

1 RL: What's Changed Since MDPs?

1. Recall the Bellman Equation we used in MDPs to determine the value of a given state:

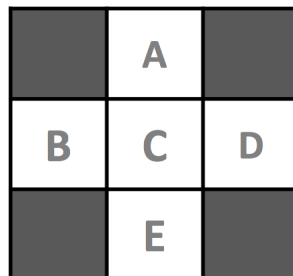
$$V^*(s) = \max_a \sum_{s'} T(s, a, s') [R(s, a, s') + \gamma V(s')]$$

What information do we no longer have direct access to in the transition to RL?

2. What is the difference between online and offline learning? Which type of learning does MDP use? How about RL?

2 Temporal Difference Learning and Q-Learning

Consider the Gridworld example that we looked at in lecture. We would like to use TD learning to find the values of these states.



Suppose we use some policy and observe the following (s, a, r, s') transitions and rewards:

$(B, \text{East}, 2, C)$, $(C, \text{South}, 4, E)$, $(C, \text{East}, 6, A)$, $(B, \text{East}, 2, C)$

The initial value of each state is 0. Let $\gamma = 1$ and $\alpha = 0.5$.

1. What are the learned values for each state from TD learning after all four observations?

2. In class, we presented the following two formulations for TD-learning:

$$V^\pi(s) \leftarrow (1 - \alpha)V^\pi(s) + (\alpha)sample$$

$$V^\pi(s) \leftarrow V^\pi(s) + \alpha(sample - V^\pi(s))$$

Mathematically, these two equations are equivalent. However, they represent two conceptually different ways of understanding TD value updates. How could we intuitively explain each of these equations?

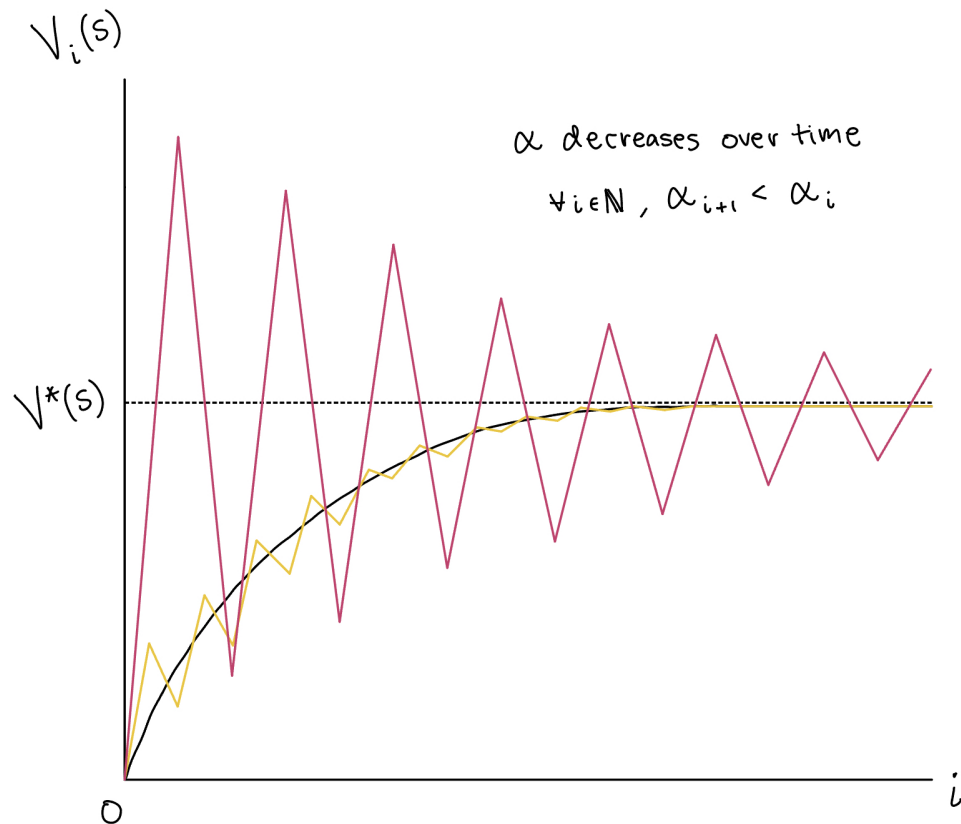
3. What are the learned Q-values from Q-learning after all four observations? Use the same $\alpha = 0.5, \gamma = 1$ as before.

Transitions	$(B, East)$	$(C, South)$	$(C, East)$
(initial)			
$(B, East, 2, C)$			
$(C, South, 4, E)$			
$(C, East, 6, A)$			
$(B, East, 2, C)$			

3 RL: Conceptual Questions

Recall that in Q-learning, we continually update the values of each Q-state by learning through a series of episodes, ultimately converging upon the optimal policy.

1. What's the main shortcoming of TD learning that Q-learning resolves?
2. We are given two runs of TD-learning using the same sequence of samples but different α values depicted in the plot below. Assume the dashed horizontal line represents the optimal value for a specific state s and the black curve represents the smoothest transition to the optimal value given this sequence of samples. In both runs α decreases over time (or iterations), but one run has α values larger than the other run at any point in time. Which run (red or yellow) corresponds to the smaller values of α ? How do the relative sizes of α affect the rate of convergence to the optimal value?



3. We are given a pre-existing table of current estimate of Q-values (and its corresponding policy), and asked to perform ϵ -greedy Q-learning. Individually, what effect does setting each of the following hyperparameters to 0 have on this process?
 - (a) α

(b) γ (c) ϵ

Remember that in ϵ -greedy Q-learning, we follow the following formulation for choosing our action:

$$\text{action at time } t = \begin{cases} \arg \max_{Q(s,a)} & \text{with probability } 1 - \epsilon \\ \text{any action } a & \text{with probability } \epsilon \end{cases}$$

4. Consider a variant of the ϵ -greedy Q-learning algorithm that is changed such that instead of using the policy extracted from our current Q-values, we use a fixed policy instead. We still perform exploration with probability ϵ . If this fixed policy happens to be optimal, how does the performance of this algorithm compare to normal ϵ -greedy Q-learning?

5. Recall the count exploration function used in the modified Q-update:

$$f(u, n) = u + \frac{k}{n + 1}$$

$$Q(s, a) \leftarrow Q(s, a) + \alpha[r + \gamma \max_{a'} f(Q(s', a'), N(s', a')) - Q(s, a)]$$

Remember that k is a hyperparameter that the designer chooses, and $N(s', a')$ is the number of times we've visited the (s', a') pair. What is the effect of increasing or decreasing k ?

6. Contrast the following pairs of reinforcement learning terms:

(i) Off-policy vs. on-policy learning

(ii) Model-based vs. model-free

(iii) Passive vs. active