops required when travelling from Malvern Centre to Union Station is 25.	4	Objective	Reasoning, Instruction following, Truthfulness, Visual Understanding	Yes
The model outlines a path listing stations from Malvern Centre to Union Station that follows the specified route, allowing equivalent station names.	3	Objective	Reasoning, Instruction following, Truthfulness, Visual Understanding	No

Multi-Turn, Sports, Progressive Visual Reasoning

Turn 1

Prompt: A badminton tournament during the 2016 Rio Olympics between Lin Dan (in red) and Lee Chong Wei (in blue) is shown, which is a best-of-three-game series. The middle of a round is currently in play. Who is currently hitting/about to hit the shuttlecock, and how do you know? Further, if the shuttlecock is missed by the opponent after being hit by the player from the answer in the first question, what numbered point on the ground labelled in red will be the farthest point away from the receiver (player who missed), such that the shuttlecock is not considered out?



Golden Answer: Who is **hitting/about to hit the shuttlecock**? Since the player in blue is jumping, this is a clear indication that he is currently hitting/about to hit the shuttlecock. Further, if one looks more carefully, a white shuttlecock can be identified near the top right of the blue player's racket. Thus, the blue player (Lee Chong Wei) is currently hitting/about to hit the shuttlecock.

What point labelled in red is farthest away from the receiving player if he misses the shot made by Lee Chong Wei, not considered out? Since the game shown is between two players (singles), any shots where the shuttlecock lands in the farthest left or farthest right rectangles are considered out (everything else is considered in, including the backmost rectangle). Thus, this eliminates the points numbered 1, 3, 5, 8, 10, and 12 as possible answers. Of the remaining points, the point numbered 6 is furthest away from Lin Dan (red player). Thus, the point labelled 6 is the farthest point away from the receiving player (Lin Dan) considered in, assuming he misses the shot made by Lee Chong Wei.

Reference Tool Trajectory:

1. `python-image-processing`: Flip the image horizontally.
2. `python-image-processing`: Adjust the brightness and contrast of the image.

Rubrics:

Description	Weight	Type	Category	Critical
The model correctly identifies Lee Chong Wei (in blue) as the player currently hitting or about to hit the shuttlecock.	5	Objective	Instruction following, Truthfulness, Visual Understanding, Reasoning	Yes
The model justifies its identification by referencing either the player's posture/position or the visible shuttlecock location near the racket.	4	Objective	Visual Understanding, Reasoning	Yes
The model must reason that the game shown is a singles match between two players.	3	Objective	Reasoning	No
The model must state that the farthest most left and farthest most right rectangles are considered out if the shuttlecock lands there.	3	Objective	Reasoning	No
Identify the 12 points labelled in red (1 through 12) on Lin Dan's side of the court.	1	Objective	instruction following, Visual Understanding	No
State that the shuttlecock is out if it lands in any of the labelled points 1, 3, 5, 8, 10, or 12.	2	Objective	instruction following, reasoning, truthfulness	No
State that point number 6 is the farthest allowed point from Lin Dan still considered in..	5	Objective	Instruction following, Truthfulness, Visual Understanding, Reasoning	Yes
Response includes explicit reasoning rather than only providing a ground truth answer.	3	Objective	Reasoning, Presentation	No

Turn 2

Prompt: Suppose now, instead of hitting the shuttlecock to the point labelled 6 on the ground, Lee Chong Wei instead hits the shuttlecock to the point labelled 8 on the ground. Assuming Lin Dan does not receive (hit) the shuttlecock and lets it hit the ground, what will the scoreboard read after this point? State the player's name associated with the number of points as well.

Golden Answer: Current score according to the scoreboard: From the scoreboard, it reads 16-16, where the left number is the score associated with the team from China (from the flag shown), whereas the right number is the score associated with the team from Malaysia (from the flag shown).

New score, assuming the shuttlecock lands at the point labelled 8: Since the game is a singles match, the outer left and outer right-most rectangles are considered out. Since the point labelled 8 is in the outer right-most rectangle, this shot made by Lee Chong Wei would be considered out. Since Lin Dan did not receive the shot as stated by the prompt, Lin Dan will win this round, increasing the score to 17-16. (Here, it is important to note that the left number on the scoreboard has increased, and not the right). Further, Lin Dan will have 17 points, while Lee Chong Wei will still have 16 points.

Reference Tool Trajectory:

1. `python.image_processing`: Identify the current match score as 16-16 (middle top right of the image), where the left number is associated with the score of the team from China, whereas the right number is associated with the score of the team from Malaysia, according to the flags.

2. `web_search`: Look up and determine that Lin Dan represents the team from China, whereas Lee Chong Wei represents the team from Malaysia.

Rubrics:

Description	Weight	Type	Category	Critical
Identify the current score as 16-16.	3	Objective	Truthfulness, Visual Understanding	Yes
Associate the left number with the team from China and the right number with the team from Malaysia based on the flags.	1	Objective	Truthfulness, Visual Understanding	No
Model correctly assigns Lin Dan to the left scoreboard number (team from China) and Lee Chong Wei to the right scoreboard number (team from Malaysia).	2	Objective	Reasoning, Visual Understanding	No
Model correctly reasons that because the shuttlecock landing point labelled 8 is out in singles matches and Lin Dan does not hit it back, Lin Dan receives the point	5	Objective	Reasoning, Visual Understanding	Yes
The model states the new score after Lin Dan receives the point as: 17-16 (17 for Lin Dan, and 16 for Lee Chong Wei).	5	Objective	Reasoning	Yes
The model states the new score after Lin Dan receives the point as: 17-16 (17 for Lin Dan, and 16 for Lee Chong Wei).	5	Objective	Truthfulness, Instruction Following	Yes
The reasoning explains why point 8 is considered out.	4	Objective	Reasoning, Presentation	Yes
The reasoning explains why the left number on the scoreboard is associated with Lin Dan and not Lee Chong Wei.	1	Objective	Reasoning, Presentation	No

Turn 3

Prompt: Now, since the tournament is a best-of-three-game series, if Lee Chong Wei in fact does hit the shuttlecock out at the point labelled 8 on the ground as discussed, how many additional points will he need to win the match (assuming the score does not tie at 20-20)? If Lee Chong Wei wins this match/game, is he guaranteed to win the series? If so, what were the scores of the previous two matches played against each other?

Golden Answer: How many points does Lee Chong Wei need to win the match after shooting the shuttlecock out? The new score after Lee Chong Wei hits the shuttlecock out is 17-16 (17 for Lin Dan, 16 for Lee Chong Wei). Since a standard badminton game is up to 21 points, and assuming the score does not tie to 20-20 (in which case you can only win if you get two points back-to-back), Lee Chong Wei thus needs $21-16=5$ points more to win.

Scores of the previous two matches: One can deduce that the previous two games have already been finished from the scoreboard, where it states that Lin Dan won one game with a final score of 21-15, whereas Lee Chong Wei won the other game with a score of 11-21. Thus, the game being played currently is the final game to determine the winner of the three-game series. Hence, if Lee Chong Wei wins this match, he is guaranteed to win the series.

Reference Tool Trajectory:

1. `python_image_processing`: Identify the scores of the previous two matches/games in the middle top right of the image as: 21-15 (for Lin Dan), and 11-21 (for Lee Chong Wei).
2. `web_search`: Lookup the amount of points needed to win a standard badminton game as 21 (assuming a 20-20 point tie does not occur).

Rubrics:

Description	Weight	Type	Category	Critical
Identifies the final scores of the previous two matches as 21-15 (Lin Dan) and 11-21 (Lee Chong Wei).	3	Objective	Instruction following, Truthfulness, Visual Understanding	No
The model reasons that Lee Chong Wei must reach a total point count of 21 to win the match (assuming a 20-20 tie does not occur).	4	Objective	Reasoning	Yes
The model correctly calculates that Lee Chong Wei needs 5 additional points to win the current match.	5	Objective	Reasoning, Instruction following, Truthfulness	Yes
The model correctly states that if Lee Chong Wei wins this match, he is guaranteed to win the best-of-three series and provides the scores of the previous two matches.	5	Objective	Reasoning, Presentation	Yes
The response provides intermediate reasoning that addresses each question posed by the prompt: (1) how many additional points Lee Chong Wei needs to win the match, and (2) whether he is guaranteed to win the series and what the scores of the previous two matches were.	2	Objective	Reasoning, Presentation	No