

Logic Dynamic Movement Primitives for Long-horizon Manipulation Tasks in Dynamic Environments

Yan Zhang^{1,2}, Teng Xue^{1,2,*}, Amirreza Razmjoo^{1,2,*}, Sylvain Calinon^{1,2}

Abstract—Learning from Demonstration (LfD) stands as an efficient framework for imparting human-like skills to robots. Nevertheless, designing an LfD framework capable of seamlessly imitating, generalizing, and reacting to disturbances for long-horizon manipulation tasks in dynamic environments remains a challenge. To tackle this challenge, we present Logic Dynamic Movement Primitives (Logic-DMP), which combines Task and Motion Planning (TAMP) with an optimal control formulation of DMP, allowing us to incorporate motion-level via-point specifications and to handle task-level variations or disturbances in dynamic environments. We conduct a comparative analysis of our proposed approach against several baselines, evaluating its generalization ability and reactivity across three long-horizon manipulation tasks. Our experiment demonstrates the fast generalization and reactivity of Logic-DMP for handling task-level variants and disturbances in long-horizon manipulation tasks.

I. INTRODUCTION

Learning from Demonstrations (LfD) has proven to be an effective approach for enabling robots to tackle complex manipulation tasks by imitating expert demonstrations [1]. Recent advancements in LfD have extended its applicability to long-horizon manipulation tasks, such as table rearrangement or kitchen activities, by segmenting demonstrations into sub-tasks [2], keyframes [3], or skills [4]. However, existing works often focus solely on reproducing demonstrations, assuming a sequential execution of learned skills can achieve task goals, which is not always applicable in real-world scenarios [5]. While traditional LfD methods like Dynamic Movement Primitives (DMP) [6], Task-parameterized Gaussian Mixture Model (TP-GMM) [7], or Dynamical Systems (DS) [8] can generalize demonstrated trajectories for new tasks or react to disturbances at the motion level, long-horizon planning necessitates not only motion-level generalization but also the ability to handle task-level variations and disturbances. Consequently, designing an LfD framework that can imitate, generalize, and react to disturbances while solving multi-step manipulation tasks remains a challenge [2, 5].

¹Idiap Research Institute, Martigny, Switzerland; ²Ecole Polytechnique Fédérale de Lausanne (EPFL), Switzerland; Authors with * contributed equally. E-mail:firstname.lastname@idiap.ch

This work was supported by the State Secretariat for Education, Research and Innovation in Switzerland for participation in the European Commission’s Horizon Europe Program through the INTELLIMAN project (<https://intelliman-project.eu/>), HORIZON-CL4-Digital-Emerging Grant 101070136 and the SESTOSENSE project (<http://sestosenso.eu/>), HORIZON-CL4-Digital-Emerging Grant 101070310). We also acknowledge support from the China Scholarship Council (Grant No. 202106230104) and the SWITCH project (<https://switch-project.github.io/>), funded by the Swiss National Science Foundation.

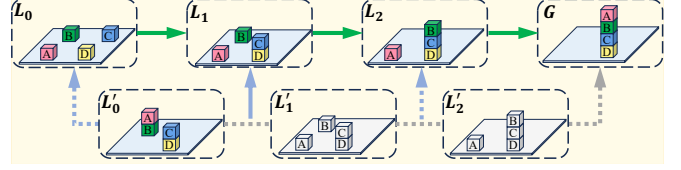


Fig. 1: Overview of Logic-DMP. Arrows refer to action primitives encoded with the proposed DMP variant (LQT-CP). The template task starts from $\mathcal{L}_0$ and ends at goal G . $\mathcal{L}_0 \rightarrow \mathcal{L}_1 \rightarrow \mathcal{L}_2 \rightarrow G$ illustrates the task-level demonstration. For a new task starting from $\mathcal{L}'_0$, a linear execution of actions primitives encoded by DMPs (green arrows) cannot transition from $\mathcal{L}'_0$ to the goal state G . Typical TAMP solvers find the action skeleton from $\mathcal{L}'_0$ to G from scratch (grey dashed long arrow). Instead, Logic-DMP tries to reach not only the goal G , but also all states on the demonstration in parallel (blue arrows) and finds a feasible plan $\mathcal{L}'_0 \rightarrow \mathcal{L}_1$ connecting $\mathcal{L}'_0$ to the demonstration trajectory within the fewest time. It then merges the new plan with the corresponding segmentation of the demonstration $\mathcal{L}_1 \rightarrow \mathcal{L}_2 \rightarrow G$, thus accomplishing the new task faster than classical TAMP solvers.

In the domain of long-horizon planning, Task and Motion Planning (TAMP) has emerged as a powerful tool for solving multi-step manipulation tasks [9, 10]. TAMP involves searching over a set of abstracted actions, with parameters often determined through sampling- or gradient-based motion planning. While theoretically capable of solving any feasible long-horizon manipulation task, TAMP relies on accurate dynamics modeling [11] and combinatorial searching of task instances and valid motions [2]. This limits the application of TAMP methods in real-world manipulation scenarios, involving model uncertainty and external disturbances.

In this work, we propose Logic-DMP, an LfD approach that generalizes one multi-step demonstration to solve new similar multi-step tasks, by combining the advantages of DMP and TAMP solvers. Logic-DMP leverages an optimal control formulation of DMP for motion modulation and extends it to incorporate via-point specifications for solving contact-rich manipulation sub-tasks, like pulling a cube with a hook in Figure. 3, while deploying TAMP solvers. Combining LfD with TAMP solvers can mitigate the challenges associated with modeling complex contact dynamics. Besides, we show that the demonstration can be formulated into multi-goal specifications for TAMP solver so that it can try to achieve all the states on the demonstrated task-and-motion trajectory in parallel, which significantly accelerates the planning process while generalizing the demonstration for solving new tasks. We further extend our approach to introduce a reactive TAMP framework by deploying Logic-DMP in a closed-loop fashion for real-world tasks in dynamic environments.

Notably, our choice of DMP as the motion modulation

block is motivated by its extrapolation generalization ability based on a single demonstration. However, conventional DMP suffers from poor via-points encoding ability, limiting its applicability for tasks with specific via-points requirements that can facilitate collision avoidance or contact-rich manipulation tasks [12, 13]. In our prior work [14], we propose Linear Quadratic Tracking with Control Primitives (LQT-CP), an optimal control formulation of DMP that can freely modulate the way in which the system reacts to perturbations. In this paper, we explore the flexible motion modulation abilities of LQT-CP and further extend it to incorporate via-points tracking and generalization capability, addressing the limitation inherent in classical DMP.

In summary, our work proposes the following technical advancements:

- Extend LQT-CP to handle long-horizon manipulation tasks and further develop its capability to incorporate via-point tracking and generalization for addressing contact-rich manipulation sub-tasks;
- Propose Logic-DMP, a novel integration of TAMP solver with DMP. This integration improves the generalization ability and reactivity of DMP in multi-step manipulation tasks;
- Develop a Reactive TAMP approach, leveraging the fast generalization capability of Logic-DMP. This approach facilitates the solution of long-horizon manipulation tasks in dynamic environments.

II. RELATED WORK

A. LfD for long-horizon manipulation tasks

Previous LfD methods typically address long-horizon manipulation tasks by segmenting demonstrations and composing movement primitives trained with the segmented demonstration [15, 16]. *Schwenkel et al.* and *Jaquier et al.* propose methods to find the best sequence of primitives based on demonstrations for the target tasks [17, 18]. Our Logic-DMP shares similarities in composing skill primitives for complex long-horizon manipulation tasks. However, unlike other methods, it does not assume a sequential execution of those primitives (namely, without generalization or replanning). We show in experiments that the removal of this assumption can improve the generalization ability and reactivity of the system to handle task variants and disturbances.

Under the umbrella of combining LfD with TAMP, *Mandlekar et al.* propose human-in-the-loop TAMP where human teleoperation assists TAMP solvers to find feasible motions for complex contact-rich manipulation tasks without complicated dynamics modeling [11]. *Mcdonald et al.* and *Dalal et al.* [19, 20] propose to train a policy imitating a TAMP planner with a large-scale database of demonstrations collected offline, achieving good performance on generalization and reactivity. In contrast, we show that our Logic-DMP can achieve better performance by relying on a single demonstration.

B. Reactive Task and Motion Planning

Our closed-loop Logic-DMP operates within the domain of reactive TAMP, where rapid replanning is essential to respond to disturbances in dynamic environments. In recent studies [5] [21], temporal logic-based reactive synthesis has been combined with Dynamical Systems (DS) or behavior tree-based control strategies for reactive action selection and plan switching to address task-level disturbances in multi-step manipulation tasks. In contrast, our approach delves into a PDDL-based (Planning Domain Definition Language) TAMP planning framework, which includes logical state and action abstraction, thus facilitates the integration of action primitives learned from demonstration.

In [22], *Migimatsu et al.* propose to execute motions of the TAMP plan in object-centric Cartesian coordinates, demonstrating efficiency in reacting to motion-level disturbances. However, their method does not extend to handling task-level disturbances, setting it apart from our approach. In [23], *Xue et al.* introduce the Dynamic Logic-Geometric Program (D-LGP), showcasing impressive time efficiency for TAMP. However, D-LGP relies on keyframe-based action skeleton planning, making it incapable of handling disturbances within the keyframes. Receding-Horizon TAMP (RH-TAMP) approaches iteratively solve a reduced planning problem over a receding window of shorter action skeleton, accelerating the TAMP planning process for handling interventions in changing environments [24, 25]. Nevertheless, a shorter prediction horizon increases the undesired infeasibility of action skeletons. Our Logic-DMP also accelerates TAMP by reducing the planning problem into a partial problem. However, it does not adopt a receding-horizon fashion, thus avoiding the influence on the feasibility.

In [26], *Harris et al.* propose a Feasibility-based Control Chain Coordination (FC³) approach, showcasing impressive reactivity by constructing offline a set of possible action plans and by switching between them in an online manner. However, the success of FC³ heavily relies on the constructed plan library and switching strategy. The closest work to ours is Robust Logical-Dynamical Systems (RLDS) [27], which modifies the initial action plan by jumping backward or forward. However, RLDS assumes task disturbances can be handled without action plan switching, significantly limiting its reactivity under substantial task perturbations. In several following experiments, our method demonstrates superior reactivity to various levels of task disturbances than RLDS, which is achieved by online generalization of the demonstrated nominal plan, without the need of constructing an action plan library offline as in [26].

III. PRELIMINARY

A. Planning Domain Definition Language (PDDL)

In this section, we present essential concepts related to the Planning Domain Definition Language (PDDL) using the block stacking task illustrated in Figure 1 as an example. PDDL serves as an interface for defining the world by specifying a set of facts that describe relationships among

objects in the environment. For instance, in $\mathcal{L}_1$ of Figure 1, the cube C is positioned on the cube D, and this relationship is represented by the fact (*on C D*). Throughout this paper, we use italicized symbols to denote these facts and the term *scene graph* to collectively represent the entire set of facts.

Among these facts, those that remain constant throughout the planning process are termed static facts (e.g., (*cube C*)). Conversely, facts that vary during the process are referred to as fluents (e.g., (*on C D*), which changes with the robot's operation).

PDDL also provides the concept of action abstraction. An action is defined by a set of parameters, preconditions that must be satisfied before executing the action, and effects that occur after the action. Taking the action *pick* as an example, its parameters may include specifying which robot(s) pick which object(s). If, for instance, we have (*robot panda*) and (*cube C*) as the parameters, one precondition for executing this action should be (*clear C*), indicating that there are no cubes on C, making it possible to grasp. The resultant effect is denoted by (*inhand panda C*), signifying that the cube C is now in the hand of the robot arm panda.

IV. METHOD

A. LQT-CP: an optimal control formulation of DMP

DMP is constructed in two parts to converge to the final goal position while tracking the acceleration profiles for mimicking the shape of a demonstrated trajectory. Similarly, Linear Quadratic Tracking (LQT), the most basic form of optimal control, is described by a cost typically composed of several parts acting at different state and control levels, with references either in the form of full trajectory or final goal positions. In particular, the cost can be specified to track a reference velocity or acceleration profile while reaching a target state at the end of the movement, with weight matrices balancing the importance of tracking the desired profiles and reaching the final goal position. With this similarity, we can reformulate DMP into a LQT fashion with the acceleration profile of the demonstrated trajectory as reference trajectory and the attractor as the goal to be reached at the end.

Similarly to DMP, our system is a virtual point mass driven by a linear integrator system to describe the evolution of the system. As shown in [14], the control profile can be encoded with basis functions, resulting in a Control Primitive (CP) formulation of LQT that estimates superposition weights instead of the full list of control commands. The resulting LQT-CP yields a DMP by minimizing the cost function:

$$c = (\boldsymbol{\mu} - \mathbf{x})^\top \mathbf{Q}(\boldsymbol{\mu} - \mathbf{x}) + \mathbf{u}^\top \mathbf{R} \mathbf{u}, \quad (1)$$

where $\boldsymbol{\mu} = [\boldsymbol{\mu}_0^\top, \boldsymbol{\mu}_1^\top, \dots, \boldsymbol{\mu}_T^\top]^\top$ indicates the concatenated vector of the position, velocity, acceleration profiles of the demonstrated trajectory and $\boldsymbol{\mu}_T$ includes the goal position $\mathbf{g}$ to be reached at the end. $\mathbf{x} = [\mathbf{x}_0^\top, \mathbf{x}_1^\top, \dots, \mathbf{x}_T^\top]^\top$ represents the concatenated system state variables. The matrices $\mathbf{Q} = \text{diag}(\mathbf{Q}_0, \mathbf{Q}_1, \dots, \mathbf{Q}_T)$ and $\mathbf{R} = \text{diag}(\mathbf{R}_0, \mathbf{R}_1, \dots, \mathbf{R}_T)$ are a block-diagonal matrices showing the evolution of weight matrices $\mathbf{Q}_t$ and $\mathbf{R}_t$ for variable $\mathbf{x}$ and command $\mathbf{u} =$

$[\mathbf{u}_0^\top, \mathbf{u}_1^\top, \dots, \mathbf{u}_T^\top]^\top$, correspondingly. The evolution of the state is described by a linear integrator $\mathbf{x}_{t+1} = \mathbf{A}_t \mathbf{x}_t + \mathbf{B}_t \mathbf{u}_t$, yielding $\mathbf{x} = \mathbf{S}_x \mathbf{x}_0 + \mathbf{S}_u \mathbf{u}$ at trajectory level, where $\mathbf{S}_x$ and $\mathbf{S}_u$ are composed of $\mathbf{A}_t$ and $\mathbf{B}_t$, see [14] for details.

Similarly as in DMP, where basis functions are used as movement primitives to represent the non-linear forcing term, basis functions can be used to represent the corresponding control commands. By leveraging a least squares formulation of recursive LQR, we showed in [14] that the control commands can be expressed as $\mathbf{u} = -\mathbf{F} \tilde{\mathbf{x}}_0 = -\boldsymbol{\Psi} \mathbf{W} \tilde{\mathbf{x}}_0$ where $\boldsymbol{\Psi}$ are the basis functions stored in matrix form, $\mathbf{W}$ is the weight matrix, and $\tilde{\mathbf{x}}_0$ is the augmented initial state for transferring the original LQT problem into an LQR problem with control primitives:

$$\begin{aligned} \min_{\mathbf{W}} \quad & \tilde{\mathbf{x}}^\top \tilde{\mathbf{Q}} \tilde{\mathbf{x}} + (\boldsymbol{\Psi} \mathbf{W} \tilde{\mathbf{x}}_0)^\top \mathbf{R} (\boldsymbol{\Psi} \mathbf{W} \tilde{\mathbf{x}}_0), \\ \text{s.t.} \quad & \tilde{\mathbf{x}} = (\tilde{\mathbf{S}}_x - \tilde{\mathbf{S}}_u \boldsymbol{\Psi} \mathbf{W}) \tilde{\mathbf{x}}_0, \end{aligned} \quad (2)$$

where $\tilde{\mathbf{x}}$ is the concatenation of augmented state $\tilde{\mathbf{x}}_t = [\mathbf{x}_t^\top, 1]^\top$, $\tilde{\mathbf{Q}}$ is the concatenation of augmented weight matrices computed as

$$\tilde{\mathbf{Q}}_t = \begin{bmatrix} \mathbf{I} & \mathbf{0} \\ -\boldsymbol{\mu}_t^\top & 1 \end{bmatrix} \begin{bmatrix} \mathbf{Q}_t & \mathbf{0} \\ \mathbf{0} & 1 \end{bmatrix} \begin{bmatrix} \mathbf{I} & -\boldsymbol{\mu}_t \\ \mathbf{0} & 1 \end{bmatrix},$$

$\tilde{\mathbf{S}}_x$ and $\tilde{\mathbf{S}}_u$ are state transfer matrices composed of $\tilde{\mathbf{A}}_t = \begin{bmatrix} \mathbf{A}_t & \mathbf{0} \\ \mathbf{0} & 1 \end{bmatrix}$ and $\tilde{\mathbf{B}}_t = \begin{bmatrix} \mathbf{B}_t \\ 0 \end{bmatrix}$ that define the augmented system dynamics $\tilde{\mathbf{x}}_{t+1} = \tilde{\mathbf{A}}_t \tilde{\mathbf{x}}_t + \tilde{\mathbf{B}}_t \mathbf{u}_t$.

Solving (2) results in the optimal weight matrix estimation:

$$\hat{\mathbf{W}} = (\boldsymbol{\Psi}^\top \tilde{\mathbf{S}}_u^\top \tilde{\mathbf{Q}} \tilde{\mathbf{S}}_u \boldsymbol{\Psi} + \boldsymbol{\Psi}^\top \mathbf{R} \boldsymbol{\Psi})^{-1} \boldsymbol{\Psi}^\top \tilde{\mathbf{S}}_u^\top \tilde{\mathbf{Q}} \tilde{\mathbf{S}}_x, \quad (3)$$

The above solution can be used as a recursive formulation to generate a feedback controller that can handle external perturbation in real-time execution, providing a feedback controller $\mathbf{u}_t = -\tilde{\mathbf{K}}_t \tilde{\mathbf{x}}_t$ where $\tilde{\mathbf{K}}_t$ is the feedback gain matrix at time step t . We can therefore recursively calculate the feedback gains of our feedback controller LQT-CP with

$$\begin{aligned} \tilde{\mathbf{K}}_t &= \boldsymbol{\Psi}_t \hat{\mathbf{W}} \mathbf{P}_t, \\ \mathbf{P}_t &= \mathbf{P}_{t-1} (\tilde{\mathbf{A}}_{t-1} - \tilde{\mathbf{B}}_{t-1} \tilde{\mathbf{K}}_{t-1})^{-1}, \end{aligned} \quad (4)$$

where $\tilde{\mathbf{K}}_0 = \boldsymbol{\Psi}_0 \mathbf{W}$, $\mathbf{P}_0 = \mathbf{I}$, see [14] for details.

LQT-CP provides additional flexibility in the design of the precision matrix $\mathbf{Q}$ to encode the importance of tracking different profiles of the demonstrated trajectory as well as capturing their correlation constraints. Therefore, we can design $\mathbf{Q}$ to track the acceleration profiles and goal position of the demonstrated trajectory to generate an optimal control formulation of DMP. In this paper, we further propose incorporating via-point specifications into the reference trajectory $\boldsymbol{\mu}$ and adjusting the weight matrix $\mathbf{Q}$ to define these specifications. This extension allows us to address contact-rich manipulation sub-tasks while concurrently extending LQT-CP to handle long-horizon manipulation tasks. We illustrate this modification in Section V-C with the example of pulling a cube using a hook.

B. Logic-DMP

In this section, we introduce Logic-DMP for Task and Motion Planning (TAMP). The primary enhancement involves augmenting the classical demonstration representation, typically consisting of geometrical state sequences (e.g., position trajectories), with logical state and action abstractions derived from PDDL, and more importantly integrating TAMP solver with the proposed LQT-CP.

Consider a long-horizon manipulation task with N sub-tasks to be solved sequentially. The demonstration, denoted as $\mathcal{P} = \{\{\mathcal{L}_i, \mathbf{a}_i\}_{i=1,\dots,N}, \{\tau_k\}_{k=1,\dots,K}\}$, encompasses both task- and motion-level trajectories. Here, $\{\mathcal{L}_i\}_{i=1,\dots,N}$ represents the logical state sequence from the initial to the goal scene graph, akin to a task-level trajectory $\mathcal{L}$. $\mathbf{a}_i$ denotes the abstracted action, defining the transition from $\mathcal{L}_i$ to $\mathcal{L}_{i+1}$, and specifies which robot arm(s) executes the motion-level trajectory for manipulating which particular object(s). In contrast to the straightforward propagation of geometrical states through Euclidean or unit quaternion operations ($\oplus, \ominus$), the logical state propagation is confined to a feasible action set $\mathcal{A} = \{a_k\}_{k=1,\dots,K}$ for the target task. Consequently, the corresponding demonstration for Logic-DMP should include not only the logical and geometrical state sequences, but also the corresponding action sequence $\{\mathbf{a}_i\}_{i=1,\dots,N}$ indicating the feasible propagation of logical states.

Given the demonstration $\mathcal{P}$, we offline train the feedback gains of a set of LQT-CPs for each action primitive in $\mathcal{A}$ with the corresponding motion-level demonstrations $\{\tau_k\}_{k=1,\dots,K}$. Replacing τ_k in $\mathcal{P}$ with the corresponding LQT-CP $_k$, by exploring the generalization ability of LQT-CP in motion level, we are able to generate an initial plan $\mathcal{P}_I$ which can handle new long-horizon manipulation tasks that only require the generalization of geometrical states. However, this limited generalization ability cannot address new tasks that require generalization in both task and motion levels.

Therefore, we further propose a new integration of TAMP solver with the proposed LQT-CP method. This integration results in a Logic-DMP that quickly generalizes the demonstration $\mathcal{P}$ to solve new long-horizon tasks that share the same task goal state G but have different initial states than the template task, thus requiring generalization in both task and motion levels.

The overview of the whole framework is depicted in Figure 1, shown with a four-block stacking scenario. Supposing the new task starts with a new logical state $\mathcal{L}'_0$, Logic-DMP efficiently solves this new task with a two-step planning process and only needs to solve a partial TAMP problem. The first step involves finding a feasible task-and-motion trajectory $\mathcal{P}^{\text{new}}$ that transitions $\mathcal{L}'_0$ to its closest logical state $\mathcal{L}_1 \in \mathcal{L}$ in the demonstration $\mathcal{P}$, then concatenating $\mathcal{P}^{\text{new}}$ with the corresponding segmentation of the initial plan $\mathcal{L}_1 \rightarrow \mathcal{L}_2 \rightarrow G$ to establish a feasible task plan transitions $\mathcal{L}'_0$ to G within the smallest time cost. Secondly, Logic-DMP generates the feasible motion-level trajectories by exploring the generalization ability of the proposed LQT-

Algorithm 1: Reactive TAMP with Logic-DMP

```

1 Given: Environment  $env$ , Demonstration  $\mathcal{P}$ , Task
   Goal  $G$ , Initial Plan  $\mathcal{P}_I$ 
2 Output: next action  $\mathbf{a}$  and motion  $\tau$ 
3 Initialization:
4  $\mathcal{G} = \text{MultiGoalsSpecification}(\mathcal{P}, G)$ 
5  $\mathcal{S}_C = \text{SceneGraph}(env)$ 
6  $\mathcal{P}_C \leftarrow \mathcal{P}_I$ 
7 while  $G \notin \mathcal{S}_C$  do
8    $flag, id = \text{LogicIn}(\mathcal{S}_C, \mathcal{G})$ 
9   if  $flag$  is True then
10     $\mathbf{a} = \mathcal{P}_I[id]$ 
11   else
12     $\mathcal{P}^{\text{new}} = \text{PDDLStream}(\mathcal{S}_C, \mathcal{G})$ 
13     $\mathbf{a} = \mathcal{P}^{\text{new}}[0]$ 
14    $\tau = \text{LQT\_CP}(\mathbf{a})$ 
15    $\mathcal{S}_C = \text{SceneGraph}(env)$ 

```

CP. The state in $\mathcal{P}^{\text{new}}$ closest to the new initial state $\mathcal{L}'_0$ is indirectly determined by reformulating the demonstrated task-level trajectory $\mathcal{L} = \{\mathcal{L}_i\}_{i=1,\dots,N}$ as multiple goals $\mathcal{G}$ and running the *PDDLStream* [9] algorithm to reach these goals concurrently, with the one reached in the shortest time considered as the target closest state. It is worth noting that our Logic-DMP framework treats the TAMP solver as a black-box tool, allowing for the integration of any TAMP solver into the framework, demonstrating its generality.

C. Reactive TAMP with Logic-DMP

Logic-DMP ensures fast discovery of a feasible task and motion plan. Since the closest state is usually reached much earlier than the goal state, Logic-DMP often resolves only a partial TAMP planning process while reacting to task-level disturbances. Considering that the planning time of a TAMP solver tends to increase significantly with the length of the final action skeleton, solving only a partial TAMP problem allows Logic-DMP to perform notably faster than the original TAMP solver. Therefore, we further extend the Logic-DMP into a close-loop TAMP framework (Algorithm 1) which promptly reacts to task-and-motion-level disturbances while executing the nominal plan in dynamic environments.

It assumes the task goal G , demonstration $\mathcal{P}$, and an initial plan $\mathcal{P}_I$ are given. In the initialization process, multi-goal specifications $\mathcal{G}$ are generated based on the demonstration and the task goal with designed *MultiGoalsSpecification* function in Line 4. This functions filters out the static facts in $\mathcal{L}$ and keeps the sequence of fluents in $\mathcal{P}$. The designed *SceneGraph* function generates the facts about current environment based on captured information about objects and robots in the environment and outputs the filtered key facts.

To solve the target task in closed-loop, we continuously estimate the key logical states $\mathcal{S}_C$ after executing each action $\mathbf{a}$ in the current plan $\mathcal{P}_C$. If the target state G is not a subset of $\mathcal{S}_C$, indicating the task is not achieved yet. We

then invoke the Logic-DMP to swiftly generate a new plan. With the *LogicIn* function in Line 8, we check if current state $\mathcal{S}_C$ is a subset of the multi-goal specifications $\mathcal{G}$. If it is, we also generate its corresponding index id in the sequence. Then, the action $\mathcal{P}_I[id]$ in the initial plan will be a feasible action that transitions current state $\mathcal{S}_C$ to the goal state $\mathcal{G}$. If $\mathcal{S}_C$ is not a subset of $\mathcal{G}$, it indicates a severe task-level disturbance in the environment, and *PDDLStream* with multi-goal specification $\mathcal{G}$ is applied to find a new feasible plan $\mathcal{P}^{\text{new}}$ that converges to the demonstration $\mathcal{P}$. The first action in $\mathcal{P}^{\text{new}}$ is executed as the feasible action for the current state. The action $\mathbf{a}$ not only shows which action should be executed, but also indicates which robot arm(s) should execute that action for manipulating which object(s). The LQT-CP is then applied in an object-centric manner to generalize the reference trajectory for generating the corresponding motion-level trajectory τ^{new} to execute action $\mathbf{a}$.

V. EXPERIMENTS

In this section, we present a comparative analysis between Logic-DMP and several baselines across three long-horizon manipulation planning problems. Figure 2 shows the experiment setups and the demonstrated tasks for the benchmarks. For each benchmark, we demonstrated the robot with an initial task and motion plan based on available action sets. Subsequently, we introduced distinct initial setups or disturbances to assess the performance of Logic-DMP against the selected baselines.

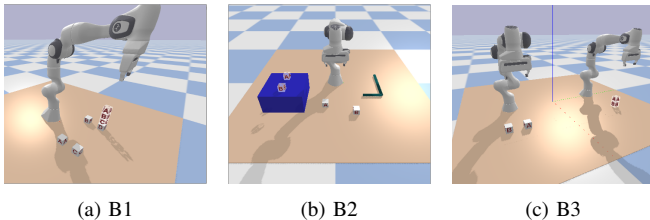


Fig. 2: Experimental setups for the three benchmarks. *B1*: Tower Construction, *B2*: Workspace Reach, *B3*: Dual-arm Box Transportation. Each sub-figure illustrates the demonstrated task, with the initial and task goal state depicted in grey and white color, respectively.

A. Benchmarks

1) *Tower Construction (B1)*: This experiment involves controlling a robotic arm to manipulate a set of blocks using actions from the set $\mathcal{A} = \{\text{pick}, \text{place}, \text{stack}, \text{unstack}\}$ to arrange them in a specific order. Using the example of constructing a four-block tower, the template task begins with an initial state where all four blocks are placed on the table (i.e., without stacked blocks). The task goal state is achieved when the blocks are stacked in the sequence *A on B on C on D*. During the interaction between the robot and blocks, both task- and motion-level constraints should be considered. For instance, the action *pick* for block *A* is only permissible if it is *clear* (no blocks on it) and if a *collision-free* trajectory exists.

2) *Workspace Reach (B2)*: As an extension of the Tower Construction benchmark, this scenario involves the robot arm manipulating blocks with tools to address more complex long-horizon manipulation tasks. Blocks may be positioned beyond the reach of the robot arm. The template task begins with an initial state where a few blocks and one hook are placed on the table within the robot arm’s workspace. It ends with the task goal state, where all blocks are placed on the shelf on the table. Actions *stack* and *unstack* are excluded to simplify the definition of the *pull* action, which describes instances where the hook is grasped by the robot arm to extend its workspace and pull the block within its original workspace for picking or placing.

3) *Dual-arm Block Transportation (B3)*: This benchmark further extends the Tower Construction experiment to validate the proposed algorithm in more complex multi-arm scenarios. The feasible action set $\mathcal{A} = \{\text{pick}, \text{place}, \text{stack}, \text{unstack}\}$ remains the same, but with augmented parameters introducing two *arms*. This complexity leads to more instantiations of abstracted actions and results in a more intricate long-horizon manipulation planning task. The template task begins with an initial state where all blocks are on the left table without any blocks on top and concludes with the task goal state, where blocks are expected to be stacked in a given sequence as in B1.

B. Disturbances

To comprehensively evaluate the generalization capability and reactivity of Logic-DMP across all benchmarks, we apply variants on initial states or disturbances at any states at different levels.

- L1 *Motion-level Variant/Disturbance*: Blocks are subjected to disturbances in different positions, while logical states remain consistent with the original or expected conditions;
- L2 *Slight Task-level Variant/Disturbance*: Blocks undergo disturbances, resulting in a different logical state. This logical state aligns with that seen in the demonstration for the template task;
- L3 *Severe Task-level Variant/Disturbance*: Blocks are disturbed to a novel and previously unseen logical state in the demonstration;
- L4 *Extreme Task-level Variant/Disturbance*: A new block is introduced into the scene that significantly influences the execution of the initial plan to achieve the target goal.

C. Comparison between LQT-CP and DMP

In this section, we demonstrate the flexible motion modulation capabilities of LQT-CP, addressing two crucial sub-tasks in B2. The experimental results are depicted in Figure 3. In the initial picking sub-task, the robot arm is tasked with picking up a hook potentially placed at three distinct positions. Both LQT-CP and DMP exhibit comparable performance, effectively generalizing the reference trajectory to new goal positions, as showcased in Figure 3-Top.

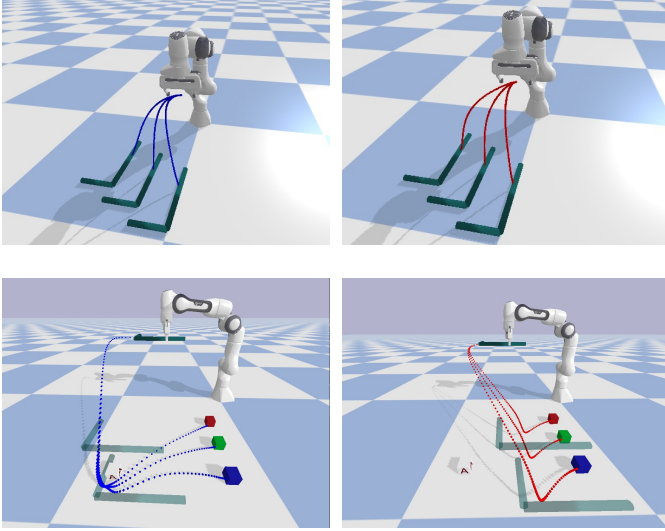


Fig. 3: Comparison between LQT-CP and standard DMP for two sub-tasks in the Workspace Reach benchmark (B2). *Top*: pick the hook. *Bottom*: pull cube A with the hook. The trajectories generated from LQT-CP are illustrated with blue lines in the left figures, the ones for DMP are shown with red lines in the right. In the pulling task, the hook is expected to pass through two crucial via-points, the top and left-down corner of the cube A, for successfully accomplishing the task. In this figure, only the via-points for pulling cube A to the blue goal are shown with transparent hooks in the *Bottom* figures.

In the subsequent pulling sub-task, the robot arm is required to pull cube A using the hook to three different goal positions represented by red, green, and blue cubes in Figure 3-*Bottom*. This task demands not only the generalization of the reference pulling trajectory to new goal positions but also the precise control of the hook passing through two crucial via-points (the top and the left-down corner of cube A) for successful pulling. Consequently, it necessitates more flexible motion modulation concerning via-points tracking and generalization. Figure 3-*Bottom* illustrates the actual position of the hook when expected to be at the via-points. Notably, LQT-CP adeptly generalizes the reference trajectory and via-points, successfully pulling cube A to the goal positions, while standard DMP encounters challenges, which highlights the flexibility of the proposed LQT-CP formulation of DMP for motion modulation.

D. Generalization Ability of Logic-DMP

In this section, we assess the generalization ability of Logic-DMP in comparison to two baseline methods: *Linear* and *PDDLStream*. Here, *Linear* refers to the linear execution of a fixed sequence of DMPs without adapting the action sequence when addressing new tasks. For a comprehensive evaluation, each benchmark comprises tasks that share the same task goal but vary in their initial states from the template task. For B1, we solved a four-block tower construction problem, while tasks are simplified by utilizing only two blocks in B2 and B3. Moreover, LQT-CP serves as the motion modulation method for all baselines to ensure a fair comparison. We analyzed the overall success rate and

computation time of the three methods in solving 100 tasks with random initial states for each benchmark across various levels (L1-L4). The summarized experimental results are presented in Table I.

In summary, both Logic-DMP and PDDLStream successfully solved all the validation tasks, while the linear execution of DMPs achieves success in only approximately 25% of the tasks, primarily those featuring motion-level variants in the initial states. The results show that Logic-DMP can enhance the generalization ability beyond linear execution. Moreover, Logic-DMP exhibits faster planning capabilities than PDDLStream across all benchmarks, with a 30% to 40% improvement in B2 and B3, and a remarkable 70% improvement in B1. The substantial acceleration observed in B1 can be attributed to the inherently longer action sequence in the task, with which the planning time of N/A increases significantly. Since Logic-DMP only needs to partially address the TAMP problem with a shorter action skeleton, it notably decreases the whole planning time, showing significant acceleration in the experimental results. This reinforces the efficacy of Logic-DMP in handling long-horizon manipulation tasks and motivates us to formulate the reactive TAMP framework in Algorithm 1 by employing Logic-DMP in a closed-loop manner.

E. Reactivity of TAMP with Logic-DMP

In this section, we conduct a comparative analysis of the reactivity exhibited by the proposed closed-loop Logic-DMP framework and two baseline approaches: *Linear* and *RLDS*. The *Linear* baseline involves the linear execution of a predetermined sequence, without closed-loop feedbacks. *RLDS* is a reproduction of the algorithm outlined in [27], with the incorporation of the proposed LQT-CP for motion modulation. The experimental setup aligns with that detailed in Section V-D, ensuring consistency across benchmarks. We assess the performance of each method in handling disturbances at different levels (L1-L4) and report the average execution times based on 10 trials for each disturbed task in Table II.

Our analysis reveals that the proposed closed-loop Logic-DMP and RLDS demonstrate comparable reactivity in the presence of motion-level and slight task-level disturbances. Specifically, under L1 and L2 disturbances, both closed-loop Logic-DMP and RLDS effectively respond to the disturbances, and achieves task goals within comparable execution times. However, RLDS relies on backward checking of the planned action sequence without plan modification, leading to challenges in handling disturbances at L3 and L4 across all benchmarks. In contrast, Logic-DMP exhibits the ability to generate a new feasible action plan when faced with those disturbances, enabling the TAMP approach to accomplish tasks even under severe or extreme disturbances. Notably, Logic-DMP proves to be faster in deriving the new action plan compared to the straightforward application of TAMP solvers, as demonstrated in the comparison results between PDDLStream and Logic-DMP in Section V-D, showing the

	Linear / Sequential			PDDLStream			Logic-DMP		
	B1	B2	B3	B1	B2	B3	B1	B2	B3
Success Rate	27%	31%	24%	100%	100%	100%	100%	100%	100%
Time [s]	N/A	N/A	N/A	0.53±0.10	0.53±0.23	0.23±0.09	0.16±0.07	0.37±0.27	0.14±0.06

TABLE I: Generalization ability comparison of Logic-DMP with baselines when solving tasks with various initial states in benchmarks. ‘N/A’ means no feasible values for the corresponding elements. Logic-DMP demonstrates superior generalization ability, achieving a higher success rate than the *Linear* execution of DMPs, and incurs significantly smaller computation time compared to PDDLStream.

	Linear / Sequential			RLDS			Logic-DMP		
	B1	B2	B3	B1	B2	B3	B1	B2	B3
L1	13.23±0.58	12.50±0.13	24.63±0.05	13.70±1.23	12.59±0.19	24.64±0.04	13.23±0.35	12.50±0.13	24.62±0.03
L2	N/A	N/A	N/A	18.07±0.95	19.08±0.01	30.83±0.04	17.62±0.67	19.08±0.01	30.79±0.02
L3	N/A	N/A	N/A	N/A	N/A	N/A	23.05±0.49	26.03±0.78	30.23±0.08
L4	N/A	N/A	N/A	N/A	N/A	N/A	19.00±1.29	12.60±0.15	31.57±0.04

TABLE II: Comparison of average execution time [s] between Logic-DMP and baselines in three benchmarks. ‘N/A’ denotes no successful trials. Closed-loop Logic-DMP shows superior reactivity to various disturbances in all benchmarks. Please refer to the supplementary video through the project webpage for the illustration of simulation experiments for each element in the table.

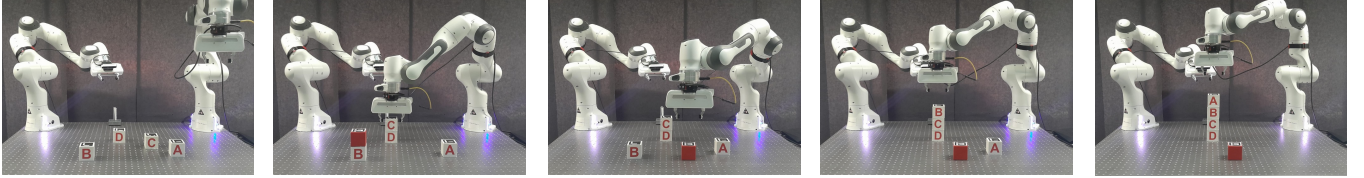


Fig. 4: Closed-loop Logic-DMP under extreme task-level disturbance (placing the red block on cube B after stacking cube C on D) in a real-world four-block stacking task. Logic-DMP reacts to the disturbance by unstacking the red block from cube B and place it on the table, then continuing the original plan for achieving the goal state, as shown in the last figure.

significant reactivity inherent to the proposed closed-loop Logic-DMP approach.

F. Real-world Experiments

We assessed the reactivity of the proposed Logic-DMP in closed-loop scenarios for handling various task-level disturbances within the four-block Tower Construction scenario, as shown in Figure 4. All experiments were conducted using a 7-axis Franka Emika robot arm equipped with a RealSense D435 camera for object localization. The entire planning algorithm demonstrates the capability to respond to severe or extreme task-level disturbances at the frequency of 1Hz, providing corresponding actions and generalized motion sequences for rearranging the objects into the desired goal configuration. Additional demonstrations of the reactive behaviors under various task-level disturbances can be found in the supplementary video of the project webpage.

VI. CONCLUSION

In summary, we introduced Logic-DMP, an Learning from Demonstration (LfD) approach tailored for long-horizon manipulation tasks in dynamic environments. We leveraged an optimal control formulation of Dynamical Movement Primitive (DMP), Linear Quadratic Tracking with Control Primitives (LQT-CP), which naturally extends DMP to incorporate via-point specifications, proving beneficial for handling contact-rich manipulation sub-tasks. We demonstrated that Logic-DMP shows superior generalization ability than DMP and superior efficiency than TAMP solver in handling task variants during long-horizon planning. Consequently, we further extended this framework into a reactive TAMP system to quickly react to disturbances in dynamic environments. We validated the proposed methods through various simulation and real-world experiments, affirming the performances in

term of skill transfer, generalization ability, and reactivity to disturbances in long-horizon manipulation tasks.

In this paper, we integrated DMP with TAMP solver because of its exceptional extrapolation generalization ability with one single demonstration. Notably, this integration is not confined to DMP alone. It can be extended to other LfD algorithms. An interesting avenue involves further extending Logic-DMP by integrating with diffusion models [28] to tackle more complex long-horizon contact-rich manipulation tasks. Moreover, we aim to extend the proposed method to other real-world scenarios characterized by partially observable environments, which necessitate rapid replanning based on new observations.

One drawback of Logic-DMP is its potential to find slightly longer solutions when generalizing to new tasks compared to directly applying TAMP solver. This is attributed to the consideration of goals in multi-goal specifications with the same priority. When multiple feasible solutions exist, Logic-DMP selects the one reached in the fastest way, but it doesn’t guarantee that the chosen goal is the closest to the task goal state among all feasible solutions. Future work should explore strategies for assigning priorities to goals within the multi-goal specifications and the integration with optimization-based TAMP solvers [10, 23] to solve long-horizon manipulation tasks where the optimality of metrics (e.g. time and energy efficiency) are important.

REFERENCES

- [1] A. G. Billard, S. Calinon, and R. Dillmann, “Learning from humans,” in *Handbook of Robotics*, B. Siciliano and O. Khatib, Eds. Secaucus, NJ, USA: Springer, 2016, ch. 74, pp. 1995–2014, 2nd Edition.
- [2] A. Mandelkar, D. Xu, R. Martín-Martín, S. Savarese, and L. Fei-Fei, “Learning to generalize across long-horizon tasks from human demonstrations,” *arXiv:2003.06085*, 2020.

- [3] C. Pérez-D'Arpino and J. A. Shah, "C-Learn: Learning geometric constraints from demonstrations for multi-step manipulation in shared autonomy," in *IEEE International Conference on Robotics and Automation (ICRA)*, 2017, pp. 4058–4065.
- [4] G. Konidaris, S. Kuindersma, R. Grupen, and A. Barto, "Robot learning from demonstration by constructing skill trees," *The International Journal of Robotics Research*, vol. 31, no. 3, pp. 360–375, 2012.
- [5] Y. Wang, N. Figueroa, S. Li, A. Shah, and J. Shah, "Temporal Logic Imitation: Learning plan-satisficing motion policies from demonstrations," *arXiv:2206.04632*, 2022.
- [6] A. J. Ijspeert, J. Nakanishi, H. Hoffmann, P. Pastor, and S. Schaal, "Dynamical movement primitives: learning attractor models for motor behaviors," *Neural computation*, vol. 25, no. 2, pp. 328–373, 2013.
- [7] S. Calinon, "A tutorial on task-parameterized movement learning and retrieval," *Intelligent service robotics*, vol. 9, pp. 1–29, 2016.
- [8] A. Billard, S. Mirrazavi, and N. Figueroa, *Learning for Adaptive and Reactive Robot Control: A Dynamical Systems Approach*. Mit Press, 2022.
- [9] C. R. Garrett, T. Lozano-Pérez, and L. P. Kaelbling, "PDDLStream: Integrating symbolic planners and blackbox samplers via optimistic adaptive planning," in *Proceedings of the International Conference on Automated Planning and Scheduling*, vol. 30, 2020, pp. 440–448.
- [10] M. Toussaint, "Logic-Geometric Programming: An optimization-based approach to combined task and motion planning," in *International Joint Conference on Artificial Intelligence*, 2015, pp. 1930–1936.
- [11] A. Mandlekar, C. R. Garrett, D. Xu, and D. Fox, "Human-in-the-loop task and motion planning for imitation learning," in *Conference on Robot Learning*. PMLR, 2023, pp. 3030–3060.
- [12] Y. Zhou, J. Gao, and T. Asfour, "Learning via-point movement primitives with inter-and extrapolation capabilities," in *IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2019, pp. 4301–4308.
- [13] S. Mghames, M. Hanheide, and A. Ghalamzan, "Interactive movement primitives: Planning to push occluding pieces for fruit picking," in *IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2020, pp. 2616–2623.
- [14] S. Calinon, "Programming industrial robots from few demonstrations," in *Human-Robot Collaboration: Unlocking the potential for industrial applications*. Institution of Engineering and Technology (IET), 2023, pp. 9–37.
- [15] M. Saveriano, F. Franzel, and D. Lee, "Merging position and orientation motion primitives," in *IEEE International Conference on Robotics and Automation (ICRA)*, 2019, pp. 7041–7047.
- [16] S. Niekum, S. Chitta, A. G. Barto, B. Marthi, and S. Osentoski, "Incremental semantically grounded learning from demonstration," in *Robotics: Science and Systems*, vol. 9. Berlin, Germany, 2013, pp. 10–15 607.
- [17] L. Schwenkel and M. Guo, "Optimizing sequences of probabilistic manipulation skills learned from demonstration."
- [18] N. Jaquier, Y. Zhou, J. Starke, and T. Asfour, "Learning to sequence and blend robot skills via differentiable optimization," *IEEE Robotics and Automation Letters*, vol. 7, no. 3, pp. 8431–8438, 2022.
- [19] M. J. McDonald and D. Hadfield-Menell, "Guided imitation of task and motion planning," in *Conference on Robot Learning*. PMLR, 2022, pp. 630–640.
- [20] M. Dalal, A. Mandlekar, C. Garrett, A. Handa, R. Salakhutdinov, and D. Fox, "Imitating task and motion planning with visuomotor transformers," *arXiv:2305.16309*, 2023.
- [21] S. Li, D. Park, Y. Sung, J. A. Shah, and N. Roy, "Reactive task and motion planning under temporal logic specifications," in *2021 IEEE International Conference on Robotics and Automation (ICRA)*, 2021, p. 12618–12624.
- [22] T. Migimatsu and J. Bohg, "Object-centric task and motion planning in dynamic environments," *IEEE Robotics and Automation Letters*, vol. 5, no. 2, p. 844–851, 2020.
- [23] T. Xue, A. Razmjoo, and S. Calinon, "D-LGP: Dynamic logic-geometric program for reactive task and motion planning," *arXiv preprint arXiv:2312.02731*, 2023.
- [24] N. Castaman, E. Pagello, E. Menegatti, and A. Pretto, "Receding horizon task and motion planning in changing environments," *Robotics and Autonomous Systems*, vol. 145, p. 103863, 2021.
- [25] C. V. Braun, J. Ortiz-Haro, M. Toussaint, and O. S. Oguz, "RHH-LGP: Receding horizon and heuristics-based logic-geometric programming for task and motion planning," *arXiv:2110.03420*, 2022.
- [26] J. Harris, D. Driess, and M. Toussaint, "FC³: Feasibility-based control chain coordination," *arXiv:2205.04362*, 2022.
- [27] C. Paxton, N. Ratliff, C. Eppner, and D. Fox, "Representing robot task plans as robust logical-dynamical systems," in *2019 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2019, p. 5588–5595.
- [28] C. Chi, S. Feng, Y. Du, Z. Xu, E. Cousineau, B. Burchfiel, and S. Song, "Diffusion policy: Visuomotor policy learning via action diffusion," *arXiv preprint arXiv:2303.04137*, 2023.