

1135
1136
1137
1138
1139
1140
1141
1142
1143
1144
1145
1146
1147
1148
1149
1150
1151
1152
1153
1154
1155
1156
1157
1158
1159
1160
1161
1162
1163
1164
1165
1166
1167
1168
1169
1170

Table of Contents

A	MIKASA-Robo Implementation Details	27
B	MIKASA-Robo Datasets for Offline RL	28
C	MIKASA-Base Implementation Details	28
D	MIKASA-Robo setup for VLA baselines	28
E	Memory Mechanisms in RL	29
F	Classic baselines performance on the MIKASA-Robo benchmark	29
G	Experiments Reproducing and Compute Resources	32
H	MIKASA-Robo Detailed Tasks Description	32
H.1	ShellGame-v0	34
H.2	RememberColor-v0	35
H.3	RememberShape-v0	36
H.4	RememberShapeAndColor-v0	37
H.5	Intercept-v0	38
H.6	InterceptGrab-v0	39
H.7	RotateLenient-v0	40
H.8	RotateStrict-v0	41
H.9	TakeItBack-v0	42
H.10	SeqOfColors-v0	43
H.11	BunchOfColors-v0	44
H.12	ChainOfColors-v0	45
I	MIKASA-Base Benchmark Tasks Description	46
I.1	Memory Cards	46
I.2	Numpad	46
I.3	BSuite MemoryLength	47
I.4	Minigrid-Memory	47
I.5	Ballet	47
I.6	Passive T-Maze	47
I.7	ViZDoom-Two-Colors	47
I.8	Memory Maze	47
I.9	MemoryGym Mortar Mayhem	48
I.10	MemoryGym Mystery Path	48
I.11	POPGym environments	48

1172 A MIKASA-Robo Implementation Details

1173 An example of running the environment from the MIKASA-Robo benchmark is shown
 1174 in [Code 1](#). For ease of debugging, we also added various wrappers (found in
 1175 `mikasa_robo_suite/utils/wrappers/`) that display useful information about the episode
 1176 on the video ([Code 2](#)). Thus, `RenderStepInfoWrapper()` displays the current step in the envi-
 1177 ronment; `DebugRewardWrapper()` displays information about the full reward at the current step
 1178 in the environment; `DebugRewardWrapper()` displays information about each component that
 1179 generates the reward function at the current step. In addition, we also added task-specific wrappers
 1180 for each environment. For example, `RememberColorInfoWrapper()` displays the target color
 1181 of the cube in the `RememberColor-v0` task, and `ShellGameRenderCupInfoWrapper()`
 1182 displays which mug the ball is actually under in the `ShellGame-v0` task.

Code 1: Getting started with MIKASA-Robo using the `RememberColor9-v0` environment.

```
1183 # pip install mikasa_robo_suite
1184 import mikasa_robo_suite
1185 from mikasa_robo_suite.utils.wrappers import
1186     ↳ StateOnlyTensorToDictWrapper
1187 from tqdm.notebook import tqdm
1188 import torch
1189 import gymnasium as gym
1190
1191 # Create the environment via gym.make()
1192 # obs_mode="rgb" for modes "RGB", "RGB+joint", "RGB+oracle" etc.
1193 # obs_mode="state" for mode "state"
1194 episode_timeout = 90
1195 env = gym.make("RememberColor9-v0", num_envs=512 obs_mode="rgb",
1196     ↳ render_mode="all")
1197 env = StateOnlyTensorToDictWrapper(env) # * always use this wrapper!
1198
1199 obs, _ = env.reset(seed=42)
1200 print(obs.keys())
1201 for i in tqdm(range(episode_timeout)):
1202     action = torch.from_numpy(env.action_space.sample())
1203     obs, reward, terminated, truncated, info = env.step(action)
1204
1205 env.close()
```

Code 2: MIKASA-Robo wrappers system.

```
1208 import mikasa_robo_suite, torch
1209 from mikasa_robo_suite.dataset_collectors.get_mikasa_robo_datasets
1210     ↳ import env_info
1211 import gymnasium as gym
1212 from mani_skill.utils.wrappers import RecordEpisode
1213 from IPython.display import Video
1214
1215 env = gym.make("RememberColor9-v0", num_envs=512, obs_mode="rgb",
1216     ↳ render_mode="all")
1217 state_wrappers_list, episode_timeout = env_info("RememberColor9-v0")
1218 for wrapper_class, wrapper_kwargs in state_wrappers_list:
1219     env = wrapper_class(env, **wrapper_kwargs)
1220 env = RecordEpisode(env, f"./videos", max_steps_per_video=
1221     ↳ episode_timeout)
1222
1223 obs, _ = env.reset(seed=42)
1224 for i in range(episode_timeout):
1225     action = torch.from_numpy(env.action_space.sample())
1226     obs, reward, terminated, truncated, info = env.step(action)
1227
1228 Video(f"./videos/0.mp4", embed=True, width=640)
1229 env.close()
```

1232 B MIKASA-Robo Datasets for Offline RL

1233 To train Offline RL baselines on camera images (in “RGB” mode) with sparse rewards (success
1234 condition), we collected datasets for each of the 32 MIKASA-Robo tasks. Datasets were collected
1235 using a PPO-MLP agent trained to SR=100% in “state” mode (i.e., with full information about the
1236 task being solved) with sparse rewards (success condition). Thus, each dataset is represented by 1000
1237 successful trajectories, where each trajectory consists of:

- 1238 1. “rgb” (shape: $(T, 128, 128, 6)$) - two RGB images (view from above and from the gripper)
- 1239 2. “joints” (shape: $(T, 25)$) - Tool Center Point (TCP) position and rotation, and joint positions
1240 and velocities
- 1241 3. “action” (shape: $(T, 8)$) - action (8-dimensional vector)
- 1242 4. “reward” (shape: $(T,)$) - (dense) reward for each step
- 1243 5. “success” (shape: $(T,)$) - (sparse) success flag for each step
- 1244 6. “done” (shape: $(T,)$) - done flag for each step

1245 These datasets are available for download from the project website. We have also published the
1246 weights of the PPO-MLP agent used to collect the dataset, as well as scripts for collecting datasets of
1247 any size, to our repository.

1248 C MIKASA-Base Implementation Details

1249 An example of running an environment from the MIKASA-Base benchmark is shown in [Code 3](#).
1250 MIKASA-Base supports the standard Gymnasium API and is fully compatible with all its
1251 wrappers. This allows users to leverage various functionalities, including parallelization using
1252 AsyncVectorEnv. MIKASA-Base provides a predefined set of environments with different levels
1253 of difficulty. However, users can customize the environment parameters by passing specific arguments
1254 (see [Code 3](#)).

Code 3: Example code for running MemoryLength-v0 environment.

```
1255 import mikasa_base
1256 import gymnasium as gym
1257
1258
1259 # use pre-defined env
1260 # env_id = "MemoryLengthEasy-v0"
1261 # env_kwargs = None
1262
1263 # create env using custom parameters
1264 env_id = "MemoryLength-v0"
1265 env_kwargs = {"memory_length": 10, "num_bits": 1}
1266 seed = 123
1267
1268 env = gym.make(env_id, env_kwargs)
1269
1270 obs, _ = env.reset(seed=seed)
1271
1272 for i in range(11):
1273     action = env.action_space.sample()
1274     next_obs, reward, terminations, truncations, infos = env.step(
1275         ↪ action)
1276 env.close()
1277
```

1278 D MIKASA-Robo setup for VLA baselines

1279 For experiments involving Vision-Language-Action (VLA) models, we focused on a representative
1280 subset of spatial and object memory tasks from MIKASA-Robo. For each task, we generated a
1281 dataset of 250 episodes using an oracle PPO policy with full access to the environment state. At

Table 5: Tasks configurations for fine-tuning VLA models. The table lists the task ID, number of evaluation steps (T), and the associated language instruction

Task	T	Language instruction
RememberColor3/5/9-v0	60	Remember the color of the cube and then pick the matching one
ShellGameTouch-v0	90	Memorize the position of the cup covering the ball, then pick that cup
InterceptMedium-v0	90	Track the ball’s movement, estimate its velocity, then aim the ball at the target

every timestep, the policy recorded two synchronized RGB frames (one from the static “base” camera and one from the robot’s wrist camera) along with the corresponding end-effector control actions (`pd_ee_delta_pose` controller from [87]). Each task was also paired with a concise language instruction (see Table 5).

All VLA baselines were trained for 50000 iterations and evaluated independently on each task. Complete training/evaluation scripts, language instruction templates, and detailed model hyperparameter settings are provided in the accompanying supplementary code.

E Memory Mechanisms in RL

In RL, memory mechanisms are techniques or models used to enable agents to retain and recall information from past interactions with the environment.

There are several approaches to incorporating memory into RL, including recurrent neural networks (RNNs) [76, 37, 14] which uses hidden states to store information from previous steps [93, 34], state-space models (SSMs) [28, 83, 27] which uses system state to store historical information [31, 77], transformers [92] which uses attention mechanism to capture sequential dependencies inside the context window [72, 54, 68], graph neural networks (GNNs) [98] which uses graphs to store information [99, 44] etc. Popular agents with memory mechanisms are summarized in Table 2.

F Classic baselines performance on the MIKASA-Robo benchmark

In this section, we present a comprehensive evaluation of PPO-MLP and PPO-LSTM baselines on our MIKASA-Robo benchmark. Our experiments with PPO-MLP in `state` mode using dense rewards demonstrate perfect performance across all tasks, consistently achieving 100% success rate, as shown in Figure 7 and Figure 8. This remarkable performance serves as a crucial validation of our benchmark design: when an agent has access to complete state information and receives dense rewards, it can master these tasks completely. Therefore, any performance degradation in `RGB+joints` mode observed with other algorithms or training configurations must stem from the algorithmic limitations or learning challenges rather than any inherent flaws in the task design. This empirical evidence confirms that our environments are well-calibrated and properly designed, establishing a solid foundation for evaluating memory-enhanced algorithms. All results are presented as mean $\pm$ standard error of the mean (SEM), where the mean is computed across three independent training runs, and each trained agent is evaluated on 16 different random seeds to ensure robust performance assessment.

The performance evaluation of PPO-MLP and PPO-LSTM with dense rewards in the `RGB+joints` mode is presented in Figure 9. This mode specifically tests the agents’ memory capabilities, as it requires remembering and utilizing historical information to solve the tasks. Our results demonstrate a clear distinction between memory-less and memory-enhanced architectures, while also revealing the limitations of conventional memory mechanisms.

Consider the `RememberColor-v0` environment as an illustrative example. In its simplest configuration with three cubes, the memory-less PPO-MLP achieves only 25% success rate. In contrast, PPO-LSTM, leveraging its memory mechanism, achieves perfect performance with 100% success rate. However, as task complexity increases to five or nine cubes, even the LSTM’s memory capabilities prove insufficient, with performance degrading significantly.

These results validate two key aspects of our benchmark: first, its effectiveness in distinguishing between memory-less and memory-enhanced architectures, and second, its ability to challenge even sophisticated memory mechanisms as task complexity increases. This demonstrates that

1325 MIKASA-Robo provides a competitive yet meaningful evaluation framework for developing and
 1326 testing advanced memory-enhanced agents.

1327 Our evaluation of PPO-MLP and PPO-LSTM baselines under sparse reward conditions in
 1328 RGB+joints mode reveals the true challenge of our benchmark tasks. As shown in Figure 10,
 1329 both architectures – even the memory-enhanced LSTM – consistently fail to achieve any meaningful
 1330 success rate across nearly all considered environments. This striking result underscores the extreme
 1331 difficulty of memory-intensive manipulation tasks when only terminal rewards are available, high-
 1332 lighting the substantial gap between current algorithms and the level of memory capabilities required
 1333 for real-world robotic applications.

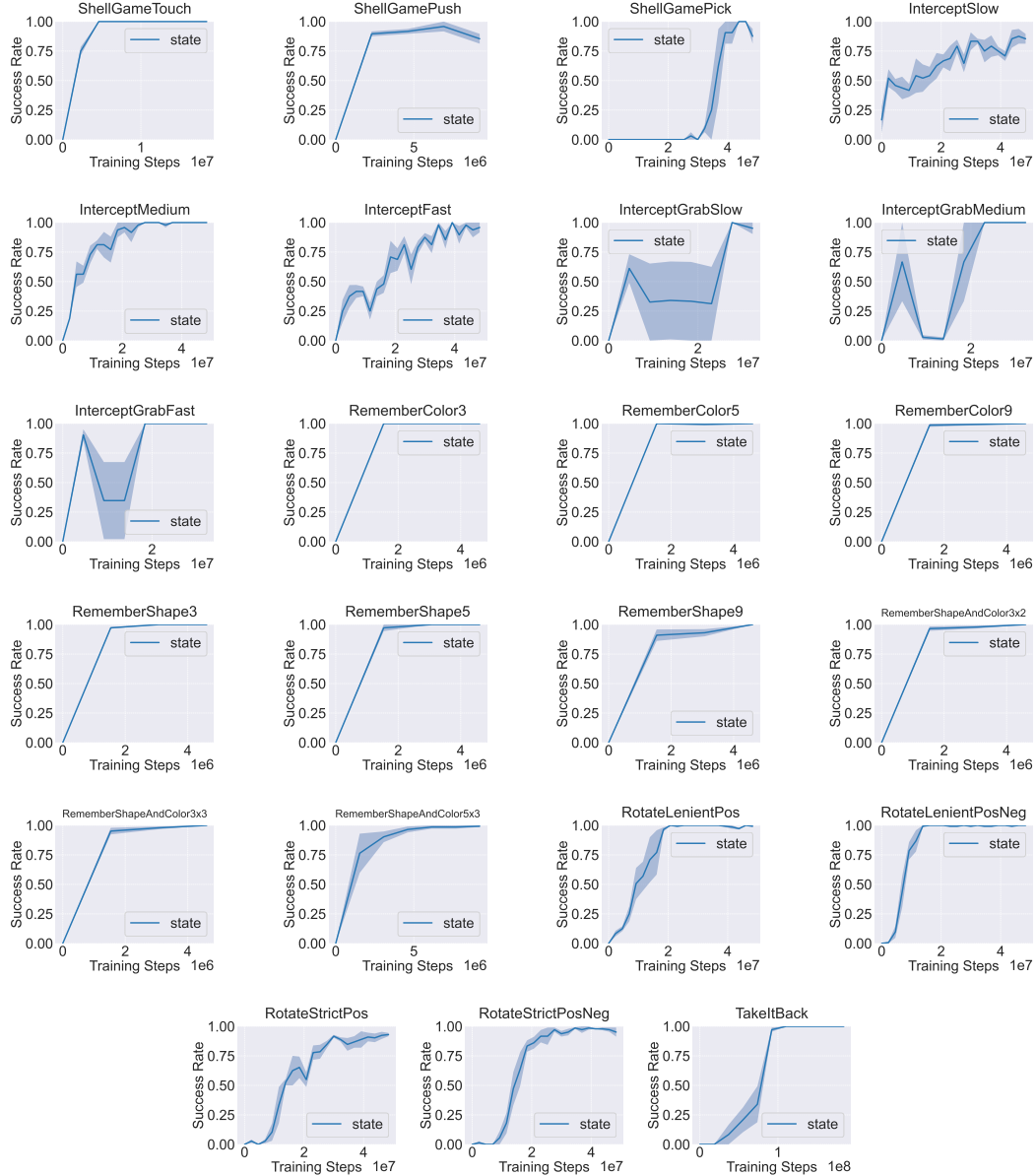


Figure 7: Demonstration of PPO-MLP performance on MIKASA-Robo benchmark when trained with oracle-level `state` information. In this learning mode, MDP problem formulation is considered, i.e. memory is not required for successful problem solving. At the same time, the obtained results show that it is possible to solve these problems and obtain 100% Success Rate.

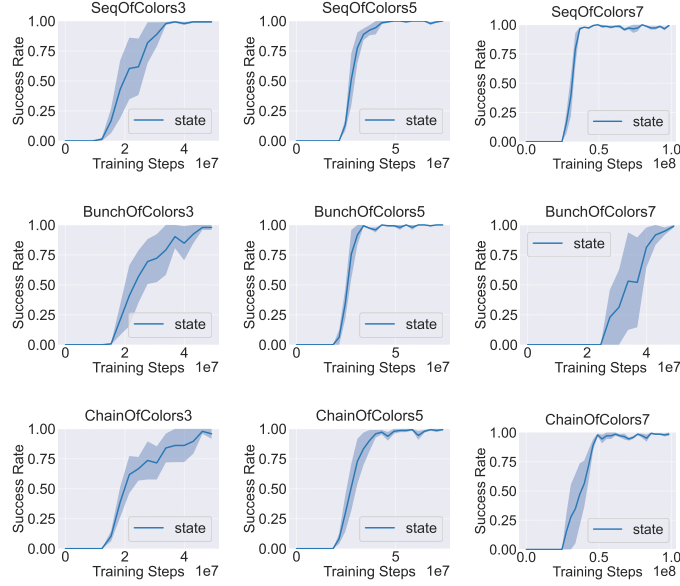


Figure 8: Demonstration of PPO-MLP performance on MUKASA-Robo benchmark when trained with oracle-level state information. Results are shown for memory capacity (SeqOfColors[3,5,7]-v0, BunchOfColors[3,5,7]-v0) and sequential memory (ChainOfColors[3,5,7]-v0).

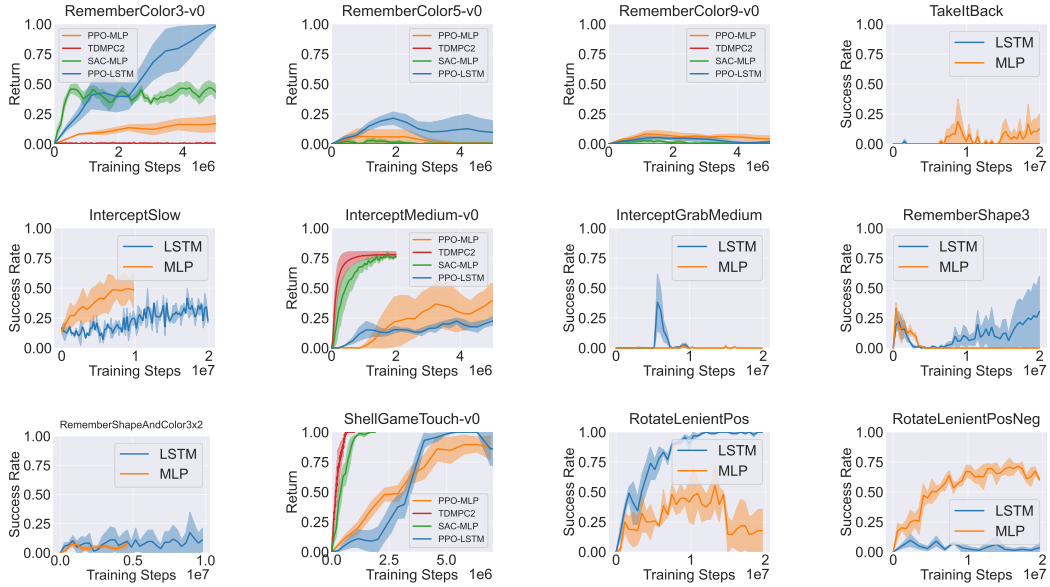


Figure 9: Performance evaluation of PPO-MLP and PPO-LSTM on the MUKASA-Robo benchmark using the “RGB+joints” training mode with dense reward function, where the agent only receives images from the camera (from above and from the gripper) and information about the state of the joints (position and velocity). The results demonstrate that numerous tasks pose significant challenges even for PPO-LSTM agents with memory, establishing these environments as effective benchmarks for evaluating advanced memory-enhanced architectures.

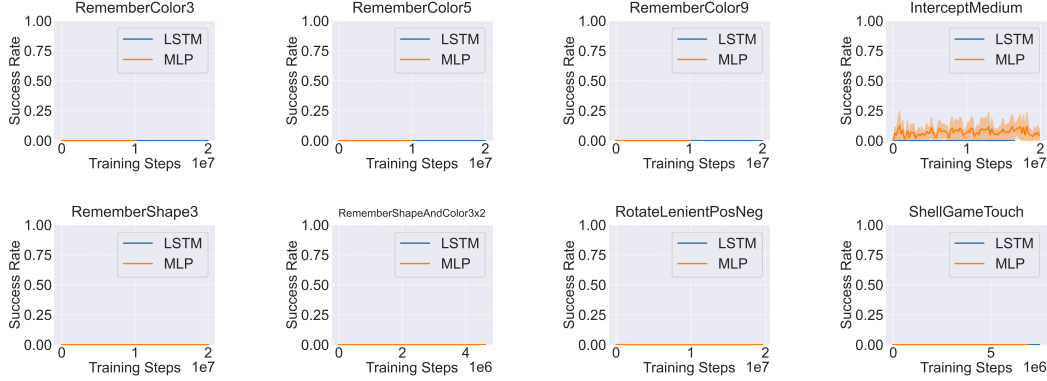


Figure 10: Performance evaluation of PPO-MLP and PPO-LSTM on the MIKASA-Robo benchmark using the “RGB+joints” with sparse reward function training mode, where the agent only receives images from the camera (from above and from the gripper) and information about the state of the joints (position and velocity). This training mode with sparse reward function causes even more difficulty for the agent to learn, making this mode even more challenging for memory-enhanced agents.

G Experiments Reproducing and Compute Resources

All baselines were trained and evaluated under a reproducible standardized setup. Training for every algorithm was performed on a single NVIDIA A100 GPU. For evaluation, each task was run for 100 independent episodes, with environment and agent random seeds ranging from 1 to 100. We report performance metrics as the mean success rate $\pm$ the standard error of the mean (SEM) over these 100 trials.

H MIKASA-Robo Detailed Tasks Description

In this section, we provide comprehensive descriptions of the 32 memory-intensive tasks that comprise the MIKASA-Robo benchmark. Each task is designed to evaluate specific aspects of memory capabilities in robotic manipulation, ranging from object tracking and spatial memory to sequential decision-making. For each task, we detail its objective, memory requirements, observation space, reward structure, and success criteria. Additionally, we explain how task complexity increases across different variants and discuss the specific memory challenges they present. The following subsections describe each task category and its variants in detail.

Each of the proposed environment supports multiple observation modes:

- **State:** Full state information including ball position
- **RGB+joints:** Two camera views (top-down and gripper) plus robot joint states
- **RGB:** Only visual information from two cameras

In the case of `RotateLenient-v0` and `RotateStrict-v0`, the prompt information available at each step is additionally added to each observation.

Table 6: **Results for Offline RL baselines.** The table shows comparison of transformer-based baselines (RATE, DT), behavior cloning (BC), classic Offline RL baselines (CQL), and Diffusion Policy (DP) on all 32 tasks from the MIKASA-Robo benchmark. Results are presented as mean $\pm$ sem across the three runs, where each run is averaged over 100 episodes and sem is the standard error of the mean. Training was performed using only RGB observations (two cameras: top view and gripper view) and using sparse rewards (success once condition). The results show that even models with memory (RATE, DT) are not able to solve most of the benchmark problems, which makes it challenging and promising for further validation of the algorithm.

#	Environment	RATE	DT	BC	CQL	DP
1	ShellGameTouch-v0	0.92 $\pm$ 0.01	0.53 $\pm$ 0.07	0.28 $\pm$ 0.01	0.16 $\pm$ 0.04	0.18 $\pm$ 0.02
2	ShellGamePush-v0	0.78 $\pm$ 0.06	0.62 $\pm$ 0.14	0.27 $\pm$ 0.01	0.25 $\pm$ 0.01	0.22 $\pm$ 0.03
3	ShellGamePick-v0	0.02 $\pm$ 0.01	0.00 $\pm$ 0.00	0.01 $\pm$ 0.01	0.00 $\pm$ 0.00	0.01 $\pm$ 0.00
4	InterceptSlow-v0	0.23 $\pm$ 0.02	0.40 $\pm$ 0.02	0.37 $\pm$ 0.06	0.25 $\pm$ 0.01	0.33 $\pm$ 0.05
5	InterceptMedium-v0	0.32 $\pm$ 0.02	0.56 $\pm$ 0.01	0.31 $\pm$ 0.14	0.03 $\pm$ 0.01	0.68 $\pm$ 0.02
6	InterceptFast-v0	0.30 $\pm$ 0.04	0.36 $\pm$ 0.04	0.03 $\pm$ 0.02	0.02 $\pm$ 0.02	0.21 $\pm$ 0.05
7	InterceptGrabSlow-v0	0.09 $\pm$ 0.03	0.00 $\pm$ 0.00	0.28 $\pm$ 0.18	0.03 $\pm$ 0.00	0.03 $\pm$ 0.01
8	InterceptGrabMedium-v0	0.09 $\pm$ 0.03	0.00 $\pm$ 0.00	0.11 $\pm$ 0.02	0.08 $\pm$ 0.04	0.03 $\pm$ 0.01
9	InterceptGrabFast-v0	0.14 $\pm$ 0.03	0.11 $\pm$ 0.03	0.09 $\pm$ 0.02	0.08 $\pm$ 0.03	0.18 $\pm$ 0.02
10	RotateLenientPos-v0	0.11 $\pm$ 0.04	0.01 $\pm$ 0.01	0.15 $\pm$ 0.03	0.16 $\pm$ 0.02	0.11 $\pm$ 0.02
11	RotateLenientPosNeg-v0	0.29 $\pm$ 0.03	0.05 $\pm$ 0.02	0.22 $\pm$ 0.01	0.12 $\pm$ 0.02	0.14 $\pm$ 0.05
12	RotateStrictPos-v0	0.03 $\pm$ 0.02	0.05 $\pm$ 0.04	0.01 $\pm$ 0.00	0.03 $\pm$ 0.01	0.06 $\pm$ 0.02
13	RotateStrictPosNeg-v0	0.08 $\pm$ 0.01	0.05 $\pm$ 0.03	0.04 $\pm$ 0.02	0.04 $\pm$ 0.02	0.15 $\pm$ 0.01
14	TakeItBack-v0	0.42 $\pm$ 0.24	0.08 $\pm$ 0.04	0.33 $\pm$ 0.10	0.04 $\pm$ 0.01	0.05 $\pm$ 0.02
15	RememberColor3-v0	0.65 $\pm$ 0.04	0.01 $\pm$ 0.01	0.27 $\pm$ 0.03	0.29 $\pm$ 0.01	0.32 $\pm$ 0.01
16	RememberColor5-v0	0.13 $\pm$ 0.03	0.07 $\pm$ 0.05	0.12 $\pm$ 0.01	0.15 $\pm$ 0.02	0.10 $\pm$ 0.02
17	RememberColor9-v0	0.09 $\pm$ 0.02	0.01 $\pm$ 0.01	0.12 $\pm$ 0.02	0.15 $\pm$ 0.01	0.17 $\pm$ 0.01
18	RememberShape3-v0	0.21 $\pm$ 0.04	0.05 $\pm$ 0.04	0.31 $\pm$ 0.04	0.20 $\pm$ 0.10	0.32 $\pm$ 0.05
19	RememberShape5-v0	0.17 $\pm$ 0.04	0.04 $\pm$ 0.04	0.18 $\pm$ 0.01	0.15 $\pm$ 0.00	0.21 $\pm$ 0.04
20	RememberShape9-v0	0.05 $\pm$ 0.00	0.05 $\pm$ 0.02	0.10 $\pm$ 0.02	0.14 $\pm$ 0.01	0.11 $\pm$ 0.02
21	RememberShapeAndColor3x2-v0	0.14 $\pm$ 0.02	0.04 $\pm$ 0.02	0.13 $\pm$ 0.02	0.11 $\pm$ 0.05	0.14 $\pm$ 0.02
22	RememberShapeAndColor3x3-v0	0.08 $\pm$ 0.03	0.06 $\pm$ 0.06	0.09 $\pm$ 0.02	0.09 $\pm$ 0.02	0.16 $\pm$ 0.01
23	RememberShapeAndColor5x3-v0	0.07 $\pm$ 0.02	0.01 $\pm$ 0.01	0.09 $\pm$ 0.01	0.09 $\pm$ 0.02	0.11 $\pm$ 0.03
24	BunchOfColors3-v0	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00
25	BunchOfColors5-v0	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00
26	BunchOfColors7-v0	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00
27	SeqOfColors3-v0	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00
28	SeqOfColors5-v0	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00
29	SeqOfColors7-v0	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00
30	ChainOfColors3-v0	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00
31	ChainOfColors5-v0	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00
32	ChainOfColors7-v0	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.00 $\pm$ 0.00

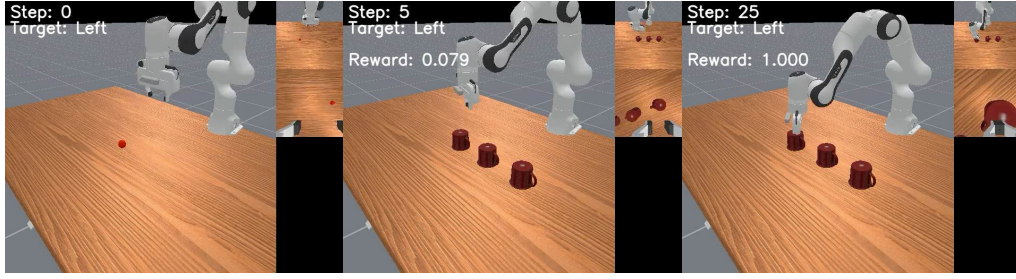


Figure 11: ShellGameTouch-v0: The robot observes a ball in front of it. next, this ball is covered by a mug and then the robot has to touch the mug with the ball underneath.

H.1 ShellGame-v0

The ShellGame-v0 task (Figure 11) is inspired by a simplified version of the classic shell game, which tests a person’s ability to remember object locations when they become occluded. This task evaluates an agent’s capacity for object permanence and spatial memory, crucial skills for real-world robotic manipulation where objects frequently become temporarily hidden from view.

Environment Description The environment consists of three identical mugs placed on a table and a red ball. The task proceeds in three phases:

1. **Observation Phase** (steps 0-4): The ball is placed at one of three positions, and the agent can observe its location.
2. **Occlusion Phase** (step 5): The ball and positions are covered by three identical mugs.
3. **Action Phase** (steps 6+): The agent must interact with the mug covering the ball’s location. The type of target interaction depends on the selected mode: Touch, Push and Pick.

Task Modes The task includes three variants of increasing difficulty:

- **Touch**: The agent only needs to touch the correct mug
- **Push**: The agent must push the correct mug to a designated area
- **Pick**: The agent must pick and lift the correct mug above a specified height

Success Criteria Success is determined by:

- **Touch**: Contact between the gripper and the correct mug
- **Push**: Moving forward the correct mug to the target zone
- **Pick**: Elevating the correct mug above 0.1m

Reward Structure The environment provides both sparse and dense reward variants:

- **Sparse**: Binary reward (1.0 for success, 0.0 otherwise)
- **Dense**: Continuous reward based on:
 - Distance between gripper and target mug
 - Robot’s motion smoothness (static reward based on joint velocities)
 - Task completion status (additional reward when the task is solved)

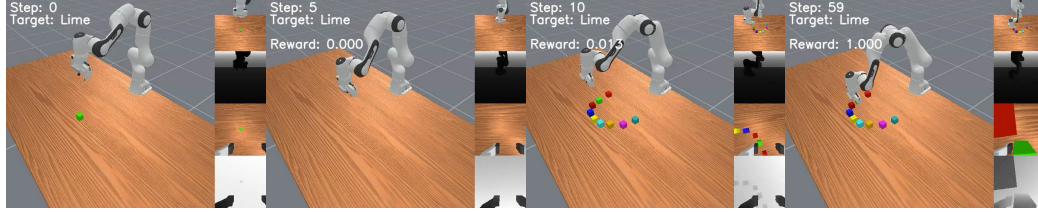


Figure 12: RememberColor9-v0: The robot observes a colored cube in front of it, then this cube disappears and an empty table is shown. Then 9 cubes appear on the table, and the agent must touch a cube of the same color as the one it observed at the beginning of the episode.

1381 H.2 RememberColor-v0

1382 The RememberColor-v0 task (Figure 12) tests an agent’s ability to remember and identify objects
 1383 based on their visual properties. This capability is essential for real-world robotics applications where
 1384 agents must recall and match object characteristics across time intervals.

1385 **Environment Description** The environment presents a sequence of colored cubes on a table. The
 1386 task proceeds in three phases:

- 1387 1. **Observation Phase** (steps 0-4): A cube of a specific color is displayed, and the agent must
 1388 memorize its color.
- 1389 2. **Delay Phase** (steps 5-9): The cube disappears, leaving an empty table.
- 1390 3. **Selection Phase** (steps 10+): Multiple cubes of different colors appear (3, 5, or 9 depending
 1391 on difficulty), and the agent must identify and interact with the cube matching the original
 1392 color.

1393 **Task Modes** The task includes three complexity levels:

- 1394 • 3 (easy): Choose from 3 different colors (red, lime, blue)
- 1395 • 5 (Medium): Choose from 5 different colors (red, lime, blue, yellow, magenta)
- 1396 • 9 (Hard): Choose from 9 different colors (red, lime, blue, yellow, magenta, cyan, maroon,
 1397 olive, teal)

1398 **Success Criteria** Success is determined by:

- 1399 • Correctly identifying and touching the cube that matches the color shown in the observation
 1400 phase
- 1401 • Maintaining contact with the correct cube for at least 0.1 seconds

1402 **Reward Structure** The environment provides both sparse and dense reward variants:

- 1403 • **Sparse:** Binary reward (1.0 for success, 0.0 otherwise)
- 1404 • **Dense:** Continuous reward based on:
 - 1405 – Distance between gripper and target cube
 - 1406 – Robot’s motion smoothness (static reward based on joint velocities)
 - 1407 – Additional reward for robot being static while touching the correct cube
 - 1408 – Task completion status (additional reward when the task is solved)

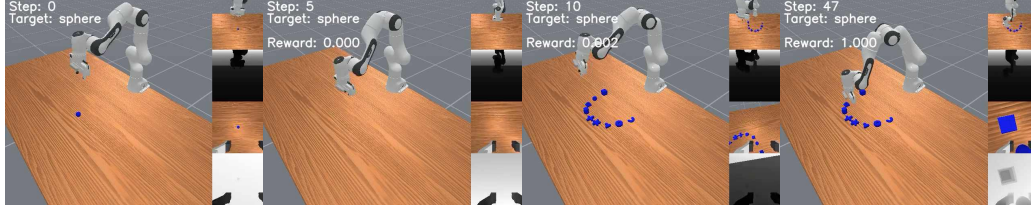


Figure 13: RememberShape9-v0: The robot observes an object with specific shape in front of it, then the object disappears and an empty table appears. Then 9 objects of different shapes appear on the table, and the agent must touch an object of the same shape as the one it observed at the beginning of the episode.

1409 H.3 RememberShape-v0

1410 The RememberShape-v0 task (Figure 13) evaluates an agent’s ability to remember and identify
 1411 objects based on their geometric properties. This capability is crucial for robotic applications where
 1412 shape recognition and recall are essential for successful manipulation.

1413 **Environment Description** The environment presents a sequence of geometric shapes on a table.
 1414 The task proceeds in three phases:

- 1415 1. **Observation Phase** (steps 0-4): A shape (cube, sphere, cylinder, etc.) is displayed, and the
 1416 agent must memorize its geometry.
- 1417 2. **Delay Phase** (steps 5-9): The shape disappears, leaving an empty table.
- 1418 3. **Selection Phase** (steps 10+): Multiple shapes appear (3, 5, or 9 depending on difficulty),
 1419 and the agent must identify and interact with the shape matching the original geometry.

1420 **Task Modes** The task includes three complexity levels:

- 1421 • 3 (Easy): Choose from 3 different shapes (cube, sphere, cylinder)
- 1422 • 5 (Medium): Choose from 5 different shapes (cube, sphere, cylinder cross, torus)
- 1423 • 9 (Hard): Choose from 9 different shapes (cube, sphere, cylinder cross, torus, star, pyramid,
 1424 t-shape, crescent)

1425 **Success Criteria** Success is determined by:

- 1426 • Correctly identifying and touching the object with the same shape shown in the observation
 1427 phase
- 1428 • Maintaining contact with the correct shape for at least 0.1 seconds

1429 **Reward Structure** The environment provides both sparse and dense reward variants:

- 1430 • **Sparse:** Binary reward (1.0 for success, 0.0 otherwise)
- 1431 • **Dense:** Continuous reward based on:
 - 1432 – Distance between gripper and target object
 - 1433 – Robot’s motion smoothness (static reward based on joint velocities)
 - 1434 – Additional reward for maintaining static position when touching correct object
 - 1435 – Task completion status (additional reward when the task is solved)

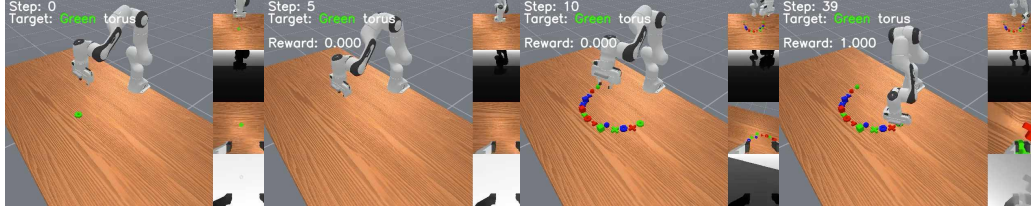


Figure 14: RememberShapeAndColor5x3-v0: An object of a certain shape and color appears in front of the agent. Then the object disappears and the agent sees an empty table. Then objects of 5 different shapes and 3 different colors appear on the table and the agent has to touch what it observed at the beginning of the episode.

H.4 RememberShapeAndColor-v0

The RememberShapeAndColor-v0 task (Figure 14) evaluates an agent’s ability to remember and identify objects based on multiple visual properties simultaneously. This task combines shape and color recognition, testing the agent’s capacity to maintain and match multiple object features across time intervals.

Environment Description The environment presents a sequence of colored geometric shapes on a table. The task proceeds in three phases:

1. **Observation Phase** (steps 0-4): An object with specific shape and color is displayed, and the agent must memorize both properties.
2. **Delay Phase** (steps 5-9): The object disappears, leaving an empty table.
3. **Selection Phase** (steps 10+): Multiple objects with different combinations of shapes and colors appear, and the agent must identify and interact with the object matching both the original shape and color.

Task Modes The task includes three complexity levels based on the number of shape-color combinations:

- 3x2 (Easy): Choose from 6 objects (3 shapes x 2 colors); shapes: cube, sphere, t-shape; colors: red, green
- 3x3 (Medium): Choose from 9 objects (3 shapes x 3 colors); shapes: cube, sphere, t-shape; colors: red, green, blue
- 5x3 (Hard): Choose from 15 objects (5 shapes x 3 colors); shapes: cube, sphere, t-shape, cross, torus; colors: red, green, blue

Success Criteria Success is determined by:

- Correctly identifying and touching the object that matches both the shape and color shown in the observation phase
- Maintaining contact with the correct object for at least 0.1 seconds

Reward Structure The environment provides both sparse and dense reward variants:

- **Sparse:** Binary reward (1.0 for success, 0.0 otherwise)
- **Dense:** Continuous reward based on:
 - Distance between gripper and target object
 - Robot’s motion smoothness (static reward based on joint velocities)
 - Additional reward for maintaining static position while touching correct object
 - Task completion status (additional reward when the task is solved)

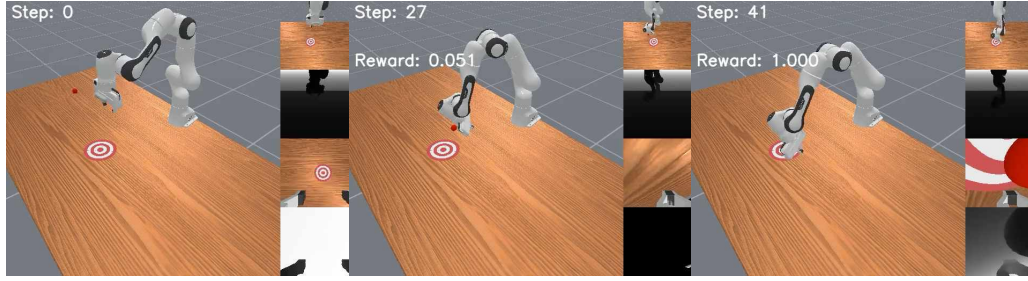


Figure 15: InterceptMedium-v0: A ball rolls on the table in front of the agent with a random initial velocity, and the agent’s task is to intercept this ball and direct it at the target zone.

H.5 Intercept-v0

The Intercept-v0 task (Figure 16) evaluates an agent’s ability to predict and intercept a moving object based on its initial trajectory. This task tests the agent’s capacity for motion prediction and spatial-temporal reasoning, which are essential skills for dynamic manipulation tasks in robotics.

Environment Description The environment consists of a red ball moving across a table and a target zone. The task requires the agent to:

1. Observe the ball’s initial position and velocity
2. Predict the ball’s trajectory
3. Guide the ball to reach a designated target zone

Task Modes The task includes three variants with increasing ball velocities:

- Slow: Ball velocity range of 0.25-0.5 m/s
- Medium: Ball velocity range of 0.5-0.75 m/s
- Fast: Ball velocity range of 0.75-1.0 m/s

Success Criteria Success is determined by:

- Guiding the ball to enter the target zone
- The ball must come to rest within the target area (radius 0.1m)

Reward Structure The environment provides both sparse and dense reward variants:

- **Sparse:** Binary reward (1.0 for success, 0.0 otherwise)
- **Dense:** Continuous reward based on:
 - Distance between gripper and ball
 - Distance between ball and target zone
 - Robot’s motion smoothness (static reward based on joint velocities)
 - Task completion status (additional reward when the task is solved)

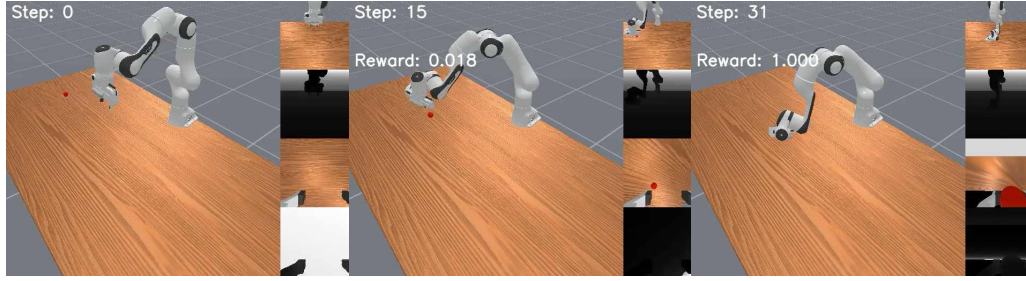


Figure 16: InterceptGrabMedium-v0: A ball rolls on the table in front of the agent with a random initial velocity, and the agent’s task is to intercept this ball with a gripper and lift it up.

H.6 InterceptGrab-v0

The InterceptGrab-v0 task (Figure 16) extends the Intercept-v0 task by requiring the agent to not only predict the trajectory of a moving object but also grasp it while in motion. This task evaluates the agent’s ability to combine motion prediction with precise manipulation timing, simulating real-world scenarios where robots must catch or intercept moving objects.

Environment Description The environment consists of a red ball moving across a table. The task requires the agent to:

1. Observe the ball’s initial position and velocity
2. Predict the ball’s trajectory
3. Position the gripper to intercept the ball’s path
4. Time the grasping action correctly to catch the ball
5. Maintain a stable grasp while bringing the ball to rest

Task Modes The task includes three variants with increasing ball velocities:

- Slow: Ball velocity range of 0.25-0.5 m/s
- Medium: Ball velocity range of 0.5-0.75 m/s
- Fast: Ball velocity range of 0.75-1.0 m/s

Success Criteria Success is determined by:

- Successfully grasping the moving ball
- Maintaining a stable grasp until the ball comes to rest
- The robot must be static with the ball firmly grasped

Reward Structure The environment provides both sparse and dense reward variants:

- **Sparse:** Binary reward (1.0 for success, 0.0 otherwise)
- **Dense:** Continuous reward based on:
 - Distance between gripper and ball
 - Grasping reward
 - Robot’s motion smoothness (static reward based on joint velocities)
 - Task completion status (additional reward when the task is solved)



Figure 17: RotateLenientPos-v0: A randomly oriented peg is placed in front of the agent. The agent’s task is to rotate this peg by a certain angle (the center of the peg can be shifted).

1518 H.7 RotateLenient-v0

1519 The RotateLenient-v0 task (Figure 17) evaluates an agent’s ability to remember and execute
 1520 specific rotational movements. This task tests the agent’s capacity to maintain and reproduce angular
 1521 information, which is crucial for manipulation tasks requiring precise orientation control. This task
 1522 tests the agent’s ability to hold information in memory about how far peg has already rotated at the
 1523 current step relative to its initial position.

1524 **Environment Description** The environment consists of a blue-colored peg on a table that must be
 1525 rotated by a specified angle. The task proceeds in one phase, but the static prompt information about
 1526 the target angle is available to the agent at each timestep:

1527 1. **Action Phase:** The agent must rotate the peg to match the target angle

1528 **Task Modes** The task includes two variants with different rotation requirements:

- 1529 • **Pos:** Rotate by a positive angle between 0 and $\pi/2$
- 1530 • **PosNeg:** Rotate by either positive or negative angle between $-\pi/4$ and $\pi/4$

1531 **Success Criteria** Success is determined by:

- 1532 • Rotating the peg to within the angle threshold (± 0.1 radians) of the target angle
- 1533 • Maintaining the final orientation in a stable position
- 1534 • The robot must be static with the peg at the correct orientation

1535 **Reward Structure** The environment provides both sparse and dense reward variants:

- 1536 • **Sparse:** Binary reward (1.0 for success, 0.0 otherwise)
- 1537 • **Dense:** Continuous reward based on:
 - 1538 – Distance between gripper and peg
 - 1539 – Angular distance to target rotation
 - 1540 – Stability of the peg’s orientation
 - 1541 – Robot’s motion smoothness (static reward based on joint velocities)
 - 1542 – Task completion status (additional reward when the task is solved)

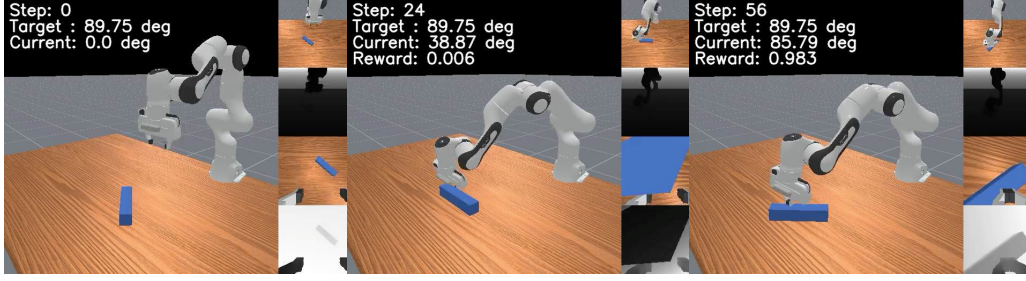


Figure 18: RotateStrictPos-v0: A randomly oriented peg is placed in front of the agent. The agent’s task is to rotate this peg by a certain angle (it is not allowed to move the center of the peg)

H.8 RotateStrict-v0

The RotateStrict-v0 task (Figure 18) extends the RotateLenient-v0 task with more stringent requirements for precise rotational control.

Environment Description The environment consists of a blue-colored peg on a table that must be rotated by a specified angle while maintaining its position. The task proceeds in one phase, but the static prompt information about the target angle is available to the agent at each timestep:

1. **Action Phase:** The agent must rotate the peg to match the target angle while keeping it centered

Task Modes The task includes two variants with different rotation requirements:

- **Pos:** Rotate by a positive angle between 0 and $\pi/2$
- **PosNeg:** Rotate by either positive or negative angle between $-\pi/4$ and $\pi/4$

Success Criteria Success is determined by:

- Rotating the peg to within the angle threshold (± 0.1 radians) of the target angle
- Maintaining the peg’s position within 5cm of its initial XY coordinates
- The robot must be static with the peg at the correct orientation
- No significant deviation in other rotation axes

Reward Structure The environment provides both sparse and dense reward variants:

- **Sparse:** Binary reward (1.0 for success, 0.0 otherwise)
- **Dense:** Continuous reward based on:
 - Distance between gripper and peg
 - Angular distance to target rotation
 - Position deviation from initial location
 - Stability of the peg’s orientation
 - Robot’s motion smoothness (static reward based on joint velocities)
 - Task completion status (additional reward when the task is solved)

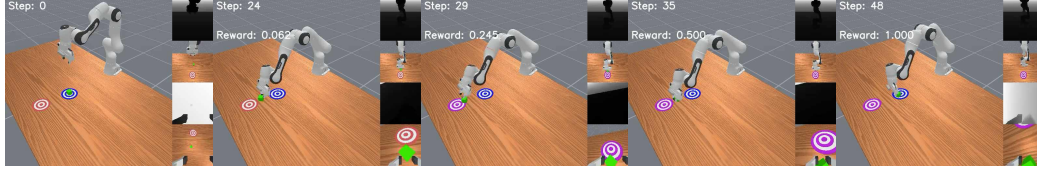


Figure 19: TakeItBack-v0: The agent observes a green cube in front of him. The agent’s task is to move the green cube to the red target, and as soon as it lights up violet, return the cube to its original position (the agent does not observe the original position of the cube).

1568 H.9 TakeItBack-v0

1569 The TakeItBack-v0 task (Figure 19) assesses the agent’s ability to perform sequential tasks and
 1570 memorize the starting position. This task tests the agent’s capacity for sequential memory and spatial
 1571 reasoning, requiring it to maintain information about past locations and achievements while executing
 1572 a multi-step plan.

1573 **Environment Description** The environment consists of a green cube and two target regions (initial
 1574 and goal) on a table. The task proceeds in two phases:

- 1575 1. **First Phase:** The agent must move the cube from its initial position to a goal region
- 1576 2. **Second Phase:** After reaching the goal, goal region change it’s color from red to magenta,
 1577 and the agent must return the cube to its original position (marked by the initial region and
 1578 invisible for the agent)

1579 **Success Criteria** Success is determined by:

- 1580 • First reaching the goal region with the cube
- 1581 • Then returning the cube to the initial region
- 1582 • Both goals must be achieved in sequence

1583 **Reward Structure** The environment provides both sparse and dense reward variants:

- 1584 • **Sparse:** Binary reward (1.0 for success, 0.0 otherwise)
- 1585 • **Dense:** Continuous reward based on:
 - 1586 – Distance between gripper and cube
 - 1587 – Distance to current target region
 - 1588 – Progress through the task sequence
 - 1589 – Stability of cube manipulation
 - 1590 – Robot’s motion smoothness (static reward based on joint velocities)
 - 1591 – Task completion status (additional reward when the task is solved)

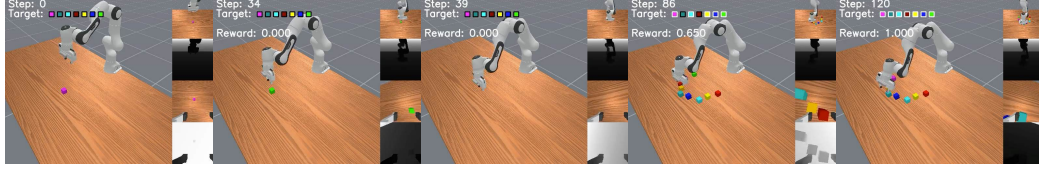


Figure 20: SeqOfColors7-v0: In front of the agent, 7 cubes of different colors appear sequentially. After the last cube is shown, the agent observes an empty table. Then 9 cubes of different colors appear on the table and the agent has to touch the cubes that were shown at the beginning of the episode in any order.

1592 H.10 SeqOfColors-v0

1593 The SeqOfColors-v0 task (Figure 20) evaluates an agent’s ability to remember and reproduce an
 1594 unordered sequence of colors. This task tests memory capacity capabilities essential for robotic tasks
 1595 that require following specific patterns or sequences.

1596 **Environment Description** The environment presents a sequence of colored cubes that must be
 1597 reproduced in any order. The task proceeds in two phases:

- 1598 1. **Observation Phase** (steps $0-(5N - 1)$): A sequence of N colored cubes is shown one at a
 1599 time, with each cube visible for 5 steps.
- 1600 2. **Delay Phase** (steps $(5N)-(5N + 4)$): All cubes disappear
- 1601 3. **Selection Phase** (steps $(5N + 5)+$): A larger set of cubes appears, and the agent must
 1602 identify and touch all previously shown cubes in any order

1603 **Task Modes** The task includes three complexity levels:

- 1604 • 3 (Easy): Remember 3 colors demonstrated sequentially
- 1605 • 5 (Medium): Remember 5 colors demonstrated sequentially
- 1606 • 7 (Hard): Remember 7 colors demonstrated sequentially

1607 **Success Criteria** Success is determined by:

- 1608 • Correctly identifying and touching all cubes from the observation phase
- 1609 • Order of selection doesn’t matter
- 1610 • Each cube must be touched for at least 0.1 seconds
- 1611 • The demonstrated set must be touched without any mistakes

1612 **Reward Structure** The environment provides both sparse and dense reward variants:

- 1613 • **Sparse:** Binary reward (1.0 for success, 0.0 otherwise)
- 1614 • **Dense:** Continuous reward based on:
 - 1615 – Distance between gripper and next target cube
 - 1616 – Number of correctly identified cubes
 - 1617 – Static reward for stable contact
 - 1618 – Robot’s motion smoothness (static reward based on joint velocities)
 - 1619 – Task completion status (additional reward when the task is solved)

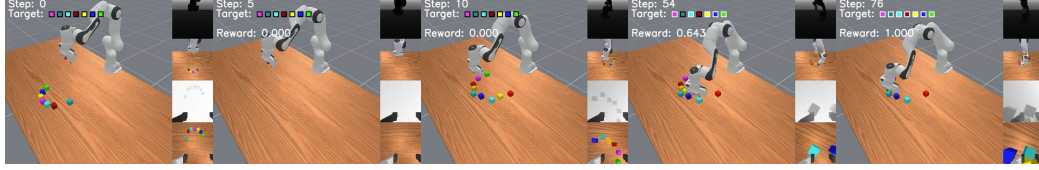


Figure 21: BunchOfColors7-v0: 7 cubes of different colors appear simultaneously in front of the agent. After the agent observes an empty table. Then, 9 cubes of different colors appear on the table and the agent has to touch the cubes that were shown at the beginning of the episode in any order.

1620 H.11 BunchOfColors-v0

1621 The BunchOfColors-v0 task (Figure 21) tests an agent’s memory capacity by requiring it to
 1622 remember multiple objects simultaneously. This capability is crucial for tasks requiring parallel
 1623 processing of multiple items.

1624 **Environment Description** The environment presents multiple colored cubes simultaneously. The
 1625 task proceeds in three phases:

- 1626 1. **Observation Phase** (steps 0-4): Multiple colored cubes are displayed simultaneously
- 1627 2. **Delay Phase** (steps 5-9): All cubes disappear
- 1628 3. **Selection Phase** (steps 10+): A larger set of cubes appears, and the agent must identify and
 1629 touch all previously shown cubes in any order

1630 **Task Modes** The task includes three complexity levels:

- 1631 • 3 (Easy): Remember 3 colors demonstrated simultaneously
- 1632 • 5 (Medium): Remember 5 colors demonstrated simultaneously
- 1633 • 7 (Hard): Remember 7 colors demonstrated simultaneously

1634 **Success Criteria** Success is determined by:

- 1635 • Correctly identifying and touching all cubes from the observation phase
- 1636 • Order of selection doesn’t matter
- 1637 • Each cube must be touched for at least 0.1 seconds
- 1638 • The demonstrated set must be touched without any mistakes

1639 **Reward Structure** The environment provides both sparse and dense reward variants:

- 1640 • **Sparse:** Binary reward (1.0 for success, 0.0 otherwise)
- 1641 • **Dense:** Continuous reward based on:
 - 1642 – Distance between gripper and next target cube
 - 1643 – Static reward for stable contact
 - 1644 – Number of correctly touched cubes
 - 1645 – Robot’s motion smoothness (static reward based on joint velocities)
 - 1646 – Task completion status (additional reward when the task is solved)



Figure 22: ChainOfColors7-v0: In front of the agent, 7 cubes of different colors appear sequentially. After the last cube is shown, the agent sees an empty table. Then 9 cubes of different colors appear on the table and the agent must unmistakably touch the cubes that were shown at the beginning of the episode, in the same strict order.

H.12 ChainOfColors-v0

The ChainOfColors-v0 task (Figure 22) evaluates the agent’s ability to store and retrieve ordered information. This task simulates scenarios where the agent must track changing relationships between objects over time.

Environment Description The environment presents an ordered sequence (chain) of colored cubes that must be followed. The task proceeds in multiple phases:

1. **Observation Phase** (steps $0-(5N - 1)$): A sequence of N colored cubes is shown one at a time, with each cube visible for 5 steps.
2. **Delay Phase** (steps $(5N)-(5N + 4)$): All cubes disappear
3. **Selection Phase** (steps $(5N + 5)+$): A larger set of cubes appears, and the agent must identify and touch all previously shown cubes in the exact order as demonstrated

Task Modes The task includes three complexity levels:

- 3 (Easy): Remember 3 colors demonstrated sequentially
- 5 (Medium): Remember 5 colors demonstrated sequentially
- 7 (Hard): Remember 7 colors demonstrated sequentially

Success Criteria Success is determined by:

- Correctly identifying and touching all cubes from the observation phase in the exact order
- Each cube must be touched for at least 0.1 seconds
- The demonstrated set must be touched without any mistakes

Reward Structure The environment provides both sparse and dense reward variants:

- **Sparse:** Binary reward (1.0 for success, 0.0 otherwise)
- **Dense:** Continuous reward based on:
 - Distance between gripper and next target cube
 - Static reward for stable contact
 - Number of correctly touched cubes
 - Robot’s motion smoothness (static reward based on joint velocities)
 - Task completion status (additional reward when the task is solved)

Table 7: Classification of environments from the MIKASA-Base benchmark according to the suggested memory-intensive tasks classification from the [Subsection 4.2](#).

Environment	Memory Task	Brief description of the task	Observation Space	Action Space
Memory Cards	Capacity	Memorize the positions of revealed cards and correctly match pairs while minimizing incorrect guesses.	vector	discrete
Numpad	Sequential	Memorize the sequence of movements and navigate the rolling ball on a 3x3 grid by following the correct order while avoiding mistakes.	image, vector	discrete, continuous
BSuite Memory Length	Object	Memorize the initial context signal and recall it after a given number of steps to take the correct action.	vector	discrete
Minigrid-Memory	Object	Memorize the object in the starting room and use this information to select the correct path at the junction.	image	discrete
Ballet	Sequential, Object	Memorize the sequence of movements performed by each uniquely colored and shaped dancer, then identify and approach the dancer who executed the given pattern.	image	discrete
Passive Visual Match	Object	Memorize the target color displayed on the wall during the initial phase. After a brief distractor phase, identify and select the target color among the distractors by stepping on the corresponding ground pad.	image	discrete
Passive-T-Maze	Object	Memorize the goal's location upon initial observation, navigate through the maze with limited sensory input, and select the correct path at the junction.	vector	discrete
VizDoom-two-colors	Object	Memorize the color of the briefly appearing pillar (green or red) and collect items of the same color to survive in the acid-filled room.	image	discrete
Memory Maze	Spatial	Memorize the locations of objects and the maze structure using visual clues, then navigate efficiently to find objects of a specific color and score points.	image	discrete
MemoryGym Mortar Mayhem	Capacity, Sequential	Memorize a sequence of movement commands and execute them in the correct order.	image	discrete
MemoryGym Mystery Path	Capacity, Spatial	Memorize the invisible path and navigate it without stepping off.	image	discrete
POPGym Repeat First	Object	Memorize the initial value presented at the first step and recall it correctly after receiving a sequence of random values.	vector	discrete
POPGym Repeat Previous	Sequential, Object	Memorize the value observed at each step and recall the value from k steps earlier when required.	vector	discrete
POPGym Autoencode	Sequential	Memorize the sequence of cards presented at the beginning and reproduce them in the same order when required.	vector	discrete
POPGym Count Recall	Object, Capacity	Memorize unique values encountered and count how many times a specific value has appeared.	vector	discrete
POPGym vectorless Cartpole	Sequential	Memorize velocity data over time and integrate it to infer the position of the pole for balance control.	vector	continuous
POPGym vectorless Pendulum	Sequential	Memorize angular velocity over time and integrate it to infer the pendulum's position for successful swing-up control.	vector	continuous
POPGym Multiarmed Bandit	Object, Capacity	Memorize the reward probabilities of different slot machines by exploring them and identify the one with the highest expected reward.	vector	discrete
POPGym Concentration	Capacity	Memorize the positions of revealed cards and match them with previously seen cards to find all matching pairs.	vector	discrete
POPGym Battleship	Spatial	Memorize the coordinates of previous shots and their HIT or MISS feedback to build an internal representation of the board, avoid repeat shots, and strategically target ships for maximum rewards.	vector	discrete
POPGym Mine Sweeper	Spatial	Memorize revealed grid information and use numerical clues to infer safe tiles while avoiding mines.	vector	discrete
POPGym Labyrinth Explore	Spatial	Memorize previously visited cells and navigate the maze efficiently to discover new, unexplored areas and maximize rewards.	vector	discrete
POPGym Labyrinth Escape	Spatial	Memorize the maze layout while exploring and navigate efficiently to find the exit and receive a reward.	vector	discrete
POPGym Higher Lower	Object, Sequential	Memorize previously revealed card ranks and predict whether the next card will be higher or lower, updating the reference card after each prediction to maximize rewards.	vector	discrete

I MIKASA-Base Benchmark Tasks Description

This section provides a detailed description of all environments included in the MIKASA-Base benchmark [Section 5](#). Understanding the characteristics and challenges of these environments is crucial for evaluating RL algorithms. Each environment presents unique tasks, offering diverse scenarios to test the memory abilities of RL agents.

I.1 Memory Cards

The Memory Cards environment [19] is a memory game environment with 5 randomly shuffled pairs of hidden cards. At each step, the agent sees one revealed card and must find its matching pair. A correct guess removes both cards; otherwise, the card is hidden again, and a new one is revealed. The game continues until all pairs are removed.

I.2 Numpad

The Numpad environment [38] consists of an $N \times N$ grid of tiles. The agent controls a ball that rolls between tiles. At the beginning of an episode, a random sequence of n neighboring tiles (excluding diagonals) is selected, and the agent must follow this sequence in the correct order. The environment is structured so that pressing the correct tile lights it up, while pressing an incorrect tile resets progress. A reward of +1 is given for the first press of each correct tile after a reset. The episode ends after a fixed number of steps. To succeed, the agent must memorize the sequence and navigate it correctly without mistakes. The ability to “jump” over tiles is not available.

1692 I.3 BSuite MemoryLength

1693 The MemoryLength environment [70] represents a sequence of observations, where at each step, the
1694 observation takes a value of either +1 or -1. The environment is structured so that a reward is given
1695 only at the final step if the agent correctly predicts the i -th value from the initial observation vector
1696 *obs*. The index of this i -th value is specified at the last step observation vector in *obs*[1]. To succeed,
1697 the agent must remember the sequence of observations and use this information to make an accurate
1698 prediction at the final step.

1699 I.4 Minigrid-Memory

1700 Minigrid-Memory [12] is a two-dimensional grid-based environment that features a T-shaped maze
1701 with a small room at the beginning of the corridor, containing an object. The agent starts at a random
1702 position within the corridor. Its task is to reach the room, observe and memorize the object, then
1703 proceed to the junction at the maze’s end and turn towards the direction where an identical object is
1704 located. The reward function is defined as $R_t = 1 - 0.9 \times \frac{t}{T}$ for a successful attempt; otherwise, the
1705 agent receives zero reward. The episode terminates when the agent makes a choice at the junction or
1706 exceeds a time limit of steps.

1707 I.5 Ballet

1708 In the Ballet environment [54] tasks take place in an 11×11 tiled room, consisting of a 9×9 central
1709 area surrounded by a one-tile-wide wall. Each tile is upsampled to 9 pixels, resulting in a 99×99
1710 pixel input image. The agent is initially placed at the center of the room, while dancers are randomly
1711 positioned in one of 8 possible locations around it. Each dancer has a distinct shape and color,
1712 selected from 15 possible shapes and 19 colors, ensuring uniqueness. These visual features serve
1713 only for identification and do not influence behavior. The agent itself is always represented as a white
1714 square. The agent receives egocentric visual observations, meaning its view is centered on its own
1715 position, which has been shown to enhance generalization.

1716 I.6 Passive T-Maze

1717 The Passive T-Maze environment [68] consists of a corridor leading to a junction that connects two
1718 possible goal states. The agent starts at a designated position and can move in four directions: left,
1719 right, up, or down. At the beginning of each episode, one of the two goal states is randomly assigned
1720 as the correct destination. The agent observes this goal location before starting its movement. The
1721 agent stays in place if it attempts to move into a wall. To succeed, the agent must navigate to the
1722 correct goal based on its initial observation. The optimal strategy involves moving through the
1723 corridor towards the junction and then selecting the correct path.

1724 I.7 ViZDoom-Two-Colors

1725 The ViZDoom-Two-Colors [84] is an environment where an agent is placed in a room with constantly
1726 depleting health. The room contains red and green objects, one of which restores health (+1 reward),
1727 while the other reduces it (-1 reward). The beneficial color is randomly assigned at the beginning
1728 of each episode and indicated by a column. The environment is structured so that the agent must
1729 memorize the column’s color to collect the correct items. Initially, the column remains visible, but in
1730 a harder variant, it disappears after 45 steps, increasing the memory requirement. To succeed, the
1731 agent must maximize survival by collecting beneficial objects while avoiding harmful ones.

1732 I.8 Memory Maze

1733 The Memory Maze environment [73] is a procedurally generated 3D maze. Each episode, the agent
1734 spawns in a new maze with multiple colored objects placed in fixed locations. The agent receives a
1735 first-person view and a prompt indicating the color of the target object. It must navigate the maze,
1736 memorize object positions, and return to them efficiently. The agent receives a reward of 1 for
1737 reaching the correct object and no reward for incorrect objects.

1738 **I.9 MemoryGym Mortar Mayhem**

1739 Mortar Mayhem [75] is a grid-based environment where the agent must memorize and execute a
1740 sequence of commands in the correct order. The environment consists of a finite grid, where the agent
1741 initially observes a series of movement instructions and then attempts to reproduce them accurately.
1742 Commands include movements to adjacent tiles or remaining in place. The challenge lies in the
1743 agent's ability to recall and execute a growing sequence of instructions, with failure resulting in
1744 episode termination. A reward of +0.1 is given for each correctly executed command

1745 **I.10 MemoryGym Mystery Path**

1746 Mystery Path [75] presents an invisible pathway that the agent must traverse without deviating. If
1747 the agent steps off the path, it is returned to the starting position, forcing it to remember the correct
1748 trajectory. The path is procedurally generated, meaning each episode introduces a new configuration.
1749 Success in this environment requires the agent to accurately recall previous steps and missteps to
1750 avoid repeating errors. The agent is rewarded +0.1 for progressing onto a previously unvisited tile

1751 **I.11 POPGym environments**

1752 The following environments are included from the POPGym benchmark [65], which is designed
1753 to evaluate RL agents in partially observable settings. POPGym provides a diverse collection of
1754 lightweight vectorized environments with varying difficulty levels.

1755 **I.11.1 POPGym Autoencode**

1756 The environment consists of a deck of cards that is shuffled and sequentially shown to the agent
1757 during the watch phase. While observing the cards, a watch indicator is active, but it disappears
1758 when the last card is revealed. Afterward, the agent must reproduce the sequence of cards in the
1759 correct order. The environment is structured to evaluate the agent's ability to encode a sequence of
1760 observations into an internal representation and later reconstruct the sequence one observation at a
1761 time.

1762 **I.11.2 POPGym Concentration**

1763 The environment represents a classic memory game where a shuffled deck of cards is placed face-
1764 down. The agent sequentially flips two cards and earns a reward if the revealed cards form a matching
1765 pair. The environment is designed in such a way that the agent must remember previously revealed
1766 cards to maximize its success rate.

1767 **I.11.3 POPGym Repeat First**

1768 The environment presents the agent with an initial value from a set of four possible values, along with
1769 an indicator signaling that this is the first value. In subsequent steps, the agent continues to receive
1770 random values from the same set but without the initial indicator. The structure requires the agent to
1771 retain the first received value in memory and recall it accurately to receive a reward.

1772 **I.11.4 POPGym Repeat Previous**

1773 The environment consists of a sequence of observations, where each observation can take one of four
1774 possible values at each timestep. The agent is tasked with recalling and outputting the value that
1775 appeared a specified number of steps in the past.

1776 **I.11.5 POPGym Stateless Cartpole**

1777 This is a modified version of the traditional Cartpole environment [6] where angular and linear
1778 position information is removed from observations. Instead, the agent only receives velocity-based
1779 data and must infer positional states by integrating this information over time to successfully balance
1780 the pole.

1781 **I.11.6 POPGym Stateless Pendulum**

1782 In this variation of the swing-up pendulum environment [18], angular position data is omitted from
1783 the agent’s observations. The agent must infer the pendulum’s position by processing velocity
1784 information and use this estimate to determine appropriate control actions.

1785 **I.11.7 POPGym Noisy Stateless Cartpole**

1786 This environment builds upon Stateless Cartpole by introducing Gaussian noise into the observations.
1787 The agent must still infer positional states from velocity information while filtering out the added
1788 noise to maintain control of the pole.

1789 **I.11.8 POPGym Noisy Stateless Pendulum**

1790 This variation extends the Stateless Pendulum environment by incorporating Gaussian noise into
1791 the observations. The agent must manage this uncertainty while using velocity data to estimate the
1792 pendulum’s position and swing it up effectively.

1793 **I.11.9 POPGym Multiarmed Bandit**

1794 The Multiarmed Bandit environment is an episodic formulation of the multiarmed bandit problem [82],
1795 where a set of bandits is randomly initialized at the start of each episode. Unlike conventional
1796 multiarmed bandit tasks, where reward probabilities remain fixed across episodes, this structure resets
1797 them every time. The agent must dynamically adjust its exploration and exploitation strategies to
1798 maximize long-term rewards.

1799 **I.11.10 POPGym Higher Lower**

1800 Inspired by the higher-lower card game, this environment presents the agent with a sequence of cards.
1801 At each step, the agent must predict whether the next card will have a higher or lower rank than the
1802 current one. Upon making a guess, the next card is revealed and becomes the new reference. The
1803 agent can enhance its performance by employing card counting strategies to estimate the probability
1804 of future values.

1805 **I.11.11 POPGym Count Recall**

1806 At each timestep, the agent is presented with two values: a next value and a query value. The agent
1807 must determine and output how many times the query value has appeared so far. To succeed, the
1808 agent must maintain an accurate count of past occurrences and retrieve the correct number upon
1809 request.

1810 **I.11.12 POPGym Battleship**

1811 A partially observable variation of the game Battleship, where the agent does not have access to
1812 the full board. Instead, it receives feedback on its previous shot, indicating whether it was a HIT or
1813 MISS, along with the shot’s location. The agent earns rewards for hitting ships, receives no reward
1814 for missing, and incurs a penalty for targeting the same location more than once. The environment
1815 challenges the agent to construct an internal representation of the board and update its strategy based
1816 on past observations.

1817 **I.11.13 POPGym Mine Sweeper**

1818 A partially observable version of the computer game Mine Sweeper, where the agent lacks direct
1819 visibility of the board. Observations include the coordinates of the most recently clicked tile and
1820 the number of adjacent mines. Clicking on a mined tile results in a negative reward and ends the
1821 game. To succeed, the agent must track previous selections and deduce mine locations based on the
1822 numerical clues, ensuring it avoids mines while uncovering safe tiles.

1823 **I.11.14 POPGym Labyrinth Explore**

1824 The environment consists of a procedurally generated 2D maze in which the agent earns rewards
1825 for reaching new, unexplored tiles. Observations are limited to adjacent tiles, requiring the agent to
1826 infer the larger maze layout through exploration. A small penalty per timestep incentivizes efficient
1827 navigation and discovery strategies.

1828 **I.11.15 POPGym Labyrinth Escape**

1829 This variation of Labyrinth Explore challenges the agent to find an exit rather than merely exploring
1830 the maze. The agent retains the same restricted observation space, seeing only nearby tiles. Rewards
1831 are only given upon successfully reaching the exit, making it a sparse reward environment where the
1832 agent must navigate strategically to achieve its goal.

1833