

LGMCTS: Language-Guided Monte-Carlo Tree Search for Executable Semantic Object Rearrangement

Haonan Chang, Kai Gao, Kowndinya Boyalakuntla, Alex Lee, Baichuan Huang,
Jingjin Yu, Abdeslam Boularias

Abstract— We present LGMCTS, a framework that uniquely combines language guidance with geometrically informed sampling distributions to effectively rearrange objects according to geometric patterns dictated by natural language descriptions. LGMCTS uses Monte Carlo Tree Search (MCTS) to create feasible action plans that ensure executable semantic object rearrangement. We present a comprehensive comparison with leading approaches that use language to generate goal rearrangements independently of actionable planning, including Structformer, StructDiffusion, and Code as policies. We also present a new benchmark, the Executable Language Guided Rearrangement (ELGR) Bench, containing tasks involving intricate geometry. With the ELGR bench, we show limitations of task and motion planning (TAMP) solutions that are purely based on Large Language Models (LLM) such as Code as Policies and Progprompt on such tasks. Our findings advocate for using LLMs to generate intermediary representations rather than direct action planning in geometrically complex rearrangement scenarios, aligning with perspectives from recent literature. Our code and supplementary materials are accessible at <https://lgmcts.github.io/>.

I. INTRODUCTION

Everyday tasks, such as “Set up the kitchen”, involve organizing objects based on verbal instructions, a process that is intuitive for humans but that presents a significant challenge for robots. The semantic rearrangement problem seeks to empower robots with the ability to reorganize a scene according to linguistic descriptions. This challenge necessitates that robots comprehend the task through natural language, and address the corresponding Task And Motion Planning (TAMP) problem effectively.

Traditionally, solving this problem requires formalizing semantic rearrangement into a symbolic representation, clearly defining the goal configuration or constraints, and using formal planners such as STRIPS [1] and PDDL [2], or search-based planners like MCTS [3] to devise a feasible plan. Although effective, this approach demands expert-level knowledge to abstract a problem into a formal representation, limiting accessibility for average users.

To overcome this challenge, numerous recent studies have sought to tackle the problem directly from linguistic inputs and RGB-D observations [4]–[6]. One approach uses multi-modality transformers to establish a correlation between verbal descriptions and object positions using data generated from the simulation. Following work such as StructDiffusion [5] further improved this method by using a diffusion

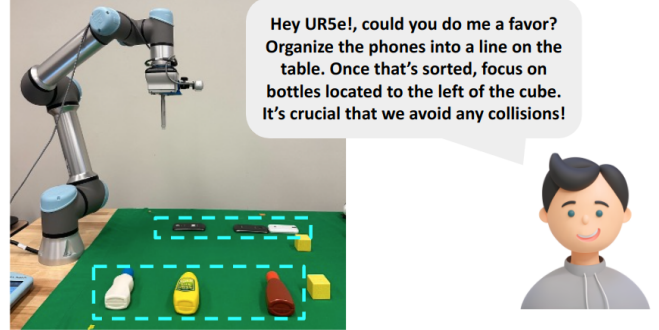


Fig. 1: Robotic Setup: a UR5e robot equipped with a RealSense D455 camera. The task is to re-arrange the objects, which are unknown to the robot, according to a natural language instruction.

model to build the multi-modality solution. However, a common drawback of these methods is that they rely on an offline training stage, which makes them applicable only to trained object categories and spatial patterns.

With the advent of Large Language Models (LLMs), models such as GPT [7] and Llama [8] have demonstrated impressive potential in understanding complex scenarios and exhibiting zero-shot planning capabilities. This has led researchers to explore the utilization of LLMs in solving language-based TAMP problems [9]–[11]. However, despite specific considerations for the feasibility of plans proposed by LLMs, it has been reported that these plans significantly lag in executability and completeness when compared to those crafted by a properly implemented traditional solver designed for the task [12]. This observation has naturally led researchers to seek methods that merge the user-friendliness of LLMs with the robustness of traditional TAMP algorithms such as PDDL, STRIPS, or MCTS. LLM-GROP [13] follows this approach in rearrangement, employing LLMs to parse user tasks from language into pairwise spatial relationship specifications and then calling a sampling-based task and motion planner [14] to generate the plan. A limitation of LLM-GROP is that it can only handle pair-wise relationships, and thus cannot perform complex rearrangement tasks. AutoTAMP [15] uses LLMs to translate natural language into formal representations and then invokes a planner to tackle the problem. AutoTAMP can solve a wide range of TAMP tasks, but it does not apply to general semantic rearrangement where the action space is not discrete and potentially large.

We present Language-Guided Monte-Carlo Tree Search (LGMCTS), a new technique for executable semantic object

rearrangement. Like its predecessors, AutoTAMP and LLM-GROP, LGMCTS leverages LLMs for generating intermediate representations and employs a planner for formulating feasible plans. A key novelty of LGMCTS is the integration of parametric geometric priors for spatial relationship representations. LGMCTS facilitates more nuanced handling of complex geometric relationships among multiple objects, addressing scenarios that require organization beyond simple pairwise interactions such as configurations in lines or rectangles. Additionally, LGMCTS takes a holistic approach by simultaneously considering task planning (goal specification) and motion planning (execution order and intermediate steps). During planning, an obstacle relocation strategy is used to handle obstacles that may block the execution. This coordination ensures that plans are not only semantically coherent but also practically executable, offering a balanced consideration of goal achievement and operational efficiency.

To assess the efficacy of LGMCTS, we introduce the Executable Language-Guided Rearrangement (ELGR) benchmark, featuring over 1,600 varied language queries and robot execution checks. Our evaluations indicate that LGMCTS performs effectively on the ELGR benchmark, especially in comparison with Code as policies and Progprompt in terms of feasibility and semantic consistency of the generated goals. LGMCTS also outperforms Structformer and StructDiffusion in goal generation on the Structformer dataset.

II. RELATED WORKS

A. Learning-based Semantic Rearrangement

The semantic rearrangement problem consists of devising a rearrangement plan that is both semantically congruent with a given language description and physically feasible [16]–[18]. In recent years, this has gained increased traction, particularly as a pivotal application in language-driven robotics. CLIPort [4] took the initial step in this direction by merging CLIP features with a Transporter network. Yet, its design is limited to basic pick-and-place tasks. Structformer [6] advanced the field using a transformer model, by simulating rearrangements with hand-crafted rules and connecting language tokens to object poses. Leveraging Structformer’s dataset, StructDiffusion [5] introduced a pose diffusion model to predict poses from language. Nonetheless, a common shortcoming amongst all these methodologies is their limitation to a single structure or pattern (e.g. circle, line) that they have been trained on, making composite patterns (e.g. rectangle + tower) a persistent challenge. Moreover, the rearrangement goals generated by these methods can be inexecutable.

B. LLM-driven Task And Motion Planning

Recent advancements in LLMs [7], [8] have showcased impressive performance across a broad spectrum of tasks. There has been a growing interest in using LLMs for TAMP [9]–[11], [13], [19]–[33], owing to their few-shot and zero-shot reasoning ability [7], [34], [35]. Grounding language into a sequence of plannable tasks/actions without retraining LLMs was initially explored in [9]. Following

this, SayCan [10] was proposed to facilitate the conversion of LLM-generated plans into robot-executable steps, though it struggled with addressing task execution failures. Inner Monologue [19] improved upon SayCan by incorporating real-time feedback to adjust plan post-execution, yet Inner Monologue is prone to generating suboptimal and infeasible plans. Instead of using LLMs to plan with predefined skills, other approaches such as Code as policies [11] and Progprompt [20] leveraged LLMs for policy code generation, showcasing their potential in behavioral common sense and sequential policy logic. However, they do not show promising results on complex object rearrangement tasks requiring more nuanced spatial context. This is due to the limitations of LLMs’ planning ability over long horizon tasks [12], [28].

Owing to the drawbacks of the aforementioned works, to offer a more reliable and interpretable planning process, recent works [12], [15], [28], [29], [31] emphasize translating natural language commands into intermediate representations that are interpretable by traditional TAMP algorithms. Text2Motion [23] uses LLMs to greedily plan a skill sequence combined with a geometric feasibility planner to ensure that the geometric dependencies are addressed. However, Text2Motion’s hybrid LLM planner is less efficient in large spaces than planners such as MCTS [3]. LLM-GROP [13] translates language instructions to symbolic spatial relationships with LLMs and employs a task and motion planner named GROP [14] to perform the rearrangement. Although GROP is optimized for efficiency and feasibility, LLM-GROP is limited by its focus on simple object rearrangements due to its treatment of multi-object semantic relationships as pairwise reasoning. AutoTAMP [15] employs LLMs for generating and validating STL representation from natural language and utilizes a formal STL planner [36] for generating optimal trajectories. Although effective across a spectrum of tasks, its applicability is limited in non-discrete action spaces, as is the case in semantic rearrangement tasks.

To address these limitations, we introduce LGMCTS, a new approach that incorporates parametric geometric priors and that is guided by MCTS’s efficient exploration to provide a feasible TAMP solution for complex rearrangement tasks.

III. PRELIMINARIES

A. Problem Formulation

Semantic rearrangement is the task of rearranging a scene according to a series of natural language descriptions. One key insight here is that a goal described by language is usually a distribution rather than a single position. For example, “Put the mug at the right side of bowl” refers to a uniform distribution among a region that is right to the bowl. If we know the position distribution for each object, we just need to sequentially sample the poses for each object. The semantic rearrangement problem is then converted to a sequential sampling problem.

With this insight, we define the task of semantic rearrangement as follows. The robot is given as input a scene with objects from a set $O_S = \{o_1, o_2, \dots, o_N\}$ and a command L , where L is a pure natural language command that implies a

desired distribution list $\mathcal{F} = \{f_i : p(o_i) \sim f_i | o_i \in O_R\}$, where $p(o_i)$ refers to the position of object o_i . Here, $O_R \subseteq O_S$ denotes the objects requiring an action based on L , and f_i indicates the desired pose distribution for each object. The objective is to identify an optimal action sequence, $A = (a_t)_{t=1}^H$, where each action a_t corresponds to moving an object o_i to a sampled position $p(o_i)$, with the objective to achieve a goal arrangement aligning L , i.e., $\prod_{o_i \in O_R} f_i(p_i) > 0$ and minimizing the number of action steps H . Noticeably, A includes not only movements of objects $o \in O_R$, but also those of distracting objects, denoted as O_D , with $O_D \subseteq O_S$.

B. Monte Carlo Tree Search (MCTS)

We provide here a brief reminder of the MCTS [3] technique. A typical MCTS algorithm iteratively builds a search tree by performing the following four operations.

- 1) **Selection.** On a fully expanded node (all the children nodes have been visited), MCTS selects to explore the branch with the highest Upper Confidence Bound (UCB),

$$\arg \max_a \left(\frac{w(f(s,a))}{n(f(s,a))} + C \sqrt{\frac{\log(n(s))}{n(f(s,a))}} \right), \quad (1)$$

where $f(s,a)$ is the child node of state s after action a , $w(\cdot)$ and $n(\cdot)$ are respectively cumulative rewards and the number of visits to a state.

- 2) **Expansion.** On a node that is not fully expanded, MCTS selects an action that has not been attempted yet.
- 3) **Simulation.** Given a node and a selected action, MCTS simulates a sequence of actions and receives a reward.
- 4) **Back-Propagation.** MCTS passes the terminal reward to ancestor nodes to update their cumulative expected rewards, which indicate the quality of the branch.

At each iteration, MCTS starts from the root node. When all the child nodes of the current node are visited, MCTS selects a child node with the UCB formula. When some child nodes of the current node are unvisited, MCTS expands by randomly selecting a new action and performing a simulation to reach a new child node. The new node returns a reward, which is back-propagated to all the ancestor nodes.

IV. METHOD

In a nutshell, LGMCTS starts by calling an LLM to parse the language description L to a list of spatial distributions $\mathcal{F} = \{f_i : p(o_i) \sim f_i | o_i \in O_R\}$. Then it uses an MCTS-based procedure to find a physically feasible action sequence A to rearrange the scene according to distributions $\mathcal{F}$.

A. Language Parsing & Object Selection

During the language parsing stage, we parse L into $\mathcal{F} = \{f_i : p(o_i) \sim f_i | o_i \in O_R\}$. Similar to previous methods [11], we conduct an automated prompt engineering to guide the LLM to perform the parsing. Fig. 2 showcases how prompt engineering is implemented. Essentially, the language model translates user requirements into structured goal configurations and constraints that guide task execution.

Consider the example depicted in Fig. 2. First comes the **system prompt** providing guidelines for the LLM to follow

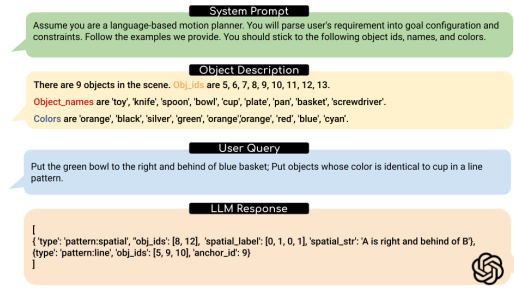


Fig. 2: An example of language parsing. We are using GPT-4 [7] in this work.

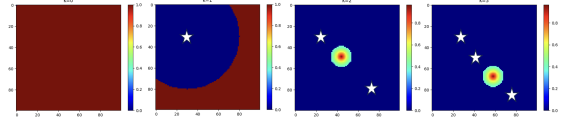


Fig. 3: Visualization of (x,y) prior for 'line' pattern. From left to right: $K = 0$, $K = 1$, $K = 2$, $K = 3$, where $K = |O_R^{sampled}|$, the number of sampled object poses. White star marks are sampled poses. When $K = 0$, the pose can be sampled anywhere. When $K = 1$, it needed to sampled outside a circle region. After that, all poses will be sampled along the line defined by the first two poses.

when interpreting user queries. Then, we need to provide an **object description** such as semantic labels, colors, and IDs for the objects in the scene. In practice, we use the Recognize Anything Model (RAM) [37], [38] for producing the semantic labels, and a color detector to determine the colors of the objects in the scene. A unique ID is assigned to each object (using for the whole planning). Finally, we provide a **user query** describing the rearrangement goal. A structured answer is returned from the LLM. LLM's answers suggest how many patterns there are and which objects should be selected to form the pattern.

B. Parametric Geometric Prior

As mentioned in Section III-A, if we know the goal position distribution for each object, the semantic rearrangement problem can be converted to a sequential sampling problem. There are multiple details we need to pay attention to in this process. (1) The position distribution is not fixed but varies with the progress of sampling. For example, if we say objects A, B, and C need to be put into a line, and we sample in the order of A, B, and C, then the distributions of A and B are actually unbounded, and C must be placed on the line defined by the positions of A and B. (2) The position distribution should be collision-free. (3) One can build a database for many different spatial distributions, but we need to have a flexible mechanism to associate LLM's inferred distribution types with items in the database.

We show in the following how to compute the pose distribution function $f_i(p_i | P_S, L)$ for each object o_i in the set O_R during the sequential pose sampling process. Here, p_i denotes the pose (x, y, θ) of an object o_i , P_S represents the list of current poses for all objects in the scene, and L is the natural language command.

We compute $f_i(p_i | P_S, L)$ as an element-wise product of

two components, a pattern prior function $f_{\text{prior}}(p_i|P_R, L)$ and a boolean function $f_{\text{free}}(p_i|P_S)$ of the workspace,

$$f_i(p_i|P_S, L) = f_{\text{prior}}(p_i|P_R, L) \times f_{\text{free}}(p_i|P_S) \quad (2)$$

In this equation, P_R is the list of current poses of all objects in O_R . We note that P_R is a subset to P_S . The function $f_{\text{free}}(p_i|P_S)$ is set to 1 if p_i is collision-free from the remaining poses in P_S , and to 0 otherwise. f_{free} is determined by running a 2D collision simulation using point cloud observations for each object $o_i \in O_S$.

Our primary focus is to determine $f_{\text{prior}}(p_i|P_R, L)$. To this end, we employ an approach akin to the one described in [10]. We maintain a database comprising a collection of predefined prior functions. Each of these functions is linked with one or more Sentence-BERT embeddings, acting as keys. The function corresponding to the best matching key is selected, as follows,

$$f_{\text{prior}}(p_i|P_R, L) = f_{\text{prior}}^K(p_i|P_R) \text{ with } K = \underset{K \in \text{database}}{\text{argmax}} (\Theta_k \cdot \Theta(L))$$

Here, Θ_k is the K^{th} key in the database, and $\Theta(L)$ is the Sentence-BERT embedding generated from the language instruction L . In summary, the most suitable prior function from the database is selected based on $\Theta(L)$.

We now delve into the definition of $f_{\text{prior}}^K(p_i|P_R)$ in the context of our work. We use a unique model for representing various distributions by employing parametric curves. A parametric curve can be expressed as $(x, y) = \gamma(t, \kappa)$, where t ranges from 0 to 1 and κ is a set of curve-defining parameters. In our work, κ is modeled as a function of two 2D positions, denoted as $\kappa(p_0, p_1)$. Therefore, for each pattern, we define two functions: γ and κ .

Given that pattern prior f_{prior} is used inside a sequential sampling process (check Section IV-C for more details), the prior distribution needs to be iteratively updated to capture the history of sampling. Consequently, we further categorize O_R based on whether the objects have been sampled in the current branch of the MCTS-Planner. Subsets O_R^{sampled} and $O_R^{\text{unsampled}}$ denote the sampled and non-sampled objects, respectively. Thus, $O_R = O_R^{\text{sampled}} \cup O_R^{\text{unsampled}}$.

The probability $f_{\text{prior}}^K(p_i|P_R)$ of sampling a pose $p_i = (x_i, y_i, \theta_i)$ for next object $o_i \in O_R^{\text{unsampled}}$ is given as follows.

- If $|O_R^{\text{sampled}}| = 0$ then $(x_i, y_i, \theta_i) \sim U$, suggesting that the first object can be placed arbitrarily.
- If $|O_R^{\text{sampled}}| = 1$ then $\sqrt{(x_i - x_0)^2 + (y_i - y_0)^2} \leq \delta$ and $(x_i, y_i, \theta_i) \sim U$, imposing that the second object must be sampled uniformly at a position that is distanced from the first by at most δ .
- If $|O_R^{\text{sampled}}| > 1$ then $(x_i, y_i) = \gamma(\frac{|O_R^{\text{sampled}}|}{|O_R|}, \kappa(p_0, p_1)) + \varepsilon$ where $\varepsilon \sim G(0, \sigma)$, here G represents a Gaussian distribution with mean zero and variance σ , and $\theta = \text{atan2}(1, \gamma'(\frac{K}{N}))$. θ represented the rotation angle of the object.

Our parametric geometrical representation enables us to model any geometric shapes that can be written into the

format of a parametric curve. In our current implementation, we defined shapes such as “line,” “circle,” “rectangle,” “tower,” “spatial:left,” “spatial:right” and so on. Due to space constraints, we refrain from elaborating on the definitions of γ and κ for all these predefined patterns. Fig. 3 illustrates an example of a parametric geometric prior. Noticeably, we divide patterns into “ordered” and “unordered” based on whether the pattern requires an execution sequence.

Note that in our work, language instruction L is typically composed of multiple instructions that deal with different subsets of objects. Specifically, L can be interpreted as a list $\{L_i\}$ of sub-instructions L_i . For example, L can be a composite instruction: $L = \{L_1, L_2\}$, where L_1 refers to placing objects A , B , and C in a line, and L_2 refers to placing object A on the left of B . Each sub-instruction $L_i \in L$ is associated with a subset of objects $O_{R_i} \subseteq O_R$. In our example, $O_{R_1} = \{A, B, C\}$ and $O_{R_2} = \{A, B\}$. In the sequential sampling process described before, we presented the case of a single language instruction for simplicity, but the same process is used for sampling poses given by a complex instruction.

C. Monte-Carlo Tree Search (MCTS) for TAMP

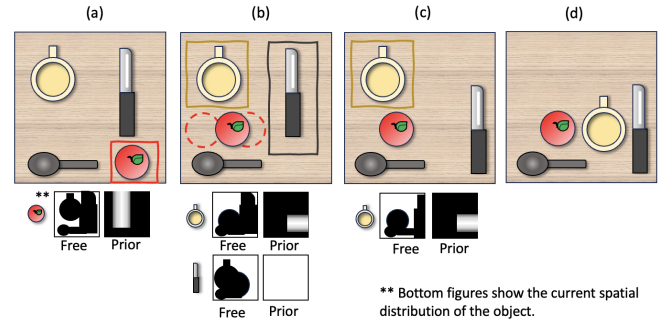


Fig. 4: A minimal example illustrates our MCTS-Planner’s aim to arrange a table. The language description provided is: “Can you please put the apple behind the spoon? And I also want the cup at the right of the apple.” The top row displays the current scene arrangement, while the bottom row shows the f_{prior} and f_{free} for the object being manipulated. $f = f_{\text{prior}} \times f_{\text{free}}$. In spatial distribution figures, black represents probability 0, and white probability 1.

As previously mentioned, our rearrangement problem can be formulated as a sequential task. In each step, we sample a pose p_i for each object $o_i \in O_R$ according to a pose distribution f_i selected by the LLM and computed as described in Section IV-B. Once we complete all samplings, all objects will have been placed in their desired locations. However, this task cannot be executed by the robot in a naive sequential order, as the rearrangements made in previous steps may obstruct subsequent sampling. Therefore, we propose a task and motion planner based on the MCTS algorithm to simultaneously arrange the task and address object relocation issues. The objective of our MCTS-Planner is to fulfill all object position requirements defined by $\mathcal{F}$, the spatial distribution list. In the MCTS-Planner, we maintain a tree where each node s in the tree comprises the current objects’ poses, $\{p_1, p_2, \dots, p_N\}$, and the remaining object spatial requirements $\mathcal{F}_r$, where $\mathcal{F}_r \subseteq \mathcal{F}$.

The MCTS-Planner operates through four phases: selection, expansion, simulation, and back-propagation, as described in Section III-B, adhering to the methodology from [39] except for the simulation stage. The reward for transitioning to a new state s is quantified by the reduction in the number of spatial requirements, that is, $|\mathcal{F}| - |\mathcal{F}_r|$. This reward structure mirrors the approach in [39], aiming to steer the search towards branches that efficiently move objects to their goal poses. The simulation phase is elaborated in Algorithm 1. This phase begins with the MCTS state s and an attempted action of placing an object o_i in a pose $p_i \sim f_i$, where f_i is the pose distribution explained in the previous section. If the targeted object o_i is not reachable, e.g. there exist some obstacles above it, a reachable obstacle is randomly selected from those above it, and a collision-free pose within the workspace is sampled for relocation (Lines 2-5). If o_i is reachable, an attempt is made to sample its pose with distribution f_i (Line 7). If the sampled pose is not collision-free, we will randomly choose an obstacle in the collision and relocate it to a collision-free pose within the workspace (Lines 9-13). Note that obstacle relocation may add previously placed objects back to the requirement list $\mathcal{F}_r$. However, this obstacle relocation strategy increases the robustness of our rearrangement planner, especially when the environment is cluttered. Although MCTS operates as an anytime search algorithm, our MCTS-Planner implementation returns the first solution it finds. Furthermore, we have proved that the MCTS-Planner is probabilistically complete within our specified framework in Proposition 4.1.

Fig. 4 presents an illustrative example of the MCTS-Planner at work. Owing to space constraints, we will only explore three simulation steps along the branch that yield a feasible solution. In this scenario, the user requests, “Can you please put the apple behind the spoon? And I also want the cup to be at the right of the apple.” In response, the LLM generates the spatial distribution list $\mathcal{F} = \{f_1, f_2\}$, where f_1 is for positioning the apple behind the spoon, and f_2 is for placing the cup to the right of the apple. Fig. 4(a) illustrates the initial arrangement of objects. Given the dependency of f_2 on the apple’s position, K actions will be sampled for f_1 , but none for f_2 in this initial setup. Fig. 4(b) depicts the outcome of an action of sampling a pose from f_1 , where the apple is relocated according to f_1 . The dashed-line circles represent the other $K - 1$ actions originating from the root node. After sampling f_1 , we are left with $\mathcal{F}_r = \{f_2\}$ as shown in Fig. 4(b), and an attempt is made to position the cup to the right of the apple. However, the goal position is in collision with the knife, so we need to relocate the knife. Fig. 4(c) demonstrates the knife’s relocation, maintaining $\mathcal{F}_r = \{f_2\}$. Ultimately, Fig. 4(d) showcases the final planning result.

Proposition 4.1: MCTS-Planner is probabilistic complete.

Proof: In the context of semantic rearrangement with distribution list $\mathcal{F}$, we consider a feasible sequence A^* that moves objects to achieve a final state A_f , where $f_i(A_f[o_i]) > 0$ for each object o_i in set O_R . We assert that as the number of iterations K increases, the probability p that MCTS finds a sequence meeting the goal approaches 1: $\lim_{K \rightarrow \infty} p = 1$.

Algorithm 1: Simulation

Input : s : an MCTS state,
 f_i : a pose distribution (place o_i in a pose $p_i \sim f_i$).
Output: (o, p) : a rearrangement action (place o in pose p).
1 **if** o_i is not reachable **then**
2 Select a reachable object o that is blocking object o_i ;
3 $p \leftarrow \text{uniformSampling}(o, s)$;
4 **if** p exists **then return** (o, p) ;
5 **else return failure**;
6 **else**
7 $p \leftarrow \text{sampling}(o_i, s, f_i)$;
8 **if** collisionFree(o_i, p, s) **then return** (o_i, p) ;
9 **else**
10 $o \leftarrow$ Randomly choose an obstacle in collision;
11 $p \leftarrow \text{uniformSampling}(o, s)$;
12 **if** p exists **then return** (o, p) ;
13 **else return failure**;

First, we prove that there is an action sequence A_0^* whose actions can all be generated by MCTS-Planner. Note that in MCTS-Planner, an action satisfies either of the two rules: **R1**: Move an object o_i to goal: it requires $o_i \in O_R$ and the new pose p satisfies goal pose requirements (i.e., $f_i(p) > 0$); **R2**: Obstacle relocation: the old pose p of the moved object makes other objects not reachable or it lies in others’ goal regions, i.e., $\exists f \in \mathcal{F}_r, s.t. f(p) > 0$ (Line 2, 10). We construct A_0^* by reordering and deleting actions in A^* as follows: (1) If an action in A^* satisfies **R1** or **R2**, we add the action to A_0^* . Otherwise, we delay the addition until it satisfies the rules; (2) If the action still cannot satisfy the rules before we examine the next action in A^* for the same object, we delete the former action. This process ensures A_0^* comprises only MCTS-Planner viable actions, leading to the desired arrangement A_f .

Next, we prove that the probability for MCTS to find an action sequence in the “neighborhood” of A_0^* approaches 1 as K increases. For each intermediate state s of A_0^* , denote $(A_0^*[s].o, A_0^*[s].p)$ the rearrangement action that A_0^* chooses at s . Let $r := (\min_{1 \leq i \leq |A_0^*|} C_i) / 2|A_0^*|$, where C_i is the minimum distance making the placement pose of the i^{th} action of A_0^* invalid (out of goal distribution or in collision). For an intermediate state s of A_0^* , let p_r^s be the probability that after K actions, there is an action moving $A_0^*[s].o$ to the r -neighborhood of $A_0^*[s].p$. We have

$$\lim_{K \rightarrow \infty} p \geq \lim_{K \rightarrow \infty} \prod_{s \in A_0^*} p_r^s \geq \lim_{K \rightarrow \infty} (1 - (1 - \frac{\pi r^2}{|W|N})^K)^{|A_0^*|}$$

where $|W|$ is the size of the workspace. Specifically, the second function indicates the probability of finding a feasible action sequence, where the intermediate state object poses are maintained in the neighborhood of those in A_0^* . In the third function, $(\pi r^2)/(|W|N)$ is the lower bound of the probability of choosing an action in s , moving $A_0^*[s].o$ to the r -neighborhood of $A_0^*[s].p$ and the base of the $|A_0^*|$ exponent is a lower bound of p_r^s . Since K is the only variable in the third function, $\lim_{K \rightarrow \infty} p = 1$. ■

Method	Line (4295)	Circle (3416)	Tower (1335)	Dinner (2440)
LFSP* [11], [20]	41.16%	51.75%	88.80%	27.05%
Structformer [6]	47.24%	62.64%	99.10%	28.36%
StructDiffusion [5]	61.49%	81.41%	98.95%	69.38%
LGMCTS (Ours)	95.99%	95.25%	100%	100%

TABLE I: Efficacy of LFSP, Structformer, StructDiffusion, and LGMCTS across diverse rearrangement tasks (task counts indicated) from the Structformer dataset. *Due to budget constraints, the LLM baseline LFSP are evaluated on 1150 (10%) randomly selected scenes of the Structformer dataset.

V. EXPERIMENTS

A. Baselines

We compare our approach with the following baselines.

Structformer [6]. It is a multi-modal transformer specifically designed for language-guided rearrangement tasks.

StructDiffusion [5]. It employs a diffusion model combined with a learning-based collision checker for pattern pose generation.

LLMs as Few-Shot Planners [11], [20]. We integrate *Code as Policies* and *Progprompt* into our evaluation pipeline, where the former generates policy code and the latter Pythonic code. As we cannot directly use the generated code, to streamline the input to our TAMP planner, our setup modifies the output as a sequence of actions (object IDs and their target poses). Initially, LLM processes complete scene details—including object names, IDs, textures, initial poses, and region boundaries for the rearrangement. We then instruct the LLM to take the natural language command and produce the optimal action sequence that contains an ordered list of object IDs and goal poses of the considered objects for rearrangement. However, for evaluating the Structformer dataset, we consider the structured goal specification pre-available in the dataset as the input to the LLM as it can infer the action plan from this intermediate representation. This approach avoids redundancy, as generating a natural language command would just restate the same specification for the goal rearrangement. In evaluating both datasets, we provide a few scenes as examples where the output format is clearly defined to the LLM with ground-truth optimal sequence. This baseline is named LLMs as Few-Shot Planners (LFSP).

Pose+MCTS. The Pose+MCTS (PMCTS) approach assumes that a collision-free and semantically aligned goal pose is provided. However, direct execution of this pose might be hindered if the target space is already occupied. To address this, we utilize MCTS to search for a viable plan to place objects in their predetermined goal poses. MCTS is only used as a motion planner. This method follows a two-step approach of using goal poses independently of task planning.

B. Structformer Dataset

We use the test set from the Structformer dataset to evaluate the goal pose generation ability. This dataset is composed of approximately 11,500 rearrangement tasks, categorized into four patterns: line, circle, tower, and dinner. A rearrangement plan is considered successful if it adheres to language constraints and is collision-free, except in the

“tower” task where collisions are inevitable. The “dinner” task is approached as a composition of patterns, involving the arrangement of items like plates, bowls, and utensils into a “tower” for plates and bowls, with other items lined up beside it. In both Structformer and StructDiffusion’s experimental setup, object selection for rearrangement is based on the object’s shape and size. Our evaluation setup does not involve object selection based on shape and size. Hence, to adapt them to our evaluation setup, we provide those two baselines with ground-truth object selection. Since the tasks already specify which objects to rearrange for the single pattern rearrangement, based on language instructions, we did not use the LLM parser in LGMCTS for this dataset.

As shown in TABLE I, LGMCTS demonstrated superior performance in all four rearrangement task categories, achieving remarkable success rates as follows: 95.99% for “line”, 95.25% for “circle”, and 100% for both “tower” and “dinner”. LFSP performs the least among the baselines due to the inability of LLMs to produce goal patterns with high geometric fidelity. While StructDiffusion showed improvement over Structformer, it did not match LGMCTS’s effectiveness. For a visual comparison, Fig. 6 illustrates LGMCTS’s success in a “circle” task scene, highlighting its more actionable goal poses that contribute to higher rearrangement success rates. Conversely, Structformer and StructDiffusion tend to generate goal arrangements with a higher collision risk, leading to lower success rates. The evaluation of this dataset strictly assesses the accuracy of collision-free geometric patterns in goal poses, disregarding the executability of plans—a notable limitation of the dataset. This limitation is addressed through our proposed ELGR-Benchmark (refer to Section V-C). PMCTS method is introduced through our benchmark to solely verify the executability of plans. As PMCTS deals with motion planning using collision-free ground-truth poses, it was not included in further comparisons on this dataset due to its inherent access to goal poses.

C. ELGR-Benchmark

Existing datasets for semantic object rearrangement, such as Structformer, are limited in that they typically feature only one pattern per scene and do not include crowded scenarios. They also fail to address the challenge of feasibility, particularly when starting configurations are infeasible like one object being placed under another. To bridge these gaps, we introduce ELGR-Bench (Executable Language-Guided Rearrangement Benchmark), which incorporates scenarios with infeasible starting configurations, including tasks that require unstacking and appropriate placement of unstacked objects before the actual rearrangement. Importantly, this new benchmark presents a novel task termed the “multi-pattern task”, which requires multiple pattern goals to be satisfied during the rearrangement process. In this benchmark, we are considering common shapes such as “line”, “circle”, “rectangle” and “spatial” (left/right, front/behind, left/right + front/behind). For each scene, we randomly compose two of the aforementioned patterns and create the multi-pattern task. Success is measured based on the executability of the

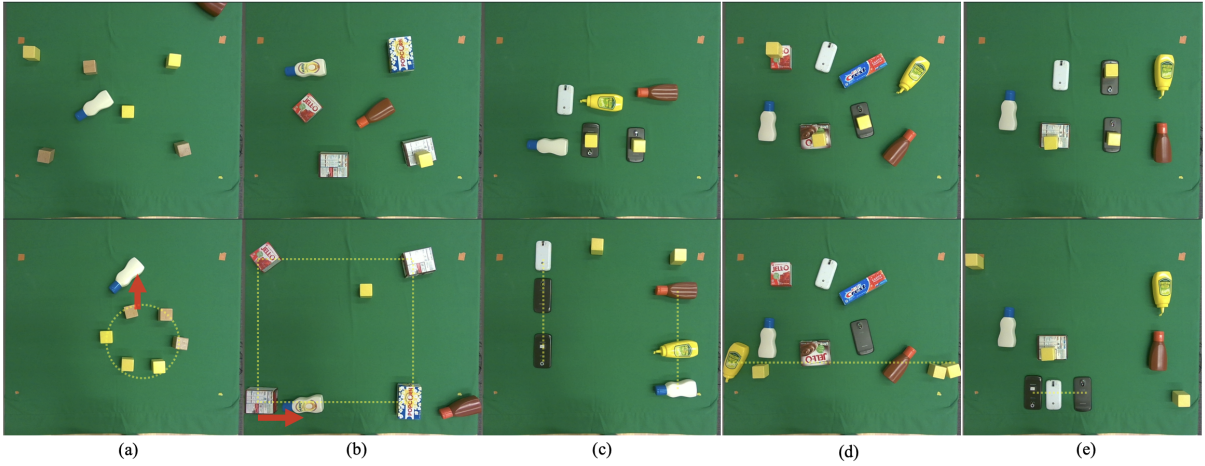


Fig. 5: Real world demonstration with a UR5e robot. The language instructions for the five scenes are: (a) “Move all blocks into a circle; while put the white bottle behind one block;” (b) “Put all boxes into a rectangle; and move the white bottle to the right of one box;” (c) “Move bottles into a line; and formulate all phones into another line;” (d) “Formulate all yellow objects into a line;” (e) “Set all phones into a line;”. The top row images show the initial scenes and the bottom ones show the results of using LGMCTS on the UR5e. Dotted lines imply a shape pattern and red arrows indicate a spatial pattern (left, right, front, back). These real robot experiments show that LGMCTS can parse complex language instructions and also deal with infeasible start configurations as well as pattern composition.



Fig. 6: Compared to Structformer and StructDiffusion, LGMCTS ensures a collision-free goal arrangement in all experiments.

Method	SR_p	SR_{ep}
LFSP [11], [20]	100%	45.2%
Structformer [6]	n.a.	n.a.
StructDiffusion [5]	n.a.	n.a.
PMCTS	82.9%	74.1%
LGMCTS (Ours)	90.9%	79.2%

TABLE II: SR_{ep} , the executable plan success rate, reflects both planning success and the success of executing these plans, indicating if the final positions of objects meet the criteria set by language-based constraints. SR_p , a part of SR_{ep} , only tracks planning success, with PMCTS and LGMCTS capped at 10,000 planning steps. If planning with MCTS exceeds the limit, often due to dense object placement in the scene, the motion planning is considered a failure.

generated plan and its adherence to semantic requirements. ELGR-Bench builds upon the VIMA-Benchmark [40].

In our benchmark, we compared LGMCTS against two baselines: LFSP and PMCTS, excluding Structformer and StructDiffusion as they cannot handle composite geometric patterns. LFSP, leveraging an LLM, plans goal poses and action sequences simultaneously, while PMCTS, follows a two-step method using a given goal pose and then using MCTS for action planning. LGMCTS uniquely combines goal generation with action planning, aiming for more executable outcomes. As shown in TABLE II, LFSP demonstrates a 100% planning success rate through LLM’s capability to generate action plans based on natural language commands and scene context. Nonetheless, over 50% of these plans are inexecutable, as indicated by the SR_{ep} scores. LGMCTS,

however, manages a 90% success rate in generating action plans, with about 80% being executable. This performance not only underlines the limitations of LLMs in direct TAMP solving but also showcases LGMCTS’s advantage over the two-step PMCTS approach, even when PMCTS is provided with accurate and feasible goal poses.

D. Real Robot Experiments

We qualitatively evaluated our system using a UR5e robot equipped with a D455 depth camera. The setup of the robot is shown in Fig. 1. We employed the Recognize-Anything-Model (RAM) [37], [38] and an HSV-based color detector to detect object semantics and colors. Selected queries and their corresponding execution outcomes are presented in Fig. 5. We considered five different language instructions involving various objects and initial configurations. For example, Fig. 5(b) illustrates the experiment with “Put all boxes into a rectangle, and move the white bottle to the right of one box.” This experiment involves pattern composition, requiring simultaneous consideration of “line” and “to the right of” constraint. Additionally, this scene presented an infeasible initial configuration, necessitating the removal of the yellow block before moving the gelatin box. Each experiment presented distinct challenges; for more details, refer to Fig. 5. These real-world robot experiments underscore the capabilities of LGMCTS in complex real-world settings.

VI. CONCLUSION

We introduced LGMCTS, a new framework for tabletop, semantic object rearrangement tasks. LGMCTS stands out by accepting free-form natural language input, accommodating multiple pattern requirements, and jointly solving goal pose generation and action planning. Its main limitation is the extended execution time for complex scenes, highlighting the need for improved tree search efficiency. Future research should focus on adapting LGMCTS to more complex rearrangement scenarios.

REFERENCES

- [1] R. E. Fikes and N. J. Nilsson, "Strips: A new approach to the application of theorem proving to problem solving," *Artificial intelligence*, vol. 2, no. 3-4, pp. 189–208, 1971.
- [2] M. Fox and D. Long, "Pddl2.1: An extension to pddl for expressing temporal planning domains," *Journal of artificial intelligence research*, vol. 20, pp. 61–124, 2003.
- [3] R. Coulom, "Efficient selectivity and backup operators in monte-carlo tree search," in *International conference on computers and games*. Springer, 2006, pp. 72–83.
- [4] M. Shridhar, L. Manuelli, and D. Fox, "Cliport: What and where pathways for robotic manipulation," in *5th Annual Conference on Robot Learning*, 2021.
- [5] W. Liu, T. Hermans, S. Chernova, and C. Paxton, "Strucdiffusion: Object-centric diffusion for semantic rearrangement of novel objects," in *Workshop on Language and Robotics at CoRL 2022*, 2022.
- [6] W. Liu, C. Paxton, T. Hermans, and D. Fox, "Structformer: Learning spatial structure for language-guided semantic rearrangement of novel objects," in *2022 International Conference on Robotics and Automation (ICRA)*. IEEE, 2022, pp. 6322–6329.
- [7] T. Brown, B. Mann, N. Ryder, M. Subbiah, J. D. Kaplan, P. Dhariwal, A. Neelakantan, P. Shyam, G. Sastry, A. Askell *et al.*, "Language models are few-shot learners," *Advances in neural information processing systems*, vol. 33, pp. 1877–1901, 2020.
- [8] H. Touvron, L. Martin, K. Stone, P. Albert, A. Almahairi, Y. Babaei, N. Bashlykov, S. Batra, P. Bhargava, S. Bhosale *et al.*, "Llama 2: Open foundation and fine-tuned chat models," *arXiv preprint arXiv:2307.09288*, 2023.
- [9] W. Huang, P. Abbeel, D. Pathak, and I. Mordatch, "Language models as zero-shot planners: Extracting actionable knowledge for embodied agents," in *International Conference on Machine Learning*. PMLR, 2022, pp. 9118–9147.
- [10] A. Brohan, Y. Chebotar, C. Finn, K. Hausman, A. Herzog, D. Ho, J. Ibarz, A. Irpan, E. Jang, R. Julian *et al.*, "Do as i can, not as i say: Grounding language in robotic affordances," in *Conference on Robot Learning*. PMLR, 2023, pp. 287–318.
- [11] J. Liang, W. Huang, F. Xia, P. Xu, K. Hausman, B. Ichter, P. Florence, and A. Zeng, "Code as policies: Language model programs for embodied control," in *2023 IEEE International Conference on Robotics and Automation (ICRA)*. IEEE, 2023, pp. 9493–9500.
- [12] K. Valmeekam, A. Olmo, S. Sreedharan, and S. Kambhampati, "Large language models still can't plan (a benchmark for LLMs on planning and reasoning about change)," in *NeurIPS 2022 Foundation Models for Decision Making Workshop*, 2022. [Online]. Available: <https://openreview.net/forum?id=wUU-7XTL5XO>
- [13] Y. Ding, X. Zhang, C. Paxton, and S. Zhang, "Task and motion planning with large language models for object rearrangement," *arXiv preprint arXiv:2303.06247*, 2023.
- [14] X. Zhang, Y. Zhu, Y. Ding, Y. Zhu, P. Stone, and S. Zhang, "Visually grounded task and motion planning for mobile manipulation," in *2022 International Conference on Robotics and Automation (ICRA)*. IEEE, 2022, pp. 1925–1931.
- [15] Y. Chen, J. Arkin, Y. Zhang, N. Roy, and C. Fan, "Autotamp: Autoregressive task and motion planning with llms as translators and checkers," *arXiv preprint arXiv:2306.06531*, 2023.
- [16] S. Tellex, T. Kollar, S. Dickerson, M. Walter, A. Banerjee, S. Teller, and N. Roy, "Understanding natural language commands for robotic navigation and mobile manipulation," in *Proceedings of the AAAI conference on artificial intelligence*, vol. 25, no. 1, 2011, pp. 1507–1514.
- [17] T. M. Howard, S. Tellex, and N. Roy, "A natural language planner interface for mobile manipulators," in *2014 IEEE International Conference on Robotics and Automation (ICRA)*. IEEE, 2014, pp. 6652–6659.
- [18] H. Chang, K. Boyalakuntla, S. Lu, S. Cai, E. Jing, S. Keskar, S. Geng, A. Abbas, L. Zhou, K. Bekris, and A. Boularias, "Context-aware entity grounding with open-vocabulary 3d scene graphs," 2023. [Online]. Available: <https://arxiv.org/abs/2309.15940>
- [19] W. Huang, F. Xia, T. Xiao, H. Chan, J. Liang, P. Florence, A. Zeng, J. Tompson, I. Mordatch, Y. Chebotar *et al.*, "Inner monologue: Embodied reasoning through planning with language models," in *Conference on Robot Learning*. PMLR, 2023, pp. 1769–1782.
- [20] I. Singh, V. Blukis, A. Mousavian, A. Goyal, D. Xu, J. Tremblay, D. Fox, J. Thomason, and A. Garg, "Progprompt: Generating situated robot task plans using large language models," in *2023 IEEE International Conference on Robotics and Automation (ICRA)*. IEEE, 2023, pp. 11523–11530.
- [21] C. H. Song, J. Wu, C. Washington, B. M. Sadler, W.-L. Chao, and Y. Su, "Llm-planner: Few-shot grounded planning for embodied agents with large language models," in *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 2023, pp. 2998–3009.
- [22] D. Driess, F. Xia, M. S. Sajjadi, C. Lynch, A. Chowdhery, B. Ichter, A. Wahid, J. Tompson, Q. Vuong, T. Yu *et al.*, "Palm-e: An embodied multimodal language model," *arXiv preprint arXiv:2303.03378*, 2023.
- [23] K. Lin, C. Agia, T. Migimatsu, M. Pavone, and J. Bohg, "Text2motion: From natural language instructions to feasible plans," in *ICRA2023 Workshop on Pretraining for Robotics (PT4R)*, 2023.
- [24] J. Wu, R. Antonova, A. Kan, M. Lepert, A. Zeng, S. Song, J. Bohg, S. Rusinkiewicz, and T. Funkhouser, "Tidybot: Personalized robot assistance with large language models," *arXiv preprint arXiv:2305.05658*, 2023.
- [25] W. Huang, C. Wang, R. Zhang, Y. Li, J. Wu, and L. Fei-Fei, "Voxposer: Composable 3d value maps for robotic manipulation with language models," in *Conference on Robot Learning*. PMLR, 2023, pp. 540–562.
- [26] Z. Wu, B. Ai, and D. Hsu, "Integrating common sense and planning with large language models for room tidying," in *RSS 2023 Workshop on Learning for Task and Motion Planning*, 2023.
- [27] T. Silver, V. Hariprasad, R. S. Shuttlesworth, N. Kumar, T. Lozano-Pérez, and L. P. Kaelbling, "Pddl planning with pretrained large language models," in *NeurIPS 2022 foundation models for decision making workshop*, 2022.
- [28] B. Liu, Y. Jiang, X. Zhang, Q. Liu, S. Zhang, J. Biswas, and P. Stone, "Llm+ p: Empowering large language models with optimal planning proficiency," *arXiv preprint arXiv:2304.11477*, 2023.
- [29] Y. Xie, C. Yu, T. Zhu, J. Bai, Z. Gong, and H. Soh, "Translating natural language to planning goals with large-language models," *arXiv preprint arXiv:2302.05128*, 2023.
- [30] K. Rana, J. Haviland, S. Garg, J. Abou-Chakra, I. Reid, and N. Suen-derhauf, "Sayplan: Grounding large language models using 3d scene graphs for scalable robot task planning," in *7th Annual Conference on Robot Learning*, 2023.
- [31] L. Guan, K. Valmeekam, S. Sreedharan, and S. Kambhampati, "Leveraging pre-trained large language models to construct and utilize world models for model-based task planning," *Advances in Neural Information Processing Systems*, vol. 36, 2024.
- [32] Z. Zhao, W. S. Lee, and D. Hsu, "Large language models as common-sense knowledge for large-scale task planning," *Advances in Neural Information Processing Systems*, vol. 36, 2024.
- [33] T. Birr, C. Pohl, A. Younes, and T. Asfour, "Autogpt+ p: Affordance-based task planning with large language models," *arXiv preprint arXiv:2402.10778*, 2024.
- [34] J. Wei, X. Wang, D. Schuurmans, M. Bosma, F. Xia, E. Chi, Q. V. Le, D. Zhou *et al.*, "Chain-of-thought prompting elicits reasoning in large language models," *Advances in Neural Information Processing Systems*, vol. 35, pp. 24824–24837, 2022.
- [35] J. Wei, Y. Tay, R. Bommasani, C. Raffel, B. Zoph, S. Borgeaud, D. Yogatama, M. Bosma, D. Zhou, D. Metzler *et al.*, "Emergent abilities of large language models," *Transactions on Machine Learning Research*, 2022.
- [36] D. Sun, J. Chen, S. Mitra, and C. Fan, "Multi-agent motion planning from signal temporal logic specifications," *IEEE Robotics and Automation Letters*, vol. 7, no. 2, pp. 3451–3458, 2022.
- [37] X. Huang, Y. Zhang, J. Ma, W. Tian, R. Feng, Y. Zhang, Y. Li, Y. Guo, and L. Zhang, "Tag2text: Guiding vision-language model via image tagging," *arXiv preprint arXiv:2303.05657*, 2023.
- [38] Y. Zhang, X. Huang, J. Ma, Z. Li, Z. Luo, Y. Xie, Y. Qin, T. Luo, Y. Li, S. Liu *et al.*, "Recognize anything: A strong image tagging model," *arXiv preprint arXiv:2306.03514*, 2023.
- [39] Y. Labbé, S. Zagoruyko, I. Kalevatykh, I. Laptev, J. Carpentier, M. Aubry, and J. Sivic, "Monte-carlo tree search for efficient visually guided rearrangement planning," *IEEE Robotics and Automation Letters*, vol. 5, no. 2, pp. 3715–3722, 2020.
- [40] Y. Jiang, A. Gupta, Z. Zhang, G. Wang, Y. Dou, Y. Chen, L. Fei-Fei, A. Anandkumar, Y. Zhu, and L. Fan, "Vima: General robot manipulation with multimodal prompts," in *NeurIPS 2022 Foundation Models for Decision Making Workshop*, 2022.