

MOKA: Open-World Robotic Manipulation through Mark-Based Visual Prompting

Kuan Fang^{*1} Fangchen Liu^{*1} Pieter Abbeel¹ Sergey Levine¹

<https://moka-manipulation.github.io>

Abstract—Open-world generalization requires robotic systems to have a profound understanding of the physical world and the user command to solve diverse and complex tasks. While the recent advancement in vision-language models (VLMs) has offered unprecedented opportunities to solve open-world problems, how to leverage their capabilities to control robots remains a grand challenge. In this paper, we introduce Marking Open-world Keypoint Affordances (MOKA), an approach that employs VLMs to solve robotic manipulation tasks specified by free-form language instructions. Central to our approach is a compact point-based representation of affordance, which bridges the VLM’s predictions on observed images and the robot’s actions in the physical world. By prompting the pre-trained VLM, our approach utilizes the VLM’s commonsense knowledge and concept understanding acquired from broad data sources to predict affordances and generate motions. To facilitate the VLM’s reasoning in zero-shot and few-shot manners, we propose a visual prompting technique that annotates marks on images, converting affordance reasoning into a series of visual question-answering problems that are solvable by the VLM. We further explore methods to enhance performance with robot experiences collected by MOKA through in-context learning and policy distillation. We evaluate and analyze MOKA’s performance on various table-top manipulation tasks including tool use, deformable body manipulation, and object rearrangement.

I. INTRODUCTION

The pursuit of open-world generalization poses a significant challenge for robotic systems: Solving various tasks commanded by humans in complex and diverse environments requires a profound understanding of the physical world. An appealing prospect for handling this challenge is to employ large pre-trained models that encapsulate extensive prior knowledge from broad data and bringing it to bear on novel problems. Recent advances in large language models (LLMs) and vision-language models (VLMs) provide particularly promising tools in this regard, with their emergent and fast-growing conceptual understanding, commonsense knowledge, and reasoning abilities [10, 50, 51, 6, 9, 1, 52, 31, 53, 32]. However, existing large models pre-trained on Internet-scale data still lack the capabilities to understand 3D space, contact physics, and robotic control, not to mention the knowledge about the embodiment and environment dynamics in each specific scenario. This creates a large gap between the promising trend in computer vision and natural language processing and applying them to robotics. It remains an open question how such tools can guide a robotic system to solve manipulation tasks in the physical world.

^{*}Equal contribution. ¹University of California, Berkeley. Correspondance to: {kuanfang, fangchenliu}@berkeley.edu

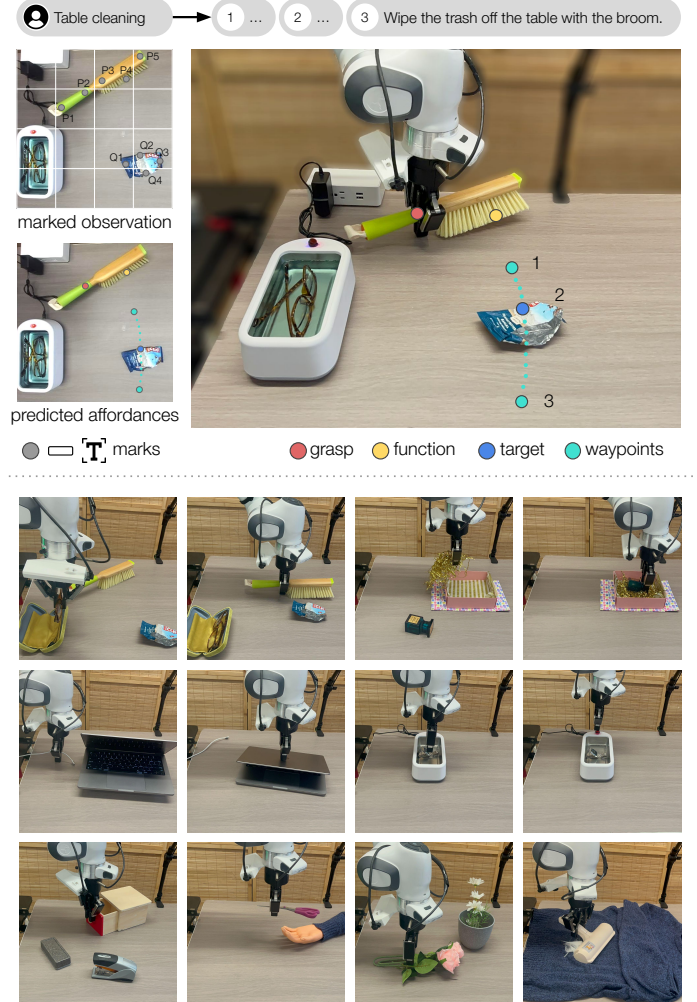


Fig. 1: To solve manipulation tasks with unseen objects and goals, MOKA employs a VLM to generate motions through a point-based affordance representation (plotted as colorful dots on the images). By annotating marks (e.g., candidate points, grids, and captions) on the observed 2D image, MOKA converts motion generation into a series of visual question-answering problems that the VLM can solve.

Recently, a growing body of research has been dedicated to utilizing pre-trained large-scale models for robotic control. To incorporate broad knowledge, these approaches either directly prompt or fine-tune the large models to generate plans [2, 22, 21, 7], rewards [38, 29, 40, 68], codes [34, 56, 61],

etc. Despite the encouraging results they have demonstrated, these approaches are subject to notable limitations. Since the advances in LLMs precede VLMs, many previous approaches first process the raw sensory inputs to obtain the language description of the environment and then query LLMs to perform reasoning and planning in the language domain. However, relying solely on high-level language descriptions may overlook the nuanced visual details of environments and objects, which are vital for accurately completing tasks. Moreover, existing approaches usually require non-trivial effort in designing in-context examples [23, 34, 56] to ensure LLMs can produce desired predictions on similar tasks. As a result, the tasks that can be solved by these approaches are largely constrained by such manual efforts.

In this work, we study how to effectively endow robots with the ability to solve novel manipulation tasks specified by free-form language instructions using VLMs. Our key insight is to find an intermediate affordance representation that connects the VLM’s prediction on images with the robot’s motion in the physical world. This affordance representation should satisfy two essential requirements. First, it should be feasible for the VLM to predict such representation based on the visual observation of the environment and the task description. Second, the representation should compactly capture the information about the relevant properties of the robot’s motion so that it can be easily executed on the robot.

To this end, we propose Marking Open-world Keypoint Affordances (MOKA), an approach that employs VLMs for robotic manipulation through mark-based visual prompting. As shown in Fig. 1, MOKA leverages a compact affordance representation consisting of a set of keypoints and waypoints, defined on open sets of objects and tasks. This point-based affordance representation is then used to specify the desired motion for the robot to solve the task. To generate the motions given the free-form language descriptions, MOKA uses hierarchical visual prompting to convert the affordance reasoning problem into a series of visual question-answering problems. Drawing inspirations from recent advances in visual question-answering [63], we use mark-based visual prompting to enable the VLM to attend to the important visual cues in the observation image and further simplify the point generation problem into multiple choice questions. As shown in the top-left part of Fig. 1, we plot candidate points on the image and query the VLM to select the optimal set of points that will result in the desired motion. The predicted keypoints and waypoints are used to specify a motion trajectory which can cover a wide range of manipulation skills such as picking, placing, pressing, tool use, etc.

We demonstrate the effectiveness of MOKA on various manipulation tasks, with robustness across the variations of instructions, objects, and initial arrangements. We further use MOKA to collect successful trajectories in each task. The collected trajectories can be used as in-context examples to bootstrap the performance of VLM, and can also be leveraged as demonstrations to train a better student policy. Our experiments show that MOKA can achieve state-of-the-art

performance on various evaluation tasks in a zero-shot manner, with consistent improvements using clean and intuitive in-context examples. To summarize, our contributions are:

- We introduce a point-based affordance representation that bridges the VLM’s prediction on RGB images and the robot’s motion in the physical world.
- We propose a mark-based visual prompting approach that converts affordance reasoning into a series of visual question-answering problems.
- We demonstrate that the proposed approach, MOKA, can effectively generate motions for open-world manipulation in zero-shot and few-shot manners.

II. RELATED WORK

Large language models and vision-language models. Recent advances in several domains and applications have been greatly influenced by the substantial progress achieved through large language models (LLMs) and vision language models (VLMs) [10, 50, 51, 6, 9, 1, 52, 31, 53, 32]. While these models can already solve various tasks in a zero-shot manner [1], well-designed prompts still serve an important role in further eliciting more advanced capabilities. As demonstrated by Brown et al. [6], few-shot prompting can match or surpass the performance of fine-tuning on LLMs. Additionally, other prompting techniques [62, 69, 28, 67] have been proposed to improve LLMs’ reasoning capabilities. Although these methodologies may initially appear enigmatic, they have been empirically validated to consistently demonstrate scalability across various models [6, 9, 1, 66]. Apart from the language prompts, recent vision models [27, 70, 30] are capable of supporting visual prompts, including points, boxes, masks, and texts. Such visual prompts exhibit greater diversity in both form and content, harnessing vision-language models for various application scenarios like perception [65], image editing [55] and reasoning [63].

Incorporating robust visual reasoning capabilities becomes imperative to build autonomous robots capable of making decisions in unstructured environments. Although the predictions by existing VLMs cannot seamlessly engage in zero-shot reasoning, they can still be effectively harnessed across diverse tasks. Prior work such as SoM [63] draws visual marks on the objects in an image using numbers and segmentation masks and demonstrates the mark-based visual prompting scheme unleashes the reasoning capabilities of GPT-4V, such as object counting and inferring spatial relationship. Inspired by SoM [63], we leverage GPT-4V for robotic manipulation. We represent manipulation skills with a set of affordance points and motion waypoints. Instead of using object segmentation masks, we use keypoints and grid cells as visual prompts as shown in Fig. 1, and then query GPT-4V to choose the affordance points and motion waypoints from the visual marks, followed by executable point-based motion plans.

Foundation models in robotics. Building upon the successes of foundation models, the field of robotics is currently witnessing a rising interest in utilizing LLMs and other foundation models in various applications. With large model

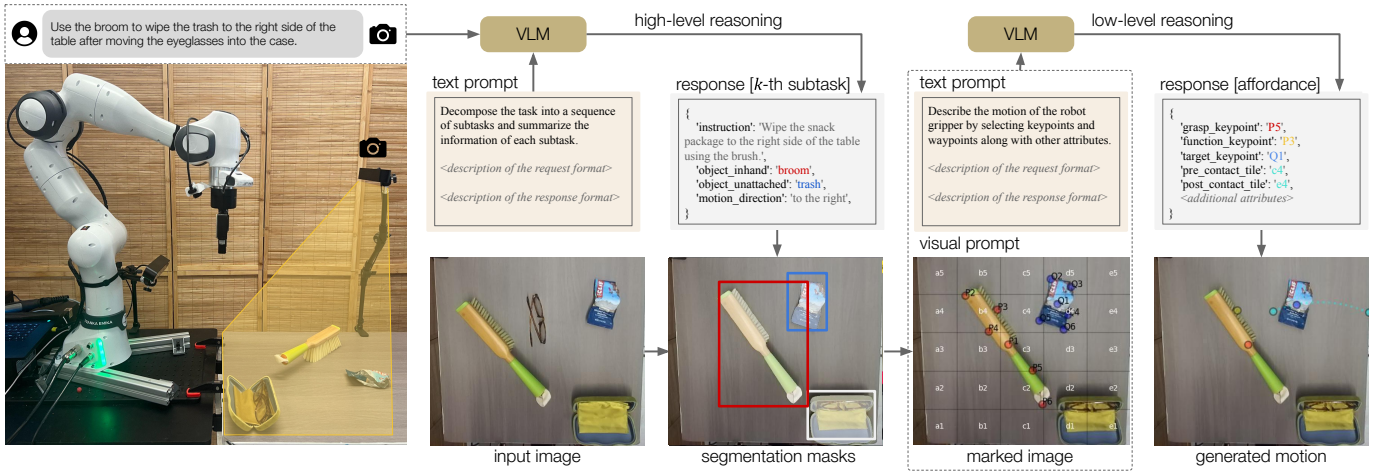


Fig. 2: **Overview of MOKA.** We propose a hierarchical approach to prompt the VLM to perform affordance reasoning based on the free-form language instruction of the task and the visual observation of the environment. On the high level, we query the VLM to decompose the free-form language description of the task into a sequence of subtasks and summarize the subtask information. On the low level, the VLM is prompted to produce the keypoints and additional attributes for the affordance representation defined in Sec. IV-A.

capacities, transformers [60] and VLMs can be trained from scratch or fine-tuned to directly predict actions through imitation learning and reinforcement learning [12, 5, 24, 37, 45]. However, training large models for open-world scenarios would usually require extensive data and supervision. Pre-trained LLMs can also be used to generate high-level task plan in natural language [2, 22, 21, 7], executable programs [34, 56, 61] or value function [23, 35] for low-level behaviours, environmental reward and feedback for reinforcement learning [38, 29, 40, 68]. However, many of these approaches necessitate the conversion of both the task and the observations into the textual form. While this can be easily accomplished in simulated environments with ground-truth object states, tasks in the real world require the utilization of robust and precise perception modules. For performing open-ended navigation and manipulation in the real world [57, 7, 16], open-world VLMs [27, 70, 30, 36, 44, 33] are often used to extract visual scene information before converting the observation to textual form. However, the process of converting an image into text may result in the loss of important details, such as shape and geometric information. With the recent advancements in VLMs, it is now conceivable that the role of state estimation combined with language models can be replaced by a single, powerful VLM, such as GPT-4V, which is leveraged by MOKA. Hu et al. [20] utilizes GPT-4V to get semantic language plans and convert each language plan to a set of pre-defined low-level skills. Different from Hu et al. [20], we represent affordances and motions using a set of points, and query GPT-4V to select the corresponding points (e.g. grasping point, function point) in each task, which can be directly translated to low-level point-based motions. While preserving the visual reasoning capabilities of VLMs, MOKA provides a framework with more general and flexible low-level motion, which doesn't require prior knowledge about a specific task compared to pre-defined skills.

Affordance reasoning for robotic control. Affordance [17] refers to the ability to perform a certain action with an object in a given environment. Much of the research on affordances focuses on understanding interactions with objects, as demonstrated by [58, 18, 3]. In the field of robot manipulation, keypoints and other representation forms are often used to provide compact information about the environment and the objects [8, 15, 59, 41, 43, 11, 42, 49, 4, 25, 26, 13, 14], representing the affordance in a structured way. Among these works, the most related one to ours is KETO [49], which predicts affordance and functional keypoints on the tool objects by learning from interactions with a trajectory optimization formulation as in Manuelli et al. [42]. In contrast to KETO [49], our keypoints selection procedure doesn't require any model training. We design an automated process to annotate keypoints as visual marks on a 2D image and leverage the broad knowledge from GPT-4V to select affordance and motion keypoints for manipulation.

III. PROBLEM STATEMENT

Our goal is to enable robots to perform manipulation tasks involving unseen objects and goals. Each task is described by a free-form language instruction l . To achieve the success, the robot needs to interact with the environment across one or multiple stages. An example task is shown in Fig. 1, in which the robot is commanded to clean the table by the instruction “Wipe the trash off the table after moving away the eyeglasses.”.

The physical interaction at each stage is referred to as a *subtask*. Each subtask can be an interaction with objects in hand (e.g., lifting up an object, opening a drawer, turning the faucet), an interaction with environmental objects unattached to the robot (e.g., pushing an obstacle, pressing a button), and tool use involving grasping a tool object to make contact with another object (e.g., poking with a stick, sweeping with

a foam, cutting with a knife). In this example above, the robot needs to perform two subtasks: First, put the eyeglasses between the broom and the snack package back into the case; and second, use the broom to sweep the trash. Since we do not assume the sequence of subtasks to be given beforehand, the robot needs to decompose the task into a sequence of feasible subtasks based on the free-form language instruction l and solve each subtask sequentially.

We consider a table-top setting with a robotic arm and RGBD cameras, as shown in Fig. 2. At each time step t , an observation s_t is received from the environment and an action a_t is selected to command the robot. The observation s_t consists of the RGBD images captured by the camera sensors as well as the proprioception of the robot. The action a_t is defined as the 6-DoF gripper pose and the finger status.

In this paper, we aim to solve such tasks by leveraging vision-language models (VLMs) with an effective visual prompting strategy. We use $\mathcal{M}$ to denote the VLM, which can take a list of language and visual inputs in a specific order, and generate text responses accordingly. The responses of $\mathcal{M}$ can be controlled by designing different input prompts, which specify information such as the problem descriptions, the input and output formats, and in-context examples [23, 34, 38, 56]. Specifically, we design text and visual prompts to enable $\mathcal{M}$ to generate desired text responses in a structured format, which can be parsed into point coordinates and other attributes to characterize the robot motion for the downstream usage.

IV. MARKING OPEN-WORLD KEYPOINT AFFORDANCES

We propose Marking Open-world Keypoint Affordances (MOKA), an approach that leverages the emergent reasoning capability of Vision-Language Models (VLMs) to guide a motion planner to solve unseen tasks specified by free-form language instructions (see Fig. 2). In this section, we will start by introducing the point-based affordance representations and how they are used for generating motions. Then, we will describe a hierarchical framework that effectively prompts VLMs for affordance reasoning. After this, we will explain a novel visual prompting approach that lifts the VLM’s reasoning capabilities by annotating a set of marks, including candidate points, grids, and labels, on 2D images. Lastly, we will investigate two ways to bootstrap the performance of MOKA by utilizing successful experiences collected in the physical world.

A. Point-Based Affordance Representations

To leverage VLMs for solving open-world manipulation tasks, there needs to be an interface that connects the predictions by the VLM and the motions performed by the robot. To achieve this goal, we design an affordance representation defined on 2D images. Produced as the end result by the VLM, the affordance representation specifies the desired motion.

By extending the definitions in Manuelli et al. [42] and Qin et al. [49], we design a point-based affordance representation that can cover a wide range of manipulation tasks. Instead of separately devising motion primitives for different pre-defined

skills, we use a unified set of keypoints and waypoints to specify the motion. These points are predicted by VLMs on 2D images and converted to poses in the $SE(3)$ space. Then a smooth motion trajectory is generated based on these poses. To perform the task, the robot interacts with the environment by following the generated motion trajectory.

We specify the robot’s motion in an object-centric manner, as shown in Fig. 2. Following the discussion in Sec. III, we would like this representation to apply to different types of interactions with objects in the environment. Therefore, we consider two types of objects, $o_{\text{in-hand}}$ (e.g., the broom) and $o_{\text{unattached}}$ (e.g., the trash), and specify the motion with a grasping phase and a manipulation phase. In the grasping phase, the robot reaches and grasps an object $o_{\text{in-hand}}$ from the environment. Then in the manipulation phase, the robot performs a motion and makes contact with another object $o_{\text{unattached}}$, either directly or using $o_{\text{in-hand}}$ as a tool. In some scenarios, only one of these two types of objects is interacted with by the robot, either $o_{\text{in-hand}}$ (e.g., unplugging a cable, opening a drawer) or $o_{\text{unattached}}$ (e.g., pressing a button), and one of the two phases can be skipped accordingly.

We now describe the definition of the keypoints and waypoints as well as how they are used to specify the motions in both phases. These points are illustrated in Fig. 2, and more examples can be found in Sec. V-C. Following the practice of Manuelli et al. [42] and Qin et al. [49], we use the **grasping keypoint** x_{grasp} to specify the position on $o_{\text{in-hand}}$ where the robot gripper should hold the object. If $o_{\text{in-hand}}$ is not involved in a task, the grasping phase will be skipped. For the manipulation phase, the robot’s gripper follows a motion trajectory specified by an additional set of points. The **function keypoint** x_{function} specifies the part of $o_{\text{in-hand}}$ that will make contact with $o_{\text{unattached}}$ in the manipulation phase. If $o_{\text{in-hand}}$ is not specified, x_{function} will be on the robot gripper, and the contact will directly be made between the robot and $o_{\text{unattached}}$. Correspondingly, the **target keypoint** x_{target} is the part of $o_{\text{unattached}}$ that will be contacted by x_{function} during the manipulation phase. We also introduce the **pre-contact waypoints** $x_{\text{pre-contact}}$ and the **post-contact waypoints** $x_{\text{post-contact}}$ defined in free space, which dictates the manipulation motion along with the keypoints defined on the objects.

During the manipulation phase, the robot moves the gripper such that x_{function} follows the path that sequentially connects the $x_{\text{pre-contact}}$, x_{target} , and $x_{\text{post-contact}}$. Besides following the path, we also require the robot gripper to follow the specified **grasping orientation** R_{grasp} and **manipulation orientation** $R_{\text{manipulate}}$ during the two phases, respectively. To better illustrate the design of our point-based motion, we provide examples of the predicted point specifications and the resultant motions from our experiments in Sec. V-C and the Appendix. In MOKA, this set of keypoints and **additional attributes** (described in Sec. IV-C) are summarized in a dictionary as the affordance representation (see Fig. 2).

B. Affordance Reasoning with Vision-Language Models

To predict the defined affordance representations, we employ the VLM $\mathcal{M}(\cdot)$ (defined in Sec. III), which is pre-trained on Internet-scale data for solving general visual question-answering (VQA) problems. Using a hierarchical prompting framework as shown in Fig. 2, MOKA converts affordance reasoning into a series of VQA problems that are solvable by the pre-trained VLM.

The hierarchical prompting framework takes as input the free-form language description l of the task and an RGB image observation of the environment s_t . MOKA examines the initial observation s_0 and decomposes the task l into a sequence of subtasks using the VLM. For each of the subtasks, the VLM is asked to provide the summary of the subtask instruction, the description of the corresponding $o_{\text{in-hand}}$, $o_{\text{unattached}}$, as well as the description of the motion (e.g., “from left to right”). On the low level, given the response from the high-level reasoning and the visual observation $s_{t(k)}$ at the beginning of k -th subtask (at the time step $t(k)$), the VLM is queried again with a different prompt to produce the affordance representation defined in Sec. IV-A. In the remainder of this section, we will describe the input formats, output formats, and prompt designs that we use to instantiate this method. Further details can be found in the Appendix.

High-level reasoning. Given the initial observation s_0 and the language description l , we first query the VLM $\mathcal{M}$ with the language prompt ρ_{high} to produce the response y_{high} :

$$y_{\text{high}} = \mathcal{M}([\rho_{\text{high}}, l, s_0]). \quad (1)$$

The representation y_{high} is a string that contains structured information for the K subtasks that are needed to solve the task. We design the prompt so as to require the VLM to produce y_{high} as a list of dictionaries. As shown in Fig. 2, each dictionary contains the language description of a subtask (e.g., “Wipe the snack package to the right side of the table using the broom.”), as well as detailed information to facilitate motion generation, including the description of $o_{\text{in-hand}}$ (e.g., “broom”), the object name of $o_{\text{unattached}}$ (e.g., “snack package”), and the description of the motion (e.g., “from left to right”). This high-level plan will be used as an intermediate result for producing the detailed affordance representation through the low-level reasoning with the VLM. The prompt ρ_{high} used for high-level reasoning contains the description of the request (the input instruction and the annotated image) and the description of the response format. The complete high-level reasoning prompts are described in TABLE V.

Low-level reasoning. Next, we prompt the VLM once again to produce the affordance representation defined in Sec. IV-A as y_{low}^k , conditioning on the high-level representation y_{high} and the visual observation $s_{t(k)}$ at the beginning of the k -th subtask. Instead of directly predicting 3D coordinates on 2D images, which is challenging and even ill-defined, we query the VLM to output 2D coordinates on the images and deproject them back to the 3D space. The three keypoints x_{grasp} , x_{function} and x_{target} are defined on the object surface, and thus we can

compute the 3D coordinates using the corresponding depth value of the 2D location based on the RGB image and camera parameters. For the waypoints in free space, we query the VLM to predict the desired height in the text. To produce such an affordance representation y_{low}^k , we query the VLM again by

$$y_{\text{low}}^k = \mathcal{M}([\rho_{\text{low}}, y_{\text{high}}^k, f(s_{t(k)})]), \quad (2)$$

where y_{high}^k is the substring corresponding to the k -th subtask extracted from y_{high} , and $f(\cdot)$ is a function that process the raw visual observation $s_{t(k)}$. The prompt ρ_{low} used for low-level reasoning contains the description of the request (the response format of the high-level reasoning and the annotated image) and the description of the response (explanation of the point-based affordance representation and the desired format of the response). The complete low-level reasoning prompts are described in TABLE V, TABLE VI, and TABLE VII.

C. Mark-Based Visual Prompting

To perform the low-level reasoning mentioned in the previous section, we need the VLM to generate keypoints and waypoints on 2D images in order to execute a specific motion for a subtask. Since VLMs are better at multiple-choice problems than directly producing continuous-valued locations, we employ a mark-based visual prompting strategy to extract the desired output from VLMs, which we will describe in this subsection.

Inspired by Yang et al. [64], MOKA uses a set of marks as visual prompts to enable VLM to apply its reasoning capability to predict the point-based affordance representation as shown in Fig. 2. Consisting of dots, grids, and text notations annotated on the visual observation, these marks play an important role in the reasoning process. Proposed by open-world object detection and segmentation algorithms, these marks facilitate visual reasoning by encouraging the VLM to attend to the target objects and other task-relevant information in the image. We annotate marks as candidate parts and regions for the VLM to choose the points from, converting the original problem of directly generating coordinates into multiple-choice questions, which is usually more tractable for existing VLMs.

To select the keypoints, which are defined on the in-hand object $o_{\text{in-hand}}$ and the unattached object $o_{\text{unattached}}$ suggested by the high-level reasoning in Sec. IV-B, we propose and plot candidate keypoints on these objects. Given the names of $o_{\text{in-hand}}$ and $o_{\text{unattached}}$, we first segment these two objects using GroundedSAM [54], which combines GroundingDINO [36] and SAM [70] to extract segmentation masks of objects specified by a text prompt. After we obtain the segmentation masks of $o_{\text{in-hand}}$ and $o_{\text{unattached}}$, we perform farthest point sampling [48] on the object contour to obtain K boundary points. Together with its geometric center, and overlay the $K + 1$ candidate keypoints on each object. Each candidate keypoint is assigned an index, which is annotated next to it as a reference. To avoid confusion, we use different colors for candidate keypoints on $o_{\text{in-hand}}$ and $o_{\text{unattached}}$ and use the caption in the format of P_i and Q_j respectively, where i and

j are integers. More implementation details can be found in Appendix B-B.

Selecting waypoints in free space involves searching over a much larger region. Instead of directly sampling points in the entire workspace, we divide the observed RGB image into an $m \times n$ grid, where m and n are integers. Both m and n are set to 5 for our evaluation tasks. The VLM is prompted to choose the tiles in which the pre-contact and post-contact keypoints are supposed to locate in, and then the exact waypoints are sampled uniformly within the tile. For this purpose, we overlay the grid along with the name of each tile on the image. The tile names follow chess notation, which uses letters to specify the columns and integers for the rows.

D. Motion Generation with Predicted Affordances

After obtaining the selected keypoints and waypoints from the VLM, we convert the affordance to an executable motion on a real robot. To achieve this goal, we need to lift up all the points from the 2D image to the $SE(3)$ space.

For the keypoints defined on the object, we can directly take the corresponding point on the registered depth image and transform it into the robot’s frame. For the waypoints that are selected from the free space, since they are not attached to any objects, their height needs to be specified in order to be deprojected to 3D space. For this purpose, we also ask the VLM to determine the height of waypoints. We only consider the cases where the waypoints are at the same height as the target point (e.g., pushing) or above the target point (e.g., placing) in most common tabletop manipulation scenarios. Additionally, we also prompt the VLM to return the orientation of $o_{\text{in-hand}}$ during the manipulation phase, which would usually affect the success of tool use tasks.

Drawing inspirations from Manuelli et al. [42], we use the vector from x_{grasp} to x_{function} to specify the orientation of the object and ask the VLM to predict the orientation from a finite set of options (e.g., forward, backward, upside, downside, left, right).

Since robust grasping relies on contact physics and gripper design, which is beyond the capability of existing pre-trained VLMs, we combine the VLM predictions with analytical approaches in robotic grasping pipelines. Similar to Manuelli et al. [42], we use a grasp sampler to propose candidate grasps based on local geometric information from the observed point cloud. Instead of directly relying on the predicted x_{grasp} , we use the position and orientation of the grasp candidate that is closest to x_{grasp} . More implementation details can be found in the Appendix.

E. Bootstrapping through Physical Interactions

We consider two ways to bootstrap the performance of MOKA through physical interactions with the real world. In both ways, we unroll the VLM policy for the target task to collect robot experiences. The success labels of the collected experiences will be annotated by the VLM or humans.

In-context learning. We use a handful of successful trajectories as in-context examples (as shown in Fig. 7) to guide the

high-level and low-level reasoning of the VLM. As discussed in prior work [6], two to three such in-context examples can significantly boost the performance of the VLM. Without changing the prompts ρ_{high} and ρ_{low} , we append the prompt with three pairs of annotated images and VLM responses from previous successful rollouts by the VLM policy.

Policy distillation. MOKA can be used for collecting demonstration datasets for real-world robot learning. We record the multi-view images and robot proprioception states when collecting successful trajectories and train a visuomotor policy using a simple behavior cloning objective. More details can be found in Sec. V-A.

V. EXPERIMENTS

The primary goal of our experiments is to validate and analyze the behavior of MOKA on open-world tasks with a wide range of objects and skills. We design our experiments to investigate the following questions:

- Can MOKA effectively reason about affordances based on the images for unseen tasks?
- After translating the prediction of MOKA into low-level motion, how well does it perform on the task we proposed?
- Can MOKA improve from real-world trials, either via in-context learning or policy distillation?

We first evaluate MOKA’s performance against baseline methods on 4 manipulation tasks in both zero-shot and in-context learning settings. Then we provide qualitative and robustness analysis in Sec. V-C.

Each of our evaluation tasks has two stages, and each task involves a variety of object interaction and tool-use scenarios. A summary of our evaluation tasks can be found in TABLE II. We also tested additional open-world tasks to demonstrate our method, which can be found in Appendix C.

A. Experimental Setup

We compare MOKA with **Code-as-Policies** [34] and **VoxPoser** [23], two baselines that also enable zero-shot execution of open-world tasks. Code-as-Policies provides a framework for language model-generated programs executed on robotic systems by prompting with code examples. For a fair comparison, we provide the two baselines with the task description in the code comments, with an additional 40 lines of code prompts providing example usage, as in the original implementation [34]. Similarly, VoxPoser [23] also provides code examples to large language models to build a 3D voxel map of value functions. For VoxPoser, we reuse the example prompts and planning pipeline in the original implementation and use the same hyper-parameters to create the voxel value map. For both baselines, we only adapt the perception modules for fair comparisons while retaining the functionality of the other components.

We evaluate MOKA in both zero-shot and in-context learning settings and refer to them as **MOKA zero-shot** and **MOKA in-context**, respectively. For the in-context learning setting, we provide two examples to GPT-4V with instructions

Methods	Table Wiping		Watch Cleaning		Gift Preparation		Laptop Packing	
	Subtask I	Subtask II	Subtask I	Subtask II	Subtask I	Subtask II	Subtask I	Subtask II
Code-as-Policies [34]	0.7	0.6	0.6	1.0	1.0	0.7	0.4	0.8
VoxPoser [23]	0.6	0	0.6	0.8	1.0	0.6	0.5	0.8
MOKA Zero-Shot (Ours)	0.6	0.6	0.7	1.0	1.0	0.7	0.5	0.8
MOKA Distilled (Ours)	1.0	0.7	0.8	0.8	1.0	0.7	1.0	1.0
MOKA In-Context (Ours)	0.9	0.9	0.9	1.0	1.0	0.9	1.0	0.9

TABLE I: Success rate of our method and baselines. Across four tasks, each consists of 2 subtasks, MOKA consistently achieves superior performances. We also demonstrated that the performance can be further bootstrapped through physical interactions in the environment using policy distillation and in-context learning.

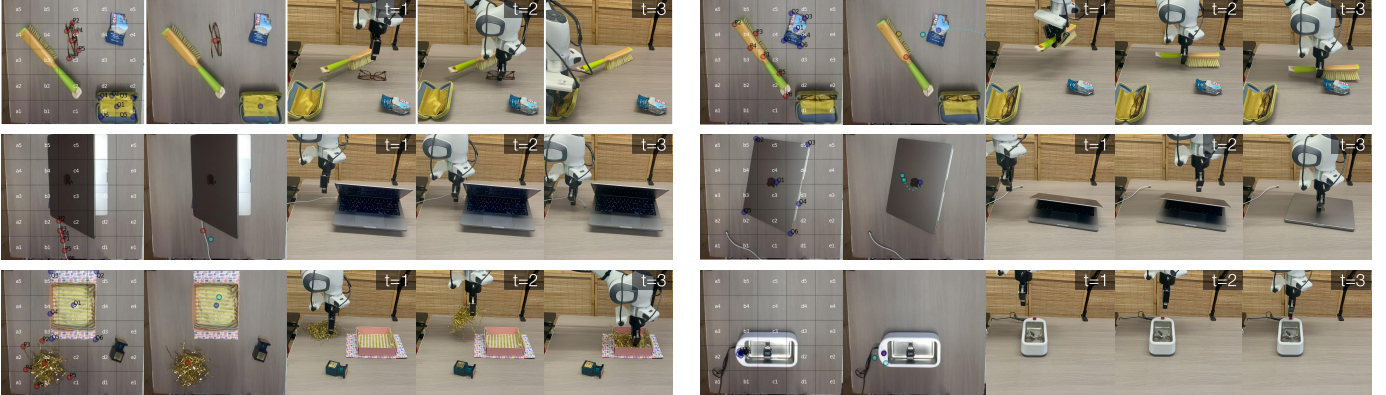


Fig. 3: Example results of the predicted keypoints and motions. On each row, two individual subtasks are displayed. For each subtask, we plot the marked image, the predicted point-based affordances, and three keyframes in the trajectory.

Table Wiping	1. Move away the glasses
	2. Sweep the trash with the broom
Watch Cleaning	1. Put the watch into the ultrasound cleaner
	2. Press the power button
Gift Preparation	1. Put the golden filler in the gift box
	2. Put the perfume in the gift box
Laptop Packing	1. Unplug the cable
	2. Close the lid of the laptop

TABLE II: The subtask instructions of each of the testing tasks. Each of the testing tasks consists of two subtasks.

and annotated images before querying information about the current task. The annotated images are collected in scenes with object and instruction variations. Some in-context prompt examples can be found in the Appendix.

The successful trajectories generated by MOKA can also serve as demonstration data for other learning-based methods. By simply training a model using the successful trajectories generated from MOKA, we can also “distill” the knowledge from a VLM to a learned policy. To achieve this goal, we employ a recent open-sourced robot foundation model Octo [46], which is pre-trained on a diverse mix of 800K robot trajectories [47] and can be easily fine-tuned to new tasks. Octo [46] supports language instructions or goal images as task specifications and takes current observations to generate actions with a transformer-based diffusion policy architecture.

It also provides code examples to easily perform fine-tuning. To construct our fine-tuning dataset, we collect 50 successful trajectories for each task with language annotations. We then adopt its latest base model checkpoint¹ and finetune the full model with the recipe provided in the official instruction². The performance of fine-tuned can be found in Sec. V-B with more implementation details in Appendix B-E. Although MOKA is a training-free method, it can be leveraged for data-driven approaches to provide training data for open-world tasks. We refer to this approach as MOKA-Distilled.

B. Quantitative Evaluation

Our quantitative evaluation results across 4 tasks are illustrated in TABLE I. For each task, we report the number of successes out of 10 trials. As shown in the table, MOKA achieves state-of-the-art performance at each subtask of the 4 tasks (total 8 subtasks), with consistent improvements using in-context learning. On most of the tasks, VoxPoser [23] has similar performance with MOKA (zero-shot), except for subtask 2 of table wiping (which is a tool-use task). Additionally, the task success rates can be sensitive to the resolution of the voxel map, which requires some hyperparameter tuning. Unlike the baselines, MOKA can work well without example prompts

¹<https://huggingface.co/rail-berkeley/octo-base>

²<https://github.com/octo-models/octo>

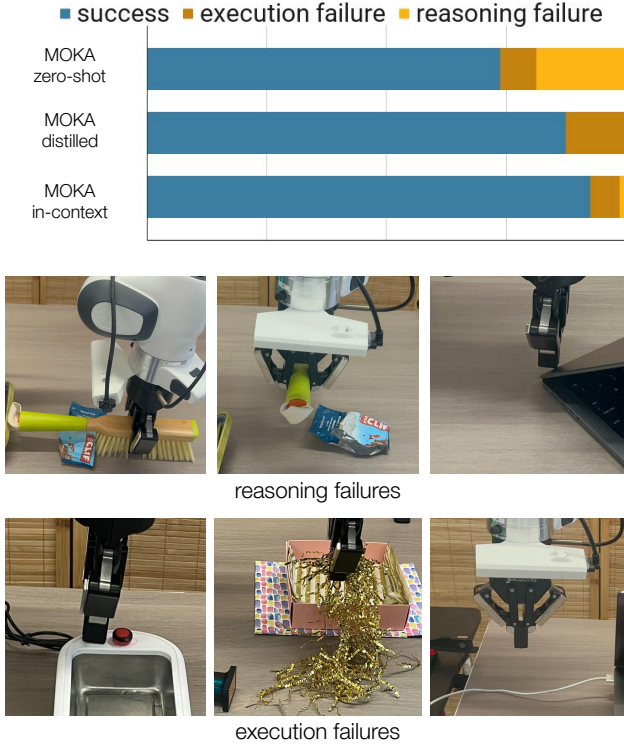


Fig. 4: **Failure analysis.** We break down the failure cases of the variants of our approach using zero-shot prediction, distilled policies, and in-context learning. We break down the failures into reasoning failures (caused by errors in the affordance prediction) and execution failures (caused by low-level motions). We demonstrate how policy distillation and in-context learning reduce failure cases.

or can achieve better performance with just two clean and intuitive example prompts.

We also observed that we can obtain an effective end-to-end policy by fine-tuning a pre-trained robot policy [46] using the successful trajectories generated by MOKA. This suggests that the generated data is of high quality and thus can be used as demonstration data for open-world tasks. The fact that distilling the successful MOKA trials actually *improves* the overall performance of the system suggests the feasibility of using MOKA as a “data generator”, where the zero-shot MOKA method is used to bootstrap a continuous improvement system. This is an exciting direction to explore in future work.

Failure breakdown We analyze the failure cases of MOKA. Trajectories with wrong predictions from the GPT-4V are counted as reasoning failures. The following failure cases are illustrated in the top row of Fig. 4, including grasping the broom upside down due to confusing the grasp point with the function point, pointing the room in the wrong direction due to the wrong target angle, pressing the hinge of the laptop due to the wrong target point. The trajectories with desired VLM prediction but fail to execute successfully are counted as execution failures. The following failure cases are illustrated in the bottom row of Fig. 4. From left to right, the failures

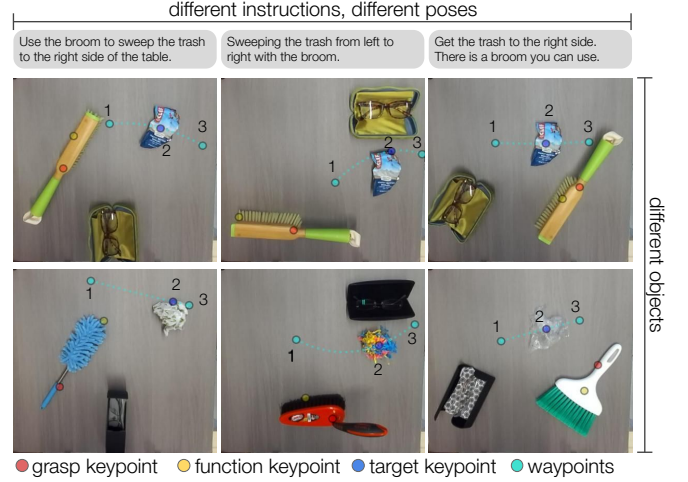


Fig. 5: **Robustness analysis.** We analyze MOKA’s robustness with respect to various instructions, initial arrangements, and objects. Each column in the image uses the same language instruction and similar initial arrangements of objects. The two rows involve different objects. MOKA demonstrates to robustly predicts the keypoints and waypoints across a wide range of variations. More examples of other tasks are presented in the Appendix.

include the gripper narrowly missing the button, the filler getting disassembled in the middle of the trajectory, and the cable slipping through the gripper fingers.

As shown in the figure, both policy distillation and in-context learning reduce the total failures. Using the distilled policy, the VLM is no longer part of the pipeline. Therefore, there are no reasoning failures in this case.

C. Qualitative Evaluation

In Fig. 3, we provide six examples of the annotated marks, generated motions, and the unrolled trajectories of MOKA. In each of these examples, based on the annotated marks, MOKA successfully selected the keypoints and waypoints for generating the desired motions. These examples demonstrate that our approach can be applied to various skills, including pressing, closing, rearranging, tool use, etc.

We analyze the robustness of MOKA with respect to instruction variations, initialization variations, and object variations. Fig. 5 provides a pictorial illustration of MOKA’s predictions under different instructions and observations. Each image is queried with the language prompt on the top within the same column. The examples in the same column share similar initialization object positions and orientations. The example in the first row uses the same set of objects, while the second row involves objects of alternative geometries, colors, and materials. Some of these objects are also deformable or transparent. Fig. 5 demonstrates consistent keypoint and waypoint predictions within rows and columns, indicating that MOKA is robust to the changes of the language instructions, objects, and initial arrangements for the same task.

VI. CONCLUSION AND DISCUSSION

In this paper, we propose MOKA, a simple and effective approach that leverages pre-trained VLMs for open-world robotic manipulation. By specifying motions with point-based affordance representations, we enable VLMs pre-trained on Internet-scale data to solve manipulation tasks in zero-shot and few-shot manners. Using the proposed mark-based visual prompting technique, MOKA converts affordance reasoning into a series of visual question-answering problems that the pre-trained VLMs can solve. Through experiments on a variety of table-top manipulation tasks, we demonstrate the effectiveness and robustness of MOKA. We also demonstrate that the trajectories generated by MOKA can be further used to bootstrap the performance through in-context learning and policy distillation. As far as we know, MOKA is the very first work that leverages visual prompting on pre-trained VLMs for open-world robot manipulation. We hope MOKA will inspire future research towards leveraging vision-language models for open-world robotic control.

MOKA is subject to limitations mostly due to the capabilities of existing VLMs and the current design of the affordance representations. The state-of-the-art VLMs still lack a profound understanding of 3D space, contact physics, and robotic control, which prevents our approach from generating more fine-grained motions in the $SE(3)$ space. The high latency of querying VLMs also makes it challenging for the robot to perform dynamic manipulation and open-loop control. Moreover, we would like to note that, while we demonstrate that a unified point-based affordance representation can be devised to support diverse tasks, the specific design of the representation in this work does not cover all useful skills that a robot can perform. To apply MOKA to more complex scenarios, e.g., bi-manual manipulation and whole-body control, we would need to extend the set of points and include additional attributes in this affordance representation. We believe these limitations will be addressed as more capable vision-language models and robotic foundation models are introduced in the future.

ACKNOWLEDGEMENT

This research was supported by the AI Institute, AFOSR FA9550-22-1-0273, and the NSF under IIS-2246811. Fangchen Liu was supported in part by NSF under the NSF AI4OPT Center. We would like to thank Karl Pertsch, Oleg Rybkin, and Zheyuan Hu for providing feedback on early drafts of the manuscript. We also want to thank Sudeep Dasari, Qiyang Li, Yide Shentu, and Homer Walke for valuable suggestions on the experiments and infrastructure.

REFERENCES

- [1] Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altenschmidt, Sam Altman, Shyamal Anadkat, et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- [2] Michael Ahn, Anthony Brohan, Noah Brown, Yevgen Chebotar, Omar Cortes, Byron David, Chelsea Finn, Chuyuan Fu, Keerthana Gopalakrishnan, Karol Hausman, et al. Do as i can, not as i say: Grounding language in robotic affordances. *arXiv preprint arXiv:2204.01691*, 2022.
- [3] Aitor Aldoma, Federico Tomba, and Markus Vincze. Supervised learning of hidden and non-hidden 0-order affordances and detection in real scenes. In *IEEE International Conference on Robotics and Automation*, pages 1732–1739. IEEE, 2012.
- [4] Shikhar Bahl, Russell Mendonca, Lili Chen, Unnat Jain, and Deepak Pathak. Affordances from human videos as a versatile representation for robotics. *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2023.
- [5] Anthony Brohan, Noah Brown, Justice Carbajal, Yevgen Chebotar, Xi Chen, Krzysztof Choromanski, Tianli Ding, Danny Driess, Avinava Dubey, Chelsea Finn, et al. Rt-2: Vision-language-action models transfer web knowledge to robotic control. *arXiv preprint arXiv:2307.15818*, 2023.
- [6] Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Nee-lakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33:1877–1901, 2020.
- [7] Boyuan Chen, Fei Xia, Brian Ichter, Kanishka Rao, Keerthana Gopalakrishnan, Michael S Ryoo, Austin Stone, and Daniel Kappler. Open-vocabulary queryable scene representations for real world planning. In *IEEE International Conference on Robotics and Automation*, pages 11509–11522. IEEE, 2023.
- [8] Changhyun Choi and Henrik I Christensen. Real-time 3d model-based tracking using edge and keypoint features for robotic manipulation. In *IEEE International Conference on Robotics and Automation*, pages 4048–4055. IEEE, 2010.
- [9] Aakanksha Chowdhery, Sharan Narang, Jacob Devlin, Maarten Bosma, Gaurav Mishra, Adam Roberts, Paul Barham, Hyung Won Chung, Charles Sutton, Sebastian Gehrmann, et al. Palm: Scaling language modeling with pathways. *Journal of Machine Learning Research*, 24 (240):1–113, 2023.
- [10] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*, 2018.
- [11] Kuan Fang, Te-Lin Wu, Daniel Yang, Silvio Savarese, and Lim J. Joseph. Demo2vec: Reasoning object affordances from online videos. *IEEE Conference on Computer Vision and Pattern Recognition*, 2018.
- [12] Kuan Fang, Alexander Toshev, Li Fei-Fei, and Silvio Savarese. Scene memory transformer for embodied agents in long-horizon tasks. *IEEE Conference on Computer Vision and Pattern Recognition*, 2019.
- [13] Kuan Fang, Patrick Yin, Ashvin Nair, and Sergey Levine.

- Planning to practice: Efficient online fine-tuning by composing goals in latent space. *IEEE/RSJ International Conference on Intelligent Robots and Systems*, 2022.
- [14] Kuan Fang, Patrick Yin, Ashvin Nair, Homer Walke, Gengchen Yan, and Sergey Levine. Generalization with lossy affordances: Leveraging broad offline data for learning visuomotor tasks. *Conference on Robot Learning (CoRL)*, 2022.
- [15] Chelsea Finn and Sergey Levine. Deep visual foresight for planning robot motion. In *IEEE International Conference on Robotics and Automation*, pages 2786–2793. IEEE, 2017.
- [16] Samir Yitzhak Gadre, Mitchell Wortsman, Gabriel Ilharco, Ludwig Schmidt, and Shuran Song. Cows on pasture: Baselines and benchmarks for language-driven zero-shot object navigation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 23171–23181, 2023.
- [17] James J Gibson. *The Ecological Approach to Visual Perception: Classic Edition*. Psychology Press, 2014.
- [18] Tucker Hermans, James M Rehg, and Aaron Bobick. Affordance prediction via learned object attributes. In *IEEE International Conference on Robotics and Automation: Workshop on Semantic Perception, Mapping, and Exploration*, pages 181–184. Citeseer, 2011.
- [19] Jonathan Ho, Ajay Jain, and Pieter Abbeel. Denoising diffusion probabilistic models. *Advances in Neural Information Processing Systems*, 33:6840–6851, 2020.
- [20] Yingdong Hu, Fanqi Lin, Tong Zhang, Li Yi, and Yang Gao. Look before you leap: Unveiling the power of gpt-4v in robotic vision-language planning. *arXiv preprint arXiv:2311.17842*, 2023.
- [21] Wenlong Huang, Pieter Abbeel, Deepak Pathak, and Igor Mordatch. Language models as zero-shot planners: Extracting actionable knowledge for embodied agents. In *International Conference on Machine Learning*, pages 9118–9147. PMLR, 2022.
- [22] Wenlong Huang, Fei Xia, Ted Xiao, Harris Chan, Jacky Liang, Pete Florence, Andy Zeng, Jonathan Tompson, Igor Mordatch, Yevgen Chebotar, et al. Inner monologue: Embodied reasoning through planning with language models. *arXiv preprint arXiv:2207.05608*, 2022.
- [23] Wenlong Huang, Chen Wang, Ruohan Zhang, Yunzhu Li, Jiajun Wu, and Li Fei-Fei. Voxposer: Composable 3d value maps for robotic manipulation with language models. *arXiv preprint arXiv:2307.05973*, 2023.
- [24] Eric Jang, Alex Irpan, Mohi Khansari, Daniel Kappler, Frederik Ebert, Corey Lynch, Sergey Levine, and Chelsea Finn. BC-Z: Zero-Shot Task Generalization with Robotic Imitation Learning. *Conference on Robot Learning*, page 12, 2021. doi: 164:991-1002.
- [25] Zhenyu Jiang, Yifeng Zhu, Maxwell Svetlik, Kuan Fang, and Yuke Zhu. Synergies between affordance and geometry: 6-dof grasp detection via implicit representations. *Robotics: Science and Systems*, 2021.
- [26] Alexander Khazatsky, Ashvin Nair, Dan Jing, and Sergey Levine. What Can I Do Here? Learning New Skills by Imagining Visual Affordances. *IEEE International Conference on Robotics and Automation*, 2021.
- [27] Alexander Kirillov, Eric Mintun, Nikhila Ravi, Hanzi Mao, Chloe Rolland, Laura Gustafson, Tete Xiao, Spencer Whitehead, Alexander C Berg, Wan-Yen Lo, et al. Segment anything. *arXiv preprint arXiv:2304.02643*, 2023.
- [28] Takeshi Kojima, Shixiang Shane Gu, Machel Reid, Yutaka Matsuo, and Yusuke Iwasawa. Large language models are zero-shot reasoners. *Advances in Neural Information Processing Systems*, 35:22199–22213, 2022.
- [29] Minae Kwon, Sang Michael Xie, Kalesha Bullard, and Dorsa Sadigh. Reward design with language models. *arXiv preprint arXiv:2303.00001*, 2023.
- [30] Feng Li, Hao Zhang, Peize Sun, Xueyan Zou, Shilong Liu, Jianwei Yang, Chunyuan Li, Lei Zhang, and Jianfeng Gao. Semantic-sam: Segment and recognize anything at any granularity. *arXiv preprint arXiv:2307.04767*, 2023.
- [31] Junnan Li, Dongxu Li, Caiming Xiong, and Steven Hoi. Blip: Bootstrapping language-image pre-training for unified vision-language understanding and generation. In *International Conference on Machine Learning*, pages 12888–12900. PMLR, 2022.
- [32] Junnan Li, Dongxu Li, Silvio Savarese, and Steven Hoi. Blip-2: Bootstrapping language-image pre-training with frozen image encoders and large language models. *arXiv preprint arXiv:2301.12597*, 2023.
- [33] Liunan Harold Li, Pengchuan Zhang, Haotian Zhang, Jianwei Yang, Chunyuan Li, Yiwu Zhong, Lijuan Wang, Lu Yuan, Lei Zhang, Jenq-Neng Hwang, et al. Grounded language-image pre-training. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 10965–10975, 2022.
- [34] Jacky Liang, Wenlong Huang, Fei Xia, Peng Xu, Karol Hausman, Brian Ichter, Pete Florence, and Andy Zeng. Code as policies: Language model programs for embodied control. In *IEEE International Conference on Robotics and Automation*, pages 9493–9500. IEEE, 2023.
- [35] Kevin Lin, Christopher Agia, Toki Migimatsu, Marco Pavone, and Jeannette Bohg. Text2motion: From natural language instructions to feasible plans. *arXiv preprint arXiv:2303.12153*, 2023.
- [36] Shilong Liu, Zhaoyang Zeng, Tianhe Ren, Feng Li, Hao Zhang, Jie Yang, Chunyuan Li, Jianwei Yang, Hang Su, Jun Zhu, et al. Grounding dino: Marrying dino with grounded pre-training for open-set object detection. *arXiv preprint arXiv:2303.05499*, 2023.
- [37] Corey Lynch and Pierre Sermanet. Language Conditioned Imitation Learning over Unstructured Data. In *Robotics: Science and Systems*. arXiv, July 2021.
- [38] Yecheng Jason Ma, William Liang, Guanzhi Wang, De-An Huang, Osbert Bastani, Dinesh Jayaraman, Yuke Zhu, Linxi Fan, and Anima Anandkumar. Eureka: Human-level reward design via coding large language models.

arXiv preprint arXiv:2310.12931, 2023.

- [39] Jeffrey Mahler, Jacky Liang, Sherdil Niyaz, Michael Laskey, Richard Doan, Xinyu Liu, Juan Aparicio Ojea, and Ken Goldberg. Dex-net 2.0: Deep learning to plan robust grasps with synthetic point clouds and analytic grasp metrics. *arXiv preprint arXiv:1703.09312*, 2017.
- [40] Parsa Mahmoudieh, Deepak Pathak, and Trevor Darrell. Zero-shot reward specification via grounded natural language. In *International Conference on Machine Learning*, pages 14743–14752. PMLR, 2022.
- [41] Jeremy Maitin-Shepard, Marco Cusumano-Towner, Jinna Lei, and Pieter Abbeel. Cloth grasp point detection based on multiple-view geometric cues with application to robotic towel folding. In *IEEE International Conference on Robotics and Automation*, pages 2308–2315. IEEE, 2010.
- [42] Lucas Manuelli, Wei Gao, Peter R. Florence, and Russ Tedrake. kpm: Keypoint affordances for category-level robotic manipulation. In *International Symposium of Robotics Research*, 2019. URL <https://api.semanticscholar.org/CorpusID:80628296>.
- [43] Stephen Miller, Mario Fritz, Trevor Darrell, and Pieter Abbeel. Parametrized shape models for clothing. In *IEEE International Conference on Robotics and Automation*, pages 4861–4868. IEEE, 2011.
- [44] Matthias Minderer, Alexey Gritsenko, Austin Stone, Maxim Neumann, Dirk Weissenborn, Alexey Dosovitskiy, Aravindh Mahendran, Anurag Arnab, Mostafa Dehghani, Zhuoran Shen, et al. Simple open-vocabulary object detection. In *European Conference on Computer Vision*, pages 728–755. Springer, 2022.
- [45] Vivek Myers, Andre He, Kuan Fang, Homer Walke, Philippe Hansen-Estruch, Ching-An Cheng, Mihai Jalobeanu, Andrey Kolobov, Anca Dragan, and Sergey Levine. Goal representations for instruction following: A semi-supervised language interface to control. *arXiv*, 2023.
- [46] Octo Model Team, Dibya Ghosh, Homer Walke, Karl Pertsch, Kevin Black, Oier Mees, Sudeep Dasari, Joey Hejna, Charles Xu, Jianlan Luo, Tobias Kreiman, You Liang Tan, Dorsa Sadigh, Chelsea Finn, and Sergey Levine. Octo: An open-source generalist robot policy. <https://octo-models.github.io>, 2023.
- [47] Abhishek Padalkar, Acorn Pooley, Ajinkya Jain, Alex Bewley, Alex Herzog, Alex Irpan, Alexander Khazatsky, Anant Rai, Anikait Singh, Anthony Brohan, et al. Open x-embodiment: Robotic learning datasets and rt-x models. *arXiv preprint arXiv:2310.08864*, 2023.
- [48] Charles R Qi, Hao Su, Kaichun Mo, and Leonidas J Guibas. Pointnet: Deep learning on point sets for 3d classification and segmentation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 652–660, 2017.
- [49] Zengyi Qin, Kuan Fang, Yuke Zhu, Li Fei-Fei, and Silvio Savarese. Keto: Learning keypoint representations for tool manipulation, 2019.
- [50] Alec Radford, Karthik Narasimhan, Tim Salimans, Ilya Sutskever, et al. Improving language understanding by generative pre-training. 2018.
- [51] Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, Ilya Sutskever, et al. Language models are unsupervised multitask learners. *OpenAI Blog*, 1(8): 9, 2019.
- [52] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. Learning transferable visual models from natural language supervision. In *International Conference on Machine Learning*, pages 8748–8763. PMLR, 2021.
- [53] Aditya Ramesh, Mikhail Pavlov, Gabriel Goh, Scott Gray, Chelsea Voss, Alec Radford, Mark Chen, and Ilya Sutskever. Zero-shot text-to-image generation. In *International Conference on Machine Learning*, pages 8821–8831. PMLR, 2021.
- [54] Tianhe Ren, Shilong Liu, Ailing Zeng, Jing Lin, Kunchang Li, He Cao, Jiayu Chen, Xinyu Huang, Yukang Chen, Feng Yan, Zhaoyang Zeng, Hao Zhang, Feng Li, Jie Yang, Hongyang Li, Qing Jiang, and Lei Zhang. Grounded sam: Assembling open-world models for diverse visual tasks, 2024.
- [55] Yujun Shi, Chuhui Xue, Jiachun Pan, Wenqing Zhang, Vincent YF Tan, and Song Bai. Dragdiffusion: Harnessing diffusion models for interactive point-based image editing. *arXiv preprint arXiv:2306.14435*, 2023.
- [56] Ishika Singh, Valts Blukis, Arsalan Mousavian, Ankit Goyal, Danfei Xu, Jonathan Tremblay, Dieter Fox, Jesse Thomason, and Animesh Garg. Progprompt: Generating situated robot task plans using large language models. In *IEEE International Conference on Robotics and Automation*, pages 11523–11530. IEEE, 2023.
- [57] Austin Stone, Ted Xiao, Yao Lu, Keerthana Gopalakrishnan, Kuang-Huei Lee, Quan Vuong, Paul Wohlhart, Brianna Zitkovich, Fei Xia, Chelsea Finn, et al. Open-world object manipulation using pre-trained vision-language models. *arXiv preprint arXiv:2303.00905*, 2023.
- [58] Jie Sun, Joshua L Moore, Aaron Bobick, and James M Rehg. Learning visual object categories for robot affordance prediction. *The International Journal of Robotics Research*, 29(2-3):174–197, 2010.
- [59] Jur Van Den Berg, Stephen Miller, Ken Goldberg, and Pieter Abbeel. Gravity-based robotic cloth folding. In *Algorithmic Foundations of Robotics IX*, pages 409–424. Springer, 2010.
- [60] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, and Illia Polosukhin. Attention is all you need, 2023.
- [61] Guanzhi Wang, Yuqi Xie, Yunfan Jiang, Ajay Mandlekar, Chaowei Xiao, Yuke Zhu, Linxi Fan, and Anima Anandkumar. Voyager: An open-ended embodied agent with large language models. *arXiv preprint arXiv:2305.16291*, 2023.
- [62] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten

- Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, et al. Chain-of-thought prompting elicits reasoning in large language models. *Advances in Neural Information Processing Systems*, 35:24824–24837, 2022.
- [63] Jianwei Yang, Hao Zhang, Feng Li, Xueyan Zou, Chunyuan Li, and Jianfeng Gao. Set-of-mark prompting unleashes extraordinary visual grounding in gpt-4v. *arXiv preprint arXiv:2310.11441*, 2023.
- [64] Jianwei Yang, Hao Zhang, Feng Li, Xueyan Zou, Chunyue Li, and Jianfeng Gao. Set-of-mark prompting unleashes extraordinary visual grounding in gpt-4v. *ArXiv*, abs/2310.11441, 2023. URL <https://api.semanticscholar.org/CorpusID:266149987>.
- [65] Lingfeng Yang, Yueze Wang, Xiang Li, Xinlong Wang, and Jian Yang. Fine-grained visual prompting. *arXiv preprint arXiv:2306.04356*, 2023.
- [66] Zhengyuan Yang, Linjie Li, Kevin Lin, Jianfeng Wang, Chung-Ching Lin, Zicheng Liu, and Lijuan Wang. The dawn of lmms: Preliminary explorations with gpt-4v(ision). *arXiv preprint arXiv:2309.17421*, 9(1):1, 2023.
- [67] Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran, Thomas L Griffiths, Yuan Cao, and Karthik Narasimhan. Tree of thoughts: Deliberate problem solving with large language models. *arXiv preprint arXiv:2305.10601*, 2023.
- [68] Wenhao Yu, Nimrod Gileadi, Chuyuan Fu, Sean Kirmani, Kuang-Huei Lee, Montse Gonzalez Arenas, Hao-Tien Lewis Chiang, Tom Erez, Leonard Hasenclever, Jan Humplik, et al. Language to rewards for robotic skill synthesis. *arXiv preprint arXiv:2306.08647*, 2023.
- [69] Denny Zhou, Nathanael Schärli, Le Hou, Jason Wei, Nathan Scales, Xuezhi Wang, Dale Schuurmans, Claire Cui, Olivier Bousquet, Quoc Le, et al. Least-to-most prompting enables complex reasoning in large language models. *arXiv preprint arXiv:2205.10625*, 2022.
- [70] Xueyan Zou, Jianwei Yang, Hao Zhang, Feng Li, Linjie Li, Jianfeng Gao, and Yong Jae Lee. Segment everything everywhere all at once. *arXiv preprint arXiv:2304.06718*, 2023.

APPENDIX A EXPERIMENT DETAILS

A. Environment

Our experiments are conducted in a real-world table-top manipulation environment, which involves a 7-DoF Franka Emika robot arm with a 2F-85 Robotiq gripper interacting with various objects on the table at 5Hz. Our tabletop environment has two fixed ZED 2.0 cameras and one ZED mini wrist camera that can take RGBD images.

The *top-down camera*, which is the primary camera used in this paper, provides RGB and the depth images for MOKA and other baseline methods.

In addition, we set up the front camera and the wrist camera for more complete observation to distill the collected robot experiences to the learned Octo policy [46]. The action space of the robot is 7-dimensional, consisting of the 6-DoF end-effector twist defined in the $SE(3)$ space with an additional dimension of gripper position. Each task can consist of multiple stages, each terminates after 100 steps.

B. Task Design

We design various table-top manipulation tasks with a diverse set of daily objects. TABLE III shows the language description of the full list of tasks.

Table Wiping	1. Move the glasses into the glasses case 2. Sweep the trash with the broom
Watch Cleaning	1. Put the watch into the ultrasound cleaner 2. Press the power button
Gift Preparation	1. Put the golden filler in the gift box 2. Put the perfume in the gift box
Laptop Packing	1. Unplug the cable 2. Close the lid of the laptop
Fur Removing	1. Grasp the fur remover 2. Sweep the fur on the sweater with the remover
Drawer Closing	1. Close the drawer
Scissor Handing	1. Grasp the scissor 2. Hand the scissor to a human
Flower Arrangement	1. Grasp the pink roses on the table 2. Insert the roses into the vase

TABLE III: The language description of all the proposed tasks. Each of the tasks consists of two stages except for the drawer closing task.

We provide the comparative evaluation on the top four tasks (Table Wiping, Laptop Packing, Gift Preparation and Ultrasound Cleaning) in the main paper, with additional results of MOKA on other four tasks in our supplementary video.

C. Success Checking

To count the failure modes and collect successful trajectories, we first query VLM with task instruction and obtain the point-based motion prediction. If the prediction is correct, it will be executed on the robot. Otherwise, it will be counted as the VLM reasoning failure as mentioned in Sec. V-B. After executing the motion prediction from VLM, the human

expert will manually check if the task succeeds or not. If the task is not successfully finished, it will be counted as an execution failure. All the successful trajectories will be saved as demonstrations for policy distillation.

APPENDIX B IMPLEMENTATION DETAILS

We introduce the overview of MOKA in Alg. 1, and the implementation details of each component in the following sections.

Algorithm 1 MOKA Pipeline

- 1: **Input:** Vision-language Model $\mathcal{M}$, Task instruction l , text prompt for high-level reasoning p_{high} , text prompt for low-level reasoning p_{low} , initial observation s_0
 - 2: Query $\mathcal{M}$ for high-level task reasoning, obtain $y_{high} = \mathcal{M}([p_{high}, l, s_0])$ which decompose the task into N sub-tasks.
 - 3: **for** subtask $k = 0 \dots N - 1$ **do**
 - 4: Get observation s_k from the top-down camera
 - 5: Propose keypoint and waypoint candidates and get annotated image $f(s_k)$
 - 6: Query $\mathcal{M}$ for low-level motion reasoning, obtain $y_{low}^k = \mathcal{M}([p_{low}, y_{high}^k, f(s_k)])$
 - 7: Execute y_{low}^k on the real robot
 - 8: **end for**
-

A. High-Level Task Reasoning

Given the initial observation image s_0 and the language task description l , we first query the VLM $\mathcal{M}$ with the language instruction to produce the decomposed subtask, which contains the structured information for the K subtasks, including descriptions of objects and desired motion. We provide the prompt for high-level reasoning in TABLE IV.

The response contains fields “object_grasped”, “object_unattached” and “motion_direction”, which will be used for later stages in our pipeline, such as keypoint proposal and motion prediction.

B. Point-based Affordance Proposal

After obtaining the high-level reasoning results, we can know the objects involved in each subtask. Since VLM cannot directly generate keypoints and waypoints, we need to propose some candidate points and let VLM select corresponding points through a visual question-answering way.

Keypoint proposal. We leverage GroundedSAM [54] to extract segmentation masks conditioned on a text prompt, which is designed to be a string of object names involved in the current subtask (e.g., “trash, broom”). Given an image of current observation and such a text prompt about the involved object names, we first employ GroundingDINO [36] to generate bounding boxes for the objects by conditioning on the text prompt. Later, the annotated boxes given by GroundingDINO [36] serve as the box prompts for SAM [27] to generate corresponding segmentation masks for the objects.

The input request contains:

- A string describes the multi-stage task.
- An image of the current table-top environment captured from a top-down camera.

The output response is a list of dictionaries in the JSON form. Each dictionary specifies the information of a subtask, following the correct order of executing the subtasks to solve the input task. Each dictionary contains these fields:

- **instruction**: A string to describe the subtask in natural language forms.
- **object_grasped**: A string to describe the name of the object that the robot gripper will hold in hand while executing the task (e.g., the object to be picked or used as a tool to interact with other objects). This field can be empty if there is no such object involved in this subtask.
- **object_unattached**: A string to describe the name of the object that the robot gripper will interact with directly or via another object without holding it in hand (e.g., the object to be touched by the tool, the target object where “object_grasped” will be moved onto). This field can be empty if there is no such object involved in this subtask.
- **motion_direction**: A string to describe the direction of the robot gripper motion while performing the task (e.g., “from right to left”, “backward”, “downward”).

TABLE IV: The high-level reasoning prompt for all the tasks. It will decompose a multi-stage task into subtasks. The specified output fields can be used for later low-level reasoning stage.



Fig. 6: The affordance proposal for the table wiping task in stage 1 and stage 2. The keypoints are plotted on corresponding objects for different stages based on the high-level reasoning results. The annotated grid cells are used for VLM to select the position of potential waypoints.

We define the keypoints of each object to be a set of boundary keypoints with a center keypoint. To extract them, we sample K points on the contour of each object using farthest point sampling [48], and also include the geometric mean point of the object segmentation mask. We then plot these $K + 1$ keypoints on the 2D image as our keypoint proposals to the VLM.

Waypoint proposal. Unlike keypoints, waypoints are often the points in the free space that are not attached to any objects. To perform waypoint selection, we evenly divide the full image into 5×5 grid tiles, which mark the column as a, b, c, d, e from left to right and rows as $1, 2, 3, 4, 5$ from bottom to top, as shown in Fig. 6. The waypoints will be sampled from the predicted tile from VLM.

C. Low-Level Motion Reasoning

After annotating the observation image using object keypoints and grid cells as described in the above section, we can query VLM about low-level motion with annotated image and text prompts. Firstly, we describe the text and image inputs to the VLM using prompt V.

Then we first explain the definition of keypoints and waypoints using prompt VI.

After that, we can specify the output format and further explain the role of each field using prompt VII. We can also

provide some step-by-step guidance about how the reasoning procedure should be done (similar to chain-of-thought prompting [62]).

D. Motion Execution

We can decompose a manipulation subtask into an **optional** grasping phase and a manipulation phase. For some tasks like tool using, the robot needs to first grasp an object before making contact with another object. For other tasks where the robot can directly interact with target objects (e.g., pushing, pressing button), the “grasp_keypoint” field will be empty. For such tasks, we will skip the execution of grasping phase.

Grasping phase. We first sample 30 antipodal 4DoF grasp proposals based on the primary camera’s depth image, which is cropped by the bounding box of **object_grasped**. Here, we use the same antipodal grasp proposal method used in DexNet 2.0 [39]. We then lift the VLM-predicted grasp point from 2D image to the robot frame based on the primary camera’s intrinsic and extrinsic parameters. After that, we can select the nearest grasp proposal based on the lifted grasp point, and execute the 4DoF grasp on the robot.

Manipulation phase. After obtaining the pre-contact tile and post-contact tile, we first sample pre-contact waypoint and post-contact waypoint from the predicted tiles. Since the waypoints are in the free space, its 3D position cannot be determined given its 2D projection on the image. So we also assign the VLM-predicted height to them. We only consider the cases where the waypoints are in same height with target point (e.g., pushing), or above the target point (e.g., placing) in most common table top manipulation scenarios.

We then sequentially move the function keypoint to the pre-contact waypoint, target keypoint and post-contact waypoint to perform a contact motion.

After both phases are successfully done, the robot need to get the motion prediction for the next subtask. In order to obtain a clean and non-occluded image to perform low-level reasoning, we first move the robot to the neural pose, get an observation image from the top-down camera, and then resume the robot. Subsequently, the robot will execute the predicted motion until the multi-stage task is finished.

Please describe the robot gripper’s motion to solve the task by selecting keypoints and waypoints.
The input request contains:

- The task information with these fields:
 - **instruction**: The task in natural language forms.
 - **object_grasped**: The object that the robot gripper will hold in hand while executing the task.
 - **object_unattached**: The object that the robot gripper will interact with either directly or via another object without holding it in hand.
 - **motion_direction**: The motion direction of the robot gripper or the in-hand object while performing the task.
- An image of the current table-top environment captured from a top-down camera, annotated with a set of visual marks:
 - **candidate keypoints on object_grasped**: Red dots marked as $P[i]$ on the image.
 - **candidate keypoints on object_unattached**: Blue dots marked as $Q[i]$ on the image.
 - **grid for waypoints**: Grid lines that uniformly divide the images into tiles. The grid equally divides the image into columns marked as a, b, c, d, e from left to right and rows marked as 1, 2, 3, 4, 5 from bottom to top.

TABLE V: The description about image and text inputs in the low-level motion reasoning stage.

The motion consists of an optional grasping phase and a manipulation phase, specified by **grasp_keypoint**, **function_keypoint**, **target_keypoint**, **pre_contact_waypoint**, and **post_contact_waypoint**.

The definitions of these points are:

- **grasp_keypoint**: The point on “object_grasped” indicates the part where the robot gripper should hold.
- **function_keypoint**: The point on “object_grasped” indicates the part that will make contact with “object_unattached”.
- **target_keypoint**: If the task is pick-and-place, this is the location where “object_grasped” will be moved to. Otherwise, this is the point on “object_unattached” indicating the part that will be contacted by “function_keypoint”.
- **pre_contact_waypoint**: The waypoint in the free space that the functional point moves to before making contact with the “target_keypoint”.
- **post_contact_waypoint**: The waypoint in the free space that the functional point moves to after making contact with the “target_keypoint”.

TABLE VI: The explanation about the definition of keypoints and waypoints.

The response should be a dictionary in JSON form, which contains:

- **grasp_keypoint**: Selected from candidate keypoints marked as $P[i]$ on the image. This will be empty if and only if object_grasped is empty.
- **function_keypoint**: Selected from candidate keypoints marked as $P[i]$ on the image. This will be empty if and only if object_grasped or object_unattached is empty.
- **target_keypoint**: Selected from keypoint candidates marked as $Q[i]$ on the image. This will be empty if and only if object_unattached is empty.
- **pre_contact_tile**: The tile that the pre-contact waypoint should be in. This is selected from candidate tiles marked on the image.
- **post_contact_tile**: The tile that post-contact waypoint should be in. This is selected from candidate tiles marked on the image.
- **pre_contact_height**: The height of pre-contact waypoint as one of the two options “same” or “above” (same or higher than the height of making contact with target keypoint).
- **post_contact_height**: The height of post-contact waypoints as one of the two options “same” or “above”.
- **target_angle**: Describe how the object should be oriented during this motion in terms of the axis pointing from the grasping point to the function point.

Think about this problem step by step and explain the reasoning steps. First, choose grasp_keypoint, function_keypoint, and target_keypoint on the correct parts of the objects. Next, describe which tile the target_keypoint is located in. Then choose pre_contact_tile, post_contact_tile, pre_contact_height, post_contact_height such that the resultant motion from pre-contact waypoint to target keypoint, then to post-contact waypoint in 3D follows the “motion_direction” input. Remember that the columns are marked as ‘a’, ‘b’, ‘c’, ‘d’, ‘e’ from left to right, and the rows are marked as 1, 2, 3, 4, 5 from bottom to top.

TABLE VII: The output format of the low-level motion reasoning, including some further explanations.

E. Policy Distillation

We employ the base checkpoint of Octo [46] and perform full-model finetuning using their official instruction. The default architecture of Octo is a transformer-based diffusion policy, which is conditioned on a task description or goal image through a task tokenizer along with tokenized pixel observations from multiple cameras, and predicts actions through a diffusion head. The diffusion head consists of a 3-layer MLP with a hidden dimension of 256 using a standard DDPM objective [19].

We mostly reuse the hyperparameters as in Octo [46]. For customized hyper-parameters we used, please refer to TABLE VIII.

Hyperparameter	Value
Learning Rate	3e-4
Warmup Steps	1000
Weight Decay	0.01
Gradient Clip Threshold	0.5
Batch Size	256
Total Gradient Steps	200000

TABLE VIII: Hyperparameters used during fine-tuning.

F. Evaluation

We evaluate MOKA in both zero-shot and in-context learning settings. For the zero-shot setting, we keep all the above prompts **unchanged** but only change the language task de-

scription (e.g., “sweep the garbage”, “pick up the perfume”, “press the button”).

For in-context learning settings, we first collect two examples from VLM’s successful predictions in different scenes with scene and object variations. As shown in Fig. 7, MOKA can be improved from such simple and intuitive in-context examples, without intricate prompt engineering.

APPENDIX C ADDITIONAL RESULTS

A. Ablation Study

We design the following ablative studies to understand the effectiveness of different design options in MOKA.

- MOKA w/o hierarchy: we can skip the high-level task reasoning but directly ask for low-level motion reasoning from GPT-4V. Here all the objects in the scene will be annotated with keypoints, and VLM is queried to generate point-based affordance directly from the annotated image.
- MOKA w/o description of points: we can remove the description of the definition of keypoints and waypoints in Tab. VI.
- MOKA w/o chain-of-thought prompting: we can remove the step-by-step guidance at the last paragraph in the prompt in Tab. VII.

The results of our ablative studies are illustrated in Tab. IX. For each task, we report the number of reasoning successes rate out of 10 trials. The results demonstrate that our method can obtain consistent improvements from all the above prompting designs. Specifically, **MOKA w/o hierarchy** decreases the performance by a large margin, which indicates the importance of subtask decomposition. **MOKA w/o keypoint description** and **MOKA w/o CoT** can preserve most of the performance, but still worse than MOKA on most of the tasks. As illustrated in Fig. 8, VLM can make various mistakes in terms of keypoint and waypoint predictions. Specifically, for MOKA w/o hierarchy, the VLM can make mistakes about subtask ordering, which can cause a complete failure of the task.

B. Additional Robustness Analysis

We provide additional results for robustness analysis. The predicted affordances for the drawer closing and fur removal tasks with various objects and initial poses are shown in Fig. 9 and Fig. 10. In these two figures, each column in the image uses similar initial arrangements of objects, while the two rows use different objects. For the drawer closing tasks, MOKA can consistently predict the waypoints and target points following the correct closing motion regardless of the poses and appearances of drawers and other distractors. For the fur removal tasks, MOKA can predict all the keypoints correctly with a feasible motion in different environments.

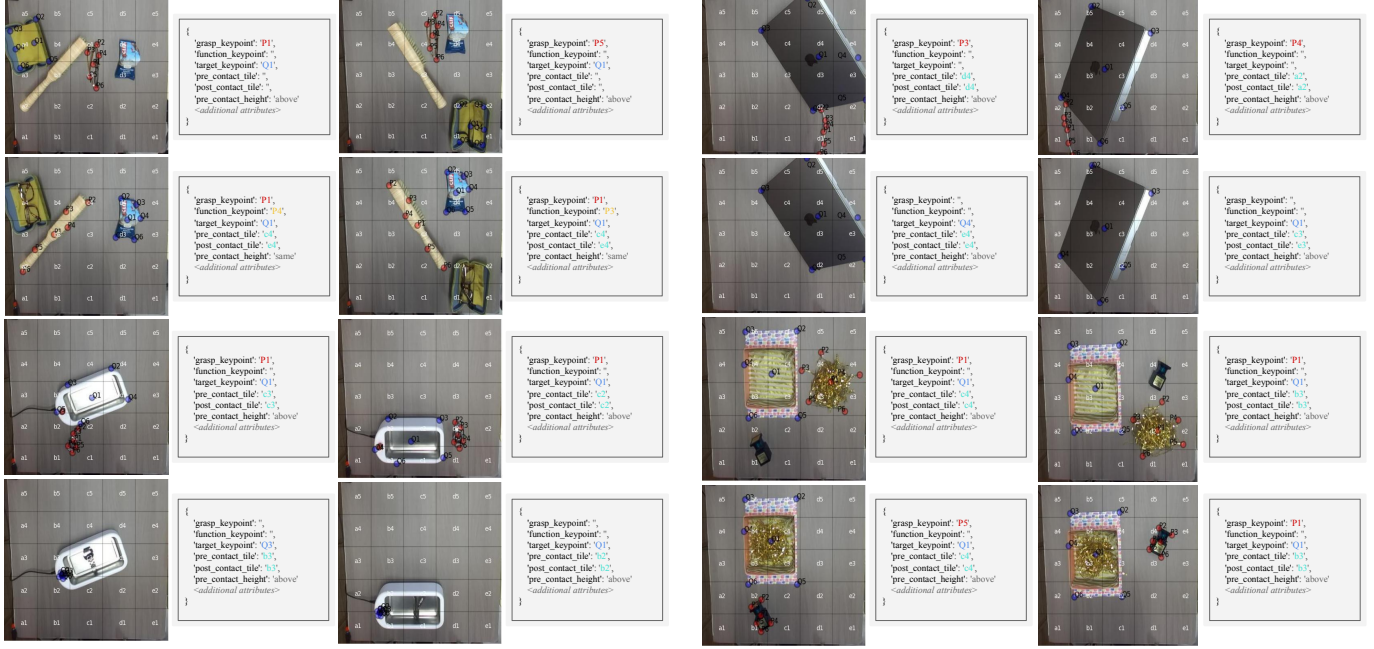


Fig. 7: The in-context examples used in MOKA which are collected under object pose variations.

	Table Wiping		Watch Cleaning		Gift Preparation		Laptop Packing	
	Subtask I	Subtask II	Subtask I	Subtask II	Subtask I	Subtask II	Subtask I	Subtask II
MOKA	0.7	0.7	0.8	1.0	1.0	0.8	0.6	0.8
MOKA w/o hierarchy	0.1	0.1	0.2	0.1	0.2	0.2	0.0	0.1
MOKA w/o keypoint description	0.5	0.6	0.7	0.8	0.9	0.6	0.4	0.8
MOKA w/o CoT	0.5	0.4	0.6	0.8	0.9	0.6	0.4	0.6

TABLE IX: Ablation studies on different prompt designs. Across 4 tasks, each consists of 2 subtasks, MOKA consistently benefits from the three prompt designs.



Fig. 8: Qualitative results of ablation studies. Without the keypoint description or chain-of-thought prompting, the model can make mistakes in keypoint prediction (left column) and waypoint prediction (right column).

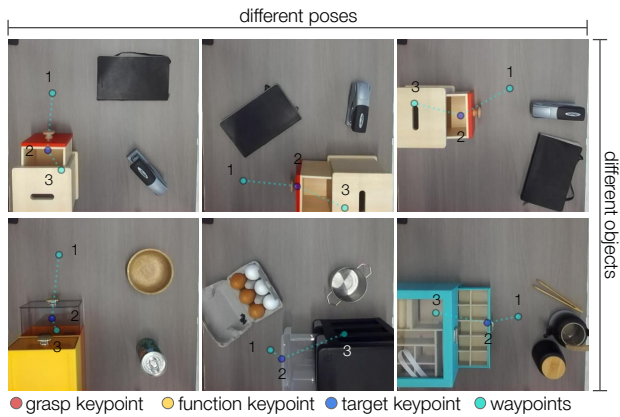


Fig. 9: Robustness analysis for the drawer closing task.



Fig. 10: Robustness analysis for the fur removal task.