

ManiSkill: Generalizable Manipulation Skill Benchmark with Large-Scale Demonstrations

Tongzhou Mu*, Zhan Ling*, Fanbo Xiang*, Derek Yang*, Xuanlin Li*,
Stone Tao, Zhiao Huang, Zhiwei Jia, Hao Su
University of California San Diego

{t3mu, z6ling, fxiang, d7yang, xul012, stao, z2huang, zjia, haosu}@ucsd.edu

<https://github.com/haosulab/ManiSkill>

Abstract

Object manipulation from 3D visual inputs poses many challenges on building generalizable perception and policy models. However, 3D assets in existing benchmarks mostly lack the diversity of 3D shapes that align with real-world intra-class complexity in topology and geometry. Here we propose SAPIEN **Manipulation Skill Benchmark** (ManiSkill) to benchmark manipulation skills over diverse objects in a full-physics simulator. 3D assets in ManiSkill include large intra-class topological and geometric variations. Tasks are carefully chosen to cover distinct types of manipulation challenges. Latest progress in 3D vision also makes us believe that we should customize the benchmark so that the challenge is inviting to researchers working on 3D deep learning. To this end, we simulate a moving panoramic camera that returns ego-centric point clouds or RGB-D images. In addition, we would like ManiSkill to serve a broad set of researchers interested in manipulation research. Besides supporting the learning of policies from interactions, we also support learning-from-demonstrations (LfD) methods, by providing a large number of high-quality demonstrations (~36,000 successful trajectories, ~1.5M point cloud/RGB-D frames in total). We provide baselines using 3D deep learning and LfD algorithms. All code of our benchmark (simulator, environment, SDK, and baselines) is open-sourced, and a challenge facing interdisciplinary researchers will be held based on the benchmark.



Figure 1: A subset of environments in ManiSkill. We currently support 4 different manipulation tasks: OpenCabinetDoor, OpenCabinetDrawer, PushChair, and MoveBucket; each features a large variety of 3D articulated objects to encourage generalizable physical manipulation skill learning.

*These authors contributed equally to this work

1 Introduction and Related Works

To automate repetitive works and daily chores, robots need to possess human-like manipulation skills. A remarkable feature of human manipulation skill is that, once we have learned to manipulate a category of objects, we will be able to manipulate even unseen objects of the same category, despite the large topological and geometric variations. Taking swivel chairs as an example, regardless of the existence of armrest or headrest, the number of wheels, or the shape of backrest, we are confident of using them immediately. We refer to such ability to interact with unseen objects within a certain category as **generalizable manipulation skills**.

Generalizable manipulation skill learning is at the nexus of vision, learning, and robotics, and poses many interesting research problems. Recently, this field has started to attract much attention across disciplines. For example, reinforcement learning and imitation learning are applied to object grasping and manipulation [1, 2, 3, 4, 5, 6, 7, 8, 9]. On the other hand, [10, 11, 12, 13, 14, 15, 16, 17] can propose novel grasp poses on novel objects based on visual inputs. To further foster synergistic efforts, it is crucial to build a benchmark that backs reproducible research and allows researchers to compare and thoroughly examine different algorithms.

However, building such a benchmark is extremely challenging. To motivate our benchmark proposal, we first analyze key factors that complicate the design of generalizable manipulation skill benchmarks and explain why existing benchmarks are still insufficient. With the motivations in mind, we then introduce design features of our SAPIEN Manipulation Skills Benchmark (abbreviated as **ManiSkill**).

Key factors that affect our benchmark design. To guide users and create concentration on algorithm design, four key factors must be considered: 1) manipulation policy structure, 2) diversity of objects and tasks, 3) targeted perception algorithms, and 4) targeted policy algorithms.

1) Manipulation policy structure: Manipulation policies have complex structures that require different levels of simulation support, and we focus on full-physics simulation. Since simulating low-level physics is difficult, many robot simulators only support abstract action space (i.e., manipulation skills already assumed) [18, 19, 20, 21, 22, 23, 24, 25, 19, 24]. It is convenient to study high-level planning in these benchmarks; however, it becomes impossible to study more challenging scenarios with high-dimensional and complex low-level physics. Some recent benchmarks [26, 27, 28, 29, 30, 31] start to leverage the latest full-physics simulators [32, 33, 34, 35] to support physical manipulation. Despite the quantity of existing environments, most of them lack the ability to benchmark object-level generalizability within categories, and lack inclusion for different methodologies in the community, while we excel in these dimensions, which is explained next.

2) Diversity of objects and tasks: To test object-level generalizability, the benchmark must possess enough intra-class variation of object topology, geometry, and appearance, and we provide such variation. Several benchmarks or environments, including robosuite [28], RLBench [31], and Meta-World [30], feature a wide range of tasks; however, they possess a common problem: lacking object-level variations. Among past works, DoorGym [29] is equipped with the best object-level variations: it is a door opening benchmark with doors procedurally generated from different knob shapes, board sizes, and physical parameters, but it still does not capture some simple real-world variations, such as multiple doors with multiple sizes on cabinets with different shapes. This is in part due to the limitations of procedural modeling. Even though procedural modeling has been used in 3D deep learning [36, 37], it often fails to cover objects with real-world complexity, where crowd-sourced data from Internet users and real-world scans are often preferred (which is our case). Finally, a single type of task like opening doors cannot cover various motion types. For example, pushing swivel chairs requires very different skills from opening doors since it involves controlling under-actuated systems through dual-arm collaboration. Therefore, it is essential to build benchmarks with *both* great asset variations and wide skill coverage.

3) Targeted perception algorithms: Benchmarks need to decide the type and format of sensor data, and we focus on 3D sensor data mounted on robots. Many existing benchmarks, such as *DoorGym*, rely on fixed cameras to capture 2D images; however, this setting greatly limits the tasks a robot can solve. Instead, robot-mounted cameras are common in the real world to allow much higher flexibility, such as Kinova MOVO [38], and autonomous driving in general; those cameras are usually designed to capture 3D inputs, especially point clouds. Moreover, tremendous progress has been achieved to build neural networks with 3D input [39, 40, 41, 42, 43, 44, 45], and these 3D networks have demonstrated strong performance (e.g., they give better performance than 2D image networks



Figure 2: ManiSkill features diverse articulated objects with complex topological and geometric variations, such as different numbers and shapes of doors and/or drawers on different shapes of cabinets. We invested significant effort to process objects from the PartNet-Mobility Dataset and integrate into our tasks, such as adjusting the size and physical parameters (e.g. friction) so that environments are solvable, along with manual convex decomposition.

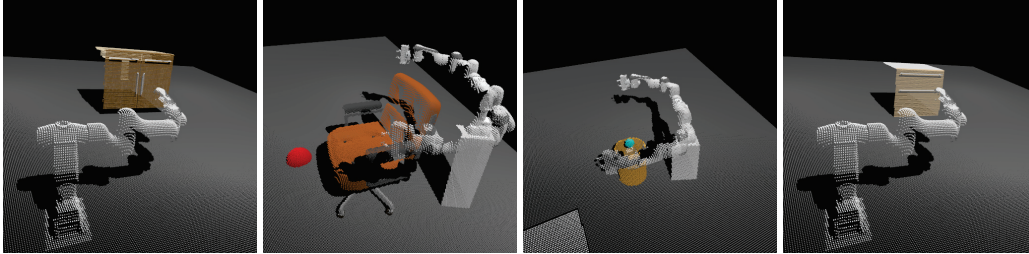


Figure 3: Rendered point clouds from our tasks. ManiSkill supports 3D visual inputs which are widely accessible in real environments, allowing various computer vision models to be applied. (For better view, we show point clouds obtained from cameras mounted in the world frame. In actual tasks, cameras are mounted on the robot head, offering an egocentric view.)

on autonomous driving datasets [46]). [47, 48, 49, 50, 51, 52] have also adopted 3D deep learning models for perceiving and identifying kinematic structures and object poses for articulated object manipulation. Our benchmark provides users with an ego-centric panoramic camera to capture point cloud / RGB-D inputs. Additionally, we present and evaluate 3D neural network-based policy learning baselines.

4) Targeted policy algorithms: Different policy learning algorithms require different training data and settings, and we provide multiple tracks to advocate for fair comparison. For example, imitation learning [53, 54, 4] and offline RL [55, 56, 57] can learn a policy purely from demonstrations datasets [58, 59], but online RL algorithms [60, 61] require interactions with environments. Therefore, a clear and meaningful split of tracks can encourage researchers with different backgrounds to explore generalizable manipulation skills and let them focus on different aspects of the challenge, e.g. network design, perception, interaction, planning, and control. While oftentimes other benchmarks are limited to a single domain of research and a single modality, our benchmark supports three different tracks for researchers from computer vision, reinforcement learning, and robotics fields.

Our benchmark. Above we discussed factors of manipulation benchmark design and mentioned the principles behind ManiSkill. Here we introduce the key features of the benchmark. ManiSkill is a large-scale open-source benchmark for physical manipulation skill learning over a diverse set of articulated objects from 3D visual inputs. ManiSkill has four main features: **First**, to support generalizable policy learning, ManiSkill provides objects of high topology and geometry variations, as shown in Fig 2. It currently includes a total of 162 objects from 3 object categories (more objects are being added) selected and manually processed from a widely used 3D vision dataset. **Second**, ManiSkill focuses on 4 object-centric manipulation tasks that exemplify household manipulation skills with different types of object motions, thereby posing challenges to distinct aspects of policy design/learning (illustrated in Fig 1 and Fig 3). As an ongoing effort, we are designing more. **Third**, to facilitate learning-from-demonstration methods, we have collected a large number of successful trajectories (~36,000 trajectories, ~1.5M 3D point cloud / RGB-D frames in total). **Fourth**, the environments feature high-quality data for physical manipulation. We take significant efforts to select, fix, and re-model the original PartNet-Mobility data [62, 63, 64], as well as design the reward generation rules, so that the manipulation task of each object can be solved by an RL algorithm.

A major challenge to build our benchmark is to collect demonstrations. Some tasks are tricky (e.g., swivel chair pushing requires dual-arm coordination), and it is difficult and unscalable to manually control the robots to collect large-scale demos. It is also unclear whether traditional motion planning pipelines can solve all tasks. Thankfully, reinforcement learning does work for individual objects, allowing for a divide-and-conquer approach to create high-quality demonstrations. With a meticulous effort on designing a shared reward template to automatically generate reward functions for all object instances of each task and executing an RL agent for each object instance, we are able to collect a large number of successful trajectories. This RL plus divide-and-conquer approach is very scalable with respect to the number of object instances within a task, and we leave cross-task RL reward design for future work. It is worth noting that, *we do NOT target at providing a **GENERIC learning-from-demonstrations benchmark** that compares methods from all dimensions*. Instead, we compare the ability of different algorithms to utilize our demonstrations to solve manipulation tasks.

Another important feature of ManiSkill is that it is completely free and built on an entirely open-source stack. Other common **physical** manipulation environments, including robosuite[28], DoorGym[29], MetaWorld[30], and RL Bench[31], depend on commercial software.

To summarize, here are the key contributions of ManiSkill Benchmark.

- The topology and geometry variation of our data allow our benchmark to compare object-level generalizability of different physical manipulation algorithms. Our data is high-quality, that every object is verified to support RL.
- The manipulation tasks we design target at distinct challenges of manipulation skills (by motion types, e.g. revolute and prismatic joint constraints, or by skill properties, e.g. requirement of dual-arm collaboration).
- ManiSkill provides a large-scale demonstrations dataset with ~36,000 trajectories and ~1.5M point cloud/RGB-D frames to facilitate learning-from-demonstrations approaches. The demonstrations are collected by a scalable RL approach with dense rewards generated by a shared reward template within each task.
- We provide several 3D deep learning-based policy baselines.

2 ManiSkill Benchmark

The goal of building ManiSkill benchmark can be best described as facilitating *learning generalizable manipulation skills from 3D visual inputs with demonstrations*. “*Manipulation*” involves low-level physical interactions and dynamics simulation between robot agents and objects; “*skills*” refer to short-horizon physics-rich manipulation tasks, which can be viewed as basic building blocks of more complicated policies; “*3D visual inputs*” are egocentric point cloud and / or RGB-D observations captured by a panoramic camera mounted on a robot; “*demonstrations*” are trajectories that solve tasks successfully to facilitate learning-from-demonstrations approaches.

In this section, we will describe the components of ManiSkill benchmark in detail, including basic terminologies and setup, task design, demonstration trajectory collection, training-evaluation protocol, and asset postprocessing with verification.

Task	# objects			Dual-arm Collaboration?	Solvable by Motion Planning?	DoF*
	All	Train	Test			
OpenCabinetDoor	52(82)	42(66)	10(16)	No	Easy	1
OpenCabinetDrawer	35(70)	25(49)	10(21)	No	Easy	1
PushChair	36	26	10	Yes	Hard	~15-25
MoveBucket	39	29	10	Yes	Medium	7

Table 1: Dataset statistics for ManiSkill. For OpenCabinetDoor and OpenCabinetDrawer, numbers outside of the parenthesis indicate the number of unique cabinets, where each cabinet may have more than one door/drawer. Numbers in the parenthesis indicate the total number of doors/drawers. * The DoF in the table indicates the DoF involved in solving a task. For OpenCabinetDoor and OpenCabinetDrawer, an agent only needs to open one designated door/drawer. For PushChair and MoveBucket, 6 extra DoF are included since chairs and buckets can move freely in 3D space.

2.1 Basic Terminologies and Setup

In ManiSkill, a **task** or a **skill** $\mathcal{T} = \{T_{o,l} : o \in \mathcal{O}, l \in \mathcal{L}_o\}$ consists of finite-horizon POMDPs (Partially Observable Markov Decision Processes) defined over a set of objects $\mathcal{O}$ of the same category (e.g., chairs) and a set of environment parameters $\mathcal{L}_o$ associated with an object $o \in \mathcal{O}$ (e.g. friction coefficients of joints on a chair). An **environment** is a set of POMDPs $\mathcal{E}_o = \{T_{o,l} : l \in \mathcal{L}_o\}$ defined over a single object o and its corresponding parameters. Each $T_{o,l}$ is a specific instance of an environment, represented by a tuple of sets (S, A, P, R, O) . Here, $s \in S$ is an environment state that consists of *robot states* (e.g. joint angles of the robot) and *object states* (e.g. object pose and the joint angles); $a \in A$ is an action that can be applied to a robot (e.g. target joint velocity of a velocity controller); $P(s'|s, a)$ is the physical dynamics; R is a binary variable that indicates if the task is successfully solved; $O(o|s)$ is a function which generates observations from an environment state, and it supports three modes in ManiSkill: `state`, `pointcloud`, and `rgbd`. In `state` mode, the observation is identical to s . In `pointcloud` and `rgbd` modes, the *object states* in s are replaced by the corresponding point cloud / RGB-D visual observations captured from a panoramic camera mounted on a robot. `state` mode is not suitable for studying generalizability, as *object states* are not available in realistic setups, where information such as object pose has to be estimated based on some forms of visual inputs that are universally obtainable (e.g. point clouds and RGB-D images).

For each task, objects are partitioned into training objects $\mathcal{O}_{train}$ and test objects $\mathcal{O}_{test}$, and environments are divided into training environments $\{T_{o,l} | o \in \mathcal{O}_{train}\}$ and test environments $\{T_{o,l} | o \in \mathcal{O}_{test}\}$. For each training environment, successful demonstration trajectories are provided to facilitate learning-from-demonstrations approaches.

We define **object-level generalizable manipulation skill** as a manipulation skill that can generalize to unseen test objects after learning on training objects where the training and test objects are from the same category. Some notable challenges of our tasks come from partial observations (i.e. point clouds / RGB-D images only covering a portion of an object), robot arms occluding parts of an object, and complex shape understanding over objects with diverse topological and geometric properties.

2.2 Tasks with Diverse Motions and Skills

Object manipulation skills are usually associated with certain types of desired motions of target objects, e.g. rotation around an axis. Thus, the tasks in ManiSkill are designed to cover different types of object motions. We choose four common types of motion constraints: revolute joint constraint, prismatic joint constraint, planar motion constraint, and no constraints, and build four tasks to exemplify each of these motion types. In addition, different tasks also feature different properties of manipulation, such as dual-arm collaboration and solvability by motion planning. Statistics for our tasks are summarized in Table 1. Descriptions for our tasks are stated below (more details in Sec B of supplementary).

OpenCabinetDoor exemplifies motions constrained by a revolute joint. In this task, a single-arm robot is required to open a designated door on a cabinet. The door motion is constrained by a revolute joint attached to the cabinet body. This task is relatively easy to solve by traditional motion planning and control pipelines, so it is suitable for comparison between learning-based methods and motion planning-based methods.

OpenCabinetDrawer exemplifies motions constrained by a prismatic joint. This task is similar to OpenCabinetDoor, but the robot needs to open a target drawer on a cabinet. The drawer motion is constrained by a prismatic joint attached to the cabinet body.

PushChair exemplifies motions constrained on a plane through wheel-ground contact. A dual-arm robot needs to push a swivel chair to a target location on the ground and prevent it from falling over. PushChair exemplifies the ability to manipulate complex underactuated systems, as swivel chairs generally have many joints, resulting in complex dynamics. Therefore, it is difficult to solve PushChair by motion planning and favors learning-based methods.

MoveBucket exemplifies motions without constraints. In this task, a dual-arm robot needs to move a bucket with a ball inside and lift it onto a platform. There are no constraints on the motions of the bucket. However, this task is still very challenging because: 1) it heavily relies on two-arm coordination as the robot needs to lift the bucket; 2) the center of mass of the bucket-ball system is consistently changing, making balancing difficult.

Note that all environments in ManiSkill are verified to be solvable, i.e., for each object, we guarantee that there is a way to manipulate it to solve the corresponding task. This is done by generating successful trajectories in each environment (details in Sec 2.4). Instead of creating lots of tasks but leaving the solvability problems to users, our tasks are constructed with appropriate difficulty and verified solvability.



Figure 4: RGB-D (RGB/Depth) and point cloud observations in ManiSkill. Left two images: RGB-D image from *one* of the three cameras mounted on the robot. The three cameras together provide an ego-centric panoramic view. Right image: visualization of fused point cloud from *all* three cameras. The center of robot body cannot be seen since the captured point cloud comes from an ego-centric view. Parts of the chair are occluded by itself (as cameras are mounted on the robot).

2.3 Robots, Actions, Visual Observations, and Rewards

All the tasks in ManiSkill use similar robots, which are composed of three parts: moving platform, Sciurus [65] robot body, and one or two Franka Panda [66] arm(s). The moving platform can move and rotate on the ground plane, and its height is adjustable. The robot body is fixed on top of the platform, providing support for the arms. Depending on the task, one or two robot arm(s) are connected to the robot body. There are 22 joints in a dual-arm robot and 13 for a single-arm robot. To match realistic robotics setups, we use PID controllers to control the joints of robots. The action space corresponds to the normalized target values of all controllers. In addition to joint space control, ManiSkill supports operational space control, i.e., directly controlling the end-effectors in the Cartesian space.

As mentioned in Sec 2.1, ManiSkill supports three observation modes: `state`, `pointcloud`, and `rgb_d`, where the latter two modes are suitable for studying object-level generalizability. The RGB-D and point cloud observations are captured from three cameras mounted on the robot to provide an ego-centric panoramic view, resembling common real-world robotics setups. The three cameras are 120° apart from each other, and the resolution of each camera is 400×160 . The observations from all cameras are combined to form a final panoramic observation. Visualizations of RGB-D / point cloud observations are shown in Fig 4. In addition, we provide some task-relevant segmentation masks for both RGB-D and point cloud observations (details in Sec B.3 of supplementary).

ManiSkill supports two kinds of rewards: `sparse` and `dense`. A sparse reward is a binary signal which is equivalent to a task-specific success condition. Learning with sparse rewards is very difficult. To alleviate such difficulty, we carefully designed well-shaped dense reward functions for each task (details in Sec C.1 of supplementary). The dense rewards are also used in demonstration collection, which will be elaborated in Sec 2.4.

2.4 RL-Based Demo Collection and MPC-Assisted Reward Template Design

Our interactive environment naturally supports methods such as reinforcement learning or classical robotics pipelines. However, to build an algorithm with object-level generalizability either by RL or by designing rules, it is probably prohibitively complex and resource-demanding for many researchers interested in manipulation learning research (e.g., vision researchers might be primarily interested in the perception module). We observe that many learning-from-demonstrations algorithms (e.g., behavior cloning) are much easier to start with and require less resources.

To serve more researchers, ManiSkill provides a public demonstration dataset with a total of $\sim 36,000$ successful trajectories and $\sim 1.5\text{M}$ frames (300 trajectories for each training object in each task). The demonstrations are provided in the format of internal environment states to save storage space, and users can render the corresponding point cloud / RGB-D frames using the provided scripts.

In order to construct this large-scale demonstration dataset, we need a **scalable** pipeline that can produce plentiful demonstrations automatically. Compared to other existing approaches (e.g. human annotation by teleoperation, motion planning), we adopted a reinforcement learning (RL)-based

pipeline, which requires significantly less human effort and can generate an arbitrary number of demonstrations at scale. Though it is common to collect demonstrations by RL [67, 68, 69, 53, 70, 71, 72, 73], directly training a single RL agent to collect demonstrations for all environments in ManiSkill is challenging because of the large number of different objects and the difficulty of our manipulation tasks. To collect demonstrations at scale, we design an effective pipeline as follows.

Our pipeline contains two stages. In the first stage, we need to design and verify dense rewards for RL agents. For each task, we design a shared reward template based on the skill definitions of this task with human prior. *Note that this reward template is shared across all the environments (objects) in a task instead of manually designed for each object.* In order to quickly verify the reward template (as our tasks are complicated and solving by RL takes hours), we use Model-Predictive Control (MPC) via Cross Entropy Method (CEM), which can be efficiently parallelized to find a trajectory within 15 minutes (if successful) from *one single initial state* using 20 CPUs. While MPC is an efficient tool to verify our reward template, it is not suitable for generating our demonstrations dataset, which should contain diverse and randomized initial states. This is because MPC has to be retrained independently each time to find a trajectory from each of the 300 initial states for each training object of each task, rendering it unscalable. Therefore, in the second stage, we train model-free RL agents to collect demonstrations. We also found that training one single RL agent on many environments (objects) of a task is very challenging, but training an agent to solve a single specific environment is feasible and well-studied. Therefore, we collect demonstrations in a divide-and-conquer way: for each environment, we train an SAC [60] agent and generate successful demonstrations. More details can be found in Sec C of supplementary.

2.5 Multi-Track Training-Evaluation Protocol

As described in Sec 2.1, agents are trained on the training environments with their corresponding demonstrations and evaluated on the test environments for object-level generalizability, under pointcloud or rgbd mode. Moreover, ManiSkill benchmark aims to encourage interdisciplinary insights from computer vision, reinforcement learning, and robotics to advance generalizable physical object manipulation. To this end, we have developed our benchmark with 3 different tracks:

No Interactions: this track requires solutions to only use our provided demonstration trajectories during training. No interactions (i.e. additional trajectory collection, online training, etc.) are allowed. For this track, solutions may choose to adopt a simple but effective supervised learning algorithm — matching the predicted action with the demonstration action given visual observation (i.e. behavior cloning). Therefore, this track encourages researchers to explore 3D computer vision network architectures for generalizable shape understanding over complex topologies and geometries.

No External Annotations: this track allows online model fine-tuning over training environments on top of No Interactions track. However, the solution must not contain new annotations (e.g. new articulated objects from other datasets). This track encourages researchers to explore online training algorithms, such as reinforcement learning with online data collection.

No Restrictions: this track allows solutions to adopt any approach during training, such as labelling new data and creating new environments. Researchers are also allowed to use manually designed control and motion planning rules, along with other approaches from traditional robotics.

The benchmark evaluation metric is the *mean success rate* on a predetermined set of test environment instances. For each task, we have defined the success condition (described in Sec B.4 of supplementary), which is automatically reported in the evaluation script provided by us. Each track should be benchmarked separately.

2.6 Asset Selection, Re-Modeling, Postprocessing, and Verification

While the PartNet-Mobility dataset (from SAPIEN [62]) provides a repository of articulated object models, the original dataset can only provide full support for vision tasks such as joint pose estimation. Therefore, we take significant efforts to select, fix, and verify the models.

First, PartNet-Mobility dataset is not free of annotation errors. For example, some door shafts are annotated at the same side as the door handles. While it does not affect the simulation, such models are unnatural and not good candidates to test policy generalizability. Thus, we first render all assets and manually exclude the ones with annotation errors.

Moreover, fast simulation requires convex decomposition of 3D assets. However, the automatic algorithm used in the original SAPIEN paper, VHACD [74], cannot handle all cases well. For example, VHACD can introduce unexpected artifacts, such as dents on smooth surfaces, which agents can take advantage of. To fix the errors, we identify problematic models by inspection and use Blender’s [75] shape editing function to manually decompose the objects.

Even with all the efforts above, some models can still present unexpected behaviors. For example, certain cabinet drawers may be stuck due to inaccurate overlapping between collision shapes. Therefore, we also verify each object by putting them in the simulator and learn a policy following Sec 2.4. We fix issues if we cannot learn a policy to achieve the task. We repeat until all models can yield a successful policy by MPC.

3 Baseline Architectures, Algorithms, and Experiments

Learning object-level generalizable manipulation skills through 3D visual inputs and learning-from-demonstrations algorithms has been underexplored. Therefore, we designed several baselines and open-sourced their implementations [here](#) to encourage future explorations in the field.

We adopted pointcloud observation mode and designed point cloud-based vision architectures as our feature extractor since previous work [46] has achieved significant performance improvements by using point clouds instead of RGB-D images. Point cloud features include position, RGB, and segmentation masks (for the details of segmentation masks, see Sec B.3 in the supplementary), and we concatenate the *robot state* to the features of each point. Intuitively, this allows the extracted feature to not only contain geometric information of objects, but also contain the relation between the robot and each individual object, such as the closest point to the robot, which is very difficult to be learned without such concatenation. In addition, we downsample the point cloud data to increase training speed and reduce the memory footprint (see Sec D.1 of supplementary).

The first point cloud-based architecture uses one single PointNet [39], a very popular 3D deep learning backbone, to extract a global feature for the entire point cloud, which is fed into the final MLP. The second architecture uses different PointNets to process points belonging to different segmentation masks. The global features from the PointNets are then fed into a Transformer [76], after which a final attention pooling layer extracts the final representations and feeds into the final MLP. We designed and benchmarked this architecture since it allows the model to capture the relation between different objects and possibly provides better performance. Details of the architectures are presented in Sec D.2 of the supplementary material, and a detailed architecture diagram of PointNet + Transformer is presented in Fig 7 of the supplementary material. While there is a great room to improve, we believe that these architectures could provide good starting points for many researchers.

For learning-from-demonstrations algorithms on top of point cloud architectures, we benchmark two approaches - Imitation Learning (IL) and Offline/Batch Reinforcement Learning (Offline/Batch RL). For imitation learning, we choose a simple and widely-adopted algorithm: behavior cloning (BC) - directly matching predicted and ground truth actions through minimizing L_2 distance. For offline RL, we benchmark Batch-Constrained Q-Learning (BCQ) [55] and Twin-Delayed DDPG[77] with Behavior Cloning (TD3+BC) [78]. We follow their original implementations and tune the hyperparameters. Details of the algorithm implementations are presented in Sec D of the supplementary material.

3.1 Single Environment Results

#Demo Trajectories	10	30	100	300	1000
#Gradient Steps	2000	4000	10000	20000	40000
PointNet, BC	0.13	0.23	0.37	0.68	0.76
PointNet + Transformer, BC	0.16	0.35	0.51	0.85	0.90
PointNet + Transformer, BCQ	0.02	0.05	0.23	0.45	0.55
PointNet + Transformer, TD3+BC	0.03	0.13	0.22	0.31	0.57

Table 2: The average success rates of different agents on **one single environment** (fixed object instance) of OpenCabinetDrawer with different numbers of demonstration trajectories. The average success rates are calculated over 100 evaluation trajectories. While network architectures and algorithms play an important role in the performance, learning manipulation skills from demonstrations is challenging without a large number of trajectories, even in one single environment.

As a glimpse into the difficulty of learning manipulation skills from demonstrations in our benchmark, we first present results with an increasing number of demonstration trajectories on **one single environment** of OpenCabinetDrawer in Table 2. We observe that the success rate gradually increases as the number of demonstration trajectories increases, which shows the agents can indeed benefit from more demonstrations. We also observe that inductive bias in network architecture plays an important role in the performance, as PointNet + Transformer is more sample efficient than PointNet. Interestingly, we did not find offline RL algorithms to outperform BC. We conjecture that this is because the provided demonstrations are all successful ones, so an agent is able to learn a good policy through BC. In addition, our robot’s high degree of freedom and the difficulty of the task itself pose a challenge to offline RL algorithms. Further discussions on this observation are presented in Sec D.3 of the supplementary material. It is worth noting that our experiment results should not discourage benchmark users to include failure trajectories and find better usage of offline RL methods, especially those interested in the No External Annotations track described in Sec 2.5.

3.2 Object-Level Generalization Results

Algorithm	BC				BCQ		TD3+BC	
Architecture	PointNet		PointNet + Transformer		PointNet + Transformer		PointNet + Transformer	
Split	Training	Test	Training	Test	Training	Test	Training	Test
OpenCabinetDoor	0.18±0.02	0.04±0.03	0.30±0.06	0.11±0.02	0.16±0.02	0.04±0.02	0.13±0.03	0.04±0.02
OpenCabinetDrawer	0.24±0.03	0.11±0.03	0.37±0.06	0.12±0.02	0.22±0.04	0.11±0.03	0.18±0.02	0.10±0.02
PushChair	0.11±0.02	0.09±0.02	0.18±0.02	0.08±0.01	0.11±0.01	0.08±0.01	0.12±0.02	0.08±0.01
MoveBucket	0.03±0.01	0.02±0.01	0.15±0.01	0.08±0.01	0.08±0.01	0.06±0.01	0.05±0.01	0.03±0.01

Table 3: Mean and standard deviation of *average success rates* on **training and test environments** of each task over 5 different runs, under the point cloud observation. Models are trained with our demonstrations dataset, with 300 demonstration trajectories per training environment. For each task, the average test success rates are calculated over the 10 test environments and 50 evaluation trajectories per environment. Obtaining one single agent capable of learning manipulation skills across multiple objects and generalizing the learned skills to novel objects is challenging.

We now present results on **object-level generalization**. We train each model for 150k gradient steps. This takes about 5 hours for BC, 35 hours for BCQ, and 9 hours for TD3+BC using the PointNet + Transformer architecture on one NVIDIA RTX 2080Ti GPU. As shown in Table 3, even with our best agent (BC PointNet + Transformer), the overall success rates on both training and test environments are low. We also observe that the training accuracy over object variations is significantly lower than the training accuracy on one single environment (in Table 2). The results suggest that existing works on 3D deep learning and learning-from-demonstrations algorithms might have been insufficient yet to achieve good performance when trained for physical manipulation skills over diverse object geometries and tested for object-level generalization. Therefore, we believe there is a large space to improve, and our benchmark poses interesting and challenging problems for the community.

4 Conclusion and Limitations

In this work, we propose ManiSkill, an articulated benchmark for generalizable physical object manipulation from 3D visual inputs with diverse object geometries and large-scale demonstrations. We expect ManiSkill would encourage the community to look into object-level generalizability of manipulation skills, specifically by combining cutting-edge research of 3D computer vision, reinforcement learning, and robotics.

Our benchmark is limited in the following aspects: 1) Currently, we provide 162 articulated objects in total. We plan to process more objects from the PartNet-Mobility dataset [62] and add them to our ManiSkill assets; 2) While the four tasks currently provided in ManiSkill exemplify distinct manipulation challenges, they do not comprehensively cover manipulation skills in household environments. We plan to add more tasks among the same skill properties (e.g, pouring water from one bucket to another bucket through two-arm coordination); 3) We have not conducted sim-to-real experiments yet, and this will be a future direction of ManiSkill.

Acknowledgement

We thank Qualcomm for sponsoring the associated challenge, Sergey Levine and Ashvin Nair for insightful discussions during the whole development process, Yuzhe Qin for the suggestions on building robots, Jiayuan Gu for providing technical support on SAPIEN, and Rui Chen, Songfang Han, Wei Jiang for testing our system.

References

- [1] D. Kalashnikov, A. Irpan, P. Pastor, J. Ibarz, A. Herzog, E. Jang, D. Quillen, E. Holly, M. Kalakrishnan, V. Vanhoucke, et al. Qt-opt: Scalable deep reinforcement learning for vision-based robotic manipulation. *arXiv preprint arXiv:1806.10293*, 2018.
- [2] L. Sergey, N. Wagener, and P. Abbeel. Learning contact-rich manipulation skills with guided policy search. In *Proceedings of the 2015 IEEE International Conference on Robotics and Automation (ICRA)*, Seattle, WA, USA, pages 26–30, 2015.
- [3] S. Levine, C. Finn, T. Darrell, and P. Abbeel. End-to-end training of deep visuomotor policies. *The Journal of Machine Learning Research*, 17(1):1334–1373, 2016.
- [4] A. Rajeswaran, V. Kumar, A. Gupta, G. Vezzani, J. Schulman, E. Todorov, and S. Levine. Learning complex dexterous manipulation with deep reinforcement learning and demonstrations. *arXiv preprint arXiv:1709.10087*, 2017.
- [5] V. Kumar, A. Gupta, E. Todorov, and S. Levine. Learning dexterous manipulation policies from experience and imitation. *arXiv preprint arXiv:1611.05095*, 2016.
- [6] M. Andrychowicz, D. Crow, A. Ray, J. Schneider, R. Fong, P. Welinder, B. McGrew, J. Tobin, P. Abbeel, and W. Zaremba. Hindsight experience replay. In I. Guyon, U. von Luxburg, S. Bengio, H. M. Wallach, R. Fergus, S. V. N. Vishwanathan, and R. Garnett, editors, *Advances in Neural Information Processing Systems 30: Annual Conference on Neural Information Processing Systems 2017, December 4-9, 2017, Long Beach, CA, USA*, pages 5048–5058, 2017. URL <https://proceedings.neurips.cc/paper/2017/hash/453fadbd8a1a3af50a9df4df899537b5-Abstract.html>.
- [7] R. Li, A. Jabri, T. Darrell, and P. Agrawal. Towards practical multi-object manipulation using relational reinforcement learning. In *2020 IEEE International Conference on Robotics and Automation (ICRA)*, pages 4051–4058. IEEE, 2020.
- [8] O. OpenAI, M. Plappert, R. Sampedro, T. Xu, I. Akkaya, V. Kosaraju, P. Welinder, R. D’Sa, A. Petron, H. P. d. O. Pinto, et al. Asymmetric self-play for automatic goal discovery in robotic manipulation. *arXiv preprint arXiv:2101.04882*, 2021.
- [9] A. Zeng, P. Florence, J. Tompson, S. Welker, J. Chien, M. Attarian, T. Armstrong, I. Krasin, D. Duong, V. Sindhwani, et al. Transporter networks: Rearranging the visual world for robotic manipulation. *arXiv preprint arXiv:2010.14406*, 2020.
- [10] J. Mahler, J. Liang, S. Niyaz, M. Laskey, R. Doan, X. Liu, J. A. Ojea, and K. Goldberg. Dex-net 2.0: Deep learning to plan robust grasps with synthetic point clouds and analytic grasp metrics. *arXiv preprint arXiv:1703.09312*, 2017.
- [11] Y. Qin, R. Chen, H. Zhu, M. Song, J. Xu, and H. Su. S4g: Amodal single-view single-shot se (3) grasp detection in cluttered scenes. In *Conference on robot learning*, pages 53–65. PMLR, 2020.
- [12] A. Mousavian, C. Eppner, and D. Fox. 6-dof graspnet: Variational grasp generation for object manipulation. In *2019 IEEE/CVF International Conference on Computer Vision, ICCV 2019, Seoul, Korea (South), October 27 - November 2, 2019*, pages 2901–2910. IEEE, 2019. doi: 10.1109/ICCV.2019.00299. URL <https://doi.org/10.1109/ICCV.2019.00299>.

- [13] H. Fang, C. Wang, M. Gou, and C. Lu. Graspnet-1billion: A large-scale benchmark for general object grasping. In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2020, Seattle, WA, USA, June 13-19, 2020*, pages 11441–11450. IEEE, 2020. doi: 10.1109/CVPR42600.2020.01146. URL <https://doi.org/10.1109/CVPR42600.2020.01146>.
- [14] H. Liang, X. Ma, S. Li, M. Görner, S. Tang, B. Fang, F. Sun, and J. Zhang. PointNetGPD: Detecting grasp configurations from point sets. In *IEEE International Conference on Robotics and Automation (ICRA)*, 2019.
- [15] A. ten Pas, M. Gualtieri, K. Saenko, and R. Platt. Grasp pose detection in point clouds. *The International Journal of Robotics Research*, 36(13-14):1455–1473, 2017.
- [16] Z. Qin, K. Fang, Y. Zhu, L. Fei-Fei, and S. Savarese. Keto: Learning keypoint representations for tool manipulation. In *2020 IEEE International Conference on Robotics and Automation (ICRA)*, pages 7278–7285. IEEE, 2020.
- [17] K. Fang, Y. Zhu, A. Garg, A. Kurenkov, V. Mehta, L. Fei-Fei, and S. Savarese. Learning task-oriented grasping for tool manipulation from simulated self-supervision. *The International Journal of Robotics Research*, 39(2-3):202–216, 2020.
- [18] E. Kolve, R. Mottaghi, W. Han, E. VanderBilt, L. Weihs, A. Herrasti, D. Gordon, Y. Zhu, A. Gupta, and A. Farhadi. AI2-THOR: An Interactive 3D Environment for Visual AI. *arXiv*, 2017.
- [19] Y. Wu, Y. Wu, G. Gkioxari, and Y. Tian. Building generalizable agents with a realistic and rich 3d environment. *arXiv preprint arXiv:1801.02209*, 2018.
- [20] X. Puig, K. Ra, M. Boben, J. Li, T. Wang, S. Fidler, and A. Torralba. Virtual-home: Simulating household activities via programs. In *2018 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2018, Salt Lake City, UT, USA, June 18-22, 2018*, pages 8494–8502. IEEE Computer Society, 2018. doi: 10.1109/CVPR.2018.00886. URL http://openaccess.thecvf.com/content_cvpr_2018/html/Puig_VirtualHome_Simulating_Household_CVPR_2018_paper.html.
- [21] F. Xia, A. R. Zamir, Z. He, A. Sax, J. Malik, and S. Savarese. Gibson env: Real-world perception for embodied agents. In *2018 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2018, Salt Lake City, UT, USA, June 18-22, 2018*, pages 9068–9079. IEEE Computer Society, 2018. doi: 10.1109/CVPR.2018.00945. URL http://openaccess.thecvf.com/content_cvpr_2018/html/Xia_Gibson_Env_Real-World_CVPR_2018_paper.html.
- [22] F. Xia, W. B. Shen, C. Li, P. Kasimbeg, M. E. Tchapmi, A. Toshev, R. Martín-Martín, and S. Savarese. Interactive gibbon benchmark: A benchmark for interactive navigation in cluttered environments. *IEEE Robotics and Automation Letters*, 5(2):713–720, 2020.
- [23] M. Savva, J. Malik, D. Parikh, D. Batra, A. Kadian, O. Maksymets, Y. Zhao, E. Wijmans, B. Jain, J. Straub, J. Liu, and V. Koltun. Habitat: A platform for embodied AI research. In *2019 IEEE/CVF International Conference on Computer Vision, ICCV 2019, Seoul, Korea (South), October 27 - November 2, 2019*, pages 9338–9346. IEEE, 2019. doi: 10.1109/ICCV.2019.00943. URL <https://doi.org/10.1109/ICCV.2019.00943>.
- [24] C. Gan, J. Schwartz, S. Alter, M. Schrimpf, J. Traer, J. De Freitas, J. Kubilius, A. Bhandwaldar, N. Haber, M. Sano, et al. Threedworld: A platform for interactive multi-modal physical simulation. *arXiv preprint arXiv:2007.04954*, 2020.
- [25] K. Ehsani, W. Han, A. Herrasti, E. VanderBilt, L. Weihs, E. Kolve, A. Kembhavi, and R. Mottaghi. Manipulathor: A framework for visual object manipulation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 4497–4506, June 2021.
- [26] G. Brockman, V. Cheung, L. Pettersson, J. Schneider, J. Schulman, J. Tang, and W. Zaremba. Openai gym. *arXiv preprint arXiv:1606.01540*, 2016.

- [27] Y. Tassa, S. Tunyasuvunakool, A. Muldal, Y. Doron, S. Liu, S. Bohez, J. Merel, T. Erez, T. Lillicrap, and N. Heess. dm_control: Software and tasks for continuous control, 2020.
- [28] Y. Zhu, J. Wong, A. Mandlekar, and R. Martín-Martín. robosuite: A modular simulation framework and benchmark for robot learning. In *arXiv preprint arXiv:2009.12293*, 2020.
- [29] Y. Urakami, A. Hodgkinson, C. Carlin, R. Leu, L. Rigazio, and P. Abbeel. Doorgym: A scalable door opening environment and baseline agent. *arXiv preprint arXiv:1908.01887*, 2019.
- [30] T. Yu, D. Quillen, Z. He, R. Julian, K. Hausman, C. Finn, and S. Levine. Meta-world: A benchmark and evaluation for multi-task and meta reinforcement learning. In *Conference on Robot Learning*, pages 1094–1100. PMLR, 2020.
- [31] S. James, Z. Ma, D. R. Arrojo, and A. J. Davison. Rlbench: The robot learning benchmark & learning environment. *IEEE Robotics and Automation Letters*, 5(2):3019–3026, 2020.
- [32] E. Todorov, T. Erez, and Y. Tassa. Mujoco: A physics engine for model-based control. In *2012 IEEE/RSJ International Conference on Intelligent Robots and Systems*, pages 5026–5033. IEEE, 2012.
- [33] E. Coumans and Y. Bai. Pybullet, a python module for physics simulation for games, robotics and machine learning. 2016.
- [34] Nvidia physx sdk. <https://developer.nvidia.com/physx-sdk>. Accessed: 2021-08-27.
- [35] E. Rohmer, S. P. N. Singh, and M. Freese. Coppeliassim (formerly v-rep): a versatile and scalable robot simulation framework. In *Proc. of The International Conference on Intelligent Robots and Systems (IROS)*, 2013. www.coppeliarobotics.com.
- [36] A. Dosovitskiy, G. Ros, F. Codevilla, A. Lopez, and V. Koltun. Carla: An open urban driving simulator. In *Conference on robot learning*, pages 1–16. PMLR, 2017.
- [37] Z. Xu, S. Bi, K. Sunkavalli, S. Hadap, H. Su, and R. Ramamoorthi. Deep view synthesis from sparse photometric images. *ACM Transactions on Graphics (ToG)*, 38(4):1–13, 2019.
- [38] Kinova. <https://www.kinovarobotics.com/en>. Accessed: 2021-08-27.
- [39] C. R. Qi, H. Su, K. Mo, and L. J. Guibas. Pointnet: Deep learning on point sets for 3d classification and segmentation. In *2017 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2017, Honolulu, HI, USA, July 21-26, 2017*, pages 77–85. IEEE Computer Society, 2017. doi: 10.1109/CVPR.2017.16. URL <https://doi.org/10.1109/CVPR.2017.16>.
- [40] C. R. Qi, L. Yi, H. Su, and L. J. Guibas. Pointnet++: Deep hierarchical feature learning on point sets in a metric space. In I. Guyon, U. V. Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017. URL <https://proceedings.neurips.cc/paper/2017/file/d8bf84be3800d12f74d8b05e9b89836f-Paper.pdf>.
- [41] H. Thomas, C. R. Qi, J.-E. Deschaud, B. Marcotegui, F. Goulette, and L. J. Guibas. Kpconv: Flexible and deformable convolution for point clouds. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 6411–6420, 2019.
- [42] B. Graham and L. van der Maaten. Submanifold sparse convolutional networks. *arXiv preprint arXiv:1706.01307*, 2017.
- [43] C. Choy, J. Gwak, and S. Savarese. 4d spatio-temporal convnets: Minkowski convolutional neural networks. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 3075–3084, 2019.
- [44] H. Zhao, L. Jiang, J. Jia, P. Torr, and V. Koltun. Point transformer. *arXiv preprint arXiv:2012.09164*, 2020.

- [45] C. R. Qi, O. Litany, K. He, and L. J. Guibas. Deep hough voting for 3d object detection in point clouds. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 9277–9286, 2019.
- [46] Y. Wang, W.-L. Chao, D. Garg, B. Hariharan, M. Campbell, and K. Q. Weinberger. Pseudo-lidar from visual depth estimation: Bridging the gap in 3d object detection for autonomous driving. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 8445–8453, 2019.
- [47] S. Chitta, B. Cohen, and M. Likhachev. Planning for autonomous door opening with a mobile manipulator. In *2010 IEEE International Conference on Robotics and Automation*, pages 1799–1806. IEEE, 2010.
- [48] F. Burget, A. Hornung, and M. Bennewitz. Whole-body motion planning for manipulation of articulated objects. In *2013 IEEE International Conference on Robotics and Automation*, pages 1656–1662. IEEE, 2013.
- [49] V. Narayanan and M. Likhachev. Task-oriented planning for manipulating articulated mechanisms under model uncertainty. In *2015 IEEE International Conference on Robotics and Automation (ICRA)*, pages 3095–3101. IEEE, 2015.
- [50] X. Li, H. Wang, L. Yi, L. J. Guibas, A. L. Abbott, and S. Song. Category-level articulated object pose estimation. In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2020, Seattle, WA, USA, June 13-19, 2020*, pages 3703–3712. IEEE, 2020. doi: 10.1109/CVPR42600.2020.00376. URL <https://doi.org/10.1109/CVPR42600.2020.00376>.
- [51] M. Mittal, D. Hoeller, F. Farshidian, M. Hutter, and A. Garg. Articulated object interaction in unknown scenes with whole-body mobile manipulation. *arXiv preprint arXiv:2103.10534*, 2021.
- [52] V. Zeng, T. E. Lee, J. Liang, and O. Kroemer. Visual identification of articulated object parts. *arXiv preprint arXiv:2012.00284*, 2020.
- [53] J. Ho and S. Ermon. Generative adversarial imitation learning. In D. D. Lee, M. Sugiyama, U. von Luxburg, I. Guyon, and R. Garnett, editors, *Advances in Neural Information Processing Systems 29: Annual Conference on Neural Information Processing Systems 2016, December 5-10, 2016, Barcelona, Spain*, pages 4565–4573, 2016. URL <https://proceedings.neurips.cc/paper/2016/hash/cc7e2b878868cbae992d1fb743995d8f-Abstract.html>.
- [54] S. Ross, G. Gordon, and D. Bagnell. A reduction of imitation learning and structured prediction to no-regret online learning. In *Proceedings of the fourteenth international conference on artificial intelligence and statistics*, pages 627–635. JMLR Workshop and Conference Proceedings, 2011.
- [55] S. Fujimoto, D. Meger, and D. Precup. Off-policy deep reinforcement learning without exploration. In K. Chaudhuri and R. Salakhutdinov, editors, *Proceedings of the 36th International Conference on Machine Learning, ICML 2019, 9-15 June 2019, Long Beach, California, USA*, volume 97 of *Proceedings of Machine Learning Research*, pages 2052–2062. PMLR, 2019. URL <http://proceedings.mlr.press/v97/fujimoto19a.html>.
- [56] A. Kumar, A. Zhou, G. Tucker, and S. Levine. Conservative q-learning for offline reinforcement learning. In H. Larochelle, M. Ranzato, R. Hadsell, M. Balcan, and H. Lin, editors, *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6-12, 2020, virtual*, 2020. URL <https://proceedings.neurips.cc/paper/2020/hash/0d2b2061826a5df3221116a5085a6052-Abstract.html>.
- [57] S. Levine, A. Kumar, G. Tucker, and J. Fu. Offline reinforcement learning: Tutorial, review, and perspectives on open problems. *arXiv preprint arXiv:2005.01643*, 2020.
- [58] S. Dasari, F. Ebert, S. Tian, S. Nair, B. Bucher, K. Schmeckpeper, S. Singh, S. Levine, and C. Finn. Robonet: Large-scale multi-robot learning. *arXiv preprint arXiv:1910.11215*, 2019.

- [59] S. Cabi, S. G. Colmenarejo, A. Novikov, K. Konyushkova, S. Reed, R. Jeong, K. Zolna, Y. Ayta, D. Budden, M. Vecerik, et al. Scaling data-driven robotics with reward sketching and batch reinforcement learning. *arXiv preprint arXiv:1909.12200*, 2019.
- [60] T. Haarnoja, A. Zhou, K. Hartikainen, G. Tucker, S. Ha, J. Tan, V. Kumar, H. Zhu, A. Gupta, P. Abbeel, et al. Soft actor-critic algorithms and applications. *arXiv preprint arXiv:1812.05905*, 2018.
- [61] J. Schulman, F. Wolski, P. Dhariwal, A. Radford, and O. Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- [62] F. Xiang, Y. Qin, K. Mo, Y. Xia, H. Zhu, F. Liu, M. Liu, H. Jiang, Y. Yuan, H. Wang, L. Yi, A. X. Chang, L. J. Guibas, and H. Su. SAPIEN: A simulated part-based interactive environment. In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2020, Seattle, WA, USA, June 13-19, 2020*, pages 11094–11104. IEEE, 2020. doi: 10.1109/CVPR42600.2020.01111. URL <https://doi.org/10.1109/CVPR42600.2020.01111>.
- [63] K. Mo, S. Zhu, A. X. Chang, L. Yi, S. Tripathi, L. J. Guibas, and H. Su. Partnet: A large-scale benchmark for fine-grained and hierarchical part-level 3d object understanding. In *IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2019, Long Beach, CA, USA, June 16-20, 2019*, pages 909–918. Computer Vision Foundation / IEEE, 2019. doi: 10.1109/CVPR.2019.00100. URL http://openaccess.thecvf.com/content_CVPR_2019/html/Mo_PartNet_A_Large-Scale_Benchmark_for_Fine-Grained_and_Hierarchical_Part-Level_3D_CVPR_2019_paper.html.
- [64] A. X. Chang, T. Funkhouser, L. Guibas, P. Hanrahan, Q. Huang, Z. Li, S. Savarese, M. Savva, S. Song, H. Su, et al. Shapenet: An information-rich 3d model repository. *arXiv preprint arXiv:1512.03012*, 2015.
- [65] Sciurus17. <https://rt-net.jp/products/sciurus17/>. Accessed: 2021-08-27.
- [66] Franka. <https://www.franka.de/>. Accessed: 2021-08-27.
- [67] J. Fu, A. Kumar, O. Nachum, G. Tucker, and S. Levine. D4rl: Datasets for deep data-driven reinforcement learning. *arXiv preprint arXiv:2004.07219*, 2020.
- [68] F. Liu, Z. Ling, T. Mu, and H. Su. State alignment-based imitation learning. *arXiv preprint arXiv:1911.10947*, 2019.
- [69] S. Levine, P. Pastor, A. Krizhevsky, J. Ibarz, and D. Quillen. Learning hand-eye coordination for robotic grasping with deep learning and large-scale data collection. *The International Journal of Robotics Research*, 37(4-5):421–436, 2018.
- [70] L. Blondé and A. Kalousis. Sample-efficient imitation learning via generative adversarial nets. In *The 22nd International Conference on Artificial Intelligence and Statistics*, pages 3138–3148. PMLR, 2019.
- [71] B. Kang, Z. Jie, and J. Feng. Policy optimization with demonstrations. In *International Conference on Machine Learning*, pages 2469–2478. PMLR, 2018.
- [72] M. Jing, X. Ma, W. Huang, F. Sun, C. Yang, B. Fang, and H. Liu. Reinforcement learning from imperfect demonstrations under soft expert guidance. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 34, pages 5109–5116, 2020.
- [73] Z. Zhu, K. Lin, B. Dai, and J. Zhou. Learning sparse rewarded tasks from sub-optimal demonstrations. *arXiv preprint arXiv:2004.00530*, 2020.
- [74] K. Mamou and F. Ghorbel. A simple and efficient approach for 3d mesh approximate convex decomposition. In *2009 16th IEEE international conference on image processing (ICIP)*, pages 3501–3504. IEEE, 2009.
- [75] B. O. Community. *Blender - a 3D modelling and rendering package*. Blender Foundation, Stichting Blender Foundation, Amsterdam, 2018. URL <http://www.blender.org>.

- [76] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser, and I. Polosukhin. Attention is all you need. In I. Guyon, U. von Luxburg, S. Bengio, H. M. Wallach, R. Fergus, S. V. N. Vishwanathan, and R. Garnett, editors, *Advances in Neural Information Processing Systems 30: Annual Conference on Neural Information Processing Systems 2017, December 4-9, 2017, Long Beach, CA, USA*, pages 5998–6008, 2017. URL <https://proceedings.neurips.cc/paper/2017/hash/3f5ee243547dee91fbd053c1c4a845aa-Abstract.html>.
- [77] S. Fujimoto, H. Hoof, and D. Meger. Addressing function approximation error in actor-critic methods. In *International Conference on Machine Learning*, pages 1587–1596. PMLR, 2018.
- [78] S. Fujimoto and S. S. Gu. A minimalist approach to offline reinforcement learning. *CoRR*, abs/2106.06860, 2021. URL <https://arxiv.org/abs/2106.06860>.
- [79] V. Mnih, K. Kavukcuoglu, D. Silver, A. A. Rusu, J. Veness, M. G. Bellemare, A. Graves, M. Riedmiller, A. K. Fidjeland, G. Ostrovski, et al. Human-level control through deep reinforcement learning. *nature*, 518(7540):529–533, 2015.
- [80] J. Zhu, D. Zhu, and J. Cai. Design of sliding mode controller with low-pass filter for robot manipulators. In *2009 Second International Workshop on Knowledge Discovery and Data Mining*, pages 296–299. IEEE, 2009.

Supplementary Material

Table of Contents

A Overview	17
B Further Details of Tasks and System	17
B.1 Environment Speed	17
B.2 Environment Parameters	18
B.3 Segmentation Masks	18
B.4 Success Conditions	18
B.5 Evaluation Kit	18
B.6 System Design	19
B.7 Controller Design	19
B.8 Environment Validation	19
B.9 Manually Processed Collision Shapes	19
C Details of Demonstration Collection	20
C.1 Dense Reward Design	20
C.2 Agent Training and Demonstration Collection	21
D Implementation Details of Baseline Architectures, Algorithms, and Experiments	22
D.1 Point Cloud Subsampling	22
D.2 Network Architectures	22
D.3 Implementation Details of Learning-from-Demonstration Algorithms	23

A Overview

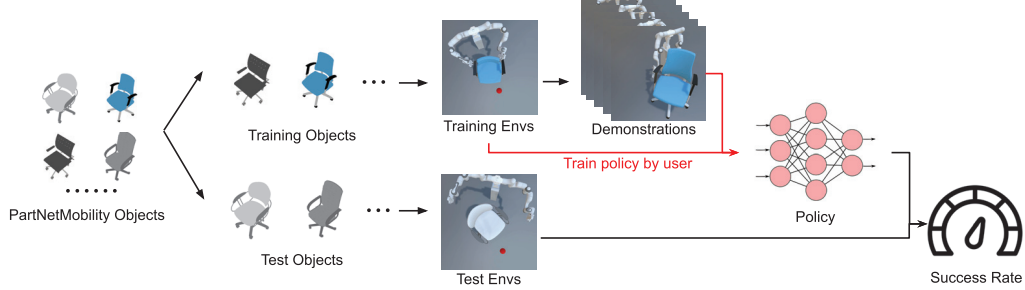


Figure 5: Overall illustration of ManiSkill. We manually re-model and postprocess objects from the PartNet-Mobility dataset, split them into training and test sets, and then build the corresponding training and test environments in the SAPIEN simulator. We then generate successful demonstration trajectories on the training environments. Users are expected to build policies based on the demonstration trajectories and the training environments, then evaluate the *mean success rate* on the test environments with the provided evaluation kit.

This supplementary material includes details on the system and task design of ManiSkill, along with implementations of baseline experiments. An overall architecture of our ManiSkill Benchmark presented in the main paper is summarized in Figure 5.

Section B and C provide more details on the system, tasks, and demonstration collection.

Section D.1 provides implementation details of point cloud subsampling in our baselines.

Section D.2 provides implementation details of our point cloud-based baseline network architectures, along with a diagram of our PointNet + Transformer model.

Section D.3 provides implementation details of learning-from-demonstrations algorithms, specifically imitation learning (Behavior Cloning) and offline RL.

B Further Details of Tasks and System

B.1 Environment Speed

Observation Mode	state	pointcloud/rgbd
OpenCabinetDoor	112 ± 4	48 ± 3
OpenCabinetDrawer	113 ± 4	47 ± 2
PushChair	53 ± 9	31 ± 4
MoveBucket	61 ± 7	34 ± 3

Table 4: Mean and standard deviation of FPS (frame per second) of the environments in ManiSkill. In `state` mode, most computations are used on physical simulation. In `pointcloud` and `rgbd` modes, most computations are used on rendering. All the numbers are tested on a single Intel i9-9960X CPU and a single NVIDIA RTX TITAN GPU.

We provide the running speeds of our environments in Table 4. We would like to note that the numbers cannot be directly compared with other environments like the MuJoCo tasks in OpenAI gym based on the following reasons.

1. In our environments, one environment step corresponds to 5 control steps. This “frame-skipping” technique can make the horizon of our tasks shorter, which is also a common practice in reinforcement learning [79]. We can make the environment FPS 5x larger by simply disabling frame-skipping, but this will make the tasks more difficult as the agent needs to make more decisions.

2. Compared to other environments, such as the MuJoCo tasks in OpenAI Gym, our articulated objects and robots are much more complicated. For example, our chairs contain up to 20 joints and tens of thousands of collision meshes. Therefore, the physical simulation process is inherently slow.
3. Many other robotics/control environments do not provide visual observations, while ManiSkill does. When generating visual observations, rendering is a very time-consuming process, especially when we are using **three** cameras simultaneously.

B.2 Environment Parameters

For each environment $\mathcal{E}_o$ and its associated object o , the environment parameters $\mathcal{L}_o$ provide randomization for our robot and for the target object.

For our robot, we randomly initialize the robot’s in-plane position (x, y) and orientation (rotation around z -axis). Joints on the robot arms are initialized to a canonical pose, which is a common practice in robotics tasks.

For the target object, several physical parameters, such as joint frictions, are randomized for all tasks. In MoveBucket and PushChair, we also randomize the initial in-plane position (x, y) and the initial orientation (rotation around z -axis) of the bucket and the chair. Detailed implementations can be found in our repository <https://github.com/haosulab/ManiSkill>.

B.3 Segmentation Masks

As mentioned in Sec 2.3, we provide task-relevant segmentation masks in pointcloud and rgb-d modes. Each mask is a binary array indicating a part or an object. Here are the details about our segmentation masks for each task:

- OpenCabinetDoor: handle of the target door, target door, robot (3 masks in total)
- OpenCabinetDrawer: handle of the target drawer, target drawer, robot (3 masks in total)
- PushChair: robot (1 mask in total)
- MoveBucket: robot (1 mask in total)

Basically, we provide the robot mask and any mask that is necessary for specifying the target. For example, in OpenCabinetDoor/Drawer environments, a cabinet might have many doors/drawers, so we provide the door/drawer mask such that users know which door/drawer to open. We also provide handle mask such that users know from which direction the door/drawer should be opened.

B.4 Success Conditions

OpenCabinetDoor and **OpenCabinetDrawer** Success is marked by opening the joint to 90% of its limit and keeping it static for a period of time afterwards. We do not constrain how the door/drawer is opened, e.g, the door can be opened by grabbing the side of the door or grasping the handle.

PushChair The task is successful if the chair (1) is close enough (within 15 centimeters) to the target location; (2) is kept static for a period of time after being close enough to the target; and (3) does not fall over.

MoveBucket The task is successful if (1) the bucket is placed on or above the platform at the upright position and kept static for a period of time, and (2) all the balls remain in the bucket.

In all tasks, the time limit for each episode is 200, which is sufficient to solve the task. An episode will be evaluated as unsuccessful if it goes beyond the time limit.

B.5 Evaluation Kit

ManiSkill provides a straightforward evaluation script. The script takes a task name, an observation mode (RGB-D or point cloud), and a solution file as input. The solution file is expected to contain a single policy function that takes observations as input and outputs an action. The evaluation kit takes

the policy function and evaluates it on the test environments. For each environment, it reports the average success rate and the average satisfactory rate for each success conditions (e.g. whether the ball is inside the bucket and whether the bucket is on or above the platform in MoveBucket).

B.6 System Design

We configure our simulation environments by a YAML-based configuration system. This system is mainly used to configure physical properties, rendering properties, and scene layouts that can be reused across tasks. It allows benchmark designers to specify simulation frequencies, physical solver parameters, lighting conditions, camera placement, randomized object/robot layouts, robot controller parameters, object surface materials, and other common properties shared across all environments. After preparing the configurations, designers can load the configurations as SAPIEN scenes and perform further specific customization with Python scripts. In our task design, after we build the environments, we manually validate them to make sure they behave as expected (see B.8 for details).

B.7 Controller Design

The joints in our robots are controlled by velocity or position controllers. For velocity controllers, we use the built-in inverse dynamics functions in PhysX to compute the balancing forces for a robot. We then apply the internal PD controllers of PhysX by setting stiffness to 0 and damping to a positive constant, where damping is used to drive a robot to a given velocity. We additionally add a first-order low-pass filter, implemented as an exponential moving average, to the input velocity signal, which is a common practice in real robotics systems [80]. Position controllers are built on top of velocity controllers: the input position signal is passed into a PID controller, which outputs a velocity signal for a velocity controller.

B.8 Environment Validation

Environment construction is not complete without testing. After modeling the environment, we need to ensure the environment has the following properties: (1) The environment is correctly modeled with realistic parameters. As the environment contains hundreds of parameters, including the friction coefficients, controller parameters, object scales, etc, its correctness needs comprehensive checking; (2) The environment is solvable. We need to check that the task can be completed within the allocated time steps, and that the physical parameters allow task completion. For example, the weight of a target object to be grasped must be smaller than the maximum force allowed for lifting the robot gripper; (3) The physical simulation is free from significant artifacts. Physical simulation faces the trade-off between stability and speed. When the simulation frequency and the contact solver iterations are too small, simulation artifacts such as jittering and interpenetration can occur; (4) There does not exist undesired exploits and shortcuts.

To inspect our environments, after modeling them, we first use the SAPIEN viewer to visually inspect the appearance of all assets. We also inspect the physical properties of crucial components in several sampled environments. Next, we design a mouse-and-keyboard based robot controller. It controls the robot gripper in Cartesian coordinates using inverse kinematics of the robot arm. We try to manually solve the tasks to identify potential problems in the environments. This manual process can identify most solvability issues and physical artifacts. However, an agent could still learn unexpected exploits, entailing us to iteratively improve the environments: we execute the demonstration collection algorithm on the current environment and record videos of sampled demonstration trajectories. We then watch the videos to identify causes of success and failure, potentially spotting unreasonable behaviors. We finally investigate and improve the environment.

B.9 Manually Processed Collision Shapes

As described in Section 4.1.3, we manually decompose the collision shapes into convex shapes. This manual process is performed on the Bucket objects. We justify this choice in Fig. 6.

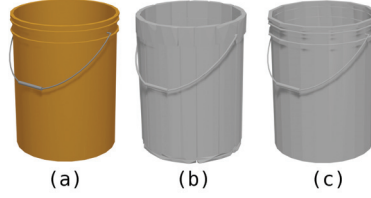


Figure 6: When decomposing a bucket (a), standard VHACD [74] algorithm (b, 2340 faces) misses details, and tends to produce artifacts, such as bumps and seams, that make visual appearances quite different from collision shapes, so we manually process the mesh (c, 1445 faces).

C Details of Demonstration Collection

C.1 Dense Reward Design

A ManiSkill task is defined on different objects of the same category. This natural structure implies an efficient way to automatically construct dense reward functions for all environments of each task. For each task, we manually design a general template, then automatically translate this template into a smooth reward for each environment. The reward templates are also open-sourced.

We use a multi-stage reward template for all of our tasks. In each stage, we guarantee that rewards in the next stage are strictly larger than rewards in the current stage to prevent RL agents from staying in an intermediate stage forever. We also carefully design our reward template at stage-transition states to ensure smoothness of our rewards.

C.1.1 OpenCabinetDoor and OpenCabinetDrawer

For environments in OpenCabinetDoor and OpenCabinetDrawer, our reward template contains three stages. In the first stage, an agent receives rewards from being close to the handle on the target link (door or drawer). To encourage contact with the target link, we penalize the Euclidean distance between the handle and the gripper. When the gripper’s distance to the target link is less than a threshold, the agent enters the second stage. In this stage, the agent gets a reward from the opening angle of the door or the opening distance of the drawer. When the agent opens the door or drawer enough, the agent enters the final stage. In the final stage, the agent receives a negative reward based on the speed of the target link to encourage the scene to be static at the end of the trajectory.

C.1.2 PushChair

The reward template for PushChair contains three stages. In the first stage, an agent receives reward from moving towards the chair. To encourage contact with the chair, we compute the distance between the robot end effectors and the chair and take its logarithm as reward. When the robot end effectors are close enough to the chair, the agent enters the second stage. In the second stage, the agent receives rewards based on the distance between the chair’s current location and the target location. The agent receives additional rewards based on the angle between the chair’s velocity vector and the vector pointing towards the target location. In our experiments, we find that this term is critical. When the chair is close enough to the target location, the agent enters the final stage. In the final stage, the agent is penalized based on the linear and angular velocity of the chair, such that the agent learns to keep the chair static. In all stages, the agent is penalized based on the chair’s degree of tilt in order to keep the chair upright.

C.1.3 MoveBucket

The reward template for MoveBucket consists of four stages. In the first stage, an agent receives rewards from moving towards the bucket. To encourage contact with the bucket, we compute the distance between the robot end effectors (grippers) and the bucket and take its log value as reward. When the robot end effectors are close enough to the bucket, the agent enters the second stage. In the second stage, the agent is required to lift the bucket to a specific height. The agent receives a position-based reward and a velocity-based reward that encourage the agent to lift the bucket. In our experiments, we find that it is very difficult for the agent to learn how to lift the bucket without any

domain knowledge. To ease the difficulty, we use the angle between the two vectors pointing from the two grippers to the bucket’s center of mass as a reward. This term encourages the agent to place the two grippers on opposite sides of the bucket. We also penalize the agent based on the grippers’ height difference in the bucket frame so that the grasp pose is more stable. Once the bucket is lifted high enough, the agent enters the third stage. In this stage, the agent receives a position-based reward and a velocity-based reward that encourages the agent to move the bucket towards the target platform. When the bucket is on top of the platform, the agent enters the final stage, and it is penalized based on the linear and angular velocity of the bucket, such that the agent learns to hold the bucket steadily. In all stages, the agent is also penalized based on the bucket’s degree of tilt to keep the bucket upright. Since it is harder to keep the bucket upright in MoveBucket than in PushChair, we take the log value of the bucket’s degree of tilt as an additional penalty term so that the reward is more sensitive at near-upright poses.

C.2 Agent Training and Demonstration Collection

# Different Cabinets	1	5	10	20
Success Rate	100%	82%	2%	0%

Table 5: The success rates of SAC [60] agents on OpenCabinetDrawer trained from scratch with 10^6 time-steps on different numbers of cabinets. The SAC agents are trained in the state mode using our designed dense rewards. Jointly training a single RL agent on a large number of environments (objects) from scratch to collect demonstrations is infeasible.

Since different environments of a task contain different objects of the same category, a straightforward idea to collect demonstrations is to train a single agent from scratch directly on all environments through trial-and-error, which seems feasible at first glance. However, as shown in Table 5, even with carefully-designed rewards, the performance of such approach drops sharply as the number of different objects increases.

While directly training one single RL agent on many environments of a task is very challenging, training an agent to solve a single specific environment is feasible and well-studied. Therefore, we collect demonstrations in a divide-and-conquer way: We train a population of SAC [60] agents on a task and ensure that each agent is able to solve a specific environment. These agents are then used to interact with their corresponding environments to generate successful trajectories. In this way, we can generate an arbitrary number of demonstration trajectories.

For all environments, we train our agents for 2.0×10^6 steps. To ensure the quality of our demonstrations, in a few cases where the success rate of an SAC agent is less than 0.3, we retrain the agent. We then uniformly sample the initial states and use the trained agents to collect 300 successful trajectories for each environment of each task.

To speed up the training process, we use 4 processes in parallel to collect samples. For better demonstration quality, during training, we remove the early-done signal when agents succeed. While this potentially lowers the success rates of agents, we find that this leads to more robust policy at near-end states. However, during demonstration collection, we stop at the first success signal. The entire training and demonstration generation process takes about 2 days in total for all environments and all tasks on 4 8-GPU and 64-CPU machines. The detailed hyperparameters of SAC can be found in Table 6.

We would like to note that our benchmark is not a generic learning-from-demonstrations benchmark, so comparing the diversity of our RL-generated demonstrations with human demonstrations is not the focus of our work (we focus on solving the tasks themselves). However, we do empirically observe that our RL-generated demonstrations are diverse, especially given that a large number of demonstrations have been provided. For example, in OpenCabinetDrawer, the demonstration trajectories show two different behaviors: opening the drawer by pulling the handle, or by pulling the side of the drawer out. One reason for such diverse behaviors could come from our reward design: While the multi-stage reward template we designed encourages an agent to grasp the handle of a drawer, we do not explicitly restrict how it opens the drawer. As long as the opening method is feasible and achieves the desired goal state, the trajectory will be considered successful.

Hyperparameters	Value
Optimizer	Adam
Learning rate	3×10^{-4}
Discount (γ)	0.95
Replay buffer size (γ)	10^6
Number of hidden layers (all networks)	3
Number of hidden units per layer	256
Number of threads for collecting samples	4
Number of samples per minibatch	1024
Nonlinearity	ReLU
Target smoothing coefficient(τ)	0.005
Target update interval	1
Gradient steps	1
Total Simulation Steps	2×10^6

Table 6: The hyperparameters of SAC for demonstration generation.

D Implementation Details of Baseline Architectures, Algorithms, and Experiments

D.1 Point Cloud Subsampling

Since point clouds are captured from the 3 cameras mounted on the robot with resolution 400×160 , as described in Section 2.3, point clouds without postprocessing (192000 points per frame) would be very memory efficient and significantly reduces the training speed. Therefore, we downsample the raw point cloud by first sampling 50 points for each segmentation mask (if there are fewer than 50 points, then we keep all the points). We then randomly sample from the rest of the points where at least one of the segmentation masks is true, such that the total number of those points is 800 (if there are fewer than 800, then we keep all of them). Finally, we randomly sample from the points where none of the segmentation masks are true and where the points are not on the ground (i.e. have positive z -coordinate value), such that we obtain a total of 1200 points.

For the convenience of researchers and benchmark users, we also made this downsampled point cloud demonstrations dataset public, so users do not need to render the demonstrations locally if they wish to use our subsampling function. Instructions for download can be found in our repository <https://github.com/haosulab/ManiSkill-Learn>.

D.2 Network Architectures

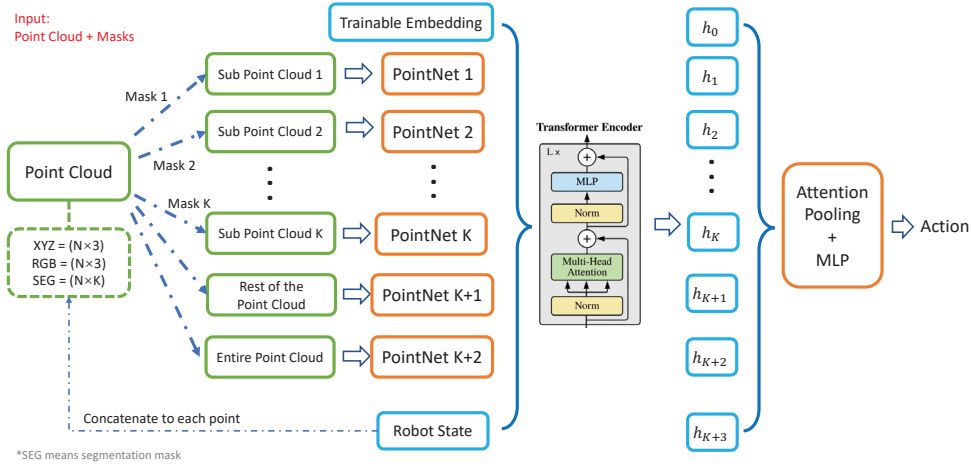


Figure 7: Architecture diagram for our “PointNet + Transformer” model.

For all of our PointNet policy network models, we concatenate the features of each point (which include position, RGB, and segmentation masks) with the *robot state* (as mentioned in Section 2.1 in our main paper) to form new point input features. For the position feature, we first calculate the mean coordinates for the point cloud / sub point cloud, then concatenate it with the original position subtracting the mean. We found such normalized position feature significantly improve performance.

In our vanilla PointNet model, we feed all point features into one single PointNet. The PointNet has hidden layer dimensions [256, 512], and the global feature is passed through an MLP with layer sizes [512, 256, action_dim] to output actions.

For our PointNet + Transformer model, we use different PointNets to process points having different segmentation masks. If the masks have dimension k , then we use $k + 2$ PointNets (one for each of the segmentation masks, one for the points without any segmentation mask, and one for the entire point cloud) with hidden dimension 256 to extract $k + 2$ global features. We also use an additional MLP to output a 256-d hidden vector for the robot state alone (i.e. the robot state is not only concatenated with the point features and fed into the PointNets, but also processed alone through this MLP). The global point features, the processed robot state vector, and an additional trainable embedding vector (serving as a bias for the task) are fed into a Transformer [76] with $d_{\text{model}} = 256$ and $d_{\text{ff}} = 1024$. We did not add position encoding to the Transformer, as we found it significantly hurts performance. The output vectors are passed through a global attention pooling to extract a representation of dimension 256, which is then fed into a final MLP with layer sizes [256, 128, action_dim] to output actions. A diagram of this architecture is presented in Figure 7. All of our models use ReLU activation.

D.3 Implementation Details of Learning-from-Demonstration Algorithms

We benchmark imitation learning with Behavior Cloning (BC), along with two offline-RL algorithms: Batch-Constrained Q-Learning (BCQ) [55] and Twin-Delayed DDPG with Behavior Cloning (TD3+BC) [78]. Different from BC, BCQ does not directly clone the demonstration actions given input, and instead uses a VAE to fit the distribution of actions in the demonstration. It then learns a Q function that estimates the reward of actions given input, and selects an action with the best reward among samples during inference. TD3+BC [78] adds a weighted BC loss to the TD3 loss to constrain the output action to the demonstration data. The original paper also normalizes the features of every state in the demonstration dataset, but this trick is not applicable in our case as our inputs are visual. There are also other offline-RL algorithms like CQL [56], and we leave them for future work.

Hyperparameters	Value
Batch size	64
Perturbation limit Φ	0.00
Action samples n during evaluation	100
Action samples n during training	10
Learning rate	5×10^{-4}
Discount (γ)	0.95
Nonlinearity	ReLU
Target smoothing coefficient(τ)	0.005

Table 7: The hyperparameters of BCQ.

Hyperparameters	Value
α	0.02
Learning rate	5×10^{-4}
Action noise	0.2
Noise clip	0.5
Discount (γ)	0.95
Nonlinearity	ReLU
Target smoothing coefficient(τ)	0.005

Table 8: The hyperparameters of TD3+BC.

For Q-networks in BCQ [55] and TD3+BC [78], when using the PointNet + Transformer model, the action is concatenated with the point features and the state vector and fed into the model. The final feature from the model is fed into an MLP with layer sizes [256, 128, 1] to output Q-values. The

α	0.00	0.02	0.2	2.5
Success Rate	0.85	0.31	0.01	0.00

Table 9: The success rates of TD3+BC trained with different values of α on one environment of OpenCabinetDrawer and 300 demonstration trajectories. The algorithm becomes equivalent to BC if $\alpha = 0$.

VAE encoder and decoder in BCQ uses similar architecture, and the dimension of the latent vector z equals 2 times the action space dimension. The hyperparameters for BCQ and TD3+BC are shown in Table 7 and Table 8.

Note that in TD3+BC, $\pi = \operatorname{argmax}_{\pi} \mathbb{E}_{(s,a) \sim \mathcal{D}} [\lambda Q(s, \pi(s)) - (\pi(s) - a)^2]$, where $\lambda = \frac{\alpha}{\frac{1}{N} \sum_{(s_i, a_i)} |Q(s_i, a_i)|}$. In the original paper, $\alpha = 2.5$, and the algorithm is equivalent to BC if $\alpha = 0$. Interestingly, as shown in Table 9, we found that when α is non-zero, the performance of TD3+BC is always worse than BC, even when α has decreased 100 times from the value in the original paper. However, in our previously reported results, we used $\alpha = 0.02$ to illustrate the performance comparison between TD3+BC and BC, since setting $\alpha = 0$ is not interesting, and does not distinguish TD3+BC from BC.