

Review

* Markov decision process (MDP)

$$\{ \mathcal{S}, A, P(s'|s, a), R(s), \gamma \}$$

states, actions, transitions, rewards, discount factor

* Policy: mapping $\pi: \mathcal{S} \rightarrow A$ of states to actions

* Value functions

$$V^\pi(s) = E^\pi \left[\sum_{t=0}^{\infty} \gamma^t R(s_t) \mid s_0 = s \right] \quad \text{state}$$

$$Q^\pi(s, a) = E^\pi \left[\sum_{t=0}^{\infty} \gamma^t R(s_t) \mid s_0 = s, a_0 = a \right] \quad \text{action}$$

* Planning - given parameters of MDP what can we compute?

1. Policy evaluation - how to compute $V^\pi(s)$?

Solve linear equations:

$$\sum_{s'} [I(s, s') - \gamma P(s'|s, \pi(s))] V^\pi(s') = R(s)$$

2. Policy improvement - how to improve on $\pi(s)$?

Greedy policy $\pi'(s) = \operatorname{argmax}_a Q^\pi(s, a)$

Thm: $V^{\pi'}(s) \geq V^\pi(s)$ for all states s .

3. Policy iteration - how to compute $\pi^*(s)$?

Algorithm

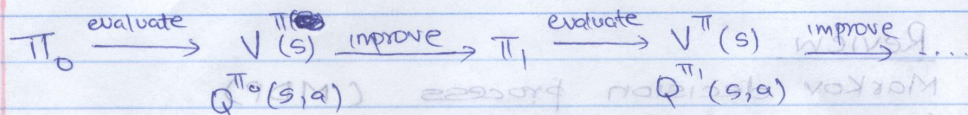
[NOT the algorithm in HW]

(1) initialize policy at random

(2) repeat until convergence

- compute state & action value function of current policy.

- derive greedy policy from action value function.



* Is policy iteration guaranteed to converge? Yes
 why? 1) # policies is finite $|A|^{|S|}$
 2) cannot cycle b/c policy improvements at every iteration.

Typically converges in far less steps than $|A|^{|S|}$

* Does it always converge to an optimal policy $\pi^*(s)$?
 Yes.

* thm: suppose $\pi'(s) = \pi(s)$ for all states s , or
 even more generally, that $V^{\pi'}(s) = V^{\pi}(s)$
 Then $V^{\pi'}(s) = V^*(s)$

Note: optimal value functions are unique, even if there are many optimal policies.

* Proof strategy: 1) Derive "Bellman optimality eqn" satisfied by $V^{\pi}(s)$ at convergence.
 2) Show that $V^{\pi}(s) \geq V^{\pi'}(s)$ for all states s and other policies π' . Hence $V^{\pi}(s) = V^*(s)$

Step 1

From Bellman eqn for $\pi'(s)$:

$$V^{\pi'}(s) = R(s) + \gamma \sum_{s'} P(s'|s, \pi'(s)) V^{\pi'}(s')$$

By assumption, $V^{\pi'}(s) = V^{\pi}(s)$ at convergence

$$V^{\pi}(s) = R(s) + \gamma \sum_{s'} P(s'|s, \pi'(s)) V^{\pi}(s')$$

By assumption, $\pi'(s)$ is greedy w.r.t. $V^\pi(s)$
 Hence:

$$V^\pi(s) = R(s) + \gamma \max_a \sum_{s'} P(s'|s, a) V^\pi(s')$$

"Bellman optimality equation" ↑ different than linear Bellman equation.

(set of n ~~linear~~ non-linear eqns for $s = 1, 2, \dots, n$
 non-linear b/c of max operation)

Step 2

Iterate RHS

$$V^\pi(s) = R(s) + \gamma \max_a \sum_{s'} P(s'|s, a) \left[R(s') + \gamma \max_{a'} \sum_{s''} P(s''|s', a') V^\pi(s'') \right]$$

$V^\pi(s')$

iterate again and again:

$$V^\pi(s) = R(s) + \gamma \max_a \sum_{s'} P(s'|s, a) \left[R(s') + \gamma \max_{a'} \sum_{s''} P(s''|s', a') \left[R(s'') + \gamma \max_{a''} \sum_{s'''} P(s'''|s'', a'') V^\pi(s''') \right] \right]$$

Now show that this iterated expression (taken out to an infinite # terms) implies optimality.

Let $\tilde{\pi}(s)$ be any other policy.

From Bellman eqn:

$$\begin{aligned} V^{\tilde{\pi}}(s) &= R(s) + \gamma \sum_{s'} P(s'|s, \tilde{\pi}(s)) V^{\tilde{\pi}}(s') \\ &\leq R(s) + \gamma \max_a \sum_{s'} P(s'|s, a) V^{\tilde{\pi}}(s') \quad \text{be greedy} \\ &= R(s) + \gamma \max_a \sum_{s'} P(s'|s, a) \left[R(s') + \gamma \sum_{s''} P(s''|s', \tilde{\pi}(s')) V^{\tilde{\pi}}(s'') \right] \\ &\leq R(s) + \gamma \max_a \sum_{s'} P(s'|s, a) \left[R(s') + \gamma \max_{a'} \sum_{s''} P(s''|s', a') V^{\tilde{\pi}}(s'') \right] \quad \text{be greedy} \end{aligned}$$

(2) V^{π} consider upper bound on $V^{\pi}(s)$ from iterating above t times.

(2) compare to equality for $V^{\pi}(s)$ after iterating t times.

As $t \rightarrow \infty$, RHS on upper bound on $V^{\pi}(s)$ converges to RHS of equality for $V^{\pi}(s)$.

Thus as $t \rightarrow \infty$

$$V^{\pi}(s) \leq \lim_{t \rightarrow \infty} [\dots] = \lim_{t \rightarrow \infty} [\dots] = V^{\pi}(s)$$

(2) V^{π} thus for all policies π and states s

we have $V^{\pi}(s) \leq V^*(s)$

or: $V^*(s) = \max_{\pi} V^{\pi}(s) = V^*(s)$

(2) V^{π} Pros / cons of policy iteration:

(+) converges quickly (in handful of steps often)

(-) each iteration requires policy evaluation $O(n^3)$

(4) Value iteration - how to compute $V^*(s)$ directly?

$$V^*(s) = \max_a Q^*(s, a)$$

$$= \max_a \left[R(s) + \gamma \max_{s'} \sum_{s'} P(s'|s, a) V^*(s') \right]$$

$$V^*(s) = R(s) + \gamma \max_a \sum_{s'} P(s'|s, a) V^*(s')$$

* n non-linear eqns for $s = 1, 2, \dots, n$

n unknowns $V^*(s)$

How to solve?

$$V^*(1) = R(1) + \gamma \max_a \sum_{s'} P(s'|1, a) V^*(s')$$

$$V^*(2) = R(2) + \gamma \max_a \sum_{s'} P(s'|2, a) V^*(s')$$

...

Algorithm: value iteration

(1) initialize $V_0(s) = 0$ for all states s

(2) iterate

$$V_{k+1}(s) = R(s) + \gamma \max_a \left[\sum_{s'} P(s'|s, a) V_k(s') \right]$$

for all $s = 1, 2, \dots, n$

Note: this algorithm works directly on value functions not policies.

But incremental policies can be computed from:

$$\pi_{k+1}(s) = \text{greedy}[V_k(s)] = \arg\max_a \left[\sum_{s'} P(s'|s, a) V_k(s') \right]$$

(3) Suppose this converges:

$$\lim_{k \rightarrow \infty} V_k(s) \rightarrow V^*(s)$$

$$\text{compute } \pi^*(s) = \arg\max_a \left[\sum_{s'} P(s'|s, a) V^*(s') \right]$$

Does algorithm converge? Yes.