

World Models for Robotic Manipulation



RILS LAB

Robot Intelligence Learning & System

임찬우

2026. 08. 14

01

DreamGen: Unlocking Generalization in Robot Learning through Video World Models

Joel Jang, Seonghyeon Ye, Zongyu Lin, Jiannan Xiang, Johan Bjorck, Yu Fang, Fengyuan Hu, Spencer Huang, Kaushil Kundalia, Lin Yen-Chen, Loic Magne, Ajay Mandlekar, Avnish Narayan, You Liang Tan, Guanzhi Wang, Jing Wang, Qi Wang, Yinzhen Xu, Xiaohui Zeng, Kaiyuan Zheng, Ruijie Zheng, Ming-Yu Liu, Luke Zettlemoyer, Dieter Fox, Jan Kautz, Scott Reed, Yuke Zhu, Linxi Fan

NVIDIA, University of Washington, KAIST, UCLA, UCSD, CalTech, NTU, University of Maryland, UT Austin

CoRL 2025
arXiv May. 2025

<https://arxiv.org/abs/2505.12705>

Source: Google Scholar (Aug. 2026) / Citations: 162

Problem

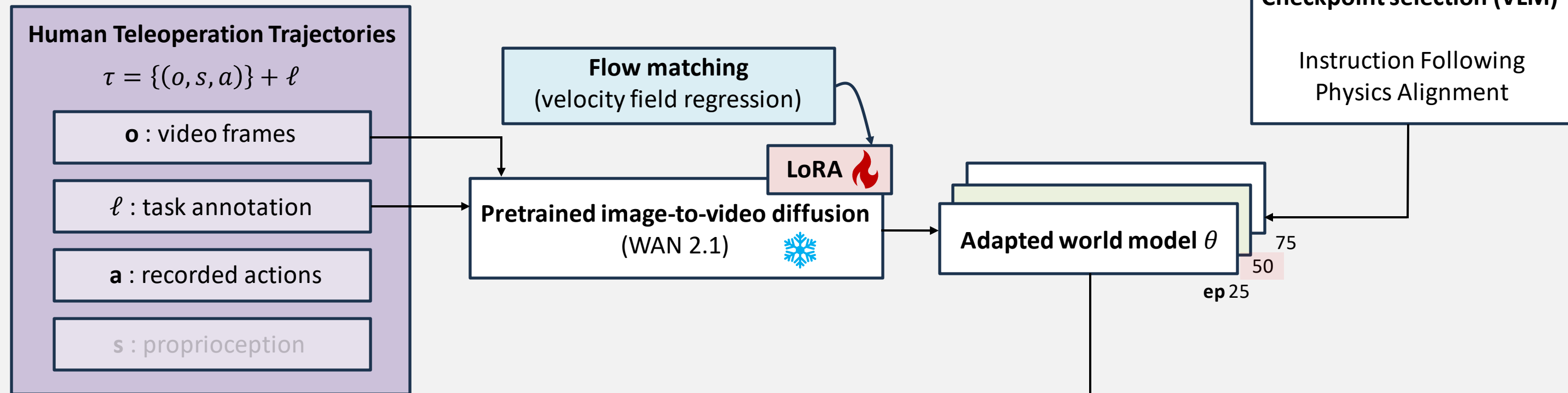


- **Teleoperation Cost** — Scales as tasks × environments
 - Every new task or environment requires **collecting data** from scratch by hand.
 - **Solution:** A **video world model** generates the data instead.
- **Simulation** — Sim-to-real gap & modeling limits
 - Simulated trajectories suffer from the **sim-to-real gap**, cannot model **liquids** or articulated objects, and are confined to **human demonstrations**.
 - **Solution:** **Bypass the simulator** and generate photorealistic video directly.
- **Missing Action Labels** — Video alone cannot train a policy
 - The world model emits pixels only. **Without actions**, behavior cloning has no regression target.
 - **Solution:** An **IDM** recovers pseudo-actions, completing the neural trajectory.

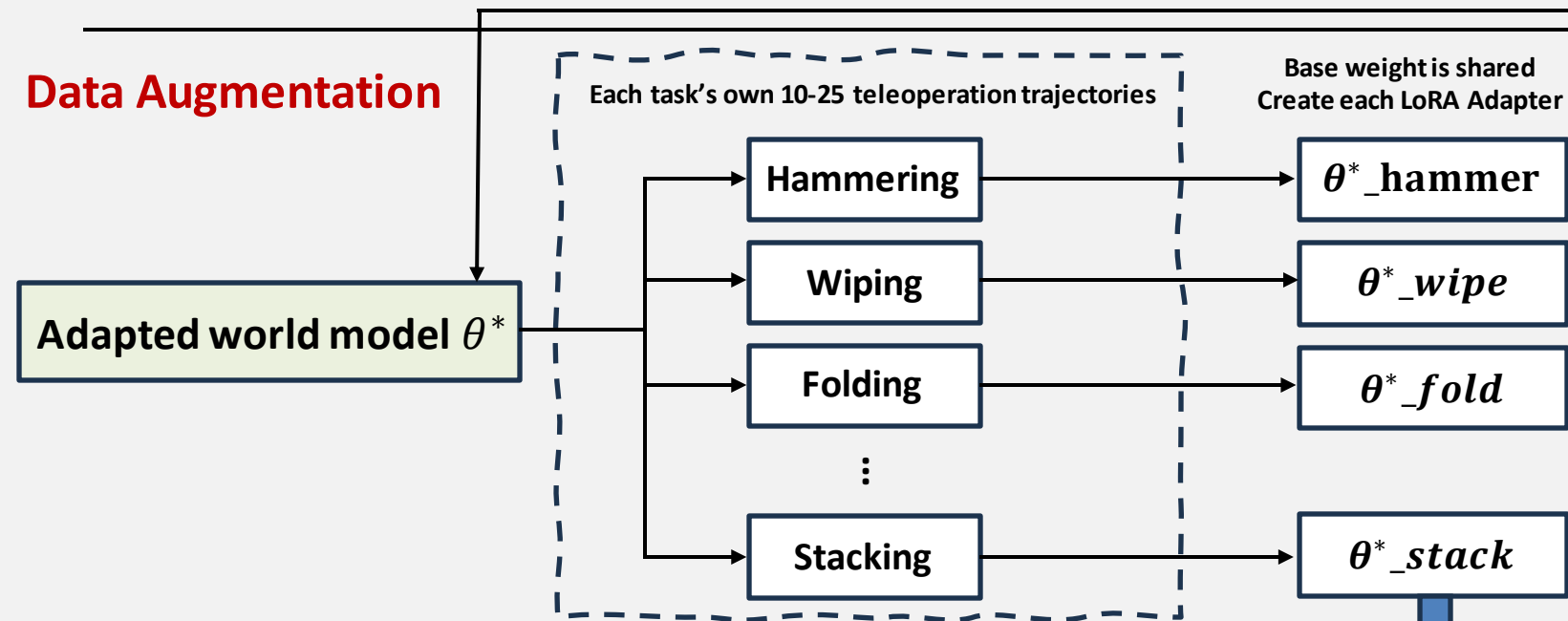
Method



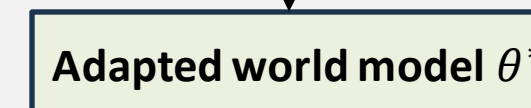
1. Video world model fine-tuning



Data Augmentation

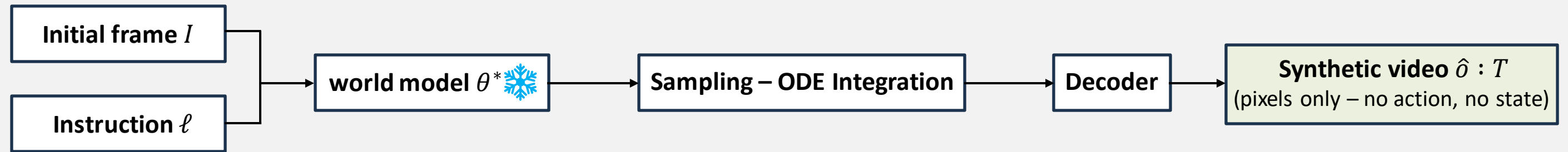


Generalization (no data)

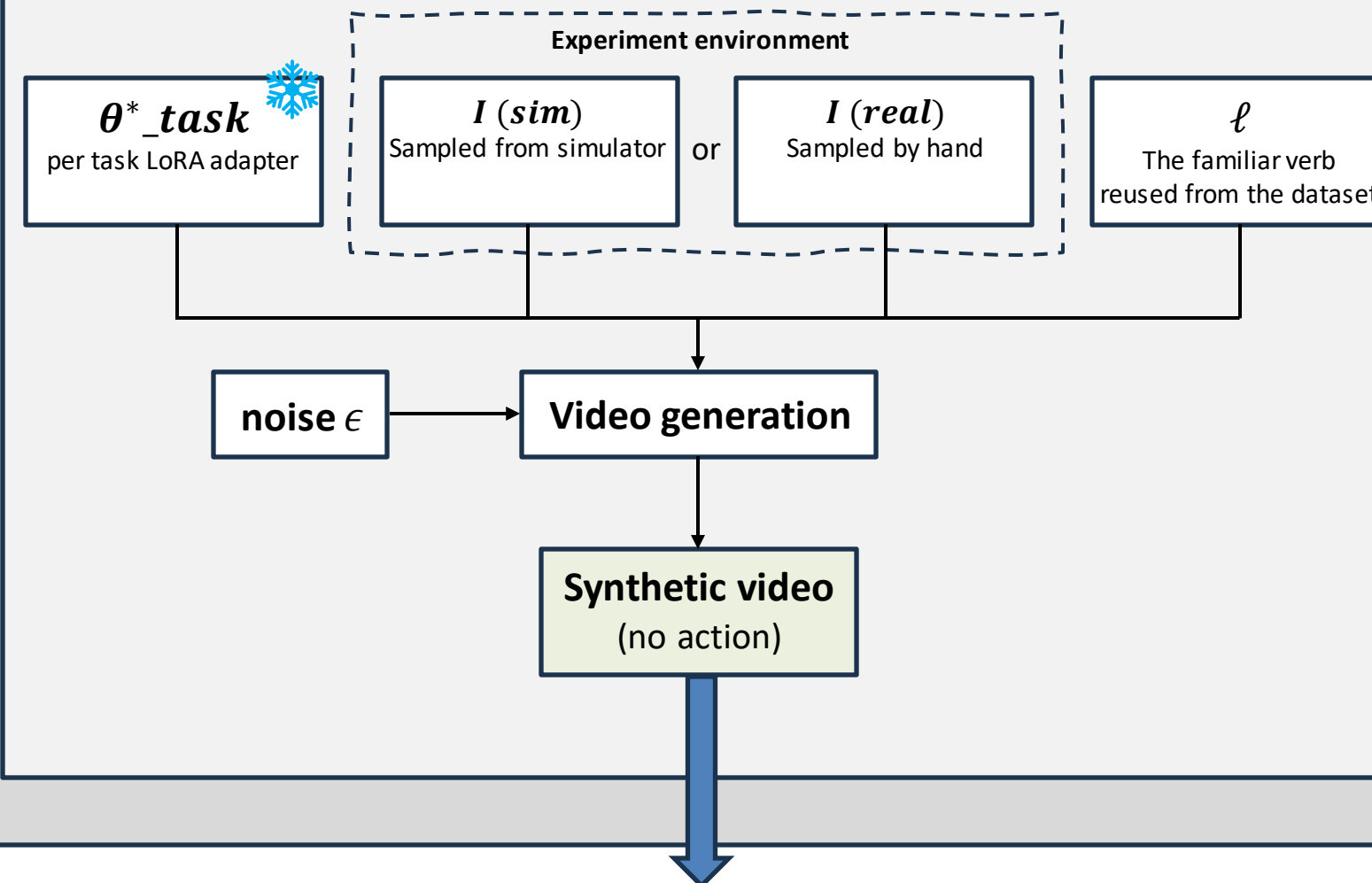


2. Video world model rollout

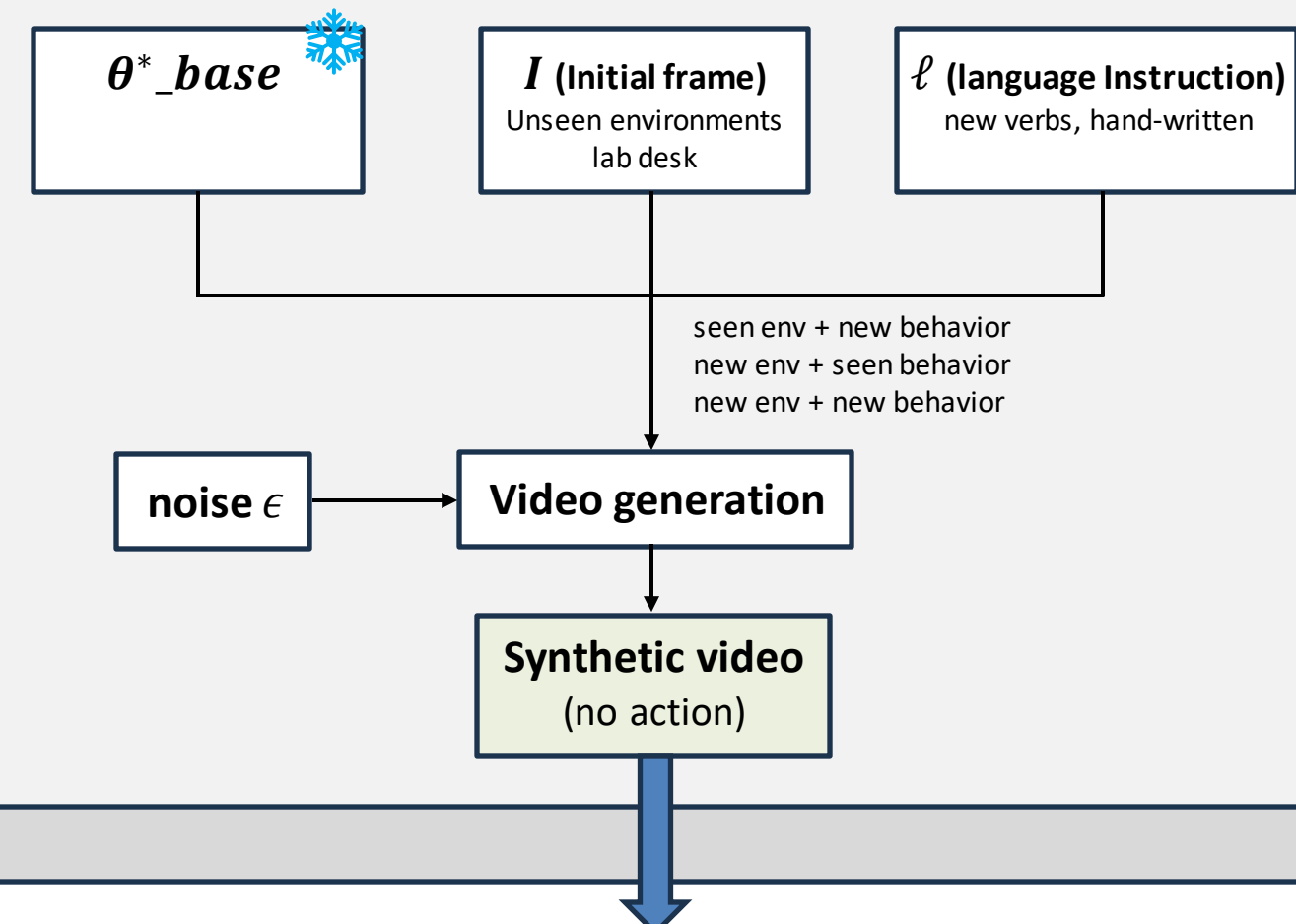
Shared mechanism



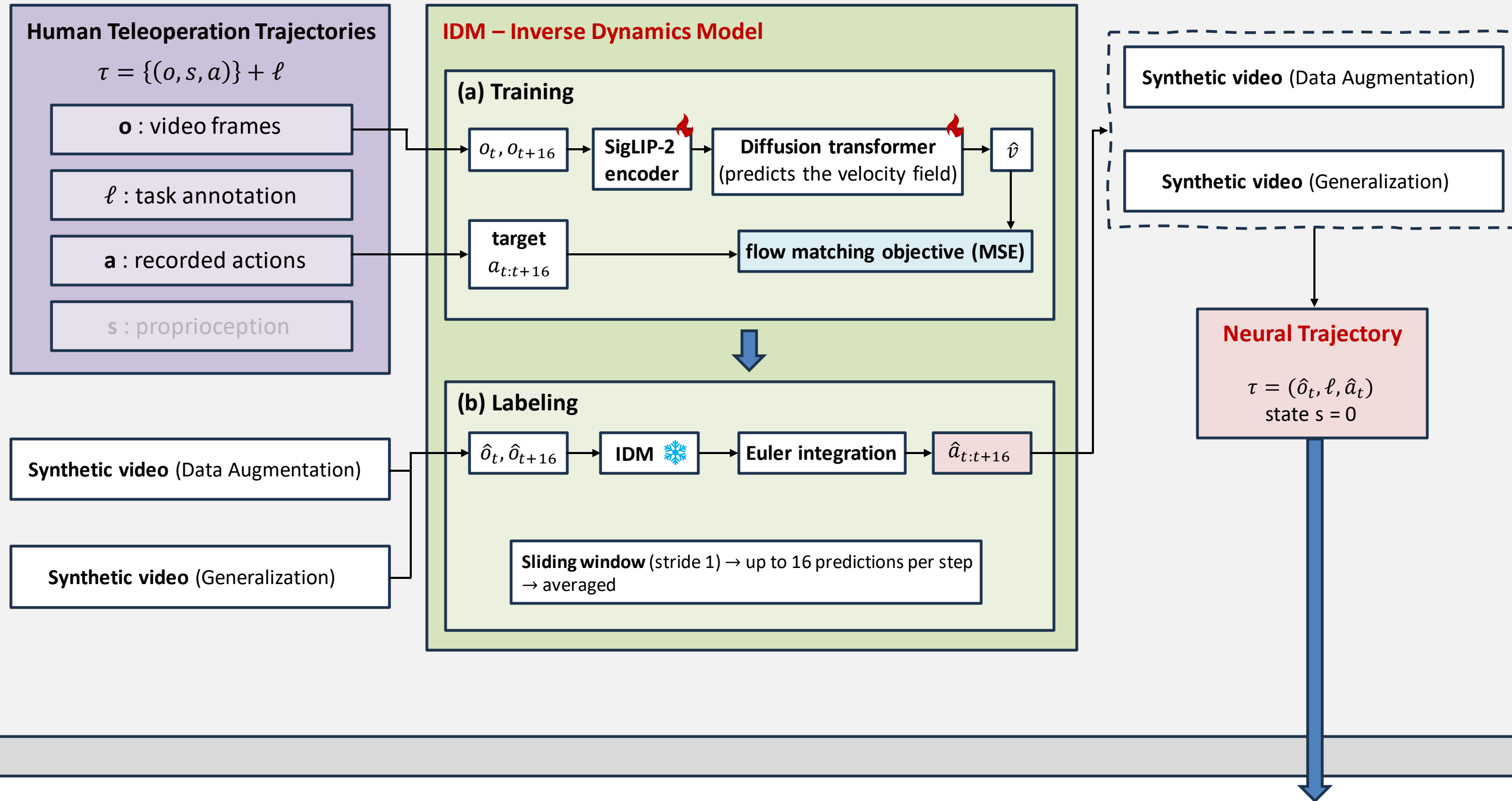
Data Augmentation



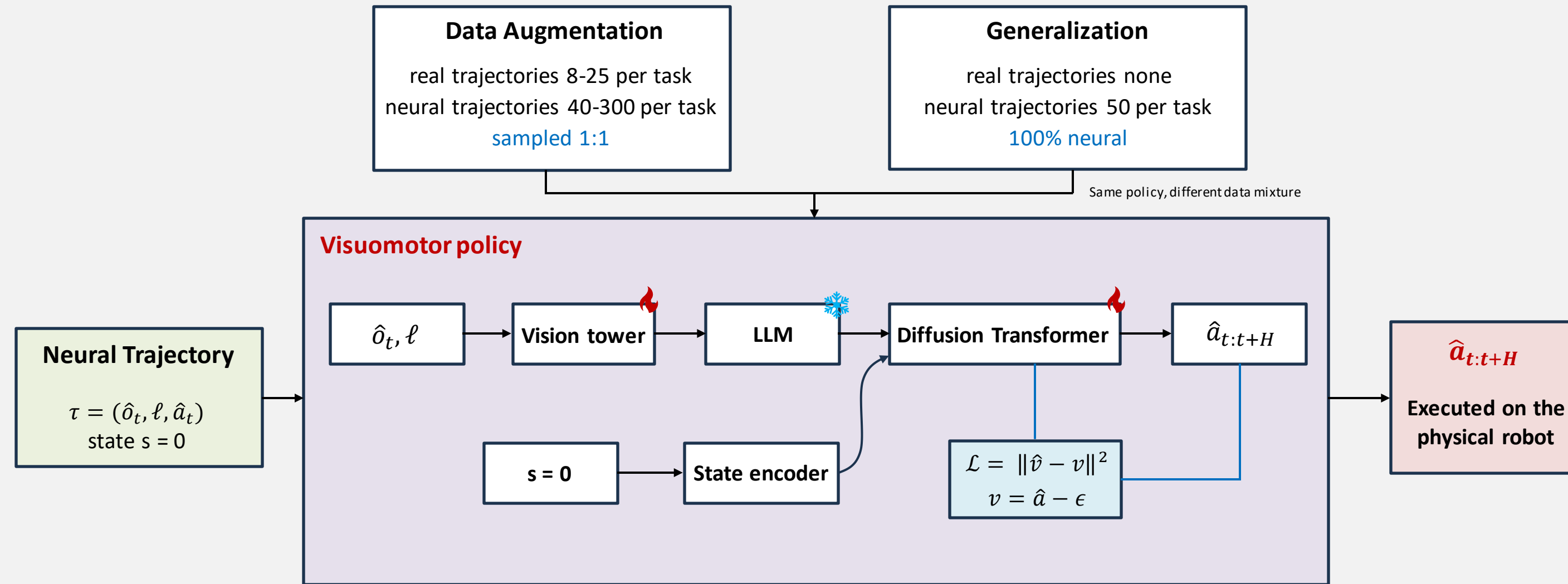
Generalization (no data)



3. Pseudo Action Labeling



4. Policy Training on Neural Trajectories



Results

GR1 (4 tasks)	37 → 46.4%
Franka (3 tasks)	23 → 37.0%
SO-100 (2 tasks)	21 → 45.5%

Seen env · new behavior (14 tasks)	11.2 → 43.2%
New env (13 tasks)	0.0 → 28.5%

02

Dream to Manipulate: Compositional World Models empowering robot Imitation Learning with Imagination

Leonardo Barcellona, Andrii Zadaianchuk, Davide Allegro, Samuele Papa, Stefano Ghidoni, Efstratios Gavves

University of Padova, University of Amsterdam, Politecnico of Torino, Archimedes/Athena RC

ICLR 2025

arXiv Dec. 2024

<https://arxiv.org/abs/2412.14957>

Source: Google Scholar (Aug. 2026) / Citations: 55

Problem



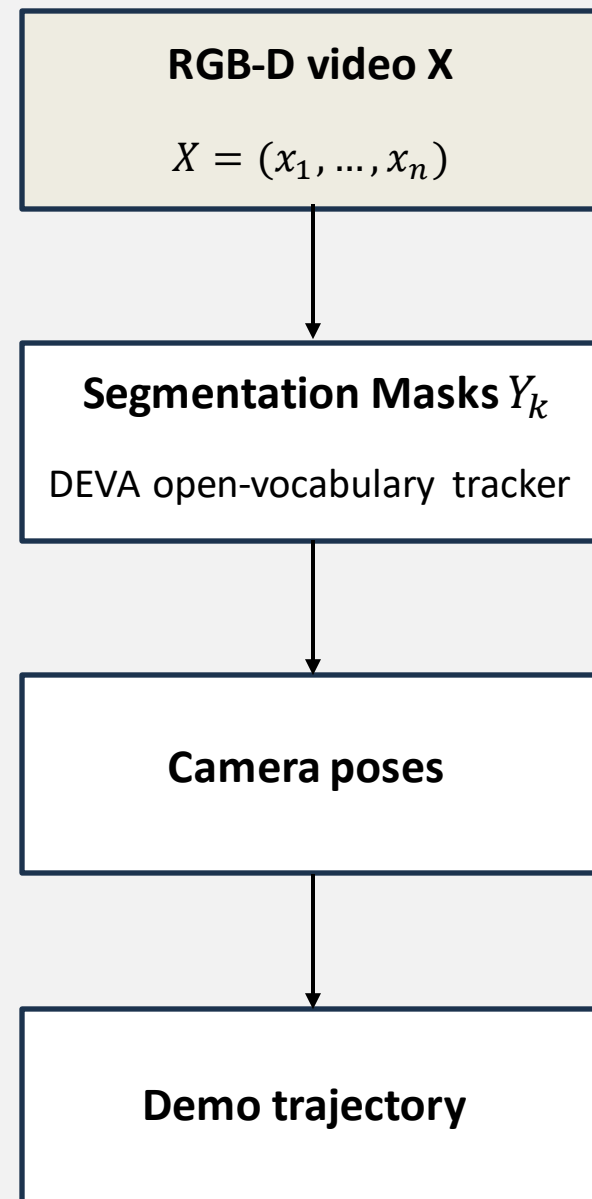
- **Demonstration Cost** — Policies need hundreds of demonstrations
 - Vision-based imitation learning still requires **hundreds of demos**, and those demos must cover enough variation.
 - **Solution:** Generate new demonstrations inside the **world model** — works with as few as one demo.
- **Learned World Models** — Break outside the training domain
 - Generative world models (UniSim, Dreamer, RoboDreamer) cannot capture **the real physics** of the robot's environment.
 - **Solution:** Do not learn the dynamics — plug a **physics engine** (PyBullet) into D directly.
- **Augmentation Validity** — Editing robot data can silently break the task
 - Changing an object's position **invalidates the action sequence**; prior methods only recombine existing trajectories.
 - **Solution:** **Equivariant transforms** move observations and actions together, then the **physics engine verifies** each result and discards invalid ones.

Method

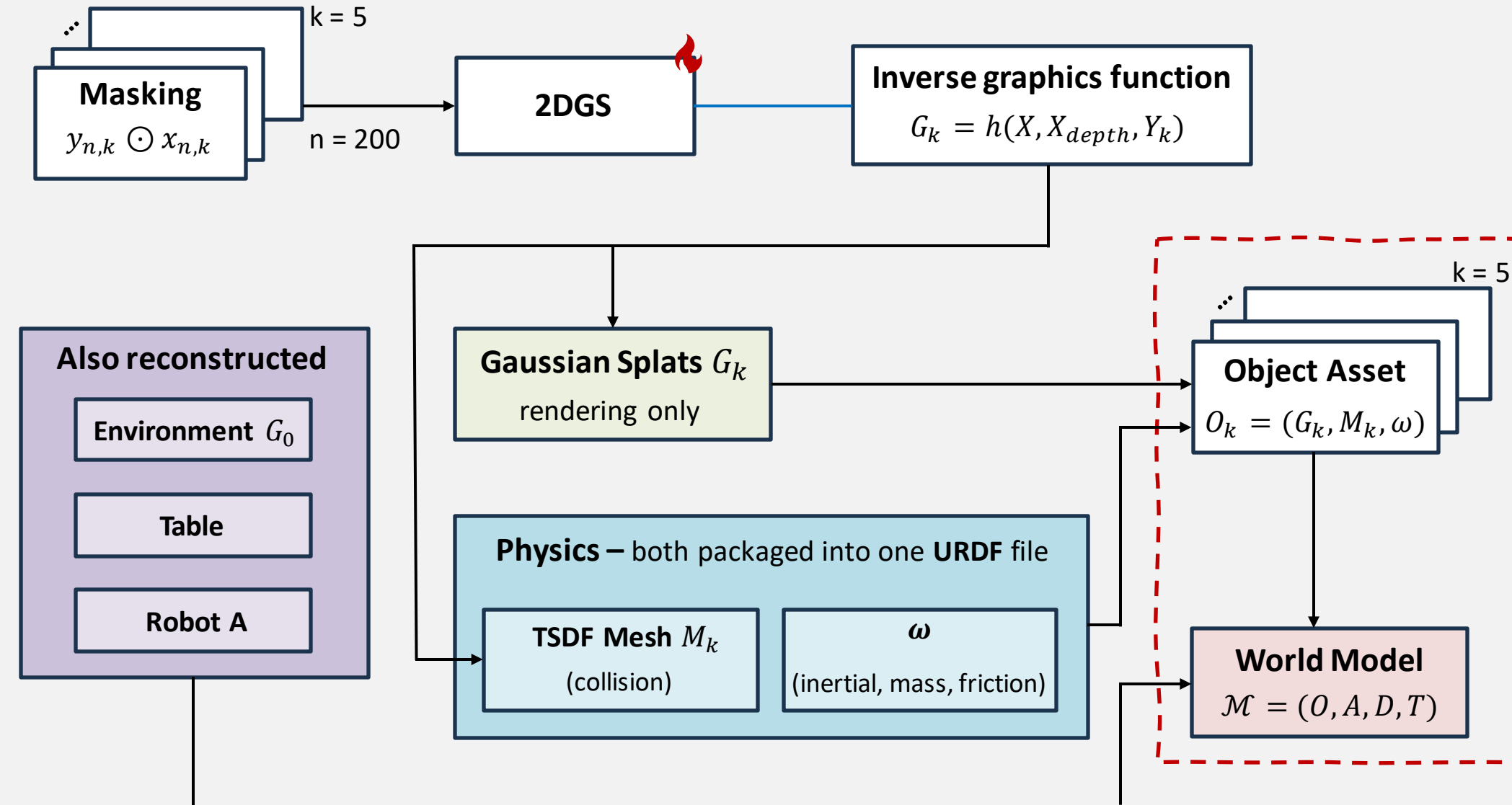


Building the world model

Input



Object-centric Reconstruction → Object Assets



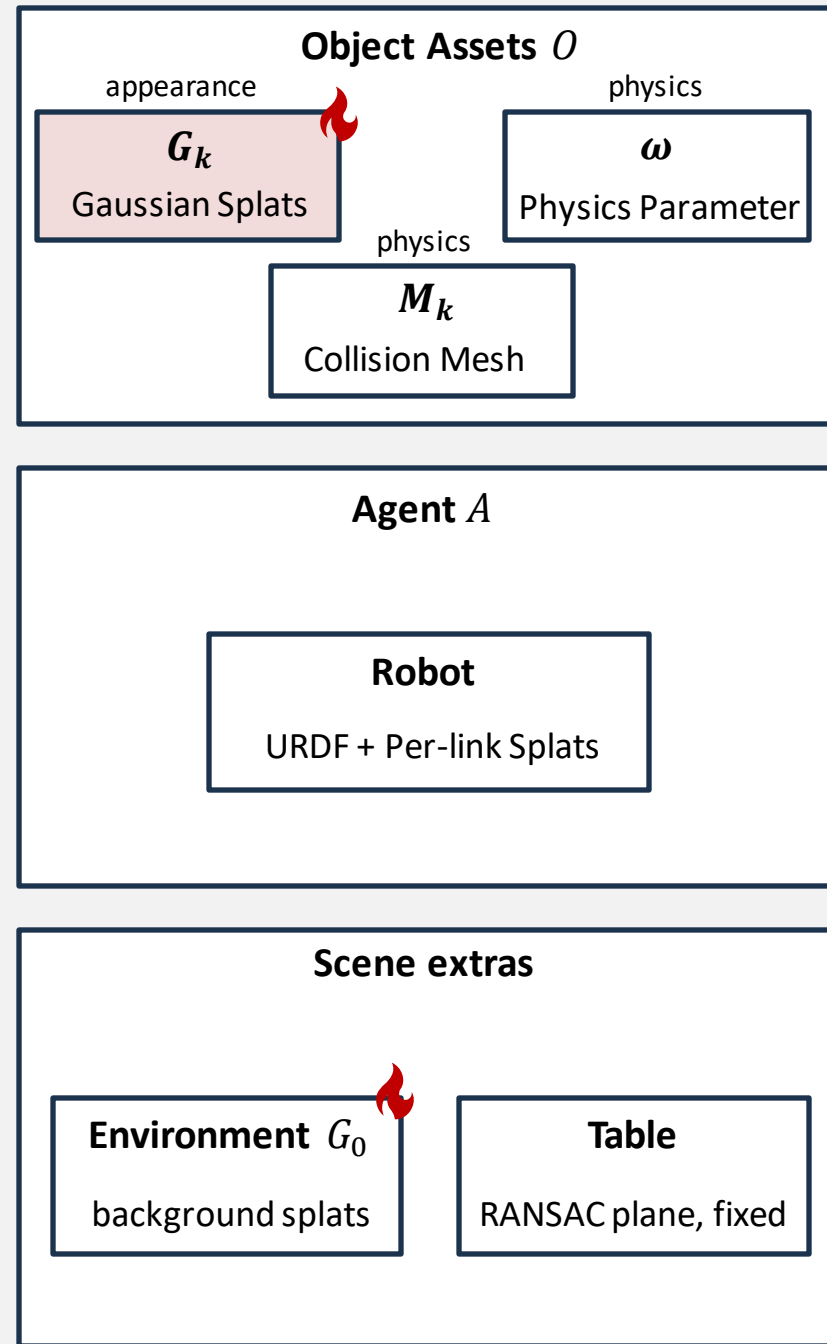
Method



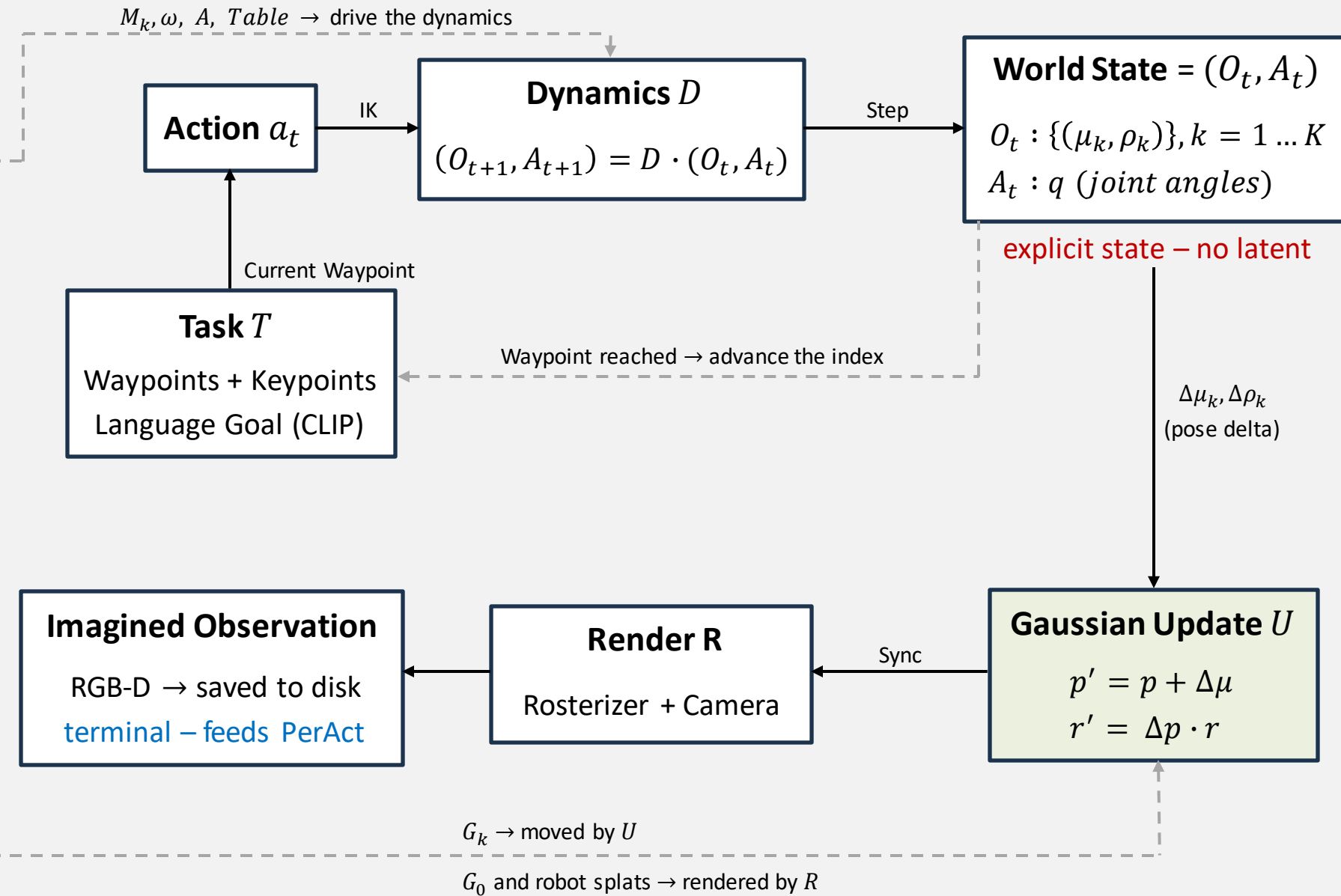
World Model : Components and Interaction

$\mathcal{M} = (O, A, D, T)$ | O : object assets, A : agent, D : dynamics operator, T : task

Static Assets (built once)

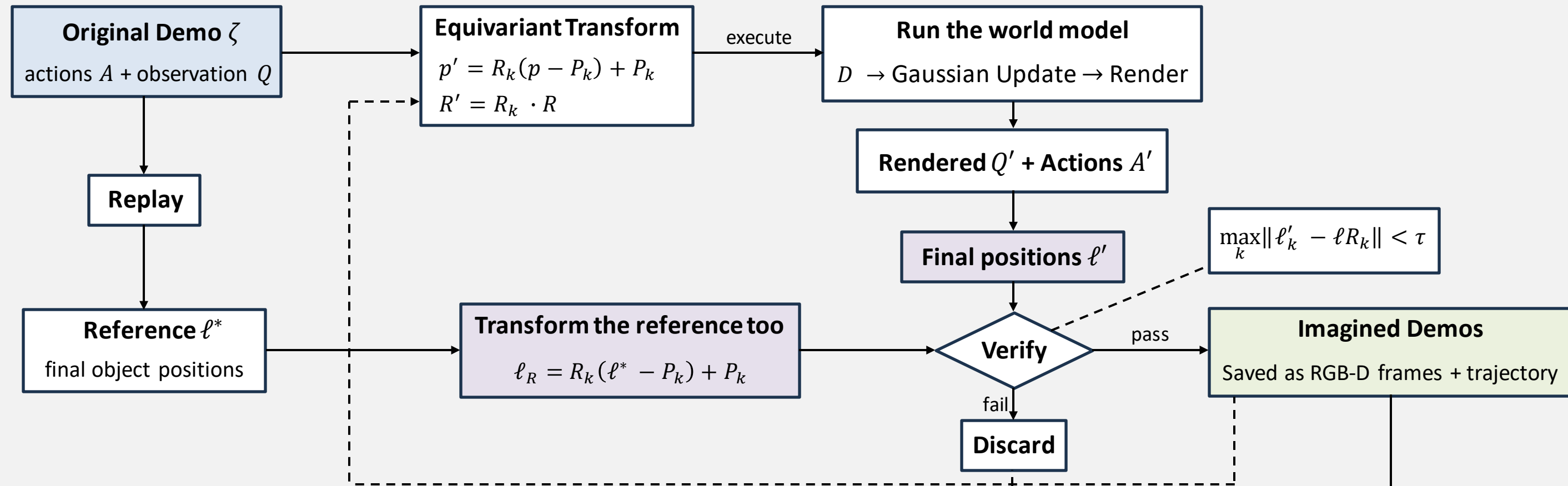


Runtime Cycle (one imagined rollout)

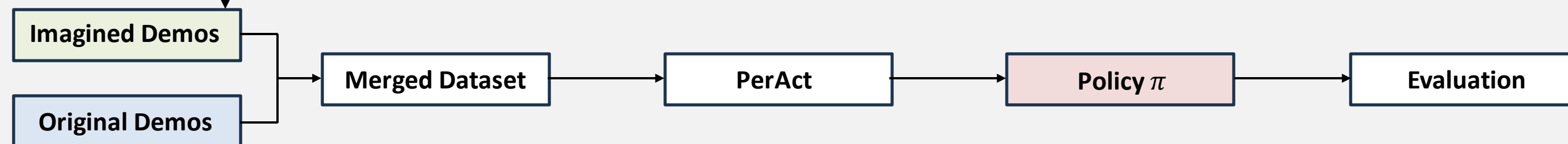


Imagination → Policy learning

Equivariant Imagination and Verification



Imitation Learning



03

PIN-WM: Learning Physics-INformed World Models for Non-Prehensile Manipulation

Wenxuan Li, Hang Zhao, Zhiyuan Yu, Yu Du, Qin Zou, Ruizhen Hu, Kai Xu

National University of Defense Technology, Wuhan University, Shenzhen University,
Guangdong Laboratory of Artificial Intelligence and Digital Economy

RSS 2025

arXiv Apr. 2025

<https://arxiv.org/abs/2504.16693>

Source: Google Scholar (Aug. 2026) / Citations: 29

- **Data Requirement** — Data-driven world models need many trajectories and break outside their training domain
 - Dreamer V2 gets only **1% push success** from 100 trajectories — prediction collapses **out of distribution**.
 - **Solution**: World model = differentiable **physics** ◦ differentiable rendering. Only 10 physical scalars are unknown → one task-agnostic push gives 97%.
- **State-Estimation Dependency** — A separate pose estimator leaks its error into the physics
 - Prior identification must **estimate object state first**, and that error **propagates** straight into the parameters.
 - **Solution**: **2DGS** renders the simulated state into an image, so a **pixel loss** is optimized directly — no estimator.
- **Residual Sim2Real Gap** — The twin is inexact, and unconstrained randomization dulls the policy
 - Partial, inaccurate observations **leave a gap**, yet wide **domain randomization** drops success to 33%.
 - **Solution**: **Digital cousins** — perturb only **physics** and **lighting** $\pm 10\%$ around the identified values → 97% in sim, 75% / 65% on the real robot, no fine-tuning.

Method

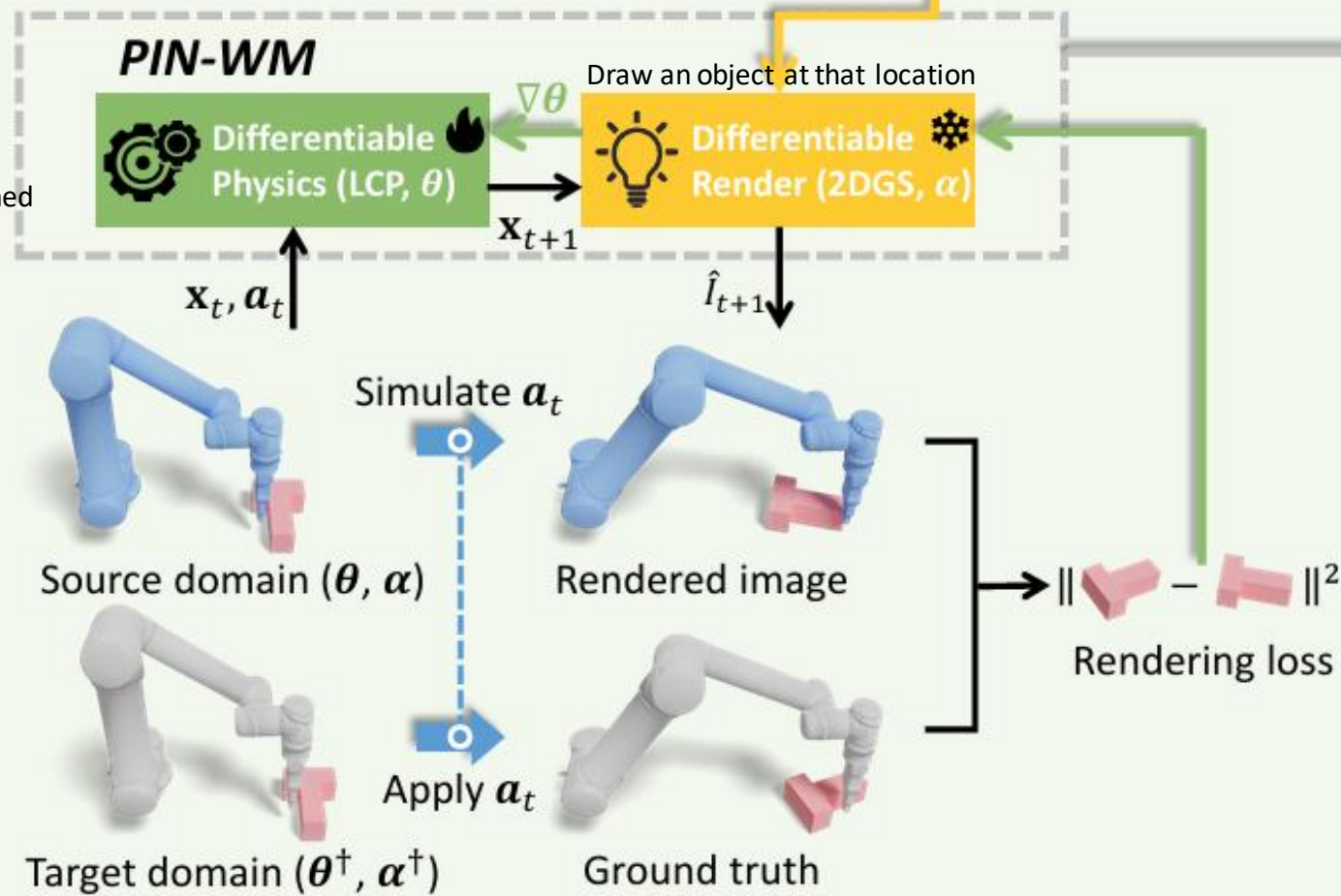
A. Real2Sim System Identification

(a) Rendering alignment



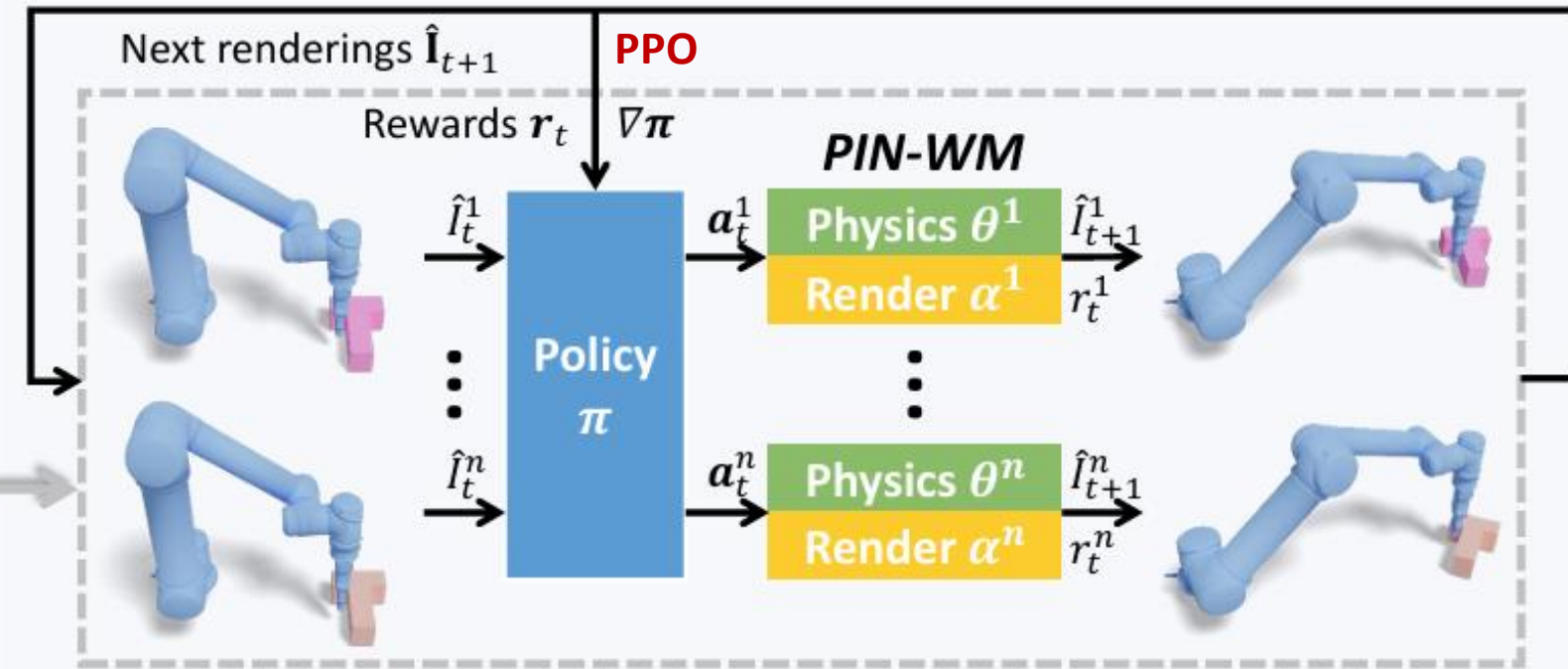
(b) Identification of physics parameters

Calculate where an object will go if pushed

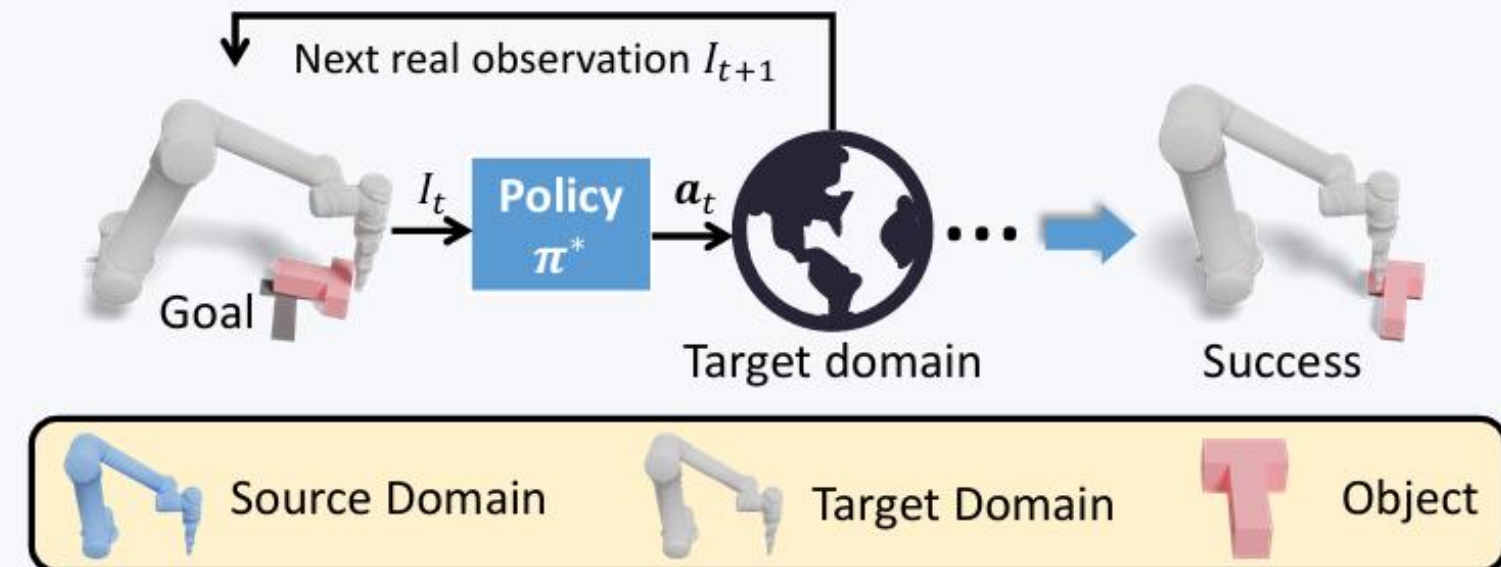


B. Sim2Real Policy Transfer

(c) Policy learning with physics-aware digital cousins



(d) Transfer to the target domain without fine-tuning



04

GWM: Towards Scalable Gaussian World Models for Robotic Manipulation

Guanxing Lu, Baoxiong Jia, Puhao Li, Yixin Chen, Ziwei Wang, Yansong Tang, Siyuan Huang,

Tsinghua University, State Key Laboratory of General Artificial Intelligence, BIGAI,
School of Electrical and Electronic Engineering, Nanyang Technological University

ICCV 2025

arXiv Aug. 2025

<https://arxiv.org/abs/2508.17600>

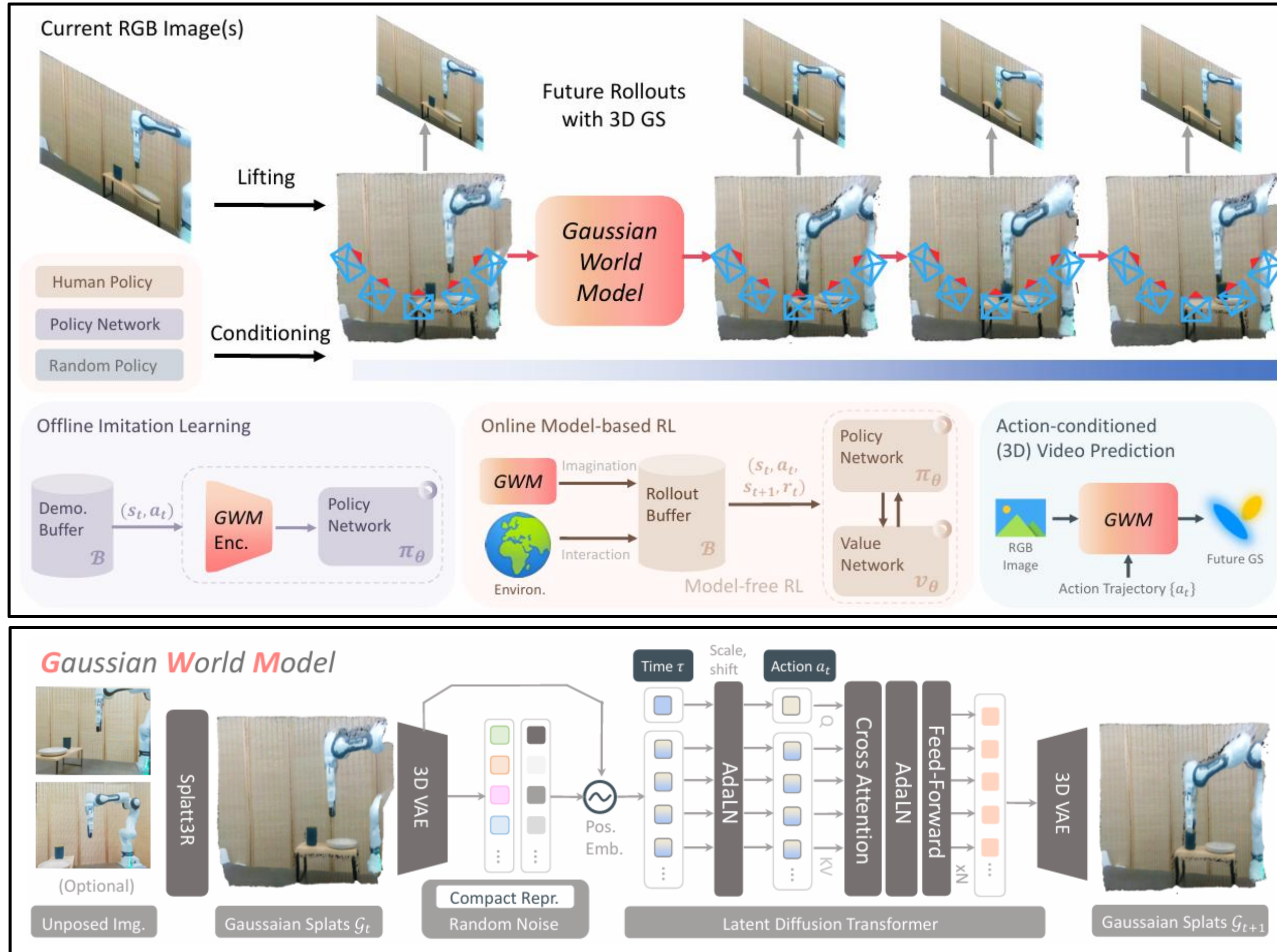
Source: Google Scholar (Aug. 2026) / Citations: 54

Problem



- **Spatial Grounding** — Image-based world models don't understand 3D
 - Pixel-space world models (iVideoGPT, UniSim) have **no geometric structure**, so they collapse under lighting, camera-pose, and texture **shifts** and cannot **localize** objects precisely.
 - **Solution**: Represent the world state as explicit **3D Gaussians** instead of pixels — the model reasons directly in metric 3D space.
- **Reconstruction Cost** — 3D representations are far too slow for a world model
 - NeRF and vanilla 3D-GS need **per-scene optimization** (~30k iterations), which cannot support model-based RL that imagines tens of thousands of rollouts.
 - **Solution**: Obtain Gaussians from a **feed-forward reconstruction model (Splatt3R)** in a single pass, then compress them with a **3D Gaussian VAE** into a fixed-length latent.

Method



05

EnerVerse: Envisioning Embodied Future Space for Robotics Manipulation

Siyuan Huang, Liliang Chen, Pengfei Zhou, Shengcong Chen, Yue Liao, Zhengkai Jiang, Yue Hu, Peng Gao, Hongsheng Li, Maoqing Yao, Guanghui Ren

SJTU, AgiBot, Shanghai AI Lab, CUHK MMLab, LV-NUS Lab

NeurIPS 2025
arXiv Jan. 2025

<https://arxiv.org/abs/2501.01895>

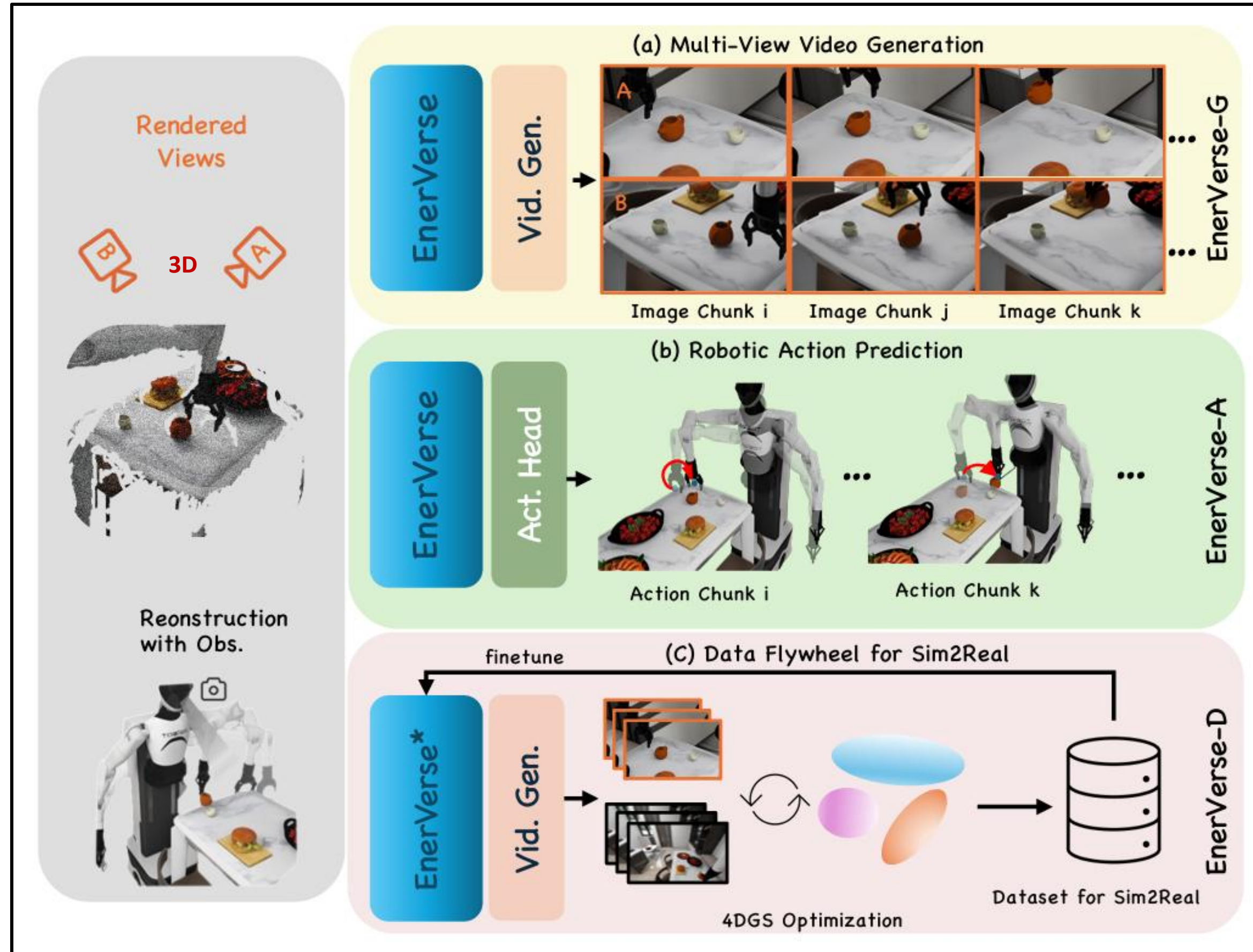
Source: Google Scholar (Aug. 2026) / Citations: 70

Problem

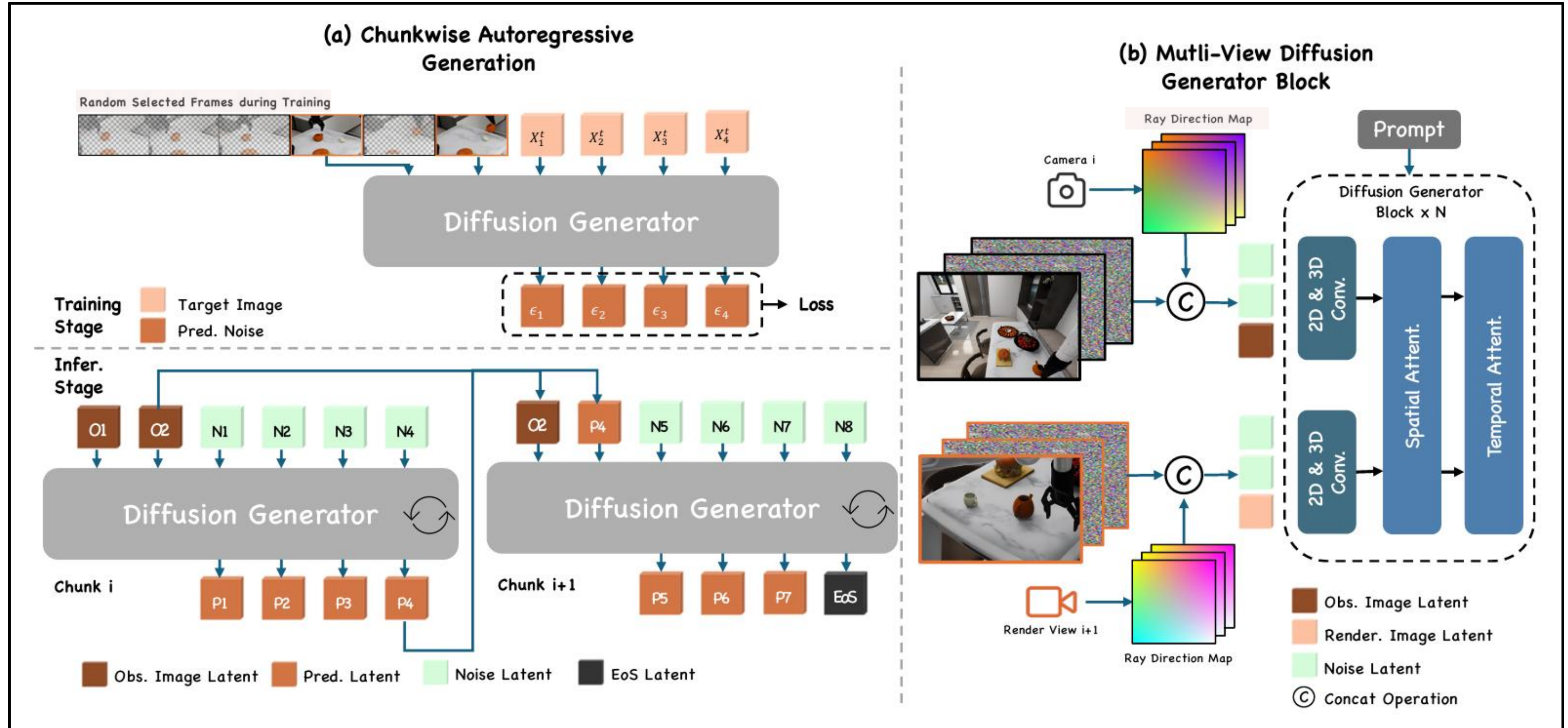


- **Long-Horizon Collapse** — Autoregressive video generation drifts and breaks down
 - Feeding **consecutive frames** as memory is redundant, so **errors accumulate** and the model collapses on long tasks.
 - **Solution: Sparse Memory** — keep only randomly sampled past frames, drop about 80%. LIBERO-Long 30.8 → 73.0.
- **3D Grounding from One Camera** — Single-view video cannot resolve depth and occlusion
 - Mounting **more cameras** raises hardware cost, I/O bandwidth, and system complexity.
 - **Solution:** Pre-train on **multi-view video** with **ray direction maps**; at deployment, depth-warp one RGB-D camera into rendered views. 79.0 → 92.1.
- **Multi-View Data Scarcity** — Calibrated multi-camera robot data barely exists
 - Public datasets are **single-view only**, and simulator data suffers the Sim2Real gap.
 - **Solution: EnerVerse-D** — a data flywheel where the **generator** imagines missing views and **4D Gaussian Splatting** enforces physical consistency. Hallucination -40%.

Method



Method



Thank you

Question & Discussion



RILS LAB