

ECE 60146 and BME 64600: Homework 11

Triparna Mahata

May 3, 2026

3 Programming Tasks

3.1 Task 1: Tabular Double Q-Learning on CartPole

3.1.1 Reproducing the baseline

Code:

```
1 import gymnasium as gym
2 import numpy as np
3 import matplotlib.pyplot as plt
4 import random
5
6 # -----
7 # Reproducibility
8 # -----
9 seed = 0
10 random.seed(seed)
11 np.random.seed(seed)
12
13 # -----
14 # Initialize parameters and environment
15 # -----
16 env = gym.make('CartPole-v1') # (A) v0 -> v1
17 # In gymnasium the action space carries its OWN RNG, separate from
18 # numpy.random and Python's random module. Must be seeded explicitly,
19 # otherwise env.action_space.sample() (used by epsilon-greedy) is
20 # nondeterministic across runs.
21 env.action_space.seed(seed)
22
23 GAMMA = 0.85
24 MAX_ITER = 1000
25 EXPLORATION_MAX = 1.0
26 EXPLORATION_MIN = 0.01
27 EXPLORATION_DECAY = 1 + np.log(EXPLORATION_MIN) / MAX_ITER
28 REP_MEM_SIZE = 100000
29 MINIBATCH_SIZE = 64
30
31 N_states = 162
32 N_actions = 2
33
34 Q1 = np.zeros((N_states, N_actions))
35 Q2 = np.zeros((N_states, N_actions))
36
37 alpha = 0.001
38 eps_rate = EXPLORATION_MAX
39
40
41 # -----
42 # Map the four-dimensional continuous state to a discrete state in [0, 161]
43 # Credit: http://www.derongliu.org/adp/adp-cdrom/Barto1983.pdf
44 # -----
45 def map_discrete_state(cs): # (B)
46     ds_vector = np.zeros(4)
```

```

47 # x (cart position)
48 if abs(cs[0]) <= 0.8:
49     ds_vector[0] = 1
50 elif cs[0] < -0.8:
51     ds_vector[0] = 0
52 elif cs[0] > 0.8:
53     ds_vector[0] = 2
54
55 # x' (cart velocity)
56 if abs(cs[1]) <= 0.5:
57     ds_vector[1] = 1
58 elif cs[1] < -0.5:
59     ds_vector[1] = 0
60 elif cs[1] > 0.5:
61     ds_vector[1] = 2
62
63 # theta (pole angle, in degrees)
64 angle = 180 / 3.1428 * cs[2]
65 if -12 < angle <= -6: ds_vector[2] = 0
66 elif -6 < angle <= -1: ds_vector[2] = 1
67 elif -1 < angle <= 0: ds_vector[2] = 2
68 elif 0 < angle <= 1: ds_vector[2] = 3
69 elif 1 < angle <= 6: ds_vector[2] = 4
70 elif 6 < angle <= 12: ds_vector[2] = 5
71
72 # theta' (pole angular velocity)
73 if abs(cs[3]) <= 50:
74     ds_vector[3] = 1
75 elif cs[3] < -50:
76     ds_vector[3] = 0
77 elif cs[3] > 50:
78     ds_vector[3] = 2
79
80 ds = int(ds_vector[0] * 54 + ds_vector[1] * 18
81         + ds_vector[2] * 3 + ds_vector[3])
82 return ds
83
84
85 # -----
86 # Epsilon-greedy action selection
87 # -----
88
89 def optimal_action(Q, state, eps_rate, env): # (C)
90     if np.random.random() < eps_rate:
91         return env.action_space.sample()
92     return int(np.argmax(Q[state, :]))
93
94
95 # -----
96 # Double Q-learning update
97 # -----
98
99 def updateQValues(state, action, reward, next_state, alpha, eps_rate, env): # (D)
100     if np.random.random() < 0.5:
101         next_action = optimal_action(Q1, next_state, eps_rate, env)
102         Q1[state][action] += alpha * (
103             reward + GAMMA * Q2[next_state][next_action] - Q1[state][action]
104         )
105     else:
106         next_action = optimal_action(Q2, next_state, eps_rate, env)
107         Q2[state][action] += alpha * (
108             reward + GAMMA * Q1[next_state][next_action] - Q2[state][action]
109         )
110     return next_action
111
112 # -----

```

```

113 # Replay memory
114 # -----
115 class ReplayMemory:                                     # (E)
116     def __init__(self, capacity):
117         self.capacity = capacity
118         self.memory = []
119
120     def push(self, s, a, r, ns):
121         self.memory.append([s, a, r, ns])
122         if len(self.memory) > self.capacity:
123             del self.memory[0]
124
125     def sample(self, MINIBATCH_SIZE):
126         return random.sample(self.memory, MINIBATCH_SIZE)
127
128     def __len__(self):
129         return len(self.memory)
130
131
132 # -----
133 # Training loop
134 # -----
135 if __name__ == "__main__":                             # (F)
136
137     duration_history = []                                # raw episode durations (timesteps)
138     epsilon_history = []
139     repmemory = ReplayMemory(REP_MEM_SIZE)
140
141     for i_episode in range(MAX_ITER):
142
143         # --- gymnasium reset returns (obs, info) ---
144         state, info = env.reset(seed=seed if i_episode == 0 else None)
145         state = map_discrete_state(state)
146         action = optimal_action(0.5 * (Q1 + Q2), state, 0, env)
147
148         # epsilon decay
149         eps_rate *= EXPLORATION_DECAY
150         eps_rate = max(EXPLORATION_MIN, eps_rate)
151
152         t = 0
153         while True:
154             # --- gymnasium step returns 5 values ---
155             next_state, reward, terminated, truncated, info = env.step(action)
156             done = terminated or truncated
157
158             if done:
159                 reward = -10
160
161             next_state = map_discrete_state(next_state)
162             repmemory.push(state, action, reward, next_state)
163
164             # one on-policy Double-Q update
165             next_action = updateQValues(state, action, reward, next_state,
166                                         alpha, eps_rate, env)
167
168             # replay updates
169             if len(repmemory) > MINIBATCH_SIZE:
170                 experiences = repmemory.sample(MINIBATCH_SIZE)
171                 for ts, ta, tr, tns in experiences:
172                     updateQValues(ts, ta, tr, tns, alpha, eps_rate, env)
173
174             state = next_state
175             action = next_action
176             t += 1
177
178         if done:

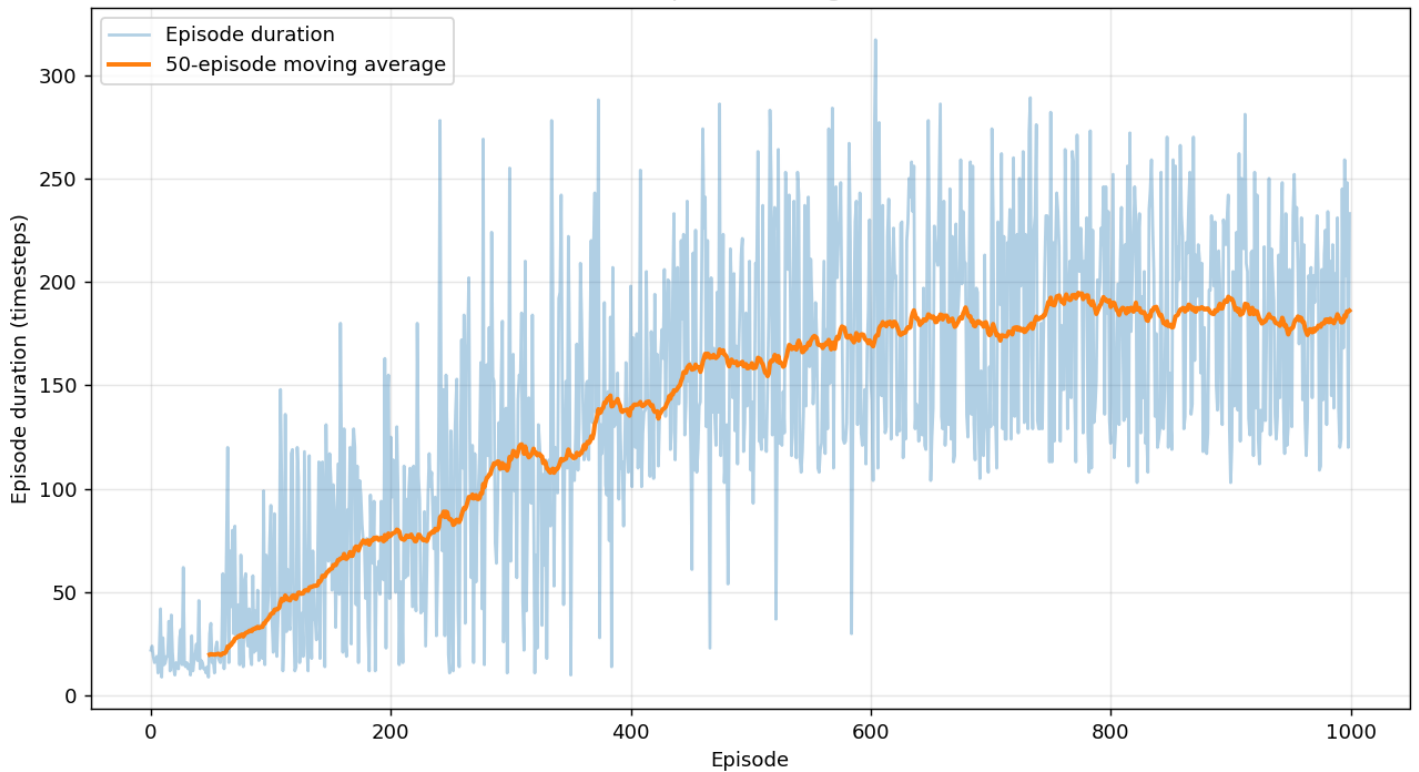
```

```

179         break
180
181     duration_history.append(t)
182     epsilon_history.append(eps_rate)
183
184     if (i_episode + 1) % 50 == 0:
185         recent = np.mean(duration_history[-50:])
186         print(f"Episode {i_episode+1:4d} | duration {t:4d} "
187               f"| 50-ep avg {recent:6.2f} | epsilon {eps_rate:.4f}")
188
189 env.close()
190
191 # -----
192 # Save Q tables
193 # -----
194 np.save('Q1.npy', Q1)
195 np.save('Q2.npy', Q2)
196
197 # -----
198 # Average episode duration over the LAST 100 training episodes
199 # -----
200 last100_avg = float(np.mean(duration_history[-100:]))
201 print("\n" + "=" * 60)
202 print(f"Average episode duration over the last 100 episodes: "
203       f"{last100_avg:.2f}")
204 print("=" * 60)
205
206 # -----
207 # Plot: episode duration vs. episode number with moving-average overlay
208 # -----
209 durations = np.asarray(duration_history)
210 window     = 50
211 moving_avg = np.convolve(durations,
212                           np.ones(window) / window,
213                           mode='valid')
214 ma_x = np.arange(window - 1, len(durations))
215
216 plt.figure(figsize=(10, 6))
217 plt.plot(durations, alpha=0.35, label='Episode duration')
218 plt.plot(ma_x, moving_avg, linewidth=2.2,
219          label=f'{window}-episode moving average')
220 plt.xlabel('Episode')
221 plt.ylabel('Episode duration (timesteps)')
222 plt.title('Tabular Double Q-Learning on CartPole-v1\n'
223           f'Last-100-episode average: {last100_avg:.2f}')
224 plt.legend(loc='upper left')
225 plt.grid(alpha=0.3)
226 plt.tight_layout()
227 plt.savefig('tabular_dq_training_curve.png', dpi=130)
228 print("\nSaved training curve to tabular_dq_training_curve.png")

```

Tabular Double Q-Learning on CartPole-v1
Last-100-episode average: 181.12



3.1.2 One ablation

Code:

```

1 import gymnasium as gym
2 import numpy as np
3 import matplotlib.pyplot as plt
4 import random
5
6
7 # -----
8 # Hyperparameters (identical across baseline and ablation)
9 # -----
10 SEED = 0
11 GAMMA = 0.85
12 MAX_ITER = 1000
13 EXPLORATION_MAX = 1.0
14 EXPLORATION_MIN = 0.01
15 EXPLORATION_DECAY = 1 + np.log(EXPLORATION_MIN) / MAX_ITER
16 REP_MEM_SIZE = 100000
17 MINIBATCH_SIZE = 64
18 ALPHA = 0.001
19 N_STATES = 162
20 N_ACTIONS = 2
21
22
23 # -----
24 # Discrete state mapping (from the lecture)
25 # -----
26 def map_discrete_state(cs):
27     ds_vector = np.zeros(4)
28
29     if abs(cs[0]) <= 0.8: ds_vector[0] = 1
30     elif cs[0] < -0.8: ds_vector[0] = 0
31     elif cs[0] > 0.8: ds_vector[0] = 2
32

```

```

33     if abs(cs[1]) <= 0.5:         ds_vector[1] = 1
34     elif cs[1] < -0.5:          ds_vector[1] = 0
35     elif cs[1] > 0.5:           ds_vector[1] = 2
36
37     angle = 180 / 3.1428 * cs[2]
38     if -12 < angle <= -6: ds_vector[2] = 0
39     elif -6 < angle <= -1: ds_vector[2] = 1
40     elif -1 < angle <= 0: ds_vector[2] = 2
41     elif 0 < angle <= 1: ds_vector[2] = 3
42     elif 1 < angle <= 6: ds_vector[2] = 4
43     elif 6 < angle <= 12: ds_vector[2] = 5
44
45     if abs(cs[3]) <= 50:         ds_vector[3] = 1
46     elif cs[3] < -50:           ds_vector[3] = 0
47     elif cs[3] > 50:            ds_vector[3] = 2
48
49     return int(ds_vector[0] * 54 + ds_vector[1] * 18
50               + ds_vector[2] * 3 + ds_vector[3])
51
52
53 def epsilon_greedy(Q, state, eps_rate, env):
54     if np.random.random() < eps_rate:
55         return env.action_space.sample()
56     return int(np.argmax(Q[state, :]))
57
58
59 class ReplayMemory:
60     def __init__(self, capacity):
61         self.capacity = capacity
62         self.memory = []
63     def push(self, s, a, r, ns):
64         self.memory.append([s, a, r, ns])
65         if len(self.memory) > self.capacity:
66             del self.memory[0]
67     def sample(self, k):
68         return random.sample(self.memory, k)
69     def __len__(self):
70         return len(self.memory)
71
72
73 # -----
74 # Generic training loop, parameterised by mode ("double" or "single")
75 # -----
76 def train(mode):
77     """
78     mode = "double" -> baseline Double Q-learning (two tables Q1, Q2)
79     mode = "single" -> ablation, standard Q-learning (one table Q)
80     """
81     assert mode in ("double", "single")
82
83     # Re-seed every RNG so the two runs share an identical starting point.
84     random.seed(SEED)
85     np.random.seed(SEED)
86
87     env = gym.make('CartPole-v1')
88     env.action_space.seed(SEED) # gymnasium action-space RNG (separate)
89
90     if mode == "double":
91         Q1 = np.zeros((N_STATES, N_ACTIONS))
92         Q2 = np.zeros((N_STATES, N_ACTIONS))
93     else:
94         Q = np.zeros((N_STATES, N_ACTIONS))
95
96     eps_rate = EXPLORATION_MAX
97     duration_history = []
98     repmemory = ReplayMemory(REP_MEM_SIZE)

```

```

99 # --- update rule ---
100 def update(s, a, r, ns, eps_rate):
101     if mode == "double":
102         if np.random.random() < 0.5:
103             na = epsilon_greedy(Q1, ns, eps_rate, env)
104             Q1[s][a] += ALPHA * (r + GAMMA * Q2[ns][na] - Q1[s][a])
105         else:
106             na = epsilon_greedy(Q2, ns, eps_rate, env)
107             Q2[s][a] += ALPHA * (r + GAMMA * Q1[ns][na] - Q2[s][a])
108         return na
109     else:
110         # Standard single Q-learning: target uses max over a' of the SAME table.
111         na = epsilon_greedy(Q, ns, eps_rate, env)
112         Q[s][a] += ALPHA * (r + GAMMA * np.max(Q[ns, :]) - Q[s][a])
113         return na
114
115 # --- training ---
116 for i_episode in range(MAX_ITER):
117     state, _ = env.reset(seed=SEED if i_episode == 0 else None)
118     state = map_discrete_state(state)
119
120     if mode == "double":
121         action = epsilon_greedy(0.5 * (Q1 + Q2), state, 0, env)
122     else:
123         action = epsilon_greedy(Q, state, 0, env)
124
125     eps_rate = max(EXPLORATION_MIN, eps_rate * EXPLORATION_DECAY)
126
127     t = 0
128     while True:
129         next_state, reward, terminated, truncated, _ = env.step(action)
130         done = terminated or truncated
131         if done:
132             reward = -10
133             next_state = map_discrete_state(next_state)
134             repmemory.push(state, action, reward, next_state)
135
136             next_action = update(state, action, reward, next_state, eps_rate)
137
138             if len(repmemory) > MINIBATCH_SIZE:
139                 for ts, ta, tr, tns in repmemory.sample(MINIBATCH_SIZE):
140                     update(ts, ta, tr, tns, eps_rate)
141
142             state, action = next_state, next_action
143             t += 1
144             if done:
145                 break
146
147     duration_history.append(t)
148
149     if (i_episode + 1) % 100 == 0:
150         recent = np.mean(duration_history[-50:])
151         print(f" [{mode:6s}] episode {i_episode+1:4d} | "
152               f"50-ep avg {recent:6.2f} | epsilon {eps_rate:.4f}")
153
154     env.close()
155     return np.asarray(duration_history)
156
157 # -----
158 # Run both experiments
159 # -----
160 if __name__ == "__main__":
161     print("Training BASELINE (Double Q-learning) ...")

```

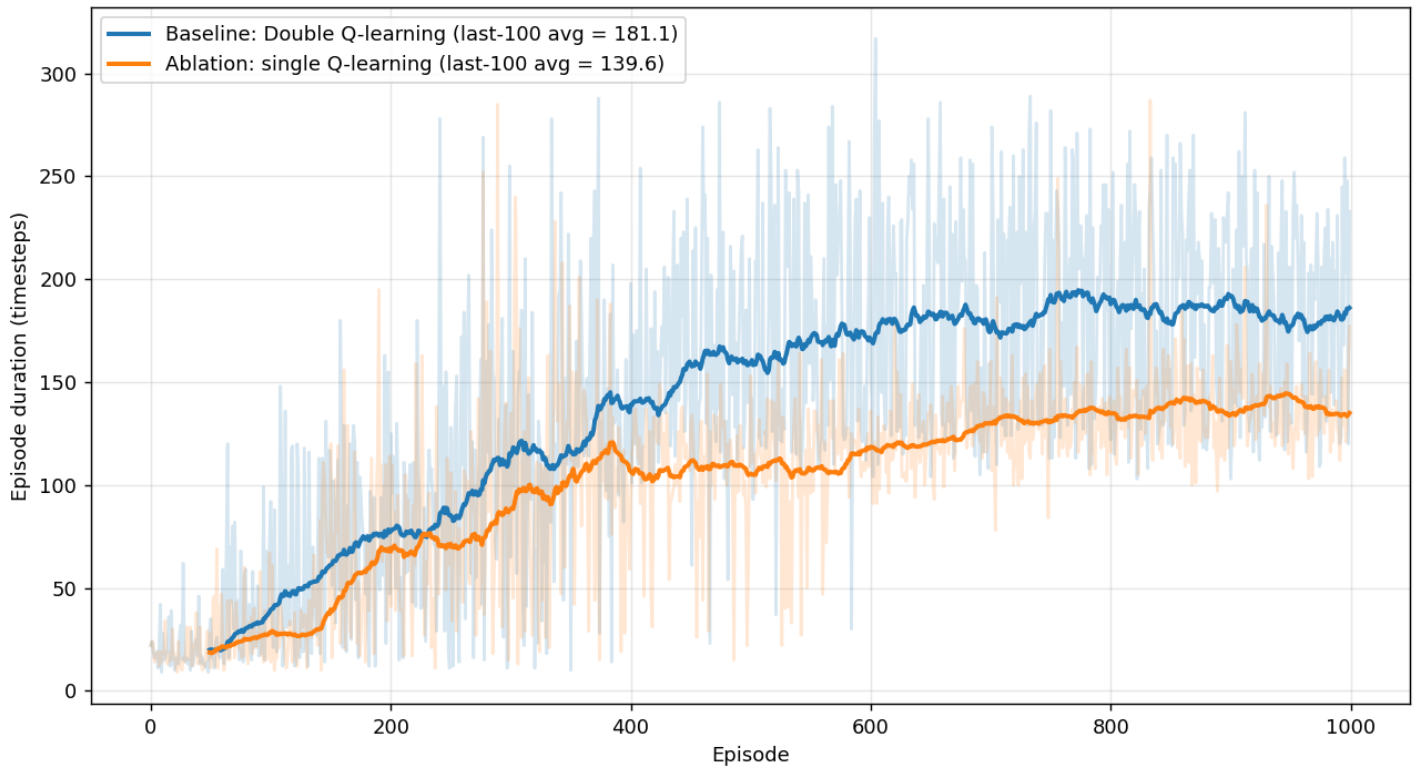
```

165 baseline = train("double")
166 print(f" -> last-100-episode average: {baseline[-100:].mean():.2f}\n")
167
168 print("Training ABLATION (single Q-learning) ...")
169 ablation = train("single")
170 print(f" -> last-100-episode average: {ablation[-100:].mean():.2f}\n")
171
172 # -----
173 # Comparison plot (50-episode moving averages on the same axes)
174 # -----
175 window = 50
176 def moving_avg(x):
177     return np.convolve(x, np.ones(window) / window, mode='valid')
178
179 x_axis = np.arange(window - 1, MAX_ITER)
180
181 plt.figure(figsize=(10, 6))
182 plt.plot(baseline, alpha=0.18, color='C0')
183 plt.plot(ablation, alpha=0.18, color='C1')
184 plt.plot(x_axis, moving_avg(baseline), color='C0', linewidth=2.2,
185         label=f'Baseline: Double Q-learning '
186              f'f'(last-100 avg = {baseline[-100:].mean():.1f})')
187 plt.plot(x_axis, moving_avg(ablation), color='C1', linewidth=2.2,
188         label=f'Ablation: single Q-learning '
189              f'f'(last-100 avg = {ablation[-100:].mean():.1f})')
190 plt.xlabel('Episode')
191 plt.ylabel('Episode duration (timesteps)')
192 plt.title('CartPole-v1: Double Q-learning vs. single Q-learning\n'
193         '(same seed, same hyperparameters, 1000 episodes)')
194 plt.legend(loc='upper left')
195 plt.grid(alpha=0.3)
196 plt.tight_layout()
197 plt.savefig('tabular_dq_ablation.png', dpi=130)
198
199 print("=" * 64)
200 print(f"Baseline (Double Q) last-100 avg: {baseline[-100:].mean():.2f}")
201 print(f"Ablation (single Q) last-100 avg: {ablation[-100:].mean():.2f}")
202 print("Saved comparison plot to tabular_dq_ablation.png")
203 print("=" * 64)

```

I replaced the Double Q-learning update with the standard single-table Q-learning update, $Q[s][a] \leftarrow Q[s][a] + \alpha(r + \gamma \max_{a'} Q[s'][a'] - Q[s][a])$, while keeping every other hyperparameter identical. The single-Q agent learns more slowly and plateaus around a 50-episode moving average of around 135 timesteps, while the Double-Q baseline continues climbing to around 185, which is about a 23% drop in final performance. The most likely reason is maximization bias: with a single table, the bootstrap target $\max_{a'} Q[s'][a']$ uses the same noisy estimates both to select the next action and to evaluate it, which produces a systematic positive bias when the Q-values still have substantial estimation noise. Double Q-learning decorrelates the two operations, i.e., one table picks the action, the other evaluates it, which removes most of that bias and yields more reliable bootstrap targets. With only 1000 episodes and $\alpha = 0.001$, the Q-tables are still quite noisy throughout training, which is exactly the regime where the double-table trick helps the most, so the gap we observe is consistent with the theory.

CartPole-v1: Double Q-learning vs. single Q-learning
(same seed, same hyperparameters, 1000 episodes)



3.2 Task 2: DQN on LunarLander

3.2.1 Background

Background provided in Homework document.

3.2.2 Implementation requirements

Code:

```
1 #!/usr/bin/env python3
2 # -*- coding: utf-8 -*-
3 """
4 dqn_lander.py
5
6 Deep Q-Network on LunarLander-v3.
7
8 Usage:
9     python dqn_lander.py                # train
10    python dqn_lander.py --mode record    # roll out trained agent to GIF
11    python dqn_lander.py --mode both     # train and then record
12
13 References:
14     Lecture slides 83-95 (Kak, RL).
15     Mnih et al., "Human-level control through deep RL", Nature 2015.
16 """
17
18 import argparse
19 import random
20 from collections import deque
21
22 import numpy as np
23 import torch
24 import torch.nn as nn
25 import torch.nn.functional as F
```

```

26 import torch.optim as optim
27 import matplotlib.pyplot as plt
28
29
30 class QNet(nn.Module):
31
32     def __init__(self, state_dim: int = 8, action_dim: int = 4):
33         super().__init__()
34         self.net = nn.Sequential(
35             nn.Linear(state_dim, 128), nn.ReLU(),
36             nn.Linear(128, 128), nn.ReLU(),
37             nn.Linear(128, action_dim),
38         )
39
40     def forward(self, s: torch.Tensor) -> torch.Tensor:
41         return self.net(s)
42
43
44 def select_action(state, model):
45     """
46     Args:
47         state : np.ndarray of shape (8,)
48         model : a QNet instance
49
50     Returns:
51         int in {0, 1, 2, 3}
52     """
53     model.eval()
54     with torch.no_grad():
55         s = torch.as_tensor(state, dtype=torch.float32).unsqueeze(0)
56         q = model(s)
57         return int(q.argmax(dim=1).item())
58
59
60 # =====
61 # Replay buffer
62 # =====
63 class ReplayBuffer:
64     """FIFO buffer storing (s, a, r, s', terminated) tuples."""
65
66     def __init__(self, capacity: int = 100_000):
67         self.buf = deque(maxlen=capacity)
68
69     def push(self, s, a, r, s_next, terminated):
70         self.buf.append((s, a, r, s_next, terminated))
71
72     def sample(self, batch_size: int):
73         batch = random.sample(self.buf, batch_size)
74         s, a, r, s_next, term = zip(*batch)
75         return (
76             torch.as_tensor(np.array(s), dtype=torch.float32),
77             torch.as_tensor(a, dtype=torch.long).unsqueeze(1),
78             torch.as_tensor(r, dtype=torch.float32).unsqueeze(1),
79             torch.as_tensor(np.array(s_next), dtype=torch.float32),
80             torch.as_tensor(term, dtype=torch.float32).unsqueeze(1),
81         )
82
83     def __len__(self):
84         return len(self.buf)
85
86
87 # =====
88 # Training
89 # =====
90 def train(args):
91     import gymnasium as gym

```

```

92 # ---- reproducibility ----
93 random.seed(args.seed)
94 np.random.seed(args.seed)
95 torch.manual_seed(args.seed)
96
97
98 env = gym.make(args.env)
99 env.action_space.seed(args.seed)          # gymnasium action-space RNG
100
101 state_dim = env.observation_space.shape[0]
102 action_dim = env.action_space.n
103 assert state_dim == 8 and action_dim == 4, \
104     "This script targets LunarLander (8-d state, 4 discrete actions)."
```

```

105
106 device = torch.device("cuda" if torch.cuda.is_available() else "cpu")
107 print(f"device           : {device}")
108 print(f"env              : {args.env}")
109 print(f"episodes          : {args.episodes}")
110 print(f"gamma              : {args.gamma}")
111 print(f"lr                 : {args.lr}")
112 print(f"buffer size        : {args.buffer_size}")
113 print(f"batch size         : {args.batch_size}")
114 print(f"target update      : Polyak tau = {args.tau} every gradient step")
115 print(f"epsilon decay      : 1.00 -> 0.05 linearly over first "
116       f"{args.eps_decay_episodes} episodes")
117 print()
118
119 # ---- networks ----
120 online = QNet(state_dim, action_dim).to(device)
121 target = QNet(state_dim, action_dim).to(device)
122 target.load_state_dict(online.state_dict())
123 for p in target.parameters():
124     p.requires_grad = False
125
126 optimizer = optim.Adam(online.parameters(), lr=args.lr)
127 buffer = ReplayBuffer(args.buffer_size)
128
129 EPS_START, EPS_END = 1.0, 0.05
130 episode_returns = []
131 best_avg100 = -float('inf')
132 total_steps = 0
133 solved = False
134 solved_episode = None
135
136 for episode in range(args.episodes):
137     frac = min(1.0, episode / max(1, args.eps_decay_episodes))
138     epsilon = EPS_START + frac * (EPS_END - EPS_START)
139
140     state, _ = env.reset(seed=args.seed if episode == 0 else None)
141     done = False
142     ep_return = 0.0
143
144     while not done:
145         # ---- epsilon-greedy action ----
146         if random.random() < epsilon:
147             action = env.action_space.sample()
148         else:
149             with torch.no_grad():
150                 s = torch.as_tensor(state, dtype=torch.float32,
151                                     device=device).unsqueeze(0)
152                 action = int(online(s).argmax(dim=1).item())
153
154     next_state, reward, terminated, truncated, _ = env.step(action)
155     done = terminated or truncated
156
157     buffer.push(state, action, reward, next_state, float(terminated))

```

158
159
160
161
162
163
164
165
166
167
168
169
170
171
172
173
174
175
176
177
178
179
180
181
182
183
184
185
186
187
188
189
190
191
192
193
194
195
196
197
198
199
200
201
202
203
204
205
206
207
208
209
210
211
212
213
214
215
216
217
218
219
220
221
222
223

```
state      = next_state
ep_return += reward
total_steps += 1

# ---- gradient step ----
if (len(buffer) >= args.warmup
    and total_steps % args.train_every == 0):

    s_b, a_b, r_b, s2_b, term_b = buffer.sample(args.batch_size)
    s_b = s_b.to(device)
    a_b = a_b.to(device)
    r_b = r_b.to(device)
    s2_b = s2_b.to(device)
    term_b = term_b.to(device)

    # current Q(s, a)
    q_sa = online(s_b).gather(1, a_b)

    # TD target y = r + gamma * (1 - terminated) * max_a' Q_tgt(s', a')
    with torch.no_grad():
        q_next = target(s2_b).max(dim=1, keepdim=True)[0]
        y = r_b + args.gamma * (1.0 - term_b) * q_next

    loss = F.smooth_l1_loss(q_sa, y)

    optimizer.zero_grad()
    loss.backward()
    nn.utils.clip_grad_norm_(online.parameters(), max_norm=10.0)
    optimizer.step()

    with torch.no_grad():
        for p, p_t in zip(online.parameters(),
                           target.parameters()):
            p_t.data.mul_(1.0 - args.tau)
            p_t.data.add_(args.tau * p.data)

# ---- end of episode ----
episode_returns.append(ep_return)
avg100 = float(np.mean(episode_returns[-100:]))

if (episode + 1) % 10 == 0:
    print(f"ep {episode+1:4d} | return {ep_return:8.2f} "
          f"| avg100 {avg100:8.2f} | eps {epsilon:.3f} "
          f"| buf {len(buffer):6d} | steps {total_steps}")

# checkpoint best running average
if len(episode_returns) >= 100 and avg100 > best_avg100:
    best_avg100 = avg100
    torch.save(online.state_dict(), args.checkpoint)

if (not solved and len(episode_returns) >= 100 and avg100 >= 200.0):
    solved = True
    solved_episode = episode + 1
    print(f"\n*** Solved at episode {solved_episode} "
          f"(avg100 = {avg100:.2f}) ***\n")
    if args.stop_when_solved:
        break

env.close()

last100 = float(np.mean(episode_returns[-100:]))
if last100 >= best_avg100:
    torch.save(online.state_dict(), args.checkpoint)
print(f"\nSaved checkpoint to {args.checkpoint}")
```

```

224 # -----
225 # Plot training curve with 100-episode moving average
226 # -----
227 returns = np.asarray(episode_returns)
228 window = 100
229 if len(returns) >= window:
230     ma = np.convolve(returns, np.ones(window) / window, mode='valid')
231     ma_x = np.arange(window - 1, len(returns))
232 else:
233     ma, ma_x = np.array([]), np.array([])
234
235 status = "SOLVED" if last100 >= 200.0 else "not solved"
236
237 plt.figure(figsize=(10, 6))
238 plt.plot(returns, alpha=0.3, label='Episode return')
239 if len(ma):
240     plt.plot(ma_x, ma, linewidth=2.2, label='100-episode moving average')
241 plt.axhline(200, linestyle='--', color='gray', alpha=0.6,
242            label='Solved threshold (+200)')
243 plt.xlabel('Episode')
244 plt.ylabel('Episode return')
245 plt.title(f'DQN on {args.env}\n'
246          f'Last-100-episode avg: {last100:.2f} ({status})')
247 plt.legend(loc='lower right')
248 plt.grid(alpha=0.3)
249 plt.tight_layout()
250 plt.savefig(args.plot, dpi=130)
251 print(f"Saved training curve to {args.plot}")
252
253 print("=" * 64)
254 print(f"Final 100-episode average return: {last100:.2f} ({status})")
255 if solved:
256     print(f"First solved at episode {solved_episode}")
257 print("=" * 64)
258
259 # =====
260 # Video recording of a trained episode
261 # =====
262
263 def record_video(args):
264     import gymnasium as gym
265     import imageio
266
267     env = gym.make(args.env, render_mode="rgb_array")
268     model = QNet().to(torch.device("cpu"))
269     model.load_state_dict(torch.load(args.checkpoint, map_location="cpu"))
270     model.eval()
271
272     frames = []
273     state, _ = env.reset(seed=args.seed)
274     done, total = False, 0.0
275     while not done:
276         action = select_action(state, model)
277         state, r, terminated, truncated, _ = env.step(action)
278         frames.append(env.render())
279         total += r
280         done = terminated or truncated
281     env.close()
282
283     imageio.mimsave(args.gif, frames, fps=30)
284     print(f"Saved trained-agent rollout to {args.gif} (return = {total:.2f})")
285
286
287 def get_parser():
288     p = argparse.ArgumentParser(description="DQN on LunarLander")
289     p.add_argument('--env', type=str, default='LunarLander-v3',

```

```

290         help='Use LunarLander-v2 if v3 is not available')
291 p.add_argument('--mode', type=str, default='train',
292               choices=['train', 'record', 'both'])
293
294 # training hyperparameters
295 p.add_argument('--episodes', type=int, default=1000)
296 p.add_argument('--gamma', type=float, default=0.99)
297 p.add_argument('--lr', type=float, default=5e-4)
298 p.add_argument('--batch_size', type=int, default=64)
299 p.add_argument('--buffer_size', type=int, default=100_000)
300 p.add_argument('--warmup', type=int, default=1_000,
301               help='Min transitions in buffer before training begins')
302 p.add_argument('--train_every', type=int, default=4,
303               help='Take a gradient step every N env steps')
304 p.add_argument('--tau', type=float, default=0.005,
305               help='Polyak averaging coefficient for target network')
306 p.add_argument('--eps_decay_episodes', type=int, default=300,
307               help='Linearly anneal eps from 1.0 to 0.05 over this many '
308               'episodes (10-30% of training)')
309 p.add_argument('--stop_when_solved', action='store_true',
310               help='Stop training as soon as the env is solved')
311
312 # bookkeeping
313 p.add_argument('--seed', type=int, default=0)
314 p.add_argument('--checkpoint', type=str, default='dqn_lander.pt')
315 p.add_argument('--plot', type=str,
316               default='dqn_lander_training_curve.png')
317 p.add_argument('--gif', type=str, default='dqn_lander.gif')
318 return p
319
320
321 if __name__ == "__main__":
322     args = get_parser().parse_args()
323     if args.mode in ('train', 'both'):
324         train(args)
325     if args.mode in ('record', 'both'):
326         record_video(args)

```

I used option (b): Polyak averaging with $\tau = 0.005$, applied after every gradient step. The code that implements it in `train()`:

```

1         with torch.no_grad():
2             for p, p_t in zip(online.parameters(),
3                               target.parameters()):
4                 p_t.data.mul_(1.0 - args.tau)
5                 p_t.data.add_(args.tau * p.data)

```

Implementation:

```
1 !python dqn_lander.py
```

Output:

```

1 device           : cuda
2 env              : LunarLander-v3
3 episodes         : 1000
4 gamma            : 0.99
5 lr               : 0.0005
6 buffer size      : 100000
7 batch size       : 64
8 target update    : Polyak tau = 0.005 every gradient step
9 epsilon decay    : 1.00 -> 0.05 linearly over first 300 episodes
10
11 ep   10 | return  -272.00 | avg100  -209.22 | eps 0.972 | buf   971 | steps 971
12 ep   20 | return  -259.54 | avg100  -191.44 | eps 0.940 | buf  1889 | steps 1889
13 ep   30 | return  -101.58 | avg100  -194.78 | eps 0.908 | buf  2892 | steps 2892
14 ...
15 ...

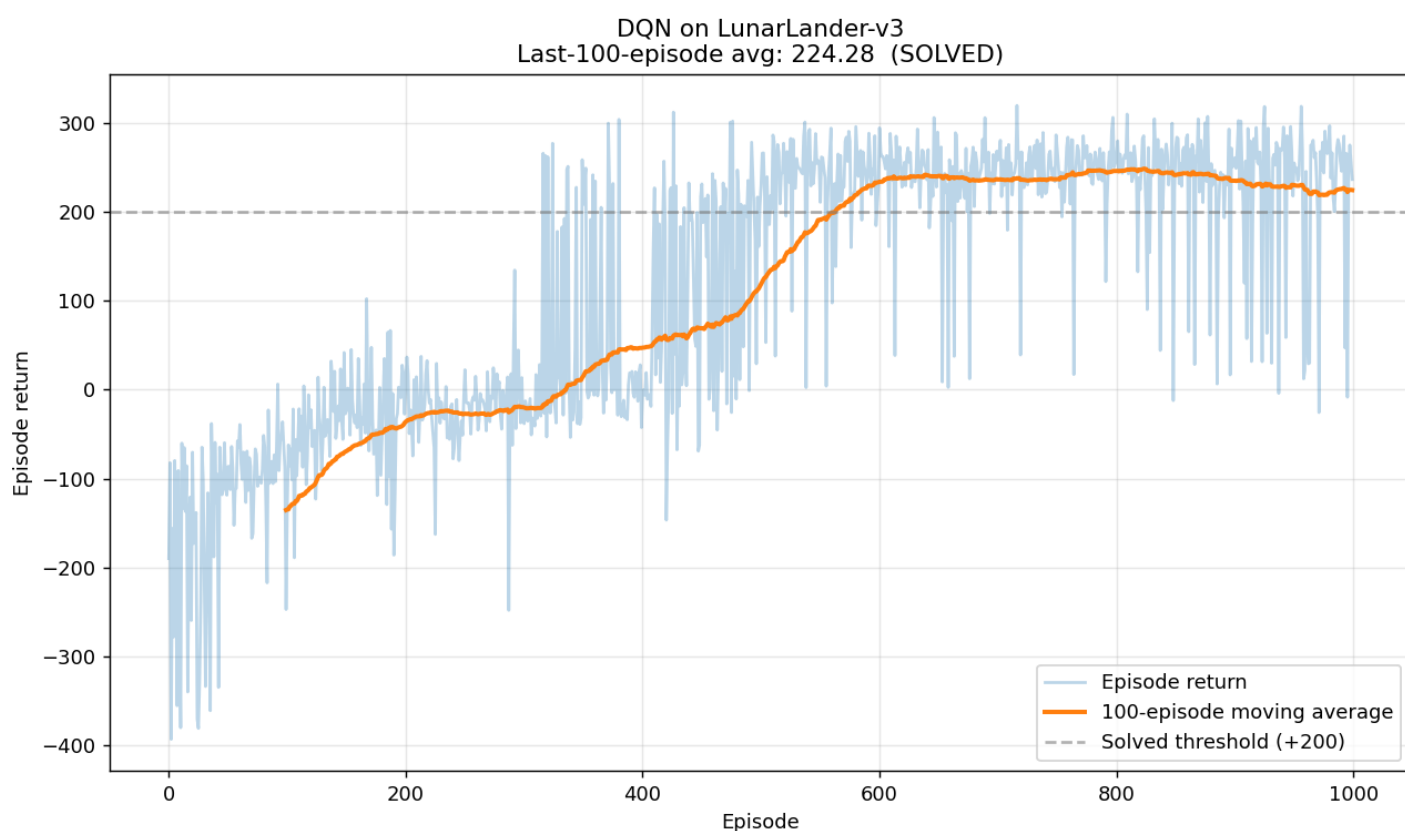
```

```

16 ep 540 | return 234.76 | avg100 175.79 | eps 0.050 | buf 100000 | steps 304355
17 ep 550 | return 255.69 | avg100 190.89 | eps 0.050 | buf 100000 | steps 308232
18 ep 560 | return 293.89 | avg100 198.03 | eps 0.050 | buf 100000 | steps 312223
19
20 *** Solved at episode 564 (avg100 = 200.91) ***
21
22 ep 570 | return 228.26 | avg100 206.80 | eps 0.050 | buf 100000 | steps 318193
23 ep 580 | return 223.63 | avg100 218.28 | eps 0.050 | buf 100000 | steps 322994
24 ...
25 ...
26 ep 990 | return 256.91 | avg100 225.43 | eps 0.050 | buf 100000 | steps 493556
27 ep 1000 | return 236.30 | avg100 224.28 | eps 0.050 | buf 100000 | steps 496811
28
29 Saved checkpoint to dqn_lander.pt
30 Saved training curve to dqn_lander_training_curve.png
31 =====
32 Final 100-episode average return: 224.28 (SOLVED)
33 First solved at episode 564
34 =====

```

Training curve showing episode return vs. episode number. Overlay a moving average over 100 episodes.



3.2.3 Brief analysis

Question: Why a target network? What would go wrong if you used the same Q_ϕ to produce both the predicted value and the TD target? Frame your answer in terms of what happens to the regression target during a single gradient step.

Answer: In a single gradient step, the regression target $y = r + \gamma \max_{a'} Q_\phi(s', a')$ is itself a function of the parameters ϕ being updated, so when the optimizer nudges $Q_\phi(s, a)$ toward y it simultaneously shifts y in (typically) the same direction, i.e. the target chases the prediction instead of standing still. This positive feedback loop amplifies any over-estimation in $Q_\phi(s', a')$, especially because the max operator already biases the bootstrap target upward by selecting the noisiest action's value. Holding a separate, slowly-updated copy Q_{ϕ^-} for the target makes the regression target a near-constant within each step, which is what gives gradient descent a stable thing to fit.

Question: ϵ -greedy decay. You annealed ϵ from 1.0 to 0.05 over training. Why not start at $\epsilon = 0.05$? Why not stay at $\epsilon = 1.0$?

Answer: Starting at $\epsilon = 0.05$ means the agent commits almost immediately to the argmax of an untrained, essentially random Q-network, so it would lock onto whatever bad action the initialization happens to favor and never visit enough of the state space (the landing pad, leg-contact transitions, the +100 success reward) to learn anything. The replay buffer would fill with garbage from one tiny region. Staying at $\epsilon = 1.0$ fixes the exploration problem but wastes the learning: every action is a coin flip regardless of how good Q has become, so the agent never actually exploits its policy and never sees the long, well-controlled landing trajectories needed to drive the 100-episode average toward +200. Annealing gives us both – broad exploration while the Q-values are still noise, then a gradual handoff to the learned greedy policy as those values become trustworthy, with the 0.05 floor preserving just enough exploration to keep escaping local optima late in training.

3.3 Task 3: Reward modeling on real human preferences

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

```
1 from lm_utils import load_preference_dataset
2
3 triples = load_preference_dataset(n_examples=500)
4 print(f"loaded {len(triples)} triples")
5
6 for i in range(3):
7     t = triples[i]
8     print(f"\n=== Example {i+1} ===")
9     print(f"POST:      {t['post']!r}")
10    print(f"CHOSEN:     {t['chosen']!r}")
11    print(f"REJECTED: {t['rejected']!r}")
```

```
1 loaded 500 triples
2
3 === Example 1 ===
4 POST:      "My boyfriend and I are long distance. We have a trip planned this summer which
5 involves me going over to him in the USA. This will be the second time I have actually been
6 with him in person. I am flying from the UK with my mum to the east coast. The original plan
7 was for me to fly over to my boyfriend in the west coast (my parents are holidaying on the
8 east coast) but because my mum was freaking out so much about me going to meet my boyfriend i
9 said we can all road trip there together. I even invited her on the trip with us. I have given
10 her all of our dates so that she can travel around with us.\n\nThe plan was for me to stay on
11 the 4th July and fly back on the 5th. Mum knew this. I told her I had booked a flight back
already from the west coast to east coast (where she would pick me up and we would fly back to
the UK together). She has gone mad at me because she can't believe I would book a flight when
she told me she didn't want me flying on my own. At the time I had booked it she told me she
wasn't gonna road trip with us. She knew the trip was happening.....how else was I to get
home if I don't fly? \n\nI am fine flying on my own it doesn't bother me at all. I feel like I
have done everything I can to make her feel comfortable with this trip and she is just trying
to sabotage it. Thoughts??"
5 CHOSEN:     ' I have made sure my mother is comfortable with my boyfriend travelling on a trip and
6 now my mother is mad because I booked it.'
7 REJECTED:   ' Mum is mad at me for not flying on my own trip to meet my boyfriend.'
8
9 === Example 2 ===
10 POST:      "My boyfriend and I are long distance. We have a trip planned this summer which
11 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 flying from the UK with my mum to the east coast. The original plan
```

```

was for me to fly over to my boyfriend in the west coast (my parents are holidaying on the
east coast) but because my mum was freaking out so much about me going to meet my boyfriend i
said we can all road trip there together. I even invited her on the trip with us. I have given
her all of our dates so that she can travel around with us.\n\nThe plan was for me to stay on
the 4th July and fly back on the 5th. Mum knew this. I told her I had booked a flight back
already from the west coast to east coast (where she would pick me up and we would fly back to
the UK together). She has gone mad at me because she can't believe I would book a flight when
she told me she didn't want me flying on my own. At the time I had booked it she told me she
wasn't gonna road trip with us. She knew the trip was happening.....how else was I to get
home if I don't fly? \n\nI am fine flying on my own it doesn't bother me at all. I feel like I
have done everything I can to make her feel comfortable with this trip and she is just trying
to sabotage it. Thoughts??"
10 CHOSEN: " mum isn't speaking to me because I booked a flight and she doesn't want me flying on
my own."
11 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.'
12
13 === Example 3 ===
14 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 flying from the UK with my mum to the east coast. The original plan
was for me to fly over to my boyfriend in the west coast (my parents are holidaying on the
east coast) but because my mum was freaking out so much about me going to meet my boyfriend i
said we can all road trip there together. I even invited her on the trip with us. I have given
her all of our dates so that she can travel around with us.\n\nThe plan was for me to stay on
the 4th July and fly back on the 5th. Mum knew this. I told her I had booked a flight back
already from the west coast to east coast (where she would pick me up and we would fly back to
the UK together). She has gone mad at me because she can't believe I would book a flight when
she told me she didn't want me flying on my own. At the time I had booked it she told me she
wasn't gonna road trip with us. She knew the trip was happening.....how else was I to get
home if I don't fly? \n\nI am fine flying on my own it doesn't bother me at all. I feel like I
have done everything I can to make her feel comfortable with this trip and she is just trying
to sabotage it. Thoughts??"
15 CHOSEN: " mum isn't speaking to me because I booked a flight and she doesn't want me flying on
my own."
16 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 chosen summaries are visibly better, particularly the one favored in Examples 2 and 3. They improve primarily in factual accuracy and hallucination reduction. The lower-ranked summaries either completely misinterpret the text (such as Example 1's chosen summary claiming the boyfriend is the one traveling) or invent entirely false details (such as Example 3's rejected summary claiming the author canceled the flight). The ultimately preferred summary ("mum isn't speaking to me because I booked a flight and she doesn't want me flying on my own") is the only one that more or less summarizes the actual core conflict without fabricating information.

3.3.6 Embed everything

```

1 from lm_utils import load_gpt2, embed_triples
2
3 tokenizer, base, device = load_gpt2()
4 embedded = embed_triples(triples, tokenizer, base, device)
5 print(f"embedded {len(embedded)} triples")
6 print(f"post embedding shape: {embedded[0]['post_emb'].shape}")
7
8 # 80/20 train/validation split
9 n_train = int(0.8 * len(embedded))
10 train_set = embedded[:n_train]
11 val_set = embedded[n_train:]
12 print(f"train set: {len(train_set)} triples")
13 print(f"val set: {len(val_set)} triples")

```

```

1 embedded 50/500
2 embedded 100/500
3 embedded 150/500

```

```

4   embedded 200/500
5   embedded 250/500
6   embedded 300/500
7   embedded 350/500
8   embedded 400/500
9   embedded 450/500
10  embedded 500/500
11 embedded 500 triples
12 post embedding shape: torch.Size([768])
13 train set: 400 triples
14 val set: 100 triples

```

3.3.7 Train the RewardNetwork

Code:

```

1  #!/usr/bin/env python3
2  # -*- coding: utf-8 -*-
3  """
4  train_reward_model.py
5
6  HW11 Task 3.3: train Prof. Kak's RewardNetwork (Slide 35 of the RL lecture)
7  on the OpenAI summarize_from_feedback preference dataset.
8  """
9
10 from __future__ import annotations
11
12 import argparse
13 import random
14
15 import numpy as np
16 import torch
17 import torch.nn as nn
18 import torch.nn.functional as F
19 import matplotlib.pyplot as plt
20
21 from lm_utils import (
22     load_preference_dataset,
23     load_gpt2,
24     embed_triples,
25 )
26
27 class RewardNetwork(torch.nn.Module):
28     def __init__(self, embedding_size):
29         torch.nn.Module.__init__(self)
30         self.linear1 = nn.Linear(3 * embedding_size, 2 * embedding_size)
31         self.linear2 = nn.Linear(2 * embedding_size, embedding_size)
32         self.linear3 = nn.Linear(embedding_size, 1)
33         self.sigmoid = nn.Sigmoid()
34
35     def forward(self, xSEPy):
36         x = self.linear1(xSEPy).clamp(min=0)
37         x = self.linear2(x).clamp(min=0)
38         x = self.linear3(x)
39         r = self.sigmoid(x)
40         return r
41
42
43 def make_input(post_emb, sep_emb, response_emb):
44     """
45     Concatenate [post, SEP, response] into a single (3 * 768,) = (2304,)
46     tensor in float32, matching the input shape RewardNetwork expects.
47     """
48     return torch.cat([post_emb, sep_emb, response_emb], dim=0).float()
49
50

```

```

51 # =====
52 # Training
53 # =====
54 def train_reward_model(train_set,
55                         embedding_size: int = 768,
56                         n_steps: int = 2000,
57                         lr: float = 1e-3,
58                         seed: int = 0):
59     random.seed(seed)
60     torch.manual_seed(seed)
61     np.random.seed(seed)
62
63     model = RewardNetwork(embedding_size)
64     optimizer = torch.optim.Adam(model.parameters(), lr=lr)
65
66     losses = []
67     for step in range(n_steps):
68         triple = random.choice(train_set)
69
70         # Build the two 2304-dim inputs.
71         x_chosen = make_input(triple['post_emb'],
72                               triple['sep_emb'],
73                               triple['chosen_emb'])
74         x_rejected = make_input(triple['post_emb'],
75                                 triple['sep_emb'],
76                                 triple['rejected_emb'])
77
78         # Forward: scalar reward for each response.
79         r_chosen = model(x_chosen).squeeze()
80         r_rejected = model(x_rejected).squeeze()
81
82         # Bradley-Terry pairwise loss (Eq. 6 specialized to two responses).
83         loss = -F.logsigmoid(r_chosen - r_rejected)
84
85         optimizer.zero_grad()
86         loss.backward()
87         optimizer.step()
88
89         losses.append(loss.item())
90
91         if (step + 1) % 200 == 0:
92             recent = float(np.mean(losses[-200:]))
93             print(f" step {step+1:4d} | loss {loss.item():.4f} "
94                   f"| 200-step avg {recent:.4f}")
95
96     return model, losses
97
98
99 # =====
100 # Validation
101 # =====
102 @torch.no_grad()
103 def evaluate(model, val_set):
104     """Return fraction of validation triples with r_chosen > r_rejected."""
105     model.eval()
106     correct = 0
107     for t in val_set:
108         x_c = make_input(t['post_emb'], t['sep_emb'], t['chosen_emb'])
109         x_r = make_input(t['post_emb'], t['sep_emb'], t['rejected_emb'])
110         if model(x_c).item() > model(x_r).item():
111             correct += 1
112     model.train()
113     return correct / len(val_set)
114
115
116 # =====

```

```

117 # Main
118 # =====
119 def main(args):
120     # ---- Load preference triples -----
121     print(f"Loading {args.n_examples} preference triples ...")
122     triples = load_preference_dataset(n_examples=args.n_examples)
123     print(f"   loaded {len(triples)} triples\n")
124
125     # ---- Embed everything with frozen distilgpt2 -----
126     print("Embedding posts and summaries with distilgpt2 ...")
127     tokenizer, base, device = load_gpt2()
128     embedded = embed_triples(triples, tokenizer, base, device)
129     print(f"   embedded {len(embedded)} triples")
130     print(f"   post embedding shape: {embedded[0]['post_emb'].shape}\n")
131
132     # ---- 80/20 split (no shuffle: keep duplicate-prompt rows together) ---
133     n_train = int(0.8 * len(embedded))
134     train_set = embedded[:n_train]
135     val_set = embedded[n_train:]
136     print(f"   train: {len(train_set)} | val: {len(val_set)}\n")
137
138     # ---- Train -----
139     print(f"Training RewardNetwork for {args.n_steps} steps "
140           f"(Adam, lr={args.lr}, batch size 1) ...")
141     model, losses = train_reward_model(
142         train_set,
143         embedding_size=embedded[0]['post_emb'].shape[0], # 768
144         n_steps=args.n_steps,
145         lr=args.lr,
146         seed=args.seed,
147     )
148
149     # ---- Validation accuracy -----
150     val_acc = evaluate(model, val_set)
151
152     # ---- Print -----
153     initial_loss = losses[0]
154     final_loss = float(np.mean(losses[-200:]))
155
156     print("\n" + "=" * 64)
157     print(f"Initial loss (step 1): {initial_loss:.4f}")
158     print(f"Final loss (last-200-step average): {final_loss:.4f}")
159     print(f"Validation accuracy (held-out 20%): {val_acc*100:.2f}% "
160           f"(random = 50%)")
161     print("=" * 64)
162
163     # ---- Save checkpoint -----
164     torch.save(model.state_dict(), args.checkpoint)
165     print(f"\nSaved reward model to {args.checkpoint}")
166
167     # ---- Loss curve -----
168     losses_arr = np.asarray(losses)
169     window = 200
170     if len(losses_arr) >= window:
171         ma = np.convolve(losses_arr, np.ones(window) / window, mode='valid')
172         ma_x = np.arange(window - 1, len(losses_arr))
173     else:
174         ma, ma_x = np.array([]), np.array([])
175
176     plt.figure(figsize=(10, 6))
177     plt.plot(losses_arr, alpha=0.3, label='Per-step loss')
178     if len(ma):
179         plt.plot(ma_x, ma, linewidth=2.2,
180                  label=f'{window}-step moving average')
181     plt.axhline(np.log(2), linestyle='--', color='gray', alpha=0.6,
182                 label=r'$\log 2 \approx 0.693$ ')

```

```

183         r'(loss when $r_{chosen}=r_{rejected}$)')
184 plt.xlabel('Gradient step')
185 plt.ylabel(r'Pairwise loss $-\log \sigma(r_{chosen}-r_{rejected})$')
186 plt.title('RewardNetwork training on summarize_from_feedback\n'
187          f'initial loss {initial_loss:.4f} ',
188          f'final-200 avg {final_loss:.4f} ',
189          f'val acc {val_acc*100:.1f}%')
190 plt.legend(loc='upper right')
191 plt.grid(alpha=0.3)
192 plt.tight_layout()
193 plt.savefig(args.plot, dpi=130)
194 print(f"Saved loss curve to {args.plot}")
195
196
197 def get_parser():
198     p = argparse.ArgumentParser(
199         description="Train RewardNetwork on summarize_from_feedback")
200     p.add_argument('--n_examples', type=int, default=500)
201     p.add_argument('--n_steps', type=int, default=2000)
202     p.add_argument('--lr', type=float, default=1e-3)
203     p.add_argument('--seed', type=int, default=0)
204     p.add_argument('--checkpoint', type=str, default='reward_model.pt')
205     p.add_argument('--plot', type=str,
206                   default='reward_model_loss.png')
207     return p
208
209
210 if __name__ == "__main__":
211     main(get_parser().parse_args())

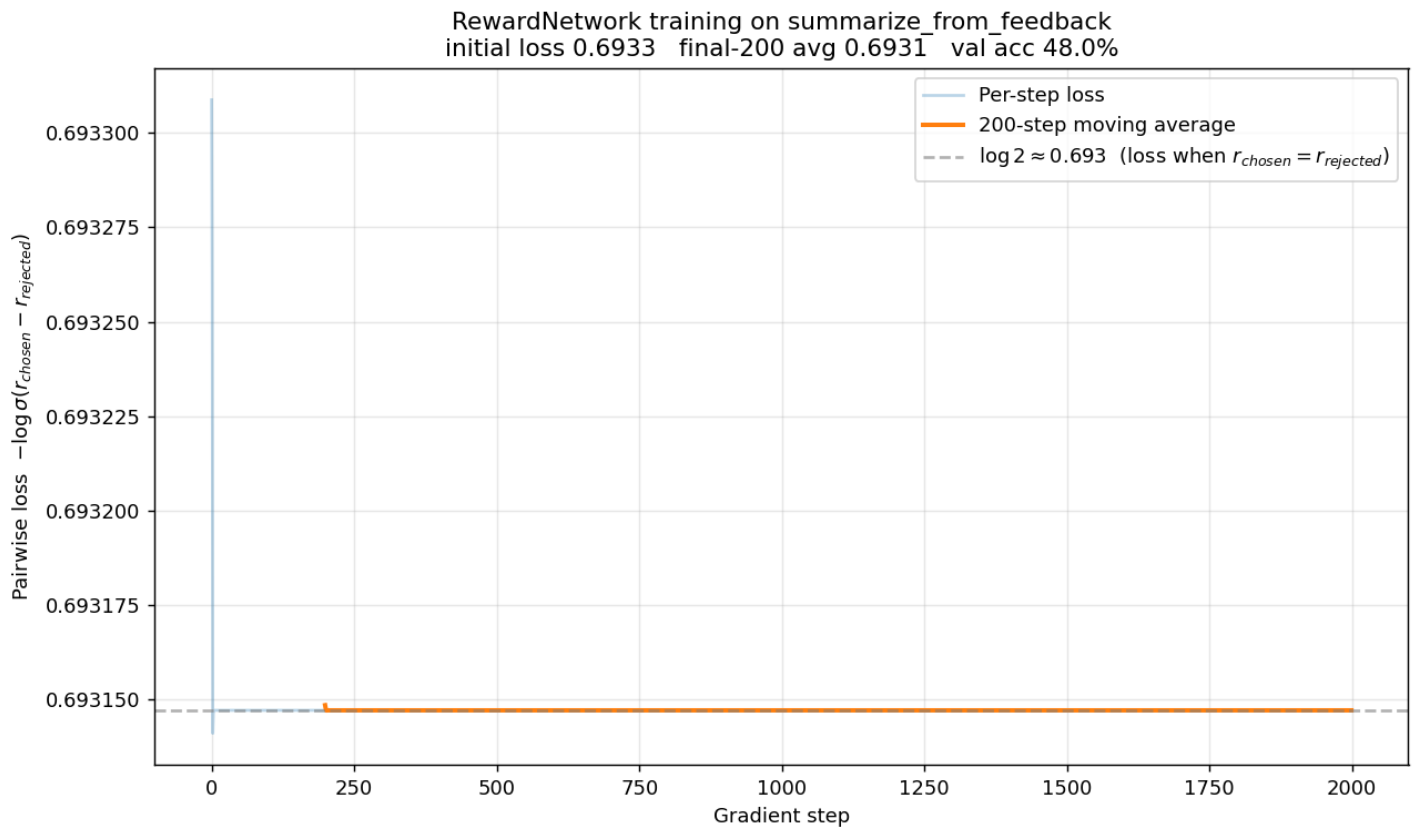
```

Output:

```

1  embedded 50/500
2  embedded 100/500
3  embedded 150/500
4  embedded 200/500
5  embedded 250/500
6  embedded 300/500
7  embedded 350/500
8  embedded 400/500
9  embedded 450/500
10 embedded 500/500
11 embedded 500 triples
12 post embedding shape: torch.Size([768])
13
14 train: 400 | val: 100
15
16 Training RewardNetwork for 2000 steps (Adam, lr=0.001, batch size 1) ...
17 step 200 | loss 0.6931 | 200-step avg 0.6931
18 step 400 | loss 0.6931 | 200-step avg 0.6931
19 step 600 | loss 0.6931 | 200-step avg 0.6931
20 step 800 | loss 0.6931 | 200-step avg 0.6931
21 step 1000 | loss 0.6931 | 200-step avg 0.6931
22 step 1200 | loss 0.6931 | 200-step avg 0.6931
23 step 1400 | loss 0.6931 | 200-step avg 0.6931
24 step 1600 | loss 0.6931 | 200-step avg 0.6931
25 step 1800 | loss 0.6931 | 200-step avg 0.6931
26 step 2000 | loss 0.6931 | 200-step avg 0.6931
27
28 =====
29 Initial loss (step 1): 0.6933
30 Final loss (last-200-step average): 0.6931
31 Validation accuracy (held-out 20%): 48.00% (random = 50%)
32 =====
33
34 Saved reward model to reward_model.pt
35 Saved loss curve to reward_model_loss.png

```



3.3.8 Evaluate and critique

1 Validation accuracy (held-out 20%): 48.00% (random = 50%)

Question: Report your validation accuracy and what happened to the training loss. How did the values you saw compare to what you would expect from a reward model that is learning to capture human preferences?

Answer: Validation accuracy was 48.0% which is indistinguishable from random guessing and the training loss only moved from 0.6933 to a final 200-step average of 0.6931, a change of approx 2×10^{-4} over 2000 gradient steps. A reward model that was actually learning human preferences would push the loss meaningfully below $\ln 2 \approx 0.693$ and reach higher than 50% validation accuracy on this dataset.

Question: Look carefully at the `RewardNetwork` class from Slide 35. Work out what values the forward pass can return: start with the three `clamp(min=0)` operations, then the final sigmoid, and describe the range of outputs. Given the pairwise loss $-\log \sigma(r_{\text{chosen}} - r_{\text{rejected}})$, how does this output range affect the loss values you can reach during training?

Answer: The two `clamp(min=0)` ops act as ReLUs and keep the hidden activations in $[0, \infty)$; `linear3` produces an unbounded scalar, and the final sigmoid squashes it to $(0, 1)$. So $r_{\text{chosen}} - r_{\text{rejected}} \in (-1, 1)$, which means the pairwise loss is bounded below by $-\log \sigma(1) \approx 0.313$ and starts near $-\log \sigma(0) = \log 2 \approx 0.693$ at initialization. The narrow window combined with sigmoid gradient saturation as the score difference approaches ± 1 leaves almost no room for the loss to descend, which is exactly why mine got pinned at approx 0.693.

Question: Suggest one modification to the `RewardNetwork` that you think would expand the range of its outputs and allow the loss to decrease further. You are not required to implement it or verify that it works; a short description is enough.

Answer: Removing the final `nn.Sigmoid()` and letting `linear3` output a raw, unbounded scalar will lead to $r_{\text{chosen}} - r_{\text{rejected}}$ ranging over all of $\mathbb{R}$, so the loss can approach 0 and gradients no longer would vanish near the bound.

4 Bonus Task: PPO on LunarLander

4.1 What to implement

4.2 Deliverables

```
1 """
2 ppo_lander.py
3
4 HW11 Bonus Task: Proximal Policy Optimization on LunarLander-v3.
5
6 Adapted from CleanRL's reference implementation 'cleanrl/ppo.py'
7 (Huang et al., JMLR 2022; https://github.com/vwxyzjn/cleanrl).
8 The structure below mirrors CleanRL closely.
9 - separate Actor and Critic networks rather than a shared trunk
10 - LunarLander-specific hyperparameters,
11 - logging is per-episode-return rather than wall-clock SPS,
12 - saves episode returns to a .npz so you can overlay the DQN training
13   curve from Task 2 on the same axes.
14
15 Citation:
16     Huang, Dossa, Raffin, Kanervisto, Wang.
17     "CleanRL: High-quality Single-file Implementations of Deep
18     Reinforcement Learning Algorithms." JMLR, 2022.
19     https://github.com/vwxyzjn/cleanrl
20
21 Usage:
22     python ppo_lander.py                # train
23     python ppo_lander.py --mode compare  # overlay DQN vs PPO curves
24     python ppo_lander.py --mode record   # roll out trained agent to GIF
25 """
26
27 from __future__ import annotations
28
29 import argparse
30 import os
31 import random
32 from dataclasses import dataclass
33
34 import numpy as np
35 import torch
36 import torch.nn as nn
37 import torch.optim as optim
38 from torch.distributions import Categorical
39 import matplotlib.pyplot as plt
40
41
42
43 # Networks
44 def layer_init(layer: nn.Linear, std: float = np.sqrt(2),
45               bias_const: float = 0.0) -> nn.Linear:
46     nn.init.orthogonal_(layer.weight, std)
47     nn.init.constant_(layer.bias, bias_const)
48     return layer
49
50
51 class Actor(nn.Module):
52     def __init__(self, state_dim: int = 8, action_dim: int = 4):
53         super().__init__()
54         self.net = nn.Sequential(
55             layer_init(nn.Linear(state_dim, 128)), nn.Tanh(),
56             layer_init(nn.Linear(128, 128)), nn.Tanh(),
57             layer_init(nn.Linear(128, action_dim), std=0.01),
58         )
59
60     def forward(self, s: torch.Tensor) -> torch.Tensor:
```

```

61         return self.net(s)                                # raw logits
62
63     def get_action_and_logprob(self, s: torch.Tensor,
64                               action: torch.Tensor = None):
65         """Sample (or evaluate) an action and return (action, logp, entropy)."""
66         logits = self.net(s)
67         dist = Categorical(logits=logits)
68         if action is None:
69             action = dist.sample()
70         return action, dist.log_prob(action), dist.entropy()
71
72
73 class Critic(nn.Module):
74     def __init__(self, state_dim: int = 8):
75         super().__init__()
76         self.net = nn.Sequential(
77             layer_init(nn.Linear(state_dim, 128)), nn.Tanh(),
78             layer_init(nn.Linear(128, 128)), nn.Tanh(),
79             layer_init(nn.Linear(128, 1), std=1.0),
80         )
81
82     def forward(self, s: torch.Tensor) -> torch.Tensor:
83         return self.net(s).squeeze(-1)
84
85
86
87 # Rollout buffer
88 @dataclass
89 class Rollout:
90     """Storage for one rollout of length T transitions."""
91     states: torch.Tensor          # (T, state_dim)
92     actions: torch.Tensor         # (T,)
93     logprobs: torch.Tensor        # (T,)
94     rewards: torch.Tensor         # (T,)
95     dones: torch.Tensor           # (T,) -- terminated only
96     values: torch.Tensor          # (T,)
97
98
99 def compute_gae(rollout: Rollout, last_value: torch.Tensor,
100                gamma: float, lam: float):
101     """
102     Generalized Advantage Estimation (Schulman 2016, Eq. 16).
103
104     
$$\begin{aligned} \delta_t &= r_t + \gamma * (1 - \text{done}_t) * V(s_{t+1}) - V(s_t) \\ A_t &= \delta_t + \gamma * \text{lam} * (1 - \text{done}_t) * A_{t+1} \\ R_t &= A_t + V(s_t) \quad \text{\# value-fn regression target} \end{aligned}$$

105
106     Returns (advantages, returns), each shape (T,).
107     """
108     T = rollout.rewards.shape[0]
109     advantages = torch.zeros_like(rollout.rewards)
110     last_gae = 0.0
111     for t in reversed(range(T)):
112         if t == T - 1:
113             next_value = last_value
114             next_nonterm = 1.0 - rollout.dones[t]
115         else:
116             next_value = rollout.values[t + 1]
117             next_nonterm = 1.0 - rollout.dones[t]
118         delta = (rollout.rewards[t]
119                 + gamma * next_value * next_nonterm
120                 - rollout.values[t])
121         last_gae = delta + gamma * lam * next_nonterm * last_gae
122         advantages[t] = last_gae
123     returns = advantages + rollout.values
124     return advantages, returns

```

```

127
128
129
130 # Training
131 def train(args):
132     import gymnasium as gym
133
134     # ---- reproducibility ----
135     random.seed(args.seed)
136     np.random.seed(args.seed)
137     torch.manual_seed(args.seed)
138     torch.backends.cudnn.deterministic = True
139
140     env = gym.make(args.env)
141     env.action_space.seed(args.seed)
142
143     state_dim = env.observation_space.shape[0]
144     action_dim = env.action_space.n
145     assert state_dim == 8 and action_dim == 4, \
146         "This script targets LunarLander (8-d state, 4 discrete actions)."
147
148     device = torch.device("cuda" if torch.cuda.is_available() else "cpu")
149     print(f"device          : {device}")
150     print(f"env              : {args.env}")
151     print(f"total timesteps      : {args.total_timesteps}")
152     print(f"rollout length       : {args.num_steps}")
153     print(f"gamma                 : {args.gamma}")
154     print(f"gae lambda           : {args.gae_lambda}")
155     print(f"clip coef            : {args.clip_coef}")
156     print(f"update epochs        : {args.update_epochs}")
157     print(f"minibatch size       : {args.minibatch_size}")
158     print(f"lr                    : {args.lr} (annealed: {args.anneal_lr})\n")
159
160     # ---- networks & optimizer ----
161     actor = Actor(state_dim, action_dim).to(device)
162     critic = Critic(state_dim).to(device)
163     optimizer = optim.Adam(
164         list(actor.parameters()) + list(critic.parameters()),
165         lr=args.lr, eps=1e-5, # CleanRL uses Adam eps=1e-5
166     )
167
168     # ---- rollout storage (reused each iteration) ----
169     states_buf = torch.zeros((args.num_steps, state_dim)).to(device)
170     actions_buf = torch.zeros((args.num_steps,), dtype=torch.long).to(device)
171     logprobs_buf = torch.zeros((args.num_steps,)).to(device)
172     rewards_buf = torch.zeros((args.num_steps,)).to(device)
173     dones_buf = torch.zeros((args.num_steps,)).to(device)
174     values_buf = torch.zeros((args.num_steps,)).to(device)
175
176     # ---- env state ----
177     obs, _ = env.reset(seed=args.seed)
178     next_state = torch.as_tensor(obs, dtype=torch.float32, device=device)
179     next_done = torch.zeros(1, device=device)
180
181     # ---- bookkeeping ----
182     episode_returns = []
183     current_return = 0.0
184     current_length = 0
185     global_step = 0
186     best_avg100 = -float('inf')
187     solved = False
188     solved_episode = None
189
190     n_iterations = args.total_timesteps // args.num_steps
191
192     for iteration in range(n_iterations):

```

```

193 # ---- optional learning-rate annealing ----
194 if args.anneal_lr:
195     frac = 1.0 - iteration / n_iterations
196     new_lr = frac * args.lr
197     for g in optimizer.param_groups:
198         g['lr'] = new_lr
199
200 # ROLLOUT -- collect num_steps transitions
201 for step in range(args.num_steps):
202     global_step += 1
203     states_buf[step] = next_state
204     dones_buf[step] = next_done
205
206     with torch.no_grad():
207         action, logp, _ = actor.get_action_and_logprob(
208             next_state.unsqueeze(0))
209         value = critic(next_state.unsqueeze(0))
210
211     actions_buf[step] = action.squeeze()
212     logprobs_buf[step] = logp.squeeze()
213     values_buf[step] = value.squeeze()
214
215     # Step env
216     obs, reward, terminated, truncated, _ = env.step(
217         int(action.item()))
218     rewards_buf[step] = float(reward)
219     done = terminated or truncated
220
221     current_return += float(reward)
222     current_length += 1
223
224     if done:
225         episode_returns.append(current_return)
226         avg100 = float(np.mean(episode_returns[-100:]))
227
228         if len(episode_returns) % 10 == 0:
229             print(f" ep {len(episode_returns):4d} | "
230                   f"step {global_step:7d} | "
231                   f"return {current_return:8.2f} | "
232                   f"avg100 {avg100:8.2f} | "
233                   f"len {current_length}")
234
235         if (len(episode_returns) >= 100
236             and avg100 > best_avg100):
237             best_avg100 = avg100
238             torch.save({'actor': actor.state_dict(),
239                       'critic': critic.state_dict()},
240                       args.checkpoint)
241
242         if (not solved and len(episode_returns) >= 100
243             and avg100 >= 200.0):
244             solved = True
245             solved_episode = len(episode_returns)
246             print(f"\n*** Solved at episode {solved_episode} "
247                   f"(global step {global_step}, "
248                   f"avg100 = {avg100:.2f}) ***\n")
249             if args.stop_when_solved:
250                 break
251
252         obs, _ = env.reset()
253         current_return = 0.0
254         current_length = 0
255
256     next_state = torch.as_tensor(obs, dtype=torch.float32,
257                                  device=device)
258     next_done = torch.tensor([float(terminated)], device=device)

```

```

259     if solved and args.stop_when_solved:
260         break
261
262
263     # GAE -- compute advantages and value targets
264     with torch.no_grad():
265         last_value = critic(next_state.unsqueeze(0)).squeeze()
266     rollout = Rollout(states_buf, actions_buf, logprobs_buf,
267                       rewards_buf, dones_buf, values_buf)
268     advantages, returns = compute_gae(rollout, last_value,
269                                       args.gamma, args.gae_lambda)
270
271     # UPDATE -- K epochs over the rollout in random minibatches
272     b_states = states_buf
273     b_actions = actions_buf
274     b_logprobs = logprobs_buf
275     b_advantages = advantages
276     b_returns = returns
277     b_values = values_buf
278
279     idxs = np.arange(args.num_steps)
280     for epoch in range(args.update_epochs):
281         np.random.shuffle(idxs)
282         for start in range(0, args.num_steps, args.minibatch_size):
283             end = start + args.minibatch_size
284             mb = idxs[start:end]
285
286             _, new_logp, entropy = actor.get_action_and_logprob(
287                 b_states[mb], b_actions[mb])
288             new_value = critic(b_states[mb])
289
290             # ---- policy loss (clipped surrogate, Schulman 2017 Eq. 7) ----
291             ratio = (new_logp - b_logprobs[mb]).exp()
292             mb_adv = b_advantages[mb]
293             mb_adv = (mb_adv - mb_adv.mean()) / (mb_adv.std() + 1e-8)
294
295             pg_loss1 = -mb_adv * ratio
296             pg_loss2 = -mb_adv * torch.clamp(ratio,
297                                             1 - args.clip_coef,
298                                             1 + args.clip_coef)
299             pg_loss = torch.max(pg_loss1, pg_loss2).mean()
300
301             # ---- value loss (clipped, CleanRL convention) ----
302             v_unclipped = (new_value - b_returns[mb]).pow(2)
303             v_clipped = b_values[mb] + torch.clamp(
304                 new_value - b_values[mb],
305                 -args.clip_coef, args.clip_coef)
306             v_clipped_loss = (v_clipped - b_returns[mb]).pow(2)
307             v_loss = 0.5 * torch.max(v_unclipped, v_clipped_loss).mean()
308
309             # ---- entropy bonus (encourages exploration) ----
310             ent_loss = entropy.mean()
311
312             loss = (pg_loss
313                   + args.vf_coef * v_loss
314                   - args.ent_coef * ent_loss)
315
316             optimizer.zero_grad()
317             loss.backward()
318             nn.utils.clip_grad_norm_(
319                 list(actor.parameters()) + list(critic.parameters()),
320                 args.max_grad_norm)
321             optimizer.step()
322
323     env.close()
324

```

```

325 # -----
326 # Save final stats
327 # -----
328 last100 = float(np.mean(episode_returns[-100:])) \
329             if len(episode_returns) >= 100 else float(np.mean(episode_returns))
330 if last100 >= best_avg100:
331     torch.save({'actor': actor.state_dict(),
332                'critic': critic.state_dict()}, args.checkpoint)
333 np.save(args.returns_file, np.asarray(episode_returns))
334 print(f"\nSaved checkpoint to {args.checkpoint}")
335 print(f"Saved per-episode returns to {args.returns_file}")
336
337 # -----
338 # Plot training curve
339 # -----
340 returns = np.asarray(episode_returns)
341 window = 100
342 if len(returns) >= window:
343     ma = np.convolve(returns, np.ones(window) / window, mode='valid')
344     ma_x = np.arange(window - 1, len(returns))
345 else:
346     ma, ma_x = np.array([]), np.array([])
347
348 status = "SOLVED" if last100 >= 200.0 else "not solved"
349
350 plt.figure(figsize=(10, 6))
351 plt.plot(returns, alpha=0.3, label='Episode return')
352 if len(ma):
353     plt.plot(ma_x, ma, linewidth=2.2, label='100-episode moving average')
354 plt.axhline(200, linestyle='--', color='gray', alpha=0.6,
355             label='Solved threshold (+200)')
356 plt.xlabel('Episode')
357 plt.ylabel('Episode return')
358 plt.title(f'PPO on {args.env}\n'
359           f'Last-100-episode avg: {last100:.2f} ({status})')
360 plt.legend(loc='lower right')
361 plt.grid(alpha=0.3)
362 plt.tight_layout()
363 plt.savefig(args.plot, dpi=130)
364 print(f"Saved training curve to {args.plot}")
365
366 print("=" * 64)
367 print(f"Final 100-episode average return: {last100:.2f} ({status})")
368 if solved:
369     print(f"First solved at episode {solved_episode}")
370 print("=" * 64)
371
372
373 # =====
374 # Comparison plot: DQN (Task 2) vs. PPO (Task 4)
375 # =====
376 def compare(args):
377     """Overlay DQN's per-episode returns and PPO's per-episode returns."""
378     if not os.path.exists(args.dqn_returns):
379         raise FileNotFoundError(
380             f"Cannot find DQN returns at {args.dqn_returns}.")
381     if not os.path.exists(args.returns_file):
382         raise FileNotFoundError(
383             f"Cannot find PPO returns at {args.returns_file}.")
384
385     dqn = np.load(args.dqn_returns)
386     ppo = np.load(args.returns_file)
387
388     window = 100
389     def ma(x):
390         if len(x) < window:

```

```

391         return np.array([]), np.array([])
392     m = np.convolve(x, np.ones(window) / window, mode='valid')
393     return np.arange(window - 1, len(x)), m
394
395     dqn_x, dqn_ma = ma(dqn)
396     ppo_x, ppo_ma = ma(ppo)
397
398     plt.figure(figsize=(10, 6))
399     plt.plot(dqn, alpha=0.15, color='C0')
400     plt.plot(ppo, alpha=0.15, color='C1')
401     if len(dqn_ma):
402         plt.plot(dqn_x, dqn_ma, color='C0', linewidth=2.2,
403                  label=f'DQN (last-100 avg = {dqn[-100:].mean():.1f})')
404     if len(ppo_ma):
405         plt.plot(ppo_x, ppo_ma, color='C1', linewidth=2.2,
406                  label=f'PPO (last-100 avg = {ppo[-100:].mean():.1f})')
407     plt.axhline(200, linestyle='--', color='gray', alpha=0.6,
408                 label='Solved threshold (+200)')
409     plt.xlabel('Episode')
410     plt.ylabel('Episode return')
411     plt.title(f'LunarLander-v3: DQN vs. PPO (100-episode moving averages)')
412     plt.legend(loc='lower right')
413     plt.grid(alpha=0.3)
414     plt.tight_layout()
415     plt.savefig(args.compare_plot, dpi=130)
416     print(f"Saved comparison plot to {args.compare_plot}")
417
418 # =====
419 # Video recording of a trained agent
420 # =====
421 def record_video(args):
422     import gymnasium as gym
423     import imageio
424
425     env = gym.make(args.env, render_mode="rgb_array")
426     actor = Actor()
427     ckpt = torch.load(args.checkpoint, map_location="cpu")
428     actor.load_state_dict(ckpt['actor'])
429     actor.eval()
430
431     frames = []
432     obs, _ = env.reset(seed=args.seed)
433     done, total = False, 0.0
434     while not done:
435         with torch.no_grad():
436             s = torch.as_tensor(obs, dtype=torch.float32).unsqueeze(0)
437             logits = actor(s)
438             action = int(logits.argmax(dim=1).item()) # greedy at eval time
439             obs, r, terminated, truncated, _ = env.step(action)
440             frames.append(env.render())
441             total += r
442             done = terminated or truncated
443     env.close()
444     imageio.mimsave(args.gif, frames, fps=30)
445     print(f"Saved trained-agent rollout to {args.gif} (return = {total:.2f})")
446
447
448 def get_parser():
449     p = argparse.ArgumentParser(description="PPO on LunarLander")
450     p.add_argument('--env', type=str, default='LunarLander-v3',
451                    help='Use LunarLander-v2 if v3 is not available')
452     p.add_argument('--mode', type=str, default='train',
453                    choices=['train', 'compare', 'record'])
454
455     # PPO hyperparameters (CleanRL defaults tuned for classic control)
456

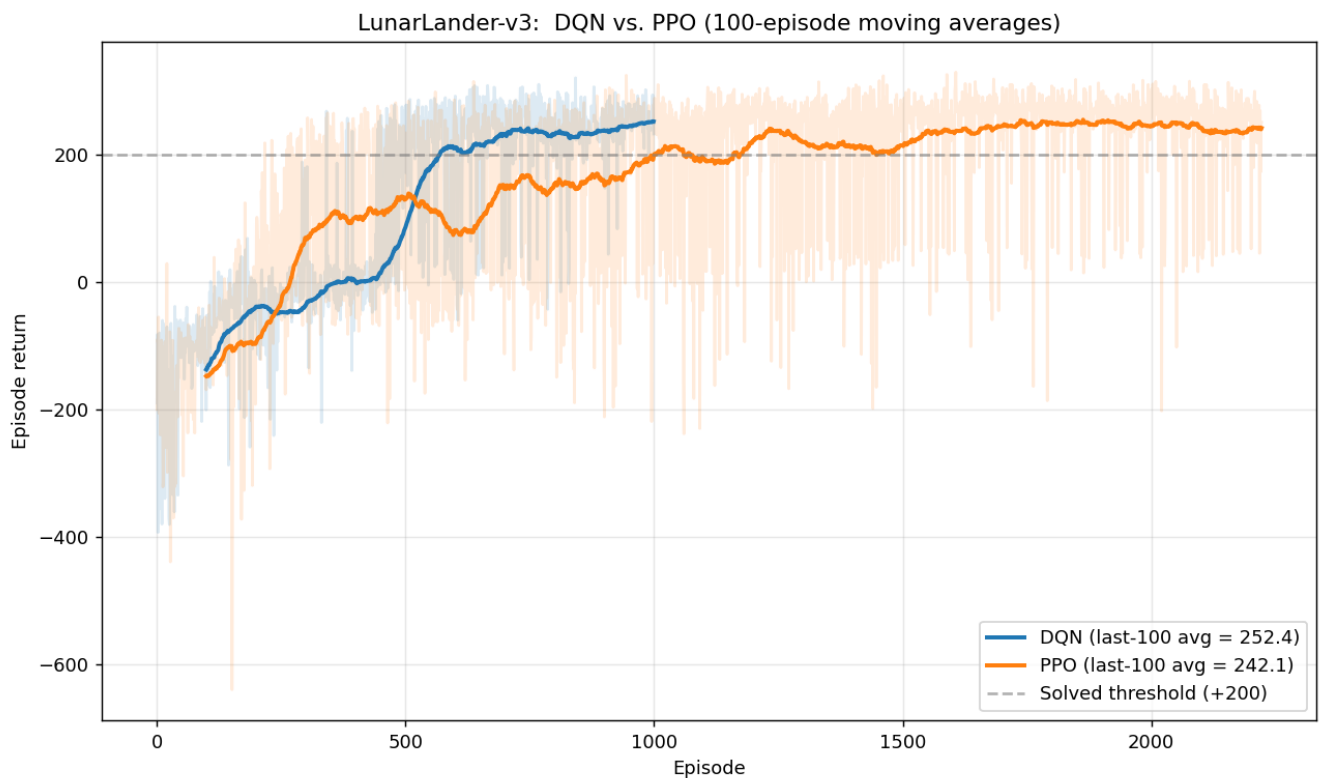
```

```

457 p.add_argument('--total_timesteps', type=int, default=600_000)
458 p.add_argument('--num_steps', type=int, default=1024,
459               help='Steps per rollout, before each PPO update.')
460 p.add_argument('--gamma', type=float, default=0.99)
461 p.add_argument('--gae_lambda', type=float, default=0.95)
462 p.add_argument('--clip_coef', type=float, default=0.2)
463 p.add_argument('--ent_coef', type=float, default=0.01)
464 p.add_argument('--vf_coef', type=float, default=0.5)
465 p.add_argument('--max_grad_norm', type=float, default=0.5)
466 p.add_argument('--update_epochs', type=int, default=10)
467 p.add_argument('--minibatch_size', type=int, default=64)
468 p.add_argument('--lr', type=float, default=3e-4)
469 p.add_argument('--anneal_lr', action='store_true', default=True)
470 p.add_argument('--stop_when_solved', action='store_true')
471
472 # Bookkeeping
473 p.add_argument('--seed', type=int, default=0)
474 p.add_argument('--checkpoint', type=str, default='ppo_lander.pt')
475 p.add_argument('--plot', type=str,
476               default='ppo_lander_training_curve.png')
477 p.add_argument('--returns_file', type=str, default='ppo_returns.npy')
478 p.add_argument('--dqn_returns', type=str, default='dqn_returns.npy')
479 p.add_argument('--compare_plot', type=str,
480               default='dqn_vs_ppo.png')
481 p.add_argument('--gif', type=str, default='ppo_lander.gif')
482 return p
483
484
485 if __name__ == "__main__":
486     args = get_parser().parse_args()
487     if args.mode == 'train':
488         train(args)
489     elif args.mode == 'compare':
490         compare(args)
491     elif args.mode == 'record':
492         record_video(args)

```

Training curve comparison:



Average return over the last 100 episodes for PPO:

```
1 =====
2 Final 100-episode average return: 242.07 (SOLVED)
3 First solved at episode 1001
4 =====
```

DQN solved the environment significantly faster (episode 569 vs. PPO's 1001) and achieved a slightly higher final average return (252.4 vs. 242.1). However, because DQN heavily relies on environment-specific hyperparameters such as the hardcoded 300-episode epsilon decay schedule and target network updates, PPO is generally considered easier to tune and is the standard default choice for new problems.