

OrbiSim: World Models as Differentiable Physics Engines for Embodied Intelligence

Jiajian Li* Jingyuan Huang* Junru Gong*
 Qi Wang Xiaokang Yang Yunbo Wang[†]

MoE Key Lab of Artificial Intelligence, AI Institute, School of Computer Science
 Shanghai Jiao Tong University

May 19, 2026

Abstract

We present OrbiSim, a novel robotic simulation paradigm that redefines world models as a fully differentiable physics engine for embodied intelligence. Unlike prior world models that focus on unconstrained imagination in latent or visual domains, OrbiSim establishes a unified, physically-grounded pathway that bridges structured scene assets, neural dynamics, and downstream reinforcement learning. By enabling end-to-end differentiability throughout the entire simulation loop—spanning from explicit state transitions to visual observation generation—OrbiSim supports tasks traditionally intractable for classical simulators, such as differentiable contact modeling, gradient-based policy optimization under sparse rewards, and intuitive physical inference. Empirical results demonstrate that OrbiSim significantly outperforms state-of-the-art world models in both predictive fidelity and control performance. Furthermore, its consistent responsiveness to asset configurations and physical parameters suggests its potential as a differentiable tool for enhancing robot simulation and policy training. To view detailed demonstrations, please visit [our website](#).

Table 1: **Comparison of OrbiSim with traditional physics engines and world models.**

Features	Classical Physics Engines (<i>e.g.</i> , Isaac Sim [27])	Generative World Models (<i>e.g.</i> , Motus [3])	OrbiSim (Ours)
Primary Focus	High-Fidelity Simulation	Video Synthesis	Full-Stack Diff. Sim
Core Mechanism	Numerical Solvers	Neural Networks	Decoupled Neural Networks
Asset Control	✓Full	✗	✓Basic Geometry + Physics
Differentiability	✗	✓Partial	✓End-to-End
Output Modality	State + Rendered Pixels	Rendered Pixels	State + Rendered Pixels

1 Introduction

The advancement of embodied intelligence is intrinsically tied to the fidelity of simulation software [38, 44, 27], with classical physics engines such as MuJoCo [38], PhysX [29], and Bullet [8] serving as the core execution backends of these platforms. Despite their success, these engines are fundamentally constrained by their non-differentiable or only partially differentiable execution pipelines, rigid modeling assumptions, and the difficulty of integrating perception, dynamics, and control within a unified optimization framework.

In parallel, recent years have seen rapid advances in world models that learn predictive representations of environment dynamics directly from data [4, 12, 32]. These models have demonstrated impressive capabilities in visual perception, video prediction, and counterfactual reasoning for planning and

Preprint. This work is part of the NeoWorld Project.

*Equal contribution.

[†]Corresponding author: Yunbo Wang <yunbow@sjtu.edu.cn>.

reinforcement learning (RL). Nevertheless, most existing world models act more as visual predictors than as functional simulators, lacking the structural integrity required to replace traditional engines in a closed-loop training pipeline. They typically do not expose explicit physical states, do not accept structured scene assets or physical parameters as inputs, and do not provide a fully differentiable execution pathway that connects environment configuration to downstream control objectives.

To bridge this gap, we introduce *OrbiSim*, a novel paradigm that redefines world models as a fully differentiable physics engine for robotics. Unlike prior models that operate on unconstrained latent spaces, OrbiSim treats scene geometry and physical parameters as first-class inputs within a unified, physically-grounded pathway. As illustrated in Figure 1, the core strengths of OrbiSim are threefold:

- *General-purpose world representation*: OrbiSim adopts an asset-conditioned representation interface that supports heterogeneous object types through appropriate state and geometry encodings, rather than being limited to a task-specific design.
- *Decoupled dynamics and vision*: By decoupling the neural architecture into interlinked dynamics and rendering modules, OrbiSim simultaneously predicts precise physical states and high-fidelity visual observations and enables seamless integration with existing simulation platforms.
- *End-to-end differentiability*: The differentiable pipeline facilitates Real-to-Sim system identification over scene parameters and gradient-based policy optimization for downstream control.

We evaluate OrbiSim across a diverse suite of physics-rich benchmarks. The results show that OrbiSim achieves competitive physical consistency compared to classical physics solvers while substantially outperforming state-of-the-art world models in video prediction fidelity and long-horizon rollout stability. Beyond standard benchmarking, we further show that OrbiSim supports real-to-simulation through visual inference of physical states and hidden properties, extends naturally to articulated and deformable objects under a unified asset-conditioned framework, and enables analytical policy gradients for sparse-reward robotic manipulation.

2 Related Work

2.1 Traditional Physics Engines for Robotics

Simulators such as MuJoCo [38], PyBullet [9], Gazebo [20], robosuite [44], and Isaac Gym/Lab [23, 26, 27] are the standard for robotic learning due to their high physical fidelity and scalability. However, their reliance on non-differentiable numerical solvers introduces discontinuities that prevent gradient propagation through environment transitions. Differentiable simulators, including Brax [10], Warp [25], DiffTaichi [18], and GradSim [28], seek to address this limitation by redesigning simulation pipelines for automatic differentiation. While these methods enable gradient-based optimization, they often trade off physical accuracy, relying on simplified or approximate dynamics. They struggle to faithfully capture complex contact interactions, limiting their applicability in high-fidelity robotic settings.

2.2 Generative World Models

Recent generative world models [1, 12, 19, 33] excel at synthesizing visually coherent futures but often lack explicit physical attributes (*e.g.*, mass, friction) necessary for precise control. While state-of-the-art embodied world models such as DreamDojo [11], Lingbot-World [36], MOTUS [3], and LAWM [37] leverage world modeling to improve VLA policies, they do not function as full physical simulators. Similarly, navigation-centric models [17, 30, 24, 42] focus on infinite scene generation rather than object-centric interactivity. Fundamentally, these paradigms entangle visual synthesis with temporal transitions, encoding physics implicitly within massive parameter spaces. This entanglement hampers reliable gradient propagation, making them ill-suited for the differentiable robot simulation.

2.3 Latent World Models

Latent world models have transformed model-based RL by enabling compact, controllable representations of environment dynamics [13]. Building on this foundation, methods such as PlaNet [15],

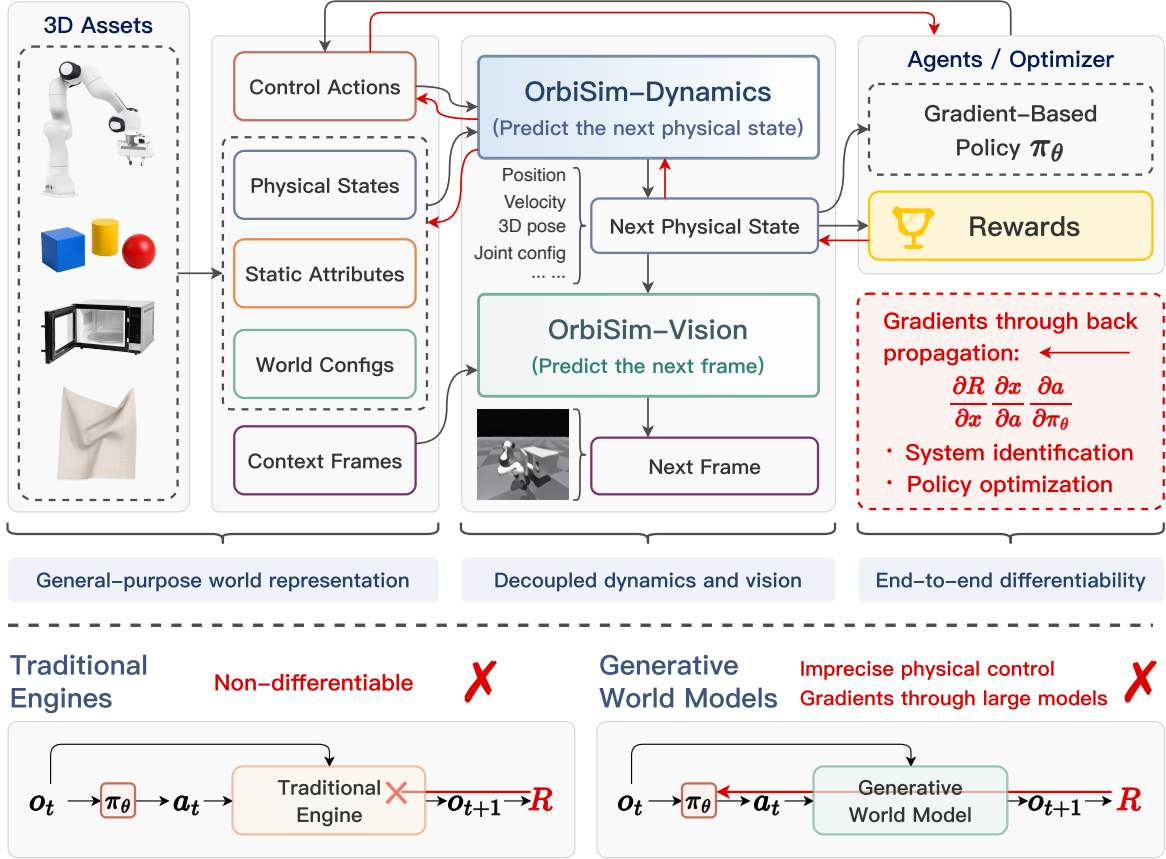


Figure 1: **Overview of OrbiSim.** Unlike prior generative world models, OrbiSim is designed as a *differentiable physics engine* that seamlessly integrates with modern robotic simulation software, enabling asset control and joint state-pixel generation. OrbiSim provides analytical gradients for direct policy optimization, which is a critical capability that distinguishes it from traditional simulators.

Dreamer [16], and subsequent extensions [31, 21, 35, 5, 7] achieve efficient policy optimization by performing imagination and rollout entirely in latent space. These advances have been further adopted in robotics for manipulation and offline evaluation (*e.g.*, IRASim [43], PIN-WM [22], and WorldGym [32]), often outperforming model-free baselines. Despite their empirical success, latent world models remain poorly aligned with simulation-oriented settings. They provide limited support for customizable assets, such as variable object geometries or compositional environments. Moreover, their reliance on abstract latent representations deprioritizes high-fidelity visual rendering, restricting their applicability in scenarios that require photorealistic simulation or rendering-based evaluation.

3 OrbiSim

We aim to learn a differentiable world model that functions as an execution engine for embodied manipulation tasks. Unlike monolithic world models that entangle physics and pixels, OrbiSim adopts a decoupled neural architecture comprising two specialized modules: *OrbiSim-Dynamics* and *OrbiSim-Vision*. As illustrated in Figure 1, this decomposition mirrors the separation of physics and rendering in classical simulation software, and it provides flexible, state-visual dual interfaces for downstream robot learning, supporting both MDP and POMDP formulations.

We first introduce a unified world representation for the robot and scene objects, which accommodates multiple object types and their corresponding geometry and physics descriptors (Sec. 3.1). At each step, OrbiSim-Dynamics autoregressively predicts physical states from historical trajectories, actions, and scene descriptors (Sec. 3.2). OrbiSim-Vision renders the next observation by grounding latent diffusion in the predicted states and scene context (Sec. 3.3). This modularity decouples physical evolution

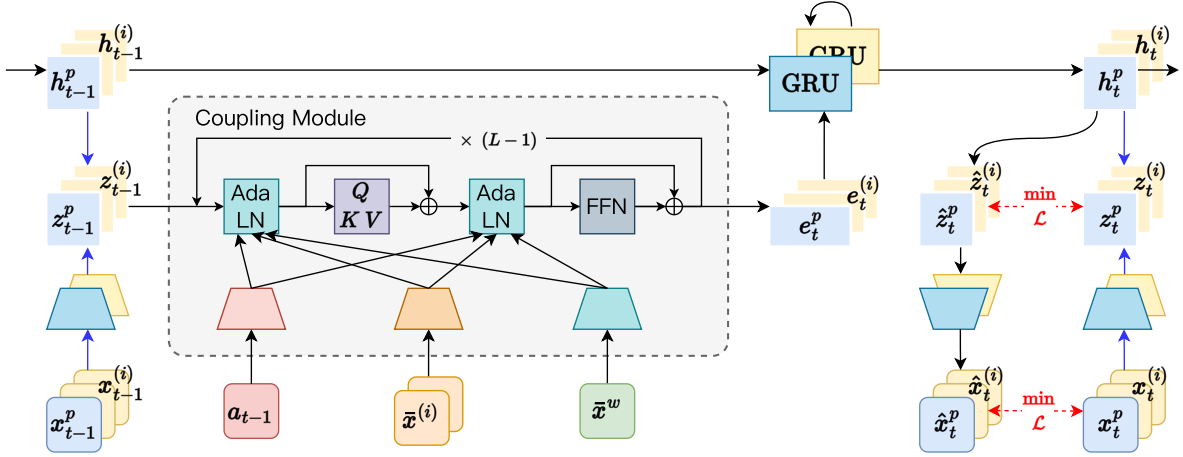


Figure 2: **OrbiSim-Dynamics**. An object-centric recurrent dynamics core that predicts next-step physical states. The model utilizes a Transformer-based coupling module to capture multi-entity interactions, with dynamics modulated by control actions and physical attributes via AdaLN.

from visual complexity, while maintaining a fully differentiable pipeline that enables real-to-sim system identification (Sec. 3.4) and downstream gradient-based policy learning (Sec. 3.5).

3.1 Unified Asset and World Representation

We model embodied manipulation as a physically grounded simulation process in which a robot interacts with N objects over discrete time steps. Each task instance is defined by a scene configuration specifying object geometry, physical properties, and global environment parameters. As a drop-in replacement for a classical physics engine, OrbiSim represents the system through three components:

- *Physical state* (x_t): The time-dependent physical state x_t is defined as the union of the robot’s physical state x_t^p and the states of N interacting objects $\{x_t^{(i)}\}_{i=0}^{N-1}$. Each state vector contains the quantities required for dynamics prediction, such as pose, velocities, joint configurations, or other time-varying state variables depending on the asset type.
- *Visual observation* (o_t): Given x_t , the world-model engine produces a high-dimensional visual observation in the form of an RGB image $o_t \in \mathbb{R}^{H \times W \times 3}$.
- *Static scene descriptors* ($\bar{x}$): We use $\bar{x}$ to represent time-invariant properties that characterize the neural simulator’s physical bounds. We categorize these into:
 - Asset attributes ($\bar{x}^{(i)}$), which are intrinsic properties of each object i , such as geometry, mass, inertia, surface friction, or other asset-specific descriptors.
 - World configurations ($\bar{x}^w$), which capture global environment parameters shared across all objects, such as gravity or contact properties.

A key advantage of this representation is that OrbiSim is designed as a general-purpose, asset-conditioned framework rather than a task-specific predictor. For simple objects, physical states and static descriptors can be represented directly using low-dimensional attributes. For objects with more complex geometry, such as *articulated* or *deformable* assets, their shape descriptions (*e.g.*, meshes or point clouds) can instead be encoded by the corresponding geometric encoder into fixed-shape object-level or part-level tokens, while the overall representation interface remains unchanged. Additional implementation details for complex objects are provided in Appendix B.4.

3.2 OrbiSim-Dynamics: Object-Centric Physics Core

OrbiSim-Dynamics models the evolution of physical states as a deterministic, history-aware recurrent process. As shown in Figure 2, given the past states $x_{0:t-1}$, the previous action a_{t-1} , and static scene

descriptors $\bar{x}$, the dynamics core predicts the subsequent state:

$$\hat{x}_t = f_\phi^{\text{dyn}}(x_{0:t-1}, a_{t-1}, \bar{x}). \quad (1)$$

To capture the complex multi-object interactions, OrbiSim-Dynamics employs an *object-centric paradigm*. Each entity (the robot and individual objects) is treated as a discrete token with its own latent and recurrent memory. Static object attributes, world configurations, and control actions are incorporated as conditioning signals via adaptive normalization, providing structured modulation of the dynamics. We employ a recurrent state space model structured as follows:

$$\begin{cases} z_t = f_\phi^{\text{enc}}(h_t, x_t), & h_t = f_\phi^{\text{rec}}(h_{t-1}, e_t), & \hat{z}_t = f_\phi^{\text{tra}}(h_t), & \hat{x}_t = f_\phi^{\text{dec}}(\hat{z}_t), \\ \text{Coupling module: } e_t = f_\phi^{\text{cp}}(z_{t-1}, a_{t-1}, \bar{x}). \end{cases} \quad (2)$$

The execution begins by mapping the physical state $x_{t-1} = \{x_{t-1}^p, x_{t-1}^{(0)}, \dots, x_{t-1}^{(N-1)}\}$ and corresponding recurrent hidden states $h_{t-1} = \{h_{t-1}^p, h_{t-1}^{(0)}, \dots, h_{t-1}^{(N-1)}\}$ into a sequence of latent tokens $z_{t-1} = \{z_{t-1}^p, z_{t-1}^{(0)}, \dots, z_{t-1}^{(N-1)}\}$, where the first token is for the robot. These tokens are processed by a *Transformer-based Coupling Module* utilizing global self-attention to model inter-entity dependencies, such as collisions and articulated constraints. We inject physical constants and control signals through Adaptive Layer Normalization (AdaLN) to modulate the dynamics flow, where the token feature u represents physical states and the condition c encodes actions and static scene descriptors:

$$\text{AdaLN}(u, c) = \gamma(c) \odot \text{LN}(u) + \beta(c), \quad (3)$$

where $\gamma(\cdot)$ and $\beta(\cdot)$ are learnable functions of the condition. Specifically, the robot token is modulated by the action a_{t-1} and world configurations $\bar{x}^w$, while object tokens are modulated by their respective attributes $\bar{x}^{(i)}$ (e.g., mass, friction) and $\bar{x}^w$. This structured modulation provides the model with the necessary inductive bias to generalize across varying physical environments.

The coupling module outputs physics-aware embeddings $e_t = \{e_t^p, e_t^{(0)}, \dots, e_t^{(N-1)}\}$, which are then combined with h_{t-1} to update h_t . From the updated recurrent states, a transition module produces the predicted latent states: $\hat{z}_t = \{\hat{z}_t^p, \hat{z}_t^{(0)}, \dots, \hat{z}_t^{(N-1)}\}$, which are finally decoded into system state predictions: $\hat{x}_t = \{\hat{x}_t^p, \hat{x}_t^{(0)}, \dots, \hat{x}_t^{(N-1)}\}$.

OrbiSim-Dynamics is trained with three loss terms, including the latent alignment loss ($\mathcal{L}_{\text{tra}}$), the encoding consistency loss ($\mathcal{L}_{\text{enc}}$), and the state reconstruction loss ($\mathcal{L}_{\text{dec}}$):

$$\mathcal{L}_{\text{tra}} = \|\hat{z}_t - \text{sg}(z_t)\|^2, \quad \mathcal{L}_{\text{enc}} = \|z_t - \text{sg}(\hat{z}_t)\|^2, \quad \mathcal{L}_{\text{dec}} = \mathcal{L}_{\text{state}}(x_t, \hat{x}_t), \quad (4)$$

where $\text{sg}(\cdot)$ denotes gradient stopping. $\mathcal{L}_{\text{state}}$ may employ different loss forms for different physical state components (see Appendix B.1). For long-horizon stability, the total objective $\mathcal{L}_{\text{dyn}}$ is optimized over T -step autoregressive rollouts over an episode $\{\mathcal{D}_t = (x_t, z_t, \hat{x}_t, \hat{z}_t)\}_{t=1}^T$.

3.3 OrbiSim-Vision: State-Guided Diffusion

OrbiSim-Vision transforms predicted physical states into high-fidelity visual observations $\hat{o}_t$ through a conditional image generation module. As shown in Figure 8, We formulate this as a *Latent Diffusion Process*, which is grounded in the scene context $\bar{x}$ and the physical state $\hat{x}_t$ produced by OrbiSim-Dynamics:

$$\hat{o}_t \sim p_\phi^{\text{vis}}(o_t \mid o_{t-K:t-1}, \hat{x}_t, \bar{x}), \quad (5)$$

where the denoising process is carried out in the latent space and explicitly conditioned on the predicted physical state. We represent the visual latent at time t and diffusion step (s) as $y_{t,(s)}$, spanning from clean ($s = 0$) to noisiest ($s = S$)

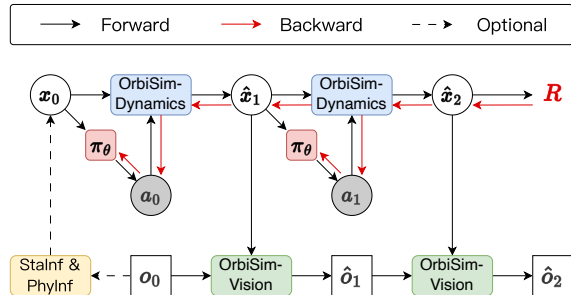


Figure 3: **Gradient pathways in OrbiSim.** In state-based tasks, gradients propagate exclusively through OrbiSim-Dynamics.

states. The preceding K observations are VAE-encoded into clean latents $y_{t-K:t-1,(0)}$, which serve as the context for the next-frame generation. We assemble a sequence $y_{t-K:t}$ by appending an initialization latent to this context, which is set to $y_{t,(0)}$ during training and a zero placeholder during inference.

Once the forward diffusion process reaches the noisy state $y_{t-K:t,(S)}$, the target frame is generated through a denoising process conditioned on the clean context latents. First, we achieve *visual anchoring* by concatenating a perturbed version of the most recent latent $y_{t-1,(0)}$ channel-wise to the denoising sequence $y_{t-K:t,(S)}$. Second, *physical grounding* is incorporated by encoding the predicted physical state $\hat{x}_t$ and static scene descriptors $\bar{x}$ into two complementary representations:

- *Spatial condition maps*: These provide spatially aligned geometric cues, such as heatmaps or depth-like priors, which are concatenated directly to the U-Net feature maps.
- *Object-level tokens*: These facilitate cross-attention modulation within the U-Net, enabling the model to attend to specific object states during denoising.

The U-Net iteratively refines the sequence for S steps, and the final clean latent $\hat{y}_{t,(0)}$ is decoded into pixel space. After the denoising process, we take the final clean latent at the last time step as the predicted latent representation $\hat{y}_{t,(0)}$. Finally, a VAE decoder maps $\hat{y}_{t,(0)}$ back to pixels to produce $\hat{o}_t$, completing a fully differentiable rendering process conditioned on physical states.

We train the conditional denoising network D_ϕ as follows, which enforces physical consistency by grounding the denoising process in the predicted states:

$$\mathcal{L}_{\text{vis}} = \mathbb{E}_{\sigma, \epsilon} \left[w(\sigma) \left\| \left(D_\phi(y_{t-K:t,(0)} + \sigma \epsilon; c_t) \right)_t - y_{t,(0)} \right\|_2^2 \right], \quad (6)$$

where c_t denotes the combined visual and physical conditioning, and $w(\sigma)$ is a noise-dependent weighting function following standard latent diffusion training. During training, we sample a noise scale σ and Gaussian noise $\epsilon \sim \mathcal{N}(0, I)$, and apply noise to the latent sequence.

3.4 Real-to-Sim System Identification

The integration of OrbiSim-Dynamics and OrbiSim-Vision provides a fully differentiable execution engine, enabling gradient-based system identification (SysID) across the entire simulation trace. In scenarios where physical states are measurable, scene descriptors $\bar{x}$ can be optimized by fixing the pretrained dynamics core and backpropagating state-prediction errors. However, direct SysID is ill-posed in visual environments where only pixel-level observations o_t are available. Because o_t is simultaneously conditioned on the unknown physical state x_t and descriptors $\bar{x}$, visual appearances are entangled with underlying poses and physical properties. This makes the joint posterior distribution of these variables intractable, making direct causal inference from pixels to physics through pure gradient descent prone to local minima or physical inconsistency.

To establish an efficient Real-to-Sim pipeline, we introduce two specialized components within OrbiSim: (i) A *state inference module (StaInf)* that maps image latents to explicit physical states and visually grounded attributes; and (ii) a *physics inference module (PhyInf)* that infers unobservable properties (*e.g.*, mass/friction) from short-horizon physical state trajectories. By decoupling the inference of x_t and $\bar{x}$, these modules provide a structured interface for mapping raw visual data into the precise configurations required by OrbiSim. Further details are provided in Appendix B.3.

3.5 Policy Optimization via Analytical Gradients

As a differentiable engine, OrbiSim enables *analytical policy gradients*. A control policy π_θ interacts with OrbiSim to produce an execution trace, and as illustrated in Figure 3, for a terminal objective $\mathcal{J}(\theta) = R(x_T)$, gradients propagate directly from the reward through the differentiable dynamics:

$$\nabla_\theta \mathcal{J}(\theta) = \frac{\partial R}{\partial x_T} \sum_{t=0}^{T-1} \left(\prod_{k=t+1}^{T-1} \frac{\partial x_{k+1}}{\partial x_k} \right) \frac{\partial x_{t+1}}{\partial a_t} \frac{\partial \pi_\theta(x_t)}{\partial \theta}, \quad (7)$$

where each Jacobian $\partial x_{k+1}/\partial x_k$ flows through *OrbiSim-Dynamics* rather than the visual renderer, avoiding the gradient degradation and training instability often induced by monolithic world models.

OrbiSim supports two policy learning regimes. When the initial physical state is available, policy optimization proceeds directly by differentiating through the dynamics rollout. When real-world data does not provide ground-truth states, the **StaInf** and **PhyInf** modules (Sec. 3.4) first infer the initial state and scene descriptors from raw observations, after which the same optimization is applied without modification. This unified formulation enables OrbiSim to serve as a differentiable policy optimizer for both simulation and real-world data.

4 Experiments

We evaluate OrbiSim as both a generative world model and a differentiable execution engine, focusing on (1) generative fidelity and physical consistency under varying configurations, and (2) the benefits of differentiable gradient pathways for downstream RL.

4.1 Experimental Configurations and Baselines

We use four physics-rich environments built upon classical simulation platforms as the learning targets of OrbiSim: • **Robosuite Push**: A UR5e manipulator must push a primary cube to collide with and relocate a target cube. Cube sizes, colors, densities, frictions, and initial positions are randomized across the dataset, requiring physics-aware, object-centric precision [44]. • **Isaac Lab Stack**: A Franka arm sequentially stacks three cubes, requiring long-horizon (200+ steps) stability and precise multi-object interaction [27]. • **AdaManip Articulated**: We use articulated-object manipulation tasks from the Isaac Gym-based AdaManip environment [39, 26], where a robot interacts with joint-constrained objects with diverse shapes and mechanisms. • **Physion Drape**: We further evaluate deformable-object dynamics using the Drape scenario from Physion [2], where a cloth falls onto rigid objects with varying geometries. More dataset details are provided in Appendix A.

We compare OrbiSim against the state-of-the-art generative world models, Vid2World [19] and AdaWorld [12]. For differentiable robot learning, we benchmark the gradient-based policies learned within OrbiSim against Soft Actor-Critic (SAC) [14], Proximal Policy Optimization (PPO) [34], PPO with Random Network Distillation (RND) [6], and Behavior Cloning (BC) learned within the original robosuite environment. In addition, DreamerV3 is also included as a model-based benchmark.

4.2 Evaluating Generative and Physical Fidelity

4.2.1 Performance on Benchmark Manipulation Tasks

As shown in Table 2, OrbiSim (Final) consistently achieves state-of-the-art performance across all metrics and horizons. We compare against **AdaWorld** [12] and **Vid2World** [19]; while these baselines produce visually plausible frames, they struggle with compounding physical drift. In contrast, OrbiSim maintains superior temporal coherence (FVD) and lower trajectory error, demonstrating a more robust alignment between physical dynamics and visual synthesis. Qualitative visualizations in Figure 4 and Figure 5 further suggest that OrbiSim remains robust across varying physics coefficients and complex stacking dynamics over a 225-step long-term simulation horizon.

Table 2: **Video-level world modeling performance on the robosuite Push task.** We report PSNR and LPIPS at different rollout horizons (10 / 100 steps), together with the overall FVD score. *TrajErr* measures the discrepancy between inferred physical states from generated videos and the corresponding true trajectories. All modes perform autoregressive rollouts from shared initial states.

Method	PSNR10 $\uparrow$	PSNR100 $\uparrow$	LPIPS10 $\downarrow$	LPIPS100 $\downarrow$	FVD $\downarrow$	TrajErr $\downarrow$
Vid2World	22.2014	17.8856	0.1312	0.2551	1750.1	0.6754
AdaWorld	26.6647	12.8346	0.1183	0.3482	1305.8	1.8597
Orbisim w/o Decoupling	27.9346	<u>19.9510</u>	0.1188	0.1799	689.1	0.8134
Orbisim w/o Random Sampling	26.6890	19.1119	0.1076	0.1669	<u>531.2</u>	0.5742
Orbisim w/o Object-Centric	25.9373	19.7581	0.1123	<u>0.1463</u>	524.5	<u>0.4687</u>
Orbisim (Final)	<u>26.7105</u>	19.9819	<u>0.1078</u>	0.1428	533.9	0.4468



(a) High friction



(b) Low friction

Figure 4: **Autoregressive simulation under varying physical parameters.** All methods are evaluated with the same initial observations and action sequences. OrbiSim better simulates system dynamics under diverse object-table friction settings.

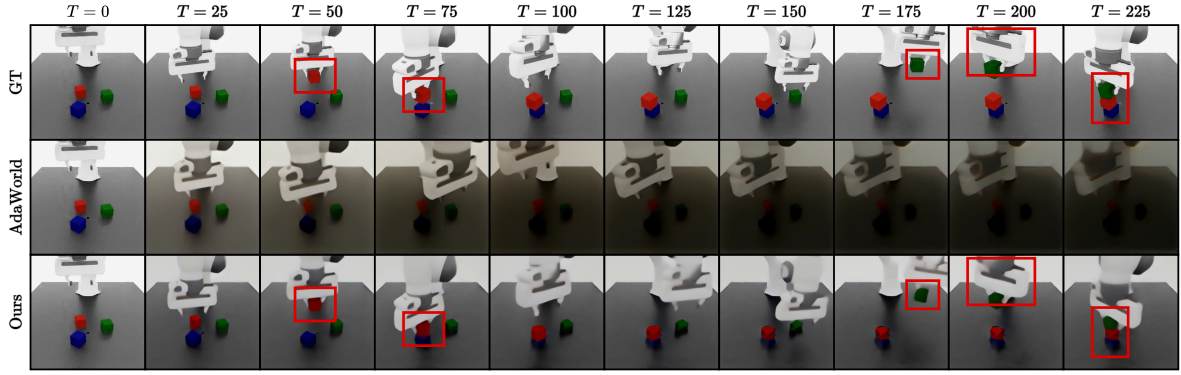


Figure 5: **Autoregressive simulation on the Isaac Lab Stack task over a 225-step horizon.** OrbiSim accurately captures intricate manipulator rotations and complex multi-object collisions, maintaining high physical fidelity and stability throughout the sequence.

4.2.2 Ablation Studies

We assess the contribution of core design choices through three variants: • *OrbiSim w/o Decoupling* removes the separation between dynamics and vision, training a single monolithic model to jointly predict states and observations. • *OrbiSim w/o Random Sampling* disables random sparsification of visual context during training, using only contiguous context frames. • *OrbiSim w/o Object-Centric* replaces object-centric dynamics modeling with a global state representation that concatenates the robot state, all object states, object attributes, and world configurations at each time step.

As shown in Table 2, removing the decoupled dynamics–vision architecture improves single-frame reconstruction (higher PSNR) but severely degrades temporal consistency and trajectory accuracy, leading to pronounced long-horizon distortions, confirming the necessity of separating dynamics modeling from visual rendering. Disabling random visual context sampling slightly degrades most metrics and increases trajectory error, indicating that context sparsification improves robustness to autoregressive noise. Replacing object-centric modeling with a global state representation yields consistent but moderate performance drops, with visualizations further revealing impaired modeling of object interactions, highlighting the importance of object-level structure for multi-body dynamics.

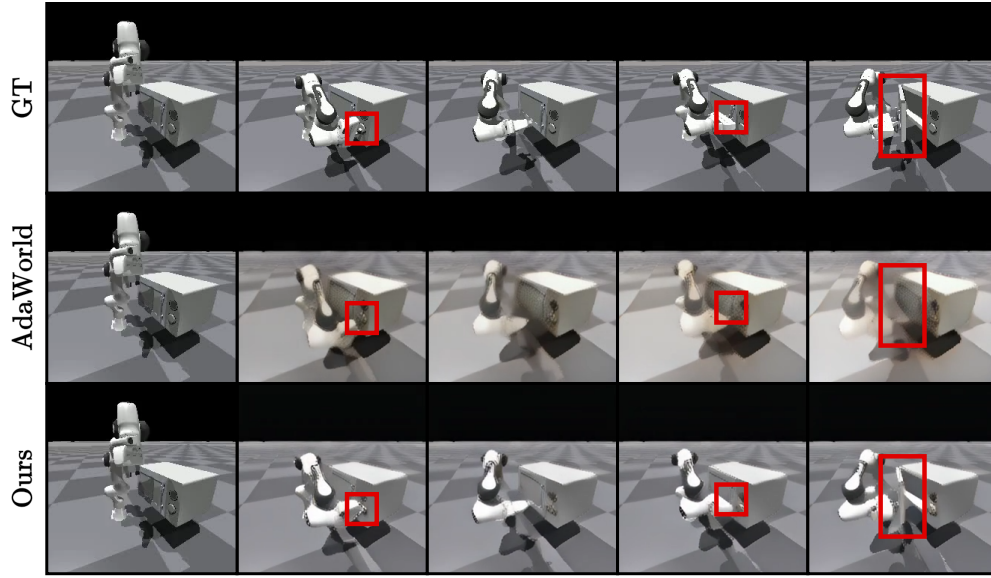
4.2.3 Generalization to Heterogeneous Assets

As discussed in Sec. 3.1, OrbiSim is designed as a general-purpose, asset-conditioned framework that can accommodate heterogeneous object types through appropriate state and asset encodings. We therefore evaluate this extensibility on two additional complex-object settings: articulated-object manipulation from AdaManip [39] and deformable cloth dynamics from the Physion Drapes scenario [2]. For articulated objects, object assets are represented by structured 3D asset files, such as URDF and mesh files, together with joint states. For deformable objects, cloth geometry is represented using point-based states. We encode these geometry-conditioned descriptors with a geometric encoder, such as Point Transformer V3 [41], and feed the resulting asset tokens into the same object-centric dynamics and state-guided vision modules. More implementation details are provided in Appendix B.4.

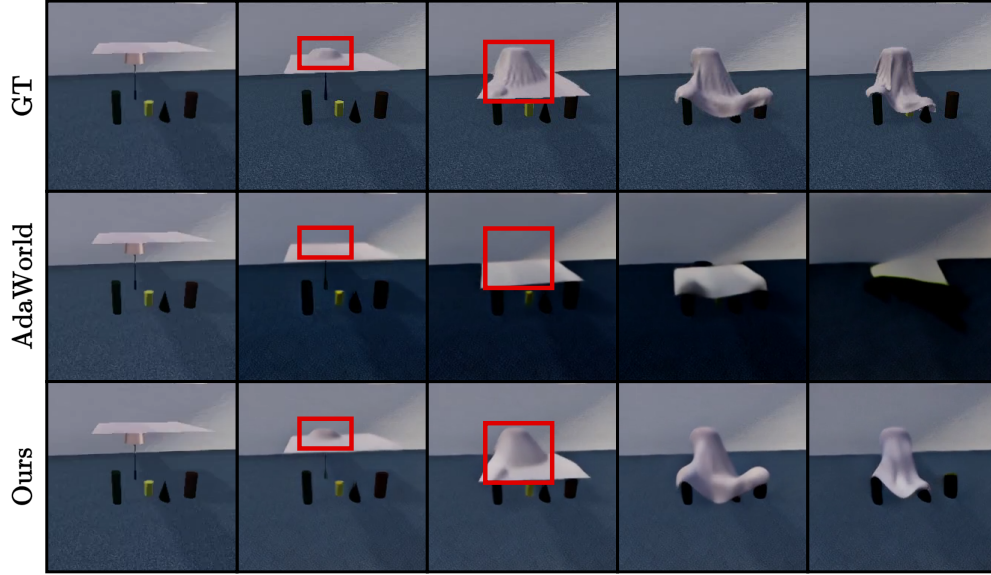
Figure 6 shows qualitative autoregressive rollouts on both settings. OrbiSim accurately follows joint-constrained part motion in articulated-object manipulation and captures the overall deformation pattern of cloth interacting with rigid objects. These results suggest that the asset-conditioned design enables OrbiSim to accommodate substantially richer object geometries and physical properties without redesigning the simulator architecture.

4.2.4 Robustness to Out-of-Distribution Scenarios

We further evaluate the generalization ability of OrbiSim under out-of-distribution (OOD) configurations. We construct three types of OOD scenarios: • *Object-count OOD*, where the model is trained on two-object scenes and evaluated on unseen three-object scenes; • *Shape OOD*, where point-cloud-based object geometries in the test split do not appear during training; and • *Physical-parameter OOD*,



(a) AdaManip Articulated



(b) Physion Drape

Figure 6: **Autoregressive simulation on complex objects.** (a) Articulated-object manipulation from AdaManip. (b) Deformable cloth dynamics from the Physion Drape scenario. OrbiSim captures both joint-constrained articulated motion and geometry-conditioned cloth deformation under the same asset-conditioned simulation framework.

Table 3: **Out-of-distribution world modeling performance.** We report PSNR and LPIPS at different rollout horizons (10 / 100 steps), together with the overall FVD score. OOD settings include unseen object counts, unseen shape distributions, and extreme physical parameters.

Method	PSNR10 $\uparrow$	PSNR100 $\uparrow$	LPIPS10 $\downarrow$	LPIPS100 $\downarrow$	FVD $\downarrow$
AdaWorld (ID)	26.6647	12.8346	<u>0.1183</u>	0.3482	1305.8
OrbiSim (ID)	<u>26.7105</u>	<u>19.9819</u>	0.1078	0.1428	533.9
AdaWorld (OOD)	26.3468	13.2215	0.1228	0.3836	1288.5
OrbiSim (OOD)	27.1867	20.1109	0.1188	<u>0.1847</u>	<u>597.3</u>

where we evaluate the model under extreme friction settings that are significantly outside the training distribution in the push-hit task.

Table 3 summarizes the video-level performance under in-distribution and OOD settings. Even under unseen object counts, shapes, and physical conditions, OrbiSim maintains strong generation quality and substantially outperforms AdaWorld, especially over long horizons. These results indicate that the asset-conditioned and object-centric design of OrbiSim provides robust generalization beyond the training distribution, rather than merely memorizing seen scene configurations.

4.3 Differentiability Evaluation on Policy Optimization

We evaluate OrbiSim as a functional execution engine for downstream RL tasks. • In the robosuite **Push** task, the robot must push one cube to hit another into a target region without either object falling off the table. A rollout is considered successful if the second cube moves more than 5 cm. We define three normalized rewards based on distance reduction to encourage the robot to approach and push the cubes (see Appendix C.8). To preserve sparse supervision, these rewards are provided only at the end of each episode during training. During evaluation, they are used to measure task completion performance. • In the Isaac Lab **Stack** task, the agent must sequentially stack three cubes, with binary success determined by whether the final stacked configuration is achieved.

While the sparse episodic reward design makes credit assignment particularly challenging, OrbiSim differs from traditional black-box simulators by exposing analytical gradient pathways that propagate task-specific reward signals directly to the policy parameters. As shown in Figure 7, OrbiSim achieves superior performance and convergence speed compared to model-free and imitation baselines.

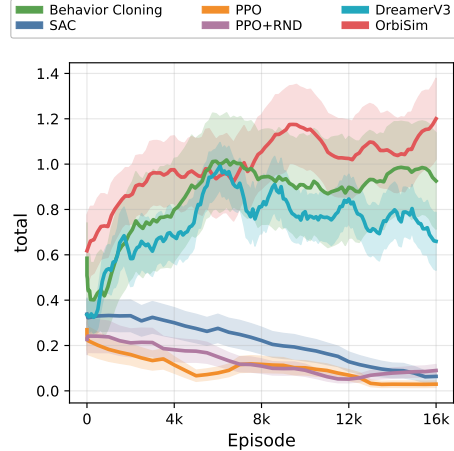


Figure 7: **Reward curve of robosuite Push task RL.** We compared OrbiSim with model-free baselines (SAC, PPO, PPO+RND), Behavior Cloning as well as model-based baseline (DreamerV3).

- *vs. Model-free RL:* We adopt SAC, PPO, and PPO with RND as model-free RL baselines. As shown in the results, model-free methods struggle under extremely sparse rewards, often failing to obtain consistent learning signals and collapsing to near-zero performance. In contrast, OrbiSim provides informative optimization signals through its differentiable dynamics, enabling the policy to reason about the physical consequences of its actions even under terminal-only supervision.
- *vs. Behavior cloning:* While BC learns the initial pushing behavior from expert demonstrations, it struggles to generalize to OOD states throughout long-horizon interactions. In contrast, OrbiSim enables the agent to refine its policy through rollouts in a physical world model.
- *vs. Model-based RL:* Model-based RL methods such as DreamerV3 predict future dynamics and optimize policies in latent space. However, DreamerV3 lacks explicit physical modeling, which limits its robustness under varying physical settings.

Detailed training reward curves and the success rate table are listed in Appendix C.8.

4.4 Additional Evaluation

In Appendix C, we provide comprehensive supplementary analyses across three primary axes:

- *Component analyses:* We investigate the isolated predictive performance of OrbiSim-Dynamics (C.1), assess various inference workflows (C.2), and validate the Real-to-Sim performance (C.3).
- *Scalability/robustness:* We evaluate OrbiSim’s scalability to multi-object/multi-task settings (C.4) and its robustness under noisy training data that simulates real-world acquisition pipelines (C.5).

- *Extended policy studies:* We offer further insights into policy optimization, including imitation learning baselines (C.6), comparisons against specialized differentiable simulators such as NVIDIA Newton (C.7), expanded RL experiments (C.8), and computational efficiency and resource usage analyses (C.9).

5 Conclusions and Limitations

In this paper, we proposed OrbiSim, a novel world model paradigm that rethinks simulation from a learning-centric perspective by treating the world model itself as a fully differentiable physics engine. Through extensive experiments on robotic simulation and control tasks, we demonstrated that OrbiSim outperforms existing world models in future prediction accuracy and control performance.

Despite these encouraging results, OrbiSim still has important limitations. Although it shows promising robustness within related simulator families and already supports heterogeneous assets, its generalization to broader robot morphologies, scene sources, and real-world domains remains underexplored, and it is still far from a complete, general-purpose industrial simulator. Nonetheless, its fully differentiable formulation and scalability in data-rich settings make OrbiSim a promising step toward practical deployment.

References

- [1] Niket Agarwal, Arslan Ali, Maciej Bala, Yogesh Balaji, Erik Barker, Tiffany Cai, Prithvijit Chattopadhyay, Yongxin Chen, Yin Cui, Yifan Ding, et al. Cosmos world foundation model platform for physical ai. *arXiv preprint arXiv:2501.03575*, 2025.
- [2] Daniel M. Bear, Elias Wang, Damian Mrowca, Felix J. Binder, H Tung, R. T. Pramod, Cameron Holdaway, Sirui Tao, Kevin A. Smith, Li Fei-Fei, Nancy G. Kanwisher, Joshua B. Tenenbaum, Daniel Yamins, and Judith E. Fan. Physion: Evaluating physical prediction from vision in humans and machines. *ArXiv*, abs/2106.08261, 2021.
- [3] Hongzhe Bi, Hengkai Tan, Shenghao Xie, Zeyuan Wang, Shuhe Huang, Haitian Liu, Ruowen Zhao, Yao Feng, Chendong Xiang, Yinze Rong, et al. Motus: A unified latent action world model. *arXiv preprint arXiv:2512.13030*, 2025.
- [4] Jake Bruce, Michael D Dennis, Ashley Edwards, Jack Parker-Holder, Yuge Shi, Edward Hughes, Matthew Lai, Aditi Mavalankar, Richie Steigerwald, Chris Apps, et al. Genie: Generative interactive environments. In *ICML*, 2024.
- [5] Maxime Burchi and Radu Timofte. Learning transformer-based world models with contrastive predictive coding. *arXiv preprint arXiv:2503.04416*, 2025.
- [6] Yuri Burda, Harrison Edwards, Amos Storkey, and Oleg Klimov. Exploration by random network distillation. *arXiv preprint arXiv:1810.12894*, 2018.
- [7] Lior Cohen, Kaixin Wang, Bingyi Kang, and Shie Mannor. Improving token-based world models with parallel observation prediction. *arXiv preprint arXiv:2402.05643*, 2024.
- [8] Erwin Coumans. Bullet physics simulation. In *SIGGRAPH*, 2015.
- [9] Erwin Coumans and Yunfei Bai. Pybullet, a python module for physics simulation for games, robotics and machine learning. <http://pybullet.org>, 2016–2021.
- [10] C Daniel Freeman, Erik Frey, Anton Raichuk, Sertan Girgin, Igor Mordatch, and Olivier Bachem. Brax—a differentiable physics engine for large scale rigid body simulation. *arXiv preprint arXiv:2106.13281*, 2021.
- [11] Shenyuan Gao, William Liang, Kaiyuan Zheng, Ayaan Malik, Seonghyeon Ye, Sihyun Yu, Wei-Cheng Tseng, Yuzhu Dong, Kaichun Mo, Chen-Hsuan Lin, et al. Dreamdojo: A generalist robot world model from large-scale human videos. *arXiv preprint arXiv:2602.06949*, 2026.

- [12] Shenyuan Gao, Siyuan Zhou, Yilun Du, Jun Zhang, and Chuang Gan. Adaworld: Learning adaptable world models with latent actions. In *ICML*, 2025.
- [13] David Ha and Jürgen Schmidhuber. World models. In *NeurIPS*, 2018.
- [14] Tuomas Haarnoja, Aurick Zhou, Pieter Abbeel, and Sergey Levine. Soft actor-critic: Off-policy maximum entropy deep reinforcement learning with a stochastic actor. In *ICML*, 2018.
- [15] Danijar Hafner, Timothy Lillicrap, Ian Fischer, Ruben Villegas, David Ha, Honglak Lee, and James Davidson. Learning latent dynamics for planning from pixels. In *ICML*, pages 2555–2565, 2019.
- [16] Danijar Hafner, Jurgis Pasukonis, Jimmy Ba, and Timothy Lillicrap. Mastering diverse control tasks through world models. *Nature*, pages 1–7, 2025.
- [17] Xianglong He, Chunli Peng, Zexiang Liu, Boyang Wang, Yifan Zhang, Qi Cui, Fei Kang, Biao Jiang, Mengyin An, Yangyang Ren, et al. Matrix-game 2.0: An open-source, real-time, and streaming interactive world model. *arXiv preprint arXiv:2508.13009*, 2025.
- [18] Yuanming Hu, Luke Anderson, Tzu-Mao Li, Qi Sun, Nathan Carr, Jonathan Ragan-Kelley, and Frédo Durand. DiffTaichi: Differentiable programming for physical simulation. *arXiv preprint arXiv:1910.00935*, 2019.
- [19] Siqiao Huang, Jialong Wu, Qixing Zhou, Shangchen Miao, and Mingsheng Long. Vid2world: Crafting video diffusion models to interactive world models. *arXiv preprint arXiv: 2505.14357*, 2025.
- [20] N. Koenig and A. Howard. Design and use paradigms for gazebo, an open-source multi-robot simulator. In *IROS*, 2004.
- [21] Jiajian Li, Qi Wang, Yunbo Wang, Xin Jin, Yang Li, Wenjun Zeng, and Xiaokang Yang. Open-world reinforcement learning over long short-term imagination. In *ICLR*, 2025.
- [22] Wenxuan Li, Hang Zhao, Zhiyuan Yu, Yu Du, Qin Zou, Ruizhen Hu, and Kai Xu. Pin-wm: Learning physics-informed world models for non-prehensile manipulation. *arXiv preprint arXiv:2504.16693*, 2025.
- [23] Jacky Liang, Viktor Makoviychuk, Ankur Handa, Nuttapong Chentanez, Miles Macklin, and Dieter Fox. Gpu-accelerated robotic simulation for distributed reinforcement learning. *arXiv preprint arXiv:1810.05762*, 2018.
- [24] Yifan Lu, Xuanchi Ren, Jiawei Yang, Tianchang Shen, Zhangjie Wu, Jun Gao, Yue Wang, Siheng Chen, Mike Chen, Sanja Fidler, et al. Infinitcube: Unbounded and controllable dynamic 3d driving scene generation with world-guided video models. *arXiv preprint arXiv:2412.03934*, 2024.
- [25] Miles Macklin. Warp: A high-performance python framework for gpu simulation and graphics. <https://github.com/nvidia/warp>, March 2022. NVIDIA GPU Technology Conference (GTC).
- [26] Viktor Makoviychuk, Lukasz Wawrzyniak, Yunrong Guo, Michelle Lu, Kier Storey, Miles Macklin, David Hoeller, Nikita Rudin, Arthur Allshire, Ankur Handa, and Gavriel State. Isaac gym: High performance gpu-based physics simulation for robot learning. *arXiv preprint arXiv:2108.10470*, 2021.
- [27] Mayank Mittal, Pascal Roth, James Tigue, Antoine Richard, Octi Zhang, Peter Du, Antonio Serrano-Muñoz, Xinjie Yao, René Zurbrugg, Nikita Rudin, et al. Isaac lab: A gpu-accelerated simulation framework for multi-modal robot learning. *arXiv preprint arXiv:2511.04831*, 2025.
- [28] J Krishna Murthy, Miles Macklin, Florian Golemo, Vikram Voleti, Linda Petrini, Martin Weiss, Breandan Considine, Jérôme Parent-Lévesque, Kevin Xie, Kenny Erleben, et al. gradsim: Differentiable simulation for system identification and visuomotor control. In *ICLR*, 2020.
- [29] NVIDIA Corporation. Physx physics sdk, 2020.

- [30] Başak Melis Öcal, Maxim Tatarchenko, Sezer Karaoğlu, and Theo Gevers. Sceneteller: Language-to-3d scene generation. In *ECCV*, pages 362–378, 2024.
- [31] Minting Pan, Xiangming Zhu, Yitao Zheng, Yunbo Wang, and Xiaokang Yang. Model-based reinforcement learning with isolated imaginations. *TPAMI*, 46(5):2788–2803, 2024.
- [32] Julian Quevedo, Ansh Kumar Sharma, Yixiang Sun, Varad Suryavanshi, Percy Liang, and Sherry Yang. WorldGym: World Model as An Environment for Policy Evaluation. *arXiv preprint arXiv:2506.00613*, 2025.
- [33] Marc Rigter, Tarun Gupta, Agrin Hilmkil, and Chao Ma. Avid: Adapting video diffusion models to world models. *arXiv preprint arXiv:2410.12822*, 2024.
- [34] John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- [35] Ruixiang Sun, Hongyu Zang, Xin Li, and Riashat Islam. Learning latent dynamic robust representations for world models. *arXiv preprint arXiv:2405.06263*, 2024.
- [36] Robbyant Team, Zelin Gao, Qiuyu Wang, Yanhong Zeng, Jiapeng Zhu, Ka Leong Cheng, Yixuan Li, Hanlin Wang, Yinghao Xu, Shuailei Ma, et al. Advancing open-source world models. *arXiv preprint arXiv:2601.20540*, 2026.
- [37] Bahey Tharwat, Yara Nasser, Ali Abouzeid, and Ian Reid. Latent action pretraining through world modeling. *arXiv preprint arXiv:2509.18428*, 2025.
- [38] Emanuel Todorov, Tom Erez, and Yuval Tassa. Mujoco: A physics engine for model-based control. In *IROS*, pages 5026–5033, 2012.
- [39] Yuanfei Wang, Xiaojie Zhang, Ruihai Wu, Yu Li, Yan Shen, Mingdong Wu, Zhaofeng He, Yizhou Wang, and Hao Dong. Adamanip: Adaptive articulated object manipulation environments and policy learning. *ArXiv*, abs/2502.11124, 2025.
- [40] Xiaoyang Wu, Daniel DeTone, Duncan P. Frost, Tianwei Shen, Christopher Xie, Nan Yang, Jakob Julian Engel, Richard A. Newcombe, Hengshuang Zhao, and Julian Straub. Sonata: Self-supervised learning of reliable point representations. *2025 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 22193–22204, 2025.
- [41] Xiaoyang Wu, Li Jiang, Peng-Shuai Wang, Zhijian Liu, Xihui Liu, Yu Qiao, Wanli Ouyang, Tong He, and Hengshuang Zhao. Point transformer v3: Simpler, faster, stronger. *2024 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 4840–4851, 2023.
- [42] Hong-Xing Yu, Haoyi Duan, Charles Herrmann, William T. Freeman, and Jiajun Wu. Wonderworld: Interactive 3d scene generation from a single image. In *CVPR*, 2025.
- [43] Fangqi Zhu, Hongtao Wu, Song Guo, Yuxiao Liu, Chiam Cheang, and Tao Kong. Irasim: A fine-grained world model for robot manipulation. In *ICCV*, pages 9834–9844, 2025.
- [44] Yuke Zhu, Josiah Wong, Ajay Mandlekar, and Roberto Martín-Martín. robosuite: A modular simulation framework and benchmark for robot learning. *arXiv preprint arXiv:2009.12293*, 2020.

Appendix

A Data Collection and Generation

In this appendix, we provide detailed descriptions of the data collection and generation procedures used in our experiments. These details complement the high-level experimental setups presented in the main paper and are included here for completeness and reproducibility.

Data Representation. Across all environments, datasets are organized as episodic trajectories consisting of structured state observations, static environment configurations, and, when applicable, control actions. For robot manipulation tasks, actions are represented as continuous end-effector commands. Each action is a 7-dimensional vector containing translational and rotational commands together with a gripper control signal. The **Physion Drape** task, however, involves passive deformable dynamics of a falling cloth and does not include actions. State observations adopt an object-centric representation that factorizes the world state into robot, rigid-object, articulated-object, or deformable-object components, depending on the task. Static configurations encode asset descriptors and physical attributes, including geometry, joint structure, friction, density, and global scene parameters. This unified episodic format allows OrbiSim to handle both action-conditioned manipulation and passive physical prediction within the same data pipeline.

A.1 Robosuite Push Task

Based on the robosuite simulation engine, we construct a trajectory dataset for a cube pushing and hitting task. In each simulation trial, several physical parameters were randomized, including cube friction, density, size, and visual appearance. The pushing velocity of the manipulator was also varied, as it significantly affects contact dynamics and task outcomes.

The task involves two cubes placed on a table and a UR5e manipulator. During an episode, the manipulator pushes the first cube, which then collides with the second cube. A rollout is considered successful if the first cube remains on the table and the second cube reaches a designated region near the target edge without falling off the other edges.

To ensure broad coverage of the task domain, the dataset was collected using a fixed multi-stage policy that generates trajectories from three distinct phases: (a) random manipulator movements, (b) pushing a cube toward a random direction, and (c) pushing a cube to hit another cube. Specifically, the manipulator first moves to a position behind the first cube, then executes a push toward a sampled direction, and finally remains stationary.

During the simulation, we recorded the full physical states of both cubes and the manipulator, the executed control actions, as well as RGB images captured from a side-view camera. In total, the dataset contains 120 trajectories from phase (a), 640 trajectories from phase (b), and 1,840 trajectories from phase (c). All trajectories feature randomized physical properties and visual appearances to promote diversity and robustness.

A.2 Robosuite Tasks with Multi-Objects

To introduce scalability and variety to the robosuite **Push** task, we subsequently construct two more tasks in robosuite.

The first task, **Push_cubes**, is an extension of the robosuite **Push** task, which also involves pushing and hitting cubes. The task is separated into four subtasks, namely, (a) randomly pushing (**random**), (b) pushing stacked cubes (**stacked**), (c) pushing toward an object (**towards_object**), and (d) pushing to hit an object (**to_hit**). The first two subtasks have randomized 3 to 4 objects, extending the number of objects. After finishing the task, the robot arm is set to randomly move with smoothed trajectories in order to extend the robot arm position coverage.

The second task is robosuite **Pick_and_place**, which requires the gripper to randomly pick cubes and place them. This task is separated into three subtasks: (a) stack cubes, (b) unstack initially stacked cubes, and (c) pick and place scattered cubes. The cubes' positions, robot arm picking order, moving keypoint trajectories, and cubes dropping heights are randomized during rollout.

The above two tasks contain most of the interaction types of robot arms in robosuite, including pushing, hitting, falling, gripper picking, etc. Each subtask of the two tasks has 420 episodes collected.

Table 4: **Articulated-object assets used in the AdaManip task.** We use nine categories with different articulated mechanisms and operable parts.

Category	# Instances	Mechanism / Operable Parts
Bottle	42	cap; rotate-and-slide interaction
Pen	26	cap; rotate-and-slide interaction
Coffee Machine	17	portafilter; rotate-and-slide interaction
Window	18	handle; lock and random rotation direction
Pressure Cooker	18	handle / lid; rotate-and-slide interaction
Lamp	24	button / knob; push-or-rotate ambiguity
Door	31	handle; lock and random rotation direction
Safe	18	door / knob; lock and switch-contact interaction
Microwave	16	door / button; lock and switch-contact interaction

A.3 Isaac Lab Stack Task

We additionally construct a trajectory dataset based on the Isaac Lab cube stacking task, where a Franka manipulator is required to sequentially stack multiple cubes in a physics-based simulation environment (**Isaac-Stack-Cube-Franka-IK-Rel**).

The dataset consists of episodic trajectories collected from interactions in the simulation, with an average episode length of approximately 230 timesteps. Each trajectory records the full sequence of states and actions throughout an episode. The final dataset includes approximately 1,000 successful trajectories and 1,200 failure trajectories, providing balanced coverage of diverse task execution outcomes.

The dataset is derived from a small set of human expert demonstrations. Starting from these demonstrations, we follow the standard Isaac Lab Mimic pipeline to automatically generate additional state-based trajectories, substantially increasing the scale and diversity of the dataset. This process enables the construction of a large collection of trajectories while relying on only a limited number of manually recorded demonstrations.

After data generation, the outputs produced by the official Isaac Lab scripts are reorganized and redirected into a unified episodic format that conforms to the dataset structure. This reorganization step preserves the original trajectory content while making the data compatible with our world model learning and downstream reinforcement learning pipelines.

For downstream reinforcement learning, we formulate a sparse episodic control task on top of the stacking environment. The reward is defined based on the final stable height of the stacked objects, measured after all object velocities have decayed to zero. This reward is provided once per episode and directly reflects task completion quality.

A.4 AdaManip Articulated Task

We construct an articulated-object manipulation dataset based on AdaManip [39]. The dataset covers nine categories of articulated objects, with a total of 210 asset instances. Each asset is represented as a structured articulated object, consisting of geometry files together with an articulation specification that defines movable links, joints, and task-relevant operable parts. The assets span diverse mechanisms, including rotating and sliding caps, lockable handles, push-or-rotate controls, and switch-contact interactions. Table 4 summarizes the asset categories used in our experiments.

For data collection, we follow the policy and data-generation scripts provided by AdaManip. For each asset instance in every object category, we collect 20 episodes, and each episode contains 301 simulation steps. This results in 4,200 articulated-object manipulation episodes in total. Each episode records the executed robot actions, structured robot and object states, articulated joint states, static asset descriptors, and rendered visual observations, and is converted into the same episodic format used by the other tasks.

During trajectory generation, the camera pose and the initial robot state are kept fixed across episodes to ensure a consistent observation and control setup. To introduce controlled variation, we randomize the initial pose of the manipulated articulated object: its position is perturbed within ± 5 cm, and its orientation angle is perturbed within $\pm 10^\circ$. This procedure preserves the underlying articulated

mechanism of each asset while providing diverse initial configurations for learning object-centric articulated dynamics.

A.5 Physion Drape Task

We adopted Physion’s point cloud dataset for arbitrarily shaped, deformable object dynamics. The Drape scene in Physion includes both irregularly shaped objects and a deformable cloth. At the start, a cloth is placed flat in the air, while several objects are placed on the floor. The objects include balls, cuboids, cones, cylinders, lamps, chairs, and so on. Then, the cloth in the air falls down, draping over the objects on the floor. During the scene, the cloth moves and deforms according to the objects’ shapes, which requires OrbiSim-Dynamics to learn the deformation dynamics. The Physion dataset does not cover physical parameters or actions.

B Implementation Details

B.1 OrbiSim-Dynamics

We train the dynamics module using a two-stage procedure to balance one-step prediction accuracy and long-horizon rollout stability. In the first stage, the model is optimized with teacher forcing, where ground-truth physical states are provided as inputs at every timestep. This stage focuses on learning accurate single-step state transitions while avoiding early error accumulation. In the second stage, we fine-tune the model under autoregressive rollouts by gradually replacing ground-truth inputs with the model’s own predictions. The transition from teacher-forced training to fully autoregressive execution is controlled by an autoregressive ratio that increases over training epochs following a cosine annealing schedule, improving robustness to compounding errors.

The dynamics model predicts explicit physical states for the robot and all manipulated objects. The robot state consists of position, linear velocity, and orientation represented as a flattened rotation matrix. Each object state includes position, linear velocity, and orientation represented as a unit quaternion.

We supervise state predictions using a composite state loss, where robot and object losses are computed separately:

$$\mathcal{L}_{\text{state}} = \mathcal{L}_{\text{robot}} + \mathcal{L}_{\text{object}}. \quad (8)$$

For translational position and velocity, we apply a Smooth- ℓ_1 loss. For rotational components, we use geometry-aware distance metrics: The robot orientation loss is defined as the geodesic distance between predicted and ground-truth rotation matrices on $\text{SO}(3)$, while object orientations are supervised using a normalized quaternion distance. Object losses are first averaged over objects and then over the batch.

During evaluation, we report a trajectory error (TrajErr) using the same state loss formulation. Specifically, an inference module maps generated video sequences to explicit physical state trajectories, which are then compared against ground-truth trajectories using $\mathcal{L}_{\text{state}}$. This metric directly measures long-horizon physical consistency of model rollouts and is aligned with the training objective of the dynamics module.

Model parameters are optimized using AdamW with a cosine-decayed learning rate schedule across all training stages.

B.2 OrbiSim-Vision

OrbiSim-Vision is trained as a latent diffusion model that generates visual observations conditioned on predicted physical states and scene descriptors. In this section, we focus on training objectives and data sampling strategies that improve robustness under different rollout conditions.

We adopt the standard diffusion denoising objective in the latent space. Each target frame is first encoded into a latent representation using a VAE encoder, and a forward diffusion process is applied to obtain noisy latents. The denoising network is trained to recover the clean latent conditioned on the predicted physical state and available context frame latents.

To improve robustness to varying rollout settings, we randomize the number of context frames during training. Specifically, for each training sample, we uniformly sample the context length $K \in \{0, 1, \dots, 6\}$

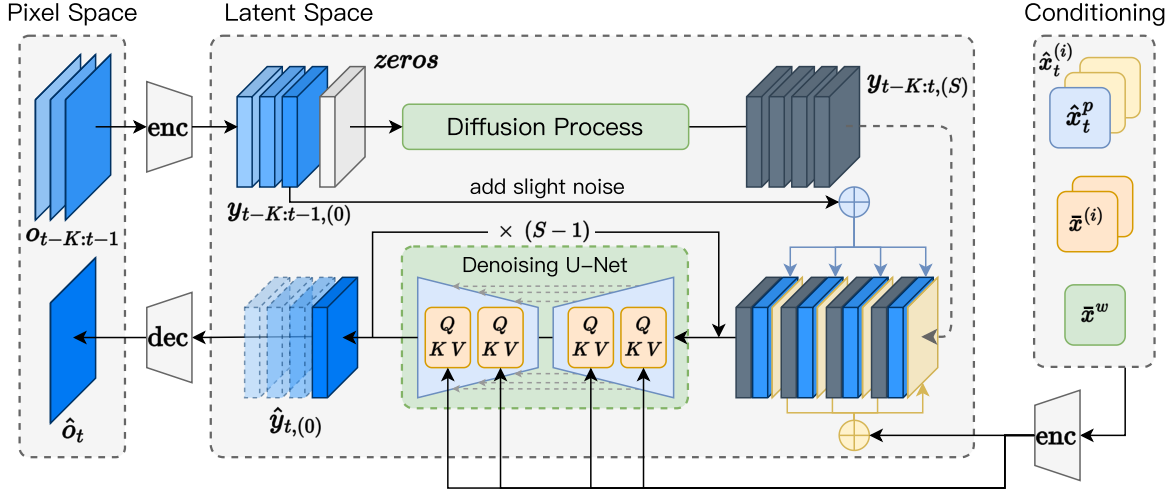


Figure 8: **OrbiSim-Vision**. The denoising process is grounded in predicted physical states through spatial condition maps and object tokens, with context frames providing supplementary visual details.

and condition the model on the corresponding past frames. Importantly, we include the case $K = 0$, which explicitly trains the model to reconstruct images solely from the given physical state without relying on visual context. This encourages stable state-to-image rendering when context frames are absent or unreliable.

In addition to varying the context length, we randomize how context frames are selected. With a certain probability, we use temporally consecutive context frames and predict the next consecutive frame. Otherwise, we randomly sample $(K+1)$ frames from a full episode at different timestamps, sort them by temporal order, and use the first K frames as context to predict the $(K+1)$ -th frame. By increasing the temporal gaps between context frames, this stochastic sampling strategy reduces over-reliance on short-term visual continuity and encourages the model to generate observations that are consistent with the provided physical state.

B.3 Real-to-Sim Module

The Real-to-Sim module converts visual observations into the structured physical inputs required by OrbiSim-Dynamics. As discussed in Sec. 3.4, directly identifying physical states and scene descriptors from pixels is challenging because visual observations entangle object pose, appearance, and hidden physical properties. We therefore decompose the inference process into two modules, as illustrated in Figure 9: a state inference module for visible scene quantities and a physics inference module for hidden physical parameters.

State inference. The state inference module predicts visible physical states and object attributes from a single RGB frame. The input frame is first encoded by the frozen OrbiSim-Vision encoder, and the resulting latent feature is passed to a lightweight prediction head. The module estimates quantities that can be inferred from one image, including the robot end-effector position and orientation, object positions, object sizes, and visual attributes such as color. These predictions provide the initial physical state and visually grounded object descriptors used by the downstream dynamics model.

Physics inference. The physics inference module estimates hidden physical parameters that are difficult to determine from a single frame. It takes as input the latent features of a short video window of length $W = 64$, extracted by the same frozen OrbiSim-Vision encoder, and predicts scene-level or object-level physical attributes such as object density and table friction. Rather than training this module only by regressing to raw simulator parameters, we optimize it through the frozen OrbiSim-Dynamics module. This design is important because different physical parameter values may produce nearly equivalent trajectories, making direct parameter regression unnecessarily restrictive. By contrast,

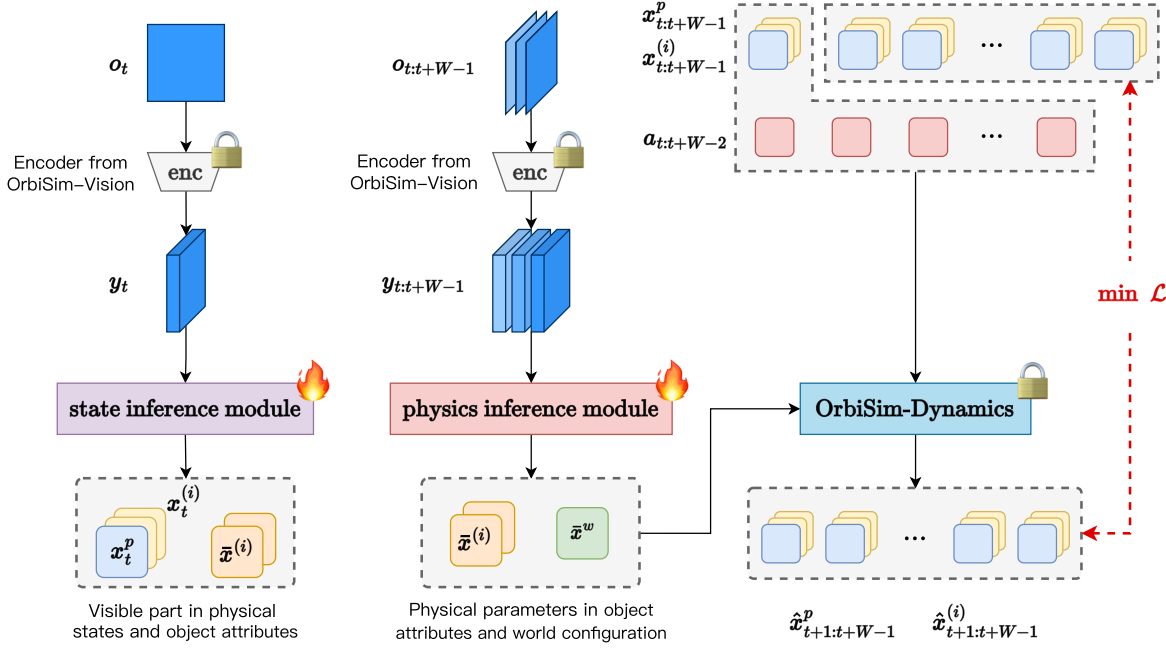


Figure 9: **Real-to-Sim inference pipeline.** The state inference module predicts visible physical states and object attributes from a single frame, while the physics inference module estimates hidden physical parameters from a 64-frame video. Both modules use latent features extracted by the frozen OrbiSim-Vision encoder. The inferred descriptors are fed into the frozen OrbiSim-Dynamics module, and the physics inference module is optimized through rollout consistency with the observed trajectory.

rollout-based supervision encourages the inferred parameters to reproduce the observed motion over the video window, which better matches the goal of system identification.

Inference pipeline. At inference time, the state inference module first recovers the visible initial state and object attributes from the initial frame. The physics inference module then infers hidden physical parameters from the observation window. The inferred state, object attributes, physical parameters, and action sequence are combined and passed into the frozen OrbiSim-Dynamics module to roll out future physical trajectories. When visual predictions are needed, the predicted states can be further rendered by OrbiSim-Vision. This pipeline provides a structured bridge from visual observations to simulation-ready physical descriptors, enabling OrbiSim to support Real-to-Sim prediction and system identification from videos.

B.4 Extension to Complex Objects

We extend OrbiSim to more complicated object types by adapting the construction of static scene descriptors and physical states while preserving the same object-centric dynamics formulation. For articulated objects, the key challenge is to encode structured multi-part geometry together with articulation priors. For deformable objects, the main difficulty is that object geometry evolves over time and therefore cannot be treated purely as a static asset attribute.

Articulated objects. For articulated objects, we represent each asset as a collection of semantic part-level tokens rather than encoding the entire object as a single rigid entity. Concretely, we first use the URDF specification to determine the part decomposition of the asset, the correspondence between each part and its associated mesh, and the articulation metadata, including joint type, axis, limits, and parent-child relations. For each part, we then recover its triangle mesh, transform it into the corresponding part-local visual frame according to the URDF scale and visual origin, and uniformly sample surface point clouds with normals. The sampled point cloud is normalized and passed through a pretrained Sonata [40]/PTV3 [41] encoder, and the resulting point features are deterministically

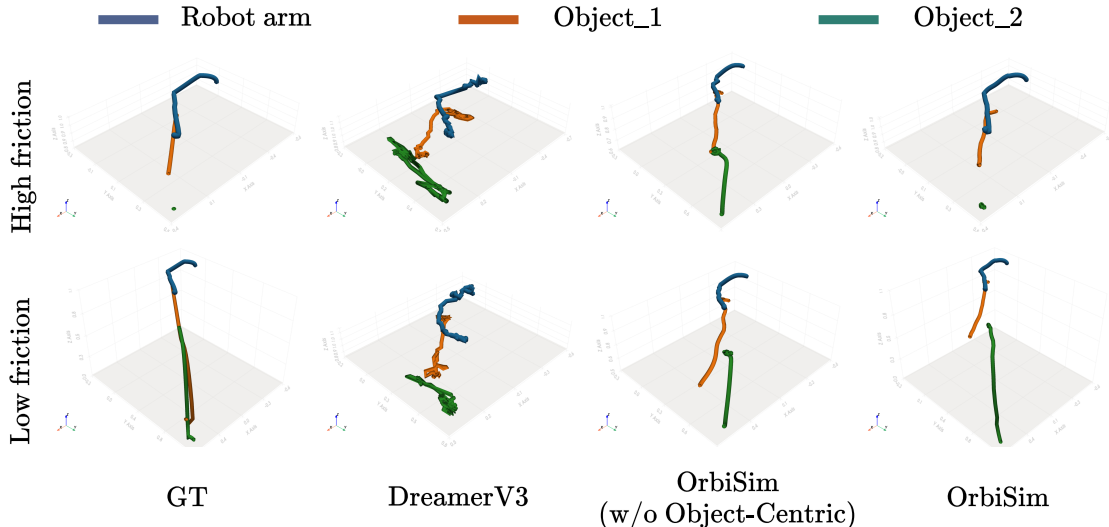


Figure 10: **Physical state trajectories by OrbiSim-Dynamics.** The autoregressive rollouts demonstrate high sensitivity to physical parameters, accurately distinguishing between low- and high-friction regimes while maintaining smoothness over long horizons.

pooled into a fixed-dimensional latent representation. The final part-level latent, together with the associated geometric and joint metadata, is stored as the asset attribute of that part. In this way, an articulated object is converted into a structured set of part-level descriptors that capture both geometry and articulation priors, and can be directly incorporated into the static scene descriptor of OrbiSim-Dynamics.

Deformable objects. For deformable objects, we adopt a different treatment because the object shape is inherently time-dependent. In the Drape [2] environment, the static objects on the floor remain fixed throughout the episode and, for simplicity, are treated as one large static object, while the cloth is modeled as a separate deformable object. We extract point-cloud-based embeddings for both the static scene object and the deformable cloth using Point Transformer V3 [41], yielding object-level features whose dimensionality depends only on the number of objects rather than the number of sampled points. This matches the object-centric encoder design of OrbiSim-Dynamics and allows deformable scenes to be represented within the same structured input interface. Unlike rigid or articulated assets, however, the shape of the cloth changes over time and therefore cannot be encoded only as a static attribute. Instead, the corresponding shape features are treated as part of the physical state and updated along the temporal rollout. This design enables OrbiSim to naturally extend from rigid-body manipulation to deformable dynamics while maintaining a unified object-centric representation.

C More Experimental Results

C.1 OrbiSim-Dynamics Evaluation

We evaluate the state-space generative performance of OrbiSim-Dynamics by benchmarking it against DreamerV3-Torch on the robosuite offline dataset. As illustrated in Figure 10, our dynamics core effectively models complex state transitions while demonstrating a high degree of sensitivity to underlying physical configurations.

Specifically, the results show that OrbiSim-Dynamics accurately distinguishes between low- and high-friction regimes, adjusting object trajectories accordingly while maintaining predictive smoothness over long-horizon rollouts. This responsiveness to varying object-table friction coefficients confirms the model’s capacity to ground its learned dynamics in explicit physical parameters, a critical requirement for robust system identification and transferable control.

Table 5: **Analysis of OrbiSim rollout variants.** (i) **GT-first:** Autoregressive rollout given true $\{o_0, x_0, \bar{x}\}$; (ii) **GT-state:** Uses true state sequences throughout to isolate rendering performance; (iii) **Recon-first:** Autoregressive rollout only from true initial state; (iv) **Recon-all:** Evaluates per-frame state-to-pixel fidelity by excluding all temporal visual context during generation, *i.e.*, $p_\phi^{\text{vis}}(o_t | \hat{x}_t, \bar{x})$. **TrajErr:** Measures the discrepancy between *inferred physical states from generated videos* (see Sec. 3.4) and the ground truth.

Method	PSNR10 $\uparrow$	PSNR100 $\uparrow$	LPIPS10 $\downarrow$	LPIPS100 $\downarrow$	FVD $\downarrow$	TrajErr $\downarrow$
OrbiSim (GT-first)	26.7105	21.8269	0.1078	0.1428	533.9	0.4468
OrbiSim (GT-state)	26.7253	20.2680	0.1087	0.1388	537.3	0.4322
OrbiSim (Recon-first)	24.6366	19.3114	0.1236	0.1509	553.6	0.4558
OrbiSim (Recon-all)	26.4910	23.3776	0.1279	0.1628	922.1	0.2927

Table 6: **Preliminary evaluation of the Real-to-Sim modules.** **StaInf** predicts visible scene states and object attributes from a single frame, while **PhyInf** infers hidden physical parameters from a 64-frame video window and is evaluated through rollout consistency under frozen OrbiSim-Dynamics.

Module	Object Metrics					Robot Metrics		
	Position $\downarrow$	Rotation $\downarrow$	Size $\downarrow$	Color(MAE) $\downarrow$	Velocity $\downarrow$	Position $\downarrow$	Velocity $\downarrow$	Rotation $\downarrow$
StaInf	27.09mm	7.37°	5.2%	0.0131	—	22.44mm	—	4.08
PhyInf	22.40mm	4.28°	—	—	28.69mm/s	12.70mm	28.63mm/s	1.02°

C.2 Workflow Analysis

To disentangle the effects of grounding, autoregressive feedback, and visual context, we evaluate four distinct execution workflows: • *GT-first* initializes rollouts with ground-truth states and observations, after which OrbiSim autoregressively predicts both future states and images. • *Recon-first* initializes rollouts from a ground-truth state only, generating the first observation from the state without visual context and then rolling out autoregressively. • *GT-state* replaces autoregressive state predictions with ground-truth state sequences, evaluating visual rendering under perfectly grounded dynamics. • *Recon-all* predicts each observation independently from the corresponding physical state, without using any temporal visual context.

As shown in Table 5, a compelling trade-off is observed in the **Recon-all** variant. It attains substantially higher PSNR due to its bias toward per-frame reconstruction accuracy; however, its degradation in perceptual similarity (LPIPS) and temporal coherence (FVD) indicates a lack of long-horizon visual consistency. In contrast, **GT-first** and **GT-state** strike a superior balance, demonstrating that OrbiSim effectively utilizes temporal context to stabilize long-horizon rollouts.

C.3 Real-to-Sim Experiments

We evaluate the proposed Real-to-Sim pipeline by examining the two inference modules introduced in Sec. 3.4: the state inference module (StaInf) and the physics inference module (PhyInf). StaInf predicts visible scene states and object attributes from a single RGB frame, while PhyInf infers hidden physical parameters from a 64-frame video window through the frozen OrbiSim-Dynamics model.

Table 6 shows that StaInf recovers visible scene quantities with high accuracy, including object position, rotation, size, and color, as well as robot pose. These results indicate that the frozen OrbiSim-Vision representation preserves sufficient geometric and appearance information for decoding structured physical states directly from visual observations.

We further evaluate PhyInf through rollout consistency, *i.e.*, whether the inferred hidden parameters are sufficient to reproduce the observed motion under the frozen OrbiSim-Dynamics model. As shown in Table 6, the resulting rollouts achieve low errors in both object and robot states, including positions, rotations, and velocities. This suggests that the inferred hidden physical parameters are informative enough to support simulation-consistent trajectory prediction and system identification from videos.

C.4 Multi Objects / Multi Tasks

Table 7: **Results on multi-object, multi-task datasets.** Single represents OrbiSim and AdaWorld trained on single-task robosuite `Push` dataset with up to 2 objects. Multi represents OrbiSim and AdaWorld trained on the extended `Push_cubes` and `Pick_and_place` dataset with up to 4 objects. For data with more than 100 frames, the results end with “100” are from all frames.

Method	PSNR10 $\uparrow$	PSNR100 $\uparrow$	LPIPS10 $\downarrow$	LPIPS100 $\downarrow$	FVD $\downarrow$
OrbiSim (single)	26.7105	19.9819	0.1078	0.1428	533.9
AdaWorld (single)	26.6647	12.8346	0.1183	0.3482	1305.8
OrbiSim (multi)	24.8676	18.4515	0.1210	0.1959	863.1
AdaWorld (multi)	24.4685	14.3235	0.1332	0.2661	1292

We conduct a scaling study focusing on predictive fidelity and memory overhead in multi-task multi-object scenarios.

For predictive fidelity, we evaluate OrbiSim and AdaWorld baseline in the extended multi-object robosuite dataset, namely `Push_cubes` and `Pick_and_place`. Dataset configuration and collection details are discussed in Appendix A.2. These task scenes contain more than 2 objects with a wide range of task settings. These tasks include not only cube pushing, hitting, and friction, but also objects falling, stacking, and robot gripper picking. Thus, these two categories are primitive and fundamental tasks for most robot arm tasks, requiring OrbiSim to learn 3-dimensional object positioning, interaction, and gripping.

The main results of multi-object and multi-task world prediction are in Table 7. Although it experiences performance degradation in multi-object, multi-task scenarios, OrbiSim still achieves higher visual quality and consistency scores.

Additionally, we measured the training memory overhead as the number of objects increased from 2 to 8. The results are shown in Figure 11. As shown in the results, VRAM usage grows approximately linearly. The VRAM usage of OrbiSim-Dynamics is nearly proportional to the number of objects, while OrbiSim-Vision’s VRAM growth is tiny compared to its original overhead. This demonstrates that the Transformer-based coupling remains computationally manageable for medium-scale scenes.

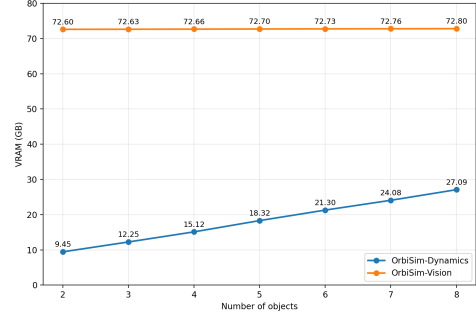


Figure 11: **Memory overhead results from object number scaling.**

C.5 Robustness

We conduct experiments on noisy states to quantify OrbiSim’s robustness. We set two levels of Gaussian noise on physical and visual states:

- Moderate Noise: Small-scale perturbations to physical states.
- Strong Noise: Significant noise levels designed to simulate high-error visual state estimation.

Table 8: **Results on noisy states rollout.**

Method	PSNR100 $\uparrow$	LPIPS100 $\downarrow$	FVD $\downarrow$
OrbiSim	19.9819	0.1428	533.9
OrbiSim with moderate noise	20.2508	0.1309	587.14
OrbiSim with strong noise	20.1911	0.1337	606.98

As shown in Table 8, OrbiSim exhibits substantial resilience to noise on both physical and visual states. Even under high levels of visual noise, OrbiSim remains highly robust, exhibiting low degradation in visual generation. This demonstrates that OrbiSim’s framework can tolerate imperfection in physical

Table 9: **Imitation learning performance on the Isaac Lab Stack task.** The imitation learning baseline is implemented using the official Isaac Lab demonstration-based example without reinforcement learning. Results highlight the sensitivity of pure imitation learning to distributional shifts in long-horizon manipulation tasks.

Method	Input	Success (%)
Isaac Lab-IL (Vision)	RGB + State	78.53
Isaac Lab-IL (Noise)	Noise + State	0.00
Isaac Lab-IL (State)	State	78.00

and visual states, which is promising for potential Real-to-Sim scenarios even where ground-truth states are unavailable.

C.6 Imitation Learning Experiments on the Isaac Lab Stack Task

We further conduct imitation learning experiments on the Isaac Lab Stack task, which is substantially more challenging than the pushing task due to its longer horizons, multi-stage object interactions, and stricter requirements on precise state coordination. The stacking task requires not only reaching intermediate configurations, but also maintaining stability under contact-rich dynamics, making it particularly sensitive to compounding errors.

As shown in Table 9, the imitation learning baseline (**Isaac Lab-IL**) can achieve reasonable performance when evaluated under in-distribution conditions that closely match the demonstration data. However, its performance degrades catastrophically under out-of-distribution (OOD) settings, such as variations in object configurations or interaction dynamics. This failure mode highlights a fundamental limitation of pure imitation learning: without an explicit mechanism for modeling environment dynamics or long-horizon credit assignment, small deviations from the demonstration manifold can rapidly accumulate and lead to irreversible task failure.

These results motivate the use of reinforcement learning, particularly when augmented with learned world models, to enable robust long-horizon decision making and recovery under distributional shifts.

C.7 Orbisim-Dynamics Compared to Newton Simulator

To prove the effectiveness of analytical gradients of OrbiSim, we conducted experiments comparing OrbiSim to differentiable physics simulators. We adopted NVIDIA Newton as a baseline, with its `diffsim` tasks simulated in Warp differentiable simulator. The `diffsim_ball` task is chosen. At the start, a ball is in the air with an initial velocity. The task requires optimization over initial velocity, making the ball bounce on the floor and a static wall, and finally reach a target position. No robot arms or actions are included in this task. In newton simulator, the optimization is achieved by the raw gradient descent method, utilizing the gradient of the squared target position error through the simulator.

For optimization in OrbiSim, we first collect a dataset by rollout in Newton with randomly-sampled initial velocity. A total of 500 episodes are collected for OrbiSim-Dynamics training. After training, we optimized initial velocity through analytical gradients from OrbiSim-Dynamics with the same method as Newton.

With a learning rate of 0.8, gradient descent in Newton simulator reaches zero loss within 60 steps. We conduct gradient descent with the same learning rate in OrbiSim-Dynamics, and the loss value is reduced to less than $1e-4$ within 800 steps.

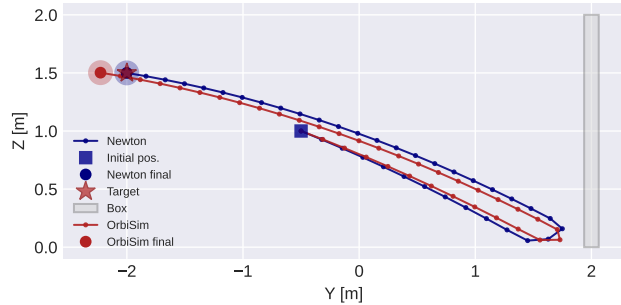


Figure 12: **Rollout trajectory in Newton simulator.** Both trajectories initialized with velocity optimized in Newton and OrbiSim-Dynamics are shown. OrbiSim-Dynamics’ result shows a tolerable difference.

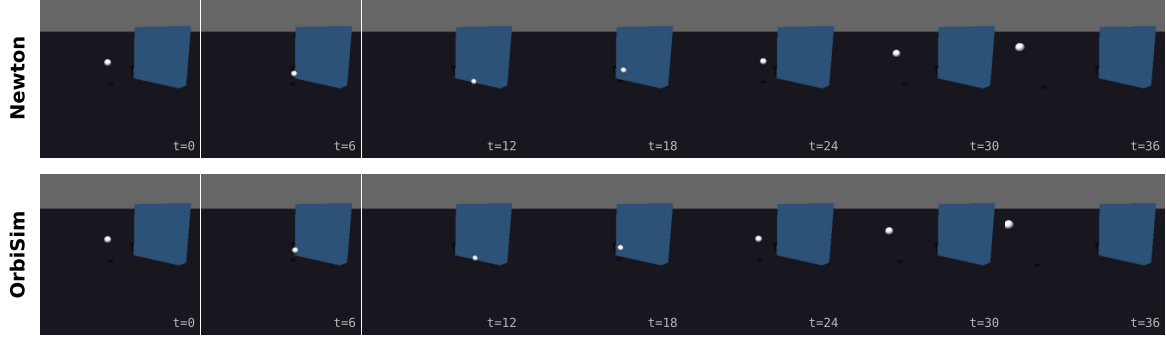


Figure 13: **Visual rollout in Newton simulator.** The two rows represent rollouts with initial velocity optimized with Newton warp and OrbiSim-Dynamics, respectively.

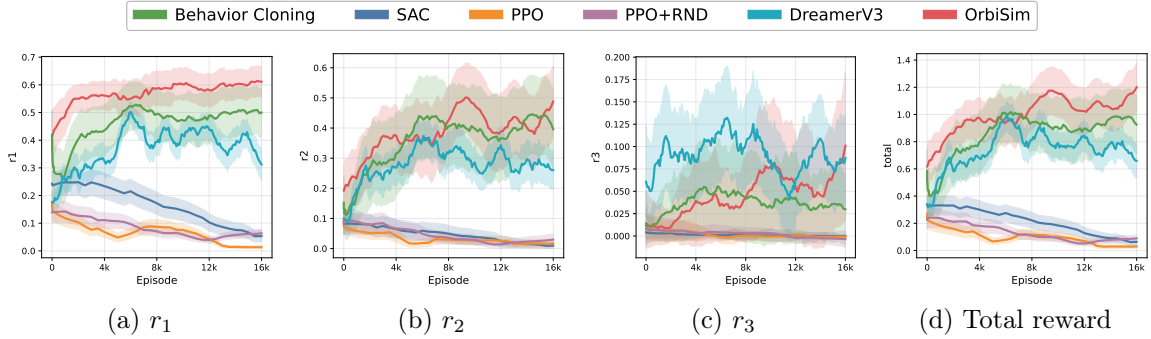


Figure 14: **Training curves on the robosuite Push task.** The x-axis denotes training episodes and the y-axis denotes normalized episode rewards defined in Sec. 4.3 and Eq. (9), namely r_1 , r_2 , r_3 , and the total reward. See reward definition in the text.

The visual results are shown in Figure 13 and Figure 12. The initial velocity optimized from newton is $[0.0, 10.639, -4.282]$, while the velocity optimized from OrbiSim-Dynamics is $[-0.07, 11.218, -4.238]$. A tolerable difference is shown in the velocity and rollout trajectories.

OrbiSim shows a slight error in dynamics estimation accuracy compared to simulators. However, differentiable simulators such as Newton DiffSim only support limited scenarios, while OrbiSim supports a wider variety of object types, scene configurations, and physics settings.

C.8 More Downstream Reinforcement Learning Experiments

In the robosuite environment, we define three normalized terminal rewards based on distance reduction: r_1 encourages approaching the first cube, r_2 promotes contact between the two cubes, and r_3 measures task progress by pushing the second cube toward the table boundary. Each reward is computed as the relative decrease from its initial distance. For r_1 and r_2 , the minimum distance over the episode is used to capture transient interactions, while r_3 uses the final distance to reflect task completion. Specifically, for an episode of T steps, let $d_{r, c1, t}$ denote the distance between the robot arm (r) and the first cube (c1) at timestep t , and the rewards are defined as:

$$r_1 = 1 - \frac{\min_{t=1}^T d_{r, c1, t}}{d_{r, c1, 0}}, \quad r_2 = 1 - \frac{\min_{t=1}^T d_{c1, c2, t}}{d_{c1, c2, 0}}, \quad r_3 = 1 - \frac{x_{\text{table_border}} - x_{c2, T}}{x_{\text{table_border}} - x_{c2, 0}}. \quad (9)$$

During training, these three rewards are provided only at episode termination, resulting in sparse learning signals that are challenging for conventional policy optimization methods. We also use these three rewards to evaluate task completion performance across different methods. Success is defined in robosuite Push task if the second cube has moved more than 5cm during rollout.

The detailed training curves are provided in Figure 14. The success rate comparison is provided in Table 10. OrbiSim performs higher task completion rewards and success rate than all baselines.

Table 10: **Success rate on the robosuite Push task.** We evaluate OrbiSim, model-based baseline (DreamerV3), Behavior Cloning (BC) and model-free baselines (SAC, PPO, PPO+RND) in the robosuite environment. Success is achieved if the second cube moves more than 5cm.

Method	OrbiSim	DreamerV3	BC	PPO+RND	PPO	SAC
Success Rate(%)	42.71	25.00	19.79	2.08	1.04	0.00

To qualitatively evaluate policy performance on the downstream task, we visualize representative rollouts produced by policies trained using model-free RL, model-based RL and Behavior Cloning (BC). For model-free RL, we include Soft Actor-Critic (SAC), Proximal Policy Optimization (PPO) and PPO with Random Network Distillation (PPO+RND). For model-based RL, we adopt DreamerV3 as a baseline.

As shown in Fig. 15, policies trained with model-free RL fail to learn effective behaviors for the downstream task, resulting in disordered motions and unsuccessful task completion. DreamerV3 policies perform more effectively but still fail in stability and generalization in unseen distributions. Behavior Cloning (BC) policies, while able to reproduce demonstrations to some extent, exhibit unstable performance and frequently deviate from desired behaviors due to the accumulation of prediction errors over long horizons. In contrast, our method consistently produces coherent and goal-directed behaviors, enabling successful task execution across different scenarios. These visual comparisons demonstrate that our approach provides more reliable policy performance on the downstream task compared with both reinforcement learning and imitation learning baselines.

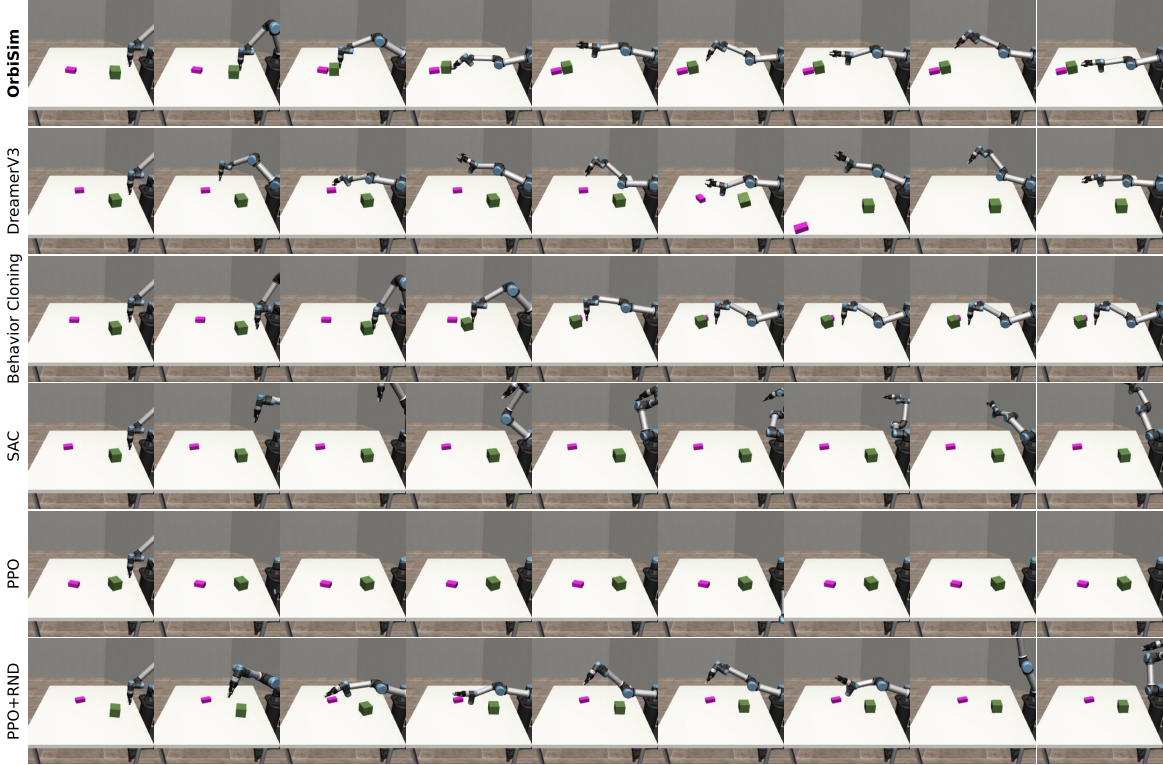


Figure 15: **Reinforcement learning results on robosuite Push task.** We show video rollouts in robosuite simulator of OrbiSim model-based policies trained with OrbiSim dynamics, DreamerV3 model-based policies trained with DreamerV3 prediction, behavior cloning trained with expert policies, and model-free baselines trained with robosuite environment.

C.9 Computational Efficiency and Resource Usage

We provide additional computational statistics for representative components and baselines used in our experiments. Table 11 reports the training throughput and GPU memory usage measured during downstream policy learning. All measurements are collected under the same simulator and training environment settings.

The results show that OrbiSim achieves substantially higher training throughput than latent world-model approaches such as DreamerV3 while maintaining moderate memory usage. Pure model-free baselines such as SAC achieve higher raw environment throughput, but fail to obtain competitive downstream control performance in our tasks. These observations support the practical benefit of structured end-to-end differentiability for efficient embodied policy optimization.

Table 11: **Training throughput and GPU memory usage of representative methods.**

Method	Steps/s	Device	GPU Memory
OrbiSim dynamics	29.41	RTX 3080	372 MB
Behavior Cloning	47.71	RTX 3080	27.14 MB
DreamerV3	4.54	RTX 4090D	931 MB
SAC	66.05	RTX 3080	33.79 MB