

03/11/14

Review:

* Markov decision process

$$\text{MDP} = \{S, A, P(s'|sa), R(s), \gamma\}$$

* Value functions

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

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

* Bellman Optimality Equation.

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

* Algorithm: Value iteration $\leftarrow$ for computing $v^*(s)$ directly.

$\rightarrow$ initialize $v_0(s) = 0 \quad \forall s$

$\rightarrow$ iterate $v_{k+1}(s) = R(s) + \gamma \max_a \left[\sum_{s'} p(s'|s, a) v_k(s') \right]$

To prove: The algorithm converges

$\rightarrow v^*(s)$ is a fixed point of iteration

$\rightarrow$ There are no other fixed points.

$\rightarrow$ The algorithm always reaches $v^*(s)$

Before going to the proof of converges, we will prove a lemma

Lemma: for any functions $f(a)$ and $g(a)$

$$|\max_a f(a) - \max_a g(a)| \leq \max_a |f(a) - g(a)|$$

Proof of lemma:

$$f(a) - \max_{a'} g(a') \leq f(a) - g(a), \text{ for all } a$$

$$\begin{aligned} \max_a f(a) - \max_{a'} g(a') &\leq \max_a (f(a) - g(a)), \text{ max over } a \\ &\leq \max_a |f(a) - g(a)| \rightarrow \textcircled{1} \end{aligned}$$

By symmetry (exchanging f & g)

$$\max_a g(a) - \max_a f(a) \leq \max_a |f(a) - g(a)| \rightarrow \textcircled{2}$$

↳ same as $|f(a) - g(a)|$

The L.H.S of eq① or eq②, must correspond to absolute value of $\max_a f(a) - \max_a g(a)$

thus, we get

$$|\max_a f(a) - \max_a g(a)| \leq \max_a |f(a) - g(a)|$$

Using this lemma, we will prove

Theorem: value iteration converges.

$$\Rightarrow \lim_{k \rightarrow \infty} [v_k(s)] \rightarrow v^*(s) \text{ for all states } s$$

Proof: let $\Delta_k = \max_s |v_k(s) - v^*(s)|$

$$\Delta_{k+1} = \max_s |v_{k+1}(s) - v^*(s)|$$

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

$$= \gamma \max_s \left| \underbrace{\max_a \sum_{s'} P(s'|s,a) v_k(s')}_{f(a)} - \underbrace{\max_a \sum_{s'} P(s'|s,a) v^*(s')}_{g(a)} \right|$$

using the lemma

$$\leq \gamma \max_s \max_a \left| \sum_{s'} P(s'|s,a) [v_k(s') - v^*(s')] \right|$$

$$\leq \gamma \max_s \max_a \left| \left[\sum_{s'} P(s'|s,a) \right] \left[\max_{s''} |v_k(s') - v^*(s'')| \right] \right|$$

$$= \gamma \max_s \max_a |\Delta_k|$$

$$= \gamma \Delta_k \quad [\Delta_k \text{ does not depend on } a \text{ or } s]$$

Thus $\Delta_{k+1} \leq \gamma \Delta_k$, where $0 \leq \gamma < 1$

By induction, we can show

$$\boxed{\Delta_k \leq \gamma^k \Delta_0}$$

Thus $\Delta_k \rightarrow 0$ as $k \rightarrow \infty$

Assume all the rewards are bounded.

$$\Delta_0 = \max_s |V_0(s) - v^*(s)| = \max_s |v^*(s)| \quad [v_0(s) = 0]$$

$$\leq \max_s |R(s)| [1 + \gamma + \gamma^2 + \dots]$$

$$= \max_s |R(s)| \left(\frac{1}{1-\gamma} \right)$$

Thus $\Delta_k \leq \frac{\gamma^k}{1-\gamma} \max_s |R(s)| \rightarrow 0 \text{ as } k \rightarrow \infty$

The above inequality suggests that the convergence rate

→ depends on γ

→ more iterations are required as $\gamma \rightarrow 1$

FINAL CONVERGES UP TO HERE !!

REINFORCEMENT LEARNING (not for final)

* What if $P(s'/s, a)$ and $R(s)$ are not known?

Can we learn $\pi^*(s)$ or $v^*(s)$ from experience?

$$\Delta^* \gamma = \Delta$$

1) Model-based (indirect) approach

Explore world, estimate model $\hat{p}(s'|s,a) \approx p(s'|s,a)$,

Compute $\hat{\pi}^*(s)$ or $\hat{v}^*(s)$ from $\hat{p}(s'|s,a)$

* Cons: $\rightarrow$ to store $p(s'|s,a)$ is $O(n^2)$ for n states

$\rightarrow$ only care about $\pi^*(s)$ or $v^*(s)$ which are $O(n)$

$\rightarrow$ Is it really necessary to estimate a model?

* Pros: $\rightarrow$ model $p(s'|s,a)$ useful for task transfer, whose rewards $R(s)$ or discount factor γ changes, but $p(s'|s,a)$ stay the same.

2) Direct Approach: learn $\pi^*(s)$, $v^*(s)$ w/o building a model

Stochastic Approximation Theory:

* To estimate mean of random variable, x from samples $x_0, x_1, x_2, \dots, x_{T-1}$?

1) obvious - sample average

$$\mu = \frac{1}{T} [x_0 + x_1 + \dots + x_{T-1}] \rightarrow E[x] \text{ as } T \rightarrow \infty$$

Converges by law of large numbers

2) Incremental update:

→ Initialize $\mu_0 = 0$

→ update $\mu_t = (1 - \alpha_t) \mu_{t-1} + \alpha_t x_t$ for $0 < \alpha_t < 1$

This can be also written as

$$\mu_t = \mu_{t-1} + \alpha_t (\underbrace{x_t - \mu_{t-1}}_{\text{temporal difference}})$$

This ^{algorithm} equation is called Temporal Difference algorithm.

Theorem: $\mu \rightarrow E[x]$ as $t \rightarrow \infty$ if

i) α_t doesn't decay too fast $\rightarrow \sum_{t=1}^{\infty} \alpha_t = \infty$

ii) α_t doesn't decay too slow $\rightarrow \sum_{t=1}^{\infty} \alpha_t^2 < \infty$ is finite

One possible choice of α_t :

choose $\alpha_t = \frac{1}{t}$. ~~also satisfy~~ satisfies the conditions.

This actually gives the same estimate as the running average.