

RH20T-P: A Primitive-Level Robotic Dataset Towards Composable Generalization Agents

Zeren Chen^{1,2*}, Zhelun Shi^{1,2*}, Xiaoya Lu^{1,5*}, Lehan He^{1,6*},
 Sucheng Qian^{1,3}, Hao Shu Fang³, Zhenfei Yin^{1,4†}, Wanli Ouyang^{1,4},
 Jing Shao^{1‡}, Yu Qiao¹, Cewu Lu^{3‡}, Lu Sheng^{2‡}
¹ Shanghai AI Laboratory, ² School of Software, Beihang University,
³ Shanghai Jiao Tong University, ⁴ University of Sydney,
⁵ University of Electronic Science and Technology of China,
⁶ Nanjing University of Posts and Telecommunications
 {czt1604, shizhelun, lsheng}@buaa.edu.cn,
 {yinzhenfei, shaojing}@pjlab.org.cn

Abstract: The ultimate goals of robotic learning is to acquire a comprehensive and generalizable robotic system capable of performing both seen skills within the training distribution and unseen skills in novel environments. Recent progress in utilizing language models as high-level planners has demonstrated that the complexity of tasks can be reduced through decomposing them into primitive-level plans, making it possible to generalize on novel robotic tasks in a composable manner. Despite the promising future, the community is not yet adequately prepared for composable generalization agents, particularly due to the lack of primitive-level real-world robotic datasets. In this paper, we propose a primitive-level robotic dataset, namely **RH20T-P**, which contains about 33000 video clips covering 44 diverse and complicated robotic tasks. Each clip is manually annotated according to a set of meticulously designed primitive skills, facilitating the future development of composable generalization agents. To validate the effectiveness of RH20T-P, we also construct a potential and scalable agent based on RH20T-P, called **RA-P**. Equipped with two planners specialized in task decomposition and motion planning, RA-P can adapt to novel physical skills through composable generalization. Our website and videos can be found at <https://sites.google.com/view/rh20t-primitive/main>. Dataset and code will be made available soon.

Keywords: Robots, Vision Language Models, Primitive Skills

1 Introduction

Robotic manipulation tasks are designed to enable robotic systems to comprehend environmental observations and open-world task instructions like language, guiding their actuators to execute specific actions in real-world applications. Previous attempts on imitation learning [1, 2, 3, 4] or reinforcement learning [5, 6, 7, 8] in robotic tasks can specialize in tasks that have been seen during training. However, accumulating training data for all the tasks in the real world is impossible, therefore crafting intelligent robotic agents is insufficient with these imitation- or feedback-based learning methods.

Recent advancements in large language models (LLMs) [9, 10, 11, 12, 13] have shown impressive potential in understanding instructions and interpreting contextual cues. By further incorpo-

*Equal contribution.

†Project leader.

‡Corresponding author.

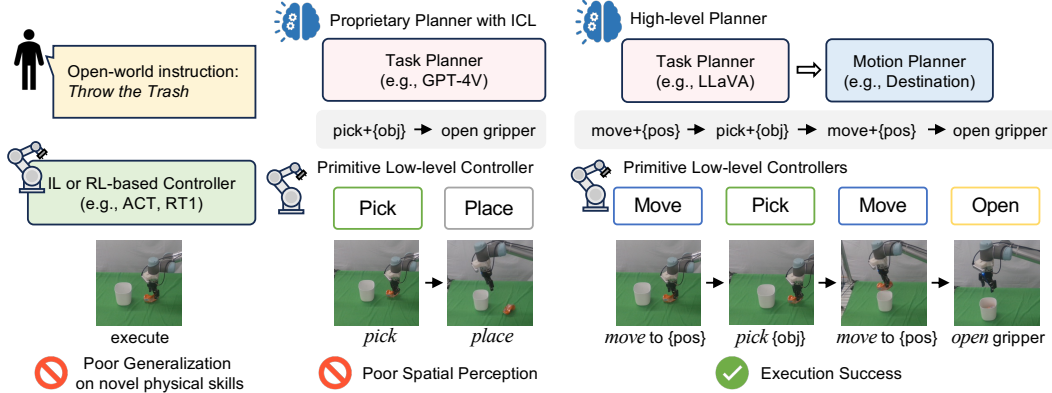


Figure 1: **Comparison on imitation-/feedback-based methods and CGAs.** Note that the low-level controllers with the same border color in the diagram represent the same controller.

rating visual knowledge and semantics through instruction tuning [14], vision language models (VLMs) [15, 16, 17, 18, 19] become well-suited in robotic tasks [20, 21, 22, 23], aiming to address the limitations of conventional algorithms. Some approaches [20] adopt an end-to-end training pipeline, allowing the language models to align with the distribution of action sequences in robotic tasks. While these models can generalize to some extent for tasks involving novel objects or compositions of known physical skills and objects in unfamiliar environments, their physical skills are restricted to those encountered in the training datasets. Alternatively, other methods [21, 22, 23] apply a plan-execute paradigm, employing VLMs as high-level planners to decompose tasks and re-plan them with more fine-grained primitive skills, *e.g.*, moving or picking. These step-by-step decompositions are then executed by low-level controllers without grasping the semantics of the tasks. We formulate these methods as **composable generalization agents (CGAs)** in the realm of robotic learning. By leveraging composable generalization, VLMs can endow CGAs with the ability to break down novel physical skills like “wipe”, “throw” and “receive” into constituent primitive skills, mitigating the unpredictability and intricacy when expanding the abilities of agents to novel physical skills.

Despite the huge promise of CGAs in handling novel skills, they face several challenges. Existing CGAs [22, 23, 24] tend to use larger-scale proprietary models like GPT-4V [15] to decompose the task specifications into primitive skills through in-context learning [9, 25]. The dependence on proprietary VLMs as decision-making backends results in a lack of transparency and flexibility. Furthermore, the extensive requirements for spatial perception in robotic tasks, particularly the ability to locate specific background positions, cannot align with the foreground-awareness priors in traditional localization tasks. For example, in the task of retrieving an object from a container, it is crucial to locate a background position outside the container. To accumulate the spatial knowledge for each movement of the robot arm on background positions, a robotic datasets where tasks are segmented based on primitive skills is essential. Without such a dataset, CGAs cannot provide precise spatial information for each movement, and they are compelled to delegate the spatial perception to low-level controllers. This may ultimately increase difficulties for those controllers, deviating from the original intention of CGAs.

To fundamentally address this problem, we propose **RH20T-P** dataset, a **primitive-level** robotic dataset built upon **RH20T**, as depicted in fig. 3. Considering that the robot manipulation process can be regarded as a sequence of the robot arm’s movements and changes in gripper, we rigorously define a set of composable and scalable primitive skills based on this principle. Each robotic manipulation video in RH20T-P is manually annotated accordingly (approximately 33000 video clips). Additionally, to assess the effectiveness of the RH20T-P dataset, we design a potential and scalable CGA, called **RA-P (RobotAgent-Primitive)**. In RA-P, we utilize an open-source VLM, LLaVA [16], for task decomposition, along with a pre-trained Deformable DETR [26] to assist VLM in predicting

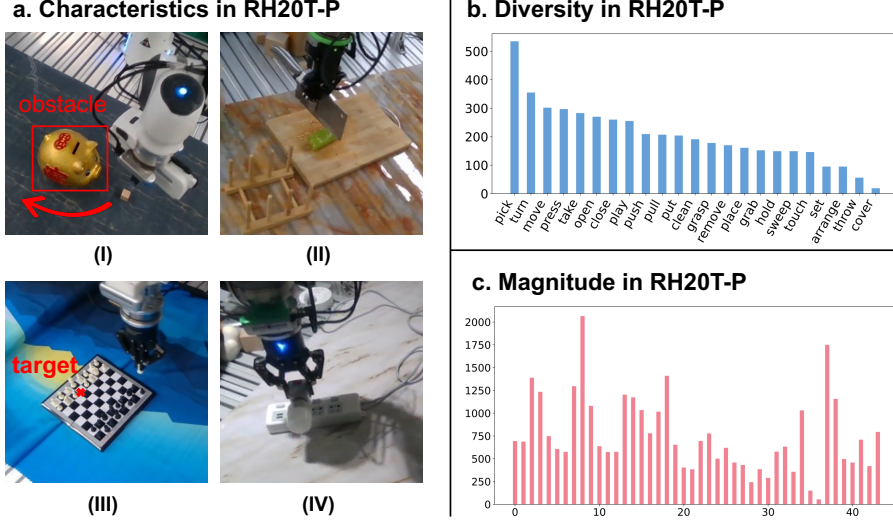


Figure 2: (a) **Characteristics in RH20T-P**: I) **Special trajectories**: pushing the block to go around the obstacle. II) **Using tools**: cutting the vegetables with a knife. III) **Complex visual reasoning**: placing a piece on chessboard to complete the setup. IV) **Long-horizon planning**: lighting up the lamp by plugging in the power cord and then turning on the socket. (b) **Diversity in RH20T-P**: the distribution of physical skills. (c) **Magnitude in RH20T-P**: the total number of annotated clips in each task.

the spatial information for each movement of the robot arm, which we refer to as motion planning. After fine-tuning on RH20T-P, RA-P demonstrates robust generalization, even on novel physical skills, adeptly decomposing tasks using the predefined primitive skills and providing helpful spatial information. We believe that the RH20T-P dataset, with its carefully designed primitive skills and diverse spatial information, will pave the way for the development of more potent CGAs in the future.

Our contributions can be summarized into three aspects:

- We propose a primitive-level robotic dataset RH20T-P for high-level decision making, breaking through the bottlenecks in current CGAs.
- We design a potential and scalable robotic agent RA-P, with two planners dedicated to task decomposition and motion planning.
- Through composable generalization, RA-P can generalize to novel physical skills, demonstrating the effectiveness of RA-P built on RH20T-P.

2 RH20T-P Dataset

We propose the first primitive-level robotic dataset, called **RH20T-P**, towards the composable generalization agents. In section 2.1, we first introduce the RH20T dataset, which served as our data source, and the sampling method used to construct RH20T-P dataset. Then, we elaborate on the principle of primitive skills we designed in RH20T-P in section 2.2. And the detailed hindsight annotation pipeline is described in section 2.3.

2.1 Preliminary of RH20T and Data Sampling

RH20T [27] encompasses a broad spectrum of real-world robot manipulation skills, with each episode featuring diverse contexts, robots, camera viewpoints, and language descriptions of the tasks. Although most tasks in RH20T are highly advanced for current robotics learning in the real world, its diversity and executability are pivot components for the development of future in-

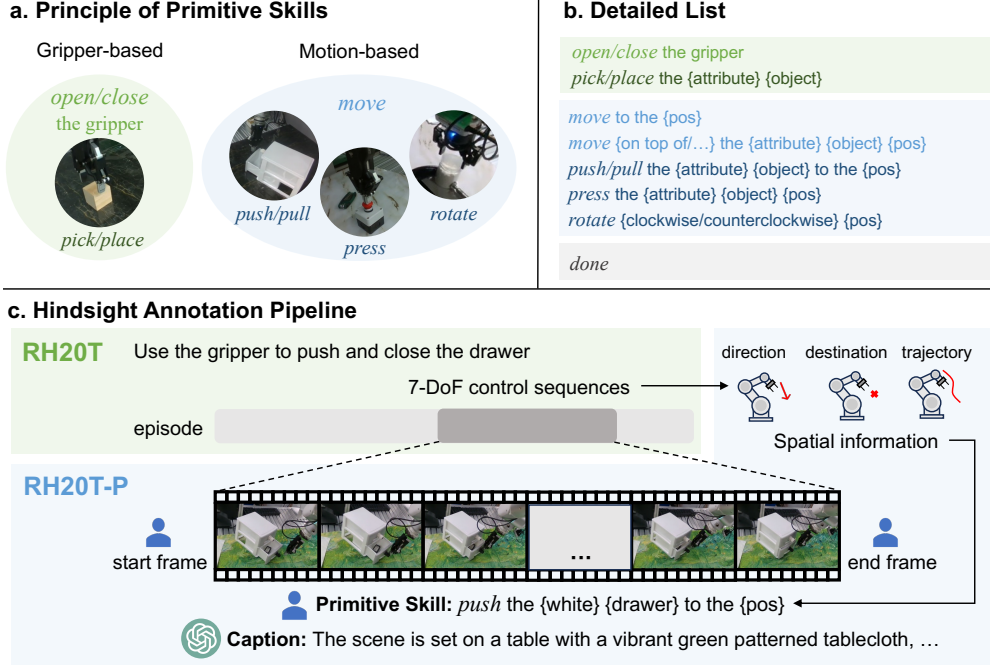


Figure 3: **Overview of RH20T-P dataset.** (a) The design principle of primitive skills in RH20T-P. (b) Detailed list of primitive skills in RH20T-P. (c) Process of hindsight primitive-level annotation.

telligent robotic agents. This motivates us to conduct primitive-level decompositions for tasks in RH20T. Specifically, we sample a subset of tasks in RH20T that are suitable for CGAs to construct a primitive-level robotic dataset RH20T-P and define a set of composable and scalable primitive skills within it. As shown in fig. 2 (a), these tasks can be characterized across 4 dimensions:

- **Special trajectories.** Executing the tasks along a specific trajectory.
- **Using tools.** Interacting with the target objects using tools and predicting action sequences from the tools’ perspective.
- **Complex visual reasoning.** Perceiving the task-related visual semantics within the image and completing the tasks through reasoning about them.
- **Long-horizon planning.** Anticipating several steps or stages into the future, accounting for numerous variables and potential changes over time.

The dataset still retains numerous complex skills beyond pick-and-place actions after sampling (see fig. 2 (b)), such as wiping a table with a cloth or arranging pieces on a chessboard to complete the initial setup.

2.2 Composable and Scalable Primitive Skills in RH20T-P

Primitive skills are crucial for RH20T-P dataset, as the way we decompose tasks and executions shapes the design of the CGAs paradigm. Applying free-form natural language [28, 29, 30] as primitive skills are overly coarse-grained, increasing difficulty for controllers to grasp task semantics, deviating from the original intention of CGAs. Formally, we define primitive skills from the perspective of the robot arms, focusing on the state changes that primarily occur in the robot arm’s motion and gripper during the manipulation process. As shown in fig. 3 (a), we can divide the primitive skills into two major high-level skills: **Motion-based** skills and **Gripper-based** skills.

Motion-based Skills. Motion-based skills are designed to illustrate a specific movement of the robot arm. To clearly characterize their behaviors, we endow them with specific spatial information. By

transforming the teleoperation records related to the robot arm’s motion in RH20T, we can achieve multiple forms of motions, such as reaching a destination or following a trajectory. In comparison, primitive skills defined in VILA [22] and SayCan [24] lack precise motion planning in their skills, requiring the robot to first move to an appropriate position before executing specific operations like picking, which significantly increases the demand for spatial perception in low-level controllers. Conversely, our design of motion-based skills, by incorporating specific information into these skills for motion planning, separates the spatial-related decisions from the subsequent executions, thereby simplifying the task for low-level controllers.

Different behaviors of motion-based skills exhibit distinct semantics in various contextual scenarios. For instance, some motions may simply represent movements to a specific position, while others involve interactions with other objects, such as using a stick to strike a drum. In our design, *move* is considered as the most basic and versatile skill, which can represent all types of motions. Based on *move*, we also define some specialized motion-based skills, which contain more complex and specific semantics, such as *pull*, *rotate*, and *press*. After translating the motions of robot arms into more specific motion-based skills with the help of VLMs, we can conduct distinct motion planning for different skills based on the context, such as pushing the object following a certain direction or rotating the robot arm along a trajectory, or assigning more proficient controllers to execute them, achieving a higher success rate. Note that we carefully define these specialized motion-based skills to prevent overlap in semantics, allowing us to further extend them in the future to accomplish more challenging tasks in the remaining parts of RH20T.

Gripper-based Skills. Gripper-based skills are intuitive. They include operations related to the gripper, such as *pick*, which requires first opening the gripper and then closing it after the object has been picked up.

Generalizable Objects in Skills. We design a template for those skills interacting with objects, e.g., “*pick* the {attribute} {object}”. The high-level planners in CGAs can perceive specific object categories and attributes with their broad knowledge. It provides strong generalization on objects and attributes, reducing the need for providing primitive skills for each type of object like SayCan [24]. Meanwhile, the placeholders in the template will be completed during annotation, serving as the ground truth in RH20T-P dataset.

Detailed List of Primitive Skills. We follow the above principle to design a set of primitive skills on the RH20T-P dataset. We provide the detailed list of primitive skills in RH20T-P in fig. 3 (b).

2.3 Hindsight Primitive-level Annotation

In RH20T, 7-DoF control information of teleoperation in each episode and the terminate signal are recorded, where the control information includes a triplet (pos_x, pos_y, pos_z) to describe the coordinates in 3D space, a quaternion $(rot_w, rot_x, rot_y, rot_z)$ to describe the rotation and a scalar describe the gripper width. We further ask the annotators to watch the complete episode and segment each episode into video clips. As shown in fig. 3 (c), each of video clip is annotated with its corresponding start frame and end frame, as well as specific primitive skills. Additionally, the placeholders in some object-related primitive skills are annotated based on the video. In this way, we can obtain a series of video clips with specific primitive skills for each episode. We then generate the spatial information for each motion-based skill using the control information in RH20T following section 2.2. Besides, we use GPT-4V [15] to caption each episode, providing detailed descriptions for each scene. Most of the hallucinations in these captions are removed after human inspection. They are also included in RH20T-P dataset as supplements.

3 RA-P: A Potential and Scalable CGA on RH20T-P

To verify the effectiveness of RH20T-P for CGAs, we build **RA-P** on the RH20T-P. In section 3.1, we first delineate the overall paradigm of RA-P to plan and then execute given robotic tasks. We

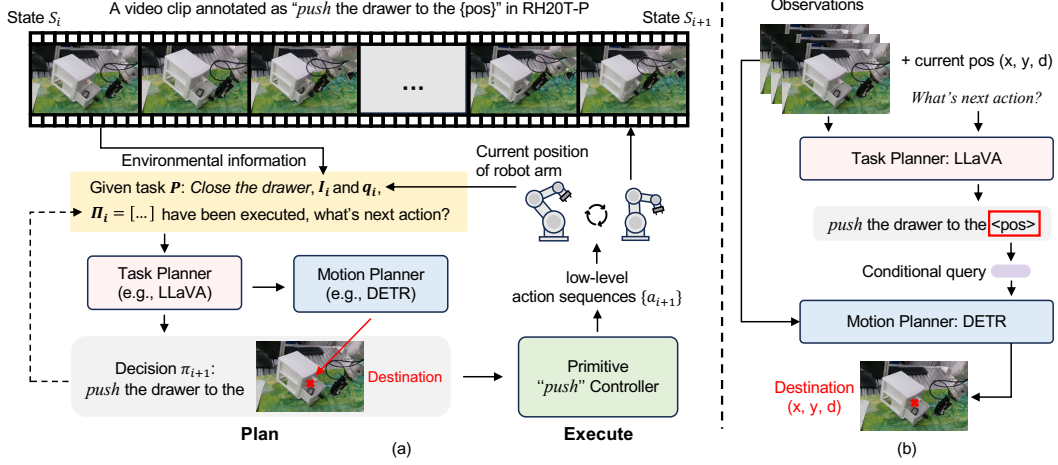


Figure 4: **Overview of RA-P.** (a) Plan-Execute paradigm of RA-P. (b) We design two planners specialized in task decomposition and motion planning in RA-P, respectively.

next elaborate on the design of the high-level planner and low-level controller in section 3.2 and section 3.3, respectively.

3.1 Overall Paradigm of RA-P

Plan-execute Paradigm. As shown in fig. 4 (a), given a language description P , we first define a state S_i , where the agent has completed a part of the task from the initial state S_0 . Based on the historical decisions $\Pi_i = \{\pi_0, \dots, \pi_i\}$ and the observation under S_i , the high-level planner in RA-P can predict the next primitive-level decision $\pi_{i+1} = \text{Planner}(P, \Pi_i, I_i, q_i)$. Here, we collect the environmental information I_i , *i.e.*, multi-frame visual features, along with the position of the robot arm q_i as the input observation. Next, the low-level controller maps the decision π_{i+1} to the low-level action sequences $\{a_{i+1}\} = \text{Controller}(\pi_{i+1}, I_i)$, which is then performed by the robot arm to transition from S_i to the next state S_{i+1} . The planner then makes the next decision under S_{i+1} . In this manner, the planner and controller continue to work alternately until the planner gives a *done* decision, indicating that the task has been completed.

Mapping the Paradigm to RH20T-P. During training, we can associate each transition $S_i \rightarrow S_{i+1}$ with a corresponding video clip from the RH20T-P dataset, based on the assumption that the part of task before S_i has been successfully executed. In each clip, we collect RGB images and the position of the robot arm from the initial frame as the input observations (I_i, q_i) , expecting the high-level planner to make decisions π_i that are consistent with the primitive skill annotated on the clip. The 7-DoF control information recorded in the clips can be regarded as the action sequences $\{a_{i+1}\}$ predicted by the low-level controller. During the inference stage, we follow the above paradigm to sequentially conduct planning and execution step by step.

3.2 High-level Planner

We follow the design of primitive skills in RH20T-P to design our high-level planner. In RA-P, the decision π_i made by the high-level planner can be broken down into a primitive skill defined in section 2.2 and its associated spatial information. Accordingly, we split the functionality of the high-level planner into two components: the task planner, which linguistically decomposes the tasks, and the motion planner, which predicts the spatial motion trends of the robot arm for motion-based skills. The whole architecture of the high-level planner in RA-P is illustrated in fig. 4 (b).

Task Planner. We employ a popular open-sourced VLM, namely LLaVA [16], as our task planner and fine-tune it through instruction tuning. In addition to the vision-language dataset used in

LLaVA, we also generate a robotic instruction-following dataset based on RH20T-P. As outlined in section 3.1, we organize the historical decisions Π_i and the observations $\{I_i, q_i\}$ under S_i in the instructions, expecting the VLMs to make the next primitive-level decisions using the next-token prediction method. More information about instructions and architecture of task planners are provided in the appendix.

Motion planner. Poor proficiency in language models in spatial perception and the challenge of linguistically expressing precise positions can sometimes hinder the task planner in making spatially-related decisions. Incorporating spatial information into motion-based primitive skills in RH20T-P allows us to solve this problem by designing an additional module for motion planning. As described in section 2.2, the various forms of spatial information in motion-based skills offer a broader range of options. In this paper, we use the destination (x, y, d) of the trajectory in motion-based skills in RA-P for simplicity and fast implementation, where x, y denote the pixel coordinates in the image, and d denotes the depth relative to the camera. Note that other forms like trajectory or direction are also available in RH20T-P. Specifically, we employ a Deformable DETR [26] as a simple motion planner to localize the next destination (x, y, d) that the robot arm should move to. Inspired by [31, 32, 33], we introduce a special token $\langle \text{pos} \rangle$ to the LLaVA vocabulary. The hidden features of the token $\langle \text{pos} \rangle$ in motion-based skills are treated as a semantic condition representing the spatial and object information. We add them to the object queries in DETR so that the DETR can localize the relevant destination based on the conditional features.

3.3 Low-level Controller

The primary objective of the low-level controller is to learn a mapping from primitive decisions π_i to action sequences $\{a_i\}$ for robotic control. In RA-P, we design two types of controllers, *i.e.*, using hard-code or policy-based method.

Hard-code Controller. Some primitive skills like *open/close* the gripper or *move* are deterministic and can be executed without any adaption to environmental changes. Therefore, we design a set of fixed codes to generate 7-DoF control parameters for these skills.

Policy-based Controller. Other primitive skills like *pick* or *press* require interactions with diverse objects in the environment. For instance, the specific operations for picking a cup and picking a cloth are significantly different. In our practice, we individually train 4 ACTs [3] as policy-based low-lever controllers for *pick*, *push*, *pull* and *press* skills, respectively. As depicted in section 3.1, we use the 7-DoF control sequences during the corresponding transition for training the policy-based controller.

4 Experiments

4.1 Experimental Setup

Evaluation Platform. As depicted in fig. 5 (a), we employ a UR-5 robotic arm equipped with a parallel Robotiq-85 gripper for object interaction. An Intel RealSense RGB-D camera is positioned in front of the robot to capture environmental information. We set up a new environment with a green tablecloth, which differs from those in the training dataset in RH20T-P.

Implementation Details. We use LLaVA-1.5-7B [16] as our task planner, training it from scratch following its two-stage pipeline. In both stages, we utilize ShareGPT4V [34] dataset in place of the original dataset used in LLaVA to improve multi-modal perception and instruction-following capabilities. As detailed in section 3.2, we further incorporate a robotic instruction dataset generated on RH20T-P during the second stage to acquire specific robotic knowledge. All parameters in LLaVA, except for multi-modal vision encoder CLIP [35], are fine-tuned through instruction tuning.

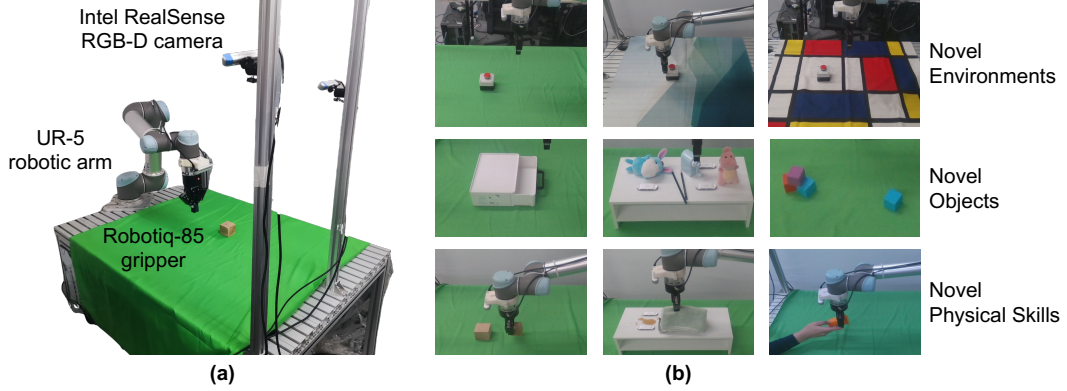


Figure 5: (a) The evaluation platform. (b) Three levels of generalization.

In RA-P, the destination of each movement of the robot arm is used as spatial information in motion-based skills. We employ a single-scale Deformable DETR [26] pre-trained on MS-COCO [36] dataset for motion planning. The visual features extracted by CLIP in LLaVA are fed to DETR instead of those extracted by CNN for better consistency in visual perception. We train a learnable object query with a length of 1 from scratch in DETR to locate the destination in each motion-based skill. We also add two heads for predicting the coordinate (x, y) and the depth d , and train them from scratch. Deformable DETR in RA-P is jointly trained with LLaVA in the second stage to ensure that it can locate the destination based on the semantic context provided by LLaVA. Besides, to adapt DETR to the new environment characterized by different sensors and novel objects, we collect a small supplementary dataset and introduce a third stage for fine-tuning while maintaining the entire LLaVA model frozen.

Evaluation. To evaluate the effectiveness and generalizability of RA-P, we select 8 unseen tasks and test them on the evaluation platform described above. Note that our task planner, LLaVA, does not encounter the distribution of the evaluation tasks during instruction tuning, showcasing its ability to generalize. To compare with agents using proprietary VLMs, we introduce GPT-4V [15] as one of our counterparts. Specifically, we use in-context learning to guide GPT-4V in selecting primitive skills identical to those in RA-P for decision making and predicting simple destination for each motion-based skill for motion planning. However, due to its insufficient spatial perception abilities and difficulties in injecting external knowledge to assist with GPT-4V, GPT-4V is asked to predict 2D coordinates related to input images as a compromise (see Appendix appendix A for more details). Additionally, we follow [27] to individually train 8 ACTs on datasets of each evaluation task collected in the evaluation platform as our baseline. During execution, we conduct 10 tests for each task to measure the success rate, which includes both basic scenarios and more challenging ones that involve longer movements or the presence of extraneous objects as distractions. We also scrutinize the decision-making process of planners. Note that we reserve the assessment of motion planning for the execution phase because a seemingly reasonable destination during the planning stage can also lead to failure of the low-level execution.

4.2 Experimental Results on Novel Tasks

As shown in table 1, our evaluation on generalizability is divided into three levels: **1.)** Generalization to novel environment, **2.)** Generalization to novel objects & compositions, **3.)** Generalization to novel physical skills. Note that these levels are incremental, suggesting that the test for level 3 also includes the assessments for levels 1 and 2.

Comparison with Imitation-based Methods. ACT exhibits consistently poor performance, even in the basic tasks such as “Pick up Block”. Notably, the success rate of ACT decreases when the initial position of the robot arm is far from the target objects. In contrast, in our proposed RA-P, by decou-

Table 1: Comparison of planning accuracy and execution success rate on novel tasks.

Tasks	Planning Accuracy (%)		Execution Success Rate (%)		
	GPT-4V	RA-P	ACT	GPT-4V	RA-P
Novel Environments					
Pick up Block	100	100	30	30	90
Press Button	100	100	50	40	70
Novel Objects & Compositions					
Take Down Object	100	80	10	20	70
Close Drawer	100	90	30	10	70
Novel Physical Skills					
Wipe Table	90	70	0	0	60
Throw Garbage	100	90	10	0	70
Stack Blocks	100	90	0	0	60
Receive Object	100	100	20	10	80

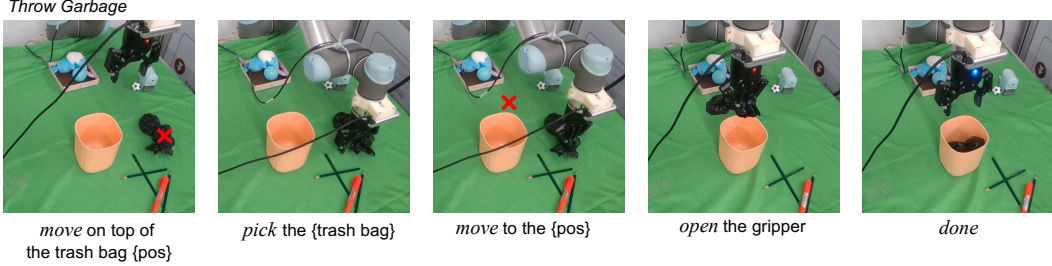


Figure 6: **Robustness on Object Distractions.** Note that the “X” in the diagram represents the destination predicted by the motion planner.

pling spatial information from subsequent execution through motion planning, the primitive-level controller can perform picking operation near the target object, resulting in a significant performance improvement. Given that we use the same setup as the baseline ACT model to train our primitive-level controllers, it demonstrates the effectiveness of motion planning. Besides, these basic tasks often serve as components of more complex tasks such as “Stack Block”. Therefore, the potential of the agents like VILA [22], which do not incorporate motion planning and rely on low-level controllers for spatial perception, is significantly limited by ACT. And for more complex tasks, the gap in the success rate between ACT and RA-P is further widened.

Comparison with Agents Using Proprietary VLMs. We also compare the performance of GPT-4V, which serves as a benchmark for the capabilities of agents using proprietary VLMs as their decision-making backends. Given that a 7B model is used in RA-P as our task planner, agents with GPT-4V achieve better visual perception and logical reasoning abilities in comparison, thus resulting in a higher planning accuracy in some tasks. However, obtaining reliable spatial information from GPT-4V through in-context learning poses great challenges, resulting in a huge disparity in the execution success rate of GPT-4V. In contrast, our RA-P achieves a higher execution success rate across all three levels through composable generalization, especially for novel physical skills, which can be attributed to the well-designed primitive skills and corresponding spatial information in RH20T-P.

4.3 Qualitative Results

Robustness on Object Distractions. We place our RA-P in an environment with various object distractions. As shown in fig. 6, we place an object that can be classified as garbage in a pile of

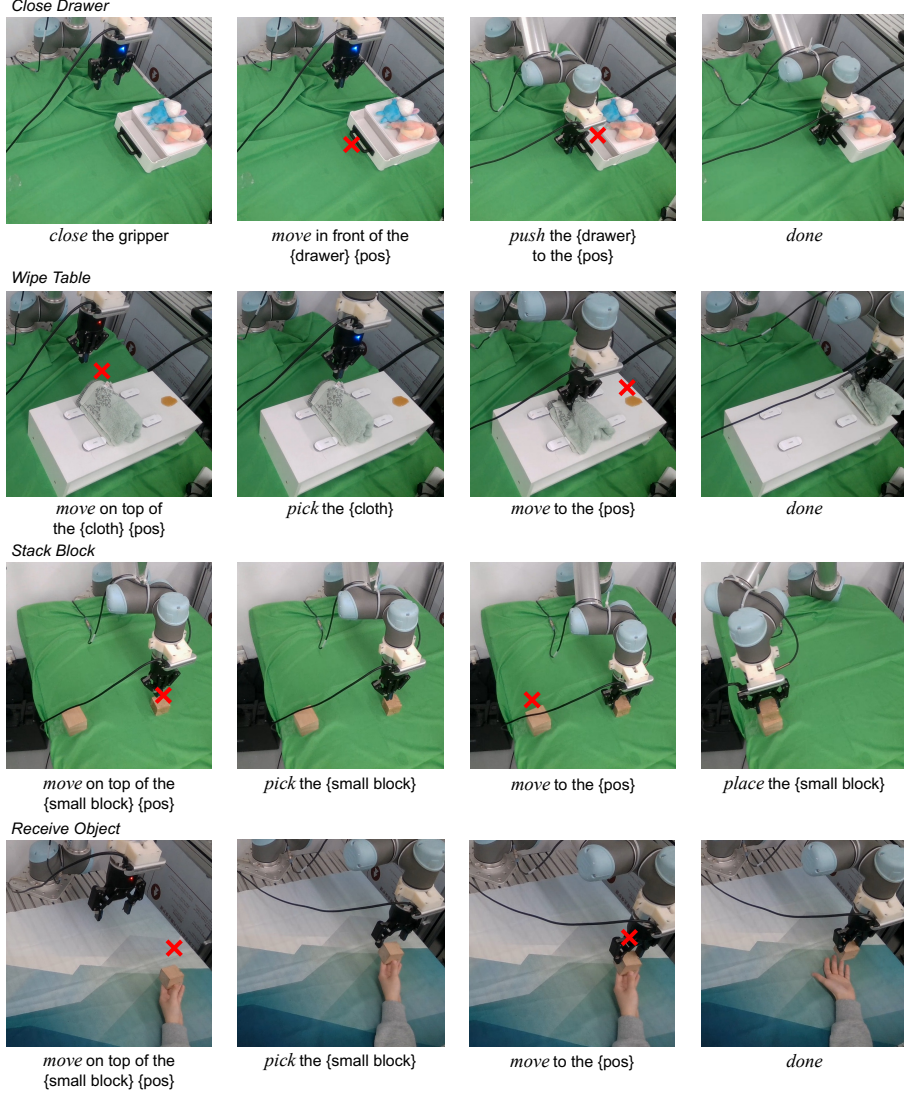


Figure 7: Visualization of RA-P execution on novel tasks.

unrelated objects and ask RA-P to conduct the “Throw garbage” task. RA-P can distinguish the correct target object from surroundings based on the given task and input observations and successfully execute the task, validating the robustness to object distractions.

Visualization of the RA-P Execution. As shown in fig. 7, we provide visualization of RA-P when executing tasks in table 1. More visualization is provided in Appendix appendix F.

5 Related Work

Vision Language Models (VLMs). Vision language models (VLMs) have gained significant attention due to their impressive capabilities in multimodal perception. Some studies [16, 17, 37, 18, 19] incorporate image semantics into language models, dedicated to understanding and reasoning on 2D images. Among these, LLaVA [16] adopts a two-stage instruction-tuning pipeline for general-purpose visual and language understanding. InstructBLIP [37] integrates the language model with an instruction-aware Q-Former [38] for extracting visual semantics that are highly relevant to the given instruction. Additionally, there are also some studies [39, 40, 41] on VLMs in 3D vision. For

example, 3D-LLM [39] introduces the 3D semantics into LLMs by rendering point clouds into 2D images, enabling the models to perform a range of 3D tasks. In this work, we primarily use 2D images as input observations in our CGA, and employ LLaVA [16] as our task planner for decomposing the tasks into predefined primitive skills.

Language Models as Task Planners. Some efforts to utilize LLMs for high-level planning [24, 42, 43, 44, 45, 46, 47] in robotic tasks have showcased their potential in task decomposition. However, without the inherent and unified visual perception abilities, LLMs rely on auxiliary visual modules such as external perception APIs [44], textual scene descriptions [43] or affordance models [24, 46] to perceive the environment, affecting the planning outcomes. Recent works attempt to use VLMs for planning instead. PaLM-E [21] develops an embodied vision-language model and trains it jointly on internet-scale general vision-language datasets and robotic datasets, but the unavailability of primitive skills results in a continuously changing granularity of output actions, confusing the low-level controllers. VILA [22] and [23] conduct in-context learning [25] with off-the-shelf VLMs, *i.e.*, GPT-4V [15], to decompose robotic tasks into primitive skills. The proprietary nature of GPT-4V leads to extensive prompt engineering and inflexibility to fully utilize other modal information in the environment, *e.g.*, depth images, or point clouds.

Primitive-level Robotic Datasets. The trend of using language models as high-level planners for task decomposition has urgently increased the demand for primitive-level robotic datasets tailored for CGAs. However, most existing robotic manipulation datasets [48, 49, 50, 1, 27] either lack textual annotations or only provide language descriptions for entire tasks, which are insufficient towards the construction of CGAs. While several datasets [28, 29, 30] feature hindsight free-form language annotations for manipulation video clips, these coarsely-grained decompositions may confuse the low-level executor, deviating from the original intention of CGAs. The absence of robotic datasets with fine-grained primitive skills hinders the development of CGAs. In this paper, we propose RH20T-P to fill this gap and provide a new direction for CGAs beyond in-context learning.

Imitation Learning. Imitation learning has been prominently explored for robots to acquire skills directly through observing expert demonstrations. Some research initiatives [2, 3, 4, 51, 1], based on behavior cloning, have made notable progress in using images as input for robot controls. For instance, RT-1 [2] incorporates natural language as an input, capable of outputting actions for various tasks. ACT [3] translates historical images and positions into action sequences in an end-to-end manner, achieving a high success rate in short-sequence tasks.

6 Conclusion

In this work, we propose RH20T-P, a primitive-level robotic dataset. With meticulously defined primitive skills and video clips manually annotated accordingly, we can clearly illustrate various motions of the robot arm in RH20T-P, facilitating the advancement of future composable generalization agents. Fine-tuned on RH20T-P, we further construct a potential and scalable agent called RA-P to validate the effectiveness of RH20T-P. Equipped with two specialized planners for task decomposition and motion planning, our RA-P can generalize to novel physical skills through composable generalization.

Limitation. First, without any adaption, the motion planner in RA-P, *i.e.*, Deformable DETR, has poor generalizability on novel objects, even if the task planner has already recognized them. Despite the acceptable cost of fine-tuning DETR alone, a more generalized motion planner is desirable. Recent progress in using direction [52, 53] or trajectory [54, 55] to guide the robot has shown great potential in generalizable motion planning. Our RH20T-P are in alignment with this vision, and we will leave these explorations of CGAs in the future. Second, limited by the scale of the language model (LLaVA-7B), RA-P still leaves room for improvement in more challenging tasks, especially the long-horizon tasks. Thus, we decide to open source the dataset and code, hoping to promote the development of CGA in the community.

References

- [1] E. Jang, A. Irpan, M. Khansari, D. Kappler, F. Ebert, C. Lynch, S. Levine, and C. Finn. Bc-z: Zero-shot task generalization with robotic imitation learning. In *Conference on Robot Learning*, pages 991–1002. PMLR, 2022.
- [2] A. Brohan, N. Brown, J. Carbajal, Y. Chebotar, J. Dabis, C. Finn, K. Gopalakrishnan, K. Hausman, A. Herzog, J. Hsu, et al. Rt-1: Robotics transformer for real-world control at scale. *arXiv preprint arXiv:22.12.06817*, 2022.
- [3] T. Z. Zhao, V. Kumar, S. Levine, and C. Finn. Learning fine-grained bimanual manipulation with low-cost hardware. *arXiv preprint arXiv:2304.13705*, 2023.
- [4] M. Shridhar, L. Manuelli, and D. Fox. Cliport: What and where pathways for robotic manipulation. In *Conference on Robot Learning*, pages 894–906. PMLR, 2022.
- [5] N. Hansen, Y. Lin, H. Su, X. Wang, V. Kumar, and A. Rajeswaran. Modem: Accelerating visual model-based reinforcement learning with demonstrations. *arXiv preprint arXiv:2212.05698*, 2022.
- [6] A. Adeniji, A. Xie, C. Sferrazza, Y. Seo, S. James, and P. Abbeel. Language reward modulation for pretraining reinforcement learning. *arXiv preprint arXiv:2308.12270*, 2023.
- [7] A. Escontrela, A. Adeniji, W. Yan, A. Jain, X. B. Peng, K. Goldberg, Y. Lee, D. Hafner, and P. Abbeel. Video prediction models as rewards for reinforcement learning. *Advances in Neural Information Processing Systems*, 36, 2024.
- [8] J. Luo, P. Dong, J. Wu, A. Kumar, X. Geng, and S. Levine. Action-quantized offline reinforcement learning for robotic skill learning. In *Conference on Robot Learning*, pages 1348–1361. PMLR, 2023.
- [9] T. Brown, B. Mann, N. Ryder, M. Subbiah, J. D. Kaplan, P. Dhariwal, A. Neelakantan, P. Shyam, G. Sastry, A. Askell, et al. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901, 2020.
- [10] OpenAI. Gpt-4 technical report, 2023.
- [11] A. Chowdhery, S. Narang, J. Devlin, M. Bosma, G. Mishra, A. Roberts, P. Barham, H. W. Chung, C. Sutton, S. Gehrmann, et al. Palm: Scaling language modeling with pathways. *Journal of Machine Learning Research*, 24(240):1–113, 2023.
- [12] H. Touvron, T. Lavril, G. Izacard, X. Martinet, M.-A. Lachaux, T. Lacroix, B. Rozière, N. Goyal, E. Hambro, F. Azhar, et al. Llama: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971*, 2023.
- [13] H. Touvron, L. Martin, K. Stone, P. Albert, A. Almahairi, Y. Babaei, N. Bashlykov, S. Batra, P. Bhargava, S. Bhosale, et al. Llama 2: Open foundation and fine-tuned chat models, 2023. URL <https://arxiv.org/abs/2307.09288>, 2023.
- [14] L. Ouyang, J. Wu, X. Jiang, D. Almeida, C. Wainwright, P. Mishkin, C. Zhang, S. Agarwal, K. Slama, A. Ray, et al. Training language models to follow instructions with human feedback. *Advances in Neural Information Processing Systems*, 35:27730–27744, 2022.
- [15] OpenAI. Gpt-4v(ision) system card, 2023.
- [16] H. Liu, C. Li, Q. Wu, and Y. J. Lee. Visual instruction tuning. *Advances in neural information processing systems*, 36, 2024.
- [17] Z. Yin, J. Wang, J. Cao, Z. Shi, D. Liu, M. Li, X. Huang, Z. Wang, L. Sheng, L. Bai, et al. Lamm: Language-assisted multi-modal instruction-tuning dataset, framework, and benchmark. *Advances in Neural Information Processing Systems*, 36, 2024.

- [18] K. Chen, Z. Zhang, W. Zeng, R. Zhang, F. Zhu, and R. Zhao. Shikra: Unleashing multimodal llm’s referential dialogue magic. *arXiv preprint arXiv:2306.15195*, 2023.
- [19] Z. Peng, W. Wang, L. Dong, Y. Hao, S. Huang, S. Ma, and F. Wei. Kosmos-2: Grounding multimodal large language models to the world. *arXiv preprint arXiv:2306.14824*, 2023.
- [20] A. Brohan, N. Brown, J. Carbajal, Y. Chebotar, X. Chen, K. Choromanski, T. Ding, D. Driess, A. Dubey, C. Finn, et al. Rt-2: Vision-language-action models transfer web knowledge to robotic control. *arXiv preprint arXiv:2307.15818*, 2023.
- [21] D. Driess, F. Xia, M. S. Sajjadi, C. Lynch, A. Chowdhery, B. Ichter, A. Wahid, J. Tompson, Q. Vuong, T. Yu, et al. Palm-e: An embodied multimodal language model. *arXiv preprint arXiv:2303.03378*, 2023.
- [22] Y. Hu, F. Lin, T. Zhang, L. Yi, and Y. Gao. Look before you leap: Unveiling the power of gpt-4v in robotic vision-language planning. *arXiv preprint arXiv:2311.17842*, 2023.
- [23] N. Wake, A. Kanehira, K. Sasabuchi, J. Takamatsu, and K. Ikeuchi. Gpt-4v (ision) for robotics: Multimodal task planning from human demonstration. *arXiv preprint arXiv:2311.12015*, 2023.
- [24] M. Ahn, A. Brohan, N. Brown, Y. Chebotar, O. Cortes, B. David, C. Finn, C. Fu, K. Gopalakrishnan, K. Hausman, et al. Do as i can, not as i say: Grounding language in robotic affordances. *arXiv preprint arXiv:2204.01691*, 2022.
- [25] J. Liu, D. Shen, Y. Zhang, B. Dolan, L. Carin, and W. Chen. What makes good in-context examples for gpt-3? *arXiv preprint arXiv:2101.06804*, 2021.
- [26] X. Zhu, W. Su, L. Lu, B. Li, X. Wang, and J. Dai. Deformable detr: Deformable transformers for end-to-end object detection. *arXiv preprint arXiv:2010.04159*, 2020.
- [27] H.-S. Fang, H. Fang, Z. Tang, J. Liu, J. Wang, H. Zhu, and C. Lu. Rh20t: A robotic dataset for learning diverse skills in one-shot. *arXiv preprint arXiv:2307.00595*, 2023.
- [28] C. Lynch and P. Sermanet. Language conditioned imitation learning over unstructured data. *arXiv preprint arXiv:2005.07648*, 2020.
- [29] S. Nair, E. Mitchell, K. Chen, S. Savarese, C. Finn, et al. Learning language-conditioned robot behavior from offline data and crowd-sourced annotation. In *Conference on Robot Learning*, pages 1303–1315. PMLR, 2022.
- [30] C. Lynch, A. Wahid, J. Tompson, T. Ding, J. Betker, R. Baruch, T. Armstrong, and P. Florence. Interactive language: Talking to robots in real time. *IEEE Robotics and Automation Letters*, 2023.
- [31] Z. Dai, B. Cai, Y. Lin, and J. Chen. Up-detr: Unsupervised pre-training for object detection with transformers. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 1601–1610, 2021.
- [32] Z. Chen, G. Huang, W. Li, J. Teng, K. Wang, J. Shao, C. C. Loy, and L. Sheng. Siamese detr. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 15722–15731, 2023.
- [33] Y. Zang, W. Li, J. Han, K. Zhou, and C. C. Loy. Contextual object detection with multimodal large language models. *arXiv preprint arXiv:2305.18279*, 2023.
- [34] L. Chen, J. Li, X. Dong, P. Zhang, C. He, J. Wang, F. Zhao, and D. Lin. Sharegpt4v: Improving large multi-modal models with better captions. *arXiv preprint arXiv:2311.12793*, 2023.
- [35] A. Radford, J. W. Kim, C. Hallacy, A. Ramesh, G. Goh, S. Agarwal, G. Sastry, A. Askell, P. Mishkin, J. Clark, et al. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, pages 8748–8763. PMLR, 2021.

- [36] T.-Y. Lin, M. Maire, S. Belongie, J. Hays, P. Perona, D. Ramanan, P. Dollár, and C. L. Zitnick. Microsoft coco: Common objects in context. In *Computer Vision—ECCV 2014: 13th European Conference, Zurich, Switzerland, September 6–12, 2014, Proceedings, Part V 13*, pages 740–755. Springer, 2014.
- [37] W. Dai, J. Li, D. Li, A. M. H. Tiong, J. Zhao, W. Wang, B. A. Li, P. Fung, and S. C. H. Hoi. Instructblip: Towards general-purpose vision-language models with instruction tuning. *arXiv preprint arXiv:2305.06500*, 2023.
- [38] J. Li, D. Li, S. Savarese, and S. Hoi. Blip-2: Bootstrapping language-image pre-training with frozen image encoders and large language models. *arXiv preprint arXiv:2301.12597*, 2023.
- [39] Y. Hong, H. Zhen, P. Chen, S. Zheng, Y. Du, Z. Chen, and C. Gan. 3d-llm: Injecting the 3d world into large language models. *Advances in Neural Information Processing Systems*, 36, 2024.
- [40] R. Xu, X. Wang, T. Wang, Y. Chen, J. Pang, and D. Lin. Pointllm: Empowering large language models to understand point clouds. *arXiv preprint arXiv:2308.16911*, 2023.
- [41] Z. Chen, Z. Wang, Z. Wang, H. Liu, Z. Yin, S. Liu, L. Sheng, W. Ouyang, Y. Qiao, and J. Shao. Octavius: Mitigating task interference in mllms via moe. *arXiv preprint arXiv:2311.02684*, 2023.
- [42] J. Wu, R. Antonova, A. Kan, M. Lepert, A. Zeng, S. Song, J. Bohg, S. Rusinkiewicz, and T. Funkhouser. Tidybot: Personalized robot assistance with large language models. *arXiv preprint arXiv:2305.05658*, 2023.
- [43] B. Li, P. Wu, P. Abbeel, and J. Malik. Interactive task planning with language models. *arXiv preprint arXiv:2310.10645*, 2023.
- [44] M. Xu, P. Huang, W. Yu, S. Liu, X. Zhang, Y. Niu, T. Zhang, F. Xia, J. Tan, and D. Zhao. Creative robot tool use with large language models. *arXiv preprint arXiv:2310.13065*, 2023.
- [45] K. Lin, C. Agia, T. Migimatsu, M. Pavone, and J. Bohg. Text2motion: From natural language instructions to feasible plans. *arXiv preprint arXiv:2303.12153*, 2023.
- [46] W. Huang, F. Xia, D. Shah, D. Driess, A. Zeng, Y. Lu, P. Florence, I. Mordatch, S. Levine, K. Hausman, et al. Grounded decoding: Guiding text generation with grounded models for robot control. *arXiv preprint arXiv:2303.00855*, 2023.
- [47] Y. Qin, E. Zhou, Q. Liu, Z. Yin, L. Sheng, R. Zhang, Y. Qiao, and J. Shao. Mp5: A multi-modal open-ended embodied system in minecraft via active perception. *arXiv preprint arXiv:2312.07472*, 2023.
- [48] A. Mandlekar, Y. Zhu, A. Garg, J. Booher, M. Spero, A. Tung, J. Gao, J. Emmons, A. Gupta, E. Orbay, et al. Roboturk: A crowdsourcing platform for robotic skill learning through imitation. In *Conference on Robot Learning*, pages 879–893. PMLR, 2018.
- [49] S. Dasari, F. Ebert, S. Tian, S. Nair, B. Bucher, K. Schmeckpeper, S. Singh, S. Levine, and C. Finn. Robonet: Large-scale multi-robot learning. *arXiv preprint arXiv:1910.11215*, 2019.
- [50] P. Sharma, L. Mohan, L. Pinto, and A. Gupta. Multiple interactions made easy (mime): Large scale demonstrations data for imitation. In *Conference on robot learning*, pages 906–915. PMLR, 2018.
- [51] Z. Mandi, F. Liu, K. Lee, and P. Abbeel. Towards more generalizable one-shot visual imitation learning. In *2022 International Conference on Robotics and Automation (ICRA)*, pages 2434–2444, 2022.

- [52] S. Nasiriany, F. Xia, W. Yu, T. Xiao, J. Liang, I. Dasgupta, A. Xie, D. Driess, A. Wahid, Z. Xu, et al. Pivot: Iterative visual prompting elicits actionable knowledge for vlms. *arXiv preprint arXiv:2402.07872*, 2024.
- [53] X. Li, M. Zhang, Y. Geng, H. Geng, Y. Long, Y. Shen, R. Zhang, J. Liu, and H. Dong. Manipllm: Embodied multimodal large language model for object-centric robotic manipulation. *arXiv preprint arXiv:2312.16217*, 2023.
- [54] W. Zhi, K. Liu, T. Zhang, and M. Johnson-Roberson. Learning orbitally stable systems for diagrammatic teaching. In *CoRL 2023 Workshop on Learning Effective Abstractions for Planning (LEAP)*, 2023.
- [55] J. Gu, S. Kirmani, P. Wohlhart, Y. Lu, M. G. Arenas, K. Rao, W. Yu, C. Fu, K. Gopalakrishnan, Z. Xu, et al. Rt-trajectory: Robotic task generalization via hindsight trajectory sketches. *arXiv preprint arXiv:2311.01977*, 2023.

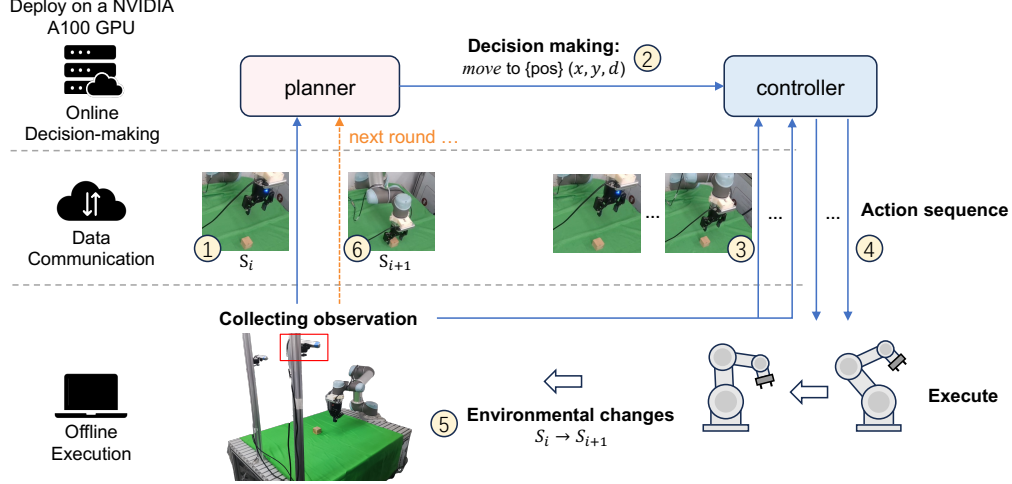


Figure 8: Deployment pipeline of RA-P.

A GPT-4V Execution Setup

Comparing with predicting 3D destinations (x, y, d) in RA-P, we make compromises in evaluating the motion planning of GPT-4V during the execution phase. Due to the lack of transparency of GPT-4V, it is challenging to incorporate the spatial priors, particularly depth information, from specific environments into its knowledge, resulting in inaccurate depth predictions. Meanwhile, GPT-4V provides only plain text responses without intermediate latent information, such as the hidden features of generated tokens, limiting the subsequent process of agents with GPT-4V. Consequently, it is impractical to design additional modules within agents to assist GPT-4V in motion planning. As a compromise, GPT-4V is asked to predict the 2D coordinates related to input images. Given a state S_i in task P , with all preceding steps successfully executed, GPT-4V will predict the next primitive skill and a 2D coordinate as the destination. If the destination is within a reasonable range, then we loosely consider that the step is successful. Only if every decision (including motion planning) at each state in the execution chain of task P is correct, we then determine the execution of P to be successful. However, even with such compromises, the execution success rate of GPT-4V remains low, due to its insufficient spatial perception ability, as shown in fig. 12.

B Deployment Pipeline of RA-P in Execution Phase

B.1 Deployment Pipeline

We design a deployment pipeline during the execution stage. As shown in fig. 8, we deploy LLaVA, Deformable DETR [26] and low-level ACT [3] controllers on a NVIDIA A100 GPU, and develop a communication module between RA-P and robot arm. We conduct the following procedures to perform the plan-execute paradigm during the evaluation stage:

1. **Collecting observation:** the sensor in the evaluation platform captures the information in the environment, and then transmits the RGB images along with the current position of the robot arm to the agent, which is deployed on an A100 GPU, through the communication module.
2. **Decision making:** the task planner will take images and position them as input and make decisions using predefined primitive skills. If any motion-based skill is chosen, the motion planner will be invoked to predict the precise coordinate of the destination (x, y, d) . The coordinate will be transformed into a 3D coordinate using camera calibration.

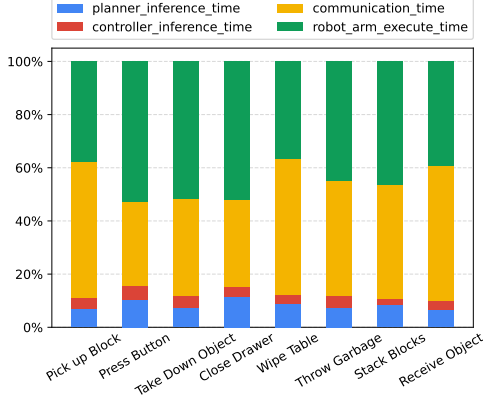


Figure 9: Inference time of RA-P.

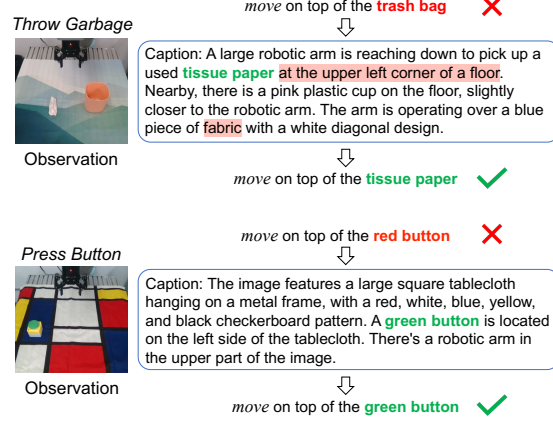


Figure 10: Chain-of-Thought inference.

3. **Mapping to action sequence and then executing:** based on the type of primitive skills, a specific low-level controller will be called to map the decision to the action sequence. For policy-based controllers, we predict action sequences based on current observations for the next 5 steps. After receiving and executing the 5-step action sequence, the robot arm collects information again and transmits it to controllers, repeating until the controllers give a terminate signal. For hard-code controllers, we interpolate a straight line between the starting position of the robot arm and the predicted destination to obtain the action sequence. The robot arm then executes the action sequence, omitting the multiple data transfers and communications like those in policy-based controllers.
4. **Iterative plan-execute process:** Once the controllers complete the decision made by the planners, we will return to step 1 to start a new round of the plan-execute process until the planner ultimately gives a *done* decision.

B.2 Inference Time

We analyze the inference time of RA-P. The result is shown in fig. 9. The time spent by the planner making decisions ($\sim 8\%$) and controller mapping decisions to action sequences ($\sim 3\%$) together accounts for a relatively low proportion of the total time, while the time spent by data communication ($\sim 40\%$) is high. Overall, using a 7B language model as a decision-making backend for inference is acceptable in terms of time, and there is still room for using a larger-scale language model. In the future, we will continue to optimize the efficiency of the entire pipeline through asynchronous communication.

C Chain-of-Thought Inference

The limited distribution of objects and scenes in our data source, *i.e.*, RH20T [27], may lead to perceptual errors in classifying object categories and attributes during decision-making. Since we fine-tune the VLM using both the original vision-language dataset and the robotic dataset, it still retains robust and generalizable scene description capabilities. Consequently, we propose to use the VLM to first describe the scene and then make decisions based on the generated descriptions. We refer it as to Chain-of-Thought inference (CoT inference). No extra adjustments are required during the training phase; instead, we can directly add the descriptions generated by VLM to the prompts for inference. Note that we only caption the scene before the initial decision, and all subsequent decisions within the same task rely on the same description, thereby increasing a minor overhead to the inference time of planning. As illustrated in the fig. 10, CoT inference can effectively reduce many perceptual errors with scene descriptions. Besides, given the scale of model we used in RA-P (7B), there is still room for improvement in captioning to assist VLM with subsequent decision-

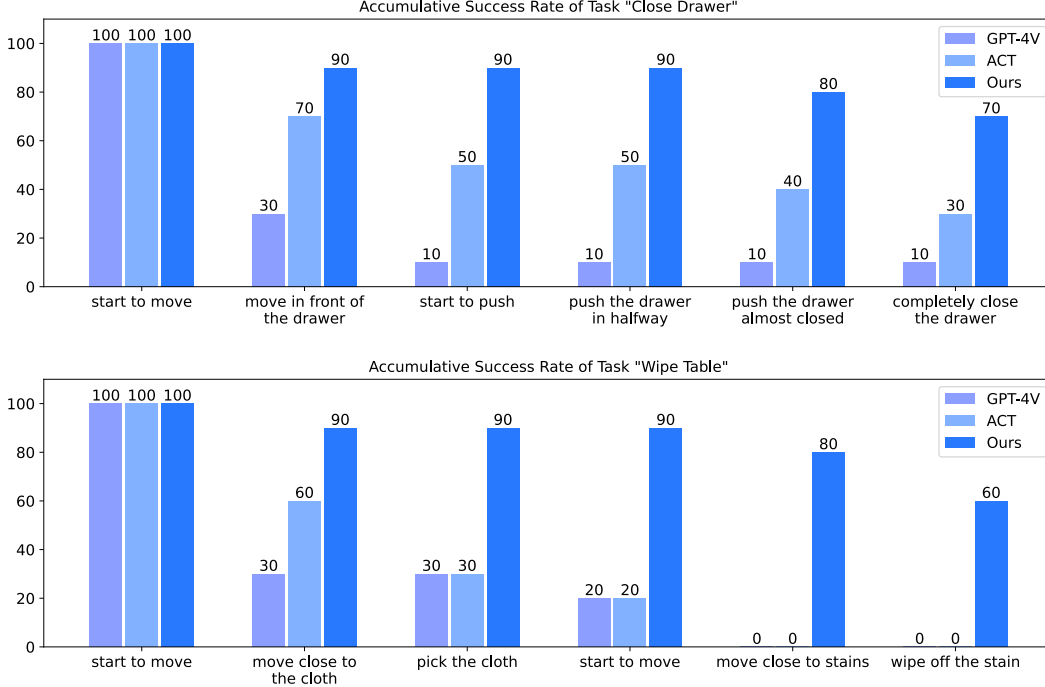


Figure 11: **Cumulative success rates for different stages of several tasks.**

making, especially in terms of hallucinations included in the descriptions (texts marked with a red background in fig. 10) and scenarios with multiple object interferences.

D Cumulative Success Rates

We show the cumulative success rates for different steps of several tasks in fig. 11. It is observed that the majority of task failures in agents with GPT-4V are attributed to insufficient localization abilities. The positions predicted by GPT-4V, which are far from the target object, result in the failure of the low-level controller’s execution. In contrast, our RA-P trained on RH20T-P can provide more reasonable spatial priors, resulting in a higher execution success rate. We also present several failure cases of RA-P in fig. 12. They primarily include failure in low-level controllers (such as failing to pick a target object), DETR localization deviation (especially in scenarios with distractions), and perception error of VLM leading to subsequent incorrect positioning.

E System Prompts and Instructions

System Prompt for Robotic tasks in RA-P. System prompt in RA-P is shown in table 2.

Instructions for Robotic tasks in RA-P. The list is shown in table 3. In these instructions, `{task_desc}`, `{historical_decisions}` and `{robot_arm_pos}` denote language specifications like “Pick Blocks”, historical decisions made by task planners before current state and the position of the robot arm, respectively. During training, we randomly select one of the instructions in the conversations.

System Prompts for agents with GPT-4V in Execution Phase. We evaluate agents with GPT-4V through in-context learning. Detailed system prompt is shown in table 4.

```
messages = [{"role": "system", "content": f""" A chat between a curious user and an
artificial intelligence assistant. The assistant is required to guide a robotic arm to perform a complex
and comprehensive robotic task. Based on the current observation, completed actions, and current
position of the robotic arm, you should give the next action for the robotic arm.
```

The given completed actions and the decision for the next command are composed of the following optional actions: [

- *move* to the {pos}
- *move* {on top of | in front of | ...} the {object}
- *push / pull* the {object} to the {pos}
- *pick / place* the {object}
- *open / close* the gripper
- *press* the {object} {pos}
- *rotate* {clockwise | counterclockwise} {pos}
- *done*

Here {object} and {pos} are the templates for the value that needs to be predicted. The position of the robotic arm is represented as [x, y, d] with floating numbers, where x and y denote the coordinates of the robotic arm range from 0 to 1, and d denotes the depth of the robot arm in the observation. """]]

Table 2: The system prompt for robotic tasks in RA-P.

F More Qualitative Results

We provide many video demos to demonstrate the advantages of RA-P over other baselines on our [project page](#), all benefiting from the training on RH20T-P. A simple comparison is provided in [fig. 12](#).

G All Tasks in RH20T-P

We provide the list of tasks in RH20T-P in [table 5](#) and [table 6](#).

H Potential Social Impact

The proposed RH20T-P dataset and RA-P model demonstrate effectiveness and generalization in robotic tasks, especially in novel physical skills, which can benefit the future development of CGAs. The violent elements (*e.g.*, using a knife) in the dataset and related knowledge learned by the robot may have some potential negative social impacts. However, considering that our data source, RH20T, has already been publicly released, these impacts are controllable.

- To accomplish the task {task_desc}, the following actions have been sequentially completed: {historical_decisions}. Based on the observation, the current position of the robotic arm is {robot_arm_pos}. What should be the next action of the robotic arm?
- Given the embodied task {task_desc}, the subsequent steps have been undertaken: {historical_decisions}. Referring to the observation, the present location of the robotic arm is indicated by {robot_arm_pos}. What would be the appropriate next action for the robotic arm?
- With the assigned task {task_desc}, the following actions have been accomplished: {historical_decisions}. As shown in the observation, the current position of the robotic arm is {robot_arm_pos}. Considering this, what are the next advisable action for the robotic arm?
- In light of the assigned task {task_desc}, the ensuing steps were carried out: {historical_decisions}. Observing the image given, the robotic arm's current position is marked as {robot_arm_pos}. In light of this, what is the recommended next step for the robotic arm?
- For the purpose of achieving the goal of {task_desc}, we have progressively executed these steps: {historical_decisions}. As indicated by the observation, the robotic arm's current location is {robot_arm_pos}. What would be the advisable subsequent action for the robotic arm?
- To achieve the goal of {task_desc}, we have progressively executed these steps: {historical_decisions}. As indicated by the observation, the robotic arm's current location is {robot_arm_pos}. What would be the advisable subsequent action for the robotic arm?
- To reach the objective of {task_desc}, we have methodically performed the following actions: {historical_decisions}. The observation shows that the robotic arm is currently positioned at {robot_arm_pos}. Given this, what should be the next action for the robotic arm?
- Aimed at accomplishing the goal of {task_desc}, we have systematically undertaken these steps: {historical_decisions}. As depicted in the supplied image, the location of the robotic arm is currently {robot_arm_pos}. Considering this, what would be the prudent next step for the robotic arm?
- With the given task {task_desc}, a sequence of actions has been diligently executed: {historical_decisions}. The current position of the robotic arm, as shown in the image, is {robot_arm_pos}. Based on this, what would be the logical next action for the robotic arm?
- Considering the ongoing {task_desc}, our approach has involved a series of progressive steps: {historical_decisions}. The image provided points out the current position of the robotic arm at {robot_arm_pos}. What is the next recommended action for the robotic arm in this scenario?

Table 3: The list of instructions used for robotic tasks in RA-P.

```
messages = [{"role": "system", "content": f""" You are placed inside an embodiment
environment with a robot arm and a gripper with it. Given a goal (task description) in language and
multi-frame RGB images depicting the scene and objects, you are required to decompose the goal
into subtasks step by step to ensure it can be accomplished eventually.
```

At each round of conversation, I will give you:

- multi-frame RGB images, time from past to present
- task description in language
- historical decisions already made
- The relative coordinates (x, y, d) in the current image where the gripper is located. The x-axis is from left to right, and the y-axis is from top to bottom, with a range of values between 0 and 1. For example, the top left corner of the image is (0,0), and the top right corner is (1,0).

Your reply should be one of the primitive skills, as follows:

- *move* to the {pos}
- *move* {on top of | in front of | ...} the {object}
- *push / pull* the {object} to the {pos}
- *pick / place* the {object}
- *open / close* the gripper
- *press* the {object} {pos}
- *rotate* {clockwise | counterclockwise} {pos}
- *done*

Remember to replace the {pos} in primitive skills with relative coordinates (x, y) as the position in the image. """}]

Table 4: The system prompt for Agents with GPT-4V in Execution Phase.

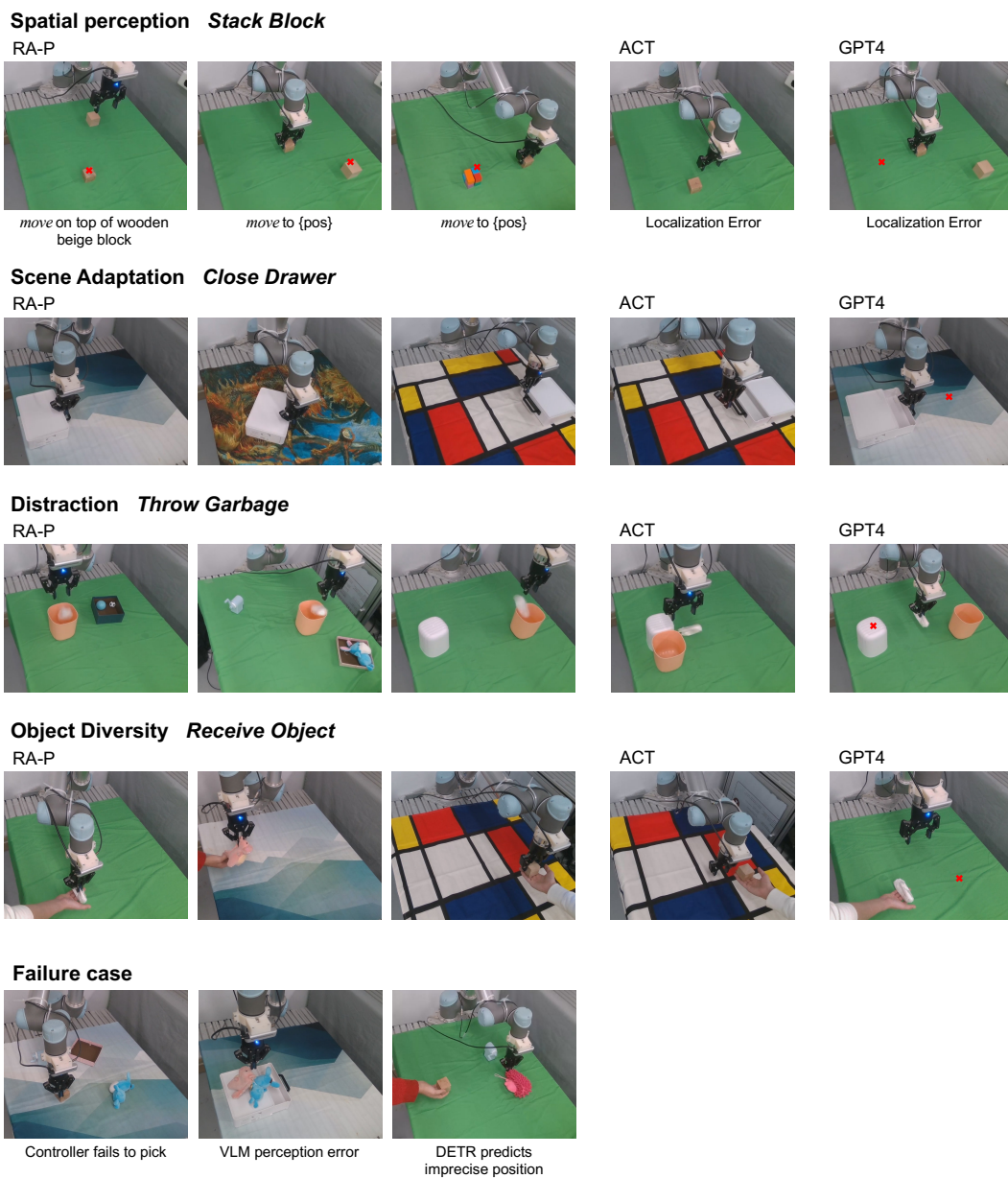


Figure 12: **More qualitative comparison on spatial perception and generalizability.**

- Turn off the desk lamp
vision reasoning
- Take the dish off the dish rack
vision reasoning
- Turn the knob to increase the volume of the speaker
- Placing a piece on the chessboard to complete the setup.
vision reasoning
- Clean the table with a cloth
vision reasoning use tools
- Throw the garbage
- Take the pencil out from the pencil sharpener
- Pull out a napkin
- Plug in the power cord of the desk lamp, turn on the socket, and light up the desk lamp
vision reasoning long-horizon
- Press three buttons from left to right in sequence
- Grab the block and place it at the designated location
- Turn the hands of a clock
special trajectory
- Remove the object from the scale
- Play the first move as black on the 3-4 points in the upper right corner of the Go board
vision reasoning
- Grasp the handle and open the drawer
- Remove the bubble ring from the assembled bubble ring and ball
long-horizon
- Pick up one small block
- Take the photo frame down from the bracket
vision reasoning
- Put the object on the shelf
- Pick up and place an object with obstacles
special trajectory
- Put the knife on the cutting board
- Press the button
- Approach and touch the side of the small block
- Hold a block with the gripper and sweep it from left to right
- Close the drawer
vision reasoning

Table 5: The list of tasks in RH20T-P.

- Take everything out of the gift box
long-horizon
- Press a button from top to bottom with obstacles
special trajectory
- Open the microwave door
- Turn on the water tap
- Place the handset of the telephone on the corresponding phone cradle
- Play the drum
use tools
- Turn on the desk lamp by pressing the button
- Turn on the power strip by pressing the button
- Put the cup on the cup rack
- Open a sliding window
vision reasoning
- Move an object from one box to another
- Cover the pot with the lid
- Push an object with obstacles
special trajectory
- Use the gripper to push the small block from left to right
- Pick up the cup
- Wipe the tabletop with a sponge
vision reasoning use tools
- Take the object down from the shelf
- Push down the lever
special trajectory
- Turn off the water tap

Table 6: The list of tasks in RH20T-P.