

CS 4700: Foundations of Artificial Intelligence

Spring 2019
Prof. Haym Hirsh

Lecture 15
February 27, 2019

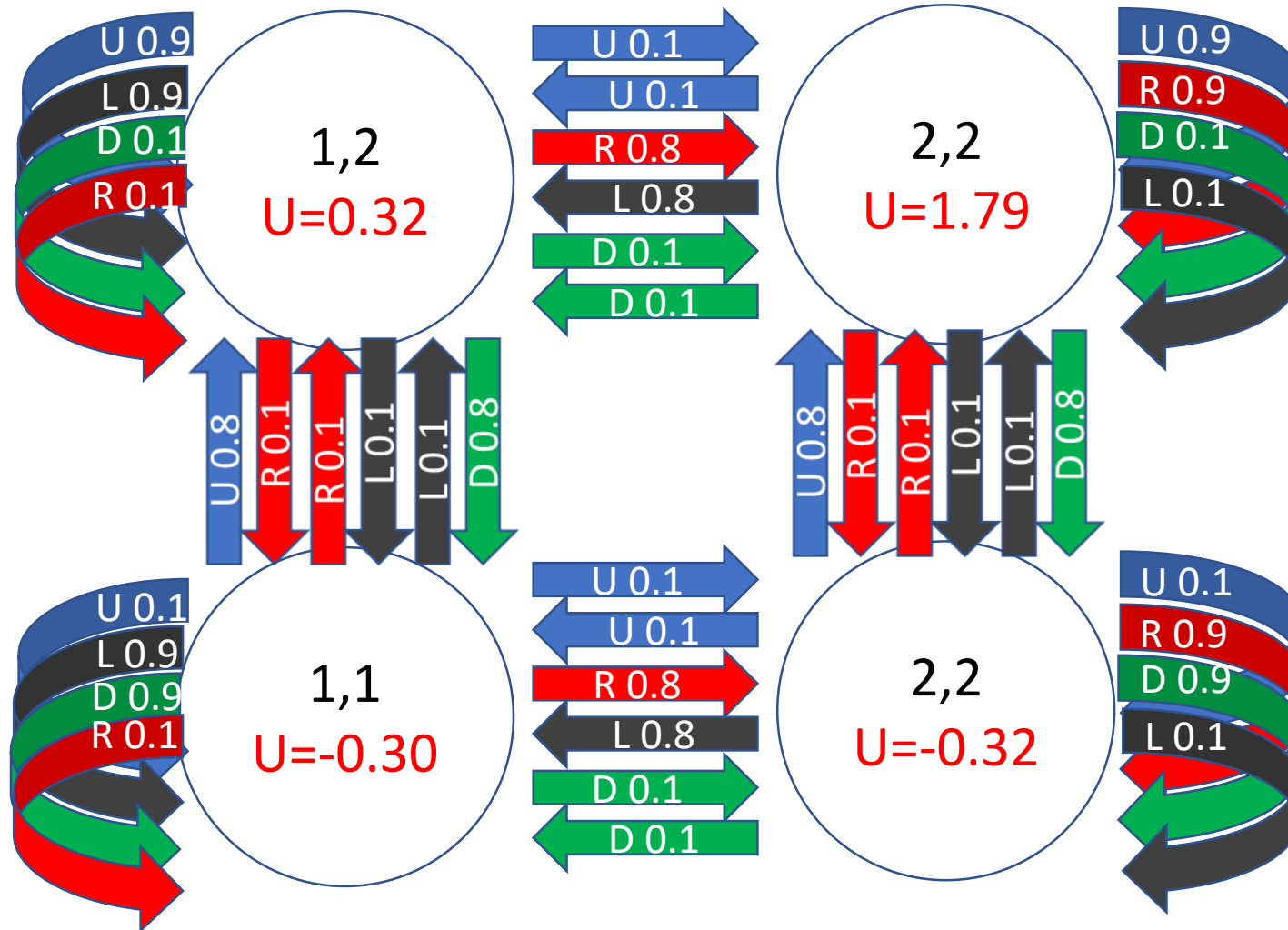
Announcements

- Karma lectures:
 - Tomorrow: 4:15, Gates G01:
“Towards Generalization and Efficiency in Reinforcement Learning”
 - Friday: 12:15, Gates 122:
“Teaching Machines Like We Teach People”
- Textbook:
 - Chapter numbers and chapter filenames don’t always match up
 - Example: Chapter 20 is in Newchap21.pdf
 - When I refer to the textbook I’ll be using Chapter NUMBERS
 - Example: Reinforcement learning is Chapter 20
- Homework 2 grades released
 - Regrade requests due March 5, 12 noon
 - Be sensible

Announcements

- Other upcoming events:
 - Tomorrow: 5:30-8:00pm G01:
Cornell CS Research Night
 - This weekend:
Microsoft Azure Hackathon – anyone know more?

Policy Π_{UR} = Go up if you're in row 1, otherwise go right
 $\gamma = 0.5$



But what about Π^* ?

But what about Π^* ?

Use Policy Iteration

But what about Π^* ?

Use Policy Iteration
(among various approaches)

Finding Π^* : Policy Iteration

Policy_Iteration(S, A, P, R, γ):

For all $s \in S$ { $\hat{\pi}_0(s) \leftarrow \text{random } a \in A$; $\hat{U}_0(s) \leftarrow 0$; $i \leftarrow 0$ }

Repeat

$i \leftarrow i+1$

 For all $s \in S$

$$\hat{U}_i(s) \leftarrow R(s) + \gamma \sum_{s' \in S} [P(s' | s, \hat{\pi}_{i-1}(s)) \hat{U}_{i-1}(s')]$$

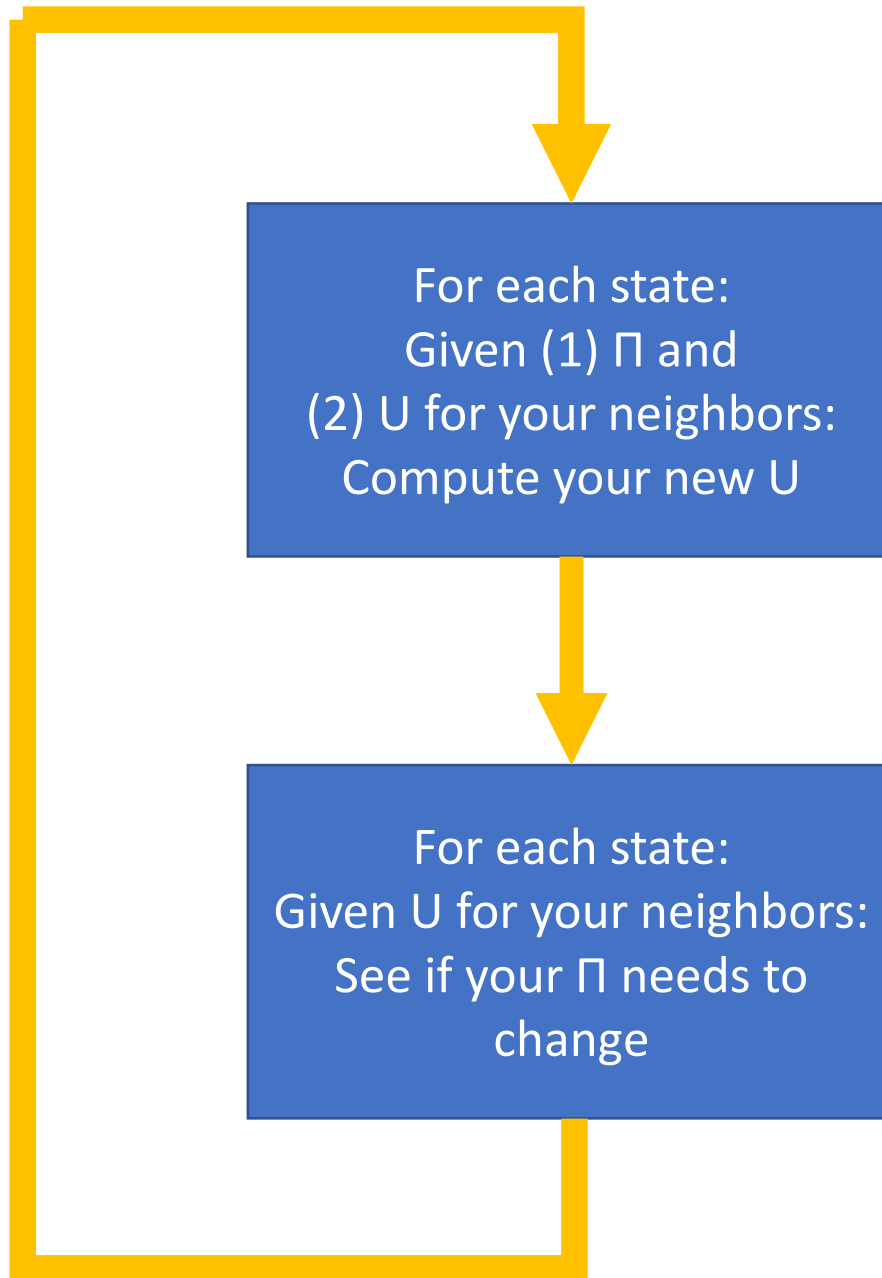
 For all $s \in S$

$$\text{If } \max_{a \in A} \sum_{s' \in S} [P(s' | s, a) \hat{U}_i(s')] > \sum_{s' \in S} [P(s' | s, \hat{\pi}_{i-1}(s)) \hat{U}_i(s')]$$

$$\text{Then } \hat{\pi}_i(s) \leftarrow \operatorname{argmax}_{a \in A} \sum_{s' \in S} [P(s' | s, a) \hat{U}_i(s')]$$

$$\text{Else } \hat{\pi}_i(s) \leftarrow \hat{\pi}_{i-1}(s)$$

Until <stopping condition> /* for example: $\hat{\pi}$ doesn't change */



Use:

$$U^*(s) = R(s) + \gamma \sum_{s' \in S} P(s'|s, \Pi^*(s)) U^*(s')$$

Use:

$$\Pi^*(s) = \operatorname{argmax}_{a \in A} \sum_{s' \in S} P(s'|s, a) U^*(s')$$

Initial values:
Random policy Π
 $U(s)=0$ for all states

For each state:
Given (1) Π and
(2) U for your neighbors:
Compute your new U

Use:

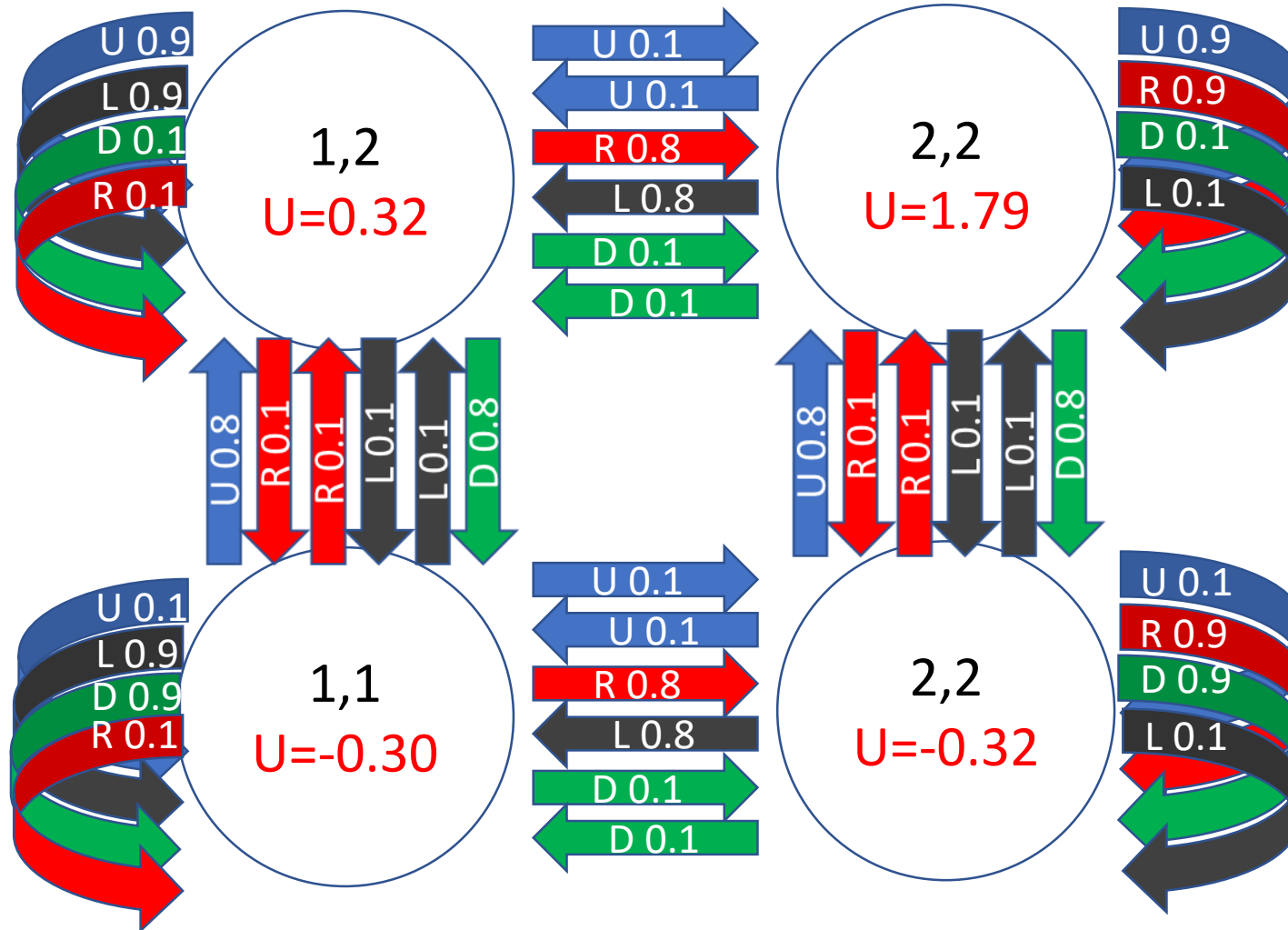
$$U^*(s) = R(s) + \gamma \sum_{s' \in S} P(s'|s, \Pi^*(s)) U^*(s')$$

For each state:
Given U for your neighbors:
See if your Π needs to
change

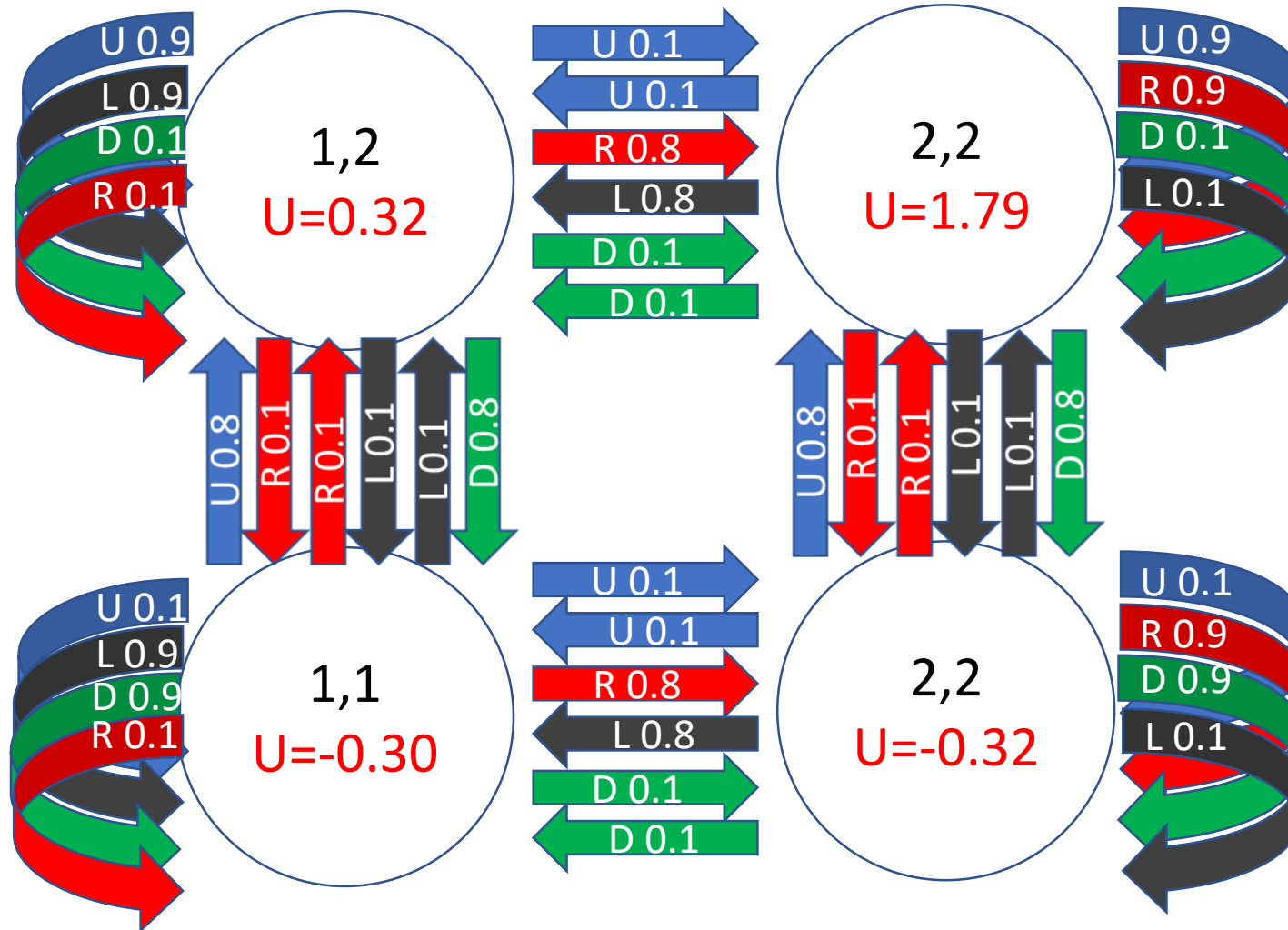
Use:

$$\Pi^*(s) = \operatorname{argmax}_{a \in A} \sum_{s' \in S} P(s'|s, a) U^*(s')$$

Policy Π_{UR} = Go up if you're in row 1, otherwise go right
 $\gamma = 0.5$



Π^* ?



Example

<https://bit.ly/2T6lYWD>

(You will need to copy it to your own spreadsheet to play with it)

What if we don't know P and R ?

What if we don't know P and R ?

Are we
learning while moving about in the world
(active reinforcement learning)?

or

learning from historical traces
(passive reinforcement learning)?

What if we don't know P and R ?

Are we

learning while moving about in the world
(active reinforcement learning)?

or

learning from historical traces
(passive reinforcement learning)?

Active Reinforcement Learning

Key concept:

You are getting (possibly bad) rewards while exploring

How do you balance learning about the world vs taking a hit for doing so?

When do you stop exploring because you've learned enough?

Active Reinforcement Learning

Key concept:

You are getting (possibly bad) rewards while exploring

How do you balance learning about the world vs taking a hit for doing so?

When do you stop exploring because you've learned enough?

Exploration vs Exploitation

Active Reinforcement Learning with Q Learning

Key idea: $Q: S \times A \rightarrow \mathbb{R}$

$Q(s,a)$: Expected value of taking action a and acting optimally thereafter

Active Reinforcement Learning with Q Learning

Key idea: $Q: S \times A \rightarrow \mathbb{R}$

$Q(s,a)$: Expected value of taking action a and acting optimally thereafter

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

$$\Pi^*(s) = \operatorname{argmax}_a Q(s, a)$$

Active Reinforcement Learning with Q Learning

Key idea: $Q: S \times A \rightarrow \mathbb{R}$

$Q(s,a)$: Expected value of taking action a and acting optimally thereafter

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

$$\Pi^*(s) = \operatorname{argmax}_a Q(s, a)$$

If I have Q I know how to act (even without having R and P)

Active Reinforcement Learning with Q Learning

Key idea: $Q: S \times A \rightarrow \mathbb{R}$

$Q(s,a)$: Expected value of taking action a and acting optimally thereafter

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

$$\Pi^*(s) = \operatorname{argmax}_a Q(s, a)$$

If I have Q I know how to act (even without having R and P)

We're going to learn Q

Q Learning

How do we learn Q?

Q Learning

How do we learn Q?

Want something analogous to this (from Policy Iteration):

$$\hat{U}_i(s) \leftarrow R(s) + \gamma \sum_{s' \in S} [P(s' | s, \hat{\pi}_{i-1}(s)) \hat{U}_{i-1}(s')]$$

$$\hat{\pi}_i(s) \leftarrow \operatorname{argmax}_{a \in A} \sum_{s' \in S} [P(s' | s, a) \hat{U}_i(s')]$$

Q Learning

How do we learn Q?

Want something analogous to this (from Policy Iteration):

$$\hat{U}_i(s) \leftarrow R(s) + \gamma \sum_{s' \in S} [P(s' | s, \hat{\pi}_{i-1}(s)) \hat{U}_{i-1}(s')]$$

$$\hat{\pi}_i(s) \leftarrow \operatorname{argmax}_{a \in A} \sum_{s' \in S} [P(s' | s, a) \hat{U}_i(s')]$$

Which are in turn derived from

$$U^*(s) = R(s) + \gamma \sum_{s' \in S} P(s' | s, \pi^*(s)) U^*(s')$$

$$\pi^*(s) = \operatorname{argmax}_{a \in A} \sum_{s' \in S} P(s' | s, a) U^*(s')$$

Q Learning

$$U^*(s) = R(s) + \gamma \sum_{s' \in S} P(s'|s, \pi^*(s)) U^*(s')$$

Q Learning

$$U^*(s) = R(s) + \gamma \sum_{s' \in S} P(s'|s, \pi^*(s)) U^*(s')$$

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

Q Learning

$$U^*(s) = R(s) + \gamma \sum_{s' \in S} P(s'|s, \pi^*(s)) U^*(s')$$

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

↓

$$U^*(s) = R(s) + \gamma \sum_{s' \in S} P(s'|s, \pi^*(s)) \max_{a' \in A_Q(s')} Q(s', a')$$

Q Learning

$$U^*(s) = R(s) + \gamma \sum_{s' \in S} P(s'|s, \pi^*(s)) U^*(s')$$

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

↓

$$U^*(s) = R(s) + \gamma \sum_{s' \in S} P(s'|s, \pi^*(s)) \max_{a' \in A_Q}(s', a')$$

↓

$$Q(s, a) = R(s) + \gamma \sum_{s' \in S} P(s'|s, a) \max_{a' \in A_Q}(s', a')$$

Q Learning

Policy iteration update rule:

$$\hat{U}_i(s) \leftarrow R(s) + \gamma \sum_{s' \in S} [P(s' | s, \hat{\pi}_{i-1}(s)) \hat{U}_{i-1}(s')]$$

Replaces $\hat{U}_i(s)$ with one-step lookahead using $\hat{U}_{i-1}(s')$

Can we do something analogous for $Q(s,a)$?