

Object-Uni: A Unified Model for Object-Centric Spatial Understanding and Controllable Generation

Mining Tan^{1,2†} Yinuo Wang^{3†} Ziqi Zhou^{2,1} Weize Quan^{2,1}
 Sifei Li^{2,1} Jingdong Chen⁴ DanDan Zheng⁴ Libin Wang⁴ Weiming Dong^{2,1 *}

¹School of Artificial Intelligence, University of Chinese Academy of Sciences

²MAIS, Institute of Automation, Chinese Academy of Sciences

³School of Vehicle and Mobility, Tsinghua University

⁴Ant Group, China

Abstract

Unified models for visual understanding and generation have made rapid progress, yet they still lack the ability to understand and manipulate the spatial states of object instances. Existing models can describe objects in natural language, but they struggle to precisely represent continuous object poses and generate geometrically consistent images under target viewpoints. To mitigate this, we propose *Object-Uni*, a unified model for object-centric spatial understanding and controllable generation. Specifically, we formulate object-centric spatial intelligence as a unified problem connecting pose perception, spatial reasoning, pose-conditioned generation, and object-centric novel view synthesis. We treat object pose as an explicit geometric variable shared by understanding and generation, rather than merely a prediction label or control signal. To make pose usable by multimodal large language models, we propose a viewpoint-based orientation abstraction that maps orientation into structured viewpoint descriptions while preserving continuous geometric supervision. We further construct an object-centric spatial benchmark (UniSpatial-80K) and train a unified model with an object-token-grounded pose anchor to associate each instance with its pose state. Experiments show that our model improves object-level pose understanding and pose-controllable generation, moving unified models from describing objects toward manipulating spatial states.

1 Introduction

The rapid development of unified models for visual understanding and generation is pushing visual intelligence from merely “understanding images” toward “understanding and generating the visual world.” Recent unified models [4, 32, 31, 7, 29, 2] have demonstrated strong capabilities in image captioning, visual question answering, image generation, and image editing, showing impressive progress in open-ended semantic understanding and high-quality visual synthesis. However, real-world visual intelligence requires more than recognizing objects and semantic content in images. A model should also understand the spatial location and orientation of object instances, as well as how an object should appear when the viewpoint changes. We refer to this capability as **object-centric spatial intelligence**: the ability to reason about object instances in terms of their spatial locations, poses, and relative relations, and to support pose-conditioned controllable generation and object-centric novel view synthesis.

Current methods typically rely on free-form textual descriptions to represent object spatial states, such as “the shoe on the left,” “a car viewed from a high angle,” or “a chair facing right.” [16,

*Corresponding author <weiming.dong@ia.ac.cn>.

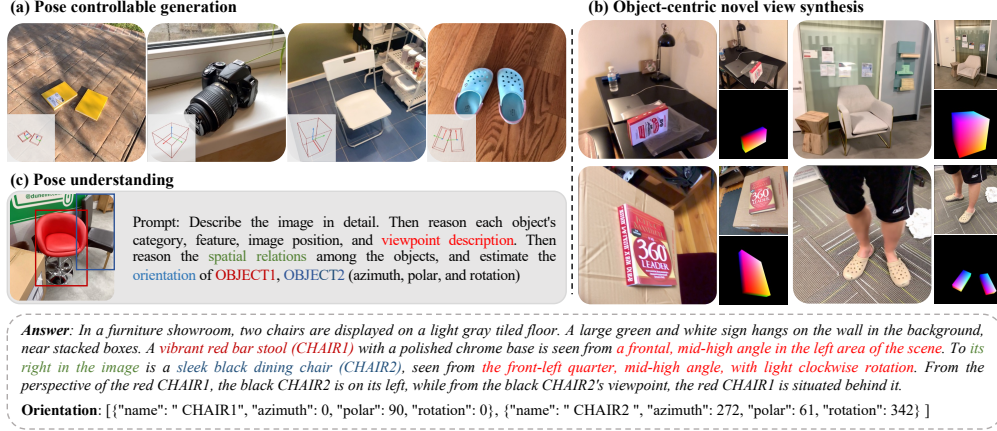


Figure 1: Versatile capabilities of Object-Uni model. It unifies (a) object-centric generation and (c) understanding, also supporting orientation estimation and (b) object-centric novel view synthesis.

14] Such descriptions are natural and flexible, but they lack stable geometric constraints and are insufficient for precisely representing continuous object poses. Meanwhile, conventional object pose estimation mainly focuses on predicting numerical pose values from images, whereas pose-conditioned generation aims to synthesize target-view images under given pose conditions. These two lines of research are usually studied as separate tasks, lacking a unified problem formulation that jointly characterizes object pose perception, spatial reasoning, and pose-conditioned generation.

In this work, we formulate object-centric spatial intelligence as a unified understanding-and-generation problem. Under this formulation, a model is required not only to infer the spatial state of objects from images but also to generate geometrically consistent visual content according to a target spatial state. Our central insight is that object pose should not be treated merely as a label for pose estimation, nor only as an external control signal for image generation. Instead, it should serve as an explicit object-centric geometric variable that unifies pose perception, spatial reasoning, and pose-conditioned generation. The key challenge is how to make object pose understandable and generable by multimodal large language models (MLLMs). Object locations and relative spatial relations can often be expressed naturally in language, such as left of, behind, or on top of. In contrast, object orientation is a multi-dimensional continuous geometric variable that is difficult to directly align with the linguistic reasoning space of MLLMs. To address this issue, we propose a viewpoint-based orientation abstraction, which maps azimuth, polar, and rotation into structured viewpoint descriptions, such as front view, high-angle view, and clockwise rotation. This abstraction does not replace continuous pose with discrete viewpoint labels. Instead, it builds a bridge between continuous geometric supervision and language-level spatial reasoning: continuous pose preserves precise geometric constraints and enables quantitative evaluation, while structured viewpoint descriptions allow the model to understand and reason about object pose in natural language.

We further construct an object-centric spatial understanding and generation benchmark with 80K+ samples. The benchmark covers complex scenarios, including street-view vehicles and pedestrians, indoor objects, and general indoor and outdoor objects. For each object instance, it organizes spatial location, relative relations, continuous pose, and viewpoint-based descriptions. Unlike conventional pose estimation datasets that focus on a single task, our benchmark is organized into a set of interconnected understanding and generation tasks, including object spatial-relation understanding, viewpoint reasoning, orientation prediction, pose-conditioned image generation, and object-centric novel view synthesis. These tasks forms a unified evaluation protocol: a model can both understand object poses from images and generate geometrically consistent images according to target poses.

Based on this benchmark, we train a unified object-centric spatial intelligence model that performs object pose understanding, viewpoint reasoning, and pose-controllable generation within a single framework. To enable stable orientation understanding in both single-object and multi-object scenes, we design an object-token-grounded pose anchor that aggregates visual features corresponding to each object instance and explicitly associates each object with its orientation state. We illustrate the versatile capabilities of our Object-Uni model in Figure 1. Our contribution can be summarized as:

- We define object-centric spatial intelligence as a unified understanding-and-generation problem. Unlike prior studies that treat object pose estimation and pose-conditioned generation as separate tasks, we model object pose as an explicit object-centric geometric variable that connects pose perception, spatial reasoning, and controllable generation.
- We propose a viewpoint-based orientation abstraction, which converts continuous 3D orientations into structured natural-language viewpoint descriptions. It allows MLLMs to understand and reason about object orientation through natural viewpoint terms.
- We construct an object-centric spatial understanding-and-generation benchmark with 80K+ samples. The benchmark covers both single-object and multi-object scenarios and unifies object spatial-relation understanding, viewpoint reasoning, orientation prediction, and pose-conditioned generation under systematic evaluation.
- We train a unified model for object-centric spatial intelligence through multi-stage training, achieving superior performance on both understanding and generation tasks. In addition, for orientation prediction during the understanding stage, we introduce an object-token-grounded pose anchor that explicitly associates each predicted orientation with its corresponding object instance, thereby improving prediction performance in multi-object scenes.

2 Related Work

Object Pose Understanding from Images Understanding object pose and orientation is a fundamental requirement for object-centric spatial intelligence. Early object-centric 3D benchmarks such as Objectron [1] provide videos with camera poses, sparse point clouds, and manually annotated 3D bounding boxes describing object position, orientation, and dimensions, while Omni3D [3] studies large-scale 3D object detection across diverse indoor and outdoor scenes. More recently, Orient Anything [25] formulates open-world object orientation estimation from a single image by rendering large-scale 3D assets with precise azimuth, polar, and in-plane rotation annotations. Orient Anything V2 [28] further extends this formulation to rotationally symmetric objects and relative rotation estimation. In parallel, recent MLLM-oriented benchmarks such as Right Side Up [16] reveal that current multimodal models still struggle with fine-grained multi-axis orientation understanding. Unlike these works, which primarily evaluate or predict object pose, our goal is to make object orientation a shared variable for both multimodal reasoning and controllable generation.

Spatially Controlled Generation Controlling object pose or viewpoint in text-to-image generation has been explored from several directions. Continuous 3D Words [5] encode continuous 3D-aware attributes, including object pose, as special tokens that can be smoothly varied during generation. Custom Diffusion 360 [11] focuses on subject-driven generation with viewpoint control for customized objects from multi-view images. Compass Control [19] addresses multi-object orientation control by introducing object-specific compass tokens and constraining their attention maps to corresponding object regions. SceneDesigner [22] further studies controllable multi-object image generation with 9-DoF pose manipulation and uses CNOCS maps as dense geometric pose conditions. The main difference between our generation module with those works is that we do not treat pose solely as a generation control signal; instead, the same object-level pose representation is also predicted and reasoned about by the MLLM side, enabling a unified understanding-and-generation framework.

Unified Models Recent unified multimodal models aim to support visual understanding and generation within a single framework. Chameleon [24], Emu2 [23], Lumina-mGPT [12], Show-O [30], and Janus [29] explore different ways to unify image-text comprehension and generation, such as mixed-modal autoregressive modeling, autoregressive-diffusion hybrid objectives, or decoupled visual encoders. These works mainly focus on architecture- or objective-level unification for general multimodal tasks. In contrast, our work focuses on object-centric spatial unification: we make object orientation a shared variable that can be described, predicted, grounded to an object token, and reused as a control signal for image generation. Therefore, Object-Uni is not intended as a new general-purpose unified architecture, but as a pose-centered interface built on a unified multimodal backbone for bridging object-level spatial understanding and controllable generation.

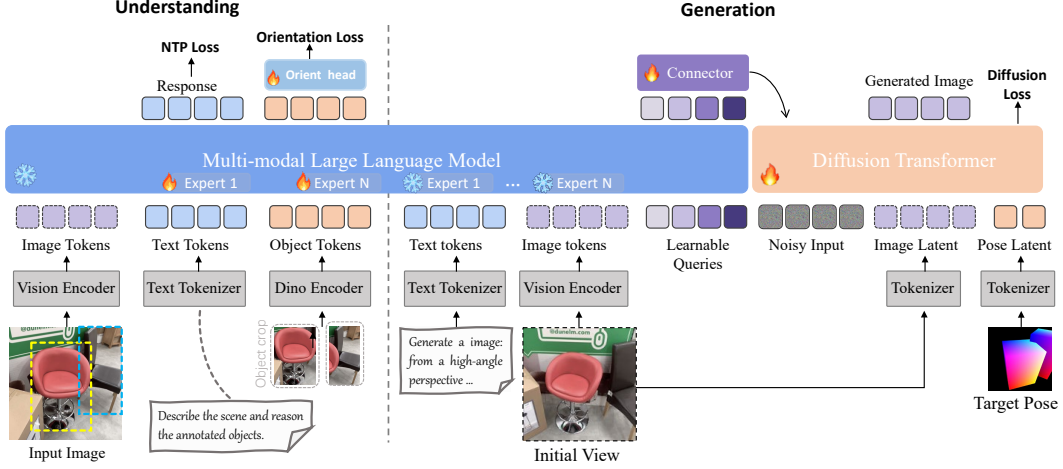


Figure 2: Overview of Object-Uni. Our model unifies object-centric spatial understanding and orientation-controllable generation. The dotted modules denote object-centric novel view synthesis.

3 Object-Centric Understanding and Controllable Generation

3.1 Overview

We study object-centric spatial intelligence in a unified understanding-and-generation framework. As shown in Figure 2, given an image, a text prompt, and object-level spatial conditions, the model is required to infer object locations, relative spatial relations, continuous orientations, and further generate images that satisfy the specified object-centric pose conditions. The main difficulty lies in the representational mismatch between object pose and multimodal language modeling. Object pose is a continuous and periodic geometric variable, while MLLMs use discrete tokens and natural-language descriptions. To bridge this gap, we introduce a viewpoint-based orientation abstraction that maps continuous orientation into structured spatial language while preserving the original continuous pose for metric supervision and controllable generation. Based on this, we construct the UniSpatial-80K benchmark, and train an object-centric unified multimodal model for pose understanding, spatial reasoning, pose-controllable generation, and object-centric novel view synthesis.

3.2 Viewpoint-based Orientation Abstraction

For each object instance obj_i in an image I , we denote its 2D bounding box as $b_i = (x_i^1, y_i^1, x_i^2, y_i^2)$ and its 9D pose as $q_i = (\mathbf{l}_i, \mathbf{s}_i, \mathbf{o}_i)$. Here, $\mathbf{l}_i \in \mathbb{R}^3$ denotes the 3D location of the object center in the camera coordinate system, $\mathbf{s}_i \in \mathbb{R}^3$ is its 3D size, and $\mathbf{o}_i = (\phi_i, \theta_i, \psi_i)$ denotes its 3D orientation from the camera view, where $\phi_i \in [0^\circ, 360^\circ)$ is the azimuth, $\theta_i \in [0^\circ, 180^\circ]$ is the polar, and $\psi_i \in [0^\circ, 360^\circ)$ is the rotation angle. We follow the orientation convention of Orient Anything [26] for defining these three angles. The azimuth describes the object-facing direction relative to the camera, the polar angle is the camera elevation with respect to the object, and the rotation angle describes the in-plane image rotation. This continuous representation is necessary for quantitative pose supervision and evaluation. However, directly exposing raw angles is hard for MLLM to understand, since numerical tokens do not naturally encode the perceptual meaning of object viewpoint. We therefore construct a dual representation of object pose: the continuous angles remain as metric targets, while a deterministic language abstraction converts them into viewpoint descriptions.

To bridge continuous object orientation and natural language descriptions, we map the three orientation angles $\mathbf{o}_i = (\phi_i, \theta_i, \psi_i)$ into coarse semantic view terms. For azimuth, we divide ϕ_i into eight canonical object views, such as front, front-right quarter, and right side. For polar angle, we map θ_i into camera-elevation descriptions. For in-plane rotation, we simply divide it into three intervals. The detailed mapping rules are summarized in Appendix A.

In addition to orientation, we describe each object using its coarse image-space location and pairwise spatial relations. The normalized bounding-box center $(\bar{x}_i, \bar{y}_i)$ is mapped to horizontal and vertical region terms, such as left, center, right, top, middle, and bottom. For multi-object scenes, we describe relations from both viewer-centric and object-centric perspectives. The former is computed from relative image positions, while the latter further incorporates the azimuth of the reference object to express relations such as left of, right of, in front of, or behind in the object’s local frame. This provides language supervision that jointly reflects image layout and object-centric geometry.

3.3 UniSpatial-80K Dataset

Data Sources We construct UniSpatial-80K, a large-scale object-centric spatial understanding-and-generation benchmark. The benchmark integrates multiple object-level 3D and pose-annotated data sources, including the KITTI [8], Cityscapes [6] and Objectron [1] subsets from OmniNOCS [10], and ImageNet3D [15]. These sources provide complementary scene distributions. KITTI and Cityscapes provide street-view data containing realistic outdoor layouts with vehicles and pedestrians. Objectron provides object-centric indoor videos with background context. ImageNet3D provides broad category coverage and general-purpose object-level 3D annotations. Together, they allow us to evaluate object-centric spatial intelligence beyond a narrow category or scene type.

Data Filtering and Balancing We filter the raw data to ensure reliable object-centric spatial supervision. First, we remove samples with excessive background clutter or too many foreground instances, as these cases make object–pose attribution ambiguous. Second, we exclude categories or samples whose canonical front direction is not well defined, because orientation supervision becomes ill posed when an object lacks a unique or stable front face. Third, to mitigate severe category imbalance, we cap overrepresented categories by sampling at most 5,000 images per category, to improve the reliability of continuous orientation labels and makes the benchmark more suitable for both understanding and generation tasks. UniSpatial-80K contains 83,252 image entries with 91,392 annotated objects across 122 training categories, the detailed distribution and statistics are provided in Appendix B.1.

Pose Conversion and Spatial Caption Construction For each retained object instance, we convert the original dataset-specific pose annotations into the unified convention. We then apply the viewpoint-based orientation abstraction described above to obtain structured spatial descriptions. Each final sample contains a general image caption, an object-centric spatial caption, and continuous pose annotations. The general caption describes the visual content and scene context, while the spatial caption explicitly describes object locations, viewpoints, rotations, and inter-object relations. The example is shown in the Figure 1 (c).

3.4 Object-Centric Unified Multimodal Model

Base Architecture We build our model upon Ming-Lite-Uni[2], a unified multimodal architecture that connects a multimodal autoregressive model with diffusion transformer blocks through multi-scale learnable query representations. The MLLM side receives text instructions and visual tokens, performs multimodal understanding and spatial reasoning, and produces hidden states for language generation. The generation side uses a DiT-based diffusion model to synthesize or edit images conditioned on text, visual features, and spatial control signals. This architecture provides a natural basis for unifying object-centric understanding and controllable generation: the MLLM handles semantic and spatial reasoning, while the diffusion model handles high-fidelity image synthesis. The overall architecture is shown in Figure 2.

Object-Token-Grounded Pose Anchor A direct way to predict orientation is to ask the MLLM to output numerical angles in text and then parse the generated values. However, this strategy is unstable, since numerical angle tokens are weakly calibrated, and in multi-object scenes, the generated text may fail to associate each angle with the correct object instance. To address this issue, we introduce an object-token-grounded pose anchor, PoseAnchor. For each object o_i , we crop the object region according to its bounding box and extract an object-level feature using a frozen vision encoder [18]. The object feature is projected into the MLLM hidden space and inserted into the input sequence as an object-specific token. The hidden state corresponding to this token serves as an instance-level

pose anchor:

$$h_i = \text{MLLM}(I, T, b_i, e_i), \quad (1)$$

where e_i denotes the projected object feature. We then attach lightweight prediction heads to h_i to estimate azimuth, polar, and rotation. This design explicitly binds each pose prediction to the corresponding object instance, which is especially important in multi-object scenes.

Pose Prediction Loss Following Orient Anything [26], we formulate object orientation prediction as a distribution fitting problem over discretized angle bins. For each valid object, the model predicts three distributions corresponding to azimuth ϕ , polar θ , and rotation ψ . The soft target distribution is constructed by placing a Gaussian-like distribution around the ground-truth angle. For azimuth and rotation, we compute the target using circular angular distance to respect periodicity, while for the polar angle we use a non-periodic distance on $[0, 180^\circ]$.

Let p_i^ϕ, p_i^θ , and p_i^ψ denote the predicted distributions, and let q_i^ϕ, q_i^θ , and q_i^ψ denote the corresponding soft target distributions. The pose loss is defined as the average soft cross-entropy over valid objects:

$$\mathcal{L}_{\text{pose}} = \frac{1}{|\mathcal{V}|} \sum_{i \in \mathcal{V}} \left[H(q_i^\phi, p_i^\phi) + H(q_i^\theta, p_i^\theta) + H(q_i^\psi, p_i^\psi) \right], \quad (2)$$

where $H(q, p) = -\sum_k q_k \log p_k$. The full understanding-stage objective is:

$$\mathcal{L}_{\text{und}} = \mathcal{L}_{\text{NTP}} + \lambda_{\text{pose}} \mathcal{L}_{\text{pose}}. \quad (3)$$

Task Formulation UniSpatial-80K is organized as a set of interconnected spatial understanding and generation tasks rather than as a single pose-estimation benchmark. For object-centric spatial understanding, we define the *Image* $\rightarrow$ *Spatial Text* + *Orient* task, in which the model generates a spatial description and predicts the continuous orientation of each object in the input image. For pose-controllable generation, we define the *Text* + *Pose* $\rightarrow$ *Image* task, where the model synthesizes an image consistent with both the textual prompt and the specified object pose conditions. To further leverage the spatial reasoning capabilities of MLLMs, we introduce the *Text* + *Pose* $\rightarrow$ *Spatial Text* task, in which the model converts textual content and pose conditions into structured spatial descriptions. This task serves as an implicit reasoning bridge between spatial understanding and generation. Finally, for object-centric novel view synthesis, we define the *Image* + *Pose* $\rightarrow$ *Image* task, where the model synthesizes a target-view image according to the specified target object pose.

Since language provides only coarse spatial constraints, we additionally encode the target object pose as a generation-side geometric condition. Following CNOCS-based 9-DoF pose control [20], we convert the object pose into a CNOCS map. The CNOCS map provides dense image-space geometric guidance and is encoded into the diffusion model through the VAE encoder.

4 Experiments

We conduct experiments to evaluate whether the proposed model can develop object-centric spatial capabilities for both understanding and generation. Specifically, we study three questions: (1) Can the unified model accurately estimate object orientation from images? (2) Can it effectively control pose for text-to-image generation? (3) Does PoseAnchor improve the orientation estimation, especially in multi-object scenes?

4.1 Implementation Details

In our implementation, the MLLM-side input image is encoded by a vision encoder, whereas the DiT-side reference image and CNOCS map are encoded by the VAE encoder. Each pose anchor `<obj>` is grounded by an object token, associated with a 2D bounding box, and inserted into the input sequence to predict the corresponding orientation tuple. Training proceeds in three stages. First, we train the MLLM-side spatial reasoning component with two understanding-oriented tasks: *Image* $\rightarrow$ *Spatial Text* + *Orient* and *Text* + *Pose* $\rightarrow$ *Spatial Text*. We update only the MLLM Mixture-of-Experts (MoE) layers, the object-token projection module, and the orientation prediction heads, while keeping the rest of the pretrained backbone frozen. Second, we train the DiT generator for orientation-controllable text-to-image generation, i.e., *Text* + *Pose* $\rightarrow$ *Image*. Finally, we train the model for object-centric novel view synthesis, i.e., *Image* + *Pose* $\rightarrow$ *Image*. The full training pipeline takes approximately 18 hours on 8 H20 GPUs.

Table 1: **Evaluation results on object orientation estimation.** The results are grouped by dataset and split. Best results in each dataset-split block are shown in bold, and second-best results are underlined. The Object-Uni w/o represents our model without object-token-grounded pose anchor.

Approach		Azimuth [in °]				Polar [in °]				Rotation [in °]				
		Error↓	AUC@5/10/30↑			Error↓	AUC@5/10/30↑			Error↓	AUC@5/10/30↑			
KITTI-Cityscapes	Single	Random	86.18	2.81	4.16	10.47	43.57	4.25	8.04	19.27	94.49	1.36	2.82	8.36
		GPT-4o	73.38	7.52	13.48	32.05	2.93	45.60	72.41	90.80	0.78	99.30	99.30	99.47
		Gemini-2.5-pro	69.32	10.14	16.21	31.78	2.92	47.87	72.64	90.88	0.80	99.26	99.63	99.88
		OriAny.V1	27.33	14.99	28.03	56.72	7.61	23.42	49.55	79.24	1.53	98.29	98.73	99.03
		OriAny.V2	29.94	15.48	27.07	55.40	2.34	59.72	79.44	93.09	3.65	31.46	65.61	88.54
		Object-Uni w/o	46.16	28.56	38.30	54.76	0.78	93.88	96.94	98.98	0.40	99.45	99.72	99.91
		Object-Uni	22.87	41.54	56.82	74.44	0.56	98.18	99.09	99.70	0.23	99.67	99.83	99.94
	Multi	Random	90.20	2.13	3.42	8.85	43.73	4.42	7.53	18.74	89.87	1.52	2.98	9.46
		GPT-4o	82.94	9.69	13.65	28.16	1.52	80.34	89.04	96.31	0.67	99.62	99.86	99.94
		Gemini-2.5-pro	48.43	24.52	31.18	42.96	2.64	58.30	76.75	92.14	0.68	99.62	99.81	99.94
		OriAny.V1	22.17	22.57	40.04	69.04	10.37	49.83	66.15	81.65	3.27	91.70	94.53	96.42
		OriAny.V2	54.93	17.03	24.12	48.33	0.99	90.39	95.04	98.31	4.35	14.62	57.15	85.70
		Object-Uni w/o	52.02	30.00	41.61	55.58	0.92	90.46	95.14	98.38	0.41	99.72	99.86	99.95
		Object-Uni	28.12	47.06	60.32	74.04	0.46	97.05	98.48	99.49	0.19	99.75	99.87	99.96
ImageNet3D	Single	Random	87.52	2.10	3.74	10.04	46.51	3.60	7.04	18.35	92.03	1.25	2.76	8.30
		GPT-4o	50.02	17.26	21.24	33.88	15.31	26.57	34.82	57.09	10.05	67.51	72.52	82.37
		Gemini-2.5-pro	53.77	11.18	14.24	28.29	13.14	22.95	36.08	64.65	10.79	65.09	71.25	81.07
		OriAny.V1	41.75	14.03	24.21	46.35	26.32	11.28	18.52	39.52	9.57	69.37	74.11	83.58
		OriAny.V2	28.25	26.02	38.76	62.36	10.10	37.31	51.47	73.99	9.80	41.95	57.90	78.04
		Object-Uni w/o	28.03	29.35	37.01	54.38	9.10	38.60	51.18	73.67	9.07	68.05	74.61	84.62
		Object-Uni	24.13	35.23	45.79	64.66	8.15	42.10	55.89	77.50	9.14	70.02	77.23	86.37
	Multi	Random	90.43	2.26	4.29	10.10	49.25	4.51	7.42	17.74	91.77	1.03	2.42	8.40
		GPT-4o	55.39	15.68	20.53	31.89	14.17	25.14	35.36	59.55	12.72	63.01	69.95	81.47
		Gemini-2.5-pro	55.00	11.58	16.83	30.58	14.49	17.90	32.26	60.95	12.80	62.12	67.86	78.17
		OriAny.V1	52.11	13.77	22.14	41.88	24.98	12.74	19.67	40.83	12.30	65.07	70.43	79.96
		OriAny.V2	34.40	22.67	35.15	59.57	10.30	36.78	51.35	73.64	10.72	43.90	59.33	78.51
		Object-Uni w/o	42.08	17.26	22.79	39.08	11.99	33.90	46.73	71.18	12.79	63.15	68.61	79.25
		Object-Uni	33.84	24.79	35.70	56.30	9.47	36.10	50.50	74.13	9.64	66.85	72.35	82.32
Objectron	Single	Random	91.66	1.17	2.35	7.80	53.95	2.39	4.71	15.09	90.15	1.51	2.89	8.58
		GPT-4o	49.50	7.57	12.98	30.48	35.16	2.06	4.26	16.95	13.01	32.95	49.31	73.31
		Gemini-2.5-pro	49.95	5.10	9.59	25.51	20.79	6.51	13.36	42.40	16.36	28.38	43.15	67.08
		OriAny.V1	37.70	15.86	27.43	49.94	31.02	14.00	23.58	59.47	15.01	33.10	48.97	71.83
		OriAny.V2	30.34	14.10	24.95	53.49	11.97	21.34	37.29	69.62	13.72	15.76	29.53	64.02
		Object-Uni w/o	28.90	17.60	29.28	54.63	5.98	36.19	56.40	82.59	11.73	32.25	49.96	75.09
		Object-Uni	22.67	26.57	41.84	65.99	4.24	47.49	67.16	87.84	8.12	46.52	63.25	82.91
	Multi	Random	90.13	1.25	2.62	8.26	58.44	2.38	4.53	14.86	92.07	1.32	2.79	8.75
		GPT-4o	69.37	4.12	6.94	17.79	45.16	0.23	0.53	5.72	23.01	16.43	27.92	53.00
		Gemini-2.5-pro	68.51	3.39	6.09	16.59	18.07	8.76	16.28	46.96	34.96	11.20	19.83	41.32
		OriAny.V1	47.26	6.99	13.94	32.66	45.34	7.24	13.64	44.40	27.41	15.12	25.33	48.70
		OriAny.V2	32.58	9.52	16.71	41.43	10.92	19.41	35.19	69.37	20.62	15.56	27.44	54.40
		Object-Uni w/o	33.95	10.23	18.01	41.27	7.42	31.21	49.36	78.16	17.71	21.82	34.65	62.29
		Object-Uni	25.89	17.62	30.43	57.01	5.49	37.79	57.59	83.06	12.70	29.26	44.16	71.26

4.2 Pose Understanding

Settings We evaluate object orientation estimation on three subsets of UniSpatial-80K: KITTI-Cityscapes, ImageNet3D, and Objectron. Each dataset is evaluated under both single-object and multi-object settings. We compare our model with four baselines: GPT-4o [17], Gemini-2.5-Pro [9], Orient Anything V1 [26], and Orient Anything V2 [27]. For GPT-4o and Gemini-2.5-Pro, we prompt the model with the full image and object bounding boxes, and ask it to estimate the orientation of each target object. The exact prompts are provided in the Appendix B.2. For Orient Anything V1/V2, we crop the target object according to the ground-truth bounding box and feed the cropped object image into the orientation estimator.

Evaluation Metrics For orientation estimation, we report the mean angular error and Area Under the Recall Curve (AUC) under thresholds of 5° , 10° , and 30° . For azimuth and rotation, we use circular angular distance: $d(\hat{\theta}, \theta) = \min(|\hat{\theta} - \theta|, 360^\circ - |\hat{\theta} - \theta|)$. For polar angle, we use the absolute error since it is defined on $[0^\circ, 180^\circ]$.

Comparison Results Table 1 reports the quantitative results. Overall, Object-Uni achieves the best performance across most datasets, splits, and orientation metrics. The gains are especially remarkable for azimuth estimation, which requires reasoning about object-facing direction and is therefore more challenging than image-plane localization. Compared with general-purpose MLLMs, Object-Uni produces more stable object-level orientation estimation, especially in multi-object scenes. It also performs competitively against specialist orientation models while remaining a unified understanding-generation model. These results demonstrate that object-centric spatial supervision and object-token grounding effectively improve orientation awareness, with the azimuth gains providing direct evidence for stronger object-centric spatial alignment.

4.3 Pose-Controllable Generation

Settings We evaluate object pose as a precise, instance-level control signal for image generation. Given a scene description and target object poses, the model must synthesize an image where each specified object is faithful to the prompt, correctly localized, and rendered in the desired orientation. This is harder than standard text-to-image generation because the model must bind geometric controls to the correct object instances, especially in multi-object scenes. We compare with SceneDesigner [20], a representative method for multi-object controllable generation with explicit pose manipulation. Experiments are conducted on the ImageNet3D, KITTI-Cityscapes, and Objectron subsets of our Unispatial dataset, covering general, street-view, and indoor object-centric scenes. For each dataset, we evaluate both single-object and multi-object settings.

Metrics We evaluate generated images along three axes: spatial grounding, pose controllability, and semantic fidelity. For spatial grounding, we first detect generated objects with Grounding DINO [13] and compare the detected boxes with the target layouts. We report mean Intersection over Union (mIoU) to measure localization overlap and spatial accuracy $Acc_{ls} = \frac{1}{N} \sum_{i=1}^N \mathbb{I}(\text{IoU}_i > 0.6)$ to measure localization success, where N denotes the number of evaluated object instances and $\mathbb{I}$ is the indicator function. For pose controllability, we estimate the orientation of each generated object with an orientation estimator [27] and compute its angular deviation from the target condition. We report the mean error and AUC@10/30, which summarize the accuracy curve within 10° and 30° error thresholds, respectively. Finally, we report the CLIP score [21] to measure text-image alignment. These metrics jointly assess whether the model can generate objects at the correct location, under the correct orientation, while preserving semantic consistency with the input prompt.

Comparison Results As shown in Table 2, Object-Uni achieves stronger overall performance than SceneDesigner across datasets and settings. For spatial grounding, our model achieves clear gains in both mIoU and Acc_{ls} , especially in multi-object scenes, indicating stronger object-level localization and layout binding. For example, on Objectron, Object-Uni improves mIoU from 40.77 to 78.23 in the single-object setting and from 27.25 to 69.08 in the multi-object setting. For pose controllability, Object-Uni generally reduces angular errors and improves AUC@10/30 for azimuth, polar, and rotation. The improvements are particularly evident on ImageNet3D and Objectron, where object appearance varies significantly with viewpoint. The advantage becomes more pronounced in multi-object scenarios, where each object must be associated with its own target orientation. Object-Uni demonstrates a stronger ability to jointly control object location, size, and orientation, while maintaining semantic consistency with the input description (see Figure 3). This indicates that our method can better understand and exploit the input scene description, corresponding to the higher CLIP scores across all settings. More visual results are provided in Appendix C.

4.4 Ablation Study

Effect of Object-Token Grounding Pose Anchor Object-Uni introduces the proposed object tokens and binds each orientation prediction to a specific annotated object. The ablation results are included in Table 1, where Object-Uni w/o denotes the base model trained with the same tasks but



Figure 3: Evaluation of pose-controllable generation in single- and multi-object scenarios. The reference image in the first row serves as the source for the input text and 9D pose annotations. Object-Uni shows better text alignment and pose consistency under diverse pose conditions.

Table 2: **Evaluation results on pose-controllable generation.** We compare SceneDesigner and our Object-Uni model on ImageNet3D, KITTI, and Objectron. The best results are marked in bold.

		Localization		Azimuth [in $^{\circ}$]			Polar [in $^{\circ}$]			Rotation [in $^{\circ}$]			Gen	
Approach		mIoU $\uparrow$	$Acc_{cls}\uparrow$	Err $\downarrow$	AUC@10/30 $\uparrow$		Err $\downarrow$	AUC@10/30 $\uparrow$		Err $\downarrow$	AUC@10/30 $\uparrow$		CLIP $\uparrow$	
ImageNet3D	Sgl.	SceneD.	66.45	68.81	62.10	13.30	24.59	15.77	35.25	58.92	12.12	41.95	70.78	0.333
		Object-Uni	74.31	83.40	34.40	30.14	50.74	10.61	47.32	71.77	9.68	55.57	78.12	0.341
	Mul.	SceneD.	39.03	20.89	67.83	13.99	25.67	20.24	32.48	54.65	16.69	46.01	69.54	0.334
		Object-Uni	54.43	44.86	50.66	22.11	39.43	12.64	42.38	67.36	11.49	58.37	77.56	0.344
KITTI	Sgl.	SceneD.	38.40	26.45	68.17	10.40	24.60	4.03	70.79	87.28	7.86	53.15	76.30	0.319
		Object-Uni	62.75	68.60	69.83	12.60	29.59	2.46	78.09	92.61	4.81	56.96	84.39	0.331
	Mul.	SceneD.	17.17	11.72	89.48	8.28	20.34	9.00	83.63	90.28	20.35	36.17	55.90	0.315
		Object-Uni	53.34	49.34	81.36	9.63	29.40	1.78	90.52	96.34	5.96	54.72	81.69	0.323
Objectron	Sgl.	SceneD.	40.77	13.56	53.44	9.12	22.68	24.67	14.92	39.36	15.83	31.67	63.09	0.326
		Object-Uni	78.23	93.12	32.10	22.61	49.52	13.63	32.03	64.98	14.11	31.68	64.55	0.335
	Mul.	SceneD.	27.25	4.53	82.27	3.65	10.26	24.89	14.39	39.98	25.88	22.90	48.01	0.319
		Object-Uni	69.08	79.37	36.26	13.93	36.34	11.74	32.48	67.48	20.91	27.39	53.82	0.334

without explicit object tokens for orientation prediction. The comparison shows that object-token grounding consistently improves orientation estimation, particularly for azimuth. The improvement is more pronounced in multi-object scenes, where the model must distinguish which object a predicted orientation belongs to.

4.5 Application: Object-Centric Novel View Synthesis

Beyond orientation prediction and text-to-image generation, our Object-Uni model can be applied to object-centric novel view synthesis. Given an input image and a target pose, the model edits the target object so that its appearance matches the desired viewpoint while maintaining its category, local texture, and surrounding context. This task requires the model to perform spatial imagination: it must infer how the object should look after rotation rather than merely copying pixels from the input image. Visual results of object-centric novel-view synthesis are shown in Figure 1 (b). The image on the left shows the generated result; the small image in the upper-right corner shows the input image from the initial viewpoint, and the small image in the lower-right corner shows the target pose. Our model remains consistent with the input appearance while accurately matching the target pose.

5 Conclusion

Object-Uni is a unified framework for object-centric spatial understanding and controllable generation. Built on UniSpatial-80K with object-centric spatial supervision, it jointly models object semantics, layout, pose, and pose-conditioned generation. Experiments show strong spatial understanding and controllable generation under object-level pose constraints.

Limitation. Although Object-Uni shows strong spatial understanding and controllable generation ability, it still has limitations in generating fine-grained text and human details. These issues are partly inherited from the underlying generative backbone and remain important directions for future improvement.

References

- [1] Adel Ahmadyan, Liangkai Zhang, Artsiom Ablavatski, Jianing Wei, and Matthias Grundmann. Objectron: A large scale dataset of object-centric videos in the wild with pose annotations. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 7822–7831, 2021.
- [2] Inclusion AI, Biao Gong, Cheng Zou, Dandan Zheng, Hu Yu, Jingdong Chen, Jianxin Sun, Junbo Zhao, Jun Zhou, Kaixiang Ji, et al. Ming-Lite-Uni: Advancements in unified architecture for natural multimodal interaction. *arXiv preprint arXiv:2505.02471*, 2025.
- [3] Garrick Brazil, Abhinav Kumar, Julian Straub, Nikhila Ravi, Justin Johnson, and Georgia Gkioxari. Omni3d: A large benchmark and model for 3d object detection in the wild. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 13154–13164, 2023.
- [4] Chameleon Team. Chameleon: Mixed-modal early-fusion foundation models. *arXiv preprint arXiv:2405.09818*, 2024.
- [5] Ta-Ying Cheng, Matheus Gadelha, Thibault Groueix, Matthew Fisher, Radomir Mech, Andrew Markham, and Niki Trigoni. Learning continuous 3d words for text-to-image generation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 6753–6762, 2024.
- [6] Marius Cordts, Mohamed Omran, Sebastian Ramos, Timo Rehfeld, Markus Enzweiler, Rodrigo Benenson, Uwe Franke, Stefan Roth, and Bernt Schiele. The cityscapes dataset for semantic urban scene understanding. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 3213–3223, 2016.
- [7] Chao Deng, Jinheng Xie, David Junhao Zhang, Yuchao Gu, Weijia Mao, Zechen Bai, Weihao Wang, Kevin Qinghong Lin, Zhijie Chen, Zhenheng Yang, and Mike Zheng Shou. Emerging properties in unified multimodal pretraining. *arXiv preprint arXiv:2505.14683*, 2025.
- [8] Andreas Geiger, Philip Lenz, and Raquel Urtasun. Are we ready for autonomous driving? the kitti vision benchmark suite. In *IEEE Conference on Computer Vision and Pattern Recognition*, pages 3354–3361. IEEE, 2012.
- [9] Google DeepMind. Gemini 2.5 pro, 2025.
- [10] Akshay Krishnan, Abhijit Kundu, Kevis-Kokitsi Maninis, James Hays, and Matthew Brown. OmninoCs: A unified nocs dataset and model for 3d lifting of 2d objects. In *European Conference on Computer Vision*, pages 127–145. Springer, 2024.
- [11] Nupur Kumari, Grace Su, Richard Zhang, Taesung Park, Eli Shechtman, and Jun-Yan Zhu. Customizing text-to-image diffusion with object viewpoint control. In *SIGGRAPH Asia 2024 Conference Papers*, pages 1–13, 2024.
- [12] Dongyang Liu, Shitian Zhao, Le Zhuo, Weifeng Lin, Yi Xin, Xinyue Li, Qi Qin, Yu Qiao, Hongsheng Li, and Peng Gao. Lumina-mgpt: Illuminate flexible photorealistic text-to-image generation with multimodal generative pretraining. *arXiv preprint arXiv:2408.02657*, 2024.
- [13] Shilong Liu, Zhaoyang Zeng, Tianhe Ren, Feng Li, Hao Zhang, Jie Yang, Qing Jiang, Chunyuan Li, Jianwei Yang, Hang Su, et al. Grounding dino: Marrying dino with grounded pre-training for open-set object detection. In *European Conference on Computer Vision*, pages 38–55. Springer, 2024.
- [14] Wufei Ma, Luoxin Ye, Celso M de Melo, Alan Yuille, and Jieneng Chen. Spatialllm: A compound 3d-informed design towards spatially-intelligent large multimodal models. In *Proceedings of the Computer Vision and Pattern Recognition Conference*, pages 17249–17260, 2025.
- [15] Wufei Ma, Guofeng Zhang, Qihao Liu, Guanning Zeng, Adam Kortylewski, Yaoyao Liu, and Alan Yuille. Imagenet3d: Towards general-purpose object-level 3d understanding. *Advances in Neural Information Processing Systems*, 37:96127–96149, 2024.

- [16] Keanu Nichols, Nazia Tasnim, Yuting Yan, Nicholas Ikechukwu, Elva Zou, Deepti Ghadiyaram, and Bryan A Plummer. Right side up? disentangling orientation understanding in mllms with fine-grained multi-axis perception tasks. *arXiv preprint arXiv:2505.21649*, 2025.
- [17] OpenAI. Introducing 4o image generation, 2025.
- [18] Maxime Oquab, Timothée Darcet, Théo Moutakanni, Huy Vo, Marc Szafraniec, Vasil Khalidov, Pierre Fernandez, Daniel Haziza, Francisco Massa, Alaaeldin El-Nouby, et al. Dinov2: Learning robust visual features without supervision. *arXiv preprint arXiv:2304.07193*, 2023.
- [19] Rishabh Parihar, Vaibhav Agrawal, Sachidanand VS, and Venkatesh Babu Radhakrishnan. Compass control: Multi object orientation control for text-to-image generation. In *Proceedings of the Computer Vision and Pattern Recognition Conference*, pages 2791–2801, 2025.
- [20] Zhenyuan Qin, Xincheng Shuai, and Henghui Ding. Scenedesigner: Controllable multi-object image generation with 9-dof pose manipulation. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems*.
- [21] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. Learning transferable visual models from natural language supervision. In *International Conference on Machine Learning*, pages 8748–8763. PmLR, 2021.
- [22] Leo Sampaio Ferraz Ribeiro, Tu Bui, John Collomosse, and Moacir Ponti. Scene designer: a unified model for scene search and synthesis from sketch. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 2424–2433, 2021.
- [23] Quan Sun, Yufeng Cui, Xiaosong Zhang, Fan Zhang, Qiying Yu, Yueze Wang, Yongming Rao, Jingjing Liu, Tiejun Huang, and Xinlong Wang. Generative multimodal models are in-context learners. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 14398–14409, 2024.
- [24] Chameleon Team. Chameleon: Mixed-modal early-fusion foundation models. *arXiv preprint arXiv:2405.09818*, 2024.
- [25] Zehan Wang, Ziang Zhang, Tianyu Pang, Chao Du, Hengshuang Zhao, and Zhou Zhao. Orient anything: Learning robust object orientation estimation from rendering 3d models. *arXiv preprint arXiv:2412.18605*, 2024.
- [26] Zehan Wang, Ziang Zhang, Tianyu Pang, Chao Du, Hengshuang Zhao, and Zhou Zhao. Orient anything: Learning robust object orientation estimation from rendering 3d models. In *International Conference on Machine Learning*, pages 65556–65574. PMLR, 2025.
- [27] Zehan Wang, Ziang Zhang, Jiayang Xu, Jialei Wang, Tianyu Pang, Chao Du, Hengshuang Zhao, and Zhou Zhao. Orient anything v2: Unifying orientation and rotation understanding. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems*.
- [28] Zehan Wang, Ziang Zhang, Jiayang Xu, Jialei Wang, Tianyu Pang, Chao Du, Hengshuang Zhao, and Zhou Zhao. Orient anything v2: Unifying orientation and rotation understanding. *arXiv preprint arXiv:2601.05573*, 2026.
- [29] Chengyue Wu, Xiaokang Chen, Zhiyu Wu, Yiyang Ma, Xingchao Liu, Zizheng Pan, Wen Liu, Zhenda Xie, Xingkai Yu, Chong Ruan, and Ping Luo. Janus: Decoupling visual encoding for unified multimodal understanding and generation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 12966–12977, 2025.
- [30] Jinheng Xie, Weijia Mao, Zechen Bai, David Junhao Zhang, Weihao Wang, Kevin Qinghong Lin, Yuchao Gu, Zhijie Chen, Zhenheng Yang, and Mike Zheng Shou. Show-o: One single transformer to unify multimodal understanding and generation. *arXiv preprint arXiv:2408.12528*, 2024.
- [31] Jinheng Xie, Weijia Mao, Zechen Bai, David Junhao Zhang, Weihao Wang, Kevin Qinghong Lin, Yuchao Gu, Zhijie Chen, Zhenheng Yang, and Mike Zheng Shou. Show-o: One single transformer to unify multimodal understanding and generation. In *International Conference on Learning Representations*, 2025.

- [32] Chunting Zhou, Lili Yu, Arun Babu, Kushal Tirumala, Michihiro Yasunaga, Leonid Shamis, Jacob Kahn, Xuezhe Ma, Luke Zettlemoyer, and Omer Levy. Transfusion: Predict the next token and diffuse images with one multi-modal model. In *International Conference on Learning Representations*, 2025.

A Orientation to viewpoint description

To bridge the gap between continuous object poses and the abstract spatial reasoning capability learned by MLLMs, we introduce an orientation-to-language abstraction as intermediate supervision for our framework. Specifically, following the Orient-Anything convention, each object orientation is represented by three angles: azimuth, polar angle, and in-plane rotation. We quantize these continuous angles and map them to coarse semantic view descriptions: (i) Azimuth: front, front-right quarter, right side, back-right quarter, back, back-left quarter, left side, and front-left quarter; (ii) Polar: overhead, high-angle, mid-high, eye-level, mid-low, low-angle, and bottom-up; (iii) Rotation: almost no in-plane rotation and clockwise/counterclockwise rotation. As quantized abstractions of object orientation, these terms are combined with object-level spatial descriptions to express 3D object poses in a linguistically accessible form. The detailed mapping relationship between the orientation angles and semantic view terms is listed in Table 3.

Table 3: Mapping from continuous orientation angles to viewpoint descriptions.

Angle type	Range	Language description
Azimuth ϕ_i	$[337.5^\circ, 360^\circ) \cup [0^\circ, 22.5^\circ)$	front
	$[22.5^\circ, 67.5^\circ)$	front-right quarter
	$[67.5^\circ, 112.5^\circ)$	right side
	$[112.5^\circ, 157.5^\circ)$	back-right quarter
	$[157.5^\circ, 202.5^\circ)$	back
	$[202.5^\circ, 247.5^\circ)$	back-left quarter
	$[247.5^\circ, 292.5^\circ)$	left side
	$[292.5^\circ, 337.5^\circ)$	front-left quarter
Polar θ_i	$[0^\circ, 15^\circ)$	overhead
	$[15^\circ, 45^\circ)$	high-angle
	$[45^\circ, 75^\circ)$	mid-high
	$[75^\circ, 105^\circ)$	eye-level
	$[105^\circ, 135^\circ)$	mid-low
	$[135^\circ, 165^\circ)$	low-angle
	$[165^\circ, 180^\circ]$	bottom-up
Rotation ψ_i	$[0^\circ, 15^\circ) \cup [345^\circ, 360^\circ)$	almost no in-plane rotation
	$[15^\circ, 180^\circ)$	counterclockwise rotation
	$[180^\circ, 345^\circ]$	clockwise rotation

B Dataset details

B.1 Dataset statistics.

UniSpatial-80K contains 83,252 image entries and 91,392 annotated objects from 122 categories, with an average of 1.10 objects per image. It integrates street-view, indoor object-centric, and general rigid-object sources, covering both single-object and multi-object spatial reasoning scenarios. The dataset exhibits a broad but moderately long-tailed category distribution: the top eight categories account for 55.16% of object instances, while the remaining categories provide diverse object geometry and scene context.

B.2 Prompt template

The prompt used to enable general models to estimate object orientation includes definitions of both the task and the orientation representation.

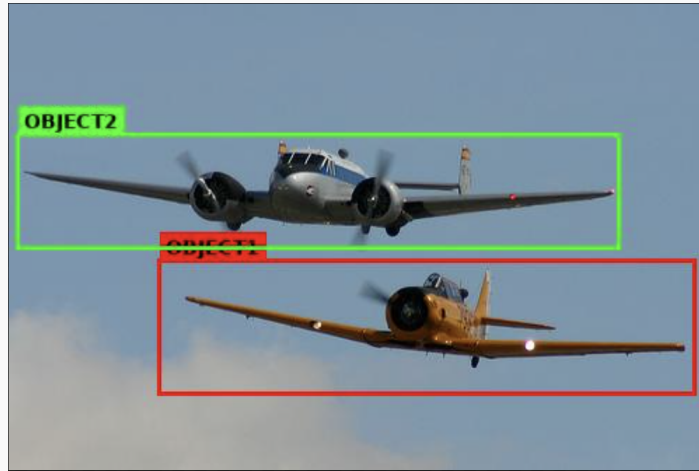
Table 4: Source and split composition of UniSpatial-80K.

Source	Train / Val	Test-Single	Test-Multi	File-level Total
KITTI-Cityscapes	1,266	363	303	1,932
ImageNet3D	39,765	1,229	146	41,140
Objectron	42,221	2,500	475	45,196
Total	83,252	4,092	924	88,268

Table 5: Top categories in UniSpatial-80K ranked by object count.

Category	Obj Count	Obj %	Img Count	Img Coverage
shoes	11,962	13.09	6,981	8.39
chair	6,058	6.63	5,579	6.70
bottle	5,744	6.29	5,726	6.88
laptop	5,585	6.11	5,479	6.58
cup	5,582	6.11	5,427	6.52
camera	5,411	5.92	5,325	6.40
books	5,050	5.53	5,004	6.01
cereal box	5,009	5.48	5,000	6.01
car	4,096	4.48	3,688	4.43

Instruction for universal orientation prediction



You are evaluating 3D object orientation in an annotated image.

The image contains bounding boxes and labels OBJECT1, OBJECT2, etc. Estimate orientation only for the listed objects, in exactly the listed order.

Objects:

- OBJECT1: category="aeroplane", original_idx=1, bbox=[108, 184, 487, 277]
- OBJECT2: category="aeroplane", original_idx=2, bbox=[7, 94, 433, 173]

Use the Orient-Anything-style angle convention:

1. Azimuth is the horizontal viewpoint angle around the object, in degrees within [0, 360).
 - 0 means front view of the object.
 - 90 means right-side view.
 - 180 means back view.
 - 270 means left-side view.
 - Front-left quarter view should be near 315 or 300; front-right quarter view should be near 45 or 60.

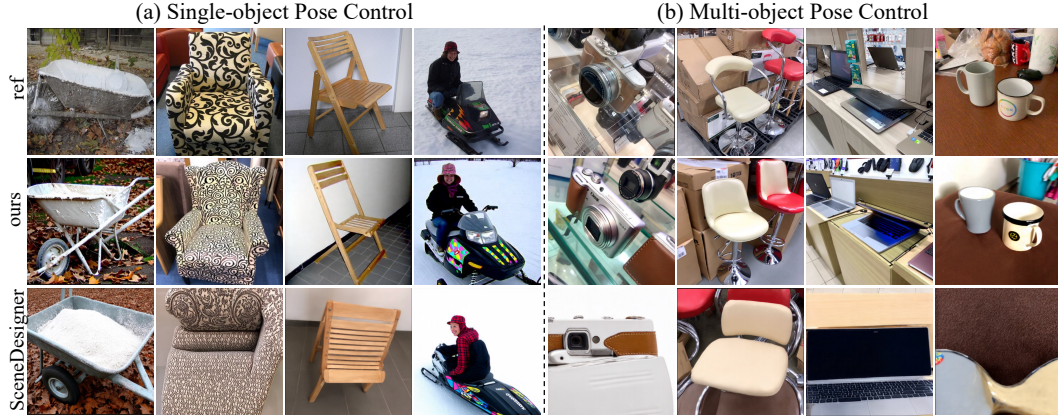


Figure 4: Additional results of pose-controllable generation in single- and multi-object scenarios. The reference image in the first row serves as the source for the input text and 9D pose annotations.

2. Polar is the vertical viewpoint angle, in degrees within $[0, 180]$.
 - 0 means top/overhead view.
 - 90 means eye-level/horizontal view.
 - 180 means bottom/upward view.
3. Rotation is the in-plane image rotation, in degrees within $[0, 360)$.
 - 0 means upright in the image plane.
 - 90 means clockwise rotation.
 - 180 means upside-down.
 - 270 means counter-clockwise rotation.

Important rules:

- Estimate the physical object viewpoint, not the 2D box position.
- Use the annotated labels in the image to identify objects.
- If the object is symmetric or ambiguous, still give the best single estimate.
- Return exactly 2 objects.
- Return only valid JSON. No markdown. No explanation.

Required JSON schema:

```
{"objects": [{ "name": "OBJECT1", "azimuth": 0, "polar": 90, "rotation": 0}]}
```

C Additional results

The additional results of pose-controllable generation are shown in Figure 4.