

MindBlock: Probing Spatial Assembly and Structure in Unified Multimodal Models

Baiqiao Yin^{1*}, Junhao Liu^{1*}, Han Yin^{1*}, Heyang Yu¹, Tingxuan Zhang¹,
Zhiheng Li², Chengzu Li³, Jihan Yang¹, Manling Li^{4†}, Chen Feng^{1†}, and
Yiming Li^{1†}

¹ New York University, Brooklyn, NY 11201, USA

² Carnegie Mellon University, Pittsburgh, PA 15213, USA

³ University of Cambridge, Cambridge, UK

⁴ Northwestern University, Evanston, IL 60208, USA

 Dataset: MindBlock-Bench  Code: [yyyybq/MindBlock](https://github.com/yyyybq/MindBlock)

Abstract. While Unified Multimodal Models (UMMs) show remarkable reasoning capabilities, their spatial intelligence remains limited to passive 2D Question-Answering (QA). In this paper, we argue that true spatial intelligence demands active construction: not only recognizing a 3D structure in pixel space, but also building and modifying it. We introduce **MindBlock**, a benchmark that challenges models’ active generative construction in pixel space across two primary axes: **Spatial Assembly**, which evaluates step-by-step compositional and causal reasoning, and **Spatial Structure**, which probes spatial equivariance through local sub-component rotations and global viewpoint transformations. To move beyond pixel-level metrics, we propose **3DGS-Eval**, a novel validation protocol using 3D Gaussian Splatting to reconstruct implicit scenes from model-generated multi-view images. This allows us to quantify structural consistency, verifying for the first time whether a model’s generative output admits a coherent internal 3D world model. Furthermore, we conduct a deep-dive diagnostic analysis into the representational grounding of spatial logic, disentangling whether structural consistency relies on textual Chain-of-Thought (CoT) as a symbolic scaffold, or emerges as a native spatial intuition within the generative latent space. Our findings reveal a significant “perception-execution” gap: while current models correctly identify the intended spatial state, they fail at active construction. They struggle to maintain spatial equivariance without explicit symbolic scaffolding. MindBlock provides a rigorous foundation for the next generation of embodied, physically-grounded multimodal AI.

Keywords: UMMs · Spatial Intelligence · Dataset and Benchmark

1 Introduction

The ability to assemble blocks into meaningful structures unfolds gradually during human development. Toddlers as young as two can stack multiple blocks,

* Equal contribution. † Equal advising.

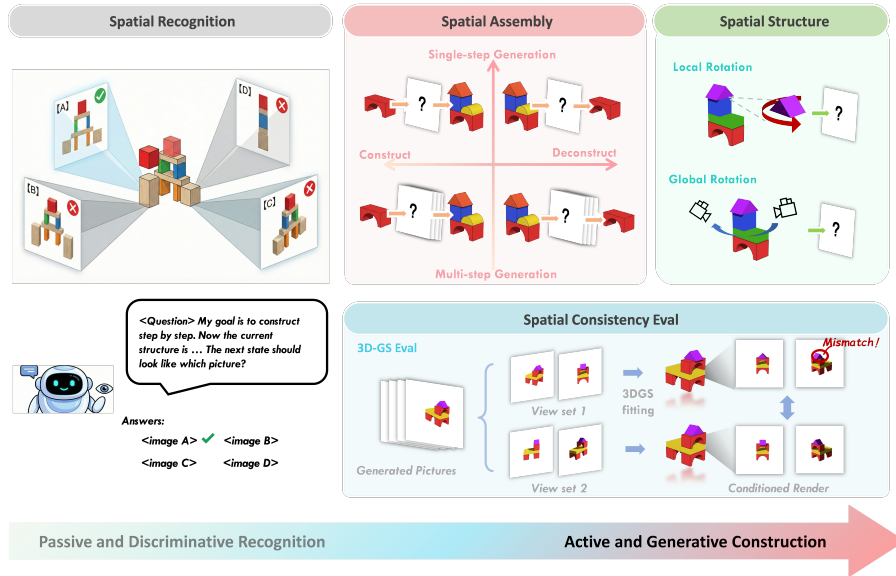


Fig. 1: Children aged four to six begin constructing complex block structures from diagrams, demonstrate awareness of enclosure and decoration, and start to grasp three-dimensional spatial relationships. Can AI agents match this level of spatial intelligence? We transform visual-spatial reasoning from **passive selection (multi-choice question answering)** to **active construction (conditional image generation)**, rigorously examining its internal three-dimensional world modeling.

but purposeful construction typically emerges around age three. By ages four to five, children begin to replicate simple structures from observation and to grasp spatial relationships such as *next to* and *on top of* [19]. The skill of *spatial mental rotation* (imagining how a structure would look from a different viewpoint) develops around ages five to six [9]. This trajectory, from simple stacking to perspective-aware, multi-step assembly, raises a research question:

Can current AI models match the spatial intelligence of a six-year-old?

Recent studies have begun to explore this question by evaluating and training foundation models in block-building tasks. **BrickGPT** [33] models the locations of LEGO bricks using spatial coordinates represented as text, employing large language models (**LLMs**) to generate physically stable structures. However, this approach operates entirely within text space, whereas humans build blocks using visual input. **PhyBlock** [29] evaluates 3D block assembly in vision-language models (**VLMs**) through visual question answering (VQA). Yet multiple-choice question answering can be hacked through statistical shortcuts: correct answers neither guarantee genuine understanding nor translate to robotic assembly. Notably, unified multimodal models (**UMMs**) [7, 44] have recently emerged to enable both multimodal inputs and outputs, typically in the form of images and

text. This makes UMMs uniquely suited for perceiving visual instructions and generating step-by-step image sequences for planning [54].

In this work, we move beyond *passive, discriminative* VQA in spatial intelligence: rather than asking models to *select* the correct answer, we ask them to *build* the correct structure. Specifically, we introduce **MindBlock**, a benchmark that probes spatial intelligence through two complementary axes, see Fig. 1.

- **Spatial Assembly** tests whether models can actively construct block structures step by step. Given an image of a target structure, the model must generate an interleaved sequence of reasoning steps and intermediate images, each showing the scene after one block is placed. This demands *compositional reasoning* (assembling parts into a coherent whole) and *causal reasoning* (a block placed in mid-air will fall; a missing base block invalidates every step above it).
- **Spatial Structure** tests whether models maintain a consistent 3D understanding across viewpoints. Given the same structure, the model must generate what it looks like from a different angle, preserving every block’s relative position and occlusion. We call this *spatial equivariance*: a model that truly understands a structure in 3D should produce a predictably transformed image when the viewpoint rotates.

Standard image-similarity metrics (*e.g.*, SSIM, LPIPS, CLIP-S) cannot capture a fundamental failure mode of modern Unified Multimodal Models (UMMs): a model may produce images that look plausible individually yet imply mutually inconsistent 3D structures. For instance, blocks may appear on different sides of the structure when rendered from different viewpoints, yielding images that are locally convincing but globally incompatible. To directly measure this phenomenon, we introduce **3DGS-Eval**. Given a set of model-generated multi-view images, we reconstruct a scene using 3D Gaussian Splatting (3DGS), a compact volumetric representation, and evaluate the geometric consistency of the reconstruction. If the generated views encode a coherent 3D structure, the reconstruction converges to a stable geometry; otherwise, contradictory views cause the volumetric fit to deteriorate, revealing hidden spatial inconsistencies that standard 2D metrics fail to detect.

Our evaluation experiments reveal that good performance under standard 2D metrics can be misleading. As shown in Table 1, several models achieve high perceptual similarity scores while producing multi-view generations that are geometrically inconsistent. Images that appear plausible in isolation often contradict each other when interpreted as a single 3D structure, exposing a fundamental *2D–3D consistency gap* in current UMMs.

To better understand the origin of these failures, we conduct a series of diagnostic analyses. First, we examine whether spatial reasoning benefits from explicit reasoning by comparing generations with and without Chain-of-Thought (CoT) explanations. CoT improves volumetric consistency under 3DGS-Eval, suggesting that language can act as a symbolic scaffold that stabilizes spatial reasoning during generation. We also extend the evaluation to multi-step construction tasks. While some models succeed in single-step edits, they often fail

to maintain consistency across longer reasoning trajectories, exhibiting a phenomenon we term *spatial drift*, where small geometric errors accumulate over time.

Finally, by comparing generative and discriminative variants of the same tasks, we identify three distinct failure modes: perceptual failures (incorrectly interpreting the input structure), planning failures (producing an incorrect sequence of operations), and execution failures (failing to realize a correct plan in image generation). This analysis reveals a pronounced *perception–execution gap*: models frequently recognize correct spatial configurations more reliably than they can construct them.

Our contributions are three-fold:

- **New Tasks.** We evaluate spatial intelligence through *active construction*: models must generate correct visual structures and reasoning trajectories, rather than merely selecting answers. What a model builds provides a stronger probe of its internal 3D understanding than what it recognizes.
- **New Benchmark.** We introduce **MindBlock**, a spatial reasoning benchmark covering six tasks that span atomic operations (single-step assembly and deletion), multi-step structural reasoning (forward assembly and reverse deconstruction), and viewpoint consistency (local and global equivariance). The benchmark data and code are publicly released. Evaluation combines conventional 2D image similarity with **3DGS-Eval**, a volumetric consistency metric that exposes hidden geometric contradictions across views.
- **New Insights.** Our analysis reveals three systematic limitations of current UMMs: (1) a *2D–3D consistency gap*, where high perceptual similarity does not imply correct geometry; (2) *spatial drift* in multi-step reasoning, where structural errors accumulate over time; and (3) a pronounced *perception–execution gap*, where models recognize correct spatial configurations more reliably than they can generate them. Together, these findings suggest that while UMMs exhibit promising spatial intuition, their internal representations remain fragile and insufficient for robust 3D reasoning.

2 Related Work

Unified Multimodal Models. Recent advances in vision-language understanding and visual synthesis have driven a trend toward unifying multimodal perception and generation in a single, parameter-shared architecture, known as Unified Multimodal Models (UMMs). Depending on their core modality-fusion backbone and state-space discretization strategy, existing UMMs can be broadly classified into three distinct paradigms: pure autoregressive models that treat all modalities as discrete sequential tokens [6], unified diffusion-based models that natively handle mixed-modality flows through parallel refinement [22, 35, 48], and hybrid frameworks that strategically fuse autoregressive reasoning with continuous or discrete diffusion processes [42, 45]. By dissolving the boundary between understanding and synthesis, UMMs handle arbitrarily interleaved text-image sequences [21, 24, 55], making them well suited to the problems studied here.

Specifically, *we investigate whether UMMs possess genuine 3D spatial intelligence by measuring their capacity for active spatial assembly, spatial equivariance across viewpoints, and their ability to bridge the critical perception-execution gap between discriminative understanding and generative construction.*

Spatial Reasoning in MLLMs. Whether artificial neural networks can exhibit human-like spatial intelligence has long been a central research question [13]. Most prior benchmarks cast spatial understanding as *passive, discriminative* VQA-style problems, focusing on object attributes or pairwise relations while placing limited emphasis on structural reasoning [8, 17, 30, 31, 40, 41, 47, 49–51]. More recent efforts introduce structure-aware tasks, yet still rely on textual or multiple-choice outputs, which can be bypassed via statistical shortcuts and limit the expressiveness of spatial representations [20, 53]. With the emergence of unified multimodal models (UMMs), researchers have shown that generative capabilities can encourage models to internalize and manipulate spatial structures [14]. Building on this insight, our work fundamentally shifts the paradigm from *passive selection* to *active construction*. *We propose a spatial reasoning framework that demands models to generate coherent, step-by-step spatial assemblies and maintain 3D spatial equivariance across viewpoints.*

Robotic Assembly. Robot assembly focuses on assembling meaningful structures in compliance with a given reference image, instructions, or 3D model [18]. The core of this task resides in both planning and execution. In terms of planning, agents need to possess an internal 3D world model. Some works neglect such modeling, instead formulating the problem as an optimization task with physical constraints [28, 32]. Some works utilize graphs to model dependencies, yet fail to capture precise positional information [29, 37]. Methods such as Bricks-GPT propose text-based frameworks, but they are constrained by components and inter-block relationships [25, 33]. Other works take visual modality as input and output assembly plans, attempting to leverage visual information [36, 39, 52]. Yet, such tasks rely primarily on passive question-answering and lack 3D-based evaluation. Existing literature highlights the lack of precise 3D world modeling for the assembly process, which is our main focus. In terms of execution, learning-based methods have become mainstream. Nevertheless, these methods still encounter challenges in complex goal-conditioned tasks [4, 11]. Adopting a slow-fast system has become a standard solution. [16]. Another viable solution is to integrate multimodal generation capabilities, such as rewards, images, videos, and visual annotations, as supervisory signals [1, 2, 23, 27, 56]. This reveals the need for powerful 3D world modeling in planning and generation. *Building on a unified multimodal model, we analyze assembly reasoning and the spatial consistency of generated images, providing a benchmark for downstream tasks.*

3 MindBlock: Probing Spatial Assembly and Structure

The MindBlock benchmark is designed to evaluate two fundamental dimensions of spatial intelligence: the capacity for *assembly*—the goal-directed construc-

tion and deconstruction of physical entities—and the mastery of *structure*—the geometric understanding of 3D space.

3.1 Evaluating Spatial Assembly

We define spatial assembly as the model’s capacity to execute precise state transitions within a 3D coordinate frame. This is evaluated through two complementary single-step tasks and a long-horizon loop.

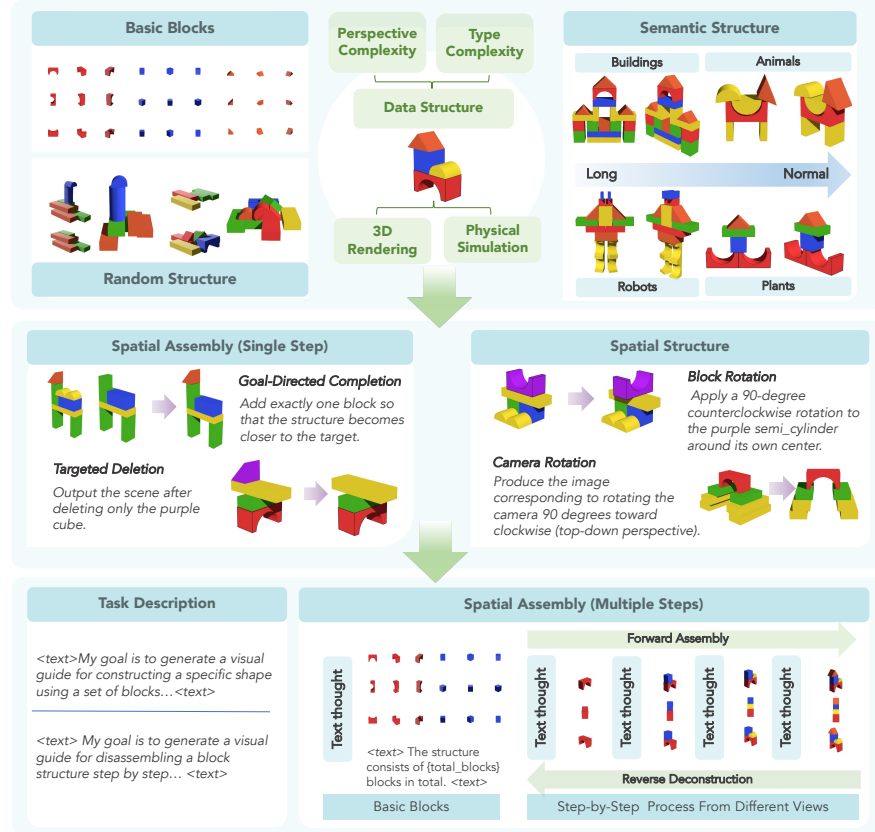


Fig. 2: Overview of the MindBlock benchmark and dataset. The dataset is constructed from basic blocks via a 3D rendering and physical simulation pipeline (top). The benchmark evaluates spatial intelligence through two primary dimensions: **Spatial Assembly** (left and bottom), covering single-step and long-horizon construction/deconstruction, and **Spatial Structure** (right), testing geometric equivariance through local and global transformations.

Single-Step Assembly We focus on the precision of atomic manipulations through two distinct paradigms:

- *Atomic Deletion: Targeted Deletion.* This task requires the model to precisely remove a specific target entity (e.g., “delete only the purple cube”) from a given scene. It serves as a rigorous test for *structural persistence* and *occlusion recovery*: the model must preserve the exact coordinates of all remaining elements and the camera perspective while hallucinating the previously occluded background. This ensures the model treats the structure as a collection of independent 3D entities rather than a monolithic image.
- *Atomic Completion: Goal-Directed Completion.* Given a partial assembly and a target state image, the model must autonomously infer and generate the single most critical next step. This evaluates the model’s proactive planning and its ability to bridge the gap between current and desired physical configurations.

Multi-Step Assembly-Deconstruction Loop To evaluate long-horizon consistency, we extend the aforementioned mappings into a bi-directional loop:

- *Iterative Completion: Forward Assembly.* The iterative execution of atomic completions, where the model must progressively build toward a target configuration across multiple steps.
- *Iterative Deletion: Reverse Deconstruction.* The continuous application of atomic deletions, requiring the model to sequentially dismantle a structure while maintaining spatial coherence at each step.

Unlike single-step tasks that probe atomic precision, the multi-step loop challenges the model to manage *structural evolution*. As the sequence progresses, the model must navigate dynamic occlusion and maintain the integrity of hierarchical support. Deconstruction further requires reasoning about physical causality, e.g., identifying load-bearing blocks before removal providing a stringent test of the internal world model.

3.2 Evaluating Spatial Structure

To probe the model’s internal 3D world model, we test for *geometric equivariance*—the principle that spatial representations should transform predictably under geometric operations.

Local Equivariance: Local Structure Manipulation: We task the model with rotating a specific sub-component (e.g., “Rotate the top triangular prism 90 degrees clockwise”) while preserving the rest of the structure. This evaluates the precision of the model’s local coordinate system and its ability to disentangle individual parts within a structured whole.

Global Equivariance: Global Structure Consistency: We assess the model’s ability to generate a globally transformed view (e.g., a 180-degree rotation) of the entire scene. This serves as a test for global 3D consistency; a model possessing a true structural understanding must maintain the relative positions and occlusions of all blocks across arbitrary viewpoint shifts, a capability we quantitatively verify using our **3DGS-Eval** protocol.

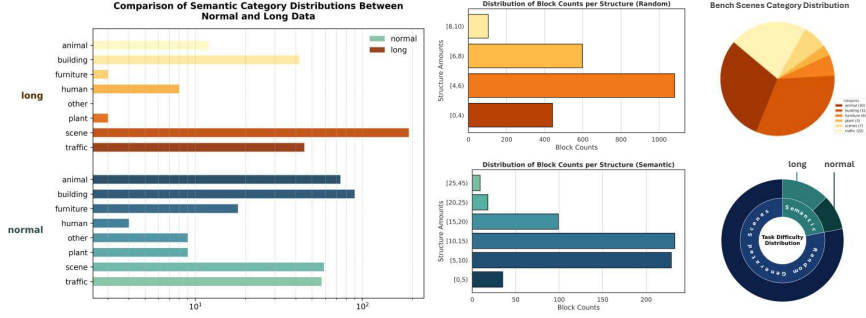


Fig. 3: Statistics of the synthesized dataset. (Left) Comparison of semantic category distributions between "Normal" and "Long" data subsets. (Middle) Distribution of block counts (assembly steps) for Random and Semantic structures, illustrating the complexity gradient. (Right) Category distribution of bench scenes and the overall task difficulty ratio (Normal vs. Long) across Semantic and Random subsets.

4 Dataset Construction and Evaluation

MindBlock bridges the gap between 2D perception and 3D execution by grounding structural assembly in a rigorous physical and geometric protocol.

4.1 Physically-Grounded Data Synthesis

We utilize the **Genesis** physics engine [10] to generate a large-scale dataset of stable brick structures with precise control over geometry, lighting conditions, and camera placement. This simulation-based pipeline enables physically consistent rendering and fine-grained supervision signals for spatial reasoning tasks. Our structure library spans two complementary categories: (i) *semantic structures* (e.g., furniture, vehicles, animals), which leverage common-sense priors about real-world objects; (ii) *random structures*, which remove semantic cues to emphasize pure geometric reasoning and compositional spatial understanding. The overall distribution of categories and task complexity is illustrated in Fig. 3.

The construction of our dataset follows a rigorous simulation-in-the-loop pipeline. For the random structures, we employ a constrained random generation process combined with physical simulation to ensure the structural integrity of each step. For the semantic structures, we utilize a hybrid approach of manual modeling and physical simulation, encompassing a diverse set of scenes including buildings, furniture, and vehicles. A portion of the semantic models is adapted from PhyBlock [29]. In total, MindBlock comprises **2,839 structures and 16,336 intermediate assembly states**. As shown in the block count distributions (Fig. 3, middle), the complexity of these structures is parameterized by their assembly sequence lengths, spanning from simple 2-block builds to

complex "Long" sequences exceeding 40 blocks, thereby providing a comprehensive benchmark for evaluating model scalability in spatial reasoning. We split the dataset at the *structure level* into **2,661 train** and **188 test** structures, ensuring no leakage during training.

For each structure, we generate a **bottom-up assembly sequence** that simulates the incremental construction process. Each construction trajectory consists of multiple discrete steps, where **exactly one block is added at each step**. This sequential formulation exposes the intermediate structural states of the scene and provides supervision for step-by-step reasoning about stability, placement, and support relationships.

At every step of the construction sequence, we render rich multimodal supervision signals:

- **Multi-view Observations:** We render images from multiple viewpoints covering a full 360-degree perspective. In practice, we sample eight pre-defined egocentric camera viewpoints around the structure, providing the multi-angle observations necessary for volumetric understanding and spatial disambiguation.
- **Structured Reasoning Traces:** We generate template-synthesized Chain-of-Thought (CoT) reasoning traces. These traces explicitly describe the intended assembly rationale, including block attributes, precise (x, y, z) coordinates, and hierarchical support relationships (e.g., "Place block A on top of block B").

4.2 3DGS-Eval: Measuring Volumetric Consistency

To move beyond pixel-level similarity, we introduce **3DGS-Eval**, which employs **3D Gaussian Splatting (3D-GS)** to reconstruct a volumetric field from generated images. Following the pipeline in MVGBench [46], we initialize 3D Gaussians using generated views and their corresponding camera poses. Each primitive—parameterized by position, opacity, anisotropic covariance, and color—is optimized via **differentiable rasterization** to minimize photometric reconstruction error against the target images.

Concretely, we partition the eight generated views into two disjoint subsets and fit an independent 3DGS reconstruction on each. We then render both reconstructions from a shared set of evaluation viewpoints and compare the resulting images, depth maps, and 3D Gaussian geometry. If the generated views encode a coherent 3D structure, the two reconstructions should agree under novel-view rendering and geometry; geometric contradictions instead lead to large discrepancies between the two reconstructions, exposing inconsistencies that may be invisible to purely 2D metrics.

To capture fine geometric details, an **adaptive density control** mechanism iteratively refines the field by splitting, merging, or pruning Gaussians. This process forces the representation to converge into a coherent 3D entity. Consequently, 3DGS-Eval acts as a **"geometric polygraph"**: if the generated views

Table 1: Quantitative Comparison on MindBlock Benchmark. We evaluate the single-step spatial intelligence of state-of-the-art Unified Multimodal Models (UMMs). The results are categorized into *Image Similarity* (2D heuristic) and our proposed *3D Gaussian Splatting* (volumetric consistency). $\uparrow$ denotes higher is better; $\downarrow$ denotes lower is better. Best results are **bolded**, and second best are underlined.

Task Model	Image Similarity 3D Gaussian Splatting (3DGS-Eval)					
	CLIP-S $\uparrow$	PSNR $\uparrow$	SSIM $\uparrow$	Depth Err. $\downarrow$	CD $\downarrow$	
<i>Spatial Assembly (Single-step Construction)</i>						
BAGEL [7]	0.77	13.83	0.73	75.97	11.72	
Blip3-o-next [3]	<u>0.81</u>	11.90	0.69	83.20	13.05	
ILLUME [15]	0.68	9.94	0.64	99.67	18.20	
Emu3.5 [5]	0.75	13.10	0.71	79.20	12.10	
Step1X-Edit [26]	0.86	15.20	0.78	68.50	9.10	
Ovis-U1 [38]	0.76	<u>14.10</u>	<u>0.75</u>	<u>73.10</u>	<u>10.40</u>	
Omni-Gen [43]	0.70	8.70	0.59	108.40	17.80	
<i>Spatial Structure (Geometric Equivariance)</i>						
BAGEL [7]	0.74	12.26	0.76	82.64	10.73	
Blip3-o-next [3]	<u>0.79</u>	10.80	0.70	90.50	12.90	
ILLUME [15]	0.68	8.90	0.64	107.33	18.58	
Emu3.5 [5]	0.72	11.50	0.73	86.70	11.80	
Step1X-Edit [26]	0.83	13.80	0.80	74.20	9.30	
Ovis-U1 [38]	0.73	<u>12.90</u>	<u>0.77</u>	<u>79.60</u>	<u>10.20</u>	
Omni-Gen [43]	0.68	7.90	0.57	118.60	19.30	

lack underlying 3D consistency, the optimization fails to form a stable structure, leading to structural collapse.

Finally, both optimized 3DGS representations are re-rendered from the same evaluation viewpoints. We compute image-level consistency between the two sets of renderings using PSNR and SSIM, compare their rendered depth maps via mean absolute depth error, and measure geometric agreement between the two Gaussian fields using Chamfer distance over sampled Gaussian points. These metrics jointly quantify whether independently reconstructed volumetric fields from disjoint view subsets agree in appearance, depth, and geometry.

5 Experiments

5.1 Evaluation Setup

Models. We evaluate a diverse set of models across three primary categories to benchmark the current landscape of spatial intelligence:

- **General-purpose UMMs:** We include state-of-the-art autoregressive and diffusion-based unified multimodal models, namely **Emu3.5** [5], **Ovis-U1** [38], **ILLUME** [15], **BAGEL** [7] and **Omni-Gen** [43], which are capable of native image-text generation.

- **Specialized and Editing Models:** We assess **Blip3-o-next** [3] and **Step1X-Edit** [26], which are optimized for conditional image manipulation and structural preservation.

Metrics. We employ a dual-pronged evaluation strategy:

- **Image Similarity (2D):** We evaluate semantic alignment between generated and ground truth images using the CLIP Similarity score (*CLIP-S*), computed with a pretrained CLIP encoder.
- **3DGS-Eval (3D):** Our core contribution. From the model-generated multi-view images, we fit two independent 3DGS reconstructions on disjoint view subsets and measure their mutual agreement. *PSNR* and *SSIM* quantify appearance consistency between the two reconstructions’ renderings, mean absolute depth error (*Depth Err.*) quantifies depth consistency, and Chamfer Distance (*CD*) measures geometric agreement between the two Gaussian fields.

Implementation Details. For 3DGS-Eval, we generate 8 views per structure. All baseline models are evaluated using their default inference hyperparameters with a temperature of 0.2.

5.2 Analysis of Single-step Assembly and Spatial Structure

As shown in Table 1, we observe several important findings regarding the spatial reasoning capabilities of current Unified Multimodal Models (UMMs).

The 2D–3D Consistency Gap. A clear discrepancy emerges between 2D perceptual similarity and volumetric consistency. For example, *Blip3-o-next* [3] achieves relatively strong 2D alignment with a CLIP-S score higher than *BAGEL* [7]. However, its 3D reconstruction metrics are substantially worse, exhibiting lower PSNR/SSIM and significantly higher depth and Chamfer errors. This indicates that although individual generated views appear semantically plausible, they fail to form a geometrically coherent 3D scene when aggregated across view-points. In contrast, models such as *Ovis-U1* [38] achieve competitive or even better 3D reconstruction quality despite having comparable or slightly lower CLIP-S scores. These results highlight that strong 2D image similarity does not necessarily imply correct spatial reasoning.

Spatial Assembly Remains Challenging but Partially Solvable. For the *Spatial Assembly* task, most models show moderate performance, with *Step1X-Edit* [26] achieving the best overall results across both 2D and 3D metrics. Models such as *Emu3.5* [5] and *Ovis-U1* [38] remain close to *BAGEL* [7], suggesting that current UMMs can partially reason about single-step object placement. However, large performance gaps remain between stronger models and weaker ones such as *ILLUME* [15] and *Omni-Gen* [43], which fail to produce geometrically stable structures.

Equivariance is a Major Bottleneck. All evaluated models show a consistent performance drop on the *Spatial Structure* task. Compared to *Spatial Assembly*, both 2D and 3D metrics deteriorate across nearly all models, indicating

Table 2: Performance during forward assembly and reverse deconstruction phases. $\uparrow / \downarrow$ indicates direction of improvement.

Model	Image Similarity	3DGS-Eval (3D)			
	CLIP-S $\uparrow$	PSNR $\uparrow$	SSIM $\uparrow$	Depth Err. $\downarrow$	CD $\downarrow$
<i>Phase I: Forward Assembly</i>					
BAGEL-zebra-CoT [7]	0.78	13.5	0.74	71.3	11.8
ThinkMorph [12]	0.76	15.1	0.79	77.4	9.6
MindBlock-Random (Ours)	<u>0.84</u>	<u>17.8</u>	<u>0.86</u>	<u>62.5</u>	<u>7.4</u>
MindBlock-Semantic (Ours)	0.90	20.4	0.91	55.8	5.9
<i>Phase II: Reverse Deconstruction</i>					
BAGEL-zebra-CoT [7]	0.73	12.4	0.71	83.6	13.2
ThinkMorph [12]	0.75	14.0	0.76	76.8	10.7
MindBlock-Random (Ours)	<u>0.82</u>	<u>16.3</u>	<u>0.83</u>	<u>67.9</u>	<u>8.5</u>
MindBlock-Semantic (Ours)	0.88	19.1	0.88	60.2	6.8

that reasoning about viewpoint transformations is substantially more difficult. Even strong models frequently produce inconsistent internal layouts when asked to imagine the same structure from a new viewpoint, leading to degraded reconstruction quality and larger geometric errors. This suggests that current UMMs still lack robust 3D-equivariant internal representations.

3D Evaluation Reveals Hidden Failures. Importantly, several models that appear competitive under 2D metrics reveal substantial failures under our 3DGS-Eval protocol. This demonstrates that traditional image-based evaluation may significantly overestimate spatial reasoning ability. By explicitly reconstructing a volumetric scene from model-generated views, our evaluation exposes inconsistencies that remain invisible under standard 2D metrics.

5.3 Multi-step Assembly: The Challenge of Interleaved Reasoning

Multi-step assembly is a more rigorous probe of spatial intelligence, as it requires the model to generate an interleaved sequence of text (reasoning/instructions) and images (intermediate states). Most contemporary UMMs are limited to single-turn image generation and fail this task entirely. We evaluate the few models capable of long-context interleaved output: *BAGEL-zebra-CoT*, *ThinkMorph*, and our trained *MindBlock* variants.

Model Specialization for Interleaved Reasoning. Since the ability to generate long-context, interleaved text-image sequences is not yet a standard feature in most off-the-shelf UMMs, we perform continual training on the *BAGEL-zebra-CoT* [7] backbone to establish a competitive baseline for multi-step evaluation. We develop two specialized variants: **MindBlock-Random** and **MindBlock-Semantic**, fine-tuned on the respective partitions of our dataset (Figure 2). During the training phase, the models are jointly supervised on two core assembly operators: *forward assembly* and *reverse deconstruction*. To enforce robust 3D structural awareness, each state transition is supervised across four views,

Table 3: Comparison of Spatial Intelligence between Direct Generation (BAGEL) and CoT-Augmented Generation (UniCoT) [34]. 3DGS-Eval metrics represent the quality of a 3D reconstruction from 8 generated views. CoT-based approach significantly reduces geometric errors (Chamfer Distance) while maintaining high semantic alignment.

Model	CLIP-S $\uparrow$	3DGS-PSNR $\uparrow$	3DGS-SSIM $\uparrow$	Chamfer Dist. $\downarrow$
BAGEL (Direct)	0.782	13.42	0.712	12.4
UniCoT [34]	0.814	20.88	0.785	6.5
<i>Gain (Δ_{SR})</i>	+0.032	+7.46	+0.073	−47.6%

requiring the model to maintain viewpoint-invariant information while progressing through the interleaved reasoning trajectory. This specialization ensures that the performance gains reported in Table 2 reflect the models’ internalized spatial logic rather than a mere failure to adhere to the interleaved format.

The results in Table 2 underscore the difficulty of maintaining structural integrity over multiple steps. General-purpose models such as *ThinkMorph* suffer from “spatial drift,” where small errors in early steps accumulate and lead to physically implausible final structures. In contrast, our *MindBlock-Semantic* model, which benefits from a two-stage curriculum over semantically meaningful structures, demonstrates stronger stability, suggesting that the model learns spatial regularities beyond simple symbolic matching.

6 Further Analyses and Discussions

6.1 Is Spatial Reasoning Carried by CoT or Intuition?

A fundamental question in unified multimodal modeling is whether spatial structural awareness is an emergent property of generative "intuition" (pixel-level patterns) or if it relies on explicit symbolic reasoning. We investigate this by comparing UniCoT [34], which generates a textual Chain-of-Thought (CoT) prior to the image, against the base BAGEL model [7] which performs direct image synthesis.

The Divergence Metric. We formalize the impact of reasoning through the *Spatial Reasoning Gain* (Δ_{SR}), measuring the improvement in 3D structural integrity when CoT is enabled. Let $\mathcal{G}_{cot}$ and $\mathcal{G}_{dir}$ denote the image generation processes for UniCoT and the baseline, respectively. The divergence is defined as:

$$\Delta_{SR} = \mathbb{E}_{s \sim \mathcal{S}} [\Phi(\mathcal{G}_{cot}(I, T | \text{CoT})) - \Phi(\mathcal{G}_{dir}(I, T))] \quad (1)$$

where Φ represents the 3DGS-Eval score, and $\mathcal{S}$ is the set of single step assembly instructions in MindBlock.

Quantitative Analysis: CoT as a Symbolic Scaffold. As shown in Table 3, direct generation (BAGEL) achieves reasonable 2D similarity (CLIP-S) but fails under 3DGS-Eval. Without the "textual anchor" of a CoT, the model often neglects occluded block faces or misinterprets the depth of "next-to" relations. In contrast, UniCoT shows improved 3D reconstructibility. This suggests that the CoT acts as a **symbolic scaffold**: by explicitly predicting the coordinates and support relations (e.g., "*Block A is placed at (x,y,z) on top of B*"), the model constrains the latent diffusion process to respect 3D geometry rather than merely hallucinating a 2D projection.

6.2 The Consistency Gap: Perception vs. Execution

To determine whether the failure of spatial reasoning originates from perceptual misunderstanding or generative incapacity, we conduct a focused diagnostic analysis on the **BAGEL** model. Specifically, we decouple the *Spatial Assembly* task into two paradigms: (1) **Discriminative Recognition**, where BAGEL must select the correct next-step image from a 4-way Multiple-Choice Question (MCQ), and (2) **Constructive Execution**, where the model generates the next-step image in pixel space.

To objectively evaluate the semantic fidelity of the generated images, we introduce a **GPT-Matching** protocol. A GPT-5 judge is presented with the model-generated image alongside the four ground-truth and distractor options from the MCQ. The judge identifies which of the four options the generated image most closely resembles. If the judge selects the ground-truth option, the execution is deemed successful.

Table 4: Diagnostic Analysis of BAGEL: The Knowledge-Action Gap. We contrast the model’s accuracy in selecting the correct spatial step against its ability to generate a semantically aligned image of that same step.

Task Setting	Perception Acc.↑	Execution Match↑	The Gap (Δ)↓
Single-step Assembly	52.5%	28.2%	24.3%

The Generative Bottleneck. As shown in Table 4, a discrepancy exists between BAGEL’s internal spatial knowledge and its constructive output. While the model correctly identifies the intended spatial state in more than half of the cases (52.5%), it successfully executes that state in pixel space only 28.2% of the time. This 24.3% gap reveals that spatial intelligence in current UMMS is often *superficial*: models may possess the discriminative "intuition" to recognize correct 3D relationships, yet they lack the **generative grounding** required to translate this symbolic understanding into consistent, physically-plausible visual structures. Our GPT-Matching results further indicate that failed generations frequently drift toward visual distractors, even when the model’s discriminative head had correctly rejected them during the MCQ phase.

7 Conclusion

In this work, we introduced **MindBlock**, a diagnostic benchmark for evaluating spatial intelligence in unified multimodal models through *active construction*. Moving beyond passive visual question answering, MindBlock requires models to generate step-by-step visual assemblies and maintain structural consistency under spatial transformations. To rigorously evaluate geometric coherence, we further proposed **3DGS-Eval**, a volumetric consistency protocol that reconstructs a 3D representation from model-generated multi-view images.

Our key findings are threefold. First, a *2D-3D consistency gap* pervades current UMMs: models that achieve high perceptual similarity scores often produce geometrically incoherent multi-view generations, a failure revealed only by our **3DGS-Eval** protocol. Second, geometric equivariance remains a critical bottleneck, with all evaluated models degrading substantially when reasoning about the same structure from a novel viewpoint. Third, a pronounced *perception-execution gap* (24.3% for BAGEL) shows that models can often understand what to build yet fail to construct it accurately, suggesting that discriminative spatial knowledge does not automatically transfer to generative spatial grounding. Explicit Chain-of-Thought reasoning partially bridges this gap, serving as a symbolic scaffold that stabilizes 3D structure during generation.

Future Work. Several directions follow naturally from these findings. **(1) Closing the perception-execution gap.** Future work could explore targeted training objectives that directly supervise the translation from spatial understanding to pixel-level generation, such as multi-view consistency losses or geometry-aware reward signals. **(2) Scaling to more complex structures.** MindBlock currently covers structures of up to 40 blocks. Extending it to larger and semantically richer environments would test whether spatial drift and equivariance failures scale predictably with structural complexity. **(3) Beyond rigid blocks.** The active construction paradigm introduced here is not limited to block structures. Applying it to articulated objects, deformable scenes, or continuous 3D environments would broaden the scope of generative spatial evaluation. **(4) Grounding in embodied action.** The perception-execution gap identified here mirrors challenges in robotics, where models must translate spatial understanding into physical manipulation. MindBlock could serve as a diagnostic stepping stone toward evaluating embodied agents in simulation.

Overall, **MindBlock** provides a controlled and interpretable framework for probing spatial reasoning in unified multimodal models, offering insights that may inform the development of more robust and physically grounded world models for embodied and interactive AI systems.

Acknowledgements

We thank Jialong Wu and Zhongang Cai for helpful discussions. The work was supported by NYU IT High Performance Computing resources. Chen Feng was supported by his NSF grant No. 2238968.

References

1. Cai, J., Cai, Z., Cao, J., Chen, Y., He, Z., Jiang, L., Li, H., Li, H., Li, Y., Liu, Y., et al.: Internvla-a1: Unifying understanding, generation and action for robotic manipulation. arXiv preprint arXiv:2601.02456 (2026)
2. Cen, J., Huang, S., Yuan, Y., Li, K., Yuan, H., Yu, C., Jiang, Y., Guo, J., Li, X., Luo, H., et al.: Rynnvla-002: A unified vision-language-action and world model. arXiv preprint arXiv:2511.17502 (2025)
3. Chen, J., Xue, L., Xu, Z., Pan, X., Yang, S., Qin, C., Yan, A., Zhou, H., Chen, Z., Huang, L., Zhou, T., Li, J., Savarese, S., Xiong, C., Xu, R.: Blip3o-next: Next frontier of native image generation (2025), <https://arxiv.org/abs/2510.15857>
4. Chen, Y., Wang, C., Fei-Fei, L., Liu, C.K.: Sequential dexterity: Chaining dexterous policies for long-horizon manipulation. Proceedings of Machine Learning Research (2023)
5. Cui, Y., Chen, H., Deng, H., Huang, X., Li, X., Liu, J., Liu, Y., Luo, Z., Wang, J., Wang, W., Wang, Y., Wang, C., Zhang, F., Zhao, Y., Pan, T., Li, X., Hao, Z., Ma, W., Chen, Z., Ao, Y., Huang, T., Wang, Z., Wang, X.: Emu3.5: Native multimodal models are world learners (2025), <https://arxiv.org/abs/2510.26583>
6. Cui, Y., Chen, H., Deng, H., Huang, X., Li, X., Liu, J., Liu, Y., Luo, Z., Wang, J., Wang, W., et al.: Emu3. 5: Native multimodal models are world learners. arXiv preprint arXiv:2510.26583 (2025)
7. Deng, C., Zhu, D., Li, K., Gou, C., Li, F., Wang, Z., Zhong, S., Yu, W., Nie, X., Song, Z., et al.: Emerging properties in unified multimodal pretraining. arXiv preprint arXiv:2505.14683 (2025)
8. Dongfang, Z., Zheng, X., Weng, Z., Lyu, Y., Paudel, D.P., Van Gool, L., Yang, K., Hu, X.: Are multimodal large language models ready for omnidirectional spatial reasoning? arXiv preprint arXiv:2505.11907 (2025)
9. Frick, A., Ferrara, K., Newcombe, N.S.: Using a touch screen paradigm to assess the development of mental rotation between 31/2 and 51/2 years of age. Cognitive processing 14(2), 117–127 (2013)
10. Genesis Authors: Genesis: A universal and generative physics engine for robotics and beyond. <https://github.com/Genesis-Embodied-AI/Genesis> (2024)
11. Gu, C., Liu, J., Chen, H., Huang, R., Wuwu, Q., Liu, Z., Li, X., Li, Y., Zhang, R., Jia, P., et al.: Manualvla: A unified vla model for chain-of-thought manual generation and robotic manipulation. arXiv preprint arXiv:2512.02013 (2025)
12. Gu, J., Hao, Y., Wang, H.W., Li, L., Shieh, M.Q., Choi, Y., Krishna, R., Cheng, Y.: Thinkmorph: Emergent properties in multimodal interleaved chain-of-thought reasoning (2026), <https://arxiv.org/abs/2510.27492>
13. Han, W., Xiang, S., Liu, C., Wang, R., Feng, C.: Spare3d: A dataset for spatial reasoning on three-view line drawings. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 14690–14699 (2020)
14. Hu, J., Zhao, S., Chen, Q.G., Qiu, X., Liu, J., Xu, Z., Luo, W., Zhang, K., Lu, Y.: Omni-view: Unlocking how generation facilitates understanding in unified 3d model based on multiview images. arXiv preprint arXiv:2511.07222 (2025)
15. Huang, R., Wang, C., Yang, J., Lu, G., Yuan, Y., Han, J., Hou, L., Zhang, W., Hong, L., Zhao, H., Xu, H.: Illume+: Illuminating unified mllm with dual visual tokenization and diffusion refinement (2025), <https://arxiv.org/abs/2504.01934>
16. Intelligence, P., Black, K., Brown, N., Darpinian, J., Dhabalia, K., Driess, D., Esmail, A., Equi, M., Finn, C., Fusai, N., et al.: $\pi_{0.5}$: a vision-language-action model with open-world generalization. arXiv preprint arXiv:2504.16054 (2025)

17. Jia, M., Qi, Z., Zhang, S., Zhang, W., Yu, X., He, J., Wang, H., Yi, L.: Omnispatial: Towards comprehensive spatial reasoning benchmark for vision language models. arXiv preprint arXiv:2506.03135 (2025)
18. Kim, J.W.: Survey on automated lego assembly construction. WSCG (2014)
19. Landau, B., Davis, E.E., Cortesa, C.S., Wang, Z., Jones, J.D., Shelton, A.L.: Young children’s copying of block constructions: Significant constraints in a highly complex task. *Cognitive Development* **71**, 101463 (2024)
20. Li, L., Bigverdi, M., Gu, J., Ma, Z., Yang, Y., Li, Z., Choi, Y., Krishna, R.: Unfolding spatial cognition: Evaluating multimodal models on visual simulations. arXiv preprint arXiv:2506.04633 (2025)
21. Li, Y., Wang, H., Zhang, Q., Xiao, B., Hu, C., Wang, H., Li, X.: Unieval: Unified holistic evaluation for unified multimodal understanding and generation. arXiv preprint arXiv:2505.10483 (2025)
22. Li, Z., Li, H., Shi, Y., Farimani, A.B., Kluger, Y., Yang, L., Wang, P.: Dual diffusion for unified image generation and understanding. In: *Proceedings of the Computer Vision and Pattern Recognition Conference*. pp. 2779–2790 (2025)
23. Liu, C., Liu, Y., Wang, T., Zhuang, Q., Liang, J.C., Yang, W., Xu, R., Wang, Q., Liu, D., Han, C.: On-the-fly vla adaptation via test-time reinforcement learning. arXiv preprint arXiv:2601.06748 (2026)
24. Liu, M., Xu, Z., Lin, Z., Ashby, T., Rimchala, J., Zhang, J., Huang, L.: Holistic evaluation for interleaved text-and-image generation. arXiv preprint arXiv:2406.14643 (2024)
25. Liu, R., Huang, P., Pun, A., Deng, K., Aggarwal, S., Tang, K., Liu, M., Ramanan, D., Zhu, J.Y., Li, J., et al.: Prompt-to-product: Generative assembly via bimanual manipulation. arXiv preprint arXiv:2508.21063 (2025)
26. Liu, S., Han, Y., Xing, P., Yin, F., Wang, R., Cheng, W., Liao, J., Wang, Y., Fu, H., Han, C., Li, G., Peng, Y., Sun, Q., Wu, J., Cai, Y., Ge, Z., Ming, R., Xia, L., Zeng, X., Zhu, Y., Jiao, B., Zhang, X., Yu, G., Jiang, D.: Step1x-edit: A practical framework for general image editing (2025), <https://arxiv.org/abs/2504.17761>
27. Liu, Z., Gu, Y., Zheng, S., Xue, X., Fu, Y.: Trivla: A triple-system-based unified vision-language-action model for general robot control. arXiv e-prints pp. arXiv–2507 (2025)
28. Luo, S.J., Yue, Y., Huang, C.K., Chung, Y.H., Imai, S., Nishita, T., Chen, B.Y.: Legolization: Optimizing lego designs. *ACM Transactions on Graphics (ToG)* **34**(6), 1–12 (2015)
29. Ma, L., Wen, J., Lin, M., Xu, R., Liang, X., Lin, B., Ma, J., Wang, Y., Wei, Z., Lin, H., Han, M., Cao, M., Chen, B., Laptev, I., Liang, X.: Phyblock: A progressive benchmark for physical understanding and planning via 3d block assembly (2025), <https://arxiv.org/abs/2506.08708>
30. Ma, W., Chen, H., Zhang, G., Chou, Y.C., Chen, J., de Melo, C., Yuille, A.: 3dsr-bench: A comprehensive 3d spatial reasoning benchmark. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision*. pp. 6924–6934 (2025)
31. Ma, W., Chou, Y.C., Liu, Q., Wang, X., de Melo, C., Xie, J., Yuille, A.: Spatial-reasoner: Towards explicit and generalizable 3d spatial reasoning. arXiv preprint arXiv:2504.20024 (2025)
32. Ono, S., André, A., Chang, Y., Nakajima, M.: Lego builder: automatic generation of lego assembly manual from 3d polygon model. *ITE Transactions on Media Technology and Applications* **1**(4), 354–360 (2013)
33. Pun, A., Deng, K., Liu, R., Ramanan, D., Liu, C., Zhu, J.Y.: Generating physically stable and buildable brick structures from text. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision*. pp. 14798–14809 (2025)

34. Qin, L., Gong, J., Sun, Y., Li, T., Yang, M., Yang, X., Qu, C., Tan, Z., Li, H.: Uni-cot: Towards unified chain-of-thought reasoning across text and vision (2026), <https://arxiv.org/abs/2508.05606>
35. Shi, Q., Bai, J., Zhao, Z., Chai, W., Yu, K., Wu, J., Song, S., Tong, Y., Li, X., Li, X., et al.: Muddit: Liberating generation beyond text-to-image with a unified discrete diffusion model. arXiv preprint arXiv:2505.23606 (2025)
36. Tang, K., Gao, J., Zeng, Y., Duan, H., Sun, Y., Xing, Z., Liu, W., Lyu, K., Chen, K.: Lego-puzzles: How good are mllms at multi-step spatial reasoning? arXiv preprint arXiv:2503.19990 (2025)
37. Thompson, R., Ghalebi, E., DeVries, T., Taylor, G.W.: Building lego using deep generative models of graphs. arXiv preprint arXiv:2012.11543 (2020)
38. Wang, G.H., Zhao, S., Zhang, X., Cao, L., Zhan, P., Duan, L., Lu, S., Fu, M., Chen, X., Zhao, J., Li, Y., Chen, Q.G.: Ovis-u1 technical report (2025), <https://arxiv.org/abs/2506.23044>
39. Wang, R., Zhang, Y., Mao, J., Cheng, C.Y., Wu, J.: Translating a visual lego manual to a machine-executable plan. In: Avidan, S., Brostow, G., Cissé, M., Farinella, G.M., Hassner, T. (eds.) Computer Vision – ECCV 2022. pp. 677–694. Springer Nature Switzerland, Cham (2022)
40. Wang, S., Pei, M., Sun, L., Deng, C., Shao, K., Tian, Z., Zhang, H., Wang, J.: Spatialviz-bench: An mllm benchmark for spatial visualization. arXiv preprint arXiv:2507.07610 (2025)
41. Wang, X., Ma, W., Zhang, T., de Melo, C.M., Chen, J., Yuille, A.: Spatial457: A diagnostic benchmark for 6d spatial reasoning of large multimodal models. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 24669–24679 (2025)
42. Wu, C., Li, J., Zhou, J., Lin, J., Gao, K., Yan, K., Yin, S.m., Bai, S., Xu, X., Chen, Y., et al.: Qwen-image technical report. arXiv preprint arXiv:2508.02324 (2025)
43. Xiao, S., Wang, Y., Zhou, J., Yuan, H., Xing, X., Yan, R., Li, C., Wang, S., Huang, T., Liu, Z.: Omnigen: Unified image generation (2024), <https://arxiv.org/abs/2409.11340>
44. Xie, J., Mao, W., Bai, Z., Zhang, D.J., Wang, W., Lin, K.Q., Gu, Y., Chen, Z., Yang, Z., Shou, M.Z.: Show-o: One single transformer to unify multimodal understanding and generation. In: The Thirteenth International Conference on Learning Representations (2025)
45. Xie, J., Yang, Z., Shou, M.Z.: Show-o2: Improved native unified multimodal models. arXiv preprint arXiv:2506.15564 (2025)
46. Xie, X., Zou, C., Karumuri, M.G., Lenssen, J.E., Pons-Moll, G.: Mvgbench: Comprehensive benchmark for multi-view generation models (2025), <https://arxiv.org/abs/2507.00006>
47. Yang, J., Yang, S., Gupta, A.W., Han, R., Fei-Fei, L., Xie, S.: Thinking in space: How multimodal large language models see, remember, and recall spaces. In: Proceedings of the Computer Vision and Pattern Recognition Conference (2025)
48. Yang, L., Tian, Y., Li, B., Zhang, X., Shen, K., Tong, Y., Wang, M.: Mmada: Multimodal large diffusion language models. arXiv preprint arXiv:2505.15809 (2025)
49. Yang, S., Yang, J., Huang, P., Brown, E., Yang, Z., Yu, Y., Tong, S., Zheng, Z., Xu, Y., Wang, M., et al.: Cambrian-s: Towards spatial supersensing in video. arXiv preprint arXiv:2511.04670 (2025)
50. Yang, S., Xu, R., Xie, Y., Yang, S., Li, M., Lin, J., Zhu, C., Chen, X., Duan, H., Yue, X., et al.: Mmsi-bench: A benchmark for multi-image spatial intelligence. arXiv preprint arXiv:2505.23764 (2025)

51. Yeh, C.H., Wang, C., Tong, S., Cheng, T.Y., Wang, R., Chu, T., Zhai, Y., Chen, Y., Gao, S., Ma, Y.: Seeing from another perspective: Evaluating multi-view understanding in mllms. arXiv preprint arXiv:2504.15280 (2025)
52. Zhang, J., Cherian, A., Rodriguez, C., Deng, W., Gould, S.: Manual-pa: Learning 3d part assembly from instruction diagrams. In: Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV). pp. 6304–6314 (October 2025)
53. Zhang, J., Bai, S., Wang, T., Guo, K., Han, K., Rao, G., Yu, K.: Ascending the infinite ladder: Benchmarking spatial deformation reasoning in vision-language models. arXiv preprint arXiv:2507.02978 (2025)
54. Zhang, J., Guo, Y., Hu, Y., Chen, X., Zhu, X., Chen, J.: Up-vla: A unified understanding and prediction model for embodied agent. In: International Conference on Machine Learning. pp. 74911–74922. PMLR (2025)
55. Zhou, P., Peng, X., Song, J., Li, C., Xu, Z., Yang, Y., Guo, Z., Zhang, H., Lin, Y., He, Y., et al.: Opening: A comprehensive benchmark for judging open-ended interleaved image-text generation. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 56–66 (2025)
56. Zhu, F., Wu, H., Guo, S., Liu, Y., Cheang, C., Kong, T.: Irasim: A fine-grained world model for robot manipulation. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 9834–9844 (2025)