

10-301/601: Introduction to Machine Learning

Lecture 22: Value and Policy Iteration

Henry Chai & Matt Gormley

11/16/22

Front Matter

- Announcements
 - HW7 released 11/11, due 11/21 at 11:59 PM
 - Please be mindful of your grace day usage
 - Apply to be a 10-301/601 TA!
 - Applications can be found at <https://www.ml.cmu.edu/academics/ta.html> and are due on Thursday, November 17th

Recall: Value Function

- Find a policy $\pi^* = \operatorname{argmax}_{\pi} V^{\pi}(s) \quad \forall s \in \mathcal{S}$
- $V^{\pi}(s) = \mathbb{E}[\text{discounted total reward of starting in state } s \text{ and executing policy } \pi \text{ forever}]$

$$= \mathbb{E}_{\underline{p(s' | s, a)}}[R(s_0 = s, \pi(s_0)) + \gamma R(s_1, \pi(s_1)) + \gamma^2 R(s_2, \pi(s_2)) + \dots]$$

$$= \sum_{t=0}^{\infty} \gamma^t \mathbb{E}_{p(s' | s, a)}[R(s_t, \pi(s_t))]$$

where $0 < \gamma < 1$ is some discount factor for future rewards

- Assumes a *stochastic* transition function and a *deterministic* reward function

Recall: Value Function

- Find a policy $\pi^* = \operatorname{argmax}_{\pi} V^{\pi}(s) \quad \forall s \in \mathcal{S}$
- $V^{\pi}(s) = \mathbb{E}[\text{discounted total reward of starting in state } s \text{ and executing policy } \pi \text{ forever}]$

$$\begin{aligned} &= R(s_0 = s, \pi(s_0)) + \gamma R(s_1 = \underline{\delta(s_0, \pi(s_0))}, \pi(s_1)) \\ &+ \gamma^2 R(s_2 = \underline{\delta(s_1, \pi(s_1))}, \pi(s_2)) + \dots \\ &= \sum_{t=0}^{\infty} \gamma^t \mathbb{E}_{p(s' | s, a)} [R(s_t, \pi(s_t))] \end{aligned}$$

where $0 < \gamma < 1$ is some discount factor for future rewards

- Assumes a *deterministic* transition function and a *deterministic* reward function

Recall: Value Function

- Find a policy $\pi^* = \operatorname{argmax}_{\pi} V^{\pi}(s) \quad \forall s \in \mathcal{S}$
- $V^{\pi}(s) = \mathbb{E}[\text{discounted total reward of starting in state } s \text{ and executing policy } \pi \text{ forever}]$

$$= \mathbb{E}_{p(s' | s, a)} [R(s_0 = s, \pi(s_0)) + \gamma R(s_1, \pi(s_1)) + \gamma^2 R(s_2, \pi(s_2)) + \dots]$$

$$= \sum_{t=0}^{\infty} \gamma^t \mathbb{E}_{\underline{p(s' | s, a)}} [R(s_t, \pi(s_t))]$$

where $0 < \gamma < 1$ is some discount factor for future rewards

- Assumes a *stochastic* transition function and a *deterministic* reward function

Value Function

$$\mathbb{E}[A] = \sum_{\text{all possible values of } A, v} \mathbb{P}(A=v) v$$

- $V^\pi(s)$ = $\mathbb{E}$ [discounted total reward of starting in state s and executing policy π forever]

$$\begin{aligned} & \downarrow \\ &= \mathbb{E}[\underbrace{R(s_0, \pi(s_0))}_{\text{immediate reward}} + \gamma R(s_1, \pi(s_1)) + \gamma^2 R(s_2, \pi(s_2)) + \dots \mid \underbrace{s_0 = s}_{\text{initial state}}] \\ &= R(s, \pi(s)) + \gamma \underbrace{\mathbb{E}[R(s_1, \pi(s_1)) + \gamma R(s_2, \pi(s_2)) + \dots \mid s_0 = s_1]}_{\text{expected future discounted reward from } s_1} \\ &= R(s, \pi(s)) + \gamma \underbrace{\sum_{s_1 \in \mathcal{S}} p(s_1 \mid s, \pi(s))}_{\text{transition probabilities}} \underbrace{(R(s_1, \pi(s_1)) + \gamma \mathbb{E}[R(s_2, \pi(s_2)) + \dots \mid s_1])}_{\text{value of next state } s_1} \end{aligned}$$

Value Function

- $V^\pi(s) = \mathbb{E}[\text{discounted total reward of starting in state } s \text{ and executing policy } \pi \text{ forever}]$
$$= \mathbb{E}[R(s_0, \pi(s_0)) + \gamma R(s_1, \pi(s_1)) + \gamma^2 R(s_2, \pi(s_2)) + \dots \mid s_0 = s]$$
$$= R(s, \pi(s)) + \gamma \mathbb{E}[R(s_1, \pi(s_1)) + \gamma R(s_2, \pi(s_2)) + \dots \mid s_0 = s]$$
$$= R(s, \pi(s)) + \gamma \sum_{s_1 \in \mathcal{S}} p(s_1 \mid s, \pi(s)) (R(s_1, \pi(s_1)) + \gamma \mathbb{E}[R(s_2, \pi(s_2)) + \dots \mid s_1])$$

Value Function

- $V^\pi(s) = \mathbb{E}[\text{discounted total reward of starting in state } s \text{ and executing policy } \pi \text{ forever}]$
$$= \mathbb{E}[R(s_0, \pi(s_0)) + \gamma R(s_1, \pi(s_1)) + \gamma^2 R(s_2, \pi(s_2)) + \dots \mid s_0 = s]$$
$$\rightarrow = R(s, \pi(s)) + \gamma \mathbb{E}[R(s_1, \pi(s_1)) + \gamma R(s_2, \pi(s_2)) + \dots \mid s_0 = s]$$
$$= R(s, \pi(s)) + \gamma \sum_{s_1 \in \mathcal{S}} p(s_1 \mid s, \pi(s)) \left(\underbrace{R(s_1, \pi(s_1)) + \gamma \mathbb{E}[R(s_2, \pi(s_2)) + \dots \mid s_1]}_{\text{Value function of } s_1} \right)$$

Value Function

- $V^\pi(s) = \mathbb{E}[\text{discounted total reward of starting in state } s \text{ and executing policy } \pi \text{ forever}]$
 $= \mathbb{E}[R(s_0, \pi(s_0)) + \gamma R(s_1, \pi(s_1)) + \gamma^2 R(s_2, \pi(s_2)) + \dots \mid s_0 = s]$
 $\rightarrow = R(s, \pi(s)) + \gamma \mathbb{E}[R(s_1, \pi(s_1)) + \gamma R(s_2, \pi(s_2)) + \dots \mid s_0 = s]$
 $= R(s, \pi(s)) + \gamma \sum_{s_1 \in \mathcal{S}} p(s_1 \mid s, \pi(s)) (R(s_1, \pi(s_1)) + \gamma \mathbb{E}[R(s_2, \pi(s_2)) + \dots \mid s_1])$

Value Function

↙ arbitrarily policy

- $V^\pi(s) = \mathbb{E}[\text{discounted total reward of starting in state } s \text{ and executing policy } \pi \text{ forever}]$

$$= \mathbb{E}[R(s_0, \pi(s_0)) + \gamma R(s_1, \pi(s_1)) + \gamma^2 R(s_2, \pi(s_2)) + \dots \mid s_0 = s]$$

$$= R(s, \pi(s)) + \gamma \mathbb{E}[R(s_1, \pi(s_1)) + \gamma R(s_2, \pi(s_2)) + \dots \mid s_0 = s]$$

$$= R(s, \pi(s)) + \gamma \sum_{s_1 \in \mathcal{S}} p(s_1 \mid s, \pi(s)) (R(s_1, \pi(s_1)) + \gamma \mathbb{E}[R(s_2, \pi(s_2)) + \dots \mid s_1])$$

$$V^\pi(s) = R(s, \pi(s)) + \gamma \sum_{s_1 \in \mathcal{S}} p(s_1 \mid s, \pi(s)) V^\pi(s_1) \quad \forall s \in \mathcal{S}$$

Bellman equations

Optimality

- Optimal value function:

$$V^*(s) = \max_{a \in \mathcal{A}} \left\{ R(s, a) + \gamma \sum_{s' \in \mathcal{S}} p(s' | s, a) V^*(s') \right\}$$

- System of $|\mathcal{S}|$ equations and $|\mathcal{S}|$ variables

- Optimal policy:

$$\pi^*(s) = \operatorname{argmax}_{a \in \mathcal{A}} \underbrace{R(s, a)}_{\text{Immediate reward}} + \gamma \underbrace{\sum_{s' \in \mathcal{S}} p(s' | s, a) V^*(s')}_{\text{Expected (Discounted) Future reward}}$$

- Insight: if you know the optimal value function, you can solve for the optimal policy!

Fixed Point Iteration

- Iterative method for solving a system of equations
- Given some equations and initial values

$$x_1 = f_1(x_1, \dots, x_n)$$

$$\vdots$$

$$x_n = f_n(x_1, \dots, x_n)$$

$$\underline{x_1^{(0)}, \dots, x_n^{(0)}}$$

- While not converged, do

$$x_1^{(t+1)} \leftarrow f_1(x_1^{(t)}, \dots, x_n^{(t)})$$

$$\vdots$$

$$x_n^{(t+1)} \leftarrow f_n(x_1^{(t)}, \dots, x_n^{(t)})$$

Fixed Point Iteration: Example

$$\begin{aligned}
 & (0)(0) + \frac{1}{2} = \frac{1}{2} \\
 & x_1 = x_1 x_2 + \frac{1}{2} \\
 & x_2 = -\frac{3x_1}{2} \\
 & x_1^{(0)} = x_2^{(0)} = 0 \\
 & \hat{x}_1 = \frac{1}{3}, \hat{x}_2 = -\frac{1}{2}
 \end{aligned}$$

Handwritten calculations in red:

$$\begin{aligned}
 & \left(\frac{1}{3}\right)\left(-\frac{1}{2}\right) + \frac{1}{2} \\
 & = -\frac{1}{6} + \frac{3}{6} \\
 & = \frac{2}{6} = \frac{1}{3} \checkmark \\
 & -\frac{3\left(\frac{1}{3}\right)}{2} = -\frac{1}{2} \checkmark
 \end{aligned}$$

t	$x_1^{(t)}$	$x_2^{(t)}$
0	0	0
1	0.5	0

Value Iteration

- Inputs: $R(s, a)$, $p(s' | s, a)$, $0 < \gamma < 1$
- Initialize $V^{(0)}(s) = 0 \forall s \in \mathcal{S}$ (or randomly) and set $t = 0$
- While not converged, do:

- For $s \in \mathcal{S}$

$$V^{(t+1)}(s) \leftarrow \max_{a \in \mathcal{A}} R(s, a) + \gamma \sum_{s' \in \mathcal{S}} p(s' | s, a) V^{(t)}(s')$$

$Q(s, a)$

- $t = t + 1$

- For $s \in \mathcal{S}$

$$\pi^*(s) \leftarrow \operatorname{argmax}_{a \in \mathcal{A}} R(s, a) + \gamma \sum_{s' \in \mathcal{S}} p(s' | s, a) V^{(t)}(s')$$

- Return π^*

Value Iteration

- Inputs: $R(s, a)$, $p(s' | s, a)$, $0 < \gamma < 1$
- Initialize $V^{(0)}(s) = 0 \forall s \in \mathcal{S}$ (or randomly) and set $t = 0$
- While not converged, do:
 - For $s \in \mathcal{S}$
 - For $a \in \mathcal{A}$
$$Q(s, a) = R(s, a) + \gamma \sum_{s' \in \mathcal{S}} p(s' | s, a) V^{(t)}(s')$$
 - $V^{(t+1)}(s) \leftarrow \max_{a \in \mathcal{A}} Q(s, a)$
 - $t = t + 1$
- For $s \in \mathcal{S}$
$$\pi^*(s) \leftarrow \operatorname{argmax}_{a \in \mathcal{A}} R(s, a) + \gamma \sum_{s' \in \mathcal{S}} p(s' | s, a) V^{(t)}(s')$$
- Return π^*

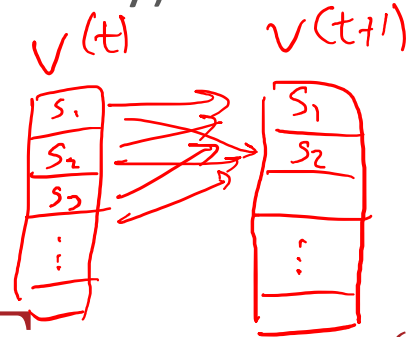
Synchronous Value Iteration

- Inputs: $R(s, a)$, $p(s' | s, a)$, $0 < \gamma < 1$
- Initialize $V^{(0)}(s) = 0 \forall s \in \mathcal{S}$ (or randomly) and set $t = 0$
- While not converged, do:

- For $s \in \mathcal{S}$
 - For $a \in \mathcal{A}$

$$Q(s, a) = R(s, a) + \gamma \sum_{s' \in \mathcal{S}} p(s' | s, a) V^{(t)}(s')$$

- $V^{(t+1)}(s) \leftarrow \max_{a \in \mathcal{A}} Q(s, a)$
- $t = t + 1$
- For $s \in \mathcal{S}$
$$\pi^*(s) \leftarrow \operatorname{argmax}_{a \in \mathcal{A}} R(s, a) + \gamma \sum_{s' \in \mathcal{S}} p(s' | s, a) V^{(t)}(s')$$
- Return π^*



Asynchronous Value Iteration

- Inputs: $R(s, a)$, $p(s' | s, a)$, $0 < \gamma < 1$
- Initialize $V^{(0)}(s) = 0 \forall s \in \mathcal{S}$ (or randomly) and set $t = 0$
- While not converged, do:

- For $s \in \mathcal{S}$

- For $a \in \mathcal{A}$

$$Q(s, a) = R(s, a) + \gamma \sum_{s' \in \mathcal{S}} p(s' | s, a) V(s')$$

- $V(s) \leftarrow \max_{a \in \mathcal{A}} Q(s, a)$

- For $s \in \mathcal{S}$

$$\pi^*(s) \leftarrow \operatorname{argmax}_{a \in \mathcal{A}} R(s, a) + \gamma \sum_{s' \in \mathcal{S}} p(s' | s, a) V(s')$$

- Return π^*



Poll Question 1: What is the runtime of one iteration of value iteration?

- Inputs: $R(s, a)$, $p(s' | s, a)$, $0 < \gamma < 1$
- Initialize $V^{(0)}(s) = 0 \forall s \in \mathcal{S}$ (or randomly) and set $t = 0$
- While not converged, do:

- For $s \in \mathcal{S}$ $|\mathcal{S}|$

- For $a \in \mathcal{A}$ $|\mathcal{A}|$

$$\star Q(s, a) = R(s, a) + \gamma \sum_{s' \in \mathcal{S}} p(s' | s, a) V(s')$$

$$\star V(s) \leftarrow \max_{a \in \mathcal{A}} Q(s, a)$$

$$O(|\mathcal{S}||\mathcal{A}||\mathcal{S}| + |\mathcal{S}||\mathcal{A}|) \\ = O(|\mathcal{S}|^2|\mathcal{A}|)$$

- For $s \in \mathcal{S}$

$$\pi^*(s) \leftarrow \operatorname{argmax}_{a \in \mathcal{A}} R(s, a) + \gamma \sum_{s' \in \mathcal{S}} p(s' | s, a) V(s')$$

- Return π^*

Value Iteration Theory

- **Theorem 1:** Value function convergence

V will converge to V^* if each state is “visited”
infinitely often (Bertsekas, 1989)

- **Theorem 2:** Convergence criterion

if $\max_{s \in \mathcal{S}} |V^{(t+1)}(s) - V^{(t)}(s)| < \epsilon$, *reverse engineer this value*

then $\max_{s \in \mathcal{S}} |V^{(t+1)}(s) - V^*(s)| < \frac{2\epsilon\gamma}{1-\gamma}$ (Williams & Baird, 1993)

- **Theorem 3:** Policy convergence

The “greedy” policy, $\pi(s) = \operatorname{argmax}_{a \in \mathcal{A}} Q(s, a)$, converges to the *how close do you want to be to optimal?*

optimal π^* in a finite number of iterations, often before
the value function has converged! (Bertsekas, 1987)

Policy Iteration

- Inputs: $R(s, a)$, $p(s' | s, a)$, $0 < \gamma < 1$

- Initialize π randomly

- While not converged, do:

- Solve the Bellman equations defined by policy π

$$V^\pi(s) = \underbrace{R(s, \pi(s))}_{\text{red underline}} + \underbrace{\gamma \sum_{s' \in \mathcal{S}} \underbrace{p(s' | s, \pi(s))}_{\text{red underline}} \underbrace{V^\pi(s')}_{\text{red underline}}}_{\text{red underline}}$$

- Update π

$$\pi(s) \leftarrow \operatorname{argmax}_{a \in \mathcal{A}} R(s, a) + \gamma \sum_{s' \in \mathcal{S}} p(s' | s, a) \underbrace{V^\pi(s')}_{\text{red underline}}$$

- Return π

convergence = policy stops for some

Poll Question 2: What is an upper bound on the number of possible policies?

- Inputs: $R(s, a)$, $p(s' | s, a)$, $0 < \gamma < 1$
- Initialize π randomly
- While not converged, do:

- Solve the Bellman equations defined by policy π

$$\left\{ \begin{array}{l} V^\pi(s) = \underbrace{R(s, \pi(s))}_{\text{reward}} + \gamma \sum_{s' \in \mathcal{S}} \underbrace{p(s' | s, \pi(s))}_{\text{transition prob}} V^\pi(s') \end{array} \right.$$

- Update π

$$\pi(s) \leftarrow \operatorname{argmax}_{a \in \mathcal{A}} R(s, a) + \gamma \sum_{s' \in \mathcal{S}} p(s' | s, a) V^\pi(s')$$

- Return π

Policy Iteration Theory

- In policy iteration, the policy improves in each iteration.
- Given finite state and action spaces, there are finitely many possible policies
- Thus, the number of iterations needed to converge is bounded!
- Policy iteration takes $O(|\mathcal{S}|^2|\mathcal{A}| + |\mathcal{S}|^3)$ time / iteration
 - However, empirically policy iteration requires fewer iterations to converge than value iteration

Two big Q's

1. What can we do if the reward and/or transition functions/distributions are unknown?
2. How can we handle infinite (or just very large) state/action spaces?

MDP and Value/Policy Iteration Learning Objectives

You should be able to...

- Compare reinforcement learning to other learning paradigms
- Cast a real-world problem as a Markov Decision Process
- Depict the exploration vs. exploitation tradeoff via MDP examples
- Explain how to solve a system of equations using fixed point iteration
- Define the Bellman Equations
- Show how to compute the optimal policy in terms of the optimal value function
- Implement value iteration and policy iteration
- Contrast the computational complexity and empirical convergence of value iteration vs. policy iteration
- Identify the conditions under which the value iteration algorithm will converge to the true value function
- Describe properties of the policy iteration algorithm