

Global Pose Control for Generative View Synthesis in Normalized Object Coordinate Space

Zhibing Li^{1*}, Amogh Gupta², Behnoosh Parsa², and Dan Casas²

¹ The Chinese University of Hong Kong, China

² Amazon, USA

<https://lizb6626.github.io/GlobalNVS/>

Abstract. Novel View Synthesis (NVS) enables the generation of unseen views of a scene from a single or multiple images, allowing users to freely explore an object from any viewpoint. Despite the recent impressive qualitative improvements of generative models for this task, existing methods struggle to provide global and intuitive control of target viewpoints because they either use input-relative camera poses or are limited to generating sparse global views. This lack of global pose control severely limits the number of downstream tasks potentially enabled by NVS. To address this limitation, we propose a novel approach for precise camera control in a customizable Normalized Object Coordinate Space (NOCS), requiring single or few unposed images. Our method operates solely on the absolute camera pose of the target view in NOCS, eliminating the need for a relative world frame or camera poses of the input images. Unlike previous methods that treat NVS as a standalone generation task, we formulate it as an image editing problem and build upon state-of-the-art editing models to leverage their superior generalization capability. Camera information is injected as dedicated camera tokens via an in-context multi-modal conditioning strategy. To alleviate the inherent ambiguity of NOCS, we incorporate text descriptions that explicitly define the object’s canonical coordinate frame, which also enhances generalization to unseen object categories. Furthermore, we curate a high-quality dataset with consistently aligned orientations and corresponding NOCS text definitions. Extensive experiments demonstrate that our method robustly generates novel views with accurate and consistent orientations from arbitrary unposed images across diverse categories, achieving state-of-the-art image quality and fidelity.

Keywords: Novel View Synthesis · Camera Pose Control

1 Introduction

Recent advances in generative diffusion models have made novel view synthesis (NVS) tasks from single or multiple images both accessible and effective [47, 56, 63, 78]. This progress in unseen image synthesis is enabling rapid construction

* Work done while Zhibing Li was intern at Amazon.

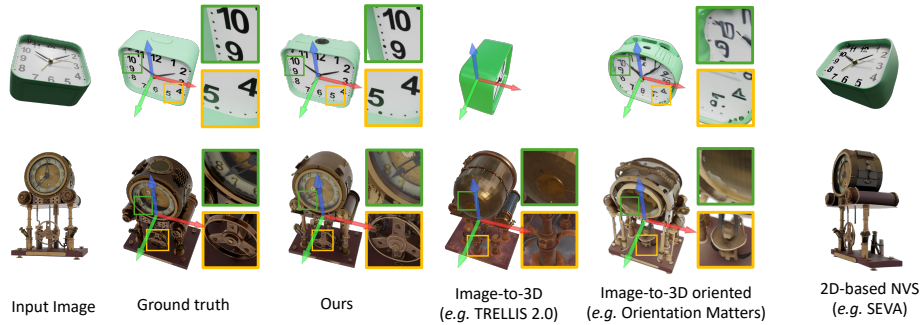


Fig. 1: Existing NVS methods fall short in providing global viewpoint control and high-quality appearance: image-to-3D methods (4th column) lack orientation awareness; orientation-aware image-to-3D (5th column) produce coarse textures; and 2D-based NVS (6th column) only provide input-relative camera control. In contrast, our method achieves both global viewpoint control and high image fidelity, with the front face consistently aligned to the canonical frame (visualized by coordinate axes).

of models for 3D reconstruction [34, 35], video manipulation [2, 43, 62, 74, 75], content creation, robotic simulation [45], and beyond.

The fundamental—and arguably most desired—feature in NVS models is their ability to provide a *global* and *intuitive* control of the target viewpoint, however, existing models either use input-relative viewpoint control [24, 31, 34, 51, 54, 72, 78] or are limited to generating sparse global views [27]. This prevents NVS models from being used for practical applications at scale, because the user cannot automatically assume a standard *canonical* coordinate frame with respect to which synthesized images are aligned to. Image-to-3D models [29, 30, 35, 53, 71] offer an attractive alternative since the output 3D mesh can be rendered from an arbitrary viewpoint, however they often neglect the mesh orientation alignment altogether, mostly due to inconsistencies in large-scale 3D datasets [9, 71], resulting in models that lack standard canonical orientations. Consequently, output 3D meshes cannot be used for tasks that require object-centric global viewpoint control.

Recent methods [19, 60] have focused on orientation alignment by fine-tuning image-to-3D models using carefully aligned 3D training sets. While they alleviate the alignment issue to some extent, they remain fundamentally constrained by the texture quality of 3D generation, which still falls far behind that of state-of-the-art 2D image synthesis, resulting in outputs that lack fine appearance details and high-fidelity visual quality.

To mitigate these limitations in NVS viewpoint control, we propose a novel 2D diffusion-based approach for global and intuitive 3D viewpoint control. We formulate the NVS task in Normalized Object Coordinate Space (NOCS) [58], directly conditioning on absolute camera poses in NOCS and thereby eliminating the need for selecting a reference view or constructing a relative world frame from input camera poses. Unlike prior methods that treat NVS as a standalone

generation task, we leverage a state-of-the-art instruction-based image editing model [67] to build an orientation-aware NVS system, inheriting its superior image quality and generalization capabilities. Specifically, we represent the target NOCS camera pose using spatially-aligned Plücker raymaps and inject them as camera tokens through an in-context learning strategy with a novel *regional attention mechanism*. This mechanism ensures that each target camera token attends exclusively to its corresponding target image tokens, preventing cross-view interference. Since our training data is orientation-aligned, this in-context regional attention enables efficient and consistent viewpoint control across generated views while preserving the advanced 2D image synthesis capabilities of the underlying diffusion model.

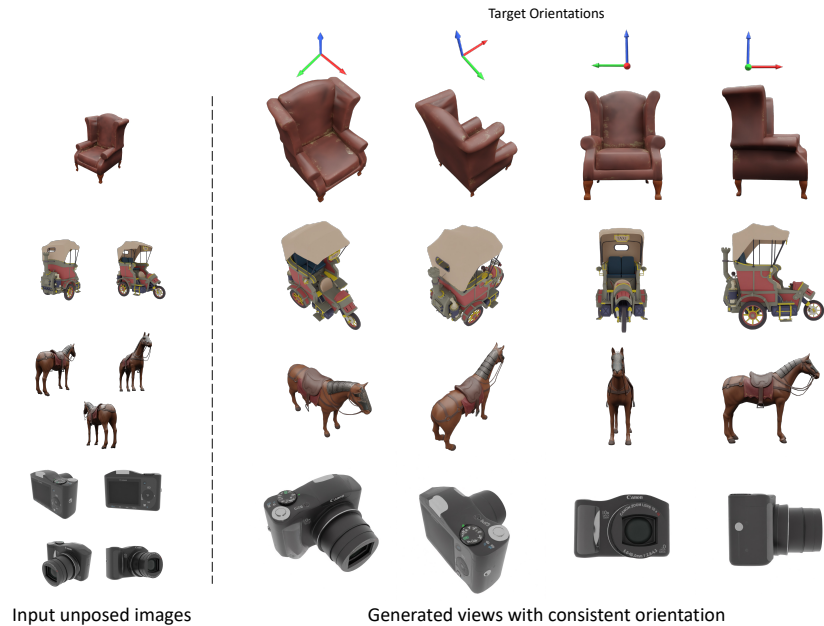


Fig. 2: Given input images captured from arbitrary viewpoints (left) and target camera poses (top-right), our model generates high-quality oriented novel views. Notice how the generated images maintain correct global orientation across different object families (e.g., the front of the object aligned with the red axis).

Furthermore, we leverage the native text-conditioning capabilities of the underlying editing model to explicitly define the canonical orientation of the input object via natural language (e.g., “the front of a camera faces the lens side”). This serves a dual purpose: (1) it disambiguates the NOCS definition, as the canonical frame and the target camera pose jointly determine the output viewpoint; and (2) it enables robust generalization to unseen object categories at test

time, by allowing users to specify arbitrary canonical orientations through text prompts without retraining.

We train our method on orientation-aligned 3D datasets, introducing a newly curated dataset of 22K objects from Objaverse with no overlap with existing resources, alongside with Objaverse-OA [36] (14K objects) and ABO [7] (7K objects). Importantly, we augment each object across all datasets with text captions that describe both the object and its front-view definition, enabling language-guided viewpoint control to foster future research.

Our experiments demonstrate that, thanks to our in-context regional attention mechanism and the use of orientation-aligned data, our model achieves state-of-the-art appearance quality for NVS tasks while significantly improving global viewpoint control compared to existing baselines.

2 Related Work

Novel view synthesis (NVS) is a foundational computer vision problem. Classical image-based rendering [14, 25] achieved photorealism but required dense multi-view capture. NVS approaches fall into two paradigms: learning intermediate 3D representations for rendering, or direct 2D view generation.

Learning 3D Representations. Generative modeling using GAN [13] was extended to learning 3D representations [3, 68]. Neural Radiance Fields (NeRF) [37] and subsequent acceleration techniques [39] advanced 3D reconstruction, though per-scene optimization requires dense views. Feed-forward methods like pixel-NeRF [73] introduced generalizable priors for sparse-view settings, while 3D Gaussian Splatting (3DGS) [23] enabled real-time rendering, spawning efficient sparse-view regressors [4, 6, 52]. The success of diffusion models [16, 48] catalyzed a new era of 3D generation, from Score Distillation Sampling [41, 57, 69] to multi-view consistent diffusion models [34, 35, 46, 47, 59]. Concurrently, Large Reconstruction Models [17] directly output 3D representations, and recent 3D native approaches like Trellis and Hunyuan3D [53, 71] encapsulate 3D priors entirely within latent spaces, bypassing intermediate multi-view generation.

Canonical Alignment in 3D Generation: A persistent challenge in 3D generative modeling is the lack of a consistent canonical coordinate space. Inspired by single-image 6DOF pose estimation [5, 40, 50, 58, 60, 65], Orientation Matters [36] addresses this by leveraging VLMs to align training data and fine-tuning 3D generators to output canonically oriented objects. However, the visual fidelity remains bottlenecked by the expressivity of the underlying 3D generator.

Direct Generative View Synthesis. Instead of explicitly modeling 3D space, several methods synthesize novel views directly via large transformers [22, 44], image-conditioned diffusion models [12, 32, 51, 54], or video diffusion models adapted for camera control [2, 15, 43, 55, 56, 62]. A critical limitation unites these approaches: they parameterize the target camera pose strictly relative to the input image, failing to establish a canonical coordinate system across semantically similar objects. Our work addresses this limitation by introducing global viewpoint control through explicit camera conditioning in a normalized object

coordinate space. Unlike relative pose parameterization, our approach establishes a shared canonical frame across instances, combining the flexibility of 2D diffusion models with geometric consistency for global pose control.

Datasets. Progress in novel view synthesis and 3D generation has been tightly coupled with 3D object datasets. Google Scanned Objects (GSO) [10] provides 1,030 high-fidelity scanned household items, Toys4K [49] offers 4,000 synthetic objects, while Objaverse [9] and Objaverse-XL [8] scaled 3D data to 800K+ and 10M+ objects, respectively. To address inconsistent orientations in Objaverse, [36] proposed Objaverse-OA, filtering 14K orientation-aligned samples from Objaverse-LVIS. However, it still contains objects with imprecise orientations, ambiguous canonical views, or overly simplified textures. To complement this effort, we curate an additional 22K high-quality objects from Objaverse with no overlap with Objaverse-OA, each with a clear front face, consistent orientation, and detailed textures.

3 Global Pose Control

3.1 Problem Definition

Given M arbitrary input view images $\mathbf{I}^{\text{src}} \in \mathbb{R}^{M \times H \times W \times 3}$, N target camera poses $\mathbf{\Pi}^{\text{tgt}} \in \text{SE}(3)^N$ defined in Normalized Object Coordinate Space (NOCS), and a text description $\mathbf{y}$ that specifies how the object is oriented within NOCS, our goal is to learn the conditional distribution of N target views $\mathbf{I}^{\text{tgt}}$:

$$\mathbf{I}^{\text{tgt}} \sim p_{\theta}(\mathbf{I}^{\text{tgt}} \mid \mathbf{I}^{\text{src}}, \mathbf{\Pi}^{\text{tgt}}, \mathbf{y}), \quad (1)$$

where p_{θ} is parameterized by a diffusion model.

We formulate this as an instruction-based image editing problem, as NVS can be naturally viewed as a spatial transformation of the input image. Accordingly, we build our method upon a state-of-the-art image editing model [67] to inherit its strong image manipulation capabilities. However, as we demonstrate in our ablation study (Section 4.4), relying solely on text instructions is insufficient for robust *global* viewpoint control, due to the inherent complexity and ambiguity of specifying precise 3D camera poses through language alone.

Figure 3 and the remainder of this section describe our architecture, which introduces: in-context camera conditioning tailored for global camera control (Section 3.2), an efficient and robust training scheme (Section 3.3), and a new aligned dataset (Section 3.4).

3.2 In-context Camera Conditioning

We introduce camera conditioning through an in-context learning method [20], which has been proven effective for DiT-based conditional generation. We describe below our camera representation, how we extend the architecture to multi-view generation, and how we ensure correct camera-view association.

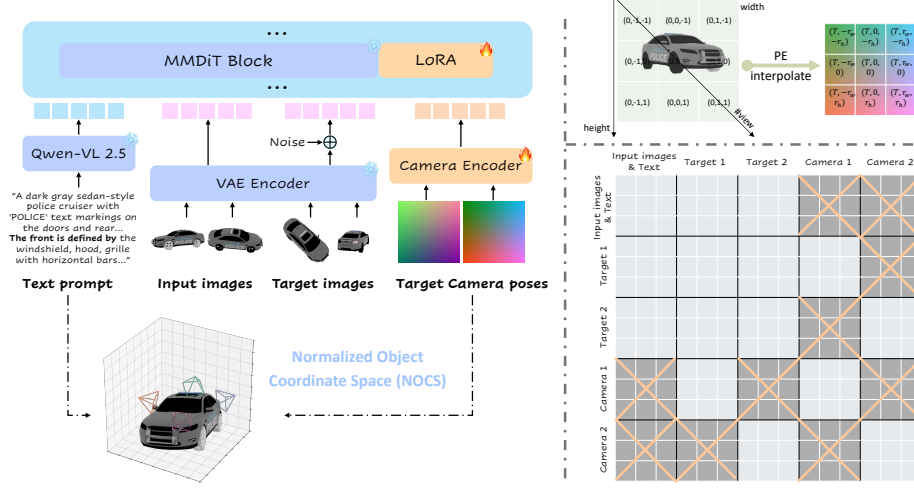


Fig. 3: Overview. *Left:* Target camera poses are encoded as Plücker ray map tokens and fed into a LoRA-adapted MMDiT alongside image tokens. A text prompt defines the object’s canonical orientation in NOCS. *Top-right:* Target images and ray map tokens are packed along the frame axis. The position embeddings of ray map tokens are interpolated to spatially align with the corresponding image tokens. *Bottom-right:* Regional attention ensures each camera token attends only to its corresponding target view, while all image tokens retain full mutual attention.

Camera Representation. We represent each target camera pose as a Plücker ray map, encoding camera rays as pixel-aligned 6D vectors that jointly capture 2D spatial correspondence and 3D ray geometry. Compared to attention-level relative encodings [24, 26, 38], this absolute, pixel-aligned representation naturally suits our NOCS-based formulation, where target cameras are specified in a global coordinate frame. Each ray map is projected into token space via a linear layer, producing camera tokens that are appended to the image tokens along the frame dimension and processed jointly by the MMDiT. We fine-tune only lightweight LoRA [18] adapters on the MMDiT layers to efficiently adapt the pretrained model while preserving its image manipulation capabilities.

Frame-axis Multi-view Packing. To generate multiple target views simultaneously, existing methods spatially concatenate outputs into a single wide image [11, 28], producing target tokens of shape $H \times NW$ against reference tokens of shape $H \times W$. This resolution mismatch disrupts the pixel-level spatial correspondence that attention mechanisms rely on. Instead, we leverage the 3D RoPE employed in Qwen-Image-Edit, which encodes positional information along three dimensions: frame, height, and width. In the original editing framework, the frame dimension serves to distinguish between the images before and after editing. We naturally extend this design by packing all target views, to-

gether with their corresponding ray maps, as separate entries along the *frame* axis while preserving the original $H \times W$ aspect ratio for each view. This naturally reuses the base model’s frame dimension and maintains consistent spatial alignment between reference and target tokens.

Compact Ray Maps with RoPE Interpolation. Although Plücker ray maps are pixel-aligned by construction, camera pose is inherently a low-dimensional signal that does not require pixel-level granularity. Encoding ray maps at full image resolution (*e.g.*, 1024×1024) would introduce unnecessarily long token sequences without commensurate benefit. We therefore fix the ray map resolution to 256×256 regardless of the target image resolution. To maintain spatial alignment with higher-resolution image tokens under the 3D RoPE, we adopt a position-aware interpolation strategy inspired by [77]. As illustrated in the top-right of Figure 3, a ray map token at grid position (i, j) is assigned RoPE coordinates $(i \cdot r_h, j \cdot r_w)$, where $r_h = H/256$ and $r_w = W/256$. This ensures that spatially co-located ray map and image tokens receive matching positional encodings, preserving their correspondence at reduced computational cost.

Regional Attention. With frame-axis packing, all entries, including multiple target views, reference images, and their corresponding ray maps, share the same $H \times W$ spatial grid and differ only along the frame dimension. This introduces an unintended side effect in the 3D RoPE attention: tokens from different entries that occupy the same spatial position have identical height and width indices, making their relative distance determined solely by the frame offset. For example, with two target views and two camera ray maps packed as $[\mathbf{z}_1^{\text{tgt}}, \mathbf{z}_2^{\text{tgt}}, \mathbf{z}_1^{\text{cam}}, \mathbf{z}_2^{\text{cam}}]$, a patch in $\mathbf{z}_1^{\text{cam}}$ may be closer in frame index to $\mathbf{z}_2^{\text{tgt}}$ than to $\mathbf{z}_1^{\text{tgt}}$, causing stronger attention between $\mathbf{z}_1^{\text{cam}}$ and the wrong target view and violating the intended camera-view correspondence. We resolve this with a regional attention mask, as shown in Figure 3 (bottom-right), that enforces correct camera-view binding: each camera token attends *only* to its corresponding target view, and conversely, each target view is restricted to attend *only* to its matched camera tokens among all camera tokens. All image tokens, both reference and target, retain full mutual attention, enabling multi-view information exchange while preventing cross-view camera leakage. This design also saves computation by masking out unrelated camera-view attention pairs.

3.3 Efficient and Robust Training

View Sampling. Each training sample consists of $T = M + N$ views in total, where M is the number of input views and N is the number of target views. Rather than fixing these counts, we dynamically sample them per iteration to expose the model to diverse task configurations and difficulty levels.

We first sample the total number of views T from a power-weighted distribution over $\{T_{\min}, \dots, T_{\max}\}$:

$$p(T=t) \propto (t - T_{\min} + \epsilon)^k, \quad (2)$$

where $k=1.5$ and $\epsilon=1$. This biases sampling toward larger T , presenting the model with more challenging multi-view configurations.

Given T , we then sample the number of input views M from a truncated exponential distribution over $\{1, \dots, T-1\}$, following [33]:

$$p(M=m) \propto e^{-\lambda m}, \quad \lambda = \ln 2, \quad (3)$$

which halves the probability as m increases. Since more reference views make the task increasingly constrained and easier, this biases toward fewer input views, encouraging the model to learn stronger priors from limited observations.

To ensure balanced workload across GPUs, we bucket training samples by their (M, N) configuration and enforce a uniform (M, N) assignment within each iteration, so that all devices process sequences of equal length.

Canonical Front View Anchoring. Our text prompt explicitly defines the front-facing orientation of each object in NOCS. To ground this description into a concrete visual reference, we include the canonical front view among the generated targets. This front view is generated jointly with all other target views in a single diffusion process, rather than as a separate preliminary step. Nevertheless, it encourages the model to first resolve the object’s canonical orientation before reasoning about the remaining viewpoints, acting as an implicit intermediate step akin to chain-of-thought prompting [64, 66]. As shown in our ablation (Section 4.4), this simple design notably improves pose accuracy for geometrically challenging objects (*e.g.*, narrow or thin shapes such as swords).

Loss Function. We adopt a flow matching objective for training. Given latent $x_0 = \mathcal{E}(\mathbf{I}^{\text{tgt}})$ encoded by the VAE and noise $x_1 \sim \mathcal{N}(0, \mathbf{I})$, the intermediate latent at timestep t is $x_t = tx_0 + (1-t)x_1$, with velocity $v_t = x_0 - x_1$. The model is trained to predict this velocity via:

$$\mathcal{L} = \mathbb{E}_{(x_0, h) \sim \mathcal{D}, x_1, t} \|v_\theta(x_t, t, h) - v_t\|^2, \quad (4)$$

where h denotes conditioning from input images, text, and camera poses.

3.4 Dataset

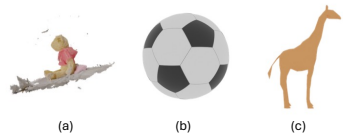


Fig. 4: Common failure modes in Objaverse-OA. (a) objects with imprecise orientations (*e.g.*, tilted or rotated off-axis), (b) objects with no recognizable or ambiguous front views, and (c) objects with over-simplified textures.

Our approach requires canonically aligned, high-quality 3D objects with front-face annotations at training time. Unfortunately, no existing dataset fully satisfies these requirements: Objaverse-OA [36] filters 14K aligned objects from Objaverse-LVIS, but we observe that many samples are tilted (yaw or pitch offset), lack a valid front face, or exhibit over-simplified textures (Figure 4). ABO [7] provides 8K models in a canonical coordinate space, but its categories are limited to household objects. We therefore propose a robust curation pipeline to assemble a larger and cleaner collection from Objaverse [9], enriched

with textual front-face definitions. To ensure high quality and canonical alignment, we combine Orient-Anything-V2 [61] with VLM-based [1] verification, prioritizing strict filtering to minimize erroneous samples while maintaining category diversity.

We start from the Objaverse subset of HY3D-Bench [21], which provides 70K objects filtered for high texture quality from Objaverse [9]. However, HY3D-Bench does not enforce canonical orientation or guarantee a valid front face. We therefore apply a multi-stage curation pipeline on top of this pool.

We first remove tilted and skewed objects using Orient-Anything-V2 [61]. Specifically, we retain only objects where the predicted elevation degree is less than 8° and the azimuth degree falls within a 5° tolerance of the canonical angles (0° , 90° , 180° , or 270°). This ensures that objects are properly upright and aligned to the principal axes. To further refine our selection, we employ a VLM to verify whether objects are properly placed or tilted, providing an additional layer of quality control beyond geometric constraints alone.

We then filter objects based on the number of valid of *frontal* faces, following the definition and implementation of Orient-Anything-V2 [61]. Objects with 0 valid faces (indicating no recognizable canonical views, for example a soccer ball) or 4 valid faces (suggesting ambiguous or symmetric geometry, for example a rectangle basket) are excluded. The resulting dataset consists of objects with 1 or 2 valid faces, which provide clear canonical orientations suitable for training.

For front face annotation, Orient-Anything-V2 successfully labels approximately 21K objects by selecting the view among the four horizontal candidates with azimuth closest to 0° and within a 5° tolerance. For the remaining 1K challenging cases, such as elongated objects like swords that may have two equally valid front faces, we employ VLM-based annotation to resolve ambiguities.

Finally, we annotate every sample with an overall object description and a textual front-face definition using the same VLM. The front-face definition specifies which side of the object constitutes the front, thereby anchoring the NOCS frame (*e.g.*, “the front is defined by the main gun barrel pointing forward”; “the front is defined by the side with the handlebars and grips facing the viewer”). Through this pipeline we curate **22K** canonically oriented objects with front-face labels from the initial 70K pool, complementing existing datasets such as Objaverse-OA and ABO while significantly broadening category coverage.

4 Experiments

4.1 Implementation Details

We build upon Qwen-Image-Edit-2509 [67] as our base model and apply LoRA [18] with rank 64 to efficiently adapt the pretrained weights. We use the AdamW optimizer with a learning rate of 1×10^{-4} and a per-GPU batch size of 1 on 8 NVIDIA H100 GPUs. We adopt a coarse-to-fine training strategy: the model is first trained at 512×512 resolution for 20k steps to learn camera control and multi-view consistency, then fine-tuned at 1024×1024 for 15k steps to recover

Table 1: Quantitative Comparison with image-to-3D methods.

	GSO [10]				Toys4k [49]			
	PSNR↑	SSIM↑	LPIPS↓	CLIP↑	PSNR↑	SSIM↑	LPIPS↓	CLIP↑
Cupid [19]	15.32	0.837	0.186	91.3	16.67	0.878	0.132	94.0
Orientation Matters [36]	16.13	0.844	0.178	89.8	16.94	0.879	0.135	91.6
TRELLIS2 [70]	15.77	0.837	0.187	90.5	15.64	0.862	0.156	91.2
HY3D-2.1 [21]	15.21	0.837	0.196	89.3	16.67	0.873	0.145	91.4
Ours	17.25	0.855	0.163	92.3	18.24	0.892	0.116	94.2

high-frequency details. Our training data comprises Objaverse-OA [36], Amazon Berkeley Objects (ABO) [7], and our custom aligned dataset (Section 3.4).

We evaluate on two publicly available benchmarks: Google Scanned Objects (GSO) [10] and Toys4k [49], selecting 100 objects with clearly identifiable front faces from each. Metrics include PSNR, SSIM, LPIPS [76] and CLIP [42] scores.

We compare against two categories of methods: (1) *image-to-3D methods*: Orientation Matters [36], Cupid [19], Hunyuan3D 2.1 [53], and Trellis-2 [70], which generate 3D assets from a single input image; and (2) *Generative NVS methods*: EscherNet [24] and SEVA [78], which we evaluate with 1, 2, 3, and 5 input views to assess performance across varying input configurations.

4.2 Comparisons with 3D methods

Table 1 presents a quantitative comparison between our method and state-of-the-art 3D generation approaches, all using a single-view image as input. Our method consistently outperforms all baselines across both benchmarks and all metrics. The improvements are particularly pronounced on Toys4k, which contains challenging categories such as vehicles and shoes with intricate geometry and texture patterns.

We attribute the limitations of existing 3D methods to two primary factors: *orientation inconsistency* and *texture degradation*. Trellis-2 and Hunyuan3D 2.1 frequently produce 3D assets with incorrect canonical orientations, leading to misaligned novel views. While Orientation Matters mitigates this by fine-tuning on orientation-aligned 3D data, and Cupid jointly predicts canonical geometry and camera pose, both still suffer from significant texture quality loss inherent to the 3D generation pipeline. By operating directly in image space, our method avoids the information bottleneck of 3D methods and preserves significantly richer visual detail, as shown in Figure 5.

4.3 Comparisons with NVS methods

We compare our method with state-of-the-art generative NVS methods, EscherNet [24] and SEVA [78], under varying numbers of input views ($M=1, 2, 3, 5$). Both methods operate in a relative coordinate frame and require known camera poses for all input images. Since our evaluation setting assumes unposed inputs,



Fig. 5: Qualitative comparison with image-to-3D methods. For each object, we show two target views: the canonical front view (left) and an arbitrary viewpoint (right). Our approach faithfully preserves fine-grained textures (e.g., clock face numerals) and structural details (e.g., shoe laces), whereas 3D methods suffer from texture loss, blurriness, or geometric distortions. Best viewed zoomed in.

Table 2: Quantitative Comparison with generative NVS methods.

		GSO				Toys4k			
	#Input Views	1	2	3	5	1	2	3	5
EscherNet [24]	PSNR $\uparrow$	14.90	17.25	17.58	17.97	15.32	15.36	15.77	16.26
	SSIM $\uparrow$	0.832	0.855	0.860	0.865	0.861	0.864	0.871	0.880
	LPIPS $\downarrow$	0.230	0.187	0.180	0.172	0.196	0.192	0.181	0.169
SEVA [78]	PSNR $\uparrow$	15.20	15.85	16.13	16.39	15.87	15.85	16.27	16.66
	SSIM $\uparrow$	0.834	0.839	0.844	0.849	0.868	0.868	0.874	0.880
	LPIPS $\downarrow$	0.188	0.181	0.176	0.170	0.147	0.152	0.143	0.134
Ours	PSNR $\uparrow$	17.25	18.11	18.52	19.07	18.24	18.77	19.17	19.68
	SSIM $\uparrow$	0.855	0.863	0.866	0.871	0.892	0.895	0.897	0.902
	LPIPS $\downarrow$	0.163	0.142	0.136	0.128	0.116	0.107	0.100	0.093

we use Orientation-Anything-V2 [61] to estimate input orientations in canonical space and convert them to the required camera poses.

Quantitative results in Table 2 show that our method consistently outperforms both baselines across all input configurations. A key reason is the *error*

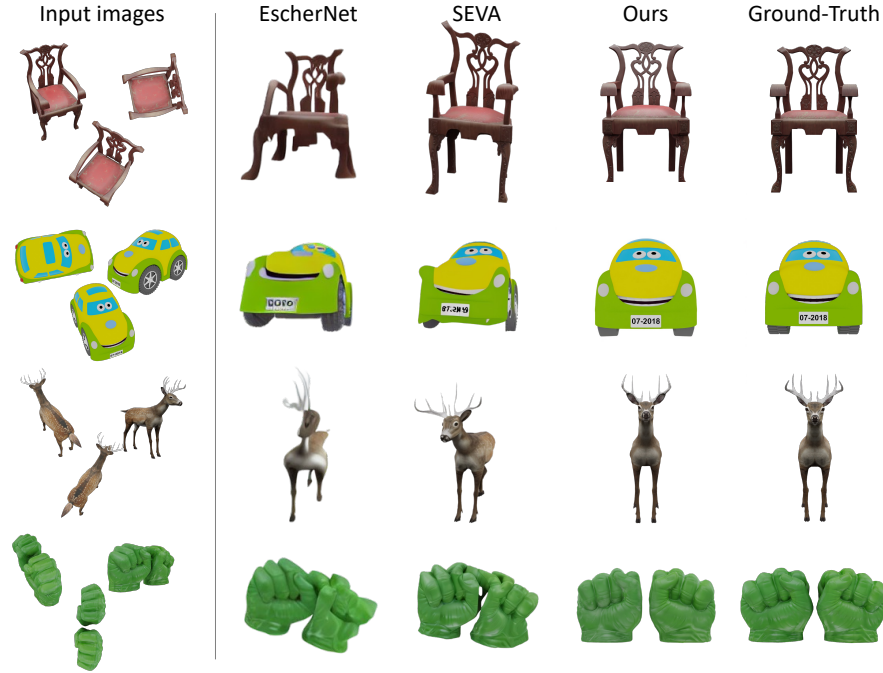


Fig. 6: Qualitative comparison with generative NVS methods. We show the generated front view for each method. Both EscherNet and SEVA produce results with noticeably incorrect orientations due to errors in the estimated input camera poses. In contrast, our method operates directly in canonical space without requiring input pose estimation, generating accurately oriented front views.

accumulation inherent to relative-pose-based methods: inaccuracies in the estimated input poses propagate directly into the generated target views, causing significant quality degradation. This issue is exacerbated by the entangled nature of Euler angles, where Orientation-Anything-V2 particularly struggles with roll estimation, compounding errors across views. By operating directly in the NOCS space, our method eliminates the dependency on input camera pose estimation entirely, making it inherently robust to this source of error.

4.4 Ablation Study

In this section we ablate different design choices and components of our method. **Camera Conditioning.** Table 3 and Figure 7 ablate camera conditioning design choices. As baseline, we use off-the-shelf Qwen-Image-Edit with structured prompts specifying NOCS definitions and camera parameters (see Supplementary). Despite explicit specification, the model fails to reliably generate multi-view outputs with correct view counts. LoRA training with structured prompts enables multi-view generation but produces substantial pose deviations. We

Table 3: Ablation study on camera conditioning.

Method	PSNR $\uparrow$	SSIM $\uparrow$	LPIPS $\downarrow$
Baseline	12.48	0.750	0.413
Baseline + LoRA	16.78	0.870	0.184
Baseline + LoRA + Camera tokens w/ Global Attention	16.87	0.870	0.180
Baseline + LoRA + Camera tokens w/ Region Attention	20.80	0.899	0.111

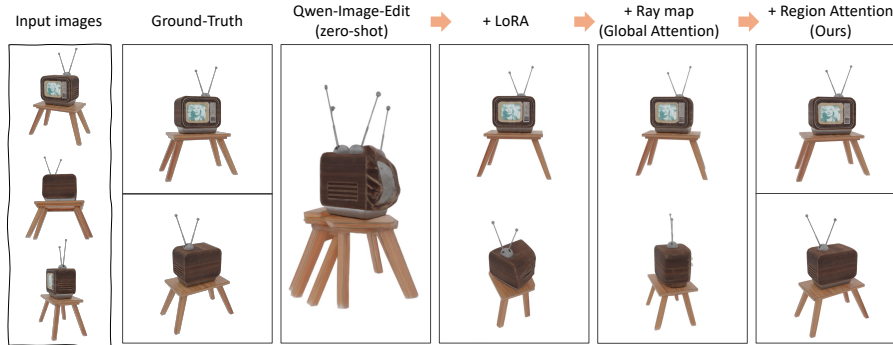


Fig. 7: Ablation study on camera conditioning. Qwen-Image-Edit cannot robustly produce multi-view layouts without fine-tuning. LoRA enables layout generation, but text-only camera descriptions lack precision. Plücker raymap tokens provide explicit geometric conditioning, yet full attention causes cross-view leakage. Regional attention restricts each raymap to its corresponding view, eliminating leakage and achieving accurate pose control.

attribute this to the mismatch between discrete text tokens and continuous SE(3) camera parameters—geometric structure is difficult to encode through text alone. Replacing text with Plücker raymaps provides geometric representation. An encoder maps raymaps to camera tokens concatenated with image tokens. However, naïve fusion fails: appending camera tokens causes cross-view entanglement and ineffective injection (Section 3.2). Regional attention explicitly binds each camera token to its spatial region, successfully injecting camera information while preventing cross-view interference. Training and evaluation use a subset of our dataset

Front Face Definition. To ablate the impact of front face definition, we train our model only on the ABO [7] dataset, which contains household objects, with and without front view definition. As shown Figure 8 (left), front view definition helps the model generalize to unseen categories such as a pigeon. Importantly, the advantages of this front face definition strategy go beyond generalization capabilities: we can also define arbitrary front views and normalized spaces as desired, and our model reacts accordingly. In Figure 8 (right), we show how redefining the front face of a dice results in correctly-aligned NVS control that

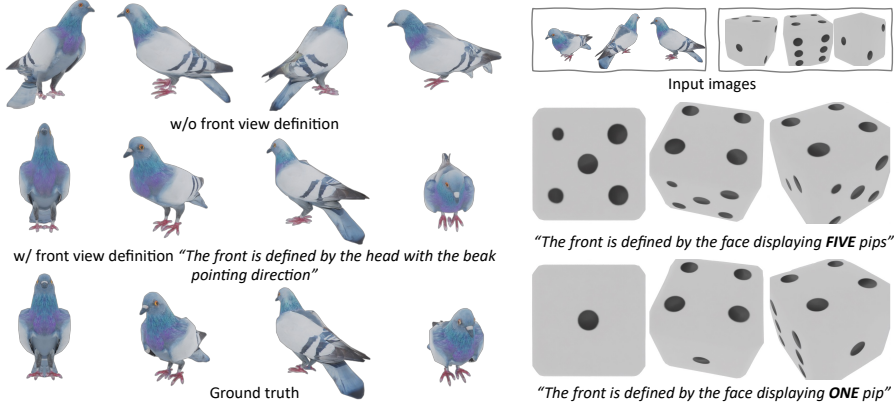


Fig. 8: Ablation study on front view definition. With front face definition, our model robustly generalizes to unseen categories (left) and enables test-time redefinition of front view (right).

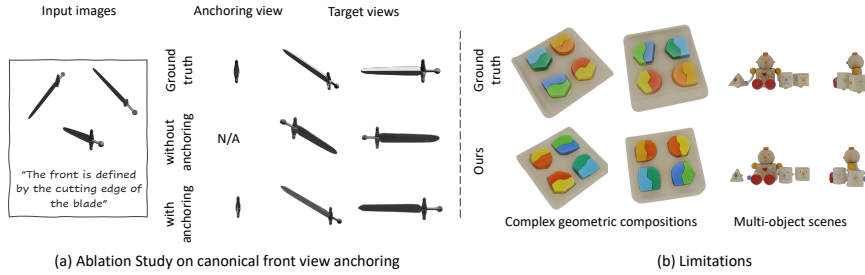


Fig. 9: Ablation study on canonical front view anchoring and Limitations.

follows the input text definition. Such test-time robust control of the global canonical frame is something existing models cannot provide.

Canonical Front View Anchoring We ablate the effect of including the canonical front view among the generated targets. As shown in Figure 9(a), this design is particularly beneficial for geometrically challenging objects such as swords, where the narrow, elongated shape makes it difficult to resolve the canonical orientation from arbitrary viewpoints alone. Without front anchoring, the model confuses the front-back orientation, producing views with incorrect pose (e.g., the blade appears flattened or reversed). By generating both the front view and the target view, the model establishes a consistent spatial reference that guides the remaining target views toward correct orientations.

5 Conclusion

We presented a novel approach for novel view synthesis with precise global view-point control in Normalized Object Coordinate Space (NOCS). By formulating NVS as an image editing problem, our method inherits the superior image quality of state-of-the-art editing models, while our in-context camera conditioning with regional attention ensures correct camera-view binding. Text-based canonical orientation descriptions disambiguate the NOCS definition and enable generalization to unseen categories without retraining. To support training, we curated a high-quality dataset of 22K canonically aligned objects with front-face text definitions, complementing existing resources. Extensive experiments demonstrate that our method achieves high image quality and accurate global viewpoint control from arbitrary unposed images, significantly outperforming both image-to-3D and generative NVS baselines.

Limitations and Future Work. Despite its strong performance, our method has several limitations. First, our current approach assumes fixed camera intrinsics across all views, limiting its flexibility for applications requiring varying focal lengths. Second, as shown in Figure 9 (left), although text descriptions help disambiguate the canonical NOCS frame, objects with complex geometric compositions remain challenging. Third, our method is primarily designed for single-object scenarios. When the input contains multiple objects with distinct orientations, performance degrades. Finally, integrating our globally controlled novel views as a drop-in module for downstream tasks such as robotic manipulation and content creation presents a promising direction for future work.

References

1. Anthropic: Claude opus 4.5. Anthropic News (2025), <https://www.anthropic.com/news/claude-sonnet-4-5>
2. Bahmani, S., Skorokhodov, I., Siarohin, A., Menapace, W., Qian, G., Vasilkovsky, M., Lee, H.Y., Wang, C., Zou, J., Tagliasacchi, A., Lindell, D.B., Tulyakov, S.: VD3d: Taming large video diffusion transformers for 3d camera control. In: International Conference on Learning Representations (ICLR) (2025), <https://openreview.net/forum?id=On4bS0R5MM>
3. Chan, E.R., Lin, C.Z., Chan, M.A., Nagano, K., Pan, B., Mello, S.D., Gallo, O., Guibas, L., Tremblay, J., Khamis, S., Karras, T., Wetzstein, G.: Efficient geometry-aware 3D generative adversarial networks. In: arXiv (2021)
4. Charatan, D., Li, S., Tagliasacchi, A., Sitzmann, V.: pixelsplat: 3d gaussian splats from image pairs for scalable generalizable 3d reconstruction. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 20912–20922 (2024)
5. Chen, X., Dong, Z., Song, J., Geiger, A., Hilliges, O.: Category level object pose estimation via neural analysis-by-synthesis. In: European Conference on Computer Vision (ECCV) (2020)
6. Chen, Y., Xu, H., Zheng, C., Zhuang, B., Pollefeys, M., Geiger, A., Cham, T.J., Cai, J.: Mvsplat: Efficient 3d gaussian splatting from sparse multi-view images. In: European Conference on Computer Vision. Springer (2024)

7. Collins, J., Goel, S., Deng, K., Luthra, A., Xu, L., Gundogdu, E., Zhang, X., Yago Vicente, T.F., Dideriksen, T., Arora, H., Guillaumin, M., Malik, J.: Abo: Dataset and benchmarks for real-world 3d object understanding. *CVPR* (2022)
8. Deitke, M., Liu, R., Wallingford, M., Ngo, H., Michel, O., Kusupati, A., Fan, A., Laforte, C., Voleti, V.S., Gadre, S.Y., VanderBilt, E., Kembhavi, A., Vondrick, C., Gkioxari, G., Ehsani, K., Schmidt, L., Farhadi, A.: Objaverse-XL: A Universe of 10M+ 3D Objects. *Conference on Neural Information Processing Systems (NeurIPS)* **36**, 35799–35813 (2023)
9. Deitke, M., Schwenk, D., Salvador, J., Weihs, L., Michel, O., VanderBilt, E., Schmidt, L., Ehsani, K., Kembhavi, A., Farhadi, A.: Objaverse: A Universe of Annotated 3D Objects. In: *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. pp. 13142–13153 (2022)
10. Downs, L., Francis, A., Koenig, N., Kinman, B., Hickman, R.M., Reymann, K., McHugh, T.B., Vanhoucke, V.: Google Scanned Objects: A High-Quality Dataset of 3D Scanned Household Items. *International Conference on Robotics and Automation (ICRA)* pp. 2553–2560 (2022)
11. Feng, J., Li, X., Lin, J., Liu, J., Liu, G., Lou, W., Ma, S., Shi, G., Wang, Q., Wang, J., Xu, Z., Yi, X., Yu, Z., Zhang, J., Zhu, Y., Chen, R., Chi, J., Du, Z., Han, L., Huang, L., Jiang, K., Li, Y., Luo, G., Wang, S., Wu, Q., Yang, F., Zhang, J., Zhang, X.: Seed3d 1.0: From images to high-fidelity simulation-ready 3d assets (2025), <https://arxiv.org/abs/2510.19944>
12. Gao*, R., Holynski*, A., Henzler, P., Brussee, A., Martin-Brualla, R., Srinivasan, P.P., Barron, J.T., Poole*, B.: Cat3d: Create anything in 3d with multi-view diffusion models. *Advances in Neural Information Processing Systems* (2024)
13. Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., Bengio, Y.: Generative adversarial networks. In: *Advances in neural information processing systems*. vol. 27, pp. 2672–2680 (2014)
14. Gortler, S.J., Grzeszczuk, R., Szeliski, R., Cohen, M.F.: The lumigraph. In: *Proceedings of the 23rd annual conference on Computer graphics and interactive techniques*. pp. 43–54 (1996)
15. He, H., Xu, Y., Guo, Y., Wetzstein, G., Dai, B., Li, H., Yang, C.: Cameractrl: Enabling camera control for text-to-video generation (2024)
16. Ho, J., Jain, A., Abbeel, P.: Denoising diffusion probabilistic models. In: *Advances in Neural Information Processing Systems*. vol. 33, pp. 6840–6851 (2020)
17. Hong, Y., Zhang, K., Gu, J., Bi, S., Yang, Y., Forsyth, D., et al.: Lrm: Large reconstruction model for single image to 3d. In: *International Conference on Learning Representations* (2024)
18. Hu, E.J., Shen, Y., Wallis, P., Allen-Zhu, Z., Li, Y., Wang, S., Wang, L., Chen, W.: LoRA: Low-rank adaptation of large language models. In: *International Conference on Learning Representations* (2022), <https://openreview.net/forum?id=nZeVKeeFYf9>
19. Huang, B., Duan, H., Zhao, Y., Zhao, Z., Ma, Y., Gao, S.: CUPID: Pose-grounded generative 3d reconstruction from a single image (2025), <https://openreview.net/forum?id=qSCjC9PFhs>
20. Huang, L., Wang, W., Wu, Z.F., Shi, Y., Dou, H., Liang, C., Feng, Y., Liu, Y., Zhou, J.: In-Context LoRA for Diffusion Transformers. *arXiv preprint arxiv:2410.23775* (2024)
21. Hunyuan3D, T., :, Zhang, B., Guo, C., Guo, D., Liu, H., Yan, H., Shi, H., Yu, J., Xu, J., Huang, J., Li, K., Wang, L., Linus, Wang, P., Lin, Q., Tang, R., Yang, X., Li, Y., Guan, Y., Zhao, Y., Yang, Y., Lai, Z., Liang, Z., Zhao, Z.: Hy3d-bench: Generation of 3d assets (2026), <https://arxiv.org/abs/2602.03907>

22. Jin, H., Jiang, H., Tan, H., Zhang, K., Bi, S., Zhang, T., Luan, F., Snavely, N., Xu, Z.: Lvsm: A large view synthesis model with minimal 3d inductive bias. In: International Conference on Learning Representations (2025)
23. Kerbl, B., Kopanas, G., Leimkühler, T., Drettakis, G.: 3d gaussian splatting for real-time radiance field rendering. *ACM Transactions on Graphics (ToG)* **42**(4), 1–14 (2023)
24. Kong, X., Liu, S., Lyu, X., Taher, M., Qi, X., Davison, A.J.: Eschnet: A generative model for scalable view synthesis. *arXiv preprint arXiv:2402.03908* (2024)
25. Levoy, M., Hanrahan, P.: Light field rendering. In: Proceedings of the 23rd annual conference on Computer graphics and interactive techniques. pp. 31–42 (1996)
26. Li, R., Yi, B., Liu, J., Gao, H., Ma, Y., Kanazawa, A.: Cameras as relative positional encoding. *arXiv preprint arXiv:2507.10496* (2025)
27. Li, X., Du, G., Zhang, R., Jin, L., Jia, Q., Lu, L., Guo, Z., Zhao, Y., Liu, H., Wang, T., Li, C., Gong, X., Li, R., Fan, B.: Droplet3D: Commonsense Priors from Videos Facilitate 3D Generation. *arXiv preprint arXiv:2508.20470* (2025)
28. Liang, Y., Luo, K., Chen, X., Chen, R., Yan, H., Li, W., Liu, J., Tan, P.: Unitek: Universal high fidelity generative texturing for 3d shapes. *arXiv preprint arXiv:2505.23253* (2025)
29. Liu, M., Shi, R., Chen, L., Zhang, Z., Xu, C., Wei, X., Chen, H., Zeng, C., Gu, J., Su, H.: One-2-3-45++: Fast Single Image to 3D Objects with Consistent Multi-View Generation and 3D Diffusion. In: IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) (2024), <https://arxiv.org/pdf/2311.07885.pdf>
30. Liu, M., Xu, C., Jin, H., Chen, L., Varma T, M., Xu, Z., Su, H.: One-2-3-45: Any single image to 3d mesh in 45 seconds without per-shape optimization. *Conference on Neural Information Processing Systems (NeurIPS)* **36** (2024)
31. Liu, R., Wu, R., Hoorick, B.V., Tokmakov, P., Zakharov, S., Vondrick, C.: Zero-1-to-3: Zero-shot One Image to 3D Object. In: International Conference on Computer Vision (ICCV) (2023)
32. Liu, R., Wu, R., Van Hoorick, B., Tokmakov, P., Zakharov, S., Vondrick, C.: Zero-1-to-3: Zero-shot one image to 3d object. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 9298–9309 (2023)
33. Liu, S., Ng, K.W., Jang, W., Guo, J., Han, J., Liu, H., Douratsos, Y., Pérez, J.C., Zhou, Z., Phung, C., et al.: Scaling sequence-to-sequence generative neural rendering. In: International Conference on Learning Representations (ICLR) (2026)
34. Liu, Y., Lin, C., Zeng, Z., Long, X., Liu, L., Komura, T., Wang, W.: Sync-Dreamer: Generating Multiview-consistent Images from a Single-view Image. In: International Conference on Learning Representations (ICLR) (2024), <https://openreview.net/forum?id=MN3yH2ovHb>
35. Long, X., Guo, Y.C., Lin, C., Liu, Y., Dou, Z., Liu, L., Ma, Y., Zhang, S.H., Habermann, M., Theobalt, C., et al.: Wonder3D: Single Image to 3D using Cross-Domain Diffusion. In: IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) (2024)
36. Lu, Y., Tian, Y., Jiang, Z., Zhao, Y., Yang, Y., Ouyang, H., Hu, H., Yu, H., Shen, Y., Liao, Y.: Orientation matters: Making 3d generative models orientation-aligned (2025), <https://arxiv.org/abs/2506.08640>
37. Mildenhall, B., Srinivasan, P.P., Tancik, M., Barron, J.T., Ramamoorthi, R., Ng, R.: Nerf: Representing scenes as neural radiance fields for view synthesis. In: European Conference on Computer Vision. pp. 405–421. Springer (2020)
38. Miyato, T., Jaeger, B., Welling, M., Geiger, A.: Gta: A geometry-aware attention mechanism for multi-view transformers. *arXiv preprint arXiv:2310.10375* (2023)

39. Müller, T., Evans, A., Schied, C., Keller, A.: Instant neural graphics primitives with a multiresolution hash encoding. *ACM Transactions on Graphics (ToG)* **41**(4), 1–15 (2022)
40. Peng, S., Liu, Y., Huang, Q.X., Bao, H., Zhou, X.: Pynet: Pixel-wise voting network for 6dof pose estimation. In: *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. pp. 4556–4565 (2018)
41. Poole, B., Jain, A., Barron, J.T., Mildenhall, B.: Dreamfusion: Text-to-3d using 2d diffusion. In: *International Conference on Learning Representations* (2023)
42. Radford, A., Kim, J.W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., et al.: Learning transferable visual models from natural language supervision. In: *International conference on machine learning*. pp. 8748–8763. PmLR (2021)
43. Ren, X., Shen, T., Huang, J., Ling, H., Lu, Y., Nimier-David, M., Müller, T., Keller, A., Fidler, S., Gao, J.: GEN3C: 3D-Informed World-Consistent Video Generation with Precise Camera Control. In: *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (2025)
44. Sajjadi, M.S.M., Meyer, H., Pot, E., Bergmann, U., Greff, K., Radwan, N., Vora, S., Lučić, M., Duckworth, D., Dosovitskiy, A., Uszkoreit, J., Funkhouser, T., Tagliasacchi, A.: Scene representation transformer: Geometry-free novel view synthesis through set-latent scene representations. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 16292–16301 (2022)
45. Shi, L., Shi, Y., Xu, X., Liu, T., Xi, J., Chen, C.: NVSPolicy: Adaptive Novel-View Synthesis for Generalizable Language-Conditioned Policy Learning (2025), <https://arxiv.org/abs/2505.10359>
46. Shi, R., Chen, H., Zhang, Z., Liu, M., Xu, C., Wei, X., Chen, L., Zeng, C., Su, H.: Zero123++: a Single Image to Consistent Multi-view Diffusion Base Model (2023)
47. Shi, Y., Wang, P., Ye, J., Mai, L., Li, K., Yang, X.: MVDream: Multi-view Diffusion for 3D Generation. In: *International Conference on Learning Representations (ICLR)* (2024), <https://openreview.net/forum?id=FUgrjq2pbB>
48. Song, Y., Ermon, S.: Generative modeling by estimating gradients of the data distribution. In: *Advances in Neural Information Processing Systems*. vol. 32 (2019)
49. Stojanov, S., Thai, A., Rehg, J.M.: Using Shape to Categorize: Low-Shot Learning with an Explicit Shape Bias. In: *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (2021)
50. Sun, J., Wang, Z., Zhang, S., He, X.H., Zhao, H., Zhang, G., Zhou, X.: OnePose: One-Shot Object Pose Estimation without CAD Models. In: *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. pp. 6815–6824 (2022)
51. Szymanowicz, S., Rupprecht, C., Vedaldi, A.: Viewset Diffusion: (0-)Image-Conditioned 3D Generative Models from 2D Data. In: *International Conference on Computer Vision (ICCV)*. pp. 8829–8839 (2023)
52. Szymanowicz, S., Rupprecht, C., Vedaldi, A.: Splatter image: Ultra-fast single-view 3d reconstruction. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 20892–20901 (2024)
53. Team, T.H.: Hunyuan3D 2.1: From Images to High-Fidelity 3D Assets with Production-Ready PBR Material (2025)
54. Tseng, H.Y., Li, Q., Kim, C., Alsisan, S., Huang, J.B., Kopf, J.: Consistent view synthesis with pose-guided diffusion models. In: *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (2023), <https://arxiv.org/abs/2303.17598>

55. Van Hoorick, B., Wu, R., Ozguroglu, E., Sargent, K., Liu, R., Tokmakov, P., Dave, A., Zheng, C., Vondrick, C.: Generative camera dolly: Extreme monocular dynamic novel view synthesis. In: European Conference on Computer Vision. Springer (2024)
56. Voleti, V., Yao, C.H., Boss, M., Letts, A., Pankratz, D., Tochilkin, D., Laforte, C., Rombach, R., Jampani, V.: SV3D: Novel Multi-view Synthesis and 3D Generation from a Single Image using Latent Video Diffusion. In: European Conference on Computer Vision (ECCV) (2024)
57. Wang, H., Du, X., Li, J., Yeh, R.A., Shakhnarovich, G.: Score jacobian chaining: Lifting pretrained 2d diffusion models for 3d generation. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 12619–12629 (2023)
58. Wang, H., Sridhar, S., Huang, J., Valentin, J., Song, S., Guibas, L.J.: Normalized object coordinate space for category-level 6d object pose and size estimation. In: IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 2642–2651 (2019)
59. Wang, P., Shi, Y.: Imagedream: Image-prompt multi-view diffusion for 3d generation. arXiv preprint arXiv:2312.02201 (2023)
60. Wang, Z., Zhang, Z., Pang, T., Du, C., Zhao, H., Zhao, Z.: Orient Anything: Learning Robust Object Orientation Estimation from Rendering 3D Models. In: International Conference on Machine Learning (ICML) (2025), <https://openreview.net/forum?id=x4yTgv2WkJ>
61. Wang, Z., Zhang, Z., Xu, J., Wang, J., Pang, T., Du, C., Zhao, H., Zhao, Z.: Orient anything v2: Unifying orientation and rotation understanding. In: The Thirty-ninth Annual Conference on Neural Information Processing Systems
62. Wang, Z., Yuan, Z., Wang, X., Li, Y., Chen, T., Xia, M., Luo, P., Shan, Y.: Motionctrl: A unified and flexible motion controller for video generation. In: ACM SIGGRAPH 2024 Conference Papers (2023)
63. Watson, D., Chan, W., Brualla, R.M., Ho, J., Tagliasacchi, A., Norouzi, M.: Novel View Synthesis with Diffusion Models. In: International Conference on Learning Representations (ICLR) (2023), <https://openreview.net/forum?id=HtoA0oT30jC>
64. Wei, J., Wang, X., Schuurmans, D., Bosma, M., Xia, F., Chi, E., Le, Q.V., Zhou, D., et al.: Chain-of-thought prompting elicits reasoning in large language models. *Advances in neural information processing systems* **35**, 24824–24837 (2022)
65. Wen, B., Yang, W., Kautz, J., Birchfield, S.T.: FoundationPose: Unified 6D Pose Estimation and Tracking of Novel Objects. In: IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 17868–17879 (2023)
66. Wiedemer, T., Li, Y., Vicol, P., Gu, S.S., Matarese, N., Swersky, K., Kim, B., Jaini, P., Geirhos, R.: Video models are zero-shot learners and reasoners. arXiv preprint arXiv:2509.20328 (2025)
67. Wu, C., Li, J., Zhou, J., Lin, J., Gao, K., Yan, K., ming Yin, S., Bai, S., Xu, X., Chen, Y., Chen, Y., Tang, Z., Zhang, Z., Wang, Z., Yang, A., Yu, B., Cheng, C., Liu, D., Li, D., Zhang, H., Meng, H., Wei, H., Ni, J., Chen, K., Cao, K., Peng, L., Qu, L., Wu, M., Wang, P., Yu, S., Wen, T., Feng, W., Xu, X., Wang, Y., Zhang, Y., Zhu, Y., Wu, Y., Cai, Y., Liu, Z.: Qwen-image technical report (2025), <https://arxiv.org/abs/2508.02324>
68. Wu, J., Zhang, C., Xue, T., Freeman, W.T., Tenenbaum, J.B.: Learning a probabilistic latent space of object shapes via 3d generative-adversarial modeling. In: *Advances in Neural Information Processing Systems (NeurIPS)*. vol. 29, pp. 82–90 (2016)

69. Wu, R., Mildenhall, B., Henzler, P., Park, K., Gao, R., Watson, D., Srinivasan, P.P., Verbin, D., Barron, J.T., Poole, B., et al.: Reconfusion: 3d reconstruction with diffusion priors. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 21382–21392 (2024)
70. Xiang, J., Chen, X., Xu, S., Wang, R., Lv, Z., Deng, Y., Zhu, H., Dong, Y., Zhao, H., Yuan, N.J., Yang, J.: Native and compact structured latents for 3d generation. Tech report (2025)
71. Xiang, J., Lv, Z., Xu, S., Deng, Y., Wang, R., Zhang, B., Chen, D., Tong, X., Yang, J.: Structured 3D Latents for Scalable and Versatile 3D Generation. In: IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) (2025), <https://arxiv.org/abs/2412.01506>
72. Xu, Y., Shi, Z., Yifan, W., Chen, H., Yang, C., Peng, S., Shen, Y., Wetzstein, G.: GRM: Large Gaussian Reconstruction Model for Efficient 3D Reconstruction and Generation. In: European Conference on Computer Vision (ECCV) (2024)
73. Yu, A., Ye, V., Tancik, M., Kanazawa, A.: pixelnerf: Neural radiance fields from one or few images. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 4578–4587 (2021)
74. Yu, M., Hu, W., Xing, J., Shan, Y.: TrajectoryCrafter: Redirecting Camera Trajectory for Monocular Videos via Diffusion Models. In: International Conference on Computer Vision (ICCV) (2025)
75. Zhang, D.J., Paiss, R., Zada, S., Karnad, N., Jacobs, D.E., Pritch, Y., Mosseri, I., Shou, M.Z., Wadhwa, N., Ruiz, N.: Recapture: Generative video camera controls for user-provided videos using masked video fine-tuning. In: IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) (2025)
76. Zhang, R., Isola, P., Efros, A.A., Shechtman, E., Wang, O.: The unreasonable effectiveness of deep features as a perceptual metric. In: Proceedings of the IEEE conference on computer vision and pattern recognition. pp. 586–595 (2018)
77. Zhang, Y., Yuan, Y., Song, Y., Wang, H., Liu, J.: EasyControl: Adding Efficient and Flexible Control for Diffusion Transformer (2025), <https://arxiv.org/abs/2503.07027>
78. Zhou, J.J., Gao, H., Voleti, V., Vasishta, A., Yao, C.H., Boss, M., Torr, P., Rupprecht, C., Jampani, V.: Stable Virtual Camera: Generative View Synthesis with Diffusion Models. In: International Conference on Computer Vision (ICCV) (2025)