

EvoWorld: A World-Model-Centric Framework for Continuous Self-Evolution of Modular Embodied Skills

Boshi Zhang^{*1}, Sen Cui^{*,†1}, BaoHuaYin³, Youyi Kou⁴, Junyu Wu⁵, Zuo Pu¹,
Tao Xue¹, Zhikang Chen², Shanshan Wei¹, Min Zhang⁶, Miao Liu¹,
Changshui Zhang^{†1}, and Zhang Tao^{†1}

¹Tsinghua University, Beijing, China

²University of Oxford, Oxford, UK

³University of Sussex, Brighton, UK

⁴Beijing Jiaotong University, Beijing, China

⁵Xiamen University, Xiamen, China

⁶East China Normal University, Shanghai, China

zbs22@mails.tsinghua.edu.cn; cuis@mail.tsinghua.edu.cn

zcs@mail.tsinghua.edu.cn; taozhang@tsinghua.edu.cn

^{*}Equal contribution. [†]Project leader. [‡]Corresponding authors.

Abstract. Current progress in physical intelligence largely relies on scaling monolithic Vision-Language-Action (VLA) models, yet real-world policy data remain fragmented across scenes and tasks. This mismatch limits transfer, exacerbates catastrophic forgetting, and impedes continual improvement. A modular design that shares dynamics while specializing skills is therefore a promising paradigm.

We introduce EVO_WORLD (EVO_W), a world-model-centric framework for skill orchestration and iterative self-evolution. In EVO_W, VLAs form an expandable library of pluggable experts. A high-level router selects experts conditioned on scene and task, while an action-conditioned video world model provides a shared dynamics prior for rollout-based planning. The world model provides counterfactual rollouts to score candidate experts, while selected experts execute in the grounded scene to generate trajectories for verification. A vision-language evaluator delivers semantic scoring and diagnostic tags, enabling targeted updates to the world model, router memory, or specific experts rather than global retraining. This closes an automated loop that jointly improves grounding, routing, and skill refinement without manual task engineering. Experiments show that EVO_W enables automated task-to-policy synthesis with competitive success rates and consistent iterative gains in the evaluated settings, while producing valid and diverse trajectories that support evaluation and skill refinement.

Keywords: Vision-Language-Action Models · World Models · Expert Routing · Robot Manipulation · Continual Skill Refinement

1 Introduction

Physical intelligence is increasingly required to follow open-ended user intent rather than a fixed list of benchmark tasks [1, 4, 10]. To meet this demand,

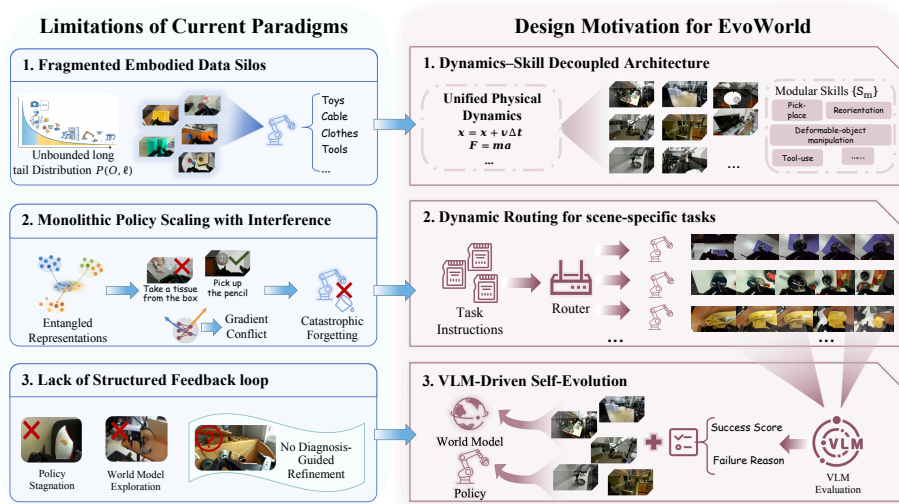


Fig. 1: Motivation and system overview of EvoW. **Left:** current embodied-learning pipelines face three structural bottlenecks: (1) fragmented embodied data silos with long-tail task/scene distributions, (2) interference and gradient conflict in monolithic policy scaling that induce catastrophic forgetting, and (3) the lack of a diagnosis-guided feedback loop for iterative refinement. **Right:** the proposed framework resolves these bottlenecks via three coupled components: (1) a dynamics-skill decoupled architecture that separates reusable physical dynamics from modular task-specialized skills, (2) scene-aware dynamic routing that maps task instructions and grounded observations to the most suitable expert policy, and (3) a VLM-driven self-evolution loop that scores outcomes, identifies failure causes, and feeds diagnostic signals back to world-model adaptation and skill updates.

the field has largely trended toward scaling monolithic Vision-Language-Action (VLA) models and diffusion-style policies to capture a broad spectrum of robotic behaviors [5, 7, 25]. Large-scale datasets such as DROID and Open X-Embodiment have further fueled this progress by providing broad task and visual coverage [8, 24]. However, monolithic scaling can obscure the mismatch between reusable physical dynamics and task-specific execution strategies, leaving pervasive real-world data silos under-addressed.

Current generalist control approaches rely predominantly on training large VLA models or flow-based policies on heterogeneous, massive data sources to achieve broad manipulation coverage [3–5, 7, 20, 25, 28, 45]. By treating diverse robotic experiences as tokens in a shared transformer-style backbone, these methods aim to absorb both physical dynamics and task-execution patterns through scale [3–5, 20, 25]. This paradigm assumes that a single, sufficiently large network can generalize across embodiments and environments via a shared latent space. Monolithic methods can face practical limitations under fragmented scene distributions and long-horizon compounding errors [2, 26]. In practice, robotic data remain distributed across labs, embodiments, and collection protocols [8, 11, 24], making continual updates to a single policy costly and vulnerable to forget-

ting when new task families are introduced. Moreover, without explicit feasibility verification during rollout, these systems can drift semantically and become physically inconsistent in open-world settings [2, 26, 30, 44].

We identify the root of these challenges in the lack of fundamental decoupling between universal physical dynamics and specialized, task-specific execution strategies [15–18, 37]. While the underlying laws of physics are largely invariant across scenarios, the "skills" required for manipulation are diverse and highly context-dependent [3, 20, 25, 28]. Current paradigms lack an explicit mapping mechanism to route task semantics to the most suitable algorithmic specialist, forcing a single network to resolve competing gradients from disparate domains [40, 41]. We argue that physical intelligence should treat the world model as a shared dynamics module while maintaining policies as modular, pluggable skills that can be updated independently [18, 37].

We propose **EvoWorld (EvoW)**, a world-model-centric framework for task-to-policy automation that enables continuous skill evolution [18, 37]. Unlike monolithic approaches [3, 20, 25, 28], **EvoW** separates shared physical grounding from task-level execution. It treats VLAs as modular components in an expandable library. The framework operates through a principled pipeline: (i) anchoring the task in a high-fidelity physical scene via semantic retrieval, (ii) performing counterfactual simulation using an action-conditioned video world model [30, 44], (iii) routing the task to the optimal VLA expert through a scene-aware scheduler, and (iv) driving a closed-loop self-evolution cycle via Vision-Language Model (VLM) diagnosis. This design allows the system to refine the world model and retrain specific skills without re-initializing the full system.

Experiments demonstrate that EvoW significantly outperforms monolithic baselines [3, 20, 25, 28] in success rates, particularly when operating across diverse and fragmented data silos. By leveraging the world-model-driven closed loop, the system achieves continuous compounding improvements across evolutionary iterations, validating the effectiveness of VLM-based failure diagnosis and expert routing. The results confirm that decoupling dynamics from skills not only mitigates catastrophic forgetting but also enables a more efficient, automated path toward open-ended physical intelligence. We summarize our key contributions as follows:

- **Paradigm:** We propose a principled world-model-centric formulation that disentangles universal physical dynamics from task-specialized skills, shifting the focus from policy scaling to world-model-driven orchestration.
- **Method:** We implement the EvoW framework, which operationalizes this decoupling through retrieval-grounded scene construction and a VLM-driven diagnostic loop that enables autonomous, closed-loop self-evolution.
- **Validation:** Extensive experiments across data silos verify the necessity of expert routing and self-evolution, demonstrating that EvoW achieves compounding performance gains without the overhead of full-system retraining.

2 Related Works

Scaling policies and data. Generalist VLA policies such as RT-1, RT-2, and OpenVLA demonstrate strong instruction-following behavior through scale [4, 5, 25]. Industrial foundation efforts (e.g., GR00T N1) further push this direction with larger model and data pipelines [28]. In parallel, large real-robot datasets, including Open X-Embodiment, DROID, and RH20T, broaden scene and task coverage [8, 11, 24]. However, most of these pipelines remain monolithic, where updates are largely global and specialist selection at deployment time is limited. Under fragmented data silos, this often leads to interference and higher continual-update cost.

Manipulation policies and structured primitives. Beyond model/data scaling, manipulation research focuses on control structure, including end-to-end visuomotor learning, reinforcement learning, and imitation learning [23, 27, 35]. Structured formulations such as Transporter Networks and CLIPort improve compositionality and sample efficiency in multi-step tasks by introducing stronger inductive bias for spatial reasoning and action decomposition [34, 42]. Recent dexterity-oriented works further extend policy capacity under richer embodiment and contact settings [43, 45]. However, these methods are typically optimized within fixed policy-training pipelines and still lack a diagnosis-guided, module-specific refinement loop that can coordinate grounding, routing, and world-model updates.

Grounding and memory for embodied reuse. Grounding and memory-based approaches improve generalization by reusing prior trajectories, grounded plans, and contextual cues [33]. Their core strength is efficient reuse under recurring task patterns. Yet, in many settings, grounding is not tightly coupled with verifiable rollout generation and failure attribution. As a result, these approaches are less suited to autonomous task-to-policy iteration when scenes and task semantics shift together in open-ended real-robot settings.

World models for long-horizon control. World models support imagined rollouts for planning and policy improvement [15–17]. Recent latent-action and video-prediction models improve rollout horizon and perceptual realism [12, 13, 39]. Other studies investigate tighter coupling between world models and control policies [6, 21, 22]. A core challenge remains long-horizon robustness: compounding prediction errors and semantic drift still accumulate over time [2, 26].

Automated task generation in simulation and our position. RoboGen shows that closed-loop task generation can scale quickly in simulation [38]. For real-robot deployment, additional constraints are necessary: real-scene grounding, expert routing at deployment, and diagnosis-guided correction. EvoW addresses this gap by combining grounded initialization, task/scene-aware routing, and targeted diagnosis-guided updates in a world-model-centric loop, with explicit decoupling between shared dynamics and task-specialized skills.

Taken together, existing lines of work provide strong components but not an integrated mechanism that jointly addresses grounded scene initialization, deployment-time expert selection, and diagnosis-driven evolution under open-ended real-robot distributions. This motivates our design of EvoW as a closed-

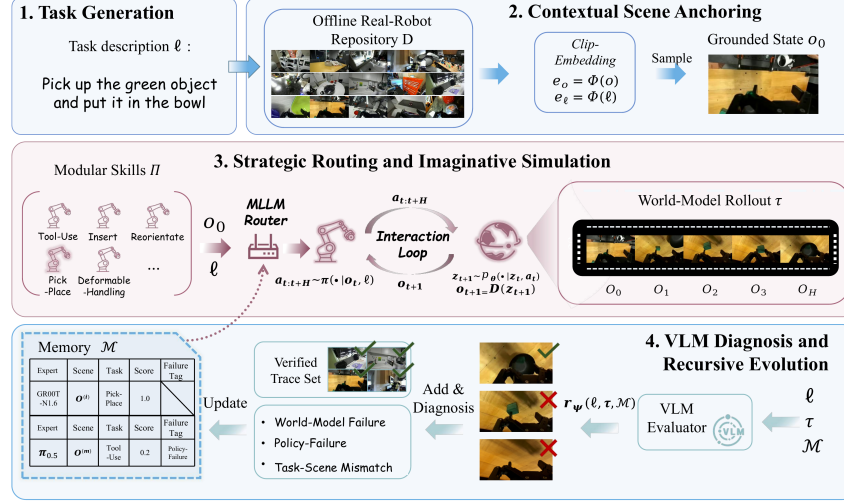


Fig. 2: EvoWorld pipeline overview. (1) **Task Generation:** construct candidate task specifications for each round. (2) **Contextual Scene Anchoring:** select a semantically aligned real scene from offline experience. (3) **Strategic Routing and Imaginative Simulation:** score candidates via world-model imagination, route to a suitable VLA expert, and execute the selected expert. (4) **VLM Diagnosis and Recursive Evolution:** use success score and failure type to trigger targeted module updates rather than global retraining.

loop system that links these modules through a unified world-model-centric objective, so that policy improvement is both targeted and scalable rather than globally retrained at every iteration.

3 Method

This section describes EvoW through three components: the problem formulation, the modular architecture, and the diagnosis-driven closed-loop optimization. The design couples Contextual Scene Anchoring, Strategic Routing, and Action-conditioned World-model Rollouts to support continuous task-to-policy evolution.

3.1 Notations and Problem Setup

We formalize open-ended task-to-policy synthesis as a language-conditioned decision process. Let ℓ denote a natural language task instruction (e.g., Pick up the green object and put it in the bowl). At each timestep t , o_t represents the visual observation and a_t the executed action. An interaction episode of horizon H is defined as a trajectory $\tau = (o_0, a_0, \dots, o_H)$, consisting of the initial grounded state and subsequent state-action pairs. To address the limitations of

monolithic scaling in fragmented environments, EvoW decomposes this process into universal dynamics and specialized control:

A Universal Video World Model (p_θ): This module captures transferable environmental dynamics across diverse scenarios. Given the action-conditioned latent transition, it is formulated as:

$$z_{t+1} \sim p_\theta(z_{t+1} \mid z_t, a_t), \quad o_{t+1} = D(z_{t+1}), \quad (1)$$

where E encodes observations into a latent space $z_t = E(o_t)$, and D decodes predicted future states back into the observation manifold.

A Modular Skill Library (Π): This library aggregates task-specific execution strategies optimized for distinct task families. It is defined as a collection of M specialized VLA experts:

$$\Pi = \{\pi_m\}_{m=1}^M, \quad a_{t:t+H} \sim \pi_m(\cdot \mid o_t, \ell), \quad (2)$$

where a high-level Strategic Router q scores experts conditioned on the instruction, grounded scene, and experience memory $\mathcal{M}$:

$$m^* = \arg \max_m q(m \mid \ell, o_0, \mathcal{M}). \quad (3)$$

This world-model-centric decomposition allows physical dynamics to be shared as a universal prior, while policy capacity is allocated through modular specialization and MLLM-driven routing. As illustrated in Figure 2, this architectural decoupling enables EvoW to bypass the constraints of real-world data silos and achieve continuous self-improvement through an introspective diagnostic loop.

3.2 Overview

The goal of EvoW is to establish a world-model-centric framework that enables autonomous skill acquisition and recursive self-improvement across fragmented data silos. Unlike monolithic VLA paradigms that entangle physical dynamics with task policies, EvoW uses a decoupled architecture where a universal Video World Model captures shared dynamics and Modular Skills act as pluggable experts. As illustrated in Figure 2, the EvoW pipeline operates through four coordinated phases:

Phase 1: Task Generation. The cycle begins with the formulation of open-vocabulary task specifications. Rather than relying on a static benchmark, EvoW adaptively constructs candidate instructions to explore the capabilities of the current skill library, ensuring that the evolution process remains open-ended and diverse.

Phase 2: Contextual Scene Anchoring. To bridge the gap between abstract instructions and physical reality, we perform Contextual Scene Anchoring (Grounding). Given a task, we leverage an offline experience manifold $\mathcal{D}$ (e.g., the DROID dataset) to extract a semantically aligned initial state o_0 . We employ a contrastive embedding space to compute the similarity between the instruction and visual candidates, sampling an initial grounded state.

Phase 3: Strategic Routing and Imaginative Simulation. Upon anchoring the scene, a high-level Skill Router queries the internal memory to select the most suitable VLA expert π_m from the skill library. The selected expert interacts with the action-conditioned world model. The world model performs rollouts $\tau = \{o_0, o_1, \dots, o_H\}$ by predicting latent dynamics $z_{t+1} \sim p_\theta(\cdot \mid z_t, a_t)$ and decoding them into observation space. This allows the system to evaluate the physical consequences of the expert’s actions before real-world execution.

Phase 4: VLM Diagnosis and Recursive Evolution. The final phase closes the loop with diagnosis. A Vision-Language Model (VLM) evaluates the simulated rollout and performs fine-grained credit assignment, categorizing failures into World-Model Failure, Policy Failure, or Task-Scene Mismatch. These diagnostic tags trigger targeted updates in the following loop.

3.3 Contextual Scene Anchoring

To bridge the gap between abstract linguistic intent and executable physical behavior, EvoW adopts a Contextual Scene Anchoring paradigm. The core objective is to initialize a semantically aligned real-world observation, thereby constraining the subsequent simulation to a physically plausible manifold rather than generating outcomes in a vacuum.

We define an offline experience manifold $\mathcal{D} = \{(o_i, \xi_i)\}_{i=1}^N$, where each entry represents a high-fidelity snapshot of real-robot interaction consisting of an observation o_i and its associated auxiliary context ξ_i . Given an open-vocabulary task specification ℓ , we first project the language instruction and the visual candidates into a shared latent embedding space:

$$e_\ell = \Phi_L(\ell), \quad e_{o_i} = \Phi_O(o_i), \quad (4)$$

where Φ_L and Φ_O denote pre-trained dual encoders optimized for cross-modal alignment. To ensure diverse yet relevant initialization, we ground ℓ by sampling an initial state o_0 from the distribution $p(i \mid \ell)$, formulated as a semantic-similarity-based Softmax:

$$o_0 \sim p(i \mid \ell) = \frac{\exp(\langle e_\ell, e_{o_i} \rangle / \tau_{\text{gnd}})}{\sum_{j=1}^N \exp(\langle e_\ell, e_{o_j} \rangle / \tau_{\text{gnd}})}, \quad (5)$$

where τ_{gnd} is a temperature parameter controlling the exploration-exploitation trade-off during grounding.

Once the initial observation o_0 is anchored, an Action-Conditioned Video World Model performs counterfactual simulation. The process begins by encoding o_0 into a compressed latent space $z_0 = E(o_0)$. The system then generates a temporal rollout $\tau = \{o_0, o_1, \dots, o_H\}$ through recursive latent transitions: $z_{t+1} \sim p_\theta(z_{t+1} \mid z_t, a_t)$ and $o_{t+1} = D(z_{t+1})$, where p_θ models the stochastic transition dynamics conditioned on the predicted action a_t from the selected VLA expert, and D decodes the latent states back into pixel space.

This hierarchical factorization explicitly disentangles Scene Grounding from Dynamics Prediction. By anchoring the simulation in a real-world experience

manifold $\mathcal{D}$, EvoW mitigates the compounding errors and semantic drift that can affect long-horizon generative models. This helps the world model provide a stable and physically consistent substrate for policy evaluation and autonomous skill evolution.

3.4 Strategic Routing and Imaginative Simulation

Given the fragmented nature of real-world robotic data, a single monolithic policy often suffers from catastrophic forgetting and performance degradation across diverse embodiments. To circumvent this, EvoW treats each specialized VLA as a modular expert π_m within a library $\{\pi_m\}_{m=1}^M$, where each expert is optimized for specific task families or physical constraints. The core challenge then shifts to Strategic Routing: mapping an open-vocabulary instruction ℓ and a grounded observation o_0 to the most compatible expert.

To facilitate this mapping, we instantiate the router with a Multi-modal Large Language Model (MLLM). The MLLM acts as a high-level reasoning module that uses semantic knowledge to evaluate compatibility between task requirements and VLA expert descriptions. To enable continuous self-improvement, the router is conditioned on a long-term Experience Memory $\mathcal{M}$, defined as:

$$\begin{aligned} m^* &= \arg \max_m q(m \mid \ell, o_0, \mathcal{M}), \\ \mathcal{M} &= \{(\ell^i, o_0^i, m^i, r_\psi^i, \kappa^i)\}_{i=1}^N. \end{aligned} \tag{6}$$

where each entry captures a historical interaction consisting of the task ℓ , the initial state o_0 , the selected expert index m , the VLM score r_ψ , and the diagnostic failure type κ .

The routing process operates on a dual-mode logic of Exploration and Exploitation:

Exploration: For unfamiliar task-scene pairs with low semantic density in $\mathcal{M}$, the router encourages diversity by sampling from a broader categorical distribution, allowing the system to discover new "scene-expert" correspondences.

Exploitation: For recurrent contexts, the MLLM synthesizes historical success/failure cues from $\mathcal{M}$ to generate memory-weighted logits, effectively biasing the selection toward experts with proven reliability in similar manifolds.

To further validate the routing decision, for each candidate expert π_m identified by the MLLM, the world model generates an imagined rollout τ_{π_m} . The router then computes an imagined-rollout score that reflects the predicted task consistency of the expert's trajectory. The final expert π_{m^*} is selected by maximizing the joint probability of semantic compatibility and simulated success, grounding execution in both high-level intent and low-level physical feasibility.

3.5 VLM Diagnosis and Recursive Evolution

Algorithm 1 EvoW closed-loop routing and self-evolution

Require: Task set $\mathcal{L}$, offline experience manifold $\mathcal{D}$, expert library $\Pi = \{\pi_m\}_{m=1}^M$, video world model p_θ , MLLM-based router q , VLM evaluator Ψ , experience memory $\mathcal{M}$, success threshold δ .

Ensure: Verified trajectory set $\hat{\mathcal{T}}$, updated skill library Π , and refined memory $\mathcal{M}$.

- 1: $\hat{\mathcal{T}} \leftarrow \emptyset$
- 2: **for** each round **do**
- 3: **(Stage 1: Task Generation)** Sample or construct task specification ℓ
- 4: **(Stage 2: Contextual Scene Anchoring)** Ground initial scene o_0 from $\mathcal{D}$ conditioned on ℓ
- 5: **(Stage 3: Strategic Routing & Imaginative Simulation)** Select $m^* \leftarrow \arg \max_m q(m \mid \ell, o_0, \mathcal{M})$ and expert π_{m^*} ; run world-model rollout for candidate scoring; execute selected expert to obtain trajectory τ
- 6: **(Stage 4: VLM Diagnosis & Recursive Evolution)** Evaluate $(r_\psi, \kappa) \leftarrow \Psi(\tau, \ell, \mathcal{M})$ and update $\mathcal{M} \leftarrow \mathcal{M} \cup \{(\ell, o_0, m^*, r_\psi, \kappa)\}$
- 7: **if** $r_\psi \geq \delta$ **then**
- 8: $\hat{\mathcal{T}} \leftarrow \hat{\mathcal{T}} \cup \{(\tau, \ell, o_0, m^*, r_\psi)\}$
- 9: **else**
- 10: **if** $\kappa = \text{TSM}$ **then**
- 11: Update grounding module
- 12: **else if** $\kappa = \text{PF}$ **then**
- 13: Update routing and/or selected expert
- 14: **else**
- 15: Update world model
- 16: **end if**
- 17: **end if**
- 18: **end for**
- 19: **return** $\hat{\mathcal{T}}, \Pi, \mathcal{M}$

To achieve continuous self-improvement while avoiding the cost and forgetting risk of global retraining, EvoW uses VLM-based diagnosis for fine-grained credit assignment. After the routed VLA expert π_m generates a trajectory τ , the evaluator Ψ compares the outcome with the task specification ℓ and the grounded initial state. Formally, the evaluation is defined as:

$$(r_\psi, \kappa) = \Psi(\tau, \ell, \mathcal{M}), \quad r_\psi \in [0, 1], \quad \kappa \in \{\text{TSM}, \text{PF}, \text{WMF}\}, \quad (7)$$

where r_ψ is a continuous success score and κ denotes the categorical failure attribution. Specifically, Task-Scene Mismatch (TSM) identifies misalignments in initial semantic grounding; Policy Failure (PF) indicates suboptimal control outputs from the VLA; and World-Model Failure (WMF) signals a divergence between predicted dynamics and real-world physical constraints. EvoW then applies targeted updates based on a predefined reliability threshold δ , so that only the identified module is refined while stable components remain fixed:

Verified Experience Accumulation ($r_\psi \geq \delta$): High-confidence trajectories are integrated into the verified trajectory set, serving as high-fidelity demonstrations for future self-imitation learning and router calibration.

Semantic Grounding Refinement ($\kappa = \text{TSM}$): If the failure stems from an irrelevant initial scene, the system updates the Router’s memory $\mathcal{M}$ to refine the semantic-to-scene grounding manifold, preventing the recurrence of analogous grounding errors.

Modular Policy Plasticity ($\kappa = \text{PF}$): For policy-induced failures, the system retrains the specific VLA expert π_{m^*} . The diagnostic feedback acts as a corrective signal for aligning the expert with successful task executions.

Dynamics Grounding Refinement ($\kappa = \text{WMF}$): When the simulation fails to capture valid physics, the world model’s dynamics parameters are updated using the discrepancy between the rollout τ and the physical ground truth, improving dynamics prediction.

This targeted evolutionary strategy allows EvoW to bypass the "one-size-fits-all" training bottleneck, achieving continuous compounding improvements in success rates through recursive, module-specific plasticity.

4 Experiments

In this section, we evaluate EvoW along four dimensions: **(1) skill-level performance**, including single-skill success rates (SR), iterative policy improvement, and compositional long-horizon tasks; **(2) task diversity and scene validity**, measuring whether generated tasks remain diverse while preserving grounded executability; **(3) generative quality**, evaluating the visual fidelity and temporal consistency of imagined trajectories under large-scale rollouts; and **(4) real-robot validation**, checking whether the gains transfer beyond imagined rollouts.

4.1 Experimental Setup

To isolate the contributions of our method, we keep VLA backbones, world-model data sources, and evaluation splits fixed across compared variants. We sample tasks from skill families and use grounded initial scenes anchored to a shared offline repository. Across all compared variants, we fix the rollout horizon, task split, scene split, and trajectory budget per task. We first compare base policies and routed variants under identical initialization, and then run iterative refinement with matched update budgets per round to analyze cumulative improvements under a controlled setting.

The benchmark skill space contains five representative manipulation families, denoted as $\mathcal{S}_1$ – $\mathcal{S}_5$: $\mathcal{S}_1$ (Pick-and-Place), $\mathcal{S}_2$ (Tool-Use), $\mathcal{S}_3$ (Reorientation), $\mathcal{S}_4$ (Deformable-Object Manipulation), and $\mathcal{S}_5$ (Articulation). Comprehensive implementation details—including pre-training configurations, VLM-evaluator fine-tuning procedures, episode statistics, and rollout curation methods—are deferred to Appendix, Experimental Details.

Table 1: Policy success rates across five representative skill families and four VLA backbones, comparing base policies with EvoW routing. Values are reported as success rates. **EvoW*** denotes our method, and Gain($\uparrow\%$) is computed relative to the corresponding base policy average.

Backbone	Setting	S_1	S_2	S_3	S_4	S_5	Avg	Gain($\uparrow\%$)
GR00T-N1.6	Base	0.38	0.24	0.18	0.28	0.20	0.256	–
	EvoW*	0.48	0.50	0.46	0.42	0.48	0.468	$\uparrow 82.8\%$
PI-0.5	Base	0.32	0.14	0.24	0.18	0.28	0.232	–
	EvoW*	0.48	0.28	0.38	0.44	0.48	0.412	$\uparrow 77.6\%$
PI-0	Base	0.40	0.22	0.26	0.32	0.22	0.284	–
	EvoW*	0.62	0.38	0.40	0.40	0.50	0.460	$\uparrow 62.0\%$
PI-0-fast	Base	0.38	0.18	0.26	0.24	0.26	0.264	–
	EvoW*	0.66	0.38	0.44	0.36	0.42	0.452	$\uparrow 71.2\%$

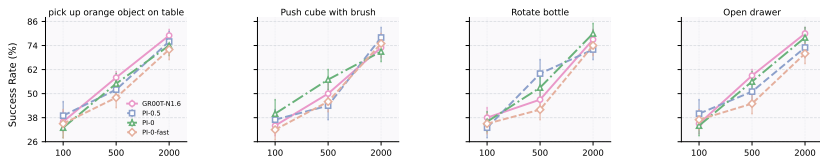


Fig. 3: Task-level refinement trends. Iterative updates consistently raise success rates, and variance bars indicate that gains are stable across runs.

4.2 Skill-Level Performance

We first investigate whether the routing mechanism independently enhances policy success under fixed model and data constraints, followed by an analysis of the cumulative benefits from iterative refinement. Table 1 details the static SR comparisons between the baseline and routed variants across four VLA backbones and the five skill families.

Ablation Results. As shown in Table 1, routing with **EvoW** consistently improves success rates across all four backbones. Averaged over backbones, SR increases from 0.259 to 0.448 (absolute +18.9 points, relative +73.0%). Per-backbone relative improvements range from +62.0% to +82.8%, indicating that the gain is stable rather than concentrated on a single backbone.

Policy Improvement Over Iterations. Figure 3 reports SR trends over refinement iterations on representative tasks. SR improves across tasks, indicating that diagnosis-conditioned updates in Stage 4 reduce recurrent failure patterns over time. The combined visualization in Figure 5 and qualitative rollouts in Figure 4 further show more stable multi-step execution in later iterations.

These results are consistent with the formulation that expert selection contributes immediate performance gains under fixed backbones, and diagnosis-conditioned refinement yields additional improvements over successive rounds by reducing recurrent execution failures.

Compositional Long-Horizon Tasks. Single-skill metrics cannot fully capture robustness under cross-skill transitions. We therefore evaluate compo-

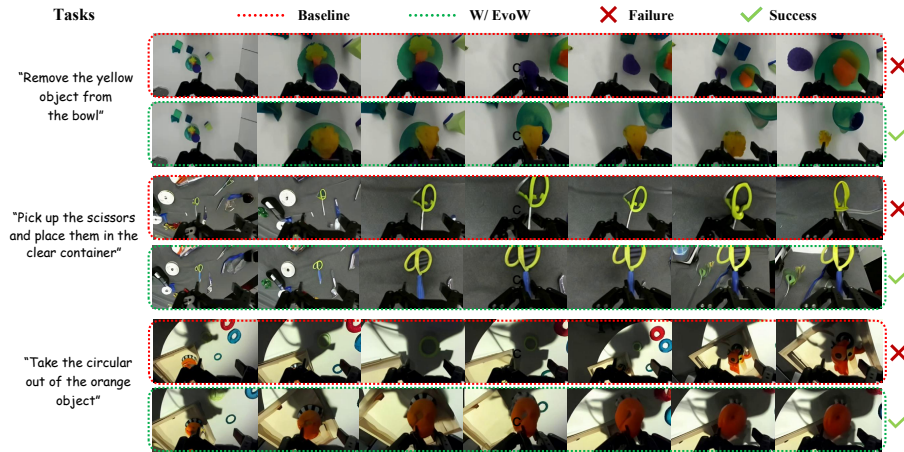


Fig. 4: Qualitative examples from our refinement loop. For each base policy shown in the top row, the bottom row presents the corresponding imagined rollouts in the grounded world model.

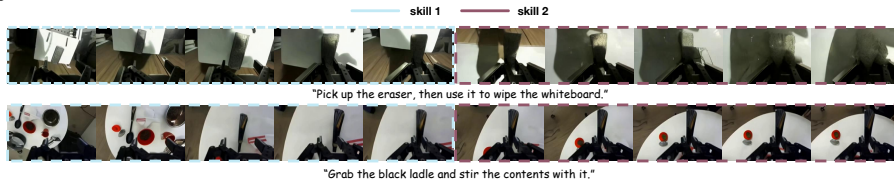


Fig. 5: Visualization of compositional long-horizon tasks. Early and late stages involve different skills, and success requires completing the full task sequence.

sitional long-horizon tasks that chain two or three skills from $\mathcal{S}_1$ – $\mathcal{S}_5$. **Results.** Figure 6 reports full-success and partial-success SR on 10 procedurally generated compositional tasks. Full-success SR requires completing all sub-goals with stable rollout–execution consistency, while partial-success SR requires completing at least one sub-goal without instability. Routed variants consistently improve the feasible SR range across backbones, especially under the full-success criterion, indicating better cross-stage coordination when tasks require skill switching. Task definitions and per-task statistics are provided in Appendix, Compositional Long-Horizon Tasks.

4.3 Task diversity and scene validity

Beyond execution success, it is critical that EvoW can scale task generation while strictly adhering to grounding constraints to preserve scene validity. Table 2 benchmarks the language and scene diversity of our generated tasks against existing simulation environments, including RoboGen [38]. We report complementary metrics widely used for text and visual diversity analysis: Self-BLEU [29], Sentence-BERT similarity [32], ViT similarity [9], and CLIP similarity [31].

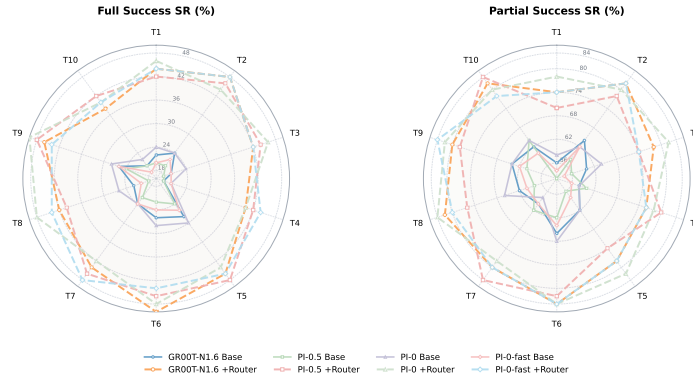


Fig. 6: Compositional long-horizon evaluation under full-completion and partial-completion criteria. Routed variants expand the performance envelope across backbones, indicating improved cross-stage robustness.

Table 2: Task and scene diversity comparison under executability constraints. The metrics indicate that EvoW preserves competitive diversity while maintaining scene validity for grounded deployment according to Self-BLEU and CLIP metrics.

Metric	EvoW	RoboGen	Behavior-100	RLBench	MetaWorld	ManiSkill2	GenSim
Number of tasks	100	106	100	106	50	20	70
Task Description – Self-BLEU ↓	0.281	0.284	0.299	0.317	0.322	0.674	0.378
Task Description – Embedding Similarity (SentenceBert) ↓	0.312	0.165	0.210	0.200	0.263	0.194	0.288
Scene Image – Embedding Similarity (ViT) ↓	0.337	0.193	0.389	0.375	0.517	0.332	0.717
Scene Image – Embedding Similarity (CLIP) ↓	0.754	0.762	0.833	0.864	0.867	0.828	0.932

As shown in Table 2, EvoW achieves the lowest Self-BLEU and CLIP similarity and remains competitive on ViT similarity. EvoW is weaker on Sentence-BERT similarity, which is expected because grounded executability constraints reduce free-form language variation. This trade-off is expected in our setting: the generated tasks are intended for physically realizable rollouts rather than unconstrained textual diversity.

4.4 Generative quality

We assess whether EvoW’s architectural improvements translate to higher-fidelity rollout generation in terms of both visual appearance and temporal consistency, compared with representative world-model baselines such as WorldGym (WPE), IRASim, and Ctrl-World [14, 30, 44]. We use FID [19] and FVD [36] as standard image- and video-level fidelity metrics.

Ablation Results. Figure 7 shows that EvoW achieves lower FID and FVD than the compared baselines. The larger relative gain on FVD indicates stronger temporal consistency over long-horizon rollouts. Under view-conditioned settings, EvoW remains better than Ctrl-World for both third-view and wrist-view comparisons, consistent with improved rollout stability.

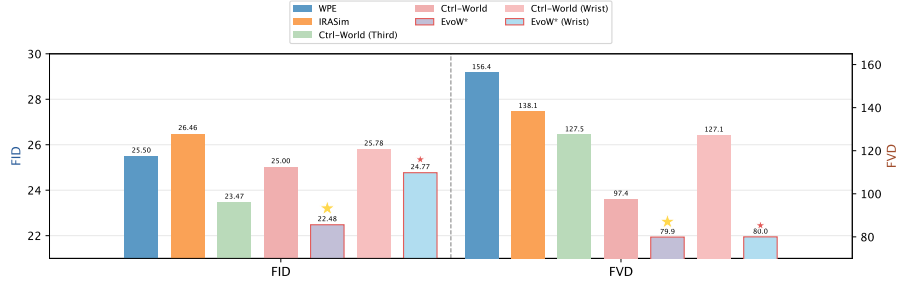


Fig. 7: Generative-quality comparison under a view-conditioned setting. EvoW achieves lower FID and FVD than the compared baselines in this setup. The larger reduction on FVD is consistent with improved temporal coherence, and view-matched expert outputs are associated with more stable rollout dynamics. Red and yellow stars mark EvoW in different views, with yellow indicating the best value under each metric.



Fig. 8: Qualitative temporal-consistency comparison for long-horizon rollouts. EvoW better preserves object identity, contact continuity, and motion progression across frames, consistent with stronger long-horizon world-model stability.

Qualitative evidence in Figure 8 is consistent with these quantitative trends: EvoW better preserves object identity and contact progression across time, while baseline rollouts more often show collapse behaviors such as object disappearance and broken motion continuity. Overall, EvoW improves both video-level fidelity and long-horizon temporal stability.

4.5 Real-world validation

To examine whether the observed gains transfer beyond imagined rollouts, we conduct a real-robot study on a Franka arm using the GR00T-N1.6 backbone. The evaluation covers five representative manipulation tasks with 20 hardware trials per task under matched task and scene initialization.

As shown in Figure 9, EvoW improves the average real-world SR from 29% to 70% over the monolithic Base, and all evaluated tasks show a positive gain. These hardware results are consistent with the simulated evaluations and indicate that scene-aware routing and diagnosis-driven updates support grounded deployment beyond world-model rollouts.

5 Conclusion

We present EvoW, a framework that maps open-ended task semantics to verified training trajectories and executable manipulation policies. EvoW grounds each

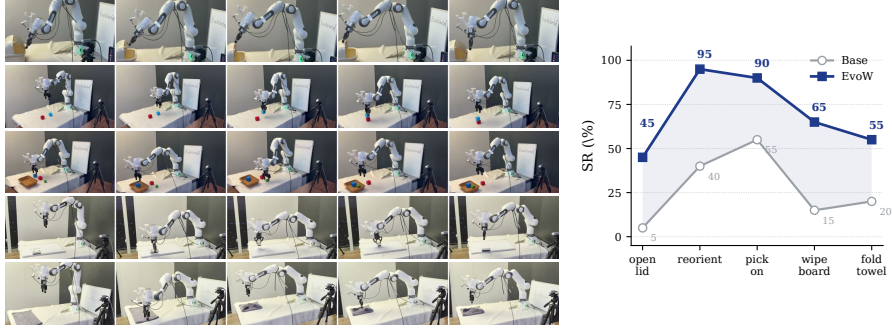


Fig. 9: Real-world validation on a Franka arm. Left: five-frame rollouts for open lid, pick on, pick two, wipe board, and fold towel. Right: SR over 20 trials per task; EvoW improves every task over Base.

task in task-consistent real-world scenes, synthesizes scene-consistent episodes with a decoupled world model, and closes the loop via VLM-based routing and verification. This design decouples shared dynamics from task-specialized skills and enables targeted updates through failure diagnosis.

Our results indicate that scene grounding and automated VLM feedback are complementary for scaling open-ended manipulation, yielding reliable task execution and continual improvement over iterations. Future work includes (i) expanding and curating the anchor repository to better match open-ended tasks, (ii) improving long-horizon world-model fidelity with uncertainty-aware rollout and filtering, and (iii) strengthening routing, verification, and failure attribution with task-aware critics and calibrated confidence signals.

Acknowledgements

This work was supported by the National Science and Technology Major Project under Grant 2022ZD0114903.

The authors also acknowledge support from the Tsinghua University Initiative Scientific Research Program (Student Academic Research Advancement Program) and the Beijing Natural Science Foundation Undergraduate “Qiyang Program” (Grant No. 26QY0429). Sen Cui additionally acknowledges support from the Science and Technology Project of Beijing Municipal Science & Technology Commission (Grant No. Z251100008125030) and the Shuimu Scholar Program of Tsinghua University.

References

1. Ahn, M., Brohan, A., Brown, N., Chebotar, Y., Cortes, O., David, B., Finn, C., Fu, C., Gopalakrishnan, K., Hausman, K., Herzog, A., Ho, D., Hsu, J., Ibarz, J., Ichter, B., Irpan, A., Jang, E., Ruano, R.J., Jeffrey, K., Jesmonth, S., Joshi, N.J., Julian, R., Kalashnikov, D., Kuang, Y., Lee, K.H., Levine, S., Lu, Y., Luu, L., Parada, C., Pastor, P., Quiambao, J., Rao, K., Rettinghouse, J., Reyes, D., Sermanet, P., Sievers, N., Tan, C., Toshev, A., Vanhoucke, V., Xia, F., Xiao, T., Xu, P., Xu, S., Yan, M., Zeng, A.: Do as i can, not as i say: Grounding language in robotic affordances. In: CORL (2022)
2. Asadi, K., Misra, D., Kim, S., Littman, M.L.: Combating the compounding-error problem with a multi-step model (2019), <https://arxiv.org/abs/1905.13320>
3. Black, K., Brown, N., Driess, D., Esmail, A., Equi, M., Finn, C., Fusai, N., Groom, L., Hausman, K., Ichter, B., Jakubczak, S., Jones, T., Ke, L., Levine, S., Li-Bell, A., Mothukuri, M., Nair, S., Pertsch, K., Shi, L.X., Tanner, J., Vuong, Q., Walling, A., Wang, H., Zhilinsky, U.: π_0 : A vision-language-action flow model for general robot control (2026), <https://arxiv.org/abs/2410.24164>
4. Brohan, A., Brown, N., Carbajal, J., Chebotar, Y., Chen, X., Choromanski, K., Ding, T., Driess, D., Dubey, A., Finn, C., Florence, P., Fu, C., Arenas, M.G., Gopalakrishnan, K., Han, K., Hausman, K., Herzog, A., Hsu, J., Ichter, B., Irpan, A., Joshi, N., Julian, R., Kalashnikov, D., Kuang, Y., Leal, I., Lee, L., Lee, T.W.E., Levine, S., Lu, Y., Michalewski, H., Mordatch, I., Pertsch, K., Rao, K., Reymann, K., Ryoo, M., Salazar, G., Sanketi, P., Sermanet, P., Singh, J., Singh, A., Soricut, R., Tran, H., Vanhoucke, V., Vuong, Q., Wahid, A., Welker, S., Wohlhart, P., Wu, J., Xia, F., Xiao, T., Xu, P., Xu, S., Yu, T., Zitkovich, B.: Rt-2: Vision-language-action models transfer web knowledge to robotic control. In: CORL (2023)
5. Brohan, A., Brown, N., Carbajal, J., Chebotar, Y., Dabis, J., Finn, C., Gopalakrishnan, K., Hausman, K., Herzog, A., Hsu, J., Ibarz, J., Ichter, B., Irpan, A., Jackson, T., Jesmonth, S., Joshi, N.J., Julian, R., Kalashnikov, D., Kuang, Y., Leal, I., Lee, K.H., Levine, S., Lu, Y., Malla, U., Manjunath, D., Mordatch, I., Nachum, O., Parada, C., Peralta, J., Perez, E., Pertsch, K., Quiambao, J., Rao, K., Ryoo, M., Salazar, G., Sanketi, P., Sayed, K., Singh, J., Sontakke, S., Stone, A., Tan, C., Tran, H., Vanhoucke, V., Vega, S., Vuong, Q., Xia, F., Xiao, T., Xu, P., Xu, S., Yu, T., Zitkovich, B.: Rt-1: Robotics transformer for real-world control at scale. In: RSS (2023)
6. Cen, J., Yu, C., Yuan, H., Jiang, Y., Huang, S., Guo, J., Li, X., Song, Y., Luo, H., Wang, F., Zhao, D., Chen, H.: Worldvla: Towards autoregressive action world model (2025), <https://arxiv.org/abs/2506.21539>
7. Chi, C., Xu, Z., Feng, S., Cousineau, E., Du, Y., Burchfiel, B., Tedrake, R., Song, S.: Diffusion policy: Visuomotor policy learning via action diffusion. In: RSS (2023)
8. Collaboration, E., O'Neill, A., Rehman, A., Gupta, A., Maddukuri, A., Gupta, A., Padalkar, A., Lee, A., Pooley, A., Gupta, A., Mandlekar, A., Jain, A., Tung, A., Bewley, A., Herzog, A., Irpan, A., Khazatsky, A., Rai, A., Gupta, A., Wang, A., Kolobov, A., Singh, A., Garg, A., Kembhavi, A., Xie, A., Brohan, A., Raffin, A., Sharma, A., Yavary, A., Jain, A., Balakrishna, A., Wahid, A., Burgess-Limerick, B., Kim, B., Schölkopf, B., Wulfe, B., Ichter, B., Lu, C., Xu, C., Le, C., Finn, C., Wang, C., Xu, C., Chi, C., Huang, C., Chan, C., Agia, C., Pan, C., Fu, C., Devin, C., Xu, D., Morton, D., Driess, D., Chen, D., Pathak, D., Shah, D., Büchler, D., Jayaraman, D., Kalashnikov, D., Sadigh, D., Johns, E., Foster, E., Liu, F., Ceola, F., Xia, F., Zhao, F., Frujeri, F.V., Stulp, F., Zhou, G., Sukhatme, G.S., Salhotra,

- G., Yan, G., Feng, G., Schiavi, G., Berseth, G., Kahn, G., Yang, G., Wang, G., Su, H., Fang, H.S., Shi, H., Bao, H., Amor, H.B., Christensen, H.I., Furuta, H., Bharadhwaj, H., Walke, H., Fang, H., Ha, H., Mordatch, I., Radosavovic, I., Leal, I., Liang, J., Abou-Chakra, J., Kim, J., Drake, J., Peters, J., Schneider, J., Hsu, J., Vakil, J., Bohg, J., Bingham, J., Wu, J., Gao, J., Hu, J., Wu, J., Wu, J., Sun, J., Luo, J., Gu, J., Tan, J., Oh, J., Wu, J., Lu, J., Yang, J., Malik, J., Silvério, J., Hejna, J., Booher, J., Tompson, J., Yang, J., Salvador, J., Lim, J.J., Han, J., Wang, K., Rao, K., Pertsch, K., Hausman, K., Go, K., Gopalakrishnan, K., Goldberg, K., Byrne, K., Oslund, K., Kawaharazuka, K., Black, K., Lin, K., Zhang, K., Ehsani, K., Lekkala, K., Ellis, K., Rana, K., Srinivasan, K., Fang, K., Singh, K.P., Zeng, K.H., Hatch, K., Hsu, K., Itti, L., Chen, L.Y., Pinto, L., Fei-Fei, L., Tan, L., Fan, L.J., Ott, L., Lee, L., Weihs, L., Chen, M., Lepert, M., Memmel, M., Tomizuka, M., Itkina, M., Castro, M.G., Spero, M., Du, M., Ahn, M., Yip, M.C., Zhang, M., Ding, M., Heo, M., Srirama, M.K., Sharma, M., Kim, M.J., Irshad, M.Z., Kanazawa, N., Hansen, N., Heess, N., Joshi, N.J., Suenderhauf, N., Liu, N., Palo, N.D., Shafiullah, N.M.M., Mees, O., Kroemer, O., Bastani, O., Sanketi, P.R., Miller, P.T., Yin, P., Wohlhart, P., Xu, P., Fagan, P.D., Mitrano, P., Sermanet, P., Abbeel, P., Sundaresan, P., Chen, Q., Vuong, Q., Rafailov, R., Tian, R., Doshi, R., Martín-Martín, R., Baijal, R., Scalise, R., Hendrix, R., Lin, R., Qian, R., Zhang, R., Mendonca, R., Shah, R., Hoque, R., Julian, R., Bustamante, S., Kirmani, S., Levine, S., Lin, S., Moore, S., Bahl, S., Dass, S., Sonawani, S., Tulsiani, S., Song, S., Xu, S., Haldar, S., Karamcheti, S., Adebola, S., Guist, S., Nasiriany, S., Schaal, S., Welker, S., Tian, S., Ramamoorthy, S., Dasari, S., Belkhale, S., Park, S., Nair, S., Mirchandani, S., Osa, T., Gupta, T., Harada, T., Matsushima, T., Xiao, T., Kollar, T., Yu, T., Ding, T., Davchev, T., Zhao, T.Z., Armstrong, T., Darrell, T., Chung, T., Jain, V., Kumar, V., Vanhoucke, V., Guizilini, V., Zhan, W., Zhou, W., Burgard, W., Chen, X., Chen, X., Wang, X., Zhu, X., Geng, X., Liu, X., Liangwei, X., Li, X., Pang, Y., Lu, Y., Ma, Y.J., Kim, Y., Chebotar, Y., Zhou, Y., Zhu, Y., Wu, Y., Xu, Y., Wang, Y., Bisk, Y., Dou, Y., Cho, Y., Lee, Y., Cui, Y., Cao, Y., Wu, Y.H., Tang, Y., Zhu, Y., Zhang, Y., Jiang, Y., Li, Y., Li, Y., Iwasawa, Y., Matsuo, Y., Ma, Z., Xu, Z., Cui, Z.J., Zhang, Z., Fu, Z., Lin, Z.: Open x-embodiment: Robotic learning datasets and rt-x models (2025), <https://arxiv.org/abs/2310.08864>
9. Dosovitskiy, A., Beyer, L., Kolesnikov, A., Weissenborn, D., Zhai, X., Unterthiner, T., Dehghani, M., Minderer, M., Heigold, G., Gelly, S., Uszkoreit, J., Houlsby, N.: An image is worth 16x16 words: Transformers for image recognition at scale. In: ICLR (2021)
 10. Driess, D., Xia, F., Sajjadi, M.S.M., Lynch, C., Chowdhery, A., Ichter, B., Wahid, A., Tompson, J., Vuong, Q., Yu, T., Huang, W., Chebotar, Y., Sermanet, P., Duckworth, D., Levine, S., Vanhoucke, V., Hausman, K., Toussaint, M., Greff, K., Zeng, A., Mordatch, I., Florence, P.: Palm-e: An embodied multimodal language model. In: ICML (2023)
 11. Fang, H.S., Fang, H., Tang, Z., Liu, J., Wang, C., Wang, J., Zhu, H., Lu, C.: Rh20t: A comprehensive robotic dataset for learning diverse skills in one-shot (2023), <https://arxiv.org/abs/2307.00595>
 12. Gao, S., Zhou, S., Du, Y., Zhang, J., Gan, C.: Adaworld: Learning adaptable world models with latent actions. In: ICML (2025)
 13. Garrido, Q., Nagarajan, T., Terver, B., Ballas, N., LeCun, Y., Rabbat, M.: Learning latent action world models in the wild (2026), <https://arxiv.org/abs/2601.05230>

14. Guo, Y., Shi, L.X., Chen, J., Finn, C.: Ctrl-world: A controllable generative world model for robot manipulation. In: ICLR (2026), <https://iclr.cc/virtual/2026/poster/10011332>
15. Hafner, D., Lillicrap, T., Ba, J., Norouzi, M.: Dream to control: Learning behaviors by latent imagination. In: ICLR (2020)
16. Hafner, D., Lillicrap, T., Fischer, I., Villegas, R., Ha, D., Lee, H., Davidson, J.: Learning latent dynamics for planning from pixels. In: ICML (2019)
17. Hafner, D., Lillicrap, T., Norouzi, M., Ba, J.: Mastering atari with discrete world models. In: ICLR (2021)
18. Hafner, D., Pasukonis, J., Ba, J., Lillicrap, T.: Mastering diverse domains through world models (2024), <https://arxiv.org/abs/2301.04104>
19. Heusel, M., Ramsauer, H., Unterthiner, T., Nessler, B., Hochreiter, S.: Gans trained by a two time-scale update rule converge to a local nash equilibrium. In: NeurIPS (2017), https://proceedings.neurips.cc/paper_files/paper/2017/file/8a1d694707eb0fefe65871369074926d-Paper.pdf
20. Intelligence, P., Black, K., Brown, N., Darpinian, J., Dhabalia, K., Driess, D., Esmail, A., Equi, M., Finn, C., Fusai, N., Galliker, M.Y., Ghosh, D., Groom, L., Hausman, K., Ichter, B., Jakubczak, S., Jones, T., Ke, L., LeBlanc, D., Levine, S., Li-Bell, A., Mothukuri, M., Nair, S., Pertsch, K., Ren, A.Z., Shi, L.X., Smith, L., Springenberg, J.T., Stachowicz, K., Tanner, J., Vuong, Q., Walke, H., Walling, A., Wang, H., Yu, L., Zhilinsky, U.: $\pi_{0.5}$: a vision-language-action model with open-world generalization (2025), <https://arxiv.org/abs/2504.16054>
21. Jang, J., Ye, S., Lin, Z., Xiang, J., Bjorck, J., Fang, Y., Hu, F., Huang, S., Kundalia, K., Lin, Y.C., Magne, L., Mandlekar, A., Narayan, A., Tan, Y.L., Wang, G., Wang, J., Wang, Q., Xu, Y., Zeng, X., Zheng, K., Zheng, R., Liu, M.Y., Zettlemoyer, L., Fox, D., Kautz, J., Reed, S., Zhu, Y., Fan, L.: Dreamgen: Unlocking generalization in robot learning through video world models (2025), <https://arxiv.org/abs/2505.12705>
22. Jiang, Y., Gu, Y., Tsang, I.W., Shou, M.Z.: Olaf-world: Orienting latent actions for video world modeling (2026), <https://arxiv.org/abs/2602.10104>
23. Kalashnikov, D., Irpan, A., Pastor, P., Ibarz, J., Herzog, A., Jang, E., Quillen, D., Holly, E., Kalakrishnan, M., Vanhoucke, V., Levine, S.: Qt-opt: Scalable deep reinforcement learning for vision-based robotic manipulation (2018), <https://arxiv.org/abs/1806.10293>
24. Khazatsky, A., Pertsch, K., Nair, S., Balakrishna, A., Dasari, S., Karamcheti, S., Nasiriany, S., Srirama, M.K., Chen, L.Y., Ellis, K., Fagan, P.D., Hejna, J., Itkina, M., Lepert, M., Ma, Y.J., Miller, P.T., Wu, J., Belkhale, S., Dass, S., Ha, H., Jain, A., Lee, A., Lee, Y., Memmel, M., Park, S., Radosavovic, I., Wang, K., Zhan, A., Black, K., Chi, C., Hatch, K.B., Lin, S., Lu, J., Mercat, J., Rehman, A., Sanketi, P.R., Sharma, A., Simpson, C., Vuong, Q., Walke, H.R., Wulfe, B., Xiao, T., Yang, J.H., Yavary, A., Zhao, T.Z., Agia, C., Baijal, R., Castro, M.G., Chen, D., Chen, Q., Chung, T., Drake, J., Foster, E.P., Gao, J., Guizilini, V., Herrera, D.A., Heo, M., Hsu, K., Hu, J., Irshad, M.Z., Jackson, D., Le, C., Li, Y., Lin, K., Lin, R., Ma, Z., Maddukuri, A., Mirchandani, S., Morton, D., Nguyen, T., O’Neill, A., Scalise, R., Seale, D., Son, V., Tian, S., Tran, E., Wang, A.E., Wu, Y., Xie, A., Yang, J., Yin, P., Zhang, Y., Bastani, O., Berseth, G., Bohg, J., Goldberg, K., Gupta, A., Gupta, A., Jayaraman, D., Lim, J.J., Malik, J., Martín-Martín, R., Ramamoorthy, S., Sadigh, D., Song, S., Wu, J., Yip, M.C., Zhu, Y., Kollar, T., Levine, S., Finn, C.: Droid: A large-scale in-the-wild robot manipulation dataset. In: RSS (2024)

25. Kim, M.J., Pertsch, K., Karamcheti, S., Xiao, T., Balakrishna, A., Nair, S., Rafailov, R., Foster, E., Lam, G., Sanketi, P., Vuong, Q., Kollar, T., Burchfiel, B., Tedrake, R., Sadigh, D., Levine, S., Liang, P., Finn, C.: Openvla: An open-source vision-language-action model. In: CORL (2024)
26. Lambert, N., Pister, K., Calandra, R.: Investigating compounding prediction errors in learned dynamics models (2022), <https://arxiv.org/abs/2203.09637>
27. Levine, S., Finn, C., Darrell, T., Abbeel, P.: End-to-end training of deep visuomotor policies. *JMLR* **17**(39), 1–40 (2016)
28. NVIDIA, Bjorck, J., Castañeda, F., Cherniadev, N., Da, X., Ding, R., Fan, L.J., Fang, Y., Fox, D., Hu, F., Huang, S., Jang, J., Jiang, Z., Kautz, J., Kundalia, K., Lao, L., Li, Z., Lin, Z., Lin, K., Liu, G., Llontop, E., Magne, L., Mandlekar, A., Narayan, A., Nasiriany, S., Reed, S., Tan, Y.L., Wang, G., Wang, Z., Wang, J., Wang, Q., Xiang, J., Xie, Y., Xu, Y., Xu, Z., Ye, S., Yu, Z., Zhang, A., Zhang, H., Zhao, Y., Zheng, R., Zhu, Y.: Gr00t n1: An open foundation model for generalist humanoid robots (2025), <https://arxiv.org/abs/2503.14734>
29. Papineni, K., Roukos, S., Ward, T., Zhu, W.J.: BLEU: a method for automatic evaluation of machine translation. In: *Proceedings of ACL* (2002)
30. Quevedo, J., Sharma, A.K., Sun, Y., Suryavanshi, V., Liang, P., Yang, S.: Worldgym: World model as an environment for policy evaluation (2025), <https://arxiv.org/abs/2506.00613>
31. Radford, A., Kim, J.W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., Krueger, G., Sutskever, I.: Learning transferable visual models from natural language supervision. In: *ICML* (2021)
32. Reimers, N., Gurevych, I.: Sentence-BERT: Sentence embeddings using siamese BERT-networks. In: *Conf. Empirical Methods Natural Lang. Process.* (2019)
33. Sarch, G., Wu, Y., Tarr, M.J., Fragkiadaki, K.: Open-ended instructable embodied agents with memory-augmented large language models. In: *EMNLP* (2023)
34. Shridhar, M., Manuelli, L., Fox, D.: Cliport: What and where pathways for robotic manipulation. In: *CORL* (2021)
35. Torabi, F., Warnell, G., Stone, P.: Behavioral cloning from observation (2018), <https://arxiv.org/abs/1805.01954>
36. Unterthiner, T., van Steenkiste, S., Kurach, K., Marinier, R., Michalski, M., Gelly, S.: Towards accurate generative models of video: A new metric & challenges (2018), <https://arxiv.org/abs/1812.01717>
37. Wang, Y., Zhang, F., Zhan, D.C., Zhao, L., Wang, K., Bian, J.: Co-evolving latent action world models (2025), <https://arxiv.org/abs/2510.26433>
38. Wang, Y., Xian, Z., Chen, F., Wang, T.H., Wang, Y., Fragkiadaki, K., Erickson, Z., Held, D., Gan, C.: Robogen: Towards unleashing infinite data for automated robot learning via generative simulation. In: *ICML* (2024)
39. Wang, Z., Shi, C., Hu, J., Rohling, K., Martín-Martín, R., Zhang, A., Stone, P.: Factored latent action world models (2026), <https://arxiv.org/abs/2602.16229>
40. Wu, S., Luo, X., Zhang, J., Xie, J., Song, J., Shen, H.T., Gao, L.: Policy contrastive decoding for robotic foundation models (2025), <https://arxiv.org/abs/2505.13255>
41. Yang, J., Lin, K., Li, J., Zhang, W., Lin, T., Wu, L., Su, Z., Zhao, H., Zhang, Y.Q., Chen, L., Luo, P., Yue, X., Li, H.: Rise: Self-improving robot policy with compositional world model (2026), <https://arxiv.org/abs/2602.11075>
42. Zeng, A., Florence, P., Tompson, J., Welker, S., Chien, J., Attarian, M., Armstrong, T., Krasin, I., Duong, D., Wahid, A., Sindhwani, V., Lee, J.: Transporter networks: Rearranging the visual world for robotic manipulation. In: *CORL* (2020)

43. Zhao, T.Z., Tompson, J., Driess, D., Florence, P., Ghasemipour, K., Finn, C., Wahid, A.: Aloha unleashed: A simple recipe for robot dexterity. In: CORL (2024)
44. Zhu, F., Wu, H., Guo, S., Liu, Y., Cheang, C., Kong, T.: Irasim: A fine-grained world model for robot manipulation (2024), <https://arxiv.org/abs/2406.14540>
45. Zhu, M., Zhu, Y., Li, J., Wen, J., Xu, Z., Liu, N., Cheng, R., Shen, C., Peng, Y., Feng, F., Tang, J.: Scaling diffusion policy in transformer to 1 billion parameters for robotic manipulation. In: ICRA (2025)