

HUMAN VIDEOS ARE NOT JUST FOR IMITATION

Robot Self-Improvement via Human-Video Dynamics Models

PREPRINT 2026

Hanzhi Chen^{*,1,2,4}, Anran Zhang^{*,1,2,4}, Simon Schaefer^{1,2,4}, Kejia Chen², Shi Chen¹,
Daniel Cremers^{2,4}, Oier Mees^{†,3,1}, Stefan Leutenegger^{†,1}

^{*}Equal Contribution [†]Equal Advising

¹ETH Zürich ²Technical University of Munich ³Microsoft ⁴MCML

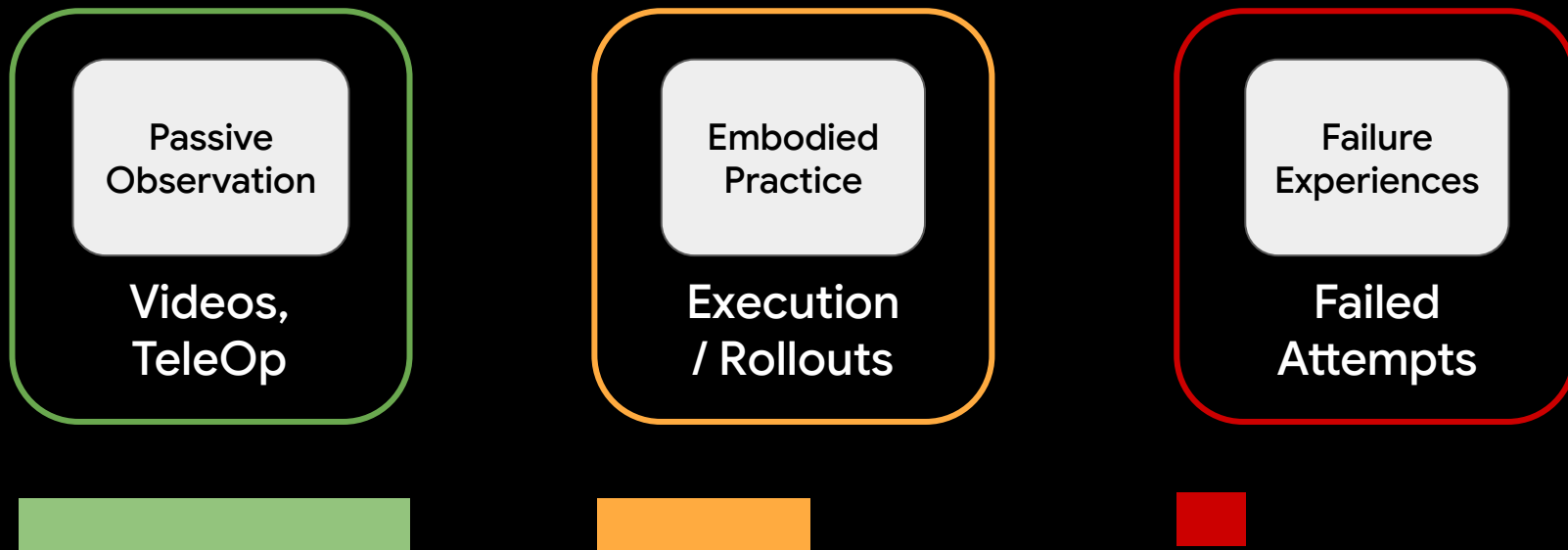
 arXiv

 Video

 Code (Coming Soon)

Presented by: [Jishnu Jaykumar P](#) | 9/4/26
Reading Group | [IRVL @ UTD](#)

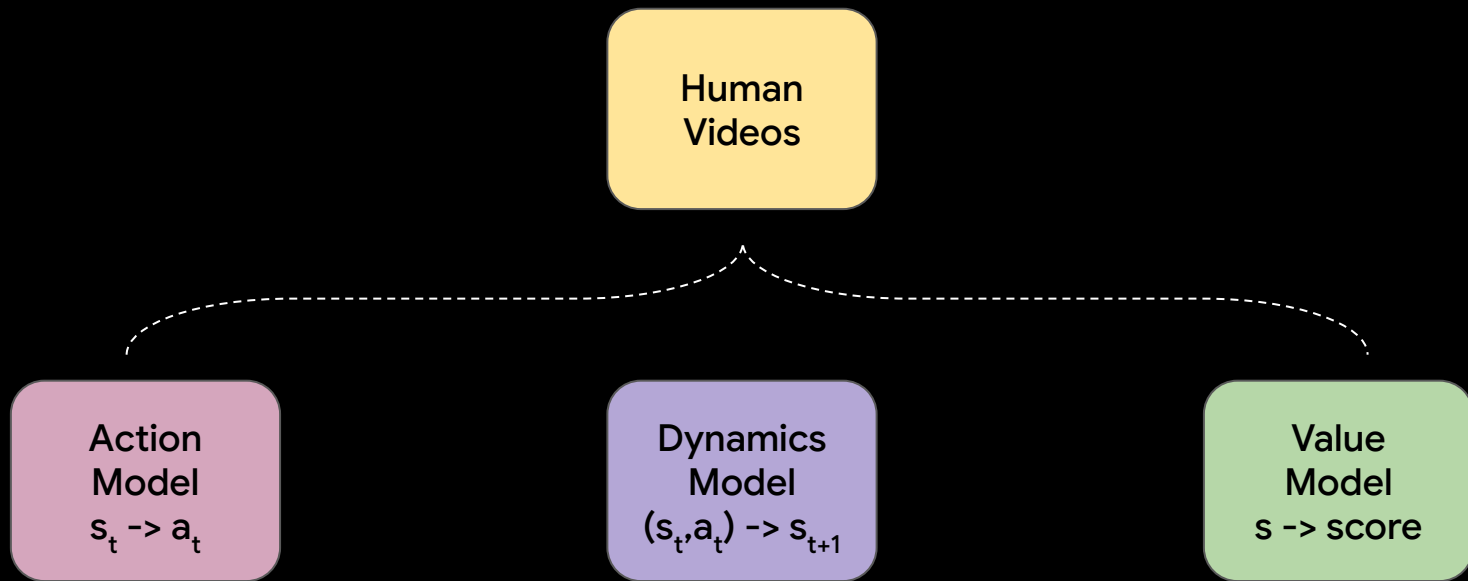
Acquire skills from the kinds of data that humans learn from



Sequence enables self improvement



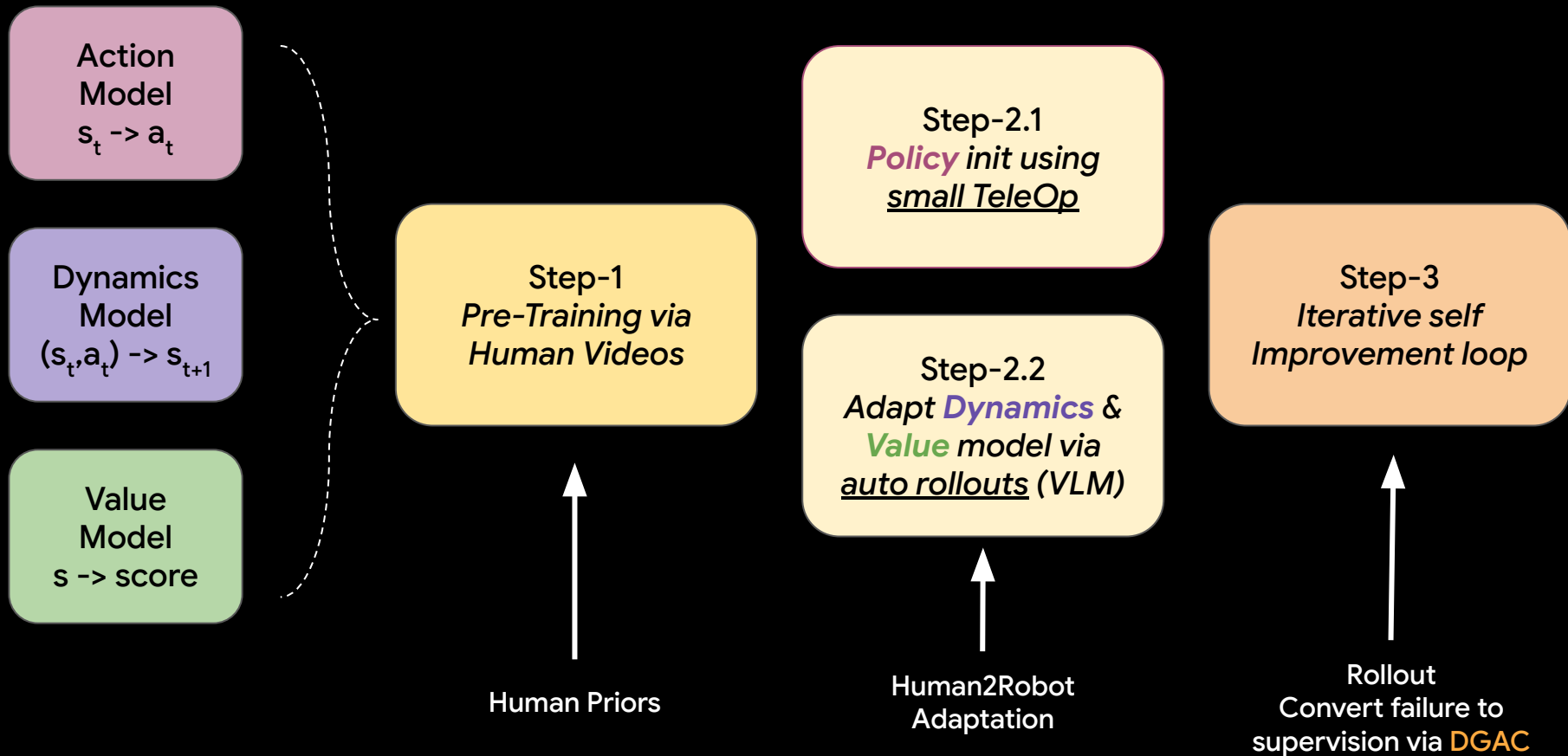
This work



Claim: Representation enables X-embodiment

Goal: Autonomously improve from own rollouts and failures

Pipeline



Step-1

Pre-Training using Human Videos

Step-1: Pre-Training using Human Videos

$$\mathbf{a}_t = [\xi_t, c_t]$$

absolute 6-DoF wrist motion and $c_t \in [0,1]$

Visual Tokens

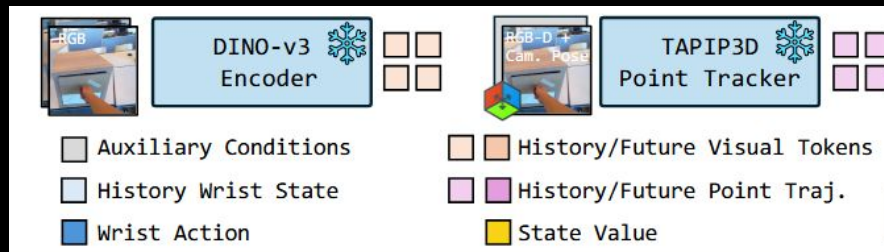
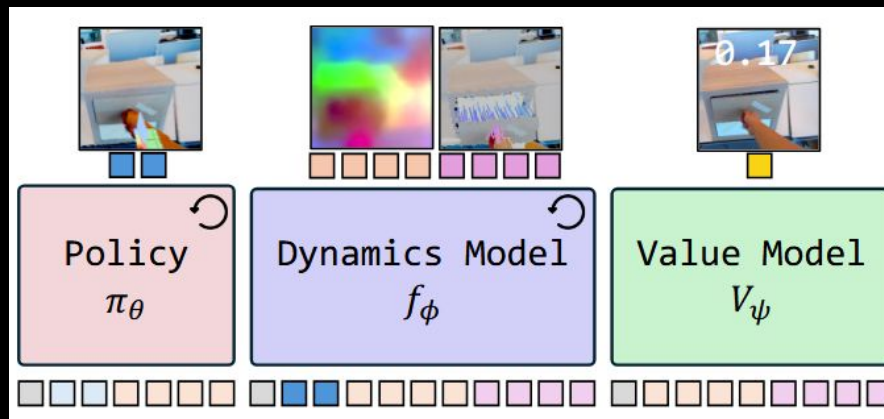
Point Traj.

State Value

Action
Model

Dyn Model

Value
Model



Policy Model: $\hat{\mathbf{a}}_{t:t+H-1} = \pi_{\theta}(\mathbf{s}_{t-H'+1:t}, \mathbf{z}_{t-H'+1:t}; \mathbf{m}_t^{\pi})$,

$$\mathcal{L}_{\pi} = \|\hat{\mathbf{a}}_{t:t+H-1}^{\tau} - \mathbf{a}_{t:t+H-1}\|_2^2$$

Dynamics Model: $\hat{\mathbf{o}}_{t+H-1} = f_{\phi}(\mathbf{o}_{t-H'+1:t}, \mathbf{a}_{t:t+H-1}; \mathbf{m}_t^f)$,

$$\mathcal{L}_f = \|\hat{\mathbf{o}}_{t+H-1}^{\tau} - \mathbf{o}_{t+H-1}\|_2^2$$

Value Model: $\hat{V}_t = V_{\psi}(\mathbf{o}_t; \mathbf{m}_t^V)$.

$$\mathcal{L}_V = \|\hat{V}_t - V_t\|_2^2$$

$\sum_{t'=t}^T \gamma^{t'-t} r_{t'} = \gamma^{T-t} r_T$ where γ is the discount factor, T is the terminal timestep, all intermediate rewards are 0, and the terminal reward $r_T \in \{1, -1\}$ denotes success or failure.

$\mathbf{m}_t^{\pi}$, $\mathbf{m}_t^f$, and $\mathbf{m}_t^V$

Aux conditions.

E.g. language

Human Datasets:

HOI4D, Arti4D, EgoDex

Step-2

Human 2 Robot Adaptation

2.1 Policy

2.2 Dynamics and Value

Step-2.1: Policy init

Action
Model Π_{ref}
From Human Vids

01 / Data Collection

- Collect a small amount of teleoperation data $\mathbf{D}_{\text{tele}}$
- Anchor the policy to the target environment
- Support multiple interaction modes

02 / Policy Init

- Start from human-pretrained policy
- Initialize robot policy with same architecture
- Train on $\mathbf{D}_{\text{tele}}$ using action loss $\mathcal{L}$

03 / Conditioning

- Condition model on gripper-view observations $\mathbf{m}_t$
- Set **improvement indicator** $I_t = 1$ for all initialization samples

$$I_t \in \{0, 1\}$$

$$\hat{\mathbf{a}}_{t:t+H-1} = \Pi(\mathbf{s}_{t-H'+1:t}, \mathbf{z}_{t-H'+1:t}; \mathbf{m}_t, \mathbf{I}_t)$$

$$\mathcal{L}_\pi = \|\hat{\mathbf{a}}_{t:t+H-1}^\tau - \mathbf{a}_{t:t+H-1}\|_2^2$$

$$\hat{A}(\mathbf{o}_t, \mathbf{a}_{t:t+H-1}) = \sum_{t'=t}^{t+H-1} r_{t'} + V_\psi(\mathbf{o}_{t+H}) - V_\psi(\mathbf{o}_t)$$

$$\hat{\pi}(\mathbf{a}_{t:t+H-1} \mid \mathbf{o}_t) \propto \pi_{\text{ref}}(\mathbf{a}_{t:t+H-1} \mid \mathbf{o}_t) p(I_t \mid \hat{A}(\mathbf{o}_t, \mathbf{a}_{t:t+H-1}))^\beta$$

$\mathbf{1}(A(\mathbf{o}_t, \mathbf{a}_{t:t+H-1}, \ell) > \epsilon)$ with a task-specific threshold ϵ

Step-2.2: Adapt Dynamics & Value

Dynamics
Model

Value
Model

01 / Why Adapt?

- Models don't transfer perfectly because robot bodies are different
- Mistakes and biases only show up when testing on the real robot

02 / How to Adapt

- Use real robot practice data to ground the AI models
- Tune predictions to fit the specific robot's unique actions

03 / Autonomous Scale

- Skip the high cost of humans guiding the robot manually
- Let AI suggest tasks while the robot practice-runs them on its own
- GPT-4 proposes plausible atomic interaction instructions (e.g., close the drawer), which the frozen human-pretrained policy then executes on the robot.
- These rollouts expose the models to robot states and natural failures

For each scene, we first use the pretrained affordance model from VidBot [8] to detect plausible contact regions. We then select a target contact point and plan a pre-contact trajectory that moves the gripper to the corresponding region while avoiding collisions. Once the gripper reaches the contact region, we query the human-pretrained policy to generate the post-contact interaction trajectory. The robot follows the predicted end-effector poses using its low-level controller, while the gripper is smoothly adjusted along the trajectory to improve execution stability. This strategy separates coarse contact acquisition from post-contact interaction, allowing the robot to explore semantically meaningful behaviors even when the human-pretrained policy does not provide precise contact poses. The collected rollouts contain both successful and failed interactions caused by robot-specific contacts, dynamics, and scene configurations. We use these autonomous rollouts to adapt the dynamics and value models before task-level self-improvement.

Step-3

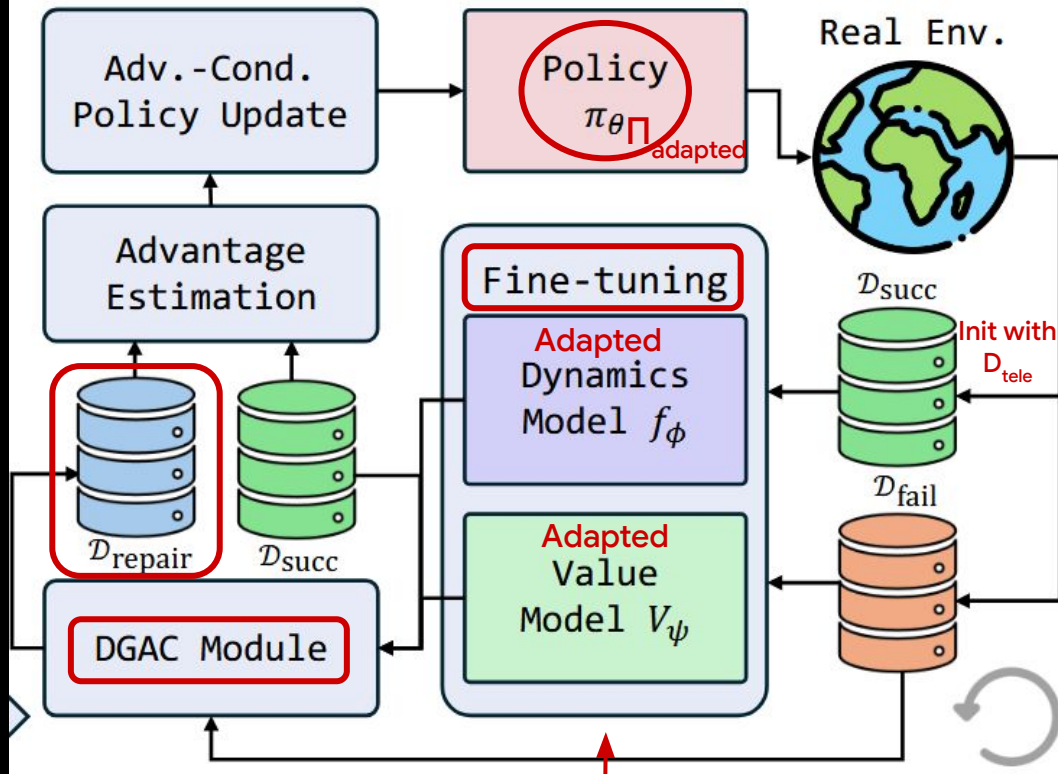
Iterative self improvement

Iterative Robot Self-Improvement Loop

$$\hat{\pi}(\mathbf{a}_{t:t+H-1} \mid \mathbf{o}_t)$$

$$\pi_{\text{ref}}(\mathbf{a}_{t:t+H-1} \mid \mathbf{o}_t) p(I_t \mid \hat{A}(\mathbf{o}_t, \mathbf{a}_{t:t+H-1}))^\beta$$

$$\hat{A}(\mathbf{o}_t, \mathbf{a}_{t:t+H-1}) = \sum_{t'=t}^{t+H-1} r_{t'} + V_\psi(\mathbf{o}_{t+H}) - V_\psi(\mathbf{o}_t)$$

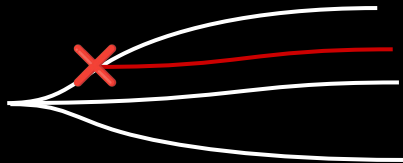


This column models are from step-2: adaptation

Only place where human is needed for determining the outcome Succ / Fail

Keep on doing this. Here authors have used 2 rounds

DGAC: Dynamics Guided Action Correction



Naive copy
or repair ✓

Algorithm 1: Dynamics-Guided Action Correction (DGAC)

Require: Policy π_θ , dynamics model f_ϕ , value model V_ψ , thresholds δ_v , δ_s

Input: Failed state $\mathbf{o}_t$, successful dataset $\mathcal{D}_{\text{succ}}$

Output: Corrected action $\mathbf{a}_{t:t+H-1}^*$ or $\emptyset$

Find $\mathcal{O}_t^{(k)} \subset \mathcal{D}_{\text{succ}}$ and select $\tilde{\mathbf{o}}_t$ by Eq. 4

if $|V_\psi(\tilde{\mathbf{o}}_t) - V_\psi(\mathbf{o}_t)| > \delta_v$ **or** $\text{sim}(\mathbf{o}_t, \tilde{\mathbf{o}}_t) < \delta_s$ **then**
return $\emptyset$

Sample action candidates $\{\mathbf{a}_{t:t+H-1}^{(n)}\}_{n=1}^N$ via Eq. 5

for $n = 1$ **to** N **do**

$\hat{\mathbf{o}}_{t+H-1}^{(n)} \leftarrow f_\phi(\mathbf{o}_{t-H'+1:t}, \mathbf{a}_{t:t+H-1}^{(n)})$

$V_{t+H-1}^{(n)} \leftarrow V_\psi(\hat{\mathbf{o}}_{t+H-1}^{(n)})$

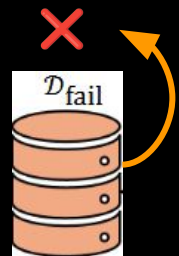
$\mathbf{a}_{t:t+H-1}^* \leftarrow \text{argmax}_n V_{t+H-1}^{(n)}$

return $\mathbf{a}_{t:t+H-1}^*$



retrieve the
steering
reference from
 $\mathcal{D}_{\text{succ}}$
in two
stages

Failed state $\mathbf{o}_t$



$$|V_\psi(\tilde{\mathbf{o}}_t) - V_\psi(\mathbf{o}_t)| \leq \delta_v \quad \text{Nearest task progress candidates}$$

$$\text{sim}(\mathbf{o}_t, \tilde{\mathbf{o}}_t) \geq \delta_s \quad \text{Nearest similar candidate}$$

$$\mathcal{O}_t^{(k)} = \text{TopK}_{\mathbf{o}' \in \mathcal{D}_{\text{succ}}} (-|V_\psi(\mathbf{o}') - V_\psi(\mathbf{o}_t)|)$$

$$\tilde{\mathbf{o}}_t = \arg \max_{\mathbf{o}' \in \mathcal{O}_t^{(k)}} \text{sim}(\mathbf{o}_t, \mathbf{o}')$$

We know the timestep.
Query the velocity field over
noisy action chunks

$$\mathbf{v}_\theta(\mathbf{a}_{t:t+H-1}^\tau, \tau; \cdot)$$

Compose $\mathbf{o}_t$ and retrieved $\mathbf{o}_t$
Velocity fields

$$w_n \sim \mathcal{U}(0.1, 1)$$

$$\mathbf{v}_\theta^{(n)} = \mathbf{v}_\theta(\mathbf{a}_{t:t+H-1}^\tau, \tau; \mathbf{h}_t) + w_n [\mathbf{v}_\theta(\mathbf{a}_{t:t+H-1}^\tau, \tau; \tilde{\mathbf{h}}_t) - \mathbf{v}_\theta(\mathbf{a}_{t:t+H-1}^\tau, \tau; \mathbf{h}_t)]. \quad (5)$$

Dynamics Guided Action Correction:
Failure becomes supervision signal now for
self improvement

DGAC: Dynamics Guided Action Correction

Sample action candidates $\{\mathbf{a}_{t:t+H-1}^{(n)}\}_{n=1}^N$ via Eq. 5

for $n = 1$ to N do

$$\hat{\mathbf{o}}_{t+H-1}^{(n)} \leftarrow f_{\phi}(\mathbf{o}_{t-H'+1:t}, \mathbf{a}_{t:t+H-1}^{(n)})$$

$$V_{t+H-1}^{(n)} \leftarrow V_{\psi}(\hat{\mathbf{o}}_{t+H-1}^{(n)})$$

$$\mathbf{a}_{t:t+H-1}^* \leftarrow \operatorname{argmax}_n V_{t+H-1}^{(n)}$$

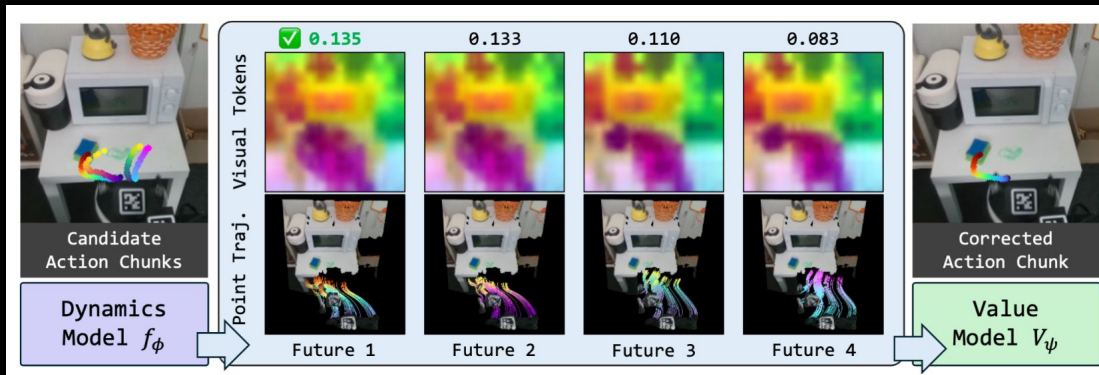
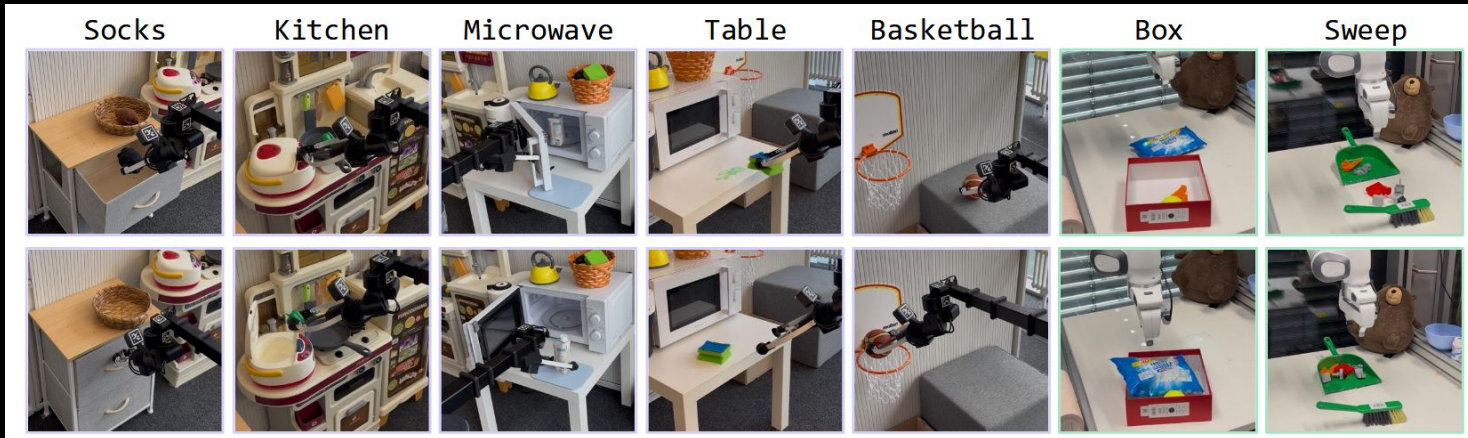


Figure 5: DGAC repairs failure-inducing actions from the original policy (top) by steering predicted futures toward successful references (middle), and training on these corrections leads to successful rollouts (bottom).

Results



Method	T1	T2	T3	T4	T5	Avg.
Expert BC	33.3	46.7	46.7	40.0	40.0	41.3
SWIM [62]	53.3	40.0	46.7	46.7	60.0	49.3
LPB [63]	53.3	40.0	66.7	80.0	60.0	60.0
AWR [61]	53.3	53.3	80.0	60.0	53.3	60.0
RECAP [†] [2]	60.0	53.3	73.3	53.3	66.7	61.3
RISE [†] [3]	86.7	60.0	73.3	80.0	80.0	76.0
Ours	86.7	80.0	80.0	93.3	86.7	85.3

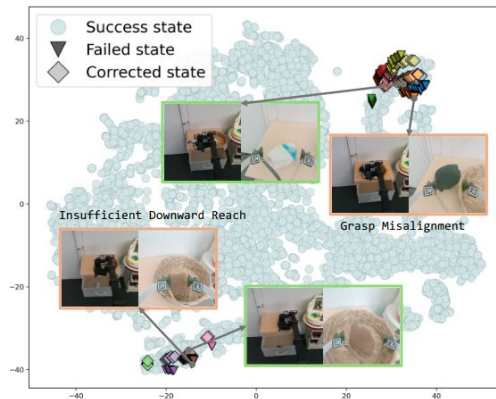
Table 1: Benchmark evaluated on success rate (%). T1: Socks, T2: Kitchen, T3: Microwave, T4: Table, T5: Basketball. [†]: no human-intervened corrections.

Method	T1	T2	T3	T4	T5	Avg.
$\pi_{0.5} + \text{SFT}$	60.0	66.7	80.0	40.0	66.7	62.7
$\pi_{0.5} + \text{RECAP}^{\dagger}$	66.7	66.7	73.3	53.3	80.0	68.0
$\pi_{0.5} + \text{DGAC}$	93.3	86.7	86.7	86.7	86.7	88.0
Ours	86.7	80.0	80.0	93.3	86.7	85.3

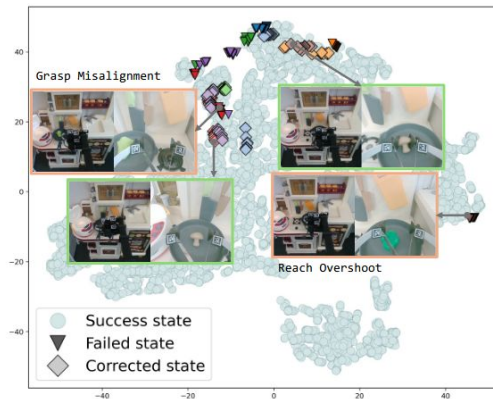
Table 3: Success rates (%) of the baseline policy $\pi_{0.5}$ under different policy improvement schemes. [†]: no human-intervened corrections.

t-SNE Dynamics Viz

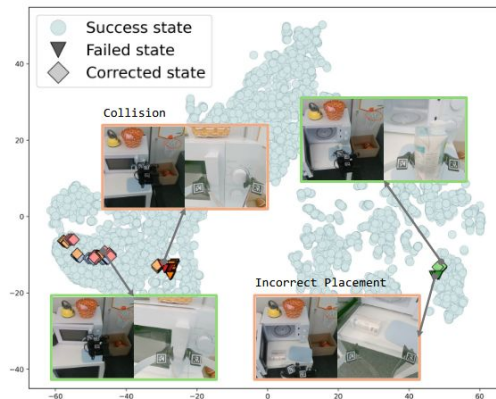
Socks



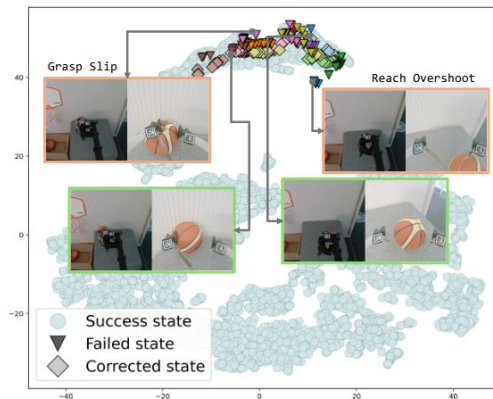
Kitchen



Microwave



Basketball



Ablations

For module specific ablations.
Please refer to the paper.

Limitations & Future Work

01

Human Supervision

- Relies on human annotations for rollout outcomes
- Needs foundation models for automated feedback

02

Horizon Limits

- Operates locally over short-horizon action chunks
- Requires extending to longer-horizon dynamics

03

Pretraining & Scale

- Scale pretraining to larger, diverse datasets
- Improves generalization across environments

Fin.

Q/A