
Supplementary information

Outracing champion Gran Turismo drivers with deep reinforcement learning

In the format provided by the
authors and unedited

This document is supplementary material to “Outracing champion Gran Turismo drivers with deep reinforcement learning.”

Approximate Python Code

Here we provide approximate Python code that replicates the most critical aspects of our work. For the purposes of illustration, the code presented here is implemented in less computationally efficient ways than our actual implementation; has abstracted away many of the system implementation details such as remote process communication and check-pointing; and is written more directly than our actual implementation which supported wide configuration capabilities that otherwise detract from readability. Our implementation uses TensorFlow [1]; all references to ‘tf.x’ refer to an operation/class in the TensorFlow API, and references to ‘keras.x’ refers to an operation/class in the TensorFlow Keras API. A high-level overview of the training and rollout processes is provided in Figure 1.

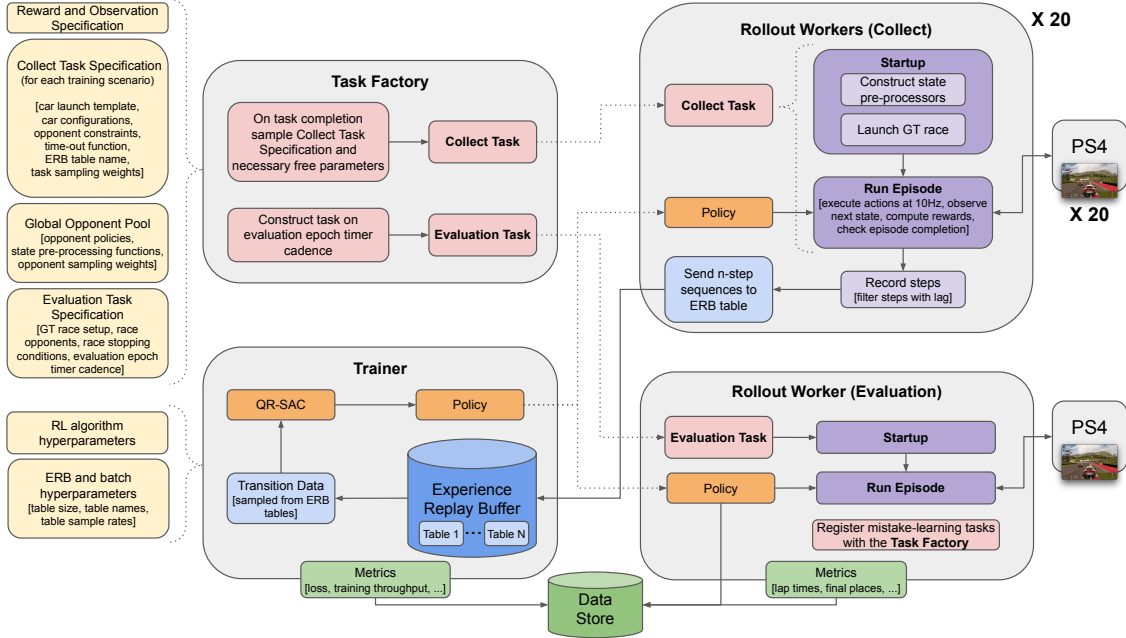


Figure 1: Diagram of modules in the GT Sophy training process

```

1 def train_loop():
2     models = Models(
3         policy=make_policy_network(),
4         af1=make_qr_q_network(),
5         af2=make_qr_q_network(),
6         af1_target=make_qr_q_network(),
7         af2_target=make_qr_q_network(),
8     )
9     # Sync prediction and target networks before start
10    copy_weights(models.af1.trainable_params, models.af1_target.trainable_params)
11    copy_weights(models.af2.trainable_params, models.af2_target.trainable_params)
12
13    opts = Optimizers(
14        policy=keras.optimizers.Adam(learning_rate=2.5e-5),
15        critic=keras.optimizers.Adam(learning_rate=5.0e-5, global_clipnorm=10.0),
16    )
17
18    param_server = start_parameter_server(models.policy)
19    replay_server = start_experience_replay_server()
20    table_datasets = [
21        make_dataset_for_table(replay_server, table_name)
22        for table_name in TABLE_NAMES
23    ]
24    task_server = start_task_factory_server(warm_up=True)
25
26    # Wait for MINIMUM_STEPS to be recorded to replay before we start any training
27    replay_server.block_and_wait(MINIMUM_STEPS)
28
29    while True:
30        for _ in range(BATCHES_PER_EPOCH):
31            transition_data = sample_data(table_datasets)
32            # see QR-SAC code listing for implementation
33            losses = step_qr_sac(models, opts, transition_data)
34            write_metrics(losses)
35            param_server.update_params(models.policy)
36            # warmup over after first epoch
37            task_server.warm_up = False
38

```

```

39 def sample_data(table_datasets: List[tf.data.Dataset]) -> Transition:
40     # For simplicity, this code omits handling the edge case when only a subset of
41     # the tables have data to sample.
42     # Before concat, each dataset has a different batch size based on the table
43     # weights
44     sub_traj_batches = [next(iter(dataset)) for dataset in table_datasets]
45     sub_traj_batches = nested_concat(sub_traj_batches, axis=0)
46     # sub_traj_batches has fields obs, action, reward, done
47     # Each tensor has two batch dimensions NT and 1 channel dimension, where N is
48     # batch and T is time. The obs T dim is one longer than other T dims.
49
50     discounting = tf.pow(DISCOUNT, tf.range(N_STEP, dtype=tf.float32))
51     discounting = tf.reshape(discounting, (1, N_STEP, 1))
52     discounted_reward = tf.math.reduce_sum(
53         sub_traj_batches.reward * discounting,
54         axis=1,
55     ) # (N, 1)
56     # Resulting transition removes T dimension, selecting relevant positions
57     return Transition(
58         obs=sub_traj_batches.obs[:, 0, :], # first obs
59         action=sub_traj_batches.action[:, 0, :], # first action
60         reward=discounted_reward, # discounted cumulative reward
61         next_obs=sub_traj_batches.obs[:, -1, :], # last obs
62         done=sub_traj_batches.done[:, -1, :], # last done
63     )

```

Listing 1: Training Loop Code

```

1 def step_qr_sac(m: Models, opts: Optimizers, t: Transition,) -> Losses:
2
3     critic_params = m.af1.trainable_params + m.af2.trainable_params
4     target_critic_params = (
5         m.af1_target.trainable_params + m.af2_target.trainable_params
6     )
7
8     # these values do not need gradients traced through them
9     policy_head_next = m.policy.policy_head(t.next_obs)
10    actions_next = m.policy.sample_from_head(policy_head_next)
11    log_prob_next = m.policy.log_prob(policy_head_next, actions_next)
12
13    # update policy
14    with tf.GradientTape() as tape:
15        policy_head = m.policy.policy_head(t.obs)
16        sampled_actions = m.policy.sample_from_head(policy_head)
17        log_prob = m.policy.log_prob(policy_head, sampled_actions)
18        policy_loss = policy_loss(
19            log_prob_samples=log_prob,
20            q1_sampled=m.af1.action_value(t.obs, sampled_actions),
21            q2_sampled=m.af2.action_value(t.obs, sampled_actions),
22        )
23    opts.policy.minimize(policy_loss, m.policy.trainable_params, tape)
24
25    # update critic networks
26    with tf.GradientTape() as tape:
27        critic_loss = critic_loss(
28            q1_quantile_observed=m.af1.quantile(t.obs, t.action),
29            q2_quantile_observed=m.af2.quantile(t.obs, t.action),
30            q1_quantile_sampled_next=m.af1_target.quantile(t.next_obs, actions_next),
31            q2_quantile_sampled_next=m.af2_target.quantile(t.next_obs, actions_next),
32            log_prob_next=log_prob_next,
33            reward=t.reward,
34            done=t.done,
35        )
36    opts.critic.minimize(critic_loss, critic_params, tape)
37
38    # move target network params toward prediction networks
39    for pred_param, target_param in zip(critic_params, target_critic_params):
40        target_param.assign(
41            SMOOTH_FACTOR * pred_param + (1.0 - SMOOTH_FACTOR) * target_param
42        )
43
44    return Losses(policy_loss, critic_loss)
45

```

```

46 def policy_loss(
47     log_prob_samples: tf.Tensor, # (N, 1)
48     q1_sampled: tf.Tensor, # (N, 1)
49     q2_sampled: tf.Tensor, # (N, 1)
50 ) -> tf.Tensor:
51     min_q = tf.minimum(q1_sampled, q2_sampled)
52     return tf.math.reduce_mean(ALPHA * log_prob_samples - min_q)
53
54
55 def critic_loss(
56     q1_quantile_observed: tf.Tensor, # (N, M)
57     q2_quantile_observed: tf.Tensor, # (N, M)
58     q1_quantile_sampled_next: tf.Tensor, # (N, M)
59     q2_quantile_sampled_next: tf.Tensor, # (N, M)
60     log_prob_next: tf.Tensor, # (N, 1)
61     reward: tf.Tensor, # (N, 1)
62     done: tf.Tensor, # (N, 1)
63 ) -> tf.Tensor:
64     target = tf.stop_gradient(qr_sac_bootstrap_target(
65         q1_quantile_sampled_next,
66         q2_quantile_sampled_next,
67         log_prob_next,
68         reward,
69         done,
70         DISCOUNT**N_STEP, # adjust for N-step transition
71         ALPHA,
72     ))
73
74     q1_loss = tf.math.reduce_mean(quantile_loss(q1_quantile_observed, target))
75     q2_loss = tf.math.reduce_mean(quantile_loss(q2_quantile_observed, target))
76     return q1_loss + q2_loss
77
78
79 def quantile_loss(
80     quantile: tf.Tensor, # (N, M)
81     quantile_target: tf.Tensor, # (N, M)
82 ) -> tf.Tensor:
83     batch_size = quantile.shape[0]
84     num_quantiles = quantile.shape[1]
85
86     y = tf.reshape(quantile_target, [batch_size, -1, 1]) # (N, M, 1)
87     x = tf.reshape(quantile, [batch_size, 1, -1]) # (N, 1, M)
88
89     # compute huber loss
90     diff = y - x # (N, M, M)
91     huber_loss = tf.where(
92         tf.math.abs(diff) < 1.0, 0.5 * diff ** 2, tf.math.abs(diff) - 0.5
93     ) # (N, M, M)
94
95     # sample taus
96     steps = tf.range(num_quantiles, dtype=tf.float32) # (M,)
97     taus = tf.reshape((steps + 1) / num_quantiles, [1, 1, -1]) # (1, 1, M)
98     taus_minus = tf.reshape((steps / num_quantiles), [1, 1, -1]) # (1, 1, M)
99     taus_hat = (taus + taus_minus) / 2.0 # (1, 1, M)
100
101     # compute quantile loss
102     delta = tf.cast(tf.stop_gradient(diff) < 0.0, dtype=tf.float32) # (N, M, M)
103     element_wise_loss = tf.math.abs(taus_hat - delta) * huber_loss # (N, M, M)
104
105     return tf.reduce_mean(tf.reduce_sum(element_wise_loss, axis=2), axis=1) # scalar

```

```

106 def qr_sac_bootstrap_target(
107     q1_quantile_sample_next: tf.Tensor, # (... , M)
108     q2_quantile_sample_next: tf.Tensor, # (... , M)
109     next_log_prob_samples: tf.Tensor, # (... , 1)
110     reward: tf.Tensor, # (... , 1)
111     done: tf.Tensor, # (... , 1)
112     discount: float, # scalar
113     alpha: float, # scalar
114 ) -> tf.Tensor:
115     quantile_next = minimum_quantile(
116         q1_quantile_sample_next, q2_quantile_sample_next
117     )
118     entropy = alpha * next_log_prob_samples
119     return reward + discount * (1.0 - done) * (quantile_next - entropy)
120
121 def minimum_quantile(
122     quantile1: tf.Tensor, # (N, M)
123     quantile2: tf.Tensor, # (N, M)
124 ) -> tf.Tensor:
125     batch_size = quantile1.shape[0]
126
127     # compute Q value
128     q1 = tf.reduce_mean(quantile1, axis=1, keepdims=True) # (N, 1)
129     q2 = tf.reduce_mean(quantile2, axis=1, keepdims=True) # (N, 1)
130
131     # pick quantiles based on minimum prediction
132     stacked_q = tf.concat([q1, q2], axis=1) # (N, 2)
133     min_indices = tf.cast(tf.math.argmin(stacked_q, axis=1), dtype=tf.int32) # (N,)
134     indices = tf.stack([tf.range(batch_size), min_indices], axis=1) # (N, 2)
135
136     stacked_quantile = tf.stack([quantile1, quantile2], axis=1) # (N, 2, M)
137     minimum_quantile = tf.gather_nd(stacked_quantile, indices) # (N, M)
138
139     return minimum_quantile
140

```

Listing 2: QR-SAC Code

```

1 def get_task() -> Task:
2     # The returned task defines which cars are agent controlled and for each opponent
3     # whether to control it with a pre-trained policy, PID, or the built-in AI
4     task_server = get_task_server_info()
5     if task_server.warm_up:
6         # During warm up, run 20 cars across the track at once all selecting
7         # uniformly random actions
8         return Task(
9             cars_launch=uniform_jittered_launch(num_cars=20, mph_range=(0, 60)),
10            random_policy=True,
11            agent_controlled_car_ids=list(range(20)),
12            car_id_to_opp_policy={},
13            car_id_to_opp_pid={},
14            builtin_ai_car_ids=[],
15            table_name="1v0",
16            time_out_fn=lambda steps, obs: steps >= MAX_STEPS,
17            is_eval=False,
18            builtin_ai_params=None,
19        )
20     elif task_server.needs_eval:
21         # Time for periodic evaluation
22         task_server.needs_eval = False
23         return Task(
24             cars_launch=EVAL_LAUNCH,
25             random_policy=False,
26             agent_controlled_car_ids=[0], # one car controlled by current policy
27             car_id_to_opp_policy=EVAL_OPP_POLICIES,
28             car_id_to_opp_pid={},
29             builtin_ai_car_ids=EVAL_OPP_BUILTIN_IDS,
30             table_name=None, # Do not collect training data in evals
31             time_out_fn=lambda steps, obs: steps >= MAX_STEPS,
32             is_eval=True,
33             builtin_ai_params=None,
34        )
35     elif (
36         # limited slots for mistake learning to prevent regular tasks from starving
37         len(task_server.mistake_learning_tasks) > 0
38         and task_server.mistake_slots_available()
39     ):
40         return task_server.pop_mistake_learning_task()
41     else:
42         # Standard data collection setting
43         # First we sample a task specification (e.g., 1v0, slipstream, etc.)
44         # The TaskSpec is almost a Task, but the launch, opponent policies, and
45         # builtin-ai params need to be sampled from a distribution
46         task_spec = np.random.choice(TASK_SPECS, p=TASK_SPEC_PROBS)
47         return Task(
48             cars_launch=task_spec.sample_launch(),
49             random_policy=False,
50             agent_controlled_car_ids=task_spec.agent_ids,
51             car_id_to_opp_policy=task_spec.sample_opp_policies(),
52             car_id_to_opp_pid=task_spec.pids,
53             builtin_ai_car_ids=task_spec.builtin_ids,
54             table_name=task_spec.table_name,
55             time_out_fn=task_spec.time_out_fn,
56             is_eval=False,
57             builtin_ai_params=task_spec.sample_builtin_ai_params(),
58        )

```

Listing 3: The task factory is run in a separate process and is remotely queried by rollout workers.

```

1 def rollout_loop():
2     policy = make_policy_network()
3     param_server = get_parameter_server()
4     ps_server = get_playstation_server()
5     replay_server = get_experience_replay_server()
6     task_server = get_task_server()
7     while True:
8         task = get_task()
9         if task.random_policy:
10             agent_policy = UniformRandomPolicy()
11         else:
12             load_weights(policy, param_server.get_params()) # Get latest policy
13             agent_policy = policy
14
15         ps_server.launch_race( # Start the race on the PlayStation
16             TRACK_NAME, CAR_INFO, task.cars_launch, task.builtin_ai_car_ids
17         )
18         obs = ps_server.wait_for_race_start()
19         p_obs = preprocess(obs) # Preprocess creates an obs batch for all cars
20         trajs = [
21             start_trajectory(initial_obs=p_obs[agent_id])
22             for agent_id in task.agent_car_ids()
23         ]
24
25         for step in range(MAX_STEPS):
26             actions = compute_all_actions(
27                 task, agent_policy, p_obs, is_exploit=task.is_eval
28             )
29             ps_server.set_and_hold(actions)
30             # Get the next observation 6 frames later = 0.1 seconds @60fps
31             next_obs = ps_server.wait_for_obs_at_frame(obs.frame_count + 6)
32             rewards = compute_reward(obs, next_obs) # list of rewards for all cars
33             dones = [False] * task.cars_launch.num_cars # Continuing task
34
35             obs = next_obs
36             p_obs = preprocess(next_obs)
37
38             for agent_id, traj in zip(task.agent_car_ids(), trajs):
39                 traj.append(
40                     obs=p_obs[agent_id],
41                     action=actions[agent_id],
42                     reward=rewards[agent_id],
43                     done=dones[agent_id],
44                 )
45                 if task.is_eval and should_retry(p_obs[agent_id]):
46                     task_server.add_mistake_learning_task(traj)
47
48             # create N-step subtrajectory if this task collects data
49             if step + 1 >= N_STEP and task.table_name:
50                 for traj in trajs:
51                     # this writes one additional observation time step
52                     replay_server.write(task.table_name, traj[-N_STEP:])
53
54             if task.time_out_fn(step, obs):
55                 break
56
57         if task.is_eval:
58             write_metrics(trajs)
59
60

```

```

61 def compute_all_actions(
62     task: Task,
63     agent_policy: Policy,
64     p_obs: tf.Tensor,
65     is_exploit: bool,
66 ) -> List[Optional[tf.Tensor]]:
67     # Actions for built-in AI cars remain None
68     all_actions = [None for _ in range(task.cars_launch.num_cars)]
69     # Compute actions for agent-controlled cars
70     for car_id in task.agent_car_ids():
71         if is_exploit:
72             a = agent_policy.exploit_action(p_obs[car_id])
73         else:
74             a = agent_policy.sample_action(p_obs[car_id])
75         all_actions[car_id] = a
76     # Compute actions for opponent cars using previously trained policies
77     for car_id, opp_policy in task.car_id_to_opp_policy.items():
78         if is_exploit:
79             a = opp_policy.exploit_action(p_obs[car_id])
80         else:
81             a = opp_policy.sample_action(p_obs[car_id])
82         all_actions[car_id] = a
83     # Compute actions for opponent cars using a PID controller
84     for car_id, pid in task.car_id_to_opp_pid.items():
85         # PID controllers always deterministically select actions
86         a = pid.exploit_action(p_obs[car_id])
87         all_actions[car_id] = a
88
89     return all_actions

```

Listing 4: Rollout worker code. 21 instances, each with their own PlayStation connection, are run in parallel.

Hyperparameters

Here we provide a set of tables that cover the hyperparameters used for the training experiments. We focus on the parameters used to train GT Sophy for competitive racing.

RL Algorithm	
n-step	7
discount factor	0.9896
entropy regularization coefficient, α	0.01
number of quantiles	32
parameter smoothing factor	0.005
policy learning rate	2.5e-5
critic learning rate	5.0e-5
optimizer	Adam
global clip norm (critic optimizer only)	10.0

Table 1: Table of parameters for QR-SAC RL algorithm.

Neural Network	
number of hidden layers	4
number of units per hidden layer	2048
activation	relu
dropout (policy network only)	0.1

Table 2: Table of parameters for the neural network architecture.

Trainer	
mini-batch size	1024
total replay buffer capacity	10e+6
number of buffer tables	8 to 10
training steps per epoch	6000
minimum transitions to take before training	40000

Table 3: Table of parameters for the trainer.

References

- [1] Martín Abadi, Ashish Agarwal, Paul Barham, Eugene Brevdo, Zhifeng Chen, Craig Citro, Greg S. Corrado, Andy Davis, Jeffrey Dean, Matthieu Devin, Sanjay Ghemawat, Ian Goodfellow, Andrew Harp, Geoffrey Irving, Michael Isard, Yangqing Jia, Rafal Jozefowicz, Lukasz Kaiser, Manjunath Kudlur, Josh Levenberg, Dandelion Mané, Rajat Monga, Sherry Moore, Derek Murray, Chris Olah, Mike Schuster, Jonathon Shlens, Benoit Steiner, Ilya Sutskever, Kunal Talwar, Paul Tucker, Vincent Vanhoucke, Vijay Vasudevan, Fernanda Viégas, Oriol Vinyals, Pete Warden, Martin Wattenberg, Martin Wicke, Yuan Yu, and Xiaoqiang Zheng. TensorFlow: Large-scale machine learning on heterogeneous systems, 2015. Software available from tensorflow.org.