

Guide, Think, Act: Interactive Embodied Reasoning in Vision-Language-Action Models

Yiran Ling^{2,3,7}^{*,†}, Qing Lian¹^{*}, Jinghang Li^{1,4}, Qing Jiang^{3,5}, Tianming Zhang⁶, Xiaoke Jiang^{3,6}, Chuanxiu Liu⁶, Jie Liu^{2,7}[†], and Lei Zhang^{1,3,6}[†]

¹ Futian Laboratory

² Faculty of Computing, Harbin Institute of Technology

³ International Digital Economy Academy (IDEA)

⁴ School of Robotics, Hunan University

⁵ South China University of Technology

⁶ Visincept

⁷ National Key Laboratory of Smart Farm Technologies and Systems

Abstract. In this paper, we propose **GTA-VLA** (**Guide, Think, Act**), an interactive Vision-Language-Action (VLA) framework that enables spatially steerable embodied reasoning by allowing users to guide robot policies with explicit visual cues. Existing VLA models learn a direct "Sense-to-Act" mapping from multimodal observations to robot actions. While effective within the training distribution, such tightly coupled policies are brittle under out-of-domain (OOD) shifts and difficult to correct when failures occur. Although recent embodied Chain-of-Thought (CoT) approaches expose intermediate reasoning, they still lack a mechanism for incorporating human spatial guidance, limiting their ability to resolve visual ambiguities or recover from mistakes. To address this gap, our framework allows users to optionally guide the policy with spatial priors, such as affordance points, boxes, and traces, which the subsequent reasoning process can directly condition on. Based on these inputs, the model generates a unified spatial-visual Chain-of-Thought that integrates external guidance with internal task planning, aligning human visual intent with autonomous decision-making. For practical deployment, we further couple the reasoning module with a lightweight reactive action head for efficient action execution. Extensive experiments demonstrate the effectiveness of our approach. On the in-domain SimplerEnv WidowX benchmark, our framework achieves a state-of-the-art 81.2% success rate. Under OOD visual shifts and spatial ambiguities, a single visual interaction substantially improves task success over existing methods, highlighting the value of interactive reasoning for failure recovery in embodied control.

Keywords: Vision-language-action models · Embodied reasoning · Interactive robot learning

^{*} Equal contribution. Emails: 25b903029@stu.hit.edu.cn, lianqing1997@gmail.com

[†] Internship at Futian Laboratory. [‡] Corresponding author



Fig. 1: Conventional direct VLA policies can fail under spatial ambiguity or imprecise grounding, since they lack an explicit mechanism for interactive correction. **GTA-VLA** resolves this by using one-shot spatial guidance (affordance points, boxes, or traces) to correct grounding and enable accurate execution.

1 Introduction

The pursuit of robust generalist robotic agents for open-world environments is a central goal of embodied AI. A major step toward this vision is the emergence of Vision-Language-Action (VLA) models [3, 4, 6, 8–10, 18, 21, 24, 29, 37, 38], which leverage large pre-trained vision language models to scale robot learning across diverse tasks and embodiments. Despite this progress, most existing VLAs still operate through an implicit direct “Sense-to-Act” mapping from multimodal observations to robot actions. While effective within the training distribution, such tightly coupled policies often become brittle under visual and semantic shifts, and provide little transparency when failures occur. When perception fails due to clutter, lighting variation, or unseen objects, humans have no explicit interface to re-ground the robot’s attention or provide targeted corrective guidance.

Recent work has begun to move beyond direct “Sense-to-Act” policies through Embodied Chain-of-Thought (CoT) reasoning [2, 11, 13, 20, 27, 35], shifting toward a more structured “Sense, Think, and Act” paradigm. By explicitly predicting intermediate representations, such as task decomposition, grounding cues, or motion plans, these methods improve interpretability and expose part of the policy’s decision-making process. However, in existing systems, the reasoning process remains largely self-contained: although intermediate reasoning is visible, it is still generated from the model’s internal belief alone and cannot be easily corrected when that internal grounding is wrong. This limitation is especially pronounced in out-of-domain (OOD) settings, where early mistakes in object grounding, contact localization, or motion targeting can propagate into coherent but incorrect plans. In such cases, human users can often resolve the ambiguity immediately

through simple spatial cues, such as pointing to a target, marking a grasp region, or sketching a desired path. Compared with language-only correction, these signals are more precise and more natural for communicating spatial intent, motivating the need for an interaction interface that allows human guidance to directly condition the policy’s reasoning process.

To address this gap, we propose **GTA-VLA**, (**Guide, Think, Act**), an interactive VLA framework that makes embodied reasoning explicitly steerable under human spatial guidance. The key idea is to treat affordances, boxes, and trajectories not as post-hoc corrections but as optional priors that the model’s reasoning process can directly condition on. With or without such guidance, the model produces a unified spatial-visual Chain-of-Thought that integrates external spatial intent with internal task understanding, visual grounding, affordance prediction, and action planning. As a result, the policy remains autonomous by default while becoming naturally correctable when failures or ambiguities arise.

To make this capability trainable at scale, we build an automated data pipeline that synthesizes large-scale interactive annotations from existing robot datasets, without requiring manual collection of human intervention traces. To mitigate the latency of autoregressive reasoning in practical control, we further decompose policy execution into a slow VLM reasoning module and a fast downstream action head. This asynchronous design allows high-level spatial-visual reasoning to run at a lower frequency, while a lightweight action module executes responsive low-level control at a higher frequency. To evaluate interactive embodied reasoning under distribution shift, we introduce *SimplerEnv-Plus*, an extension of *SimplerEnv* with more challenging OOD conditions, including camera variation, lighting changes, unseen objects, and language perturbations, while also supporting human spatial intervention during execution. Experiments in both simulation and the real world show that our framework improves not only autonomous task performance, but also failure recovery through minimal human interaction.

In summary, our contributions are three-fold:

1. We propose GTA-VLA (*Guide, Think, Act*), an interactive VLA framework that unifies explicit human spatial guidance with embodied Chain-of-Thought reasoning, enabling more steerable and interpretable robot policies.
2. We develop a scalable data generation pipeline for synthesizing interaction-style supervision from existing robot datasets, making guided spatial reasoning trainable at scale.
3. We introduce *Simpler-Plus*, a benchmark for evaluating interactive embodied policies under OOD conditions. Experiments in simulation and the real world demonstrate strong autonomous performance as well as substantial gains in failure recovery from minimal human intervention.

2 Related Work

Vision-Language-Action Models. Vision-Language-Action (VLA) models learn policies that map visual observations and language instructions to robot

actions, and have shown strong generalization across diverse scenes, tasks, and embodiments. Early end-to-end approaches [9, 18, 21, 29, 38] demonstrated the effectiveness of scaling imitation learning for real-world embodied manipulation. More recent dual-system architectures [3, 4, 6, 8, 10, 24, 37] further improve performance by pairing a vision-language backbone with a dedicated continuous-control module. Despite these advances, current VLA systems still rely heavily on large-scale behavioral data [16, 25, 32] and predominantly learn implicit action policies from demonstrations. As a result, while they are effective at direct policy execution, they offer limited support for explicit task understanding, interactive correction, and user-guided control when failures or ambiguities arise.

Visual and Embodied Reasoning and Chain of Thought. Recent work has explored explicit intermediate reasoning as a way to improve the generalization and interpretability of VLA policies. Several approaches [2, 11, 13] formulate robotic control as a structured reasoning process rather than direct action prediction. For example, ECoT [35] and Mind2Hand [27] introduce embodied Chain-of-Thought (CoT) reasoning to improve high-level planning and policy interpretability, while MolmoAct [20] further grounds intermediate reasoning in explicit spatial representations. A practical challenge in this line of work is that autoregressively generating explicit reasoning tokens can introduce substantial inference latency, which limits responsiveness in manipulation tasks. To improve efficiency, more recent methods have explored compact reasoning representations. For instance, Fast-ThinkAct [12] replaces explicit token generation with latent planning states to reduce inference cost. While such designs improve efficiency, they may also reduce the transparency and fine-grained controllability of the reasoning process compared with explicit spatially grounded intermediate representations.

Interactive Perception and Visual Prompting. Recent work on interactive perception has substantially improved the spatial grounding ability of foundation models. In computer vision, the SAM family [5, 19, 26] demonstrates that simple geometric prompts can support strong zero-shot segmentation, while subsequent systems such as T-Rex2 [15] and Rex-Omni [14] extend this paradigm to open-vocabulary and interactive object detection with both text and visual prompts. In parallel, multimodal large language models (MLLMs) [1, 7, 28] have also become increasingly capable of fine-grained visual grounding. Works such as Ferret [33, 34] and Set-of-Mark (SoM) show that language models can be conditioned on points, boxes, and marks to support precise spatial reference and coordinate-level reasoning. While these advances have significantly strengthened interactive 2D perception and grounding, extending such capabilities to embodied control remains nontrivial. Most existing visual-prompting systems are designed primarily for pixel-level perception tasks, such as segmentation, detection, or spatial reference, rather than for generating continuous motor actions. As a result, they do not directly address the temporal dynamics, control interfaces, or action generation requirements needed for robotic manipulation.

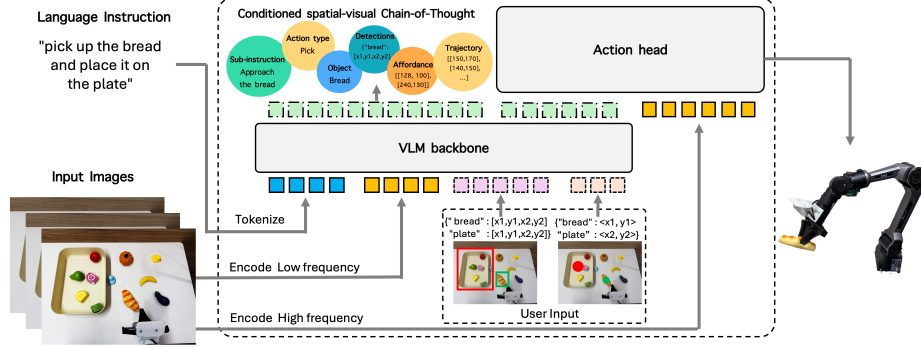


Fig. 2: Overview of GTA-VLA (Guide, Think, Act). The framework consists of three stages. **Guide:** the model receives the primary image, the language instruction, and optional spatial priors (e.g., affordance points, boxes, or traces). **Think:** the VLM backbone generates a conditioned spatial-visual reasoning sequence and the corresponding latent reasoning states $H_{\text{reasoning}}$. **Act:** a downstream Flow-Matching action head consumes the latest reasoning states together with high-frequency control observations to produce continuous action chunks. This design decouples slow autoregressive reasoning from fast closed-loop control.

3 Methodology

3.1 Preliminaries and Overall Architecture

Preliminaries. We formulate robotic manipulation as a conditional sequence modeling problem. At time step t , a standard Vision-Language-Action (VLA) policy π receives multi-view RGB observations $\mathcal{I}_t = \{I_t^{(v)}\}_{v=1}^V$, a natural language instruction L , and the robot proprioceptive state s_t , and predicts a future action chunk

$$A_t = [a_t, a_{t+1}, \dots, a_{t+k-1}], \quad (1)$$

where k denotes the size of the action chunk. The standard policy is therefore written as

$$\pi : (\mathcal{I}_t, L, s_t) \rightarrow A_t. \quad (2)$$

In our implementation, the multi-view observations consist of a primary external view and a wrist-mounted view.

We extend this formulation by introducing an optional spatial prior P_{spatial} , which provides sparse geometric guidance in image space and may be supplied either by a human user or by an expert annotation pipeline during training. Concretely, P_{spatial} may take the form of an affordance point, a bounding box, or a trace. The policy is thus extended to

$$\pi : (\mathcal{I}_t, L, s_t, P_{\text{spatial}}) \rightarrow A_t, \quad (3)$$

where P_{spatial} is optional and may be absent during fully autonomous execution.

Overall Architecture. Our framework is built on top of a vision-language backbone and a downstream continuous control module. We use Qwen3-VL-2B [1] as

the core VLM backbone due to its strong multimodal understanding and spatial grounding capabilities. Given the augmented input, the backbone first generates a structured spatial-visual reasoning sequence C , and we use the hidden states associated with these reasoning tokens as the latent reasoning representation, denoted by $H_{\text{reasoning}}$. In our implementation, the reasoning branch consumes only the primary image, the language instruction, and the optional spatial prior, while proprioceptive inputs and the wrist-view image are introduced only in the downstream fast control branch. These latent reasoning states are then consumed by a downstream action model to generate continuous control actions.

For action generation, we adopt a Flow-Matching action head [3, 24, 37], which models action chunks in continuous space and is well-suited for complex multi-modal action distributions. Concretely, the action branch takes the current control observation together with the latest reasoning states, combining the primary image, the wrist-view image, proprioceptive inputs, and $H_{\text{reasoning}}$ to predict continuous action chunks. In our implementation, actions are primarily parameterized as end-effector poses, although the framework is not restricted to this choice. While our experiments focus on single-arm manipulation, the same formulation can be extended to dual-arm settings by enlarging the action space. **The *Guide-Think-Act* Paradigm.** Based on this formulation, we decompose policy inference into three stages, as illustrated in Fig. 2:

- **Guide (Sec. 3.2):** We incorporate optional spatial priors P_{spatial} into the observation stream, allowing human users to provide sparse geometric guidance alongside the language instruction.
- **Think (Sec. 3.3):** Instead of directly predicting actions from observations, the VLM generates a structured spatial-visual reasoning sequence C conditioned on the current observation and the optional spatial prior.
- **Act (Sec. 3.4):** A downstream action head consumes the latest latent reasoning states $H_{\text{reasoning}}$ and produces continuous action chunks for control. To support responsive execution, the reasoning module and the action head operate asynchronously at different frequencies.

3.2 The “Guide” Phase: Multimodal Spatial Priors

Spatial Prior Interface. We introduce an optional spatial prior P_{spatial} as an additional input interface for sparse human guidance. The role of P_{spatial} is not to replace the language instruction L , but to provide complementary geometric constraints in image space when the task is ambiguous or when targeted correction is needed. In practice, P_{spatial} is specified on the primary camera view and is consumed jointly with the visual observation and language instruction.

Source and Timing of Spatial Priors. The guidance interface is source-agnostic: spatial priors may come from scripted simulator annotations, external perception models such as Grounding-DINO or Gemini, or direct user intervention. They can also be provided under two timing regimes. An *up-front* prior is attached at the episode start as additional task context, while a *mid-episode* prior is injected when a human user or failure detector observes mis-grounding,

an incorrect affordance, or an undesired motion path. In our batch evaluation of ambiguous scenarios, spatial priors are generated from simulator state for controlled comparison; in interactive settings, the intervention trigger is provided by the human user. Automatic triggering based on confidence, uncertainty, or failure detection is left for future work.

Spatial Formulations. Our framework supports three levels of spatial guidance:

- **Affordance Guide (P_{point}):** A single 2D image coordinate (x, y) indicating a task-relevant affordance location, most commonly a grasp point, contact point, or interaction anchor on the target object. This is the lightest-weight form of intervention and is useful for rapidly specifying where the robot should interact.
- **Box Guide (P_{bbox}):** A bounding box $(x_{\min}, y_{\min}, x_{\max}, y_{\max})$ specifying a target region. This is useful when object identity or coarse localization is the bottleneck, although boxes may still include nearby distractors and therefore provide less precise affordance-level supervision than points.
- **Trace Guide (P_{trace}):** An ordered sequence of 2D points

$$P_{\text{trace}} = [(x_1, y_1), (x_2, y_2), \dots, (x_m, y_m)], \quad (4)$$

representing a coarse image-space path. This form of guidance is useful for expressing directional preferences, motion style, or obstacle-avoidance cues.

Serialization and Tokenization. To integrate P_{spatial} into the VLM backbone, we serialize spatial priors into the model’s coordinate token space and concatenate them with the textual instruction. Qwen3-VL natively supports point- and box-based localization in relative coordinate space. We therefore encode P_{point} and P_{bbox} using the same coordinate representation as the backbone’s native grounding interface. For trace guidance, we represent the path as a short ordered sequence of point coordinates, using the same coordinate tokenization scheme for each waypoint.

Fusion with the Observation Stream. After serialization, the spatial prior is provided together with the language instruction and visual observation, allowing the VLM to jointly attend to semantic content and human-provided geometric cues. This design preserves a unified inference interface: when P_{spatial} is absent, the model operates fully autonomously; when it is present, the same backbone conditions its subsequent reasoning on the provided spatial prior without requiring a separate interaction-specific branch.

3.3 The “Think” Phase: Conditioned Spatial-Visual CoT

Structured Reasoning Sequence. Given the augmented input tuple $(\mathcal{I}_t, L, s_t, P_{\text{spatial}})$, our model does not directly map observations to motor actions. Instead, the VLM first generates a structured spatial-visual reasoning sequence C in an autoregressive manner. For both exposition and supervision, we organize this sequence into three functional segments,

$$C = [C_{\text{task}}, C_{\text{vision}}, C_{\text{robot}}], \quad (5)$$

which correspond to task decomposition, visual grounding, and robot-oriented motion reasoning, respectively.

Reasoning Decomposition.

1. **Task CoT (C_{task}):** The model first produces a high-level semantic rationale that decomposes the instruction L into executable sub-tasks and identifies the relevant objects and interactions required for completion.
2. **Vision CoT (C_{vision}):** Conditioned on the observation and the preceding task rationale, the model predicts visually grounded intermediate targets, including target regions and task-relevant affordance locations in image space. This step anchors the semantic plan to concrete visual entities.
3. **Robot CoT (C_{robot}):** Based on the grounded visual targets, the model further predicts a coarse image-space motion sketch for the end-effector, represented as a sequence of 2D waypoints that summarize the intended manipulation trajectory.

Conditioning on Human Spatial Guidance. The key property of this phase is that the reasoning process is explicitly conditioned on the optional spatial prior:

$$P(C \mid \mathcal{I}_t, L, P_{\text{spatial}}). \quad (6)$$

When P_{spatial} is absent, the model reasons autonomously from the visual observation and language instruction alone. When a spatial prior is provided, it serves as an additional geometric constraint on the reasoning process. In particular, an affordance guide or box guide can anchor the model’s visual grounding to a user-specified interaction point or region, while a trace guide can bias the predicted motion sketch toward a desired path. As a result, the same reasoning backbone supports both fully autonomous execution and interaction-conditioned correction under spatial ambiguity.

Latent Reasoning States. Let $H_{\text{reasoning}}$ denote the hidden states associated with the generated reasoning tokens in C . These states provide a dense latent representation of the model’s task understanding, visual grounding, and motion intent, and are passed to the downstream action head in the subsequent *Act* phase.

3.4 The “Act” Phase: Asynchronous Flow-Matching

Motivation: Decoupling Slow Reasoning from Fast Control. Autoregressive VLM reasoning is significantly slower than the control frequency typically required for closed-loop manipulation. If the model were forced to regenerate the full reasoning sequence C at every control step, action execution would be limited by the decoding speed of the VLM, leading to delayed feedback and unstable behavior in dynamic interaction. To mitigate this mismatch, we separate high-level reasoning from low-level action generation.

Asynchronous Execution Architecture. Our *Act* phase adopts an asynchronous slow-fast design with two coupled modules:

- **Slow Reasoning Module:** The VLM backbone executes the *Guide* and *Think* phases at a lower update frequency. Given the current multimodal input, it produces the structured reasoning sequence C and the corresponding latent reasoning states

$$H_{\text{reasoning}} \in \mathbb{R}^{N \times D}, \quad (7)$$

where N is the number of reasoning tokens and D is the hidden dimension. These states summarize the current task decomposition, visual grounding, and motion intent, and are stored as the latest cached reasoning memory.

- **Fast Action Module:** A downstream Flow-Matching action head operates at a higher control frequency. At each control step, it receives the current observation together with the latest cached reasoning states $H_{\text{reasoning}}^{\text{latest}}$, and predicts continuous action chunks conditioned on this reasoning context. In our implementation, the action head accesses $H_{\text{reasoning}}^{\text{latest}}$ through cross-attention, allowing the control module to reuse the most recent reasoning output between VLM updates.

Flow-Matching Action Generation. We parameterize action generation with a flow-matching policy head [3, 24, 37], which models continuous action chunks by learning a time-dependent vector field over actions. In our asynchronous design, the VLM and the action head consume different input streams at different update frequencies.

At a lower frequency, the VLM takes the primary image, the language instruction, and the optional spatial prior as input, i.e.,

$$(\mathcal{I}_t^{\text{main}}, L, P_{\text{spatial}}) \longrightarrow C, H_{\text{reasoning}}. \quad (8)$$

This stage produces the structured reasoning sequence C and its corresponding latent reasoning states $H_{\text{reasoning}}$.

At a higher control frequency, the action head consumes the current control observation together with the latest cached reasoning states. Specifically, it takes the primary image, the wrist image, the proprioceptive state, and the most recent reasoning memory as input, and predicts

$$v_{\theta}(x, \tau \mid \mathcal{I}_t^{\text{main}}, \mathcal{I}_t^{\text{wrist}}, s_t, H_{\text{reasoning}}^{\text{latest}}), \quad (9)$$

where x denotes the action variable and τ denotes the flow time. Integrating this vector field yields a continuous action chunk conditioned on both the current control observation and the latest available reasoning context.

This separation allows semantic and spatial reasoning to be updated at a lower frequency, while the action head continues to generate responsive actions at a higher frequency using the latest cached reasoning states. In this way, the policy preserves rich reasoning capacity without requiring autoregressive VLM decoding at every control step.

3.5 Data Construction and Training Recipe

To train the framework at scale, we construct **Interact-306K**, a multi-embodiment dataset for guided spatial reasoning. As shown in Fig.3, it is

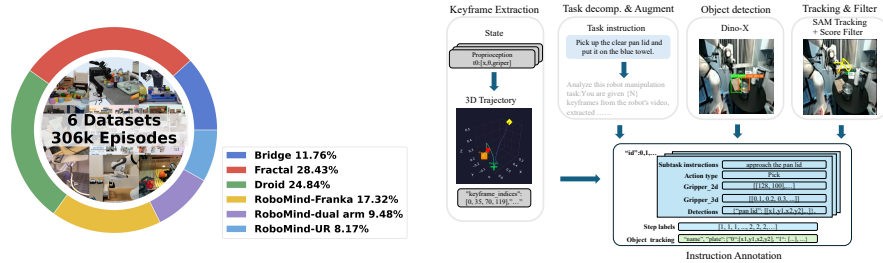


Fig. 3: Interact-306K and automatic instruction annotation. Left: Dataset composition: 306K episodes collected from six manipulation sources (*e.g.*, Bridge [30], Fractal [38], Droid [16], and RoboMind variants [32]). Right: Automatic annotation pipeline: keyframe extraction and task decomposition from trajectories, followed by open-vocabulary grounding and tracking to produce structured subtask instructions with temporally consistent object annotations.

built from approximately 306K real-world manipulation trajectories collected from Open X-Embodiment (OXE) [25], DROID [16], RoboMind [32], and our own data, and augmented with automatically generated spatial-reasoning annotations.

Automated Spatial-CoT Supervision. Since raw robot demonstrations do not contain explicit reasoning traces, we automatically construct supervision for both the *Guide* and *Think* phases. For each trajectory, we generate a structured reasoning target

$$C = [C_{\text{task}}, C_{\text{vision}}, C_{\text{robot}}], \quad (10)$$

aligned with the three-part decomposition used by our model: task decomposition is inferred from keyframes and language, visual grounding is obtained by localizing and tracking task-relevant objects in image space, and robot-centric supervision is derived by projecting end-effector motion into the primary view to produce affordance locations and coarse 2D motion sketches. To better match inference-time interaction, we further perturb the generated spatial annotations with stochastic noise, producing synthetic affordance points, boxes, and traces for training the *Guide* interface.

Training Recipe. We train the model in two stages. In Stage 1, we train the VLM backbone on Interact-306K to learn the *Guide* and *Think* components, using stochastic spatial conditioning so that the model is exposed to both guided and unguided inputs. We then train the Flow-Matching action head to map the latent reasoning states $H_{\text{reasoning}}$ together with control observations to continuous action chunks. In Stage 2, we jointly fine-tune the full policy on domain-specific robot data (*e.g.*, BridgeData V2 [30]) to adapt the reasoning module and action head to the target embodiment and environment. Unless otherwise specified, reasoning generation is optimized with autoregressive token prediction on C , while the action module is optimized with the standard flow-matching objective on action chunks.

Table 1: Main Results. Success rates (%) on the LIBERO and SimplerEnv (Bridge) benchmarks. * denotes reproduced results evaluated with a maximum inference horizon of 120 steps, consistent with the setting used for other models.

Method	LIBERO					SIMPLER-Env (Bridge)				
	Spatial	Object	Goal	Long	Avg	Spoon	Carrot	Cube	Eggplant	Avg
OpenVLA [18]	84.7	88.4	79.2	53.7	76.5	4.2	0.0	8.3	45.8	14.6
OpenVLA-OFT [17]	96.2	98.3	96.2	90.7	95.3	-	-	-	-	-
π_0 [3]	96.8	98.8	95.8	85.2	94.1	50.0	41.7	29.2	70.8	47.9
GR00T-N1 [24]	94.4	97.6	93.0	90.6	93.9	64.5	65.5	5.5	93.0	57.1
$\pi_{0.5}$ [4]	98.8	98.2	98.0	92.4	96.9	-	-	-	-	-
X-VLA* [37]	98.2	98.6	97.8	97.6	98.1	95.8	75.0	62.5	70.8	76.0
ThinkAct [13]	88.3	91.4	87.1	70.9	84.4	37.5	8.7	58.3	70.8	43.8
CoT-VLA [36]	87.5	91.6	87.6	69.0	81.1	-	-	-	-	-
Uni-VLA [31]	97.0	99.0	92.6	90.8	94.8	83.3	66.7	33.3	95.8	69.8
MolmoAct [20]	87.0	95.4	87.6	77.2	86.6	-	-	-	-	-
GTA-VLA	99.0	98.8	98.4	97.6	98.6	95.8	87.5	66.7	75.0	81.2

4 Experiments

We evaluate our method along three main axes: standard benchmark performance, out-of-distribution (OOD) robustness, and the effectiveness of explicit visual guidance under spatial ambiguity. We first assess autonomous performance on established manipulation benchmarks, including LIBERO [23] and SimplerEnv [22]. We then evaluate OOD generalization using our proposed SimplerEnv-Plus benchmark, which introduces systematic perturbations across visual, robot, language, and object-centric factors. Finally, we study whether sparse visual guidance, such as affordance points and boxes, can effectively resolve ambiguity when language alone is insufficient.

4.1 Experimental Setup

We mainly evaluate our method on two simulation benchmarks: sim-to-sim Libero [23] and real-to-sim: SimplerEnv [22].

Standard Benchmark for In-Domain Evaluation. We primarily evaluate our method on two simulation benchmarks: LIBERO [23] and SimplerEnv [22]. LIBERO is a sim-to-sim benchmark covering diverse multi-task manipulation suites, including Spatial, Object, Goal, and Long, and is commonly used to evaluate multi-task generalization, instruction following, and long-horizon reasoning. SimplerEnv is a real-to-sim benchmark built on high-fidelity digital twins of real robot setups, and is designed to assess zero-shot visuomotor transfer and manipulation robustness under realistic visual conditions. In the main paper, we focus on the WidowX domain of SimplerEnv and defer additional results on Google Robot to the appendix.

SimplerEnv-Plus for OOD Evaluation. To evaluate robustness under systematic shift, we introduce **SimplerEnv-Plus**, an extended benchmark built on top of SimplerEnv. It includes four categories of perturbations:

- *Visual Shift*: We perturb low-level visual conditions through lighting variation and sensor viewpoint changes to test robustness to appearance changes.
- *Robot State Shift*: We randomize the robot’s initial state and introduce execution noise to simulate uncertainty in embodiment state and control.
- *Language Shift*: We modify task instructions through lexical variation in verbs, nouns, and attributes to evaluate robustness to instruction diversity.
- *Object Shift*: We replace standard targets with novel objects and introduce distractors to test zero-shot object generalization and robustness under perceptual ambiguity.

Visual Guidance Evaluation. To evaluate the effectiveness of explicit visual guidance under ambiguity, we consider two challenging settings:

- *Unseen Object Ambiguity*: We replace standard training objects with novel instances from multiple categories, including *Unseen Toy*, *Unseen Fruit*, and *Unseen Tool*. We then compare unguided execution against point- and box-guided execution to measure whether spatial priors improve zero-shot object grounding in unseen scenarios.
- *Distractor-based Ambiguity*: We introduce same-category distractors that create fine-grained spatial ambiguity, including *Color Distractors* (same category, different colors) and *Position Distractors* (same object type at different locations). We compare unguided execution against point- and box-guided execution to evaluate whether visual guidance can resolve ambiguity when language alone is insufficient to uniquely specify the target.

4.2 Main Results: Standard and OOD Performance

In-Domain Performance. As shown in Table 1, our method achieves strong in-domain performance on both LIBERO and SimplerEnv. On the highly competitive LIBERO benchmark, our approach reaches an average success rate of 98.6%, performing on par with or slightly above the strongest baselines. This indicates that introducing explicit spatial reasoning does not compromise the policy’s core manipulation ability or multi-task execution performance.

On the real-to-sim SimplerEnv benchmark, our method achieves an average success rate of 81.2%, outperforming the strongest reported baseline. This gain is more pronounced than on LIBERO, suggesting that explicit spatial reasoning is particularly beneficial when policies must bridge the visual and semantic gap between open-world training data and simulated robot execution. By explicitly grounding task-relevant regions and affordances before action generation, the policy is better able to align semantic understanding with executable control.

Out-of-Distribution Generalization. Table 2 shows that our method also improves robustness under systematic distribution shifts in SimplerEnv-Plus. Compared with baseline methods, our approach consistently maintains stronger performance across visual, robot-state, language, and object-centric perturbations.

Under *Visual Shift*, our model remains more robust to sensor viewpoint changes and lighting variation, indicating that explicit intermediate grounding

Table 2: OOD Generalization on SimplerEnv-Plus. Success rates (%) under distribution shifts across visual, robot-state, language, and object-centric factors.

Method	Visual		Robot	Language	Objects		Avg
	Sensor	Lighting	State	Diversity	Unseen Obj.	Distractor	
OpenVLA [18]	5.2	6.3	0.0	8.3	2.1	0.0	3.7
$\pi_{0.5}$ [4]	9.4	10.4	9.4	8.3	6.3	0.0	7.3
X-VLA [37]	27.1	68.8	68.7	66.3	36.2	46.9	52.3
GTA-VLA	39.6	76.1	79.2	68.1	58.3	50.0	61.4

reduces reliance on spurious low-level correlations. Under *Robot State Shift*, our method maintains strong performance despite randomized initial states and execution noise, suggesting that the combination of explicit reasoning and stable action generation improves robustness to embodiment uncertainty. Under *Language Shift*, our model preserves competitive performance under lexical variation, showing that introducing spatial supervision does not substantially weaken language understanding. Finally, under *Object Shift*, which includes both unseen objects and distractor-heavy scenes, our method shows the largest relative advantage, indicating that explicit task decomposition and grounded intermediate reasoning are especially helpful when target identification becomes ambiguous.

Overall, these results suggest that the proposed *Think* and *Act* design improves both in-domain execution and robustness under distribution shift, while preserving the ability to operate autonomously without external guidance.

4.3 Guidance Efficacy

Table 3 shows that explicit visual guidance consistently improves performance under both unseen-object ambiguity and distractor-based ambiguity. When language alone is insufficient to uniquely specify the correct target, sparse spatial priors provide a much stronger grounding signal than dense instruction. The two guidance forms show complementary behavior. Point guidance provides precise affordance-level supervision and is especially effective when same-category distractors create fine-grained spatial ambiguity. Box guidance provides stronger object-level localization and is more helpful when object identity is the main bottleneck, as in unseen-object settings. Overall, these results show that visual guidance is an effective mechanism for resolving ambiguity while allowing different spatial priors to target different failure modes.

Trace guidance. Trace guidance targets motion preference, such as path shape and obstacle avoidance, rather than target selection. We therefore isolate it on obstacle-avoidance and constrained-path tasks, where the trace is constructed by linearly interpolating waypoints from the source gripper pose to the target box center. On this setting, trace guidance improves success from 27.8% to 30.5%, indicating that path-level priors can provide additional control beyond point- or box-level target disambiguation.

Table 3: Effectiveness of Visual Guidance in Ambiguous Scenarios. We evaluate different input modalities for our model on challenging SimplerEnv-Plus tasks. All values are success rates (%). Visual guidance significantly outperforms even dense linguistic instructions, especially when spatial ambiguity is high.

Guidance Modality	Unseen Object Ambiguity				Distractor-based Ambiguity		
	Unseen Toy	Unseen Fruit	Unseen Tool	Avg.	Color Distractor	Pos. Distractor	Avg.
Dense Instruction ($\pi_{0.5}$)	8.3	12.5	8.3	9.7	8.3	8.3	8.3
Dense Instruction(GTA-VLA)	12.5	41.6	29.2	27.8	45.8	29.2	37.5
+ Visual Point Guide	33.3	47.9	41.6	40.9	58.3	50.0	54.2
+ Visual Box Guide	54.1	70.8	45.8	56.9	41.6	45.8	43.7

Table 4: Ablation of structured CoT fields and free-form CoT on SimplerEnv-Bridge. All values are success rates (%).

Setting	Spoon	Carrot	Cube	Eggplant	Avg	Δ
GTA-VLA (full)	95.8	87.5	66.7	75.0	81.2	–
– C_{robot}	95.8	87.5	54.2	75.0	78.1	–3.1
– C_{task}	95.8	91.7	54.2	33.3	68.8	–12.4
– C_{vision}	95.8	79.2	41.7	50.0	66.7	–14.5
Free-form CoT	100.0	91.7	50.0	20.8	65.6	–15.6

4.4 Ablation Study

Table 4 ablates the three fields in the structured Chain-of-Thought on SimplerEnv-Bridge. Removing C_{vision} causes the largest drop, showing that explicit visual grounding is central for cluttered scenes where target localization is the main bottleneck. Removing C_{task} also substantially hurts performance, confirming that semantic decomposition is important for mapping instructions to objects and interactions. Removing C_{robot} has a smaller but still measurable effect, suggesting that the action head can recover part of the low-level motion intent while still benefiting from explicit robot-oriented motion sketches.

The free-form CoT variant underperforms the structured format under the same pseudo-label supervision. This suggests that, for the current scenarios, explicitly structured fields provide a more controllable and stable way to align task decomposition, visual grounding, and motion intent with user-provided priors. More discussion of this trade-off is provided in the appendix.

4.5 Real-World Robot Deployment

The real-world setup is shown in Figure 4, with additional implementation details provided in the appendix. We evaluate the model on four real-world picking tasks defined along two axes: whether the target object is seen or unseen, and whether the scene contains a single target or multiple same-category candidates requiring disambiguation. Concretely, the tasks include: (1) a single seen target in clutter, (2) a single unseen target in clutter, (3) a referred target among

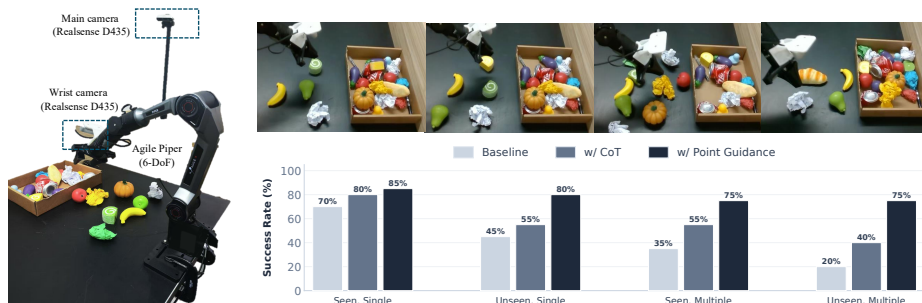


Fig. 4: Real-world robot deployment. Left: the experimental setup with the Agile Piper robot, a primary, and a wrist-mounted camera. Right: qualitative examples and success rates across four picking tasks (seen/unseen targets and single/multiple candidate objects). Explicit reasoning improves over the baseline, and point guidance provides the largest gains in unseen and reference-ambiguous settings.

multiple identical seen objects, and (4) a referred target among multiple identical unseen objects. As shown in Figure 4, our method succeeds not only in cluttered seen-object settings, but also in more challenging unseen-object and referring scenarios, indicating that the proposed guidance and reasoning mechanism transfers effectively to real-world spatial ambiguity.

5 Conclusion

We presented **GTA-VLA (Guide, Think, Act)**, an interactive Vision-Language-Action framework that enables spatially steerable embodied reasoning through human visual guidance. By allowing spatial priors, such as affordance points, boxes, and traces, to directly condition a unified spatial-visual Chain-of-Thought, GTA-VLA moves beyond passive “Sense-to-Act” policies toward robot policies that are both autonomous and naturally correctable when failures or ambiguities arise. Experiments in both simulation and the real world show that our approach achieves strong autonomous performance while substantially improving failure recovery under out-of-domain shifts and spatial ambiguity. A current limitation of our framework is that both the reasoning process and the guidance interface are primarily formulated in 2D image space. An important direction for future work is to extend both the Chain-of-Thought representation and the visual guidance cues into 3D, enabling richer geometric grounding and more general interaction in real-world embodied settings.

Acknowledgements

This work was partially supported by National Natural Science Foundation of China (Grant No. 62350710797).

References

1. Bai, S., Cai, Y., Chen, R., Chen, K., Chen, X., Cheng, Z., Deng, L., Ding, W., Gao, C., Ge, C., Ge, W., Guo, Z., Huang, Q., Huang, J., Huang, F., Hui, B., Jiang, S., Li, Z., Li, M., Li, M., Li, K., Lin, Z., Lin, J., Liu, X., Liu, J., Liu, C., Liu, Y., Liu, D., Liu, S., Lu, D., Luo, R., Lv, C., Men, R., Meng, L., Ren, X., Ren, X., Song, S., Sun, Y., Tang, J., Tu, J., Wan, J., Wang, P., Wang, P., Wang, Q., Wang, Y., Xie, T., Xu, Y., Xu, H., Xu, J., Yang, Z., Yang, M., Yang, J., Yang, A., Yu, B., Zhang, F., Zhang, H., Zhang, X., Zheng, B., Zhong, H., Zhou, J., Zhou, F., Zhou, J., Zhu, Y., Zhu, K.: Qwen3-vl technical report. arXiv preprint arXiv:2511.21631 (2025)
2. Belkhale, S., Ding, T., Xiao, T., Sermanet, P., Vuong, Q., Tompson, J., Chebotar, Y., Dwibedi, D., Sadigh, D.: Rt-h: Action hierarchies using language. arXiv preprint arXiv:2403.01823 (2024)
3. Black, K., Brown, N., Driess, D., Esmail, A., Equi, M., Finn, C., Fusai, N., Groom, L., Hausman, K., Ichter, B., et al.: π_0 : A vision-language-action flow model for general robot control. arXiv preprint arXiv:2410.24164 (2024)
4. Black, K., Brown, N., Driess, D., Esmail, A., Equi, M., Finn, C., Fusai, N., et al.: $\pi_{0.5}$: A vision-language-action model with open-world generalization. arXiv preprint arXiv:2504.16054 (2025)
5. Carion, N., Gustafson, L., Hu, Y.T., Debnath, S., Hu, R., Suris, D., Ryali, C., Alwala, K.V., Khedr, H., Huang, A., Lei, J., Ma, T., Guo, B., Kalla, A., Marks, M., Greer, J., Wang, M., Sun, P., Rädle, R., Afouras, T., Mavroudi, E., Xu, K., Wu, T.H., Zhou, Y., Momeni, L., Hazra, R., Ding, S., Vaze, S., Porcher, F., Li, F., Li, S., Kamath, A., Cheng, H.K., Dollár, P., Ravi, N., Saenko, K., Zhang, P., Feichtenhofer, C.: Sam 3: Segment anything with concepts. arXiv preprint arXiv:2511.16719 (2025)
6. Cheang, C., Chen, S., Cui, Z., Hu, Y., Huang, L., Kong, T., Li, H., Li, Y., Liu, Y., Ma, X., Niu, H., Ou, W., Peng, W., Ren, Z., Shi, H., Tian, J., Wu, H., Xiao, X., Xiao, Y., Xu, J., Yang, Y.: Gr-3 technical report. arXiv preprint arXiv:2507.15493 (2025)
7. Chen, K., Zhang, Z., Zeng, W., Zhang, R., Zhu, F., Zhao, R.: Shikra: Unleashing multimodal llm’s referential dialogue magic. arXiv preprint arXiv:2306.15195 (2023)
8. Chen, X., Chen, Y., Fu, Y., Gao, N., Jia, J., Jin, W., Li, H., Mu, Y., Pang, J., Qiao, Y., Tian, Y., Wang, B., Wang, B., Wang, F., Wang, H., Wang, T., Wang, Z., Wei, X., Wu, C., Yang, S., Ye, J., Yu, J., Zeng, J., Zhang, J., Zhang, J., Zhang, S., Zheng, F., Zhou, B., Zhu, Y.: Internvla-m1: A spatially guided vision-language-action framework for generalist robot policy. arXiv preprint arXiv:2510.13778 (2025)
9. Chi, C., Xu, Z., Feng, S., Cousineau, E., Du, Y., Burchfiel, B., Tedrake, R., Song, S.: Diffusion policy: Visuomotor policy learning via action diffusion. *The International Journal of Robotics Research* **44**(10-11), 1684–1704 (2025)
10. Deng, S., Yan, M., Wei, S., Ma, H., Yang, Y., Chen, J., Zhang, Z., Yang, T., Zhang, X., Cui, H., et al.: Graspvla: A grasping foundation model pre-trained on billion-scale synthetic action data. In: CoRL (2025)
11. Gu, J., Kirmani, S., Wohlhart, P., Lu, Y., Arenas, M.G., Rao, K., Yu, W., Fu, C., Gopalakrishnan, K., Xu, Z., et al.: Rt-trajectory: Robotic task generalization via hindsight trajectory sketches. In: ICLR (2023)
12. Huang, C.P., Man, Y., Yu, Z., Chen, M.H., Kautz, J., Wang, Y.C.F., Yang, F.E.: Fast-thinkact: Efficient vision-language-action reasoning via verbalizable latent planning. arXiv preprint arXiv:2601.09708 (2026)

13. Huang, C.P., Wu, Y.H., Chen, M.H., Wang, Y.C.F., Yang, F.E.: Thinkact: Vision-language-action reasoning via reinforced visual latent planning. In: NeurIPS (2025)
14. Jiang, Q., Huo, J., Chen, X., Xiong, Y., Zeng, Z., Chen, Y., Ren, T., Yu, J., Zhang, L.: Detect anything via next point prediction. arXiv preprint arXiv:2510.12798 (2025)
15. Jiang, Q., Li, F., Zeng, Z., Ren, T., Liu, S., Zhang, L.: T-rex2: Towards generic object detection via text-visual prompt synergy. In: ECCV (2024)
16. Khazatsky, A., Pertsch, K., Nair, S., Balakrishna, A., Dasari, S., Karamcheti, S., Nasiriany, S., Srirama, M.K., Chen, L.Y., Ellis, K., et al.: Droid: A large-scale in-the-wild robot manipulation dataset. In: Robotics: Science and Systems (2024)
17. Kim, M.J., Finn, C., Liang, P.: Fine-tuning vision-language-action models: Optimizing speed and success. arXiv preprint arXiv:2502.19645 (2025)
18. Kim, M.J., Pertsch, K., Karamcheti, S., Xiao, T., Balakrishna, A., Nair, S., Rafailov, R., Foster, E.P., Sanketi, P.R., Vuong, Q., et al.: Openvla: An open-source vision-language-action model. In: CoRL (2025)
19. Kirillov, A., Mintun, E., Ravi, N., Mao, H., Rolland, C., Gustafson, L., Xiao, T., Whitehead, S., Berg, A.C., Lo, W.Y., et al.: Segment anything. In: CVPR (2023)
20. Lee, J., Duan, J., Fang, H., Deng, Y., Liu, S., Li, B., Fang, B., Zhang, J., Wang, Y.R., Lee, S., Han, W., Pumacay, W., Wu, A., Hendrix, R., Farley, K., Vanderbilt, E., Farhadi, A., Fox, D., Krishna, R.: Molmoact: Action reasoning models that can reason in space. arXiv preprint arXiv:2508.07917 (2025)
21. Li, Q., Liang, Y., Wang, Z., Luo, L., Chen, X., Liao, M., Wei, F., Deng, Y., Xu, S., Zhang, Y., Wang, X., Liu, B., Fu, J., Bao, J., Chen, D., Shi, Y., Yang, J., Guo, B.: Cogact: A foundational vision-language-action model for synergizing cognition and action in robotic manipulation. arXiv preprint arXiv:2411.19650 (2024)
22. Li, X., Hsu, K., Gu, J., Pertsch, K., Mees, O., Walke, H.R., Fu, C., Lunawat, I., Sieh, I., Kirmani, S., Levine, S., Wu, J., Finn, C., Su, H., Vuong, Q., Xiao, T.: Evaluating real-world robot manipulation policies in simulation. arXiv preprint arXiv:2405.05941 (2024)
23. Liu, B., Zhu, Y., Gao, C., Feng, Y., Liu, Q., Zhu, Y., Stone, P.: Libero: Benchmarking knowledge transfer for lifelong robot learning. NeurIPS **36**, 44776–44791 (2023)
24. 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. arXiv preprint arXiv:2503.14734 (2025)
25. O'Neill, A., Rehman, A., Maddukuri, A., Gupta, A., Padalkar, A., Lee, A., Pooley, A., Gupta, A., Mandlekar, A., Jain, A., et al.: Open x-embodiment: Robotic learning datasets and rt-x models: Open x-embodiment collaboration 0. In: IEEE International Conference on Robotics and Automation (2024)
26. Ravi, N., Gabeur, V., Hu, Y.T., Hu, R., Ryali, C., Ma, T., Khedr, H., Rädle, R., Rolland, C., Gustafson, L., Mintun, E., Pan, J., Alwala, K.V., Carion, N., Wu, C.Y., Girshick, R., Dollár, P., Feichtenhofer, C.: Sam 2: Segment anything in images and videos. arXiv preprint arXiv:2408.00714 (2024)
27. Tang, P., Xie, S., Sun, B., Huang, B., Luo, K., Yang, H., Jin, W., Wang, J.: Mind to hand: Purposeful robotic control via embodied reasoning. arXiv preprint arXiv:2512.08580 (2025)

28. Team, B.S.: Seed1.5-vl technical report. arXiv preprint arXiv:2505.07062 (2025)
29. Team, O.M., Ghosh, D., Walke, H., Pertsch, K., Black, K., Mees, O., Dasari, S., Hejna, J., Kreiman, T., Xu, C., Luo, J., Tan, Y.L., Chen, L.Y., Sanketi, P., Vuong, Q., Xiao, T., Sadigh, D., Finn, C., Levine, S.: Octo: An open-source generalist robot policy. arXiv preprint arXiv:2405.12213 (2024)
30. Walke, H., Black, K., Lee, A., Kim, M.J., Du, M., Zheng, C., Zhao, T., Hansen-Estruch, P., Vuong, Q., He, A., Myers, V., Fang, K., Finn, C., Levine, S.: Bridge-data v2: A dataset for robot learning at scale. In: CoRL (2023)
31. Wang, Y., Li, X., Wang, W., Zhang, J., Li, Y., Chen, Y., Wang, X., Zhang, Z.: Unified vision-language-action model. arXiv preprint arXiv:2506.19850 (2025)
32. Wu, K., Hou, C., Liu, J., Che, Z., Ju, X., Yang, Z., Li, M., Zhao, Y., Xu, Z., Yang, G., et al.: Robomind: Benchmark on multi-embodiment intelligence normative data for robot manipulation. In: Robotics: Science and Systems (2025)
33. Yang, J., Zhang, H., Li, F., Zou, X., yue Li, C., Gao, J.: Set-of-mark prompting unleashes extraordinary visual grounding in gpt-4v. arXiv preprint arXiv:2310.11441 (2023)
34. You, H., Zhang, H., Gan, Z., Du, X., Zhang, B., Wang, Z., Cao, L., Chang, S.F., Yang, Y.: Ferret: Refer and ground anything anywhere at any granularity. In: ICLR (2023)
35. Zawalski, M., Chen, W., Pertsch, K., Mees, O., Finn, C., Levine, S.: Robotic control via embodied chain-of-thought reasoning. In: CoRL. pp. 3157–3181. PMLR (2025)
36. Zhao, Q., Lu, Y., Kim, M.J., Fu, Z., Zhang, Z., Wu, Y., Li, Z., Ma, Q., Han, S., Finn, C., et al.: Cot-vla: Visual chain-of-thought reasoning for vision-language-action models. In: CVPR. pp. 1702–1713 (2025)
37. Zheng, J., Li, J., Wang, Z., Liu, D., Kang, X., Feng, Y., Zheng, Y., Zou, J., Chen, Y., Zeng, J., Zhang, Y.Q., Pang, J., Liu, J., Wang, T., Zhan, X.: X-vla: Soft-prompted transformer as scalable cross-embodiment vision-language-action model. arXiv preprint arXiv:2510.10274 (2025)
38. Zitkovich, B., Yu, T., Xu, S., Xu, P., Xiao, T., Xia, F., Wu, J., Wohlhart, P., Welker, S., Wahid, A., et al.: Rt-2: Vision-language-action models transfer web knowledge to robotic control. In: CoRL (2023)