

Reasoning Grasping via Multimodal Large Language Model

Shiyu Jin¹, Jinxuan Xu², Yutian Lei¹, Liangjun Zhang¹
¹Baidu Research ²Rutgers University

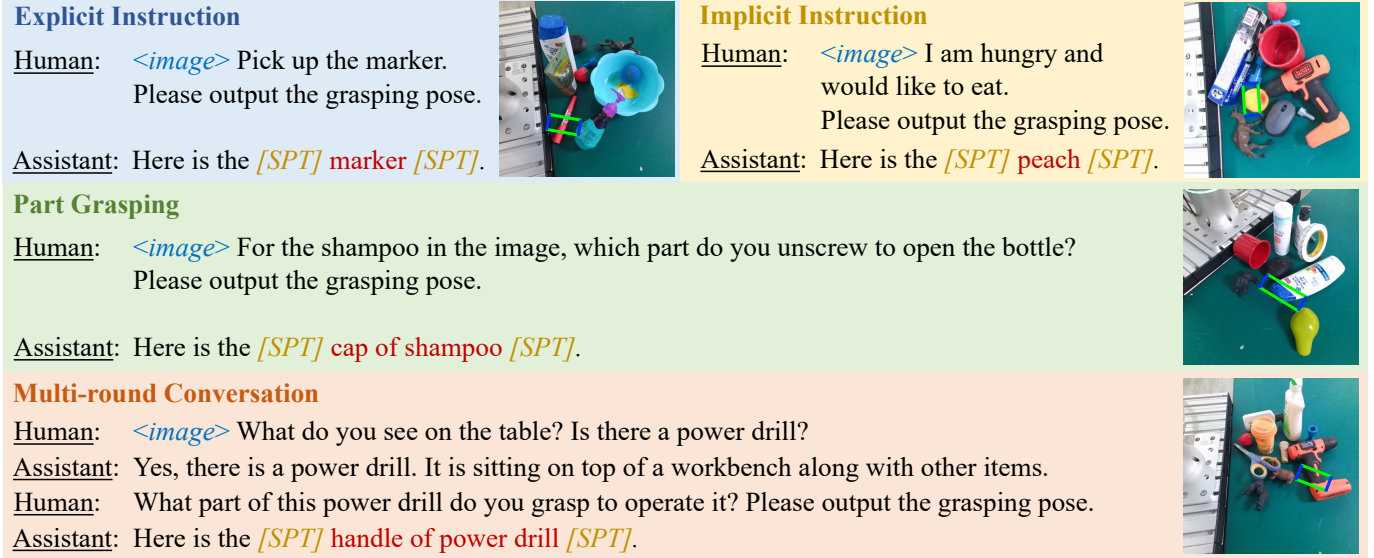


Fig. 1: Overview. We integrate the reasoning abilities of multi-modal Large Language Models with robotic grasping. The resulting model supports multi-round conversations with users, interprets complex and implicit instructions, and accurately predicts robotic grasping poses for target objects or specific parts within cluttered environments. In the textual output from the model, the grasping target is indicated by two special tokens [SPT], as demonstrated in the figure. The output grasp poses are visualized in the images with rectangles.

Abstract—Despite significant progress in robotic systems for operation within human-centric environments, existing models still heavily rely on explicit human commands to identify and manipulate specific objects. This limits their effectiveness in environments where understanding and acting on implicit human intentions are crucial. In this study, we introduce a novel task: reasoning grasping, where robots need to generate grasp poses based on indirect verbal instructions or intentions. To accomplish this, we propose an end-to-end reasoning grasping model that integrates a multi-modal Large Language Model (LLM) with a vision-based robotic grasping framework. In addition, we present the first reasoning grasping benchmark dataset generated from the GraspNet-1 billion, incorporating implicit instructions for object-level and part-level grasping, and this dataset will soon be available for public access. Our results show that directly integrating CLIP or LLaVA with the grasp detection model performs poorly on the challenging reasoning grasping tasks, while our proposed model demonstrates significantly enhanced performance both in the reasoning grasping benchmark and real-world experiments.

I. INTRODUCTION

Robotic grasping has long been a subject of extensive study. The utilization of CNN-based neural networks has demonstrated efficiency in generating high-quality grasping poses from visual input [31, 21, 18, 23, 33, 27]. One significant

limitation of these methods is the lack of scene understanding. The robot cannot identify the objects they are grasping. Recent methods have introduced language-guided grasping with instructions such as “I want a stapler” [6] or “Pick the food box in front of the ball” [42]. However, those methods will struggle when dealing with implicit instructions, where the desired object is not explicitly named. For example, if a user says “I need to drink water”, an intelligent assistant robot should identify and grasp a cup from a selection of household items, although “cup” does not appear in the user’s instruction. Such scenarios are common in real-world settings and the solutions to them remain open questions.

Recently, the advance of large language models like ChatGPT [32] has introduced new possibilities for language reasoning. Beyond text, numerous studies have leveraged multi-modal Large Language Models (LLMs) for visual understanding [50, 11, 24], enabling these models to interpret and generate responses based on both textual and visual inputs. However, these approaches focus on visual comprehension and text generation, lacking the capability to generate robotic actions. While there have been efforts to integrate LLMs with robotics [2, 16, 20, 36, 43, 17], these concentrate on high-level planning.

In this work, we introduce the novel task of Reasoning Grasping, where a robot determines grasp poses based on users’ “implicit” instructions. “Implicit” instructions are language instructions that do not specify the name of the grasp target. A formal definition of reasoning grasping and implicit instruction is given in Sec III.

To address the reasoning grasping task, we propose a model that directly output grasp poses according to image inputs and language instructions in cluttered environments. Specifically, this model leverages a multi-modal LLM to interpret both images and instructions. We employ a special token strategy to identify tokens that are relevant to the grasping target (i.e., names of target objects). Subsequently, embeddings of these identified tokens are used to generate accurate grasping poses. To train this model, we extend the GraspNet-1 billion dataset [13] with diverse implicit instructions and object part grasping.

Our contribution can be summarized as follows:

- We introduce a novel task of reasoning grasping, which directs robots to grasp objects based on implicit instructions.
- We propose a multi-modal LLM for reasoning grasping. Our experiments demonstrate the model’s promising performance.
- We have developed a unique dataset for the reasoning grasping task. It includes 64 objects, 109 parts, 1730 reasoning instructions, and around 100 million grasping poses. This dataset provides a comprehensive resource for training and evaluating models on reasoning grasping tasks.

II. RELATED WORK

A. Robotic Grasping

Robotic grasping has traditionally relied on analytical methods [28, 10, 35], which focus on understanding the geometry of objects or analyzing contact forces. Convolutional neural networks (CNNs) based methods [31, 21, 18, 23, 33, 27] utilize large datasets of labeled grasping examples [19, 8, 13] to train models to predict grasps based on visual input. However, a critical limitation of both analytical and CNN-based methods is their lack of scene understanding and inability to process language instructions. They do not inherently know what they are grasping, which restricts their application in more dynamic, human-centric environments.

Recent advancements in robotic grasping have seen the integration of language understanding, enabling robots to grasp objects based on natural language instructions [26, 38, 40, 5, 46, 37, 45, 47, 49, 1]. This shift allows robots to identify and manipulate objects as specified by users, enhancing their utility in complex scenarios. For example, Tziafas et al. focus on referring grasp synthesis, predicting grasp poses for referred objects through natural language in cluttered scenes [42]. Chen et al. proposed to learn from visual and language features jointly and predict the 2D grasp box from the RGB images [6].

B. LLMs for robotics

Recent breakthroughs in Large Language Models (LLMs) have been impressive [9, 4, 32, 41] in understanding and generating human-like text. Multi-modal LLMs can seamlessly integrate with other modalities, such as vision, enhancing their proficiency in tasks like visual understanding [24]. There are also a lot of efforts integrating LLMs with robotics. Many studies [2, 16, 20] have integrated LLMs into closed-loop planning structures, decomposing language-conditioned long-horizon tasks into small steps. Yet, the gap between language instructions and actions still remains. Furthermore, some studies [36, 43, 17] have employed program-like specifications to prompt LLMs, melding planning and action using a predefined library of action functions.

There is also an increase in research focused on utilizing LLMs specifically for robotic grasping. Tang et al. proposed GraspGPT to use LLMs to generate semantic knowledge and then selected task-oriented grasp pose from grasp candidates [39]. Mirjalili et al. proposed LAN-grasp to identify the optimal part of an object for grasping with LLM, treating the rest as obstacles, and then applied a traditional grasp planner to determine the best grasp [29]. Compared to GraspGPT and LAN-grasp, our method is distinguished by its ability to operate in cluttered scenes and its end-to-end training framework, while GraspGPT and LAN-grasp utilize modular frameworks and depend heavily on other pre-trained grasp detection models.

III. REASONING GRASPING

In this section, we define the task of **Reasoning Grasping**, which aims to generate a robotics grasp pose g , given an input textual instruction t , an input RGB image v , and, optionally, an input depth image v_d . In the reasoning grasping task, instructions t are not limited to being straightforward, they can also be “implicit”. *We define implicit instructions as language instructions that do not explicitly specify the name of the grasp target but provide relevant contextual cues or indicate users’ intent implying the grasp target.* For example, implicit instructions may include cues such as shapes, colors, or other attributes of grasp targets, or they may describe users’ needs or intent implying the grasp targets based on their functionalities.

The output grasp poses in this work are defined in alignment with previous works [21, 31]. We represent a grasp pose as $g = (x, y, \theta, w, q)$, where (x, y) represents the grasp’s central coordinates in the image plane, θ denotes the rotation angle about the z -axis in the range $[-\frac{\pi}{2}, \frac{\pi}{2}]$, and w is the grasp width of the required object in the image coordinates. The parameter q signifies the grasp quality score, assigned to each point in the image on a scale from 0 to 1, with 1 indicating a higher likelihood of successful grasping.

IV. REASONING GRASPING VIA MULTI-MODAL LLMs

In this section, we present an innovative model for reasoning grasping, which combines conventional robotic grasping in [21] and reasoning capabilities from multi-modal LLMs. Specifically, the proposed model utilizes the reasoning power

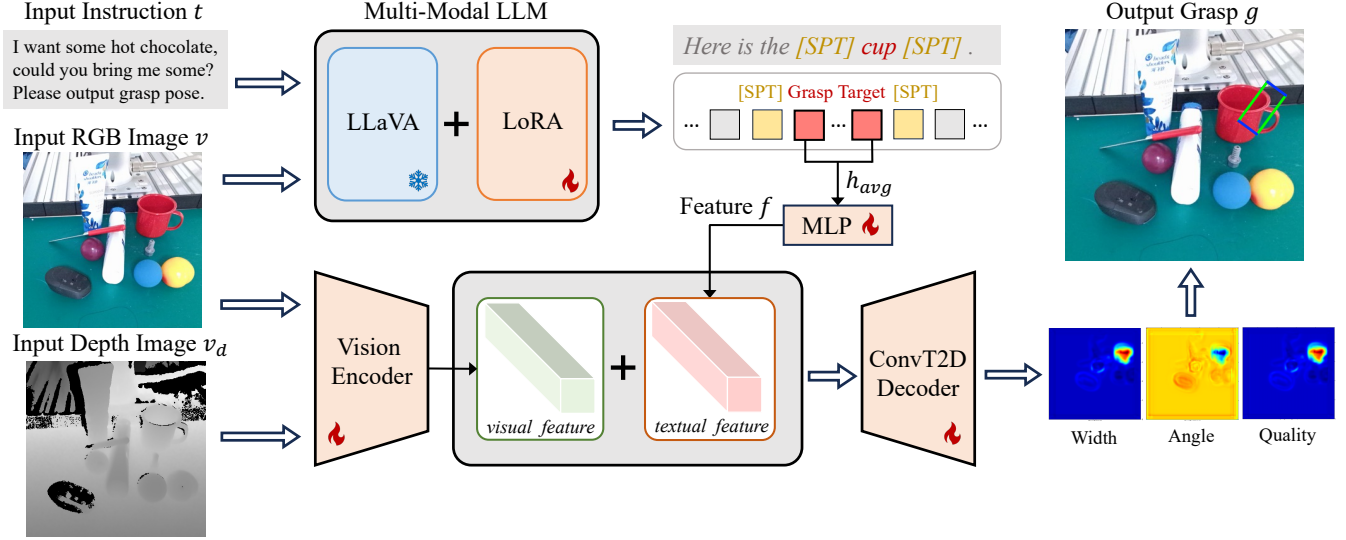


Fig. 2: Framework of the proposed reasoning grasping model. This model processes visual images v and textual instructions t to output the grasp pose g for the specified target object or part. The embeddings of the grasp target are passed to the grasping module for grasping poses detection.

of multi-modal LLMs by integrating the extracted features from its textual outputs into the robotic grasping prediction process.

This model addresses two key challenges: firstly, it addresses the adaptation of pre-trained multi-modal LLMs for efficient training in downstream tasks. Although LLMs have strong reasoning abilities, they often produce verbose and diverse outputs, making it challenging to directly integrate them into training other tasks, especially within an end-to-end framework. Secondly, the model addresses the combination of textual outputs from LLMs and image-based robotic grasping systems. This integration is crucial as it ensures that the rich linguistic context is effectively utilized to refine the subsequent grasping prediction.

A. Architecture

The proposed model framework is illustrated in Fig. 2. As for multi-modal LLMs, we utilize the pre-trained LLaVA [24] enhanced with LoRA [15] fine-tuning techniques. LLaVA integrates a CLIP [34] visual encoder with LLaMA [41], an open-source Large Language Model comparable in performance to GPT-3 [4]. And LoRA, known for its computational efficiency, is incorporated into both the projection layer and all linear layers of multi-modal LLMs.

The key role of multi-modal LLMs in the proposed model is to interpret both instruction t and image v , then accurately identify the grasping target. Such a target could be an object or, more challenging, a specific part within cluttered scenes. To address the verbosity of LLM outputs, we introduce a special token, i.e., $[SPT]$, strategically placed around the names of the grasping target in the LLM textual outputs. For instance, $[SPT]$ grasping target name $[SPT]$. The purpose of introducing this special token is to identify and isolate the crucial tokens, specifically, the names of the grasping target, from the textual

outputs generated by LLMs.

Once the grasping target is identified using the special token, we proceed to extract the last-layer embeddings of its corresponding tokens. It is important to note that there could be multiple tokens associated with the identified grasping target, and we extract embeddings for all identified tokens. These extracted embeddings encapsulate the essential linguistic information related to the grasping target. We then compute the average of these embeddings denoted as h_{avg} , which is passed through a projection network to obtain the feature f . This feature f is then fed into the image-based grasping detection process, guiding the prediction of the grasp pose.

The grasping detection in our model, rooted in the CNN-based robotic grasping framework in [21], takes n-channel images as inputs. The input image is first passed through three convolutional layers and five residual layers. The extracted visual features are concatenated with the embeddings obtained from the reasoning module. Following this, we use the convolutional transpose operation to up-sample this concatenated feature, increasing its size to match the input image size. The model subsequently generates four output images, each representing different aspects of the grasp pose g : the grasp quality score q , the rotation angle expressed in $\cos 2\theta$ and $\sin 2\theta$, and the required width w for the end-effector. The grasp central coordinate (x, y) is determined from the output image of the grasp quality score. This is obtained by selecting the point in the image with the highest quality score. Notably, the depth image v_d is an optional input modality for the grasping prediction, depending on data availability.

B. Loss Function

The proposed end-to-end model is trained by a combination of text generation loss L_{text} and grasping prediction loss L_{grasp} . The overall objective L is a weighted sum of these

two components, formulated as follows:

$$L = \lambda_t L_{text} + \lambda_g L_{grasp}, \quad (1)$$

where λ_t and λ_g are parameters that determine the relative weights of two loss terms, respectively. Here L_{text} represents the auto-regressive cross-entropy loss for the textual output generated by LLMs, which assesses the performance of identifying the correct grasping target according to input instructions. On the other hand, L_{grasp} evaluates the accuracy of the output grasping predictions, adhering to the loss definition detailed in [21].

C. Training Strategy

The proposed method leverages the pre-trained LLaVA model, with the LoRA technique for fine-tuning during the training phase. With LoRA fine-tuning, the multi-modal LLM can generate the grasp target within just a few epochs. Based on this observation, our training strategy involves freezing the parameters of the LoRA-enhanced model after the initial epochs. Specifically, from the third epoch, we set the parameter λ_t in the loss function to 0 and freeze the parameters of the LoRA-enhanced model. Our training then focuses exclusively on other parameters outside the LoRA framework. This approach is designed to reduce the training time and expedite the convergence of the model, particularly in terms of grasping prediction.

Regarding the datasets used for training, our model utilizes both the reasoning grasping dataset (detailed in Section V) and the original LLaVA Instruct 150K dataset [24] utilized in LLaVA’s training. To maintain LLaVA’s visual reasoning capabilities, we initially train the model using both datasets concurrently for the first two epochs. Starting from the third epoch, with the LoRA parameters frozen, the model’s training proceeds solely with the reasoning grasping dataset.

V. REASONING GRASPING DATASET

To train the model for reasoning grasping tasks, we introduce our reasoning grasping dataset, which builds upon the GraspNet-1 billion dataset [13]. GraspNet-1 billion consists of 190 indoor tabletop RGB-D scenes, 88 objects, and 97280 images. Its wide variety of scenes, objects, and grasp poses provides an excellent foundation. We extend the GraspNet-1B dataset with reasoning instructions as well as object parts grasping annotations. We provide some examples in our dataset in Fig. 3.

A. Reasoning Instruction Generation

As illustrated in Fig. 4, we developed a unique and detailed method to create reasoning text instructions for our dataset by creating detailed object descriptions, automated generation with GPT-4, and manual review. The generated instructions enable robots to grasp based on nuanced, context-rich queries.

1) *Description Creation*: The initial step in our instruction generation process involved creating detailed descriptions for each object and part, providing essential information such as functionalities, physical attributes, and typical uses. For instance, a pair of scissors was detailed as “A handheld cutting instrument with two crossing metal blades pivoted together, typically used for cutting paper or fabric. This particular pair has a black handle with yellow inner grips.” These descriptions served as the foundation for generating indirect questions and instructions.

2) *Automated Generation with GPT-4*: The next step in instruction generation involved utilizing a customized variant of the GPT-4 model. This model was tasked with generating indirect questions and instructions for various objects and their parts. The output consists of 10 strings: 5 indirect questions that reference the object’s functions or design, and 5 indirect instructions that describe actions involving the object or its usage in certain tasks. These are crafted to be specific, relevant, and direct for grasping tasks.

3) *Manual Review and Refinement*: Post-generation, each set of GPT-generated instructions underwent a manual review and refinement process. The manual intervention allowed us to fine-tune the language, enhancing the quality and applicability of the instructions for the targeted objects and their parts. This semi-automated approach balanced the efficiency of machine-generated content and judgment of human experts.

B. Part Annotation

1) *Part Segmentation*: We manually segmented the part point clouds of each object using a specially developed annotation tool. The segmentation process was guided by an understanding of how humans typically interact with and use various parts of an object. For instance, a knife was divided into its blade and handle, reflecting the distinct functional roles of each part.

2) *Grasping Pose Annotation*: Our dataset’s annotation process includes two key components: pose assignment and pose projection, for both 3D and 2D grasping poses, as shown in Fig. 5.

Pose assignment involves aligning the grasping poses with the nearest part of the object. For each object’s segmented parts, we determine the grasping center of each point. These grasping poses are then assigned to the nearest part, ensuring an accurate association with the specific object part they correspond to.

Pose projection transforms the grasping poses of parts into both 3D and 2D representations. This is achieved using camera matrices and object transformations. In the case of 3D, the original grasping poses are preserved in the spatial domain. For 2D projection, the 3D points of each segmented part are projected onto 2D images. This results in 2D masks that depict the spatial distribution of each part on the image plane. Grasping poses within these masks are then mapped to the corresponding parts in 2D.

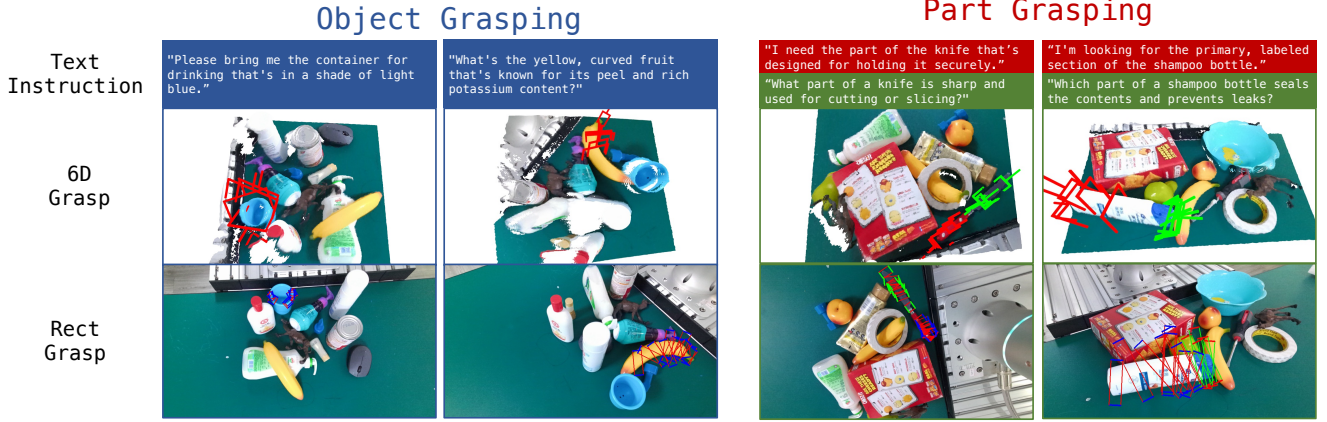


Fig. 3: Our reasoning grasping dataset enriches the GraspNet-1 billion dataset [13] with additional annotations for part-level grasping and reasoning instructions.

[OBJ]: [DESC] [PART-1]: [DESC] [PART-2]: [DESC] [Scissor]: A handheld cutting instrument with two crossing metal blades pivoted together, typically used for cutting paper or fabric. This particular pair has a black handle with yellow inner grips. [Blade]: This is the long, straight metal part converging to a point, used for cutting. It appears to be made of stainless steel and have a sharp edge on the inner side. [Handle]: The handle is the black plastic parts with yellow inner rubber grips, shaped into loops for inserting fingers when used manually. It provides a comfortable and secure grip for users.	[OBJ/Part]: [TXT INSTR] [TXT INSTR] [Scissor]: "How would you describe the tool that can alter the length of curtains?" "I'm looking for the implement with black and yellow grips to snip these threads; do you see it?" "What's the tool in the drawer with two blades that slides past each other to cut materials?" [Blade]: [TXT INSTR] [Handle]: [TXT INSTR] 	[OBJ/Part]: [TXT INSTR] [TXT INSTR] [Scissor]: "How would you describe the tool that can alter the length of curtains?" (Delete) "I'm looking for the instrument with black and yellow grips to snip these threads; do you see it?" (Accept) "What's the tool in the drawer with two blades that slides past each other to cut materials?" (Revise) [Blade]: [TXT INSTR] [Handle]: [TXT INSTR]
Object/Part Descriptions	Automated Generation with GPT-4	Manual Review and Refinements

Fig. 4: Reasoning instruction generation. The instructions generation process involves 1) Object/Part Descriptions; 2) Automated Generation with GPT-4; and 3) Manual Review and Refinements.

C. Dataset Statistics

In ensuring data quality, low-quality grasping poses were eliminated based on confidence evaluations and objects lacking semantic information were removed. In total, our enhanced dataset comprises an extensive collection of 1,730 instruction pairs, 64 objects, 109 segmented parts, and around 100 million associated grasping poses. Table I shows the statistics comparison of the previous grasping and visual reasoning datasets with ours. Unlike other grasping datasets, we offer detailed reasoning instructions and part-specific grasping information.

VI. EXPERIMENT

In this section, we report the experiment results of the proposed method on the reasoning grasping dataset and real-world experiments. In the following experiments, we utilize both explicit instructions, such as "Pick up the <object>", as

well as implicit instructions that do not directly mention the grasping target.

A. Implementation Details

For the multi-modal LLM in the proposed model, we utilize the pre-trained LLaVA-7B-v0, which is derived from the large language model LLaMA-7B. As for the LoRA fine-tuning, we choose a rank $r = 64$. During the training, we use the batch size of 8, and the learning rate is set to $5e-4$, utilizing a cosine annealing schedule for adaptive learning rate adjustment. For the relative weight parameters in the loss function 1, we set $\lambda_t = 1$ and $\lambda_g = 1$ for the initial configuration, and starting from the third epoch, the parameter λ_t is adjusted to 0. For the train-test-split, we use 90% of the first 100 scenes in GraspNet-1 billion [13] for training and 10% of the first 100 scenes for testing. We sampled 50k from the original LLaVA Instruct

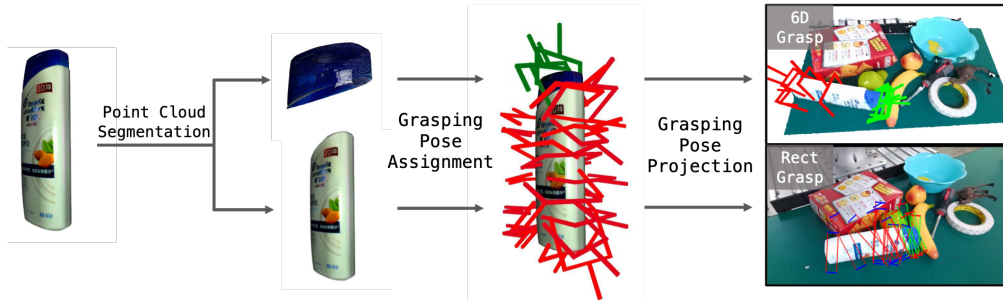


Fig. 5: Part grasping annotation process: 1) the object’s point cloud is manually segmented into distinct parts; 2) annotated grasping poses are allocated to each segmented part of the object; and 3) the grasping pose is projected to 2D and 6D.

TABLE I: Comparison on different datasets

	Grasp Label	Modality	Multi Object	Num. Objects	Grasps per Object	Num. Grasps	Num. Samples	Part Annotations	Reasoning Instructions
LISA[22]	✗	RGB	✓	N/A	N/A	N/A	1218	✗	✓
DectGPT [30]	✗	RGB	✓	N/A	N/A	N/A	5000	✗	✓
Cornell [19]	Rect.	RGB-D	✗	240	33	8019	1035	✗	✗
Jacquard [8]	Rect.	RGB-D	✗	11619	~ 20	1.1M	54,485	✗	✗
GraspNet [13]	6D	3D	✓	88	~400K	1.2B	97K	✗	✗
OCID-Grasp [3]	Rect.	RGB-D	✓	89	~ 7	75K	1763	✗	✗
MetaGraspNet [14]	6D	3D	✓	82	1-5K	N/A	217K	✗	✗
ACRONYM [12]	6D	3D	✓	8872	2000	10.5M	N/A	✗	✗
Grasp-Anything [44]	Rect.	RGB	✓	3M	200	600M	1M	✗	✗
ReasoningGrasp (ours)	6D	RGB-D	✓	64	~40K	~99.3M	97K	✓	✓

150K dataset [24] to mix with 100k grasping data to maintain the visual reasoning capabilities.

B. Baselines

For the problem of reasoning grasping, we implement two baselines to compare with our proposed model. One baseline employs a CLIP text encoder [34] to extract textual features, while the other utilizes a modular approach incorporating the latest LLaVA model.

1) *CLIP + GR-ConvNet Baseline*: In this baseline, we combine a CLIP text encoder with a CNN-based grasp detection model, GR-ConvNet [21]. We first utilize the CLIP text encoder to extract the textual feature from input instructions, and this feature is then integrated into the hidden layers of GR-ConvNet. While CLIP can effectively extract language features relevant to visual content, it may fall short in capturing implicit reasoning compared to the VLM.

2) *LLaVa → GR-ConvNet Baseline*: This is a modular baseline that integrates the latest LLaVA-1.6-34b [25] with a pre-trained grasping detection model GR-ConvNet. As LLaVA cannot directly output the grasping pose, this baseline operates in two stages to obtain the grasping pose: first, LLaVA takes an image and an instruction as the input and is prompted

to output a bounding box of the target object in the image. Then GR-ConvNet detects the optimal grasping pose within this specified bounding box. Note that LLaVA-1.6-34b has a much higher performance than LLaVa-v0-7b [24], which is the base model of our proposed reasoning grasping model. This baseline mirrors a similar process to LAN-grasp.

C. Results

1) *Text Generation*: We first evaluate the precision of text generation by the fine-tuned multi-modal LLM in the reasoning module. Accurately generating special tokens and target names is crucial, since this identified grasping target will be utilized in the subsequent grasping module. The special tokens and target names act as intermediate states and enable the model to eventually generate a precise grasping pose from given instructions. We differentiate between explicit instructions, which directly mention the object’s name, and implicit instructions, which hint at the target object without stating its name. Note that the instructions used for testing were not seen during the training phase. Results, presented in Table II, show that our model can successfully interpret 92.52% and 88.79% of implicit instructions for object grasping and part grasping,

respectively. This demonstrates our model’s good reasoning capabilities in understanding the grasping instructions.

TABLE II: Text generation accuracy.

	Explicit Instruction	Implicit Instruction
Object Grasping	99.01%	92.52%
Part Grasping	95.33%	88.79%

2) *Reasoning Grasping Result:* Table III shows the grasp prediction results on the reasoning grasping dataset. Performance is reported using the rectangle evaluation metric [18]. A grasp pose is deemed valid if it fulfills the following two conditions: 1) The Intersection over Union (IoU) score between the predicted and ground truth rectangles exceeds 25%; 2) The angular deviation between the orientations of the predicted and ground truth rectangles is less than 30 degrees. We reported $R@k$, indicating the presence of valid grasps within the top- k grasp predictions. Four distinct scenarios are evaluated, categorized by the combination of instruction types (explicit or implicit) and grasping targets (object or part). We assess both baseline models and our reasoning grasping model with and without depth input. Note that when depth information is utilized in the reasoning grasping model, it only serves as input for the grasping module and is not incorporated into the reasoning module, as shown in Fig. 2.

The results show that our reasoning grasping model outperforms the LLaVA $\rightarrow$ GR-ConvNet and CLIP + GR-ConvNet baselines across all scenarios. The first baseline, LLaVA $\rightarrow$ GR-ConvNet, shows limited effectiveness, as the direct combining of two models fails to convey sufficient information for accurate grasp detection. Although the CLIP + GR-ConvNet baseline exhibits some improvements, it still falls short, particularly in scenarios involving part grasping with implicit instructions. This indicates CLIP’s limited capability in processing implicit instructions. In contrast, our model excels in all scenarios, both with and without depth input, showcasing its superior ability to interpret implicit instructions and accurately generate the corresponding grasping poses. Note that some common failure cases are distinguishing objects with similar appearances, such as a shampoo bottle and a hair conditioner bottle, or identifying specific parts of complex objects like the uniformly colored toy camel.

3) *Visual Reasoning Ability:* Besides reasoning grasping, we also examine the visual reasoning capabilities of our model. The objective is to determine whether the model, after being fine-tuned, maintains the visual comprehension skills of the original LLaVA model. For a quantitative evaluation of performance, we employ GPT-4 to assess the quality of responses generated by our model. This is inspired by the evaluation methods in [24, 7]. For each given image and question pair, both the original LLaVA model and our fine-tuned version provide answers related to the input image. We subsequently submit the questions, and answers from both LLaVA and our model, along with the ground-truth text descriptions, to a text-only GPT-4-based judge. This judge assesses the responses

based on their helpfulness, relevance, accuracy, and level of detail in comparison to the ground-truth descriptions. GPT-4 then allocates an overall performance score ranging from 0 to 10, where higher scores denote superior performance.

The image-question-answer sets used in this evaluation are randomly chosen from the LLaVA-Instruct-150k dataset [24]. According to the judge of GPT-4, the original LLaVA-7B-v0 achieves 8.25/10 when taking ground-truth textual descriptions as references. On the other hand, our model achieves 7.43/10 under the same GPT-4 judge. This shows that our fine-tuned model still preserves a good portion of the reasoning ability of the original LLaVA. It also shows that the multi-round conversation ability is preserved, as the image-question-answer sets include sequences of multi-round conversations.

D. Real World Experiments

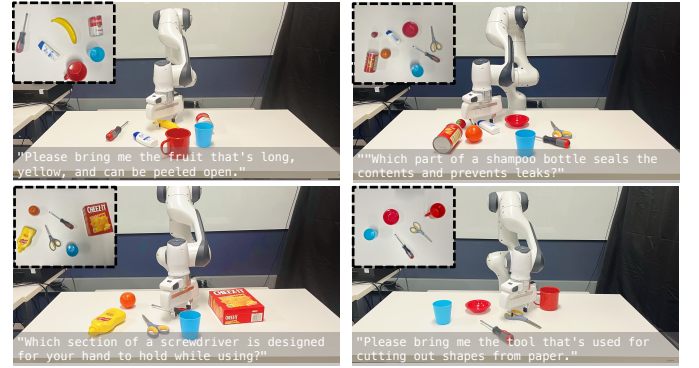


Fig. 6: Real world experiments with implicit instructions.

We conduct real-world experiments to evaluate the proposed reasoning grasping model and the CLIP + GR-ConvNet baseline. A 7-DoF Franka Emika Panda equipped with a Franka Hand parallel gripper was utilized for grasping execution. To perceive the objects, one Azure Kinect was positioned in front of the robot. A cropped 480×480 RGB image of the workspace was utilized as the input to the model. To execute the 2D grasps generated from the model, we always apply a top-down grasp with the height computed using the depth of the grasping point. In our experiments, we randomly place 4-8 objects on a tabletop and provide either explicit or implicit language instruction to the robot to grasp an object or a specific part, as shown in Fig. 6. We select objects that are either identical or similar to those in the dataset. We deliberately excluded any objects or parts that are not infeasible for grasping, such as items too wide for the robot’s gripper. The performance is evaluated using three evaluation metrics: the correctness of generating special tokens and grasping target names, the accuracy of the output grasp pose, and the success in lifting the object. The results are reported in Table IV.

Overall our proposed model can generate accurate special tokens and target names 39 out of 40 trials. This again shows the excellent reasoning capabilities of our model to interpret the various instructions. The failure case involved the model incorrectly identifying a screwdriver instead of scissors in

TABLE III: Grasp prediction accuracy on the reasoning grasping dataset.

Models	Explicit Instruction Object Grasping			Implicit Instruction Object Grasping			Explicit Instruction Part Grasping			Implicit Instruction Part Grasping		
	R@1	R@2	R@3	R@1	R@2	R@3	R@1	R@2	R@3	R@1	R@2	R@3
LLaVA → GR-ConvNet	26.80%	39.17%	44.66%	18.34%	23.76%	31.29%	15.38%	23.52%	37.93%	0.92%	3.84%	8.53%
CLIP + GR-ConvNet	50.40%	62.88%	66.38%	44.35%	57.15%	61.09%	43.78%	55.71%	59.61%	28.88%	42.34%	47.37%
Reasoning Grasping	64.49%	84.11%	86.92%	63.55%	73.83%	77.57%	59.81%	70.01%	81.31%	61.68%	71.96%	83.18%
LLaVA → GR-ConvNet (with depth)	31.88%	40.72%	45.95%	20.39%	25.77%	34.84%	20.11%	29.30%	36.63%	1.36%	2.77%	11.95%
CLIP + GR-ConvNet (with depth)	54.11%	66.50%	69.57%	45.89%	58.11%	62.50%	45.53%	59.51%	62.27%	33.06%	48.55%	52.25%
Reasoning Grasping (with depth)	60.75%	71.95%	76.63%	57.94%	70.09%	77.46%	69.16%	72.90%	77.55%	64.48%	75.70%	80.37%

TABLE IV: Real-world experiment results.

		CLIP + GR-ConvNet	Reasoning Grasping
Object Grasping with Explicit Instruction	Token Accuracy	N/A	10/10
	Pose Accuracy	5/10	7/10
	Execution Success	2/10	5/10
Object Grasping with Implicit Instruction	Token Accuracy	N/A	9/10
	Pose Accuracy	3/10	5/10
	Execution Success	2/10	3/10
Part Grasping with Explicit Instruction	Token Accuracy	N/A	10/10
	Pose Accuracy	5/10	7/10
	Execution Success	3/10	4/10
Part Grasping with Implicit Instruction	Token Accuracy	N/A	10/10
	Pose Accuracy	2/10	6/10
	Execution Success	2/10	4/10

response to a prompt requesting a tool to open a package box. We classified it as a failure solely because it did not match the specific pairing in our dataset, which serves as the basis for our evaluation metrics. This example underscores that even when the model’s response is technically incorrect, it is still reasonable. In terms of grasping performance, we assessed top-3 output grasp poses and manually selected the most feasible one for execution on the robot arm. We conducted 10 trials for each scenario, totaling 80 trials, to compare reasoning grasping with CLIP + GR-ConvNet baseline. The CLIP + GR-ConvNet baseline and our reasoning grasping model produced 15/40 and 25/40 accurate grasping poses, respectively, with the success rates for lifting objects are 9/40 and 16/40, respectively. The main reasons for accurate poses but fail to lift were due to the object slipping from the gripper or the object being too heavy. The experiment results demonstrate that our model outperforms the baseline in four distinct scenarios, showing its superior ability in reasoning grasping tasks.

VII. CONCLUSION

In this paper, we introduce a novel task of reasoning grasping, where robots need to interpret and generate grasping poses based on implicit instructions. By leveraging the reasoning ability in a multi-modal large language model, our proposed model can understand indirect verbal commands

and generate corresponding grasping poses. In addition, we also present the first dataset for reasoning grasping, derived from the GraspNet-1 billion dataset, featuring implicit human instructions alongside object and part grasping annotations. Experiment results demonstrate that our approach can effectively comprehend implicit instructions and accurately generate corresponding grasping poses. Overall, this work offers valuable insights bridging implicit human instructions and generating precise robot actions.

While the proposed reasoning grasping model shows promising results, the model is limited when generating optimal grasping poses for novel objects and refining poses through conversational interactions. A possible reason is that the variety of objects with grasping annotations is still limited. For future work, the LLaVA-v0-7b model is not the state-of-the-art Visual Language Model. Other models such as LLaVA-1.6 [25] or GPT-4V [48] have shown much better performance and stronger ability in terms of visual understanding. Our model’s architecture, with its modular reasoning and grasping components, allows for flexibility in upgrading to more sophisticated models in the future.

REFERENCES

- [1] Hyemin Ahn, Sungjoon Choi, Nuri Kim, Geonho Cha, and Songhwai Oh. Interactive text2pickup networks for natural language-based human–robot collaboration. *IEEE Robotics and Automation Letters*, 3(4):3308–3315, 2018.
- [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] Stefan Ainetter and Friedrich Fraundorfer. End-to-end trainable deep neural network for robotic grasp detection and semantic segmentation from rgb. In *2021 IEEE International Conference on Robotics and Automation (ICRA)*, pages 13452–13458. IEEE, 2021.
- [4] 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.
- [5] Chilam Cheang, Haitao Lin, Yanwei Fu, and Xiangyang Xue. Learning 6-dof object poses to grasp category-level objects by language instructions. In *2022 International Conference on Robotics and Automation (ICRA)*, pages 8476–8482. IEEE, 2022.
 - [6] Yiye Chen, Ruinian Xu, Yunzhi Lin, and Patricio A Vela. A joint network for grasp detection conditioned on natural language commands. In *2021 IEEE International Conference on Robotics and Automation (ICRA)*, pages 4576–4582. IEEE, 2021.
 - [7] Wei-Lin Chiang, Zhuohan Li, Zi Lin, Ying Sheng, Zhanghao Wu, Hao Zhang, Lianmin Zheng, Siyuan Zhuang, Yonghao Zhuang, Joseph E Gonzalez, et al. Vicuna: An open-source chatbot impressing gpt-4 with 90%* chatgpt quality. See <https://vicuna.lmsys.org> (accessed 14 April 2023), 2023.
 - [8] Amaury Depierre, Emmanuel Dellandréa, and Liming Chen. Jacquard: A large scale dataset for robotic grasp detection. In *2018 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 3511–3516. IEEE, 2018.
 - [9] 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.
 - [10] Yukiyasu Domae, Haruhisa Okuda, Yuichi Taguchi, Kazuhiko Sumi, and Takashi Hirai. Fast graspability evaluation on single depth maps for bin picking with general grippers. In *2014 IEEE International Conference on Robotics and Automation (ICRA)*, pages 1997–2004. IEEE, 2014.
 - [11] Danny Driess, Fei Xia, Mehdi SM Sajjadi, Corey Lynch, Aakanksha Chowdhery, Brian Ichter, Ayzaan Wahid, Jonathan Tompson, Quan Vuong, Tianhe Yu, et al. Palm-e: An embodied multimodal language model. *arXiv preprint arXiv:2303.03378*, 2023.
 - [12] Clemens Eppner, Arsalan Mousavian, and Dieter Fox. Acronym: A large-scale grasp dataset based on simulation. In *2021 IEEE International Conference on Robotics and Automation (ICRA)*, pages 6222–6227. IEEE, 2021.
 - [13] Hao-Shu Fang, Chenxi Wang, Minghao Gou, and Cewu Lu. Graspnet-1billion: A large-scale benchmark for general object grasping. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 11444–11453, 2020.
 - [14] Maximilian Gilles, Yuhao Chen, Tim Robin Winter, E Zhixuan Zeng, and Alexander Wong. Metagrasp-net: A large-scale benchmark dataset for scene-aware ambidextrous bin picking via physics-based metaverse synthesis. In *2022 IEEE 18th International Conference on Automation Science and Engineering (CASE)*, pages 220–227. IEEE, 2022.
 - [15] Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. Lora: Low-rank adaptation of large language models. *arXiv preprint arXiv:2106.09685*, 2021.
 - [16] 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.
 - [17] 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.
 - [18] Yun Jiang, Stephen Moseson, and Ashutosh Saxena. Efficient grasping from rgb-d images: Learning using a new rectangle representation. In *2011 IEEE International conference on robotics and automation*, pages 3304–3311. IEEE, 2011.
 - [19] Yun Jiang, Stephen Moseson, and Ashutosh Saxena. Efficient grasping from rgb-d images: Learning using a new rectangle representation. In *2011 IEEE International conference on robotics and automation*, pages 3304–3311. IEEE, 2011.
 - [20] Chuhao Jin, Wenhui Tan, Jiange Yang, Bei Liu, Ruihua Song, Limin Wang, and Jianlong Fu. Alphablock: Embodied finetuning for vision-language reasoning in robot manipulation. *arXiv preprint arXiv:2305.18898*, 2023.
 - [21] Sulabh Kumra, Shirin Joshi, and Ferat Sahin. Antipodal robotic grasping using generative residual convolutional neural network. In *2020 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 9626–9633. IEEE, 2020.
 - [22] Xin Lai, Zhuotao Tian, Yukang Chen, Yanwei Li, Yuhui Yuan, Shu Liu, and Jiaya Jia. Lisa: Reasoning segmentation via large language model. *arXiv preprint arXiv:2308.00692*, 2023.
 - [23] Ian Lenz, Honglak Lee, and Ashutosh Saxena. Deep learning for detecting robotic grasps. *The International Journal of Robotics Research*, 34(4-5):705–724, 2015.
 - [24] Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. Visual instruction tuning. *arXiv preprint arXiv:2304.08485*, 2023.
 - [25] Haotian Liu, Chunyuan Li, Yuheng Li, Bo Li, Yuanhan Zhang, Sheng Shen, and Yong Jae Lee. Llava-1.6: Improved reasoning, ocr, and world knowledge, January 2024. URL <https://llava-vl.github.io/blog/2024-01-30-llava-1-6/>.
 - [26] Yuhao Lu, Yixuan Fan, Beixing Deng, Fangfu Liu, Yali Li, and Shengjin Wang. VL-Grasp: a 6-Dof Interactive Grasp Policy for Language-Oriented Objects in Cluttered Indoor Scenes. In *2023 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 976–983. IEEE, 2023.
 - [27] 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.

- [28] 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 *2010 IEEE International Conference on Robotics and Automation*, pages 2308–2315. IEEE, 2010.
- [29] Reihaneh Mirjalili, Michael Krawez, Simone Silenzi, Yannik Blei, and Wolfram Burgard. Lan-grasp: Using large language models for semantic object grasping. *arXiv preprint arXiv:2310.05239*, 2023.
- [30] Eric Mitchell, Yoonho Lee, Alexander Khazatsky, Christopher D Manning, and Chelsea Finn. Detectgpt: Zero-shot machine-generated text detection using probability curvature. *arXiv preprint arXiv:2301.11305*, 2023.
- [31] Douglas Morrison, Peter Corke, and Jürgen Leitner. Learning robust, real-time, reactive robotic grasping. *The International journal of robotics research*, 39(2-3):183–201, 2020.
- [32] OpenAI. Chatgpt. URL chat.openai.com.
- [33] Lerrel Pinto and Abhinav Gupta. Supersizing self-supervision: Learning to grasp from 50k tries and 700 robot hours. In *2016 IEEE international conference on robotics and automation (ICRA)*, pages 3406–3413. IEEE, 2016.
- [34] 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.
- [35] Máximo A Roa and Raúl Suárez. Grasp quality measures: review and performance. *Autonomous robots*, 38: 65–88, 2015.
- [36] 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 *2023 IEEE International Conference on Robotics and Automation (ICRA)*, pages 11523–11530. IEEE, 2023.
- [37] Yaoxian Song, Penglei Sun, Pengfei Fang, Linyi Yang, Yanghua Xiao, and Yue Zhang. Human-in-the-loop robotic grasping using bert scene representation. *arXiv preprint arXiv:2209.14026*, 2022.
- [38] Qiang Sun, Haitao Lin, Ying Fu, Yanwei Fu, and Xiangyang Xue. Language Guided Robotic Grasping with Fine-Grained Instructions. In *2023 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 1319–1326. IEEE, 2023.
- [39] Chao Tang, Dehao Huang, Wenqi Ge, Weiyu Liu, and Hong Zhang. Grasppt: Leveraging semantic knowledge from a large language model for task-oriented grasping. *IEEE Robotics and Automation Letters*, 2023.
- [40] Chao Tang, Dehao Huang, Lingxiao Meng, Weiyu Liu, and Hong Zhang. Task-oriented grasp prediction with visual-language inputs. *arXiv preprint arXiv:2302.14355*, 2023.
- [41] Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, et al. Llama: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971*, 2023.
- [42] Georgios Tziafas, Yucheng Xu, Arushi Goel, Mohammadreza Kasaei, Zhibin Li, and Hamidreza Kasaei. Language-guided Robot Grasping: CLIP-based Referring Grasp Synthesis in Clutter. *arXiv preprint arXiv:2311.05779*, 2023.
- [43] Sai Vemprala, Rogerio Bonatti, Arthur Buckner, and Ashish Kapoor. Chatgpt for robotics: Design principles and model abilities. *Microsoft Auton. Syst. Robot. Res.*, 2:20, 2023.
- [44] An Dinh Vuong, Minh Nhat Vu, Hieu Le, Baoru Huang, Binh Huynh, Thieu Vo, Andreas Kugi, and Anh Nguyen. Grasp-anything: Large-scale grasp dataset from foundation models. *arXiv preprint arXiv:2309.09818*, 2023.
- [45] Kechun Xu, Shuqi Zhao, Zhongxiang Zhou, Zizhang Li, Huaijin Pi, Yifeng Zhu, Yue Wang, and Rong Xiong. A joint modeling of vision-language-action for target-oriented grasping in clutter. *arXiv preprint arXiv:2302.12610*, 2023.
- [46] Yang Yang, Yuanhao Liu, Hengyue Liang, Xibai Lou, and Changhyun Choi. Attribute-based robotic grasping with one-grasp adaptation. In *2021 IEEE International Conference on Robotics and Automation (ICRA)*, pages 6357–6363. IEEE, 2021.
- [47] Yang Yang, Xibai Lou, and Changhyun Choi. Interactive robotic grasping with attribute-guided disambiguation. In *2022 International Conference on Robotics and Automation (ICRA)*, pages 8914–8920. IEEE, 2022.
- [48] Zhengyuan Yang, Linjie Li, Kevin Lin, Jianfeng Wang, Chung-Ching Lin, Zicheng Liu, and Lijuan Wang. The dawn of llms: Preliminary explorations with gpt-4v (ision). *arXiv preprint arXiv:2309.17421*, 9(1):1, 2023.
- [49] Hanbo Zhang, Yunfan Lu, Cunjun Yu, David Hsu, Xuguang La, and Nanning Zheng. Invigorate: Interactive visual grounding and grasping in clutter. *arXiv preprint arXiv:2108.11092*, 2021.
- [50] Deyao Zhu, Jun Chen, Xiaoqian Shen, Xiang Li, and Mohamed Elhoseiny. Minigpt-4: Enhancing vision-language understanding with advanced large language models. *arXiv preprint arXiv:2304.10592*, 2023.

APPENDIX A
OBJECT AND PART STATISTICS

	Object Name	Part Name	6D Grasps per Sample	Rect Grasps per Sample	Num. Scenes
0	Cracker Box	-	10971	4672	35
1	Tomato Soup Can	-	16625	8284	34
5	Banana	Stem, Flesh	1298	2941	25
7	Red Mug	Handle, Rim, Body	3363	2045	25
8	Power Drill	Handle, Body, Batterypack, Chuck	2445	4476	26
9	Scissor	Handle, Blade	548	574	22
11	Strawberry	-	130	903	24
14	Peach	-	960	1922	25
15	Pear	-	948	2069	26
17	Orange	-	581	1466	32
18	Knife	Handle, Blade	1975	1526	29
20	Red Screwdriver	Handle, Shaft	1191	1910	29
21	Racquetball	-	443	573	6
22	Blue Cup	Rim, Body	3960	2770	29
36	Daobao Wash Soap	Cap, Body	2594	1744	25
37	Nzskincare Mouth Rinse	Cap, Body	1806	1786	24
38	Daobao Sod	Cap, Body	3256	2227	25
40	Kispa Cleanser	Cap, Body	931	1654	25
41	Darlie Toothpaste	Cap, Body	442	1181	25
43	Baoke Marker	Cap, Body	953	697	21
44	Hosjam Toothpaste Pump	Pump, Body	725	1958	22
46	Dish	Rim, Body	8	22	25
48	Camel	Legs, Head, Body	96	663	25
51	Elephant	Legs, Head, Body, Truck	218	1692	24
52	Rhinocero	Legs, Head, Body, Horn	152	1419	22
57	Weiquan Chicken Bouillon	Cap, Body	2297	2617	25
58	Darlie Toothpaste Box	-	5526	1646	25
60	Black Mouse	-	170	618	25
61	Dabao Facewash	Cap, Body	4683	2297	22
62	Pantene Shampoo	Cap, Body	502	1175	26
63	Head Shoulders Supreme	Cap, Body	777	1700	24
66	Head Shoulders Care	Cap, Body	1338	1269	24
70	Tape	-	4465	1859	25
Total	33	54	76377	64355	100

TABLE V: Selected Objects and Parts for the Training Split of the Dataset

APPENDIX B 3D OBJECT AND PART VISUALIZATION



Fig. 7: 3D Visualization of Selected Objects and its Parts

APPENDIX C
PROMPT FOR INSTRUCTION GENERATION

System Prompt for Instruction Generation

You are tasked with creating specific, indirect questions and instructions that a robot could use to identify and interact with objects based on their names or detailed descriptions provided by users. When an object is given, such as a pair of scissors described as "A handheld cutting instrument with two crossing metal blades pivoted together, typically used for cutting paper or fabric. This particular pair has a black handle with yellow inner grips," you must formulate responses that precisely hint at the object's uses or features without naming it directly. The aim is to enable the robot to deduce the correct object through these indirect cues, enhancing its ability to understand and execute tasks involving the object.

Your output should include:

- **Indirect Questions (5 Total):** Construct questions that indirectly reference the object's specific functions, design attributes, or contexts in which it is used. These questions should guide the robot to consider the essential features or tasks the object is associated with, aiding in its identification without directly mentioning the object's name.
- **Indirect Instructions (5 Total):** Develop instructions that subtly describe actions involving the object or request its utilization for particular tasks. These instructions should be carefully phrased to convey the object's use or physical characteristics indirectly, allowing the robot to infer which object is needed without explicit naming.

Please format your responses as a list of 10 strings, organized as follows: ['Indirect Question 1', 'Indirect Question 2', ..., 'Indirect Question 5', 'Indirect Instruction 1', ..., 'Indirect Instruction 5'].

In structuring your responses, prioritize:

- **Specificity and Relevance:** Use language that precisely hints at the object's attributes or uses, ensuring the robot can accurately identify and grasp the intended object.
- **Directness and Functionality:** While maintaining indirectness, your hints should be straightforward and functional, focusing on enabling the robot to understand and act upon the instructions or questions effectively.
- **Consistency and Clarity:** Ensure each hint is consistently structured and clear, avoiding overly creative or ambiguous phrasing that could confuse the robot's learning process.

This methodical approach is designed to improve the robot's ability to interpret indirect language and identify objects based on functional cues, thereby enhancing its interaction with and manipulation of objects in its environment.

Fig. 8: Example Prompts used for Initial Reasoning Instruction Generation with GPT-4