

MDP: Markov Decision Process

① "Actions are indeterminate"

"Action" in "Search" (state A) Action (state B) "Determinate"

Action in MDP: (state A) Action → (state B) $p = 0.2$
(state C) $p = 0.4$ "Indeterminate"
(state D) $p = 0.4$

② Core function depends

the starting and ending states of action"

MDP consists of

- ① a set of states
- ② a starting state
- ③ a set of actions modeled by a global probabilistic transition relation
- ④ Reward function
Reward(s, a, s'): the rewarding for (s) → (s')
- ⑤ Goal test: End(s)
- ⑥ A "discount factor": γ , $0 \leq \gamma \leq 1$. "How much we value the future"

Transition probabilities $T(s, a, s')$

1. if $T(s, a, s') \neq 0$ then s' is a possible successor of s
2. $\sum_s T(s, a, s') = 1$

Policy: A "mapping" from "state" to "action".

$\pi: S \rightarrow A$, $\pi(s_1) = a_1$, $\pi(s_2) = a_2, \dots$

How to compare Policies? $\pi_1(m_1), m_1(?) \pi_2(m_2), m_2$

Expected utilities $V_\pi(s)$

How to compute V_π ?

1. By recursion.
$$V_\pi(s) = \begin{cases} 0, & \text{if End}(s) \\ \sum_s T(s, \pi(s), s') [\text{Reward}(s, \pi(s), s') + \gamma V_\pi(s')], & \text{otherwise} \end{cases}$$

Q-value: " $Q_\pi(s, a) = \sum_s T(s, a, s') [\text{Reward}(s, a, s') + \gamma V_\pi(s')]$ "

2. By Iteration

Initialize $V_\pi^0(s) = 0$ for all s

for $t = 1 \dots \text{max-iter.}$

for each state $s \in S$.

$$V_\pi^t(s) = \sum_s T(s, \pi(s), s') [\text{Reward}(s, \pi(s), s') + \gamma V_\pi^{t-1}(s')]$$

Stop when V^t doesn't change much from V^{t-1}

Optimal Policy and Value function

Given a MDP, we want to know

the optimal policy and optimal value function

π_{opt} : for every policy π' $V_{\pi_{\text{opt}}}(s) \geq V_{\pi'}(s)$ for all s

Hill Climbing Algorithm $\Rightarrow$ optimal policy

Value iteration $\Rightarrow$ optimal policy value

Initialize π to a random policy

For $i = 1 \dots \text{Max-iter.}$

$\pi = \text{Successor}(\pi)$ "A better one"

"How to make it better?"

$V_\pi(s) = Q(s, \pi(s))$ ("Q-value": $Q_\pi(s, a) = \sum_s T(s, a, s') [\text{Reward}(s, a, s') + \gamma V_\pi(s')]$ ")

$\Rightarrow$ Find an action π' . $V_{\pi'}(s) = \max_{a \in \text{Action}(s)} Q_\pi(s, a) \Rightarrow V_{\pi'}(s) \geq V_\pi(s)$ since $V_\pi(s)$ is maximum
 $\Rightarrow \pi'$ is as good as π or better than π

$\Rightarrow$ Policy improvement: Given a $\pi \Rightarrow V_\pi(s)$, improve the policy to π_{new}

Compute $Q_\pi(s, a)$ for every s and $a \in \text{action}(s)$

$\pi_{\text{new}}(s) = \arg \max_{a \in \text{Action}(s)} Q_\pi(s, a)$

Successor(π) = π_{new}

Policy iterations

π policy evaluation $\rightarrow V_\pi$ policy improvement $\rightarrow \pi_1$ policy evaluation $\rightarrow V_{\pi_1}$ policy improvement $\rightarrow \pi_2 \dots$

Initialization: generate $\pi(s) \in \text{Action}(s)$ randomly for all $s \in S$
stable $\leftarrow$ false
while not stable:
 policy evaluation.
 policy improvement.
 if for all $s \in S$, $\pi(s)$ is unchanged,
 stable $\leftarrow$ true.

Value iteration:

Note: if we have π_{opt} , we can obviously compute $V_{\pi_{\text{opt}}}(s)$ for all s .

$\Rightarrow$ How can we compute $V_{\pi_{\text{opt}}}(s)$ directly (i.e. without knowing π_{opt})

$$\begin{cases} V_{\text{opt}}(s) = 0 & \text{if end}(s). \\ V_{\text{opt}}(s) = \max_{a \in \text{Action}(s)} Q_{\text{opt}}(s, a). \\ Q_{\text{opt}}(s, a) = \sum_s T(s, a, s') [\text{reward}(s, a, s') + \gamma V_{\text{opt}}(s')] \end{cases}$$

With recurrence above, we can compute $V_{\text{opt}}(s)$ (DP).

Initialize $V_{\text{opt}}^0(s) = 0$ for all s

for $t = 1 \dots \text{Max-iter}$

$$V_{\text{opt}}^t(s) = \max_{a \in \text{Action}(s)} \sum_s T(s, a, s') [\text{reward}(s, a, s') + \gamma V_{\text{opt}}^{t-1}(s')]$$

Here, At each step, we record the optimal action, then we can define the policy (policy extraction)

Convergence:

if either the discount factor $\gamma < 1$ or the MDP graph is acyclic

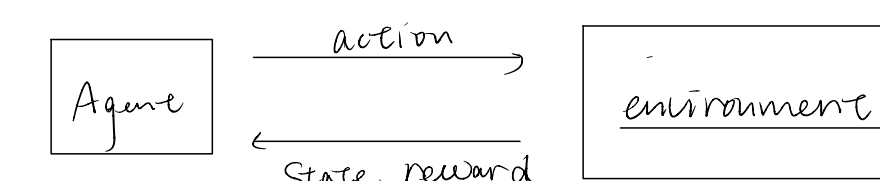
then both value iterations and policy iteration converge to the optimal solution

Reinforcement Learning

Goal: Learn a policy to maximize the rewards

Given: a set of actions to do in a state

Assuming: know the states, unlimited memory / computational resource



Initialization

Loop:

Choose an action.

Observe the Current State and record the reward.

Update parameters: (transition distribution ($T(s, a, s')$), expected values, expected Q-values)

Exploration / Exploitation

(try something new) (do the best move at current state)

$\Rightarrow$ long-term benefit $\Rightarrow$ greedy action

"How to balance them?"

Epsilon-Greedy

Epsilon ϵ : the probability of choosing to explore

($1 - \epsilon$: the probability of choosing to exploit)

Epsilon-Greedy

$$\pi(s) = \begin{cases} \arg \max_{a \in \text{Action}(s)} Q(s, a) & \text{with } p = 1 - \epsilon \text{ "exploit"} \\ \text{random } a \in \text{Action}(s) & \text{with } p = \epsilon \text{ "explore"} \end{cases}$$

Representation For Reinforcement Learning

"How to update parameters depends on how they are represented"

Tabular representation: transition relations, rewards, values, Q-values, policies.
All represented by tables

Approximation

Approximated by efficient functions such as neural networks.

Tabular representation

Main idea from Monte-Carlo Method: Average

Sample: $D = [s_1, a_1, r_1, s_2, a_2, r_2, s_3, \dots]$

$$\begin{cases} \hat{T}(s, a, s') = \frac{\#(s, a, s') \in D}{\sum_s \#(s, a, s') \in D} \\ \hat{\text{Reward}}(s, a, s') = \text{average of } \{r : (s, a, r, s') \in D\} \\ \hat{\text{Reward}}(s, a) = \text{average of } \{r : (s, a, r) \in D\} \end{cases}$$

Estimate a Model $\Rightarrow$ optimal policy / value

"Why not we directly estimate Q-value?"

Utility at s_t : $U_t = r_t + \gamma U_{t+1} + \gamma^2 U_{t+2} \dots$

$Q_\pi(s, a) = \text{average of } \{U_t \mid s_t = s, a_t = a\}$

Temporal Difference

MC Method: require an entire episode $[s_1, a_1, r_1, s_2, a_2, r_2, \dots]$,

Update Q-value at each step (Q-learning)

$$\hat{Q}(s, A) = (1 - \mu) \hat{Q}(s, A) + \underbrace{\mu}_{\text{learning rate}} \underbrace{(r + \gamma \hat{V}(s'))}_{\text{Having a new transition } (s, A, r, s')}$$

Current estimate of Q-value