

Reinforcement Learning

Returns

The agent learns a policy

- Policy at step t , π_t which is a mapping from states to action probabilities, $\pi(s, a) = \text{prob. that } a_t = a \text{ when } s_t = s.$

- Reinforcement ^{learning} ~~algo.~~ is not a single algorithm, it is a learning methods that represents how an agent changes its policy as a result of experience. That means the agent's goal is to get as much reward as it can over the long run. i.e the goal of an agent is to maximize its returns.

- Suppose the sequence of reward after step t is $r_{t+1}, r_{t+2}, r_{t+3}, \dots$ we want to maximize the return i.e total reward G_t for every step t .

Next term with related to return is important $\rightarrow$

"Episode Task":

Interaction breaks naturally into "episode" as for example "Maze Game".

$$G_t = r_{t+1} + r_{t+2} + r_{t+3} \dots + r_T \rightarrow \text{Final step (terminal state)} \\ \text{i.e. an end to the episode.}$$

However, there is a possibility of "Return for Continuing Task"

Cont. Task: no natural episodes (for interactions)

$\therefore$ Discounted return with the introduction of discounted factor γ

$$\therefore \text{The new, } G_t = r_{t+1} + \gamma r_{t+2} + \gamma^2 r_{t+3} + \dots = \sum_{k=0}^{\infty} \gamma^k r_{t+k+1}$$

where, $\gamma, 0 \leq \gamma \leq 1$; γ = discounted factor.

$\gamma \rightarrow 0$ // shortsighted | γ is a controlling factor
 $\gamma \rightarrow 1$ // far sighted

Main intention is to maximize the expected return, $E[G_t]$ for each step t .

[$\gamma = 0$ is a immediate reward problem? NO
as next state that we are going to see depends on my current actions.]

Value Functions

∴ How to maximize the expected "Returns" ? at every timestep t .

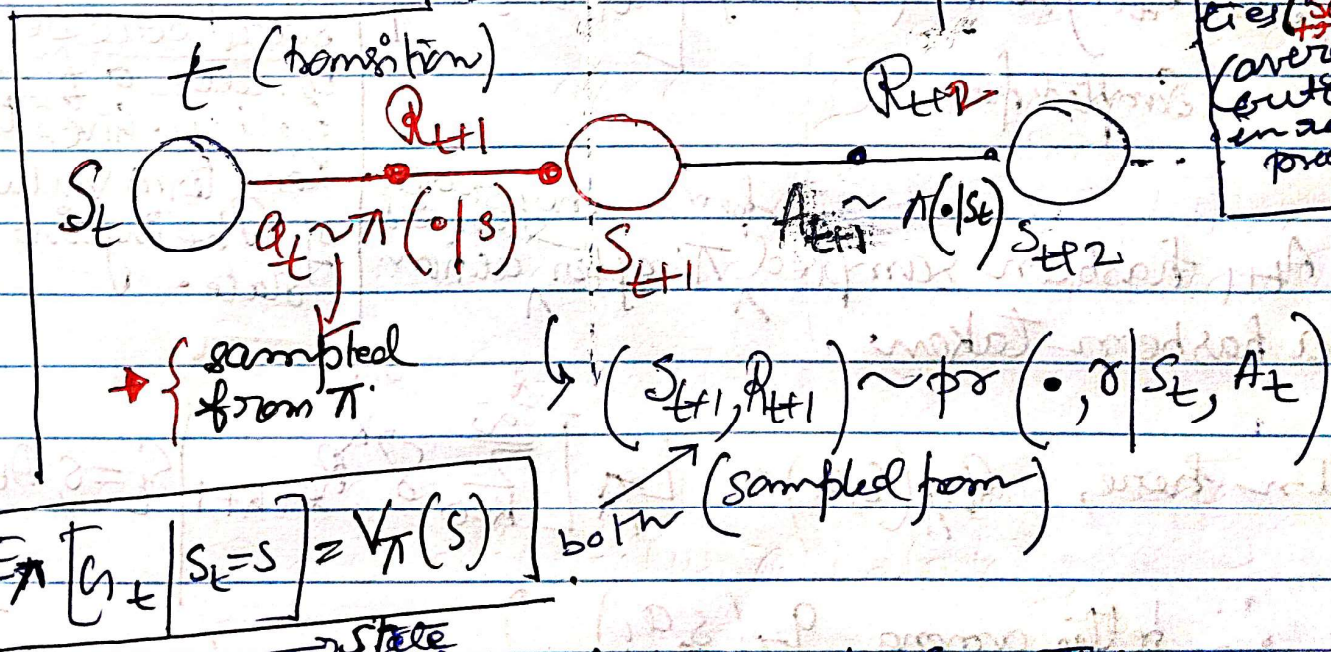
• To achieve this maximization process we are introducing "Value Function".

- Expected future rewards
- Start state (or state-action pair)
- Following policy π

State-value function
for policy π :

$$V_{\pi}(s) = E_{\pi} \left[\sum_{k=0}^{\infty} \gamma^k R_{t+k+1} \mid S_t = s \right]$$

EA
Expected value
= Average
future out-
come consid-
ering proba-
bilities of
all possible
cases (several
trajectories)
(average
outcome
in repeated
process)



Action-value function for policy π :

$$Q_{\pi}(s, a) = E_{\pi} \left[\sum_{k=0}^{\infty} \gamma^k R_{t+k+1} \mid S_t = s, A_t = a \right]$$

State
↓

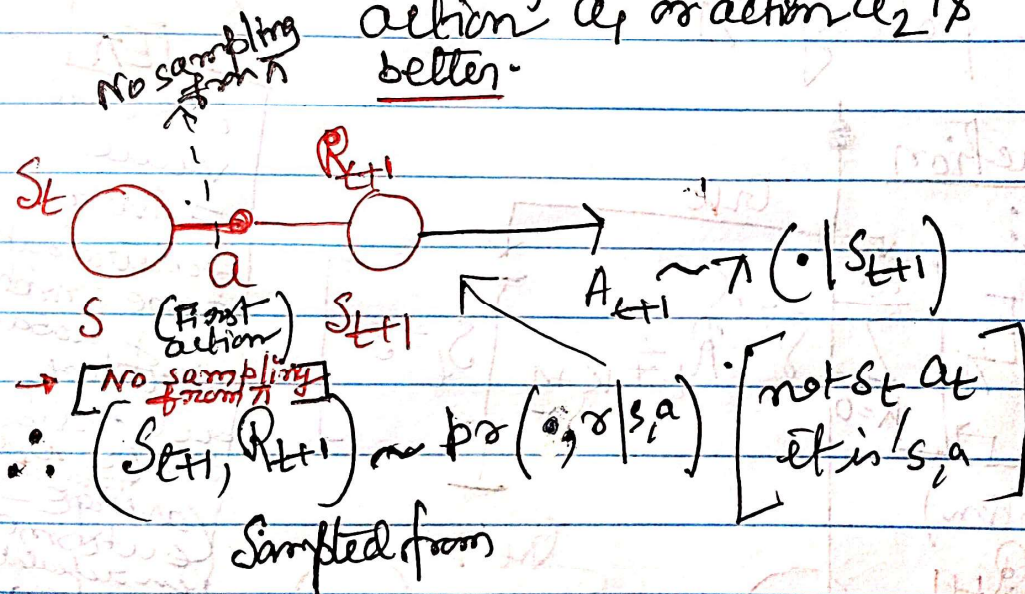
Action-value function for policy π :

$$Q_{\pi}(s, a) = E_{\pi} \left[\sum_{k=0}^{\infty} \gamma^k R_{t+k+1} \mid S_t = s, A_t = a \right]$$

↑
action

★ Action value f^n /

- What if from single state action a_1 or action a_2 is better.



A_{t+1} has been sampled π after action a has been taken

↑ from only started from S_t

★ state value f^n

- prev. one state-value
- This above is related, is called action value.
- state-value: returns expected value from state S_t following policy π there after.
- Purpose: which states are better under a fixed policy. Gives the long term value of being in a state.

In here, $Q_{\pi}(s, a) = E_{\pi} \left[\sum_{k=0}^{\infty} \gamma^k R_{t+k+1} \mid S_t = s, A_t = a \right]$

• better among $Q_{\pi}(s, a_1)$ or $Q_{\pi}(s, a_2)$ } also that policy π shall be followed

• Comparisons among actions

• change actions

$\therefore$ from a particular state a_1 or a_2 :
which one is better

$\therefore E_{\pi} [G_t | S_t = s, A_t = a] \rightarrow$ Two value fⁿ

$$= Q_{\pi}(s, a)$$

- State value fⁿ

- Action value fⁿ

- Most of time Q_{π} is better than V_{π} fⁿ

Although these two fⁿs are different, they however have a relationship

$$V_{\pi}(s) = \sum_a \pi(a|s) Q_{\pi}(s, a)$$

$$= E_{\pi} [Q_{\pi}(s, a)]$$

selecting
action
based on
policy π .

But Q_{π} is an integral of V_{π} ?

$\therefore$ from a particular state a_1 or a_2 :
which one is better

$\therefore E_{\pi} [G_t | S_t = s, A_t = a] \rightarrow$ Two value fⁿ

$$= Q_{\pi}(s, a)$$

- State value fⁿ
- Action value fⁿ

• Most of time Q_{π} is better than V_{π} fⁿ

Although these two fⁿs are different, they however have a relationship

$$V_{\pi}(s) = \sum_a \pi(a|s) Q_{\pi}(s, a)$$

$$= E_{\pi} [Q_{\pi}(s, a)]$$

selecting action based on policy π .

But Q_{π} interacts of V_{π} ?

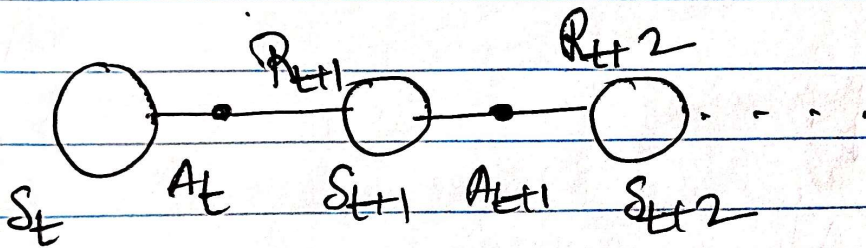
Value f^n

$$V_\pi = E_\pi [G_t | S_t = s] = E_\pi \left[\sum_{k=0}^{\infty} \gamma^k R_{t+k+1} | S_t = s \right], \forall s \in S$$

Expected value of return $\left[\sum_{k=0}^{\infty} \gamma^k R_{t+k+1} | S_t = s \right]$

For a Policy π : Bellman Equation

Let's see, $G_k = R_{t+1} + \gamma R_{t+2} + \gamma^2 R_{t+3} + \dots = R_{t+1} + \gamma [R_{t+2} + \gamma R_{t+3} + \dots]$



Return at $t \rightarrow$ start adding up from R_{t+1}
 Return at $t+1 \rightarrow$ start summing " " R_{t+2}
 Return at $t+2 \rightarrow$ " " " " R_{t+3}
 $\therefore G_t$ can be written recursively as it expresses the total future reward as the current reward + discounted future (γG_{t+1}) rewards. It means that the return at time t depends on the return at the next time step, linking all future rewards together.

$\therefore G_t = R_{t+1} + \gamma R_{t+2} + \gamma^2 R_{t+3} + \dots$
 $= R_{t+1} + \gamma [R_{t+2} + \gamma R_{t+3} + \gamma^2 R_{t+4} + \dots]$

$G_t = R_{t+1} + \gamma G_{t+1}$

Now let's see the value f^n

$$\begin{aligned}
 V_{\pi}(s) &= E_{\pi} [G_t | S_t = s] \quad \Bigg| \quad E_{\pi} = \text{Expectation} \\
 &= E_{\pi} [R_{t+1} + \gamma V_{\pi}(S_{t+1}) | S_t = s] \\
 &= \sum_a \pi(a|s) \sum_{s'} \sum_{\gamma} p(s', \gamma | s, a) [\gamma + \gamma E_{\pi} [G_{t+1} | S_{t+1} = s']] \\
 &= \sum_a \pi(a|s) \sum_{s', \gamma} p(s', \gamma | s, a) [\gamma + \gamma V_{\pi}(s')], \quad \begin{array}{c} \text{Diagram: } s \xrightarrow{a} s' \text{ with } R_{t+1} \text{ reward} \\ \text{Labels: } s \sim \pi(\cdot|s), \quad (s', R_{t+1}) \sim \Pr(\cdot, \cdot | s, a) \end{array}
 \end{aligned}$$

Grid World

	0.1	0.8	0.2
		0.7	0.2
			0.2

$$(S_{t+1}, R_{t+1}) \sim \Pr(\cdot, \cdot | s, A_t)$$

probability distribution

start state s is fixed

policy $\rightarrow \pi(\uparrow|s) = 0.5$ [up.] $\pi(\rightarrow|s) = 0.5$ [right] Two possible actions

let's say up $\Pr(\cdot, \cdot | s, \uparrow)$ [we have to sample state]

let's say current state is 7
 $s = 7$

prob. distribution

$$V_s(s) = \sum_{a \in \mathcal{A}} \pi(a|s) \sum_{s', r} p(s', r | s, a) [r + \gamma V_\pi(s')]$$

Explanation:

- $\pi(a|s)$: probability of choosing action a in state s
- $p(s', r | s, a)$: joint prob of next state s' and reward r given (s, a)

- The inner expression

$[r + \gamma V_\pi(s')]$ is the immediate reward + discounted value of the next state

- If rewards are continuous the $\sum$ becomes an integral

How the expectation is computed

1. start at s (fixed)
2. Sample $a \sim \pi(\cdot | s)$
3. Sample $(s', r) \sim p(\cdot, \cdot | s, a)$
4. The contribution is $r + \gamma V_\pi(s')$
5. Average (sum) over all possible a, s', r

Only first step of the probabilities are explicitly expanded in the Bellman equation & the remainder of the trace is

Linear system & solvability

- For a finite MDP with N states, the Bellman expectation eqⁿs form an $N \times N$ linear system whose unknowns are the N value-function entries $\{V_{\pi}(s_i)\}_{i=1}^N$
- N states & N unknown values, $V_{\pi}(s_1), V_{\pi}(s_2), \dots, V_{\pi}(s_N)$
- The Bellman Expectation equation gives you N equations one for each state
- There is a unique solⁿ for V_{π} (assumption $0 \leq \gamma < 1$)
- Convert infinite-sum-return problem to tractable linear algebra problem.