

ece60146hw11

Xinyu Chen

May 2, 2026

3 Programming Tasks

3.1 Task 1: Tabular Double Q-Learning on CartPole (15 points)

3.1.1 Reproducing the baseline (8 points)

```
1  #!/usr/bin/env python3
2  # -*- coding: utf-8 -*-
3  ## Cartpole_DQL.v2.py
4  ## Balancing a cartpole using Double Q-learning (without a
   neural network) and using
5  ## a discretized state space. The algorithm is trained using
   epsilon-greedy approach
6  ## with Replay Memory
7  ##
8  ## @author: Rohan Sarkar (sarkarr@purdue.edu)
9  import gymnasium as gym
10 import numpy as np
11 import matplotlib.pyplot as plt
12 import random
13
14 seed = 0
15 random.seed(seed)
16 np.random.seed(seed)
17
18 ## Initialize parameters and environment variables:
19 ENV_NAME = "CartPole-v1"
20 BASELINE_GAMMA = 0.85
21 ABLATION_GAMMA = 0.99
22 # MAX_ITER = 500
23 MAX_ITER = 1000
24 EXPLORATION_MAX = 1.0
25 EXPLORATION_MIN = 0.01
26 EXPLORATION_DECAY = 1 + np.log(EXPLORATION_MIN) / MAX_ITER
27 REP_MEM_SIZE = 100000
28 MINIBATCH_SIZE = 64
```

```

29 N_states = 162
30 N_actions = 2
31 Q1 = np.zeros((N_states, N_actions))
32 Q2 = np.zeros((N_states, N_actions))
33 alpha = 0.001
34 eps_rate = EXPLORATION_MAX
35 cum_reward = 0
36 train_reward_history = []
37
38
39 ## Map the four dimensional continuous state-space to
    discretized state-space:
40 ## Credit: http://www.derongliu.org/adp/adp-cdrom/Barto1983.pdf.
41 def map_discrete_state(cs): ## (B)
42     ds_vector = np.zeros(4)
43     ds = -1
44
45     # Discretize x (position)
46     if abs(cs[0]) <= 0.8:
47         ds_vector[0] = 1
48     elif cs[0] < -0.8:
49         ds_vector[0] = 0
50     elif cs[0] > 0.8:
51         ds_vector[0] = 2
52
53     # Discretize x' (velocity)
54     if abs(cs[1]) <= 0.5:
55         ds_vector[1] = 1
56     elif cs[1] < -0.5:
57         ds_vector[1] = 0
58     elif cs[1] > 0.5:
59         ds_vector[1] = 2
60
61     # Discretize theta (angle)
62     angle = 180 / 3.1428 * cs[2]
63     if -12 < angle <= -6:
64         ds_vector[2] = 0
65     elif -6 < angle <= -1:
66         ds_vector[2] = 1
67     elif -1 < angle <= 0:
68         ds_vector[2] = 2
69     elif 0 < angle <= 1:
70         ds_vector[2] = 3
71     elif 1 < angle <= 6:
72         ds_vector[2] = 4
73     elif 6 < angle <= 12:
74         ds_vector[2] = 5
75
76     # Discretize theta' (angular velocity)

```

```

77     if abs(cs[3]) <= 50:
78         ds_vector[3] = 1
79     elif cs[3] < -50:
80         ds_vector[3] = 0
81     elif cs[3] > 50:
82         ds_vector[3] = 2
83
84     ds = int(ds_vector[0] * 54 + ds_vector[1] * 18 +
85             ds_vector[2] * 3 + ds_vector[3])
86
87     return ds
88
89 ## Return the most optimal action and the corresponding Q
90 value:
91
92 def optimal_action(Q, state, eps_rate, env): ## (C)
93     p = np.random.random()
94
95     # Choose random action if in 'exploration' mode
96     if p < eps_rate:
97         action = env.action_space.sample() # choose
98         randomly between 0 and 1
99     else:
100         # Choose optimal action based on learned weights if
101         in 'exploitation' mode
102         action = np.argmax(Q[state, :])
103
104     return action
105
106 ## Update Qvalues based on the logic of the Double Q-
107 learning: ## (D)
108
109 def updateQValues(state, action, reward, next_state, alpha,
110 eps_rate, env, gamma, Q1, Q2):
111     p = np.random.random()
112
113     if p < 0.5:
114         # Update Qvalues for Table Q1
115         next_action = optimal_action(Q1, next_state,
116                                     eps_rate, env)
117         Q1[state][action] = Q1[state][action] + alpha * (
118             reward + gamma * Q2[next_state][next_action] -
119             Q1[state][action]
120         )
121     else:
122         # Update Qvalues for Table Q2
123         next_action = optimal_action(Q2, next_state,
124                                     eps_rate, env)
125         Q2[state][action] = Q2[state][action] + alpha * (
126             reward + gamma * Q1[next_state][next_action] -
127             Q2[state][action]

```

```

117         )
118
119     return next_action
120
121
122 class ReplayMemory: ## (E)
123     def __init__(self, capacity):
124         self.capacity = capacity
125         self.memory = []
126
127     def push(self, s, a, r, ns):
128         self.memory.append([s, a, r, ns])
129         if len(self.memory) > self.capacity:
130             del self.memory[0]
131
132     def sample(self, MINIBATCH_SIZE):
133         return random.sample(self.memory, MINIBATCH_SIZE)
134
135     def __len__(self):
136         return len(self.memory)
137
138
139 def run_training(gamma, label):
140     ## Learn the weights for Double Q learning:
141     random.seed(seed)
142     np.random.seed(seed)
143
144     env = gym.make(ENV_NAME) ## (A)
145     env.action_space.seed(seed)
146     Q1 = np.zeros((N_states, N_actions))
147     Q2 = np.zeros((N_states, N_actions))
148     eps_rate = EXPLORATION_MAX
149     cum_reward = 0
150     max_t = 0
151     duration_history = []
152     history = []
153     epsilon_history = []
154     train_reward_history = []
155     repmemory = ReplayMemory(REP_MEM_SIZE)
156
157     for i_episode in range(MAX_ITER):
158         ## state is a 4-tuple: [cart_pos, cart_vel,
159             pole_angle, pole_ang_vel]
160         if i_episode == 0:
161             state, info = env.reset(seed=seed)
162         else:
163             state, info = env.reset() ### state set to
164             uniformly dist random bet -0.05 and 0.05
165         done = False

```

```

165     # Initial state and action selection when
166         environment restarts
167     state = map_discrete_state(state)
168     action = optimal_action(0.5 * (Q1 + Q2), state, 0,
169         env)
170     t = 0
171
172     # Decay the exploration parameter in epsilon-greedy
173     approach.
174     eps_rate *= EXPLORATION_DECAY
175     eps_rate = max(EXPLORATION_MIN, eps_rate)
176
177     while True:
178         # env.render()
179         next_state, reward, terminated, truncated, info
180             = env.step(action)
181         done = terminated or truncated
182         if done:
183             reward = -10
184
185         next_state = map_discrete_state(next_state)
186         repmemory.push(state, action, reward, next_state
187             )
188
189         # Update Q table using Double Q learning and get
190         the next optimal action.
191         next_action = updateQValues(
192             state, action, reward, next_state, alpha,
193             eps_rate, env, gamma, Q1, Q2
194         )
195
196         # Update Q values by randomly sampling
197         experiences in Replay Memory
198         if len(repmemory) > MINIBATCH_SIZE:
199             experiences = repmemory.sample(
200                 MINIBATCH_SIZE)
201             for experience in experiences:
202                 ts, ta, tr, tns = experience
203                 updateQValues(ts, ta, tr, tns, alpha,
204                     eps_rate, env, gamma, Q1, Q2)
205
206         state = next_state
207         action = next_action
208         t += 1
209
210         if done:
211             break
212
213     history.append(t)
214

```

```

205         if i_episode > 50:
206             latest_duration = history[-50:]
207         else:
208             latest_duration = history
209
210         # print(latest_duration)
211         duration_run_avg = np.mean(latest_duration)
212         # print(duration_run_avg)
213         duration_history.append([t, duration_run_avg])
214         epsilon_history.append(eps_rate)
215         cum_reward += t
216
217         if t > max_t:
218             max_t = t
219
220         train_reward_history.append([cum_reward / (i_episode
221             + 1), max_t])
222         print(
223             "\n%s Episode: %d Episode duration: %d timesteps
224             | epsilon %f"
225             % (label, i_episode + 1, t + 1, eps_rate)
226         )
227
228         last_100_avg = np.mean(history[-100:])
229         print("%s average episode duration over last 100
230             training episodes: %.2f" % (label, last_100_avg))
231         env.close()
232         return Q1, Q2, history, duration_history,
233             epsilon_history, last_100_avg
234
235     if __name__ == "__main__": ## (F)
236         baseline_Q1, baseline_Q2, baseline_history,
237             baseline_duration_history, baseline_epsilon_history,
238             baseline_avg = run_training(
239                 BASELINE_GAMMA, "Baseline gamma=0.85"
240             )
241         ablation_Q1, ablation_Q2, ablation_history,
242             ablation_duration_history, ablation_epsilon_history,
243             ablation_avg = run_training(
244                 ABLATION_GAMMA, "Ablation gamma=0.99"
245             )
246
247         np.save("Q1.npy", ablation_Q1)
248         np.save("Q2.npy", ablation_Q2)
249
250         fig = plt.figure(1)
251         plt.clf()

```

```

246 baseline_episodes = np.arange(1, len(baseline_history) +
247                                1)
248 ablation_episodes = np.arange(1, len(ablation_history) +
249                                1)
250 baseline_moving_average = np.asarray(
251     baseline_duration_history)[: , 1]
252 ablation_moving_average = np.asarray(
253     ablation_duration_history)[: , 1]
254
255 plt.subplot(1, 2, 1)
256 plt.title("Training History")
257 plt.plot(baseline_episodes, baseline_moving_average,
258     label="gamma=0.85 moving average")
259 plt.plot(ablation_episodes, ablation_moving_average,
260     label="gamma=0.99 moving average")
261 plt.xlabel("Episodes")
262 plt.ylabel("Episode Duration")
263 plt.legend()
264
265 plt.subplot(1, 2, 2)
266 plt.title("Epsilon for Episode")
267 plt.plot(baseline_episodes, np.asarray(
268     baseline_epsilon_history), label="gamma=0.85")
269 plt.plot(ablation_episodes, np.asarray(
270     ablation_epsilon_history), label="gamma=0.99")
271 plt.xlabel("Episodes")
272 plt.ylabel("Epsilon Value")
273 plt.legend()
274
275 plt.tight_layout()
276 plt.savefig("Cartpole_DoubleQ_Learning.png")
277 plt.savefig("Cartpole_DoubleQ_Gamma_Comparison.png")
278 plt.close(fig)  ## (G)

```

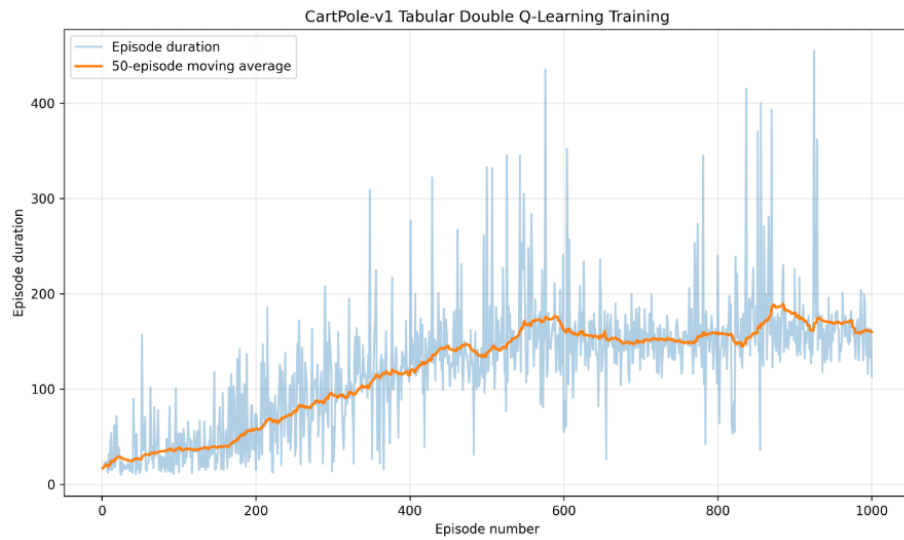


Figure 1: Loss curve

Average episode duration over last 100 training episodes: 163.13

3.1.2 One ablation (7 points)

I pick change discount factor to 0.99

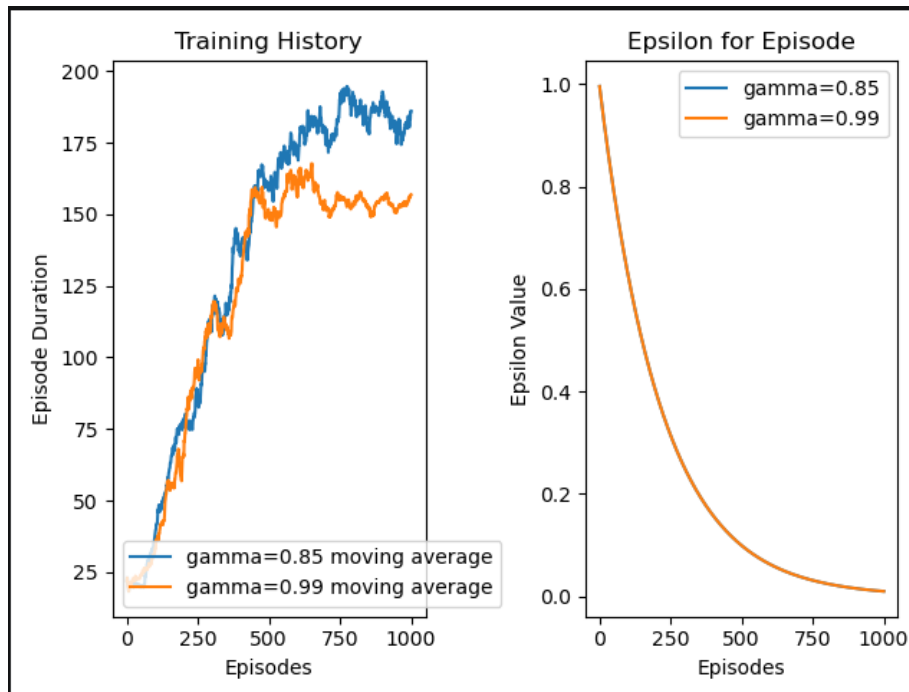


Figure 2: Comparison

For the ablation, I changed the discount factor from 0.85 to 0.99 while keeping the same random seed, number of episodes, and other hyperparameters. The training curves show that both agents improve over time, but the baseline with discount factor = 0.85 reaches a higher episode duration. With discount factor = 0.99, the moving average increases in the first half of training. This suggests that using a very large discount factor made learning less effective in this tabular CartPole setting. One possible reason is that discount factor = 0.99 puts more weight on long-term future rewards, which can make the Q-value estimates less stable when the state space is discretized. In contrast, discount factor = 0.85 focuses more on shorter term rewards, which may be easier to learn. So the ablation make the training worse.

3.2 Task 2: DQN on LunarLander (35 points)

3.2.1 Background

3.2.2 Implementation requirements (30 points)

```

1 import random
2 from collections import deque
3
4 import gymnasium as gym

```

```

5 import imageio
6 import matplotlib.pyplot as plt
7 import numpy as np
8 import torch
9 import torch.nn as nn
10 import torch.optim as optim
11
12
13
14 EPISODES = 1000
15 MAX_STEPS = 1000
16
17 GAMMA = 0.99
18 LR = 1e-3
19
20 BATCH_SIZE = 64
21 REPLAY_SIZE = 50000
22 MIN_REPLAY_SIZE = 1000
23
24 EPS_START = 1.0
25 EPS_END = 0.05
26 EPS_DECAY_EPISODES = int(0.2 * EPISODES) # first 20% of
    training
27
28 TARGET_UPDATE_FREQ = 1000 # hard update every 1000
    environment steps
29
30 CHECKPOINT_PATH = "dqn_lander.pt"
31 PLOT_PATH = "dqn_lander_training_curve.png"
32 GIF_PATH = "dqn_lander.gif"
33
34 SEED = 0
35
36
37 random.seed(SEED)
38 np.random.seed(SEED)
39 torch.manual_seed(SEED)
40
41
42
43 device = torch.device("cuda" if torch.cuda.is_available()
    else "cpu")
44
45
46 class QNet(nn.Module):
47     def __init__(self, state_dim=8, action_dim=4):
48         super().__init__()
49
50         self.net = nn.Sequential(
51             nn.Linear(state_dim, 128),

```

```

52         nn.ReLU(),
53         nn.Linear(128, 128),
54         nn.ReLU(),
55         nn.Linear(128, action_dim),
56     )
57
58     def forward(self, s):
59         return self.net(s)
60
61
62
63 class ReplayBuffer:
64     def __init__(self, capacity=50000):
65         self.buf = deque(maxlen=capacity)
66
67     def push(self, s, a, r, s_next, done):
68         self.buf.append((s, a, r, s_next, done))
69
70     def sample(self, batch_size):
71         batch = random.sample(self.buf, batch_size)
72
73         states, actions, rewards, next_states, dones = zip(*
74             batch)
75
76         states = torch.tensor(np.array(states), dtype=torch.
77             float32, device=device)
78         actions = torch.tensor(actions, dtype=torch.long,
79             device=device).unsqueeze(1)
80         rewards = torch.tensor(rewards, dtype=torch.float32,
81             device=device)
82         next_states = torch.tensor(np.array(next_states),
83             dtype=torch.float32, device=device)
84         dones = torch.tensor(dones, dtype=torch.float32,
85             device=device)
86
87         return states, actions, rewards, next_states, dones
88
89     def __len__(self):
90         return len(self.buf)
91
92
93 def get_epsilon(episode):
94     """
95     Linearly anneal epsilon from 1.0 to 0.05 over the first
96     20% of training.
97     """
98     if episode >= EPS_DECAY_EPISODES:
99         return EPS_END

```

```

95     frac = episode / EPS_DECAY_EPISODES
96     return EPS_START + frac * (EPS_END - EPS_START)
97
98
99
100 def select_action(state, model):
101     state_t = torch.tensor(state, dtype=torch.float32).
102         unsqueeze(0)
103
104     with torch.no_grad():
105         q_values = model(state_t)
106         action = torch.argmax(q_values, dim=1).item()
107
108     return action
109
110 def select_action_epsilon(q_net, state, epsilon, action_dim)
111 :
112     if random.random() < epsilon:
113         return random.randrange(action_dim)
114
115     state_t = torch.tensor(state, dtype=torch.float32,
116         device=device).unsqueeze(0)
117
118     with torch.no_grad():
119         q_values = q_net(state_t)
120         action = torch.argmax(q_values, dim=1).item()
121
122     return action
123
124 def train_step(q_net, target_net, optimizer, replay_buffer):
125     if len(replay_buffer) < MIN_REPLAY_SIZE:
126         return None
127
128     states, actions, rewards, next_states, dones =
129         replay_buffer.sample(BATCH_SIZE)
130
131     # Q_phi(s, a)
132     current_q = q_net(states).gather(1, actions).squeeze(1)
133
134     # y = r + gamma * (1 - done) * max_a' Q_target(s', a')
135     with torch.no_grad():
136         max_next_q = target_net(next_states).max(dim=1)[0]
137         target_q = rewards + GAMMA * (1.0 - dones) *
138             max_next_q
139
140     loss = nn.functional.mse_loss(current_q, target_q)

```

```

140     optimizer.zero_grad()
141     loss.backward()
142
143     # optional gradient clipping for stability
144     torch.nn.utils.clip_grad_norm_(q_net.parameters(),
145                                     max_norm=10.0)
146
147     optimizer.step()
148
149     return loss.item()
150
151
152 def moving_average(data, window=100):
153     if len(data) < window:
154         return np.array([])
155
156     return np.convolve(data, np.ones(window) / window, mode=
157                         "valid")
158
159
160 def plot_training_curve(episode_returns):
161     plt.figure(figsize=(10, 6))
162
163     plt.plot(episode_returns, label="Episode return", alpha
164             =0.5)
165
166     ma = moving_average(episode_returns, window=100)
167     if len(ma) > 0:
168         plt.plot(
169             range(99, 99 + len(ma)),
170             ma,
171             label="100-episode moving average",
172             linewidth=2,
173         )
174
175     plt.xlabel("Episode number")
176     plt.ylabel("Episode return")
177     plt.title("DQN on LunarLander")
178     plt.legend()
179     plt.grid(True)
180     plt.tight_layout()
181     plt.savefig(PLOT_PATH, dpi=300)
182     plt.show()
183
184
185 def record_gif(q_net, env_name, gif_path=GIF_PATH):
186     """

```

```

187     Record one episode using the trained policy.
188     """
189     env = gym.make(env_name, render_mode="rgb_array")
190
191     state, info = env.reset(seed=SEED)
192     frames = []
193     total_return = 0.0
194     done = False
195     steps = 0
196
197     while not done and steps < MAX_STEPS:
198         action = select_action(state, q_net)
199
200         state, reward, terminated, truncated, info = env.
201             step(action)
202         frames.append(env.render())
203         done = terminated or truncated
204
205         total_return += reward
206         steps += 1
207
208     env.close()
209
210     imageio.mimsave(gif_path, frames, fps=30)
211     print(f"Saved GIF to {gif_path}")
212     print(f"Recorded episode return: {total_return:.2f}")
213
214 def main():
215     # Try LunarLander-v3 first. If unavailable, use
216     # LunarLander-v2.
217     try:
218         env = gym.make("LunarLander-v3")
219         env_name = "LunarLander-v3"
220     except Exception:
221         env = gym.make("LunarLander-v2")
222         env_name = "LunarLander-v2"
223
224     print("Using environment:", env_name)
225
226     state_dim = env.observation_space.shape[0]
227     action_dim = env.action_space.n
228
229     print("State dim:", state_dim)
230     print("Action dim:", action_dim)
231
232     q_net = QNet(state_dim, action_dim).to(device)
233     target_net = QNet(state_dim, action_dim).to(device)
234
235     # Initial target network copy

```

```

235     target_net.load_state_dict(q_net.state_dict())
236     target_net.eval()
237
238     optimizer = optim.Adam(q_net.parameters(), lr=LR)
239     replay_buffer = ReplayBuffer(REPLAY_SIZE)
240
241     episode_returns = []
242     losses = []
243
244     global_step = 0
245     solved = False
246
247     for episode in range(EPOCHS):
248         state, info = env.reset(seed=SEED + episode)
249
250         episode_return = 0.0
251         epsilon = get_epsilon(episode)
252
253         for step in range(MAX_STEPS):
254             action = select_action_epsilon(q_net, state,
255                                           epsilon, action_dim)
256
257             next_state, reward, terminated, truncated, info
258                 = env.step(action)
259             done = terminated or truncated
260
261             replay_buffer.push(state, action, reward,
262                               next_state, done)
263
264             loss = train_step(q_net, target_net, optimizer,
265                              replay_buffer)
266             if loss is not None:
267                 losses.append(loss)
268
269             state = next_state
270             episode_return += reward
271             global_step += 1
272
273             # Hard update target network
274             if global_step % TARGET_UPDATE_FREQ == 0:
275                 target_net.load_state_dict(q_net.state_dict())
276
277             if done:
278                 break
279
280         episode_returns.append(episode_return)
281
282     avg_last_100 = np.mean(episode_returns[-100:])

```

```

280         print(
281             f"Episode {episode + 1:4d} | "
282             f"Return: {episode_return:8.2f} | "
283             f"Avg100: {avg_last_100:8.2f} | "
284             f"Epsilon: {epsilon:.3f} | "
285             f"Replay: {len(replay_buffer)}"
286         )
287
288         # Check solved condition
289         if len(episode_returns) >= 100 and avg_last_100 >=
290             200:
291             print(f"Solved at episode {episode + 1}!")
292             solved = True
293             break
294
295     env.close()
296
297     last_100_avg = np.mean(episode_returns[-100:])
298
299     print("\nTraining finished.")
300     print(f"Average return over last 100 episodes: {
301         last_100_avg:.2f}")
302     print("Solved environment:", solved)
303
304     # Save checkpoint
305     torch.save(q_net.state_dict(), CHECKPOINT_PATH)
306
307     print(f"Saved checkpoint to {CHECKPOINT_PATH}")
308
309     # Save plot
310     plot_training_curve(episode_returns)
311     print(f"Saved training curve to {PLOT_PATH}")
312
313     # Save GIF
314     record_gif(q_net, env_name, GIF_PATH)
315
316 if __name__ == "__main__":
317     main()

```

Average return over last 100 episodes: 200.49 It did solved environment.

3.2.3 Brief analysis (5 points)

- Why a target network? A target network make training more stable. If we use the same Q network to predict both current and target, the target will change every time the network update. It make model unstable. Using a separate network, the target will be fixed, so it has more stable value to learn.

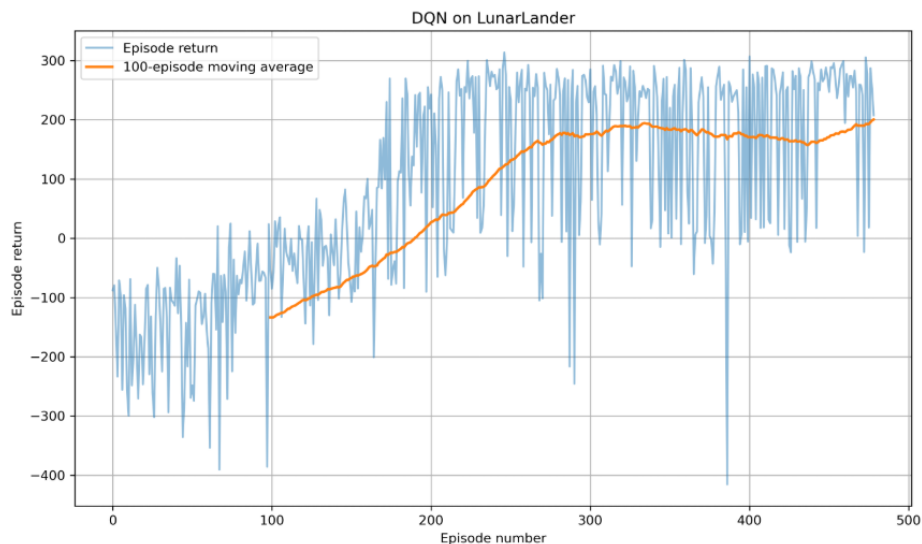


Figure 3: episode return vs. episode number

- epsilon-greedy decay : Because the agent should explore a lot at the beginning and use what it learned later. If epsilon starts at 0.05, the agent explores too little and may learn from poor early actions. If epsilon stays at 1.0, the agent keeps acting randomly and never really uses the learned policy. So changing from 1 to 0.05 can give balance.

3.3 Task 3: Reward modeling on real human preferences (50 points)

3.3.1 Background: what is reward modeling, and who is the “critic”?

3.3.2 The dataset

3.3.3 One small adaptation of Slide 20

3.3.4 Starter code

3.3.5 Load the dataset and look at it (8 points)

```

1 import sys
2 from pathlib import Path
3
4
5 STARTER_DIR = Path(__file__).resolve().parent / "
   HW11_starter_code"
6 sys.path.insert(0, str(STARTER_DIR))
7

```

```

8 from lm_utils import load_gpt2, load_preference_dataset,
   embed_triples
9
10
11 def main():
12     triples = load_preference_dataset(n_examples=500)
13     print(f"loaded {len(triples)} triples")
14
15     for i in range(3):
16         t = triples[i]
17         print(f"\n=== Example {i + 1} ===")
18         print(f"POST: {t['post'][:200]!r}...")
19         print(f"CHOSEN: {t['chosen']!r}")
20         print(f"REJECTED: {t['rejected']!r}")

```

=== Example 1 ===

POST: 'My boyfriend and I are long distance. We have a trip planned this summer which involves me going over to him in the USA. This will be the second time I have actually been with him in person. I am flyi'...

CHOSEN: 'I have made sure my mother is comfortable with my boyfriend travelling on a trip and now my mother is mad because I booked it.'

REJECTED: 'Mum is mad at me for not flying on my own trip to meet my boyfriend.'

=== Example 2 ===

POST: 'My boyfriend and I are long distance. We have a trip planned this summer which involves me going over to him in the USA. This will be the second time I have actually been with him in person. I am flyi'...

CHOSEN: "mum isn't speaking to me because I booked a flight and she doesn't want me flying on my own."

REJECTED: 'I have made sure my mother is comfortable with my boyfriend travelling on a trip and now my mother is mad because I booked it.'

=== Example 3 ===

POST: 'My boyfriend and I are long distance. We have a trip planned this summer which involves me going over to him in the USA. This will be the second time I have actually been with him in person. I am flyi'...

CHOSEN: "mum isn't speaking to me because I booked a flight and she doesn't want me flying on my own."

REJECTED: 'Mum thought I was going to road trip with my boyfriend. I cancelled the flight. Mum is annoyed because she thought she would be traveling with me. I am fine with that.'

The example come from the reddit post. The chosen summaries are only slightly better to me: it capture the main issue more directly, especially that the mother is upset because of the flight. However, the difference is not always very obvious, because some rejected summaries also contain similar information and the summaries are all quite short.

3.3.6 Embed everything (12 points)

```
1 tokenizer, base, device = load_gpt2()
2 embedded = embed_triples(triples, tokenizer, base,
3 device)
4 print(f"\nembedded {len(embedded)} triples")
5 print(f"post embedding shape: {embedded[0]['post_emb'].
6 shape}")
7
8 split_idx = int(0.8 * len(embedded))
9 train_embedded = embedded[:split_idx]
10 val_embedded = embedded[split_idx:]
11 print(f"train triples: {len(train_embedded)}")
12 print(f"validation triples: {len(val_embedded)}")
```

embedded 500 triples post embedding shape: torch.Size([768]) train triples: 400 validation triples: 100

3.3.7 Train the RewardNetwork (20 points)

```
1 #!/usr/bin/env python3
2 # -*- coding: utf-8 -*-
3
4 import random
5 import sys
6 from pathlib import Path
7
8 import matplotlib.pyplot as plt
9 import numpy as np
10 import torch
11 import torch.nn as nn
12
13
14 STARTER_DIR = Path(__file__).resolve().parent / "
15 HW11_starter_code"
16 sys.path.insert(0, str(STARTER_DIR))
17
18 from lm_utils import HIDDEN_SIZE, embed_triples, load_gpt2,
19 load_preference_dataset
20
21 SEED = 2
22 N_EXAMPLES = 500
23 TRAIN_STEPS = 2000
24 LR = 1e-3
25 LOSS_PLOT_PATH = Path(__file__).resolve().parent / "
26 reward_model_loss_curve.png"
```

```

27 random.seed(SEED)
28 np.random.seed(SEED)
29 torch.manual_seed(SEED)
30 if torch.cuda.is_available():
31     torch.cuda.manual_seed_all(SEED)
32
33
34 class RewardNetwork(torch.nn.Module):
35     def __init__(self, embedding_size):
36         torch.nn.Module.__init__(self)
37         self.linear1 = nn.Linear(3 * embedding_size, 2 *
38                                 embedding_size)
39         self.linear2 = nn.Linear(2 * embedding_size,
40                                 embedding_size)
41         self.linear3 = nn.Linear(embedding_size, 1)
42         self.sigmoid = nn.Sigmoid()
43
44     def forward(self, xSEPy):
45         x = self.linear1(xSEPy).clamp(min=0)
46         x = self.linear2(x).clamp(min=0)
47         x = self.linear3(x)
48         r = self.sigmoid(x)
49         return r
50
51 def make_input(example, response_key, device):
52     return torch.cat(
53         [
54             example["post_emb"],
55             example["sep_emb"],
56             example[response_key],
57         ]
58     ).to(device)
59
60 def pairwise_loss(reward_model, example, device):
61     chosen_input = make_input(example, "chosen_emb", device)
62     rejected_input = make_input(example, "rejected_emb",
63                                 device)
64     r_chosen = reward_model(chosen_input)
65     r_rejected = reward_model(rejected_input)
66     return -torch.nn.functional.logsigmoid(r_chosen -
67                                             r_rejected).mean()
68
69 def moving_average(values, window=200):
70     if len(values) < window:
71         return np.array([])
72     return np.convolve(values, np.ones(window) / window,
73                         mode="valid")

```

```

72
73
74 def plot_losses(losses):
75     steps = np.arange(1, len(losses) + 1)
76     avg = moving_average(losses, window=200)
77
78     plt.figure(figsize=(10, 6))
79     plt.plot(steps, losses, label="Per-step loss", alpha
80             =0.35)
81     if len(avg) > 0:
82         plt.plot(np.arange(200, 200 + len(avg)), avg, label=
83                 "Moving average (200)", linewidth=2)
84     plt.xlabel("Gradient step")
85     plt.ylabel("Loss")
86     plt.title("Reward Model Training Loss")
87     plt.legend()
88     plt.grid(True)
89     plt.tight_layout()
90     plt.savefig(LOSS_PLOT_PATH, dpi=300)
91     plt.close()
92
93 def validation_accuracy(reward_model, val_data, device):
94     correct = 0
95     reward_model.eval()
96     with torch.no_grad():
97         for example in val_data:
98             chosen_input = make_input(example, "chosen_emb",
99                                       device)
100             rejected_input = make_input(example, "
101                                         rejected_emb", device)
102             r_chosen = reward_model(chosen_input)
103             r_rejected = reward_model(rejected_input)
104             if r_chosen.item() > r_rejected.item():
105                 correct += 1
106     return correct / len(val_data)
107
108 def main():
109     triples = load_preference_dataset(n_examples=N_EXAMPLES)
110     print(f"loaded {len(triples)} triples")
111
112     tokenizer, base, embed_device = load_gpt2()
113     embedded = embed_triples(triples, tokenizer, base,
114                              embed_device)
115     print(f"embedded {len(embedded)} triples")
116     print(f"post embedding shape: {embedded[0]['post_emb'].
117           shape}")
118
119     split_idx = int(0.8 * len(embedded))

```

```

116 train_data = embedded[:split_idx]
117 val_data = embedded[split_idx:]
118 print(f"train triples: {len(train_data)}")
119 print(f"validation triples: {len(val_data)}")
120
121 del base
122 if torch.cuda.is_available():
123     torch.cuda.empty_cache()
124
125 device = torch.device("cuda" if torch.cuda.is_available
126                        () else "cpu")
127 reward_model = RewardNetwork(HIDDEN_SIZE).to(device)
128 optimizer = torch.optim.Adam(reward_model.parameters(),
129                               lr=LR)
130
131 losses = []
132 order = list(range(len(train_data)))
133 random.shuffle(order)
134
135 reward_model.train()
136 for step in range(TRAIN_STEPS):
137     if step > 0 and step % len(order) == 0:
138         random.shuffle(order)
139
140     example = train_data[order[step % len(order)]]
141     loss = pairwise_loss(reward_model, example, device)
142
143     optimizer.zero_grad()
144     loss.backward()
145     optimizer.step()
146
147     losses.append(loss.item())
148     if (step + 1) % 200 == 0:
149         print(f"step {step + 1:4d} | loss {loss.item()
150               :.4f}")
151
152 initial_loss = losses[0]
153 final_loss = float(np.mean(losses[-200:]))
154 val_acc = validation_accuracy(reward_model, val_data,
155                               device)
156 plot_losses(losses)
157
158 print(f"\ninitial loss: {initial_loss:.4f}")
159 print(f"final loss (last-200-step average): {final_loss
160       :.4f}")
161 print(f"validation accuracy: {val_acc:.4f}")
162 print(f"saved loss curve to: {LOSS_PLOT_PATH}")

```

```

161 if __name__ == "__main__":
162     main()

```

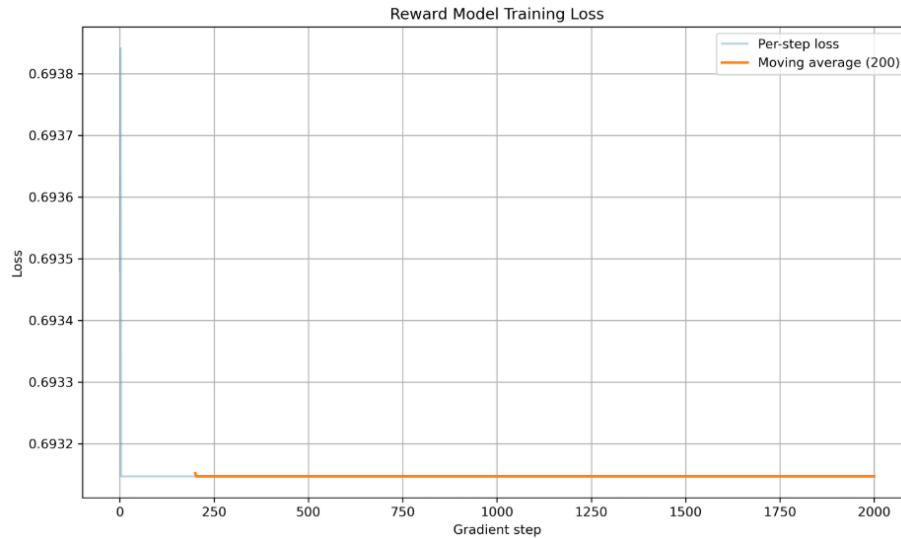


Figure 4: Loss curve

initial loss: 0.6935 final loss (last-200-step average): 0.6931

3.3.8 Evaluate and critique (10 points)

- The validation accuracy is 0.4600, The loss start with 0.6935 and the final last 200 step average is 0.6931. This is much worse than I expected, since the random guess will give 50% accuracy. So the model is learning nothing, gave the same reward.
- In the rewardnetwork, the first two layer will clipped activations to 0. And the final layer is passed to sigmoid. So the reward should be between 0 and 1. In this case, $-1 < r_{chase} - r_{rejected} < 1$. So the sigmoid prevents the model from producing large reward, which make model harder to decrease the loss.
- one possible modification is to remove the final sigmoid. And try iwth linear layer. So the $r_{chase} - r_{rejected}$ will become larger than 1, and the loss will decrease further.

3.4 Bonus Task: PPO on LunarLander (20 points)

I am using the code from CleanRL

```

1 import random
2 import re
3 import time
4 from pathlib import Path
5
6 import gymnasium as gym
7 import matplotlib.pyplot as plt
8 import numpy as np
9 import torch
10 import torch.nn as nn
11 import torch.optim as optim
12 from torch.distributions.categorical import Categorical
13
14
15 SEED = 0
16 ENV_ID = "LunarLander-v3"
17 TOTAL_TIMESTEPS = 500000
18 NUM_ENVS = 4
19 NUM_STEPS = 128
20 LEARNING_RATE = 2.5e-4
21 GAMMA = 0.99
22 GAE_LAMBDA = 0.95
23 UPDATE_EPOCHS = 4
24 NUM_MINIBATCHES = 4
25 CLIP_COEF = 0.2
26 ENT_COEF = 0.01
27 VF_COEF = 0.5
28 MAX_GRAD_NORM = 0.5
29
30 PPO_PLOT_PATH = "ppo_lander_training_curve.png"
31 COMPARISON_PLOT_PATH = "ppo_vs_dqn_comparison.png"
32 REPORT_PATH = "ppo_lander_report.txt"
33 DQN_LOG_PATH = "task2_v2.log"
34
35
36 def make_env(env_id, seed, idx):
37     def thunk():
38         env = gym.make(env_id)
39         env = gym.wrappers.RecordEpisodeStatistics(env)
40         env.action_space.seed(seed + idx)
41         env.observation_space.seed(seed + idx)
42         return env
43
44     return thunk
45
46
47 def layer_init(layer, std=np.sqrt(2), bias_const=0.0):
48     torch.nn.init.orthogonal_(layer.weight, std)
49     torch.nn.init.constant_(layer.bias, bias_const)

```

```

50     return layer
51
52
53 class Actor(nn.Module):
54     def __init__(self, state_dim=8, action_dim=4):
55         super().__init__()
56         self.net = nn.Sequential(
57             layer_init(nn.Linear(state_dim, 128)),
58             nn.ReLU(),
59             layer_init(nn.Linear(128, 128)),
60             nn.ReLU(),
61             layer_init(nn.Linear(128, action_dim), std=0.01)
62         ),
63
64     def forward(self, x):
65         return torch.softmax(self.net(x), dim=-1)
66
67
68 class Critic(nn.Module):
69     def __init__(self, state_dim=8):
70         super().__init__()
71         self.net = nn.Sequential(
72             layer_init(nn.Linear(state_dim, 128)),
73             nn.ReLU(),
74             layer_init(nn.Linear(128, 128)),
75             nn.ReLU(),
76             layer_init(nn.Linear(128, 1), std=1.0),
77         )
78
79     def forward(self, x):
80         return self.net(x)
81
82
83 class Agent(nn.Module):
84     def __init__(self, state_dim=8, action_dim=4):
85         super().__init__()
86         self.actor = Actor(state_dim, action_dim)
87         self.critic = Critic(state_dim)
88
89     def get_value(self, x):
90         return self.critic(x)
91
92     def get_action_and_value(self, x, action=None):
93         probs = self.actor(x)
94         dist = Categorical(probs=probs)
95         if action is None:
96             action = dist.sample()
97         return action, dist.log_prob(action), dist.entropy()
98         , self.critic(x)

```

```

98
99
100 def moving_average(values, window=100):
101     values = np.asarray(values, dtype=np.float32)
102     if len(values) < window:
103         return np.array([])
104     return np.convolve(values, np.ones(window) / window,
105                        mode="valid")
106
107 def parse_dqn_returns(log_path):
108     path = Path(log_path)
109     if not path.exists():
110         return []
111     text = path.read_text()
112     return [
113         float(match.group(1))
114         for match in re.finditer(r"Return:\s*([-+]?\d+(?:\.\d+)?)", text)
115     ]
116
117
118 def save_ppo_plot(ppo_returns):
119     plt.figure(figsize=(10, 6))
120     plt.plot(ppo_returns, label="PP0 episode return", alpha
121            =0.45)
122     ma = moving_average(ppo_returns, window=100)
123     if len(ma) > 0:
124         plt.plot(range(99, 99 + len(ma)), ma, label="PP0
125                100-episode moving average", linewidth=2)
126     plt.xlabel("Episode")
127     plt.ylabel("Return")
128     plt.title("PP0 on LunarLander")
129     plt.grid(True, alpha=0.3)
130     plt.legend()
131     plt.tight_layout()
132     plt.savefig(PP0_PLOT_PATH, dpi=300)
133     plt.close()
134
135 def save_comparison_plot(ppo_returns, dqn_returns):
136     plt.figure(figsize=(10, 6))
137
138     ppo_ma = moving_average(ppo_returns, window=100)
139     if len(ppo_ma) > 0:
140         plt.plot(range(99, 99 + len(ppo_ma)), ppo_ma, label=
141                "PP0 Avg100", linewidth=2)
142
143     dqn_ma = moving_average(dqn_returns, window=100)
144     if len(dqn_ma) > 0:

```

```

143         plt.plot(range(99, 99 + len(dqn_ma)), dqn_ma, label=
144                     "DQN Avg100", linewidth=2)
145
146     plt.xlabel("Episode")
147     plt.ylabel("100-episode moving average return")
148     plt.title("LunarLander Training Curve Comparison")
149     plt.grid(True, alpha=0.3)
150     plt.legend()
151     plt.tight_layout()
152     plt.savefig(COMPARISON_PLOT_PATH, dpi=300)
153     plt.close()
154
155
156
157 def main():
158     total_timesteps = TOTAL_TIMESTEPS
159     seed = SEED
160
161     random.seed(seed)
162     np.random.seed(seed)
163     torch.manual_seed(seed)
164     torch.backends.cudnn.deterministic = True
165
166     device = torch.device("cuda" if torch.cuda.is_available
167                           () else "cpu")
168
169     try:
170         envs = gym.vector.SyncVectorEnv(
171             [make_env(ENV_ID, seed, i) for i in range(
172                 NUM_ENVS)]
173         )
174         env_id = ENV_ID
175     except Exception:
176         env_id = "LunarLander-v2"
177         envs = gym.vector.SyncVectorEnv(
178             [make_env(env_id, seed, i) for i in range(
179                 NUM_ENVS)]
180         )
181
182     state_dim = int(np.array(envs.single_observation_space.
183                             shape).prod())
184     action_dim = envs.single_action_space.n
185     agent = Agent(state_dim, action_dim).to(device)
186     optimizer = optim.Adam(agent.parameters(), lr=
187                             LEARNING_RATE, eps=1e-5)
188
189     batch_size = NUM_ENVS * NUM_STEPS
190     minibatch_size = batch_size // NUM_MINIBATCHES
191     num_iterations = total_timesteps // batch_size

```

```

187     obs = torch.zeros((NUM_STEPS, NUM_ENVS) + envs.
188         single_observation_space.shape, device=device)
189     actions = torch.zeros((NUM_STEPS, NUM_ENVS), device=
190         device)
191     logprobs = torch.zeros((NUM_STEPS, NUM_ENVS), device=
192         device)
193     rewards = torch.zeros((NUM_STEPS, NUM_ENVS), device=
194         device)
195     dones = torch.zeros((NUM_STEPS, NUM_ENVS), device=device
196         )
197     values = torch.zeros((NUM_STEPS, NUM_ENVS), device=
198         device)
199
200     ppo_returns = []
201     global_step = 0
202     start_time = time.time()
203     next_obs, info = envs.reset(seed=seed)
204     next_obs = torch.tensor(next_obs, dtype=torch.float32,
205         device=device)
206     next_done = torch.zeros(NUM_ENVS, device=device)
207
208     print(f"Using environment: {env_id}")
209     print(f"Using device: {device}")
210
211     for iteration in range(1, num_iterations + 1):
212         frac = 1.0 - (iteration - 1.0) / num_iterations
213         optimizer.param_groups[0]["lr"] = frac *
214             LEARNING_RATE
215
216         for step in range(NUM_STEPS):
217             global_step += NUM_ENVS
218             obs[step] = next_obs
219             dones[step] = next_done
220
221             with torch.no_grad():
222                 action, logprob, entropy, value = agent.
223                     get_action_and_value(next_obs)
224                 values[step] = value.flatten()
225
226             actions[step] = action
227             logprobs[step] = logprob
228
229             next_obs_np, reward, terminated, truncated,
230                 infos = envs.step(action.cpu().numpy())
231             done = np.logical_or(terminated, truncated)
232             rewards[step] = torch.tensor(reward, dtype=torch
233                 .float32, device=device)
234             next_obs = torch.tensor(next_obs_np, dtype=torch
235                 .float32, device=device)
236             next_done = torch.tensor(done, dtype=torch.

```

```

float32, device=device)

225
226 if "final_info" in infos:
227     for info in infos["final_info"]:
228         if info and "episode" in info:
229             episode_return = float(info["episode
                "]["r"])
230             ppo_returns.append(episode_return)
231             avg100 = np.mean(ppo_returns[-100:])
232             print(
233                 f"Episode {len(ppo_returns):4d}
                    | Return: {episode_return:8.2
                        f} | "
234                 f"Avg100: {avg100:8.2f} | Global
                    step: {global_step}")
235             )
236
237 with torch.no_grad():
238     next_value = agent.get_value(next_obs).reshape
        (1, -1)
239     advantages = torch.zeros_like(rewards, device=
        device)
240     lastgaelam = 0
241     for t in reversed(range(NUM_STEPS)):
242         if t == NUM_STEPS - 1:
243             nextnonterminal = 1.0 - next_done
244             nextvalues = next_value
245         else:
246             nextnonterminal = 1.0 - dones[t + 1]
247             nextvalues = values[t + 1]
248             delta = rewards[t] + GAMMA * nextvalues *
                nextnonterminal - values[t]
249             advantages[t] = lastgaelam = delta + GAMMA *
                GAE_LAMBDA * nextnonterminal *
                lastgaelam
250     returns = advantages + values
251
252 b_obs = obs.reshape((-1,) + envs.
    single_observation_space.shape)
253 b_logprobs = logprobs.reshape(-1)
254 b_actions = actions.reshape(-1).long()
255 b_advantages = advantages.reshape(-1)
256 b_returns = returns.reshape(-1)
257 b_values = values.reshape(-1)
258
259 b_inds = np.arange(batch_size)
260 for epoch in range(UPDATE_EPOCHS):
261     np.random.shuffle(b_inds)
262     for start in range(0, batch_size, minibatch_size
        ):

```

```

263         mb_inds = b_inds[start : start +
264             minibatch_size]
265         _, newlogprob, entropy, newvalue = agent.
266             get_action_and_value(
267                 b_obs[mb_inds], b_actions[mb_inds]
268             )
269         logratio = newlogprob - b_logprobs[mb_inds]
270         ratio = logratio.exp()
271
272         mb_advantages = b_advantages[mb_inds]
273         mb_advantages = (mb_advantages -
274             mb_advantages.mean()) / (mb_advantages.
275                 std() + 1e-8)
276
277         pg_loss1 = -mb_advantages * ratio
278         pg_loss2 = -mb_advantages * torch.clamp(
279             ratio, 1 - CLIP_COEF, 1 + CLIP_COEF)
280         pg_loss = torch.max(pg_loss1, pg_loss2).mean()
281
282         newvalue = newvalue.view(-1)
283         v_loss = 0.5 * ((newvalue - b_returns[
284             mb_inds]) ** 2).mean()
285         entropy_loss = entropy.mean()
286         loss = pg_loss - ENT_COEF * entropy_loss +
287             VF_COEF * v_loss
288
289         optimizer.zero_grad()
290         loss.backward()
291         nn.utils.clip_grad_norm_(agent.parameters(),
292             MAX_GRAD_NORM)
293         optimizer.step()
294
295         if iteration % 10 == 0:
296             sps = int(global_step / (time.time() -
297                 start_time))
298             print(f"Iteration {iteration:4d}/{num_iterations
299                 } | SPS: {sps}")
300
301     envs.close()
302
303     save_ppo_plot(ppo_returns)
304     dqn_returns = parse_dqn_returns(DQN_LOG_PATH)
305     save_comparison_plot(ppo_returns, dqn_returns)
306     print(f"Saved PPO curve to {PPO_PLOT_PATH}")
307     print(f"Saved comparison plot to {COMPARISON_PLOT_PATH}")
308
309 if __name__ == "__main__":

```

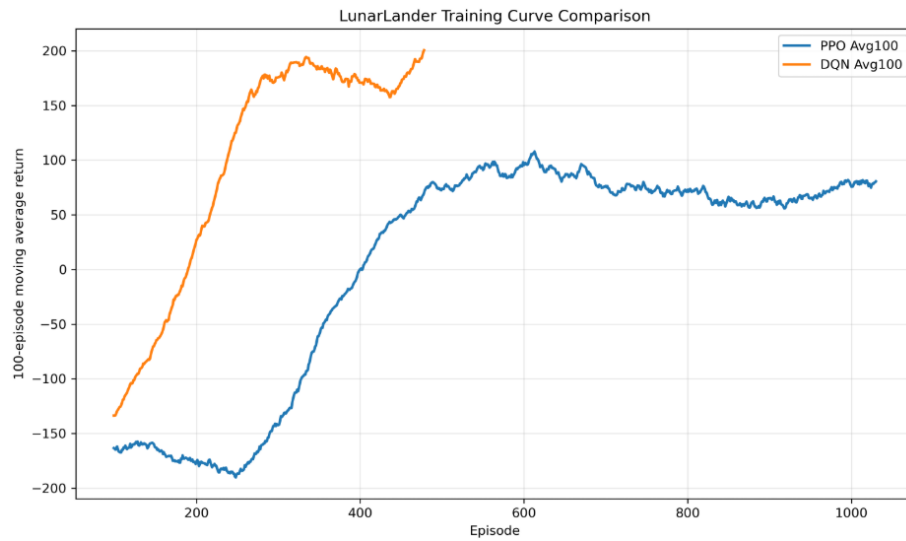


Figure 5: comparison

Average return over the last 100 episodes for PPO: 80.56.
DQN solved the environment much faster than PPO, it using around 500 episode to reach 200 return and ppo use 1000 episode. DQN appears easier to tune for this task, since it achieved strong performance with fewer episodes. I think for a new problem, I would like to test on DQN first, since it is fast.