

Figure 1: Code World Model of the Ant-v4 environment generated by GIF-MCTS with Llama 3. For reference, the official implementation of the environment can be found at the official Gymnasium GitHub repository of the Farama Foundation (at `gymnasium.envs.mujoco.Ant_v4`).

```

1 import numpy as np
2 class Environment:
3     def __init__(self,
4         healthy_reward=1.0,
5         forward_reward_weight=1.0,
6         ctrl_cost_weight=0.5,
7         contact_cost_weight=5e-4,
8         contact_force_range=10.0,
9         healthy_z_range=(0.2, 1.0),
10        terminate_when_unhealthy=True,
11        reset_noise_scale=1e-2,
12        exclude_current_positions_from_observation=True):
13        [...]
14    def _reset_state(self):
15        [...]
16    def set_state(self, state):
17        self.state = state
18
19    def step(self, action):
20        # implement dynamics model or physics engine to predict next state
21        next_state = self._predict_next_state(self.state, action)
22        observation = next_state
23        reward = self._calculate_reward(action, next_state)
24        done = self._check_episode_end(next_state)
25        return observation, reward, done
26
27    def _predict_next_state(self, state, action):
28        # implement dynamics model or physics engine to predict next state
29        # for example, using a simple integrator
30        # assume velocities are the last 14 components of the state
31        velocities = state[13:]
32        new_velocities = velocities.copy()
33        # update front left leg, front right leg, back left leg, and back right leg
34        new_velocities[:4] += 0.01 * action[:4]
35        # update corresponding body parts
36        new_velocities[4:8] += 0.01 * action[4:]
37        # Fix: slice new_velocities to match shape of state[:13]
38        new_positions = state[:13] + 0.01 * new_velocities[:13]
39        next_state = np.concatenate((new_positions, new_velocities))
40        return next_state
41
42    def _calculate_reward(self, action, next_state):
43        # implement reward calculation logic here
44        healthy_reward = self.healthy_reward
45        forward_reward = self.forward_reward_weight * \
46            (next_state[13] - self.state[13]) # x-coordinate velocity
47        ctrl_cost = self.ctrl_cost_weight * np.sum(np.square(action))
48        contact_cost = 0 # implement contact cost calculation logic here
49        reward = healthy_reward + forward_reward - ctrl_cost - contact_cost
50        return reward
51
52    def _check_episode_end(self, state):
53        # implement episode end check logic here
54        done = not (self.healthy_z_range[0] <= state[0] <= self.healthy_z_range[1])
55        return done

```