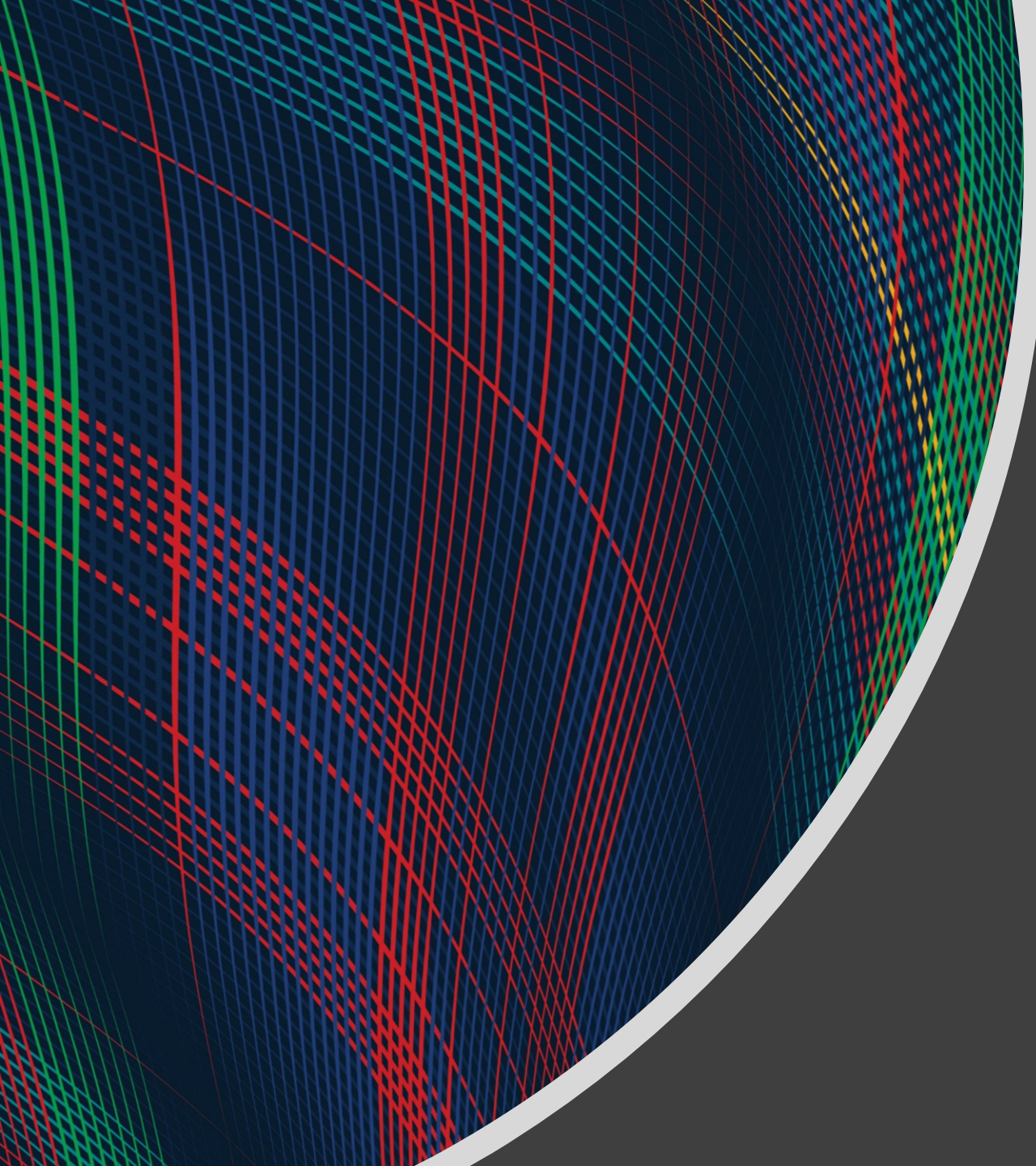


An abstract graphic on the left side of the slide, featuring a sphere-like shape composed of a dense grid of intersecting red, green, and blue lines. The lines are curved and follow the contours of the sphere, creating a complex, woven pattern. The sphere is set against a dark gray background.

07-280
AI & ML I

Reinforcement Learning

Instructors:

Pat Virtue & Nihar Shah

Reinforcement learning

What if we didn't know $P(s'|s, a)$ and $R(s, a, s')$?

Value iteration:

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

Q-iteration:

$$Q_{k+1}(s, a) = \sum_{s'} \cancel{P(s'|s, a)} [\cancel{R(s, a, s')} + \gamma \max_{a'} Q_k(s', a')], \quad \forall s, a$$

Policy extraction:

$$\pi_V(s) = \operatorname{argmax}_a \sum_{s'} \cancel{P(s'|s, a)} [\cancel{R(s, a, s')} + \gamma V(s')], \quad \forall s$$

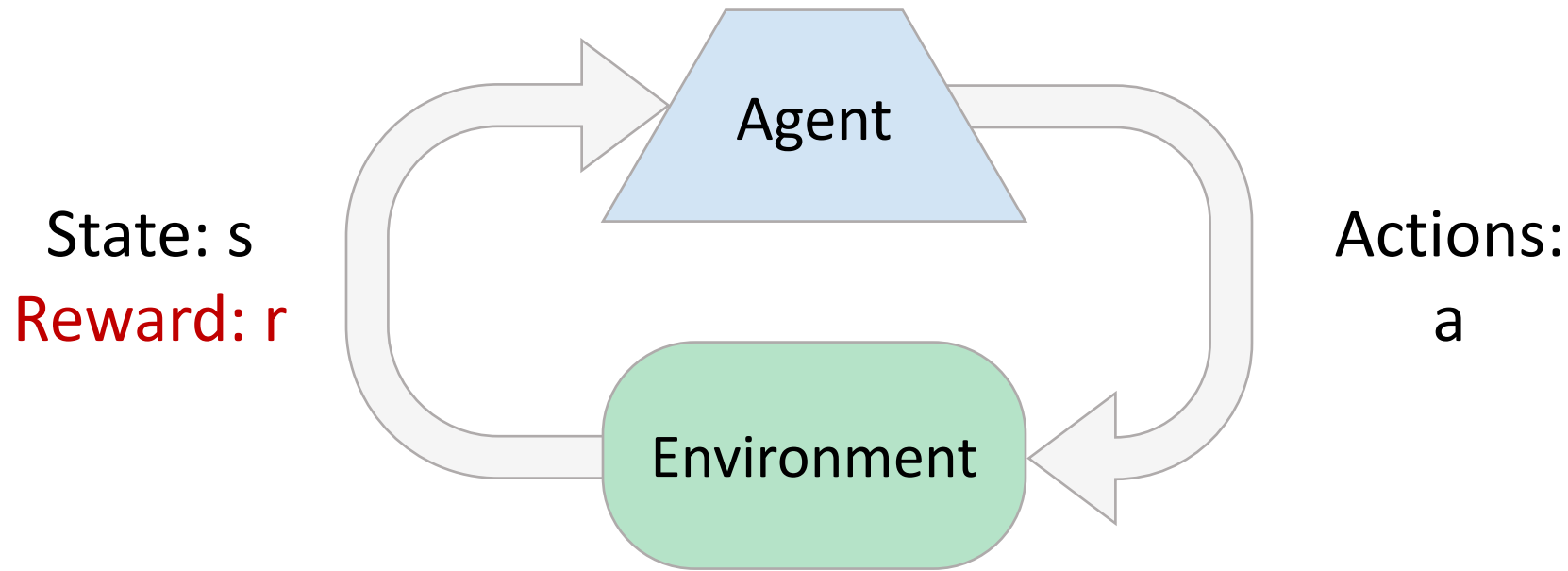
Policy evaluation:

$$V_{k+1}^\pi(s) = \sum_{s'} \cancel{P(s'|s, \pi(s))} [\cancel{R(s, \pi(s), s')} + \gamma V_k^\pi(s')], \quad \forall s$$

Policy improvement:

$$\pi_{new}(s) = \operatorname{argmax}_a \sum_{s'} \cancel{P(s'|s, a)} [\cancel{R(s, a, s')} + \gamma V^{\pi_{old}}(s')], \quad \forall s$$

Reinforcement Learning



Basic idea:

- All learning is based on observed **samples** of rewards and next states!
- Receive feedback in the form of **rewards**
- Must (learn to) act so as to **maximize expected sum of discounted rewards**

$$\pi(s) \quad \forall s$$

Reinforcement Learning

Still assume a Markov decision process (MDP):

- A set of states $s \in S$
- A set of actions (per state) A
- A model $T(s,a,s')$
- A reward function $R(s,a,s')$

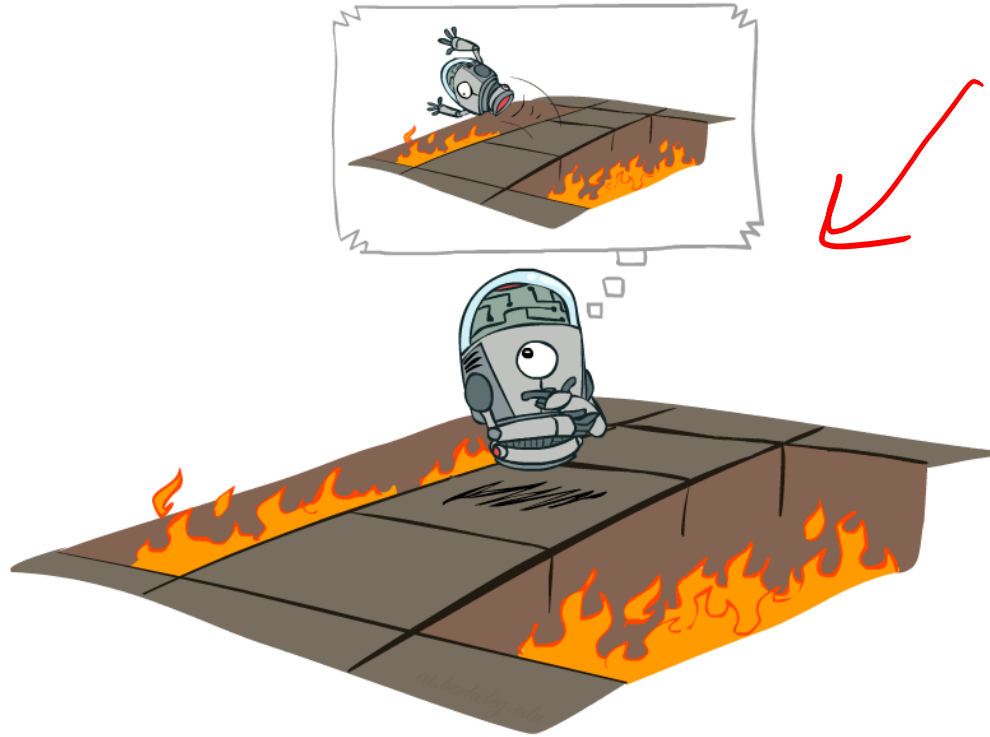
Still looking for a policy $\pi(s)$

New twist: don't know T or R

- I.e. we don't know which states are good or what the actions do
- Must actually try actions and states out to learn



Offline (MDPs) vs. Online (RL)



Offline Solution
(Known MDP)



Online Learning
(Unknown MDP)

Overview: MDPs and Reinforcement Learning

Known MDP: Offline Solution

Value Iteration / Policy Iteration

Unknown MDP: Online Learning

Model-Based

Estimate MDP $T(s,a,s')$ and $R(s,a,s')$
from samples of environment

Model-Free

Passive Reinforcement Learning

- Monte Carlo Policy Evaluation
- TD(0) Learning

Active Reinforcement Learning

- Q-Learning

Learn π

Fixed π

Why Not Use Policy Evaluation?

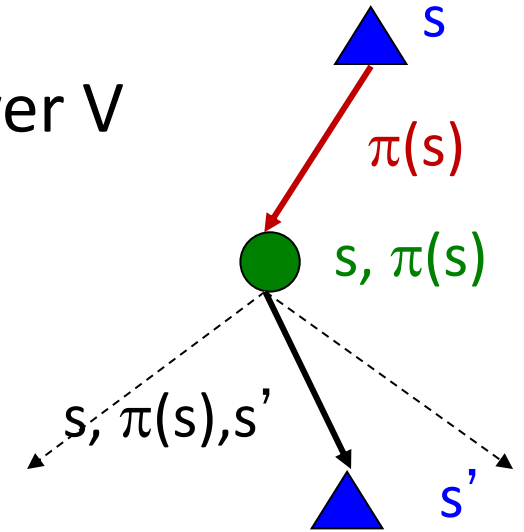
$\uparrow \rightarrow \sqrt{\pi}$

Simplified Bellman updates calculate V for a fixed policy:

- Each round, replace V with a one-step-look-ahead layer over V

$$V_0^\pi(s) = 0$$

$$V_{k+1}^\pi(s) \leftarrow \sum_{s'} T(s, \pi(s), s') [R(s, \pi(s), s') + \gamma V_k^\pi(s')]$$



- This approach fully exploited the connections between the states
- Unfortunately, we need T and R to do it!

Key question: How can we do this update to V without knowing T and R ?

- In other words, how to we take a weighted average without knowing the weights?

Online Learning (RL)

Passive Reinforcement Learning

Monte Carlo Policy Evaluation

Monte Carlo (Sample-Based) Learning

Collect a bunch of trajectories from episodes following π starting at a given state, s .

- Option 1: Collect N trajectories then evaluate π

$$V^{\pi}(s=B) = \frac{1}{N} \sum_{i=1}^N g^{(i)}$$

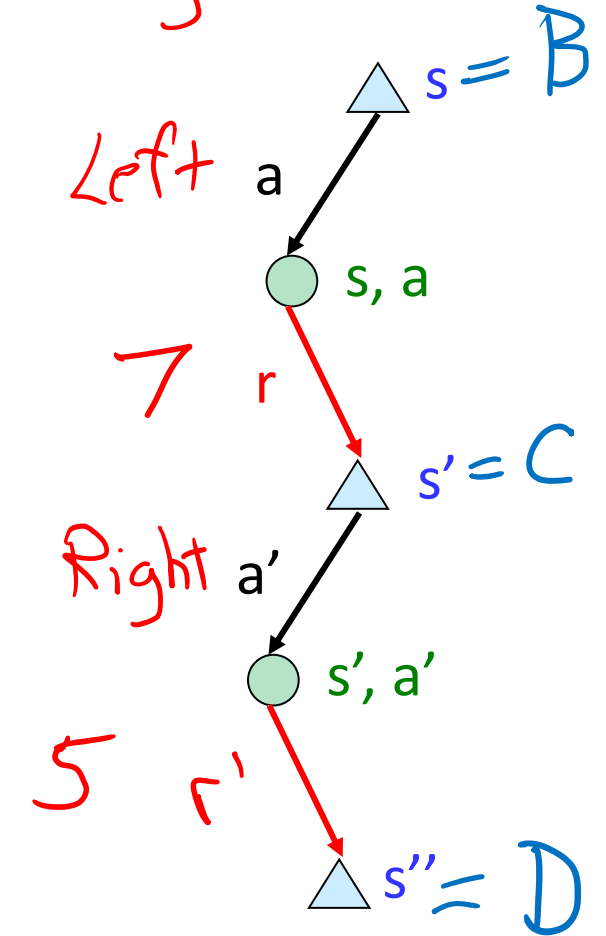
Trajectory $s, a, r, s', a', r', s'' \dots$

$B, L, 7, C, R, 5, D$

Return: sum of discounted rewards

$$G \text{ or } g = \sum_{k=0}^{\infty} \gamma^k r_{k+1} : \gamma^0 7 + \gamma^1 5 + \gamma^2 \underline{6}$$

Episode:
1 play of game $\rightarrow$ end



Monte Carlo (Sample-Based) Learning

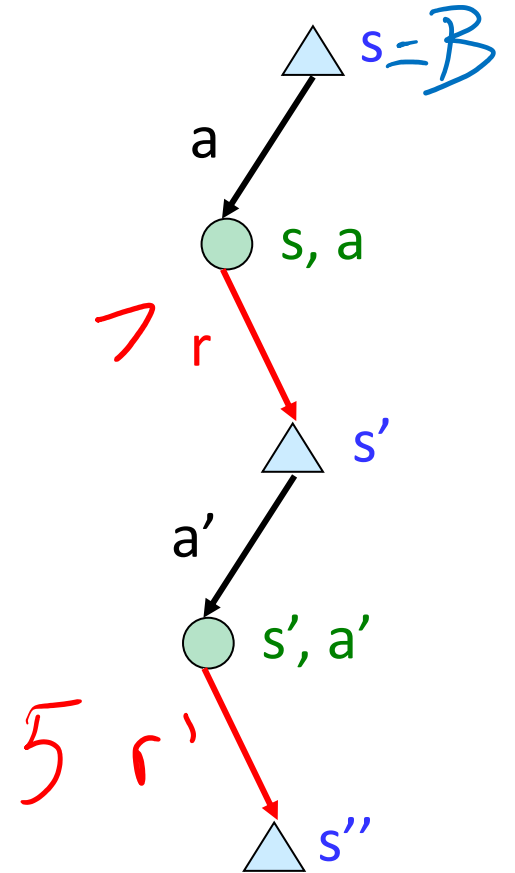
Collect a bunch of trajectories from episodes following π starting at a given state, s .

- Option 1: Collect N trajectories then evaluate π
- Option 2: Update $V^\pi(s)$ after each episode

Play episode $\rightarrow r \rightarrow g_k \quad 7 + \gamma 5$

Sample of $V(s)$: g_k

Update to $V(s)$: $V_{k+1}(s) = (1-\alpha)V_k(s) + \alpha \underbrace{g_{k+1}}_{\text{Sample}}$



Monte Carlo (Sample-Based) Learning

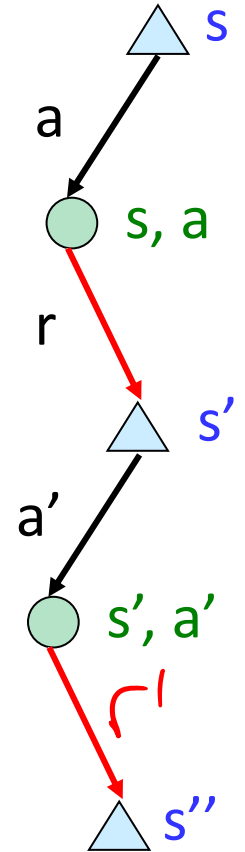
Collect a bunch of trajectories from episodes following π starting at a given state, s .

- Option 1: Collect N trajectories then evaluate π
- Option 2: Update $V^\pi(s)$ after each episode

Sample of $V(s)$: $g = \sum_{t=0}^{\infty} \gamma^t r_{t+1}$

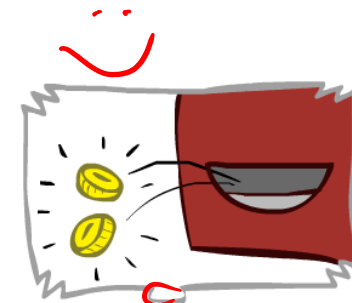
Update to $V(s)$: $V^\pi(s) \leftarrow (1 - \alpha)V^\pi(s) + (\alpha)\text{sample}$

Same update: $V^\pi(s) \leftarrow \underbrace{V^\pi(s)} + \alpha(\underbrace{\text{sample} - V^\pi(s)}_{\text{Temporal difference}})$

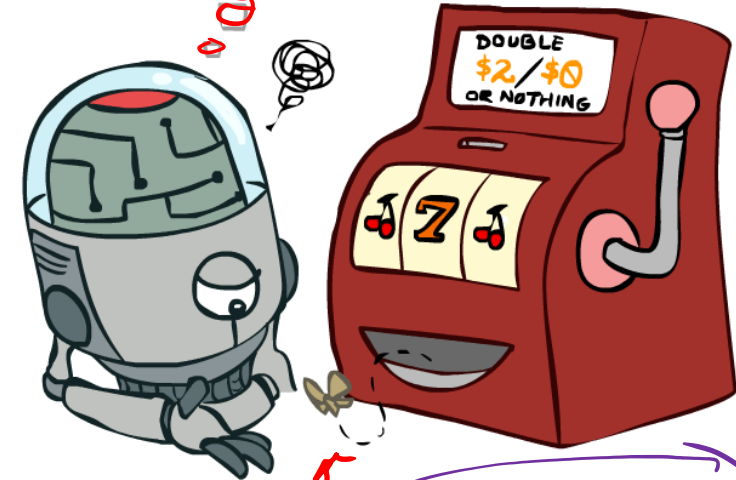


$$\mathcal{L} \frac{1}{2} (y - \hat{y})^2$$

$$\hat{y} = V_k(s)$$



$$V_{k+1}(s) = V_k(s) + \alpha (\underbrace{\text{sample}}_y - \underbrace{V_k(s)}_{\hat{y}})$$



sample ☹️

Online Learning

Passive Reinforcement Learning

Temporal Difference Learning

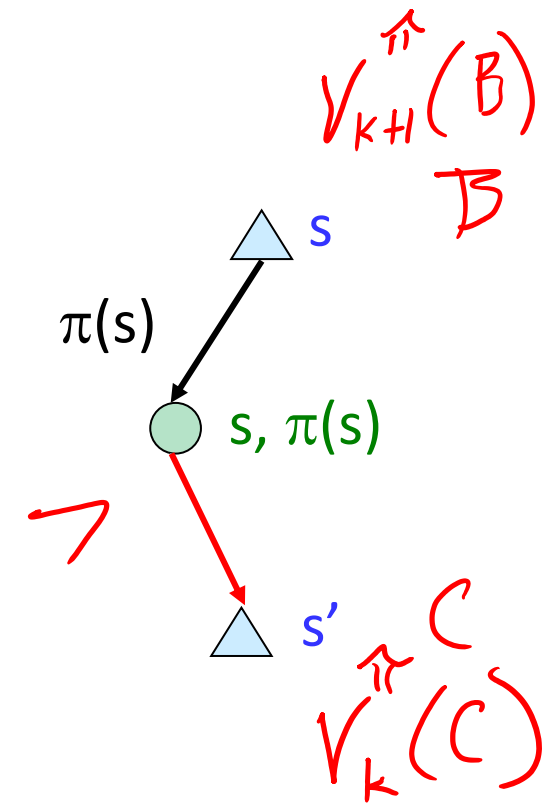
$TD(0)$

fixed $\pi \rightarrow V^\pi$

Temporal Difference Learning

Big idea: learn from every step!

- Update $V(s)$ each time we experience a transition (s, a, s', r)
- Likely outcomes s' will contribute updates more often
- Bootstrap: Leverage previous estimates of next state, $V_k^\pi(s')$ (rather than full trajectory)
- Policy still fixed, still doing evaluation!



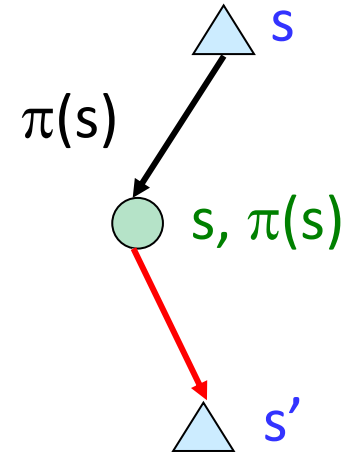
Sample of $V(s)$: sample = $r + \gamma V_k^\pi(s')$

Update to $V(s)$: $V_{k+1}^{\pi}(s) = V_k^{\pi}(s) + \alpha (\text{sample} - V_k^{\pi}(s))$

Temporal Difference Learning

Big idea: learn from every experience!

- Update $V(s)$ each time we experience a transition (s, a, s', r)
- Likely outcomes s' will contribute updates more often
- Bootstrap: Leverage previous estimates of next state, $V_k^\pi(s')$ (rather than full trajectory)
- Policy still fixed, still doing evaluation!



Sample of $V(s)$: $sample = r + \gamma V_k^\pi(s')$

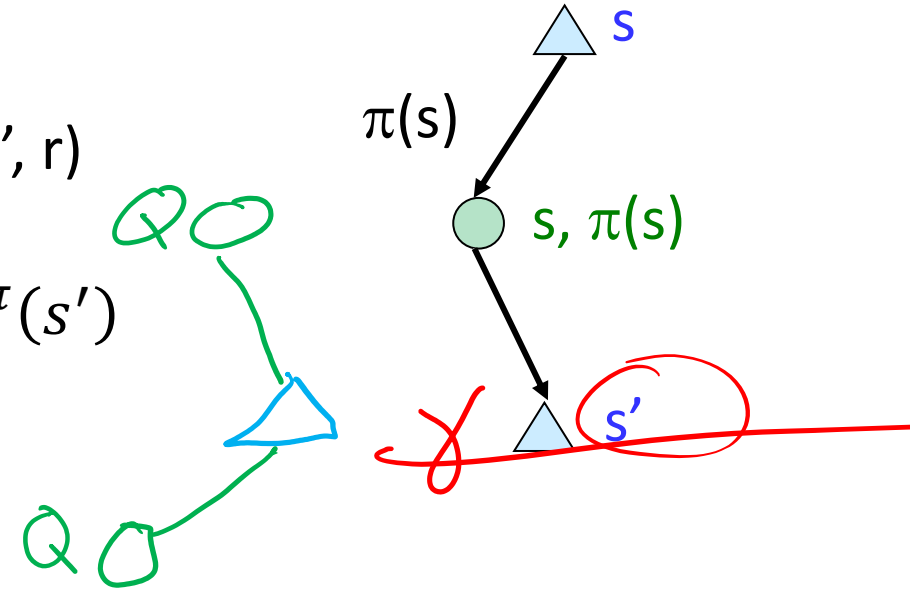
Update to $V(s)$: $V^\pi(s) \leftarrow (1 - \alpha)V^\pi(s) + (\alpha)sample$

Temporal Difference Learning

TD(0)

Big idea: learn from every experience!

- Update $V(s)$ each time we experience a transition (s, a, s', r)
- Likely outcomes s' will contribute updates more often
- **Bootstrap:** Leverage previous estimates of next state, $V_k^\pi(s')$ (rather than full trajectory)
- Policy still fixed, still doing evaluation!



Sample of $V(s)$: $sample = r + \gamma V_k^\pi(s')$

Update to $V(s)$: $V^\pi(s) \leftarrow (1 - \alpha)V^\pi(s) + (\alpha)sample$

Same update: $V^\pi(s) \leftarrow V^\pi(s) + \alpha(sample - V^\pi(s))$

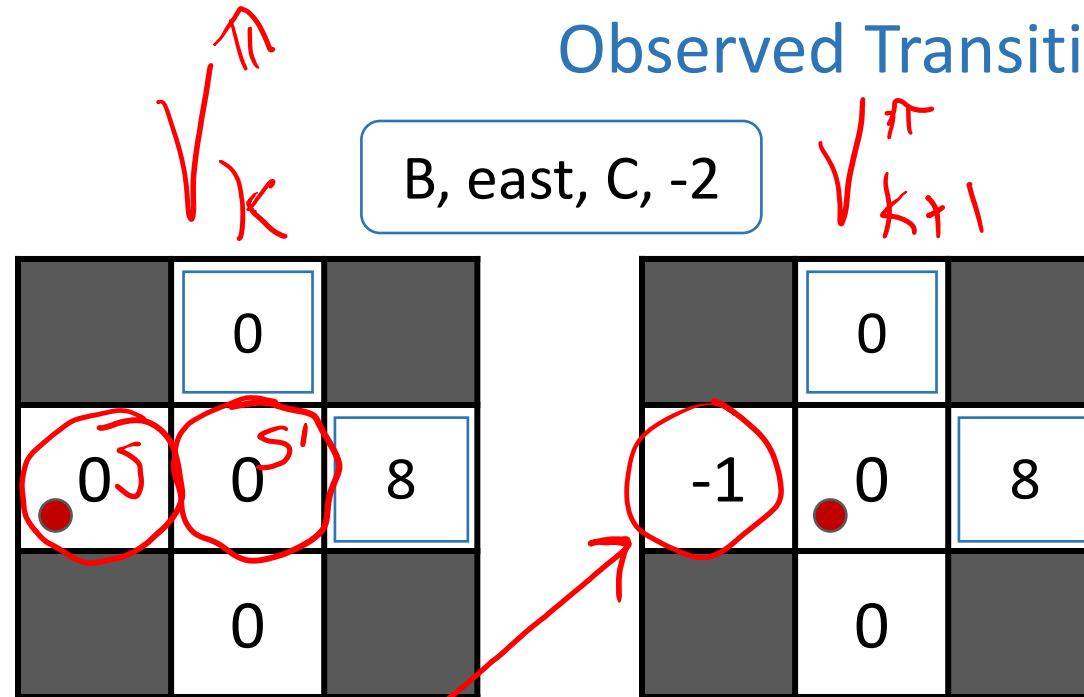
Example: Temporal Difference Learning

States

	A	
B	C	D
	E	

Assume: $\gamma = 1$,
 $\alpha = 1/2$

Observed Transitions



$$\text{sample} = R(s, \pi(s), s') + \gamma V^\pi(s')$$

$$V(\cdot) \leftarrow 0 + \frac{1}{2}(-2 + \gamma V(\cdot))$$

$$\underline{V^\pi(s)} \leftarrow \underline{V^\pi(s)} + \alpha(\text{sample} - \underline{V^\pi(s)})$$

Example: Temporal Difference Learning

States

	A	
B	C	D
	E	

Assume: $\gamma = 1$,
 $\alpha = 1/2$

Observed Transitions

B, east, C, -2

C, east, D, -2

	0	
0	0	8
	0	

	0	
-1	0	8
	0	

	0	
-1	3	8
	0	

$$\text{sample} = R(s, \pi(s), s') + \gamma V^\pi(s')$$

$$V(C) \leftarrow 0 + \frac{1}{2}(6 - 0) \quad -2 + \gamma V(D)$$

$$\underline{V^\pi(s)} \leftarrow \underline{V^\pi(s)} + \alpha(\underline{\text{sample}} - \underline{V^\pi(s)})$$

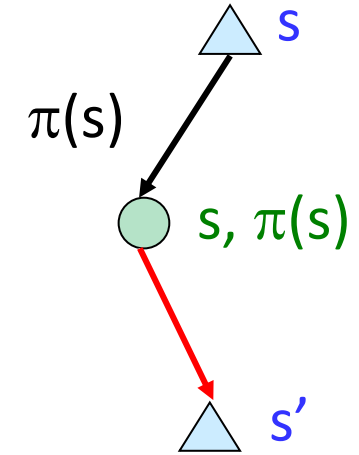
Temporal Difference Learning

Big idea: learn from every experience!

- Update $V(s)$ each time we experience a transition (s, a, s', r)
- Likely outcomes s' will contribute updates more often

Temporal difference learning of values

- Policy still fixed, still doing evaluation!
- Move values toward value of whatever successor occurs: running average



Sample of $V(s)$: $sample = r + \gamma V^\pi(s')$

Update to $V(s)$: $V^\pi(s) \leftarrow (1 - \alpha) V^\pi(s) + (\alpha) sample$

Same update: $V^\pi(s) \leftarrow V^\pi(s) + \alpha [sample - V^\pi(s)]$

Same update: $V^\pi(s) \leftarrow V^\pi(s) - \alpha \nabla Error$

$$Error = \frac{1}{2} \left(\underbrace{\gamma^{(i)} - h(x^{(i)})}_{sample - V^\pi(s)} \right)^2$$

ML detour: Quick Calculus Quiz

$$f(x) = \frac{1}{2}(y - x)^2$$

What is $\frac{df}{dx}$?

ML detour: Gradient Descent

Goal: find x that minimizes $f(x)$

1. Start with initial guess, x_0
2. Update x by taking a step in the direction that $f(x)$ is changing fastest (in the negative direction) with respect to x :

$$x \leftarrow x - \alpha \nabla_x f, \text{ where } \alpha \text{ is the step size or learning rate}$$

3. Repeat until convergence

$$f(x) = \frac{1}{2} (y - x)^2$$

$$\frac{df}{dx} = -(y - x)$$

TD goal: find value(s), V , that minimizes difference between sample(s) and V

$$V \leftarrow V - \alpha \nabla_V \text{Error}$$

$$\text{Error}(V) = \frac{1}{2} (\text{sample} - V)^2$$

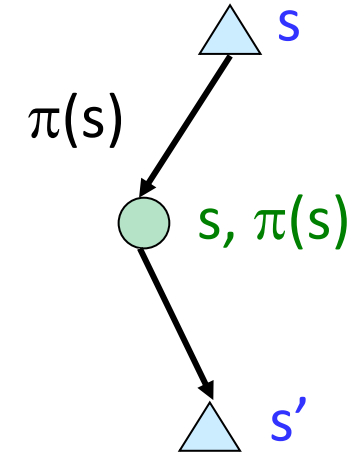
Temporal Difference Learning

Big idea: learn from every experience!

- Update $V(s)$ each time we experience a transition (s, a, s', r)
- Likely outcomes s' will contribute updates more often

Temporal difference learning of values

- Policy still fixed, still doing evaluation!
- Move values toward value of whatever successor occurs: running average



Sample of $V(s)$: $sample = r + \gamma V^\pi(s')$

Update to $V(s)$: $V^\pi(s) \leftarrow (1 - \alpha) V^\pi(s) + (\alpha) sample$

Same update: $V^\pi(s) \leftarrow V^\pi(s) + \alpha [sample - V^\pi(s)]$

Same update: $V^\pi(s) \leftarrow V^\pi(s) - \alpha \nabla Error$

$$Error = \frac{1}{2} (sample - V^\pi(s))^2$$

Poll 1

TD update:

$$V^\pi(s) = V^\pi(s) + \alpha [r + \gamma V^\pi(s') - V^\pi(s)]$$

Which converts TD values into a policy?

A) Value iteration:

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

B) Q-iteration:

$$Q_{k+1}(s, a) = \sum_{s'} P(s'|s, a) [R(s, a, s') + \gamma \max_{a'} Q_k(s', a')], \quad \forall s, a$$

C) Policy extraction:

$$\pi_V(s) = \operatorname{argmax}_a \sum_{s'} P(s'|s, a) [R(s, a, s') + \gamma V(s')], \quad \forall s$$

D) Policy evaluation:

$$V_{k+1}^\pi(s) = \sum_{s'} P(s'|s, \pi(s)) [R(s, \pi(s), s') + \gamma V_k^\pi(s')], \quad \forall s$$

E) Policy improvement:

$$\pi_{new}(s) = \operatorname{argmax}_a \sum_{s'} P(s'|s, a) [R(s, a, s') + \gamma V^{\pi_{old}}(s')], \quad \forall s$$

F) None of the above

Problems with TD Value Learning

TD value learning is a model-free way to do policy evaluation, mimicking Bellman updates with running sample averages

However, if we want to turn values into a (new) policy, we're sunk:

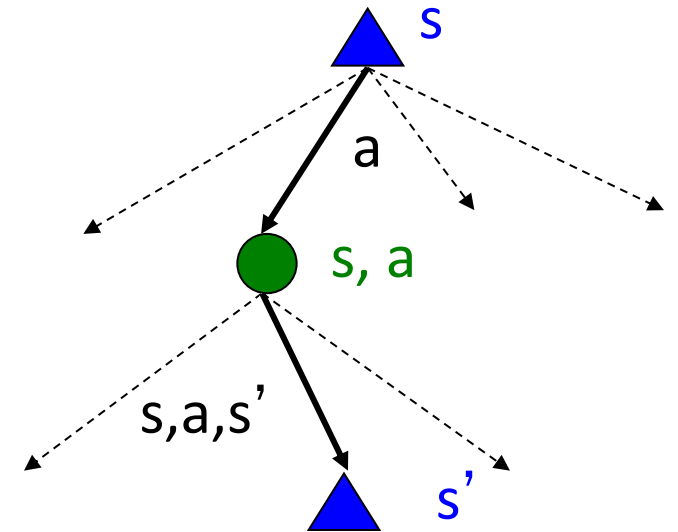
$$\pi(s) = \arg \max_a Q(s, a)$$



$$Q(s, a) = \sum_{s'} T(s, a, s') [R(s, a, s') + \gamma V(s')]$$

Idea: learn Q-values, not values

Makes action selection model-free too!



MDP/RL Notation

Standard expectimax:

$$V(s) = \max_a \sum_{s'} P(s'|s, a) V(s')$$

Bellman equations:

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

Value iteration:

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

Q-iteration:

$$Q_{k+1}(s, a) = \sum_{s'} P(s'|s, a) [R(s, a, s') + \gamma \max_{a'} Q_k(s', a')], \quad \forall s, a$$

Policy extraction:

$$\pi_V(s) = \operatorname{argmax}_a \sum_{s'} P(s'|s, a) [R(s, a, s') + \gamma V(s')], \quad \forall s$$

Policy evaluation:

$$V_{k+1}^\pi(s) = \sum_{s'} P(s'|s, \pi(s)) [R(s, \pi(s), s') + \gamma V_k^\pi(s')], \quad \forall s$$

Policy improvement:

$$\pi_{new}(s) = \operatorname{argmax}_a \sum_{s'} P(s'|s, a) [R(s, a, s') + \gamma V^{\pi_{old}}(s')], \quad \forall s$$

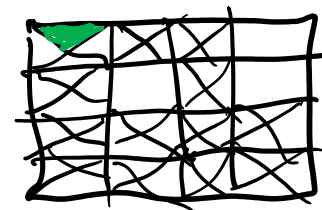
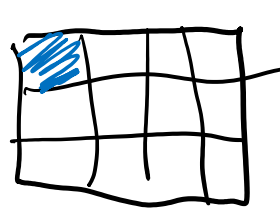
Value (TD) learning:

$$V^\pi(s) = \underline{V^\pi(s)} + \alpha [r + \gamma \underline{V^\pi(s')} - \underline{V^\pi(s)}]$$

Q-learning:

$$Q(s, a) = \underline{Q(s, a)} + \alpha [r + \gamma \max_{a'} \underline{Q(s', a')} - \underline{Q(s, a)}]$$

RL Summary

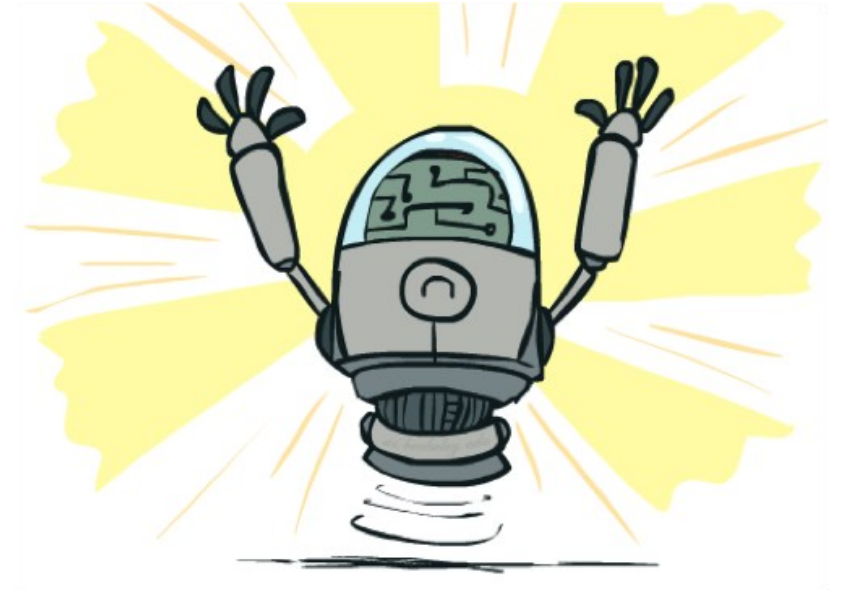


Name	Monte Carlo	TD(0)	Q-learning	Approx. Q-learning
Model	$V(s)$ table	$V(s)$ table	<u>$Q(s, a)$ table</u>	θ for $Q_\theta(s, a)$
Action selection	$a = \pi(s)$	$a = \pi(s)$	$a' = \operatorname{argmax} Q(s', a')$ (and/or exploration)	$a' = \operatorname{argmax} Q(s', a')$ (and/or exploration)
Sampling technique	<u>MC</u>	Bootstrap	Bootstrap	Bootstrap
Loss	$\frac{1}{2} (y - \hat{y})^2$	$\frac{1}{2} (y - \hat{y})^2$	$\frac{1}{2} (y - \hat{y})^2$	$\frac{1}{2} (y - \hat{y})^2$
Prediction, $\hat{y}$	$V^\pi(s)$	$V^\pi(s)$	$Q(s, a)$	$Q_\theta(s, a)$
Sample, y	$y = \sum_{t=0}^{\infty} \gamma^t r_{t+1}$	$y = r + \gamma V^\pi(s')$	$y = r + \gamma \max_{a'} Q(s', a')$	

RL Summary

Name	Monte Carlo	TD(0)	Q-learning	Approx. Q-learning
Model	$V(s)$ table	$V(s)$ table	$Q(s, a)$ table	θ for $Q_\theta(s, a)$
Action selection	$a = \pi(s)$	$a = \pi(s)$	$a' = \operatorname{argmax} Q(s', a')$ (and/or exploration)	$a' = \operatorname{argmax} Q(s', a')$ (and/or exploration)
Sampling technique	MC	Bootstrap	Bootstrap	Bootstrap
Loss	$\frac{1}{2}(y - \hat{y})^2$	$\frac{1}{2}(y - \hat{y})^2$	$\frac{1}{2}(y - \hat{y})^2$	$\frac{1}{2}(y - \hat{y})^2$
Prediction, $\hat{y}$	$V^\pi(s)$	$V^\pi(s)$	$Q(s, a)$	$Q_\theta(s, a)$
Sample, y	$ \begin{aligned} &r + \\ &\gamma r' + \\ &\gamma^2 r'' + \\ &\gamma^3 r''' + \\ &\gamma^4 r'''' + \\ &\dots \end{aligned} $	$ \begin{aligned} &r + \\ &\gamma V^\pi(s') \end{aligned} $	$ \begin{aligned} &r + \\ &\gamma \max_{a'} Q(s', a') \end{aligned} $	$ \begin{aligned} &r + \\ &\gamma \max_{a'} Q_\theta(s', a') \end{aligned} $

Online Learning
Model-free Learning
Active Reinforcement Learning
Q-Learning



Q-Learning

We'd like to do Q-value updates to each Q-state:

$$Q_{k+1}(s, a) \leftarrow \sum_{s'} T(s, a, s') \left[R(s, a, s') + \gamma \max_{a'} Q_k(s', a') \right]$$

- But can't compute this update without knowing T, R

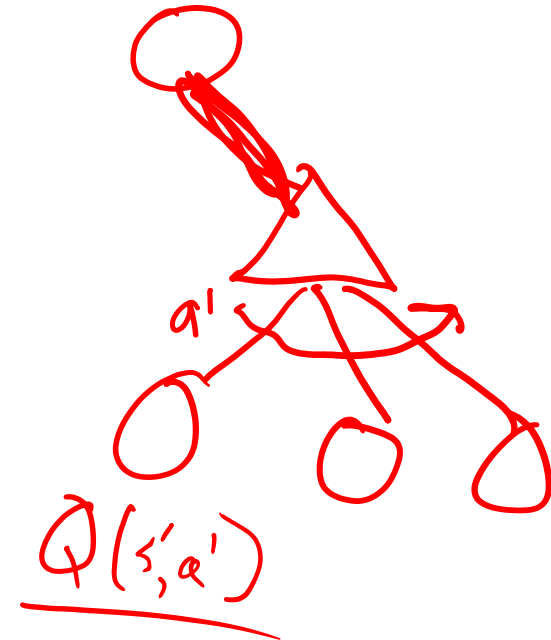
Instead, compute average as we go

- Receive a sample transition (s,a,r,s')
- The sample value is then:

$$\text{sample} = r + \gamma \max_{a'} Q(s', a')$$

- But we want to average over ~~results from (s,a)~~ (Why?)
- So keep a running average

$$Q(s, a) \leftarrow (1 - \alpha)Q(s, a) + (\alpha) \left[r + \gamma \max_{a'} Q(s', a') \right]$$



Q-Learning Properties

Amazing result: Q-learning converges to optimal policy -- even if you're acting suboptimally!

This is called off-policy learning

Caveats:

- You have to explore enough
- You have to eventually make the learning rate small enough
- ... but not decrease it too quickly
- Basically, in the limit, it doesn't matter how you select actions (!)



Overview: MDPs and Reinforcement Learning

Known MDP: Offline Solution

Value Iteration / Policy Iteration

Unknown MDP: Online Learning

Model-Based

Estimate MDP $T(s,a,s')$ and $R(s,a,s')$
from samples of environment

Model-Free

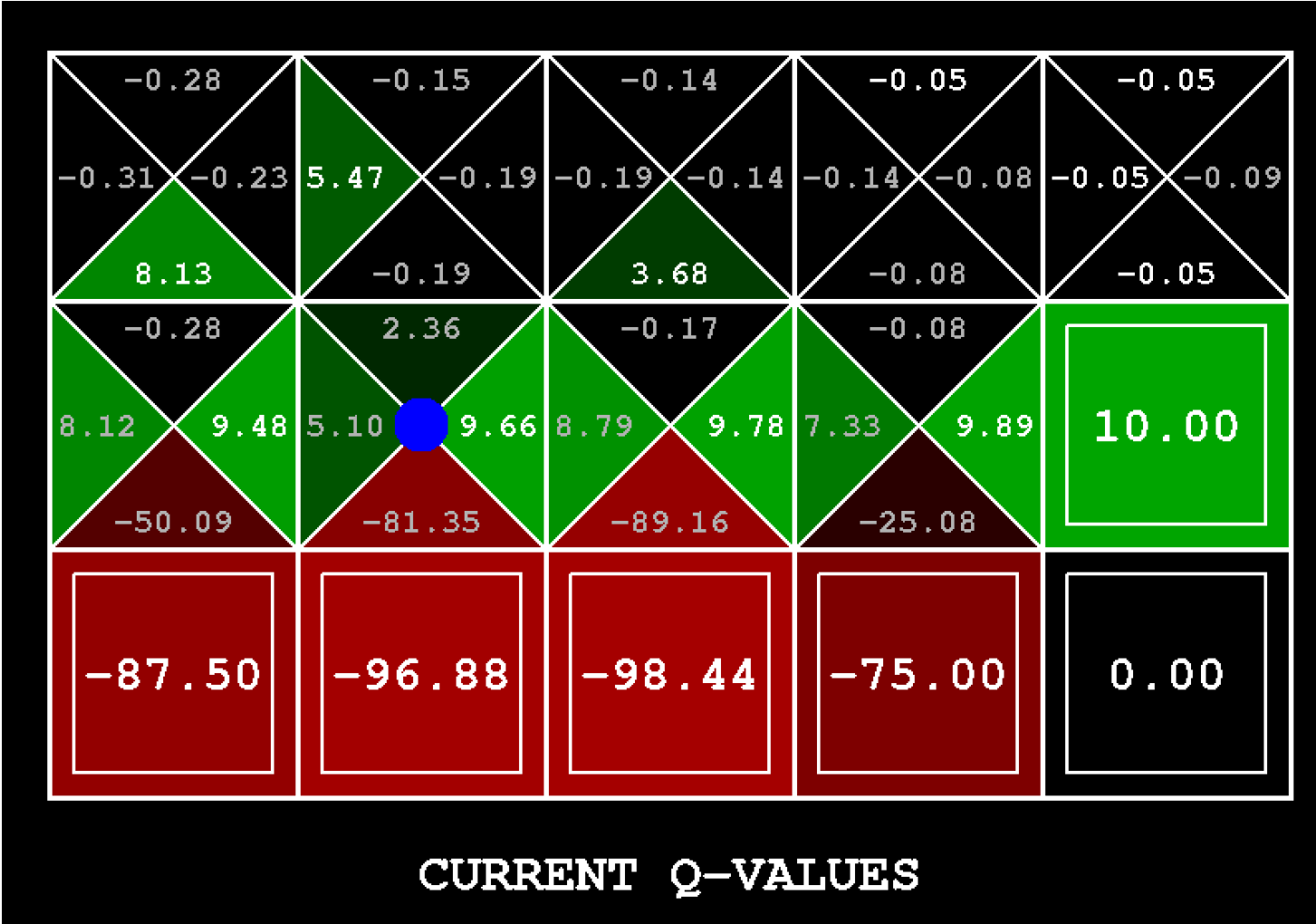
Passive Reinforcement Learning

- Monte Carlo Policy Evaluation
- TD Learning

Active Reinforcement Learning

- Q-Learning

Demo Q-Learning Auto Cliff Grid



Double Bandits

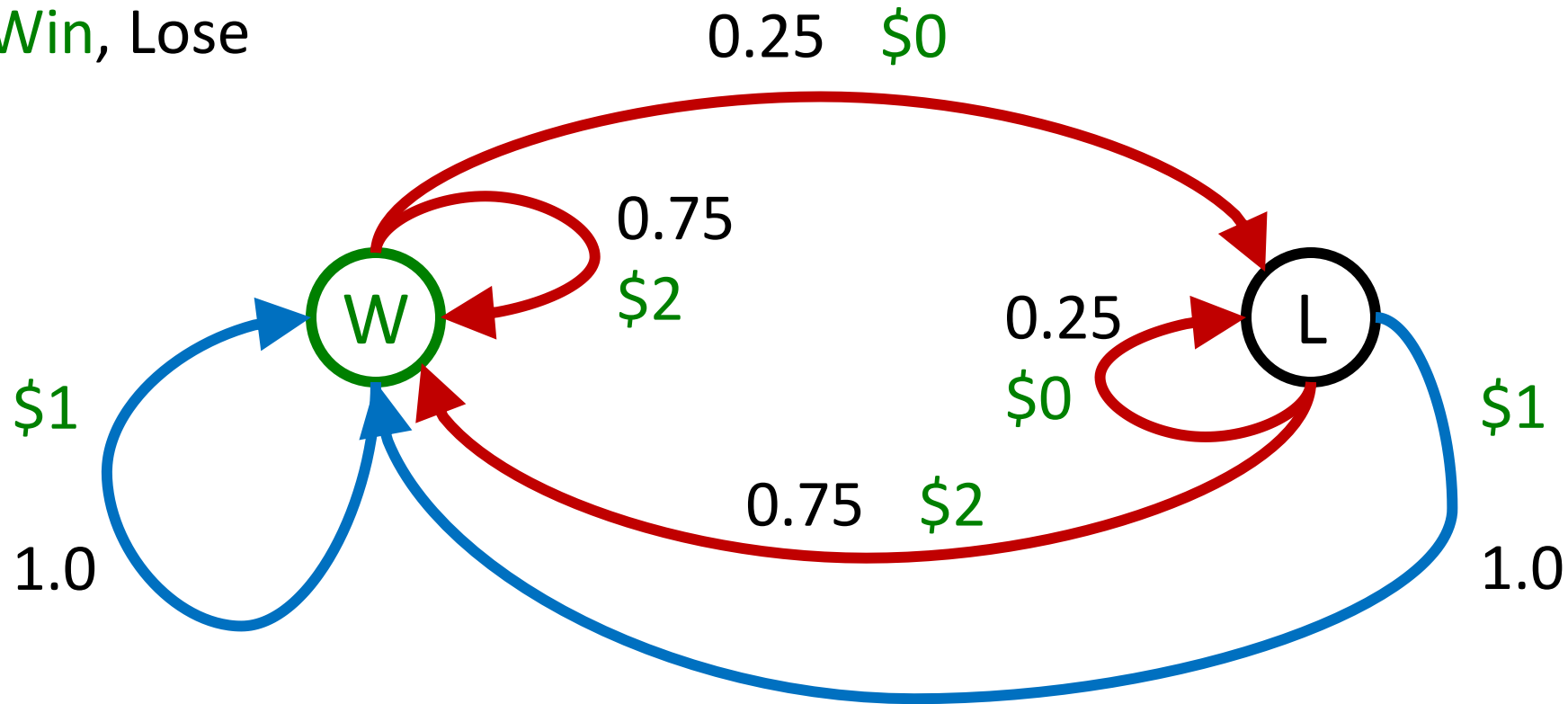


$$S = 0$$

Double-Bandit MDP

Actions: *Blue*, *Red*

States: *Win*, Lose



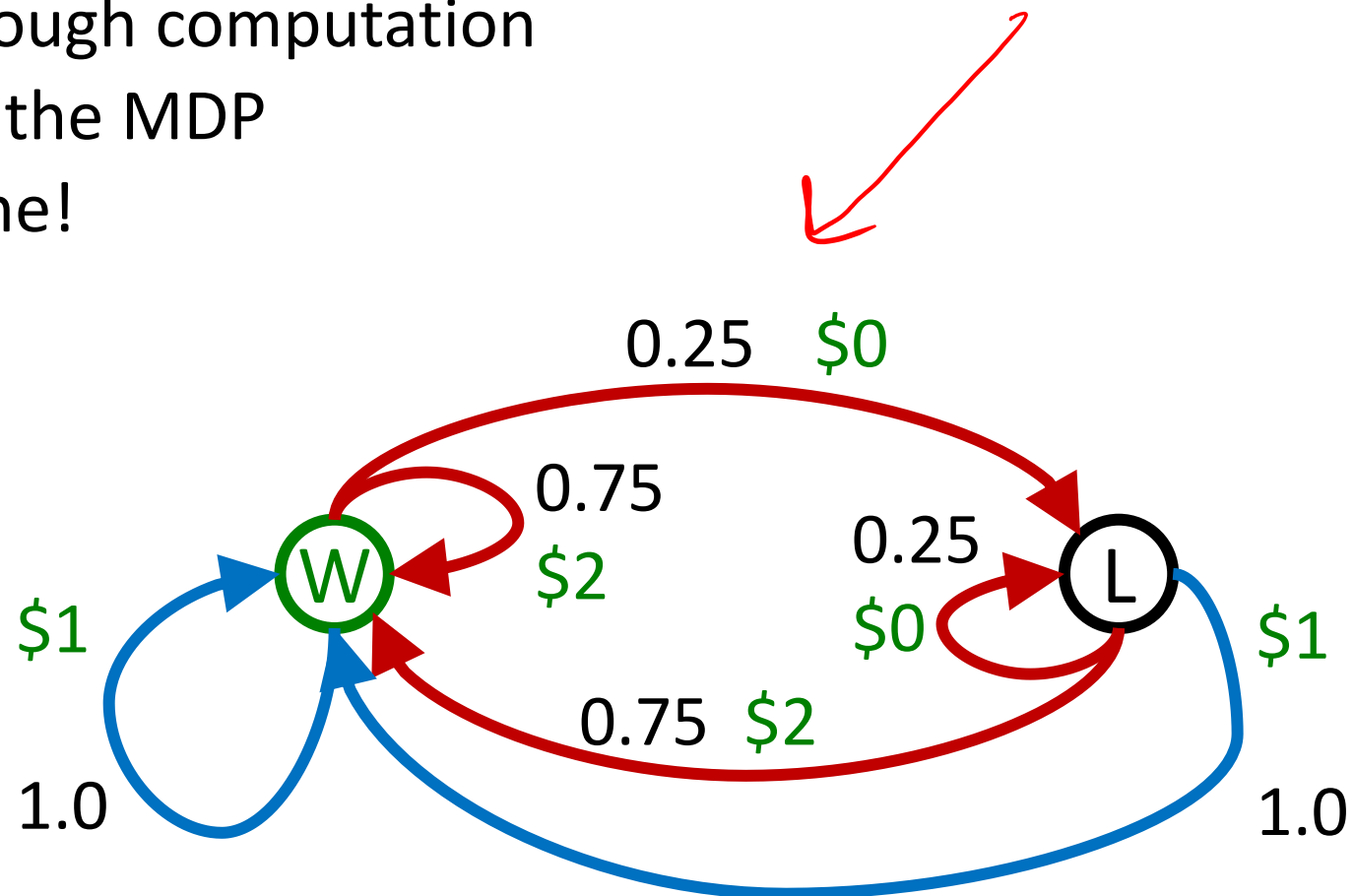
Offline Planning

No discount
100 time steps

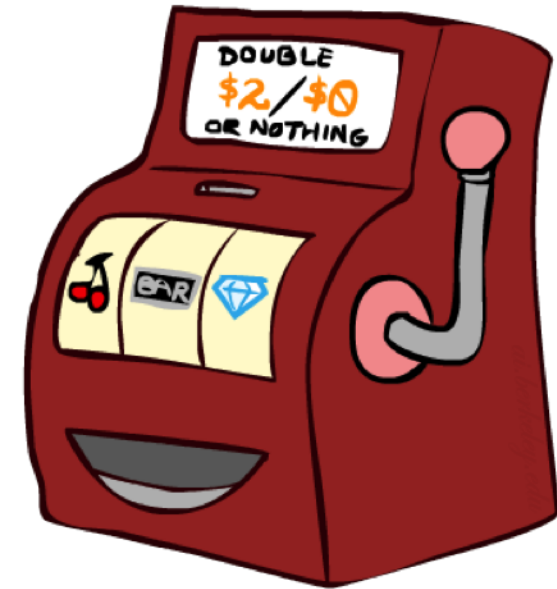
Solving MDPs is offline planning

- You determine all quantities through computation
- You need to know the details of the MDP
- You do not actually play the game!

Q	Value
Q Play Red	150
Q Play Blue	100



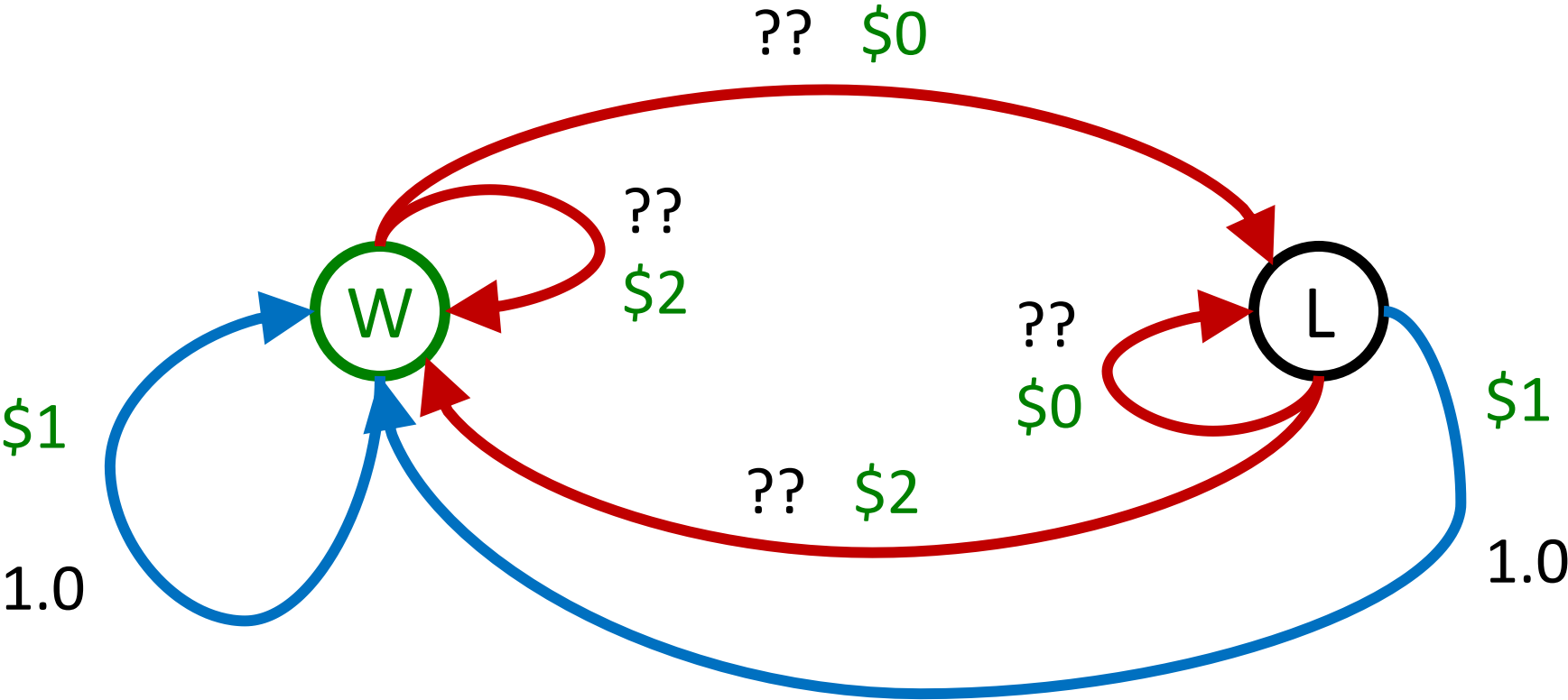
Let's Play!



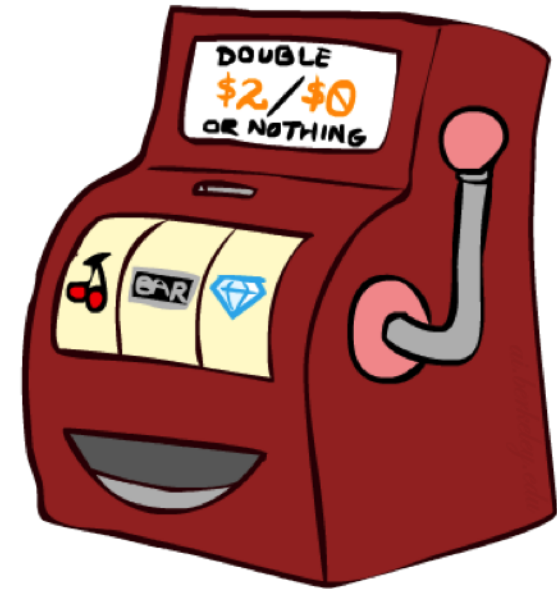
\$2 \$2 \$0 \$2 \$2
\$2 \$2 \$0 \$0 \$0

Online Planning

Rules changed! Red's win chance is different.



Let's Play!



\$0 \$0 \$0 \$2 \$0
\$2 \$0 \$0 \$0 \$0

What Just Happened?



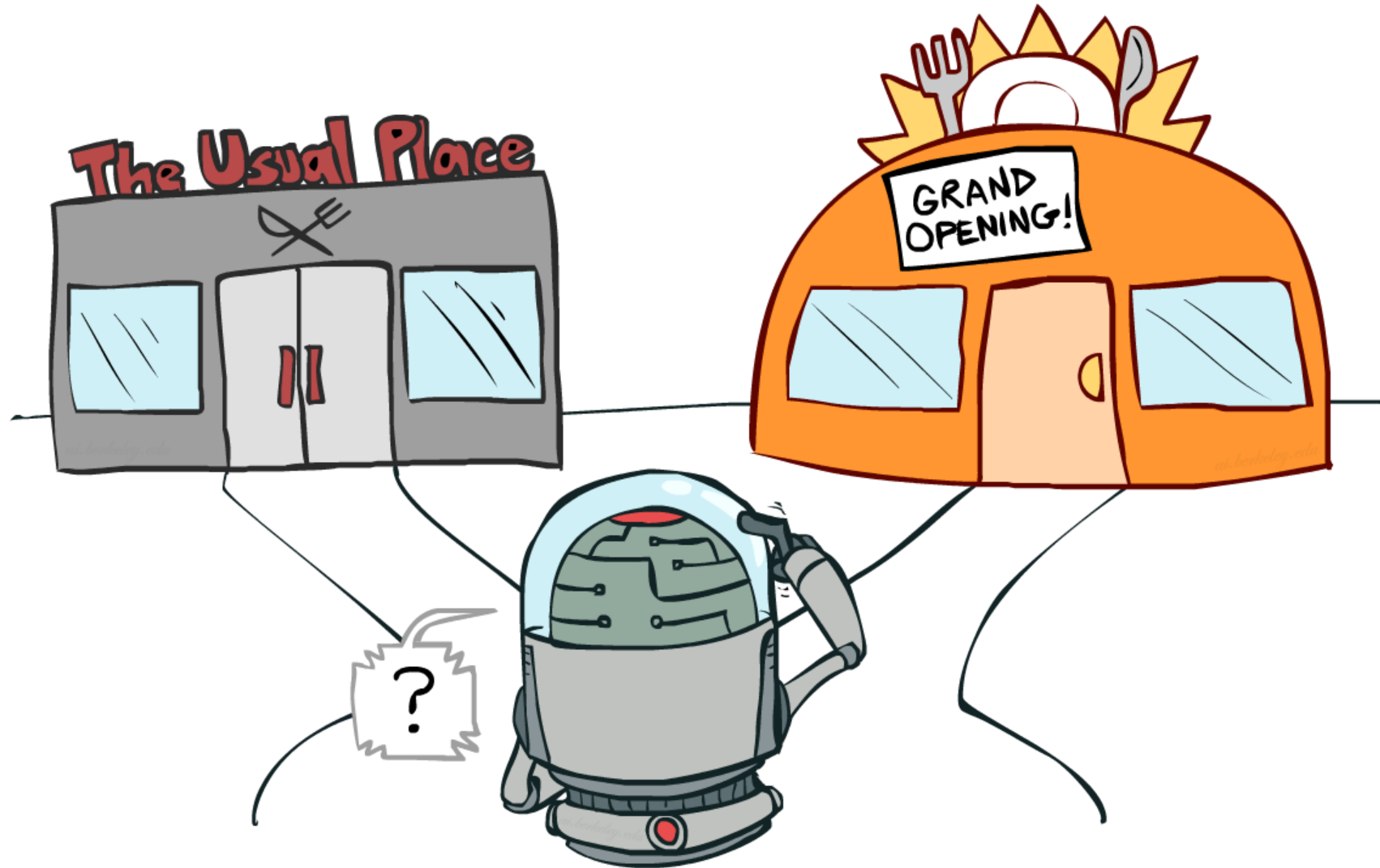
That wasn't planning, it was learning!

- Specifically, reinforcement learning
- There was an MDP, but you couldn't solve it with just computation
- You needed to actually act to figure it out

Important ideas in reinforcement learning that came up

- **Exploration**: you have to try unknown actions to get information
- **Exploitation**: eventually, you have to use what you know
- ➔ **Regret**: even if you learn intelligently, you make mistakes
- ➔ **Sampling**: because of chance, you have to try things repeatedly
- **Difficulty**: learning can be much harder than solving a known MDP

Exploration vs. Exploitation



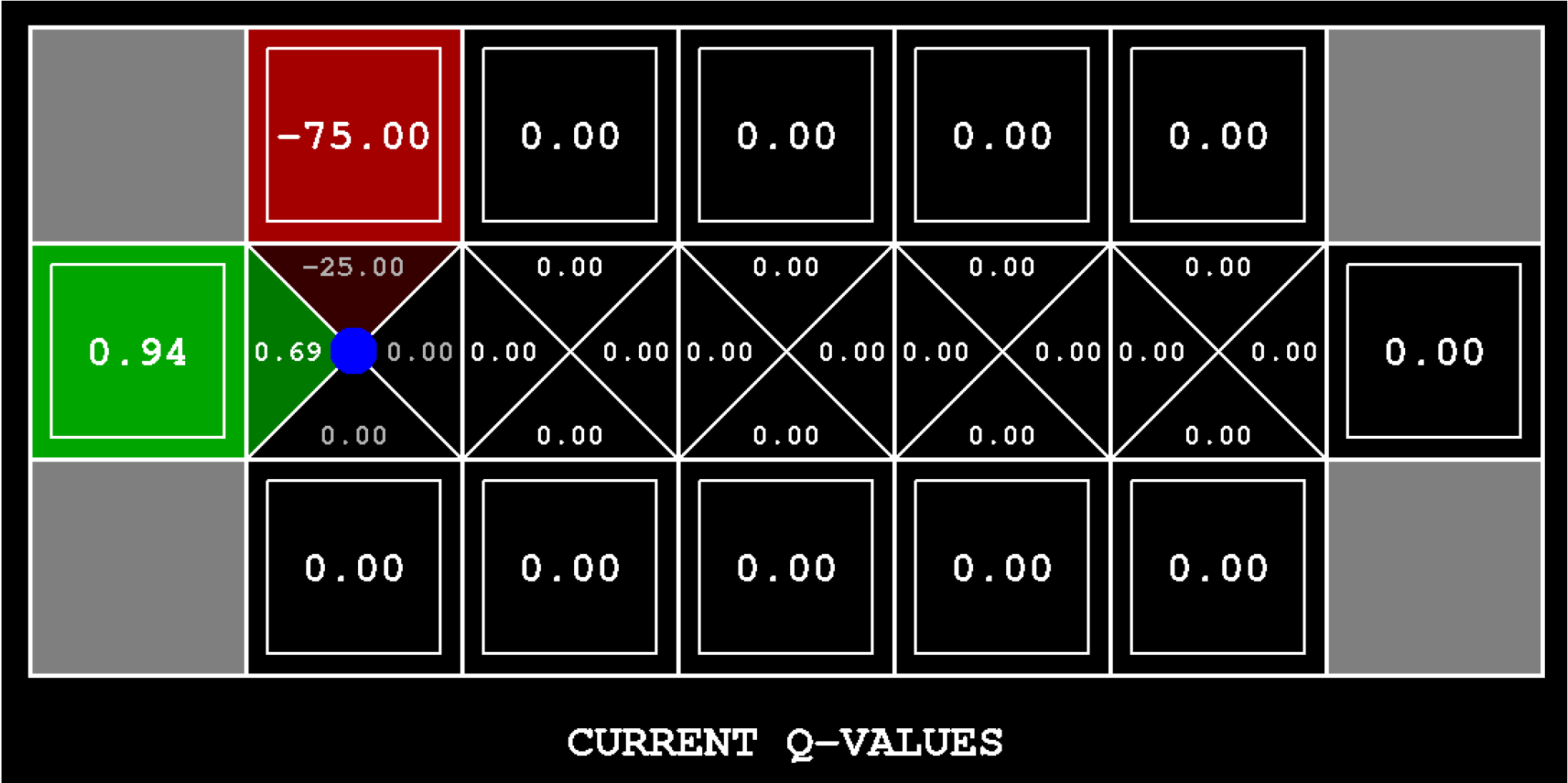
How to Explore?

Several schemes for forcing exploration

- Simplest: random actions (ϵ -greedy)
 - Every time step, flip a coin 0.2
 - With (small) probability ϵ , act randomly
 - With (large) probability $1-\epsilon$, act on current policy 0.8
- Problems with random actions?
 - You do eventually explore the space, but keep thrashing around once learning is done
 - One solution: lower ϵ over time
 - Another solution: exploration functions



Demo Q-learning – Manual Exploration – Bridge Grid



Exploration Functions

When to explore?

- Random actions: explore a fixed amount
- Better idea: explore areas whose badness is not (yet) established, eventually stop exploring

Exploration function

- Takes a value estimate u and a visit count n , and returns an optimistic utility, e.g.

$$f(u, n) = u + k/n$$

Regular Q-Update: $Q(s, a) \leftarrow_{\alpha} R(s, a, s') + \gamma \max_{a'} Q(s', a')$

Modified Q-Update: $Q(s, a) \leftarrow_{\alpha} R(s, a, s') + \gamma \max_{a'} \underbrace{f(Q(s', a'), N(s', a'))}_{\text{exploration bonus}}$

- Note: this propagates the “bonus” back to states that lead to unknown states as well!



Exploration Functions

When to explore?

- Random actions: explore a fixed amount
- Better idea: explore areas whose badness is not (yet) established, eventually stop exploring

$$\begin{array}{l}
 \text{Handwritten calculations:} \\
 \text{Red: } k=100 \\
 \text{Blue: } \frac{Q(s,a)}{10} \frac{n(s,a)}{1} \Rightarrow f(u,n) = 10 + 100/1 = 110 \\
 \text{Purple: } \frac{Q}{100} \frac{n}{100} \Rightarrow 100 + 100/100 = 101
 \end{array}$$

Exploration function

- Takes a value estimate u and a visit count n , and returns an optimistic utility, e.g.

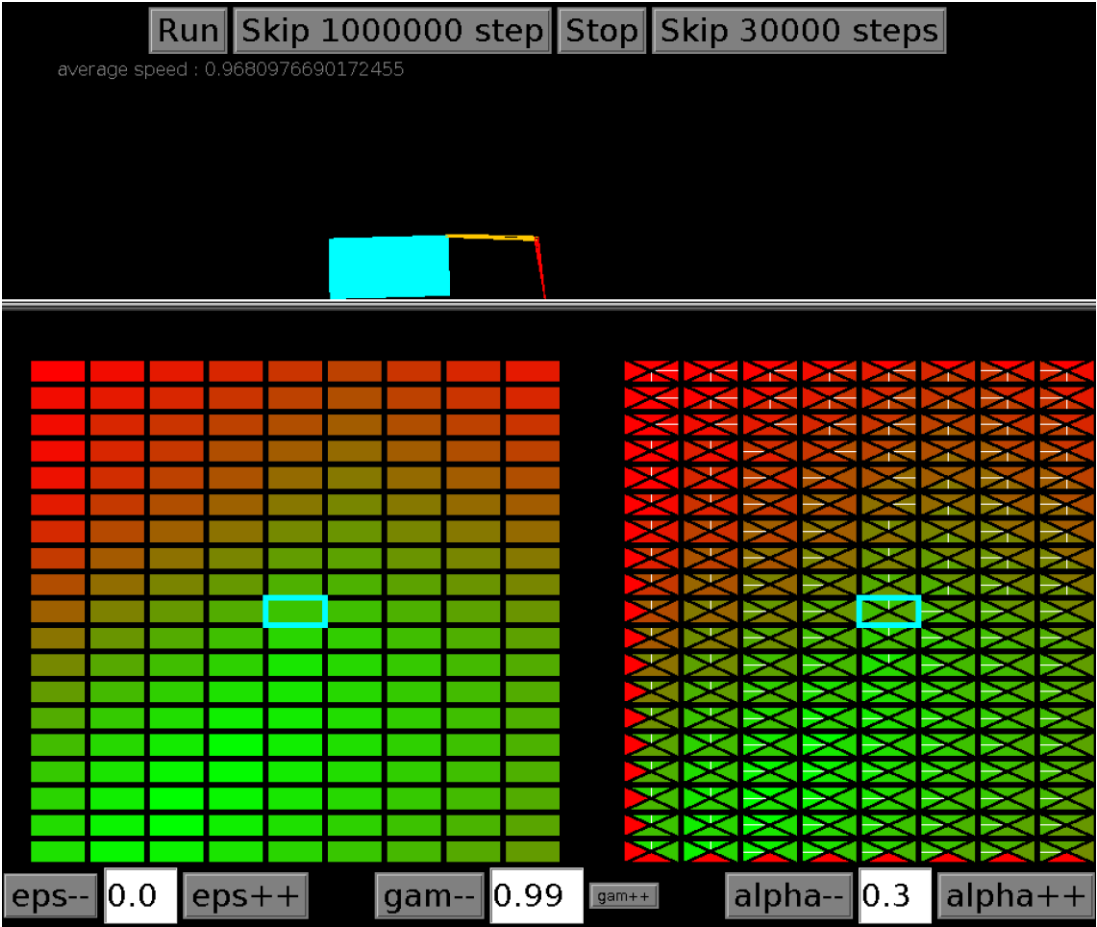
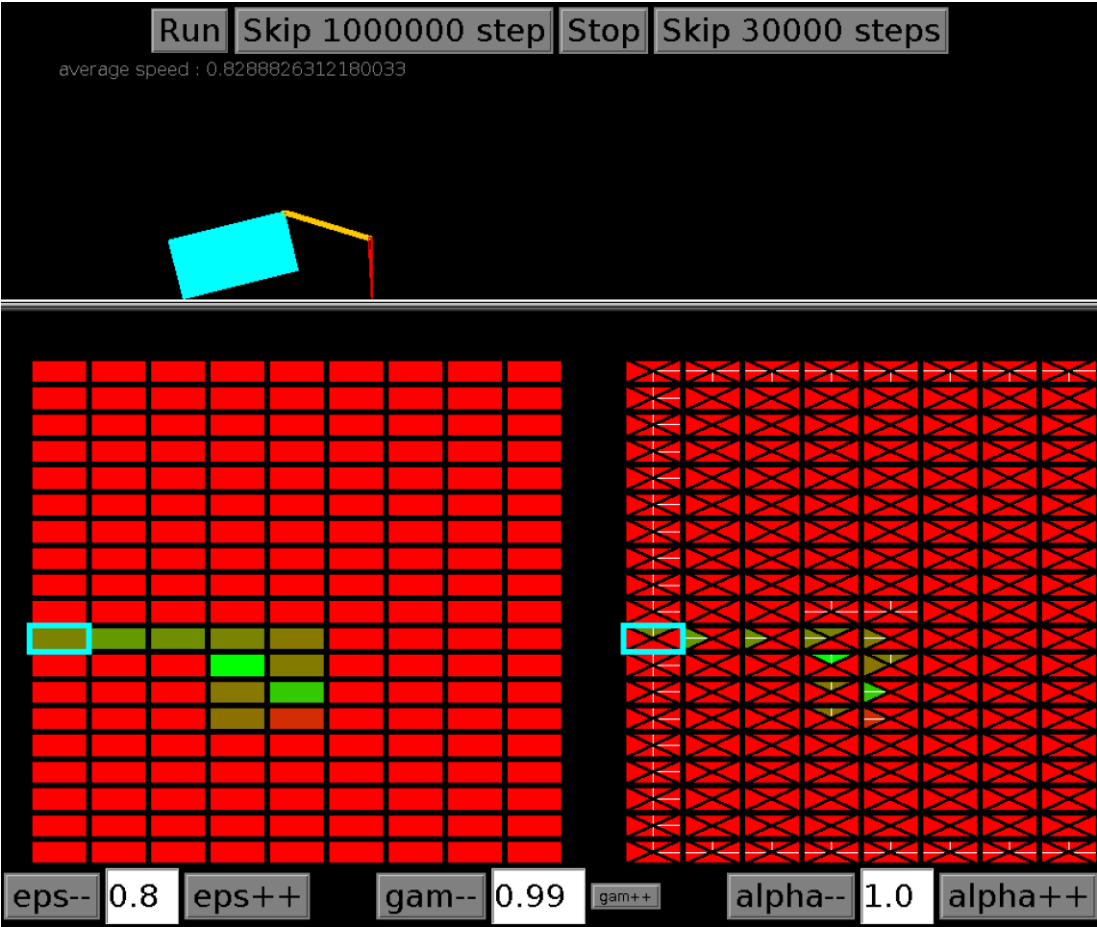
$$f(u, n) = \underline{u} + k/(n+1)$$

Regular Q-Update: $Q(s, a) \leftarrow_{\alpha} R(s, a, s') + \gamma \max_{a'} Q(s', a')$

Modified Q-Update: $Q(s, a) \leftarrow_{\alpha} R(s, a, s') + \gamma \max_{a'} \underline{f(Q(s', a'), N(s', a'))}$

- Note: this propagates the “bonus” back to states that lead to unknown states as well!

Demo Q-learning – Epsilon-Greedy – Crawler



Exploration Functions

When to explore?

- Random actions: explore a fixed amount
- Better idea: explore areas whose badness is not (yet) established, eventually stop exploring

Exploration function

- Takes a value estimate u and a visit count n , and returns an optimistic utility, e.g.

$$f(u, n) = u + k/(n + 1)$$

Regular Q-Update: $Q(s, a) = Q(s, a) + \alpha [r + \gamma \max_{a'} Q(s', a') - Q(s, a)]$

Modified Q-Update: $Q(s, a) = Q(s, a) + \alpha [r + \gamma \max_{a'} f(Q(s', a'), N(s', a')) - Q(s, a)]$

- Note: this propagates the “bonus” back to states that lead to unknown states as well!



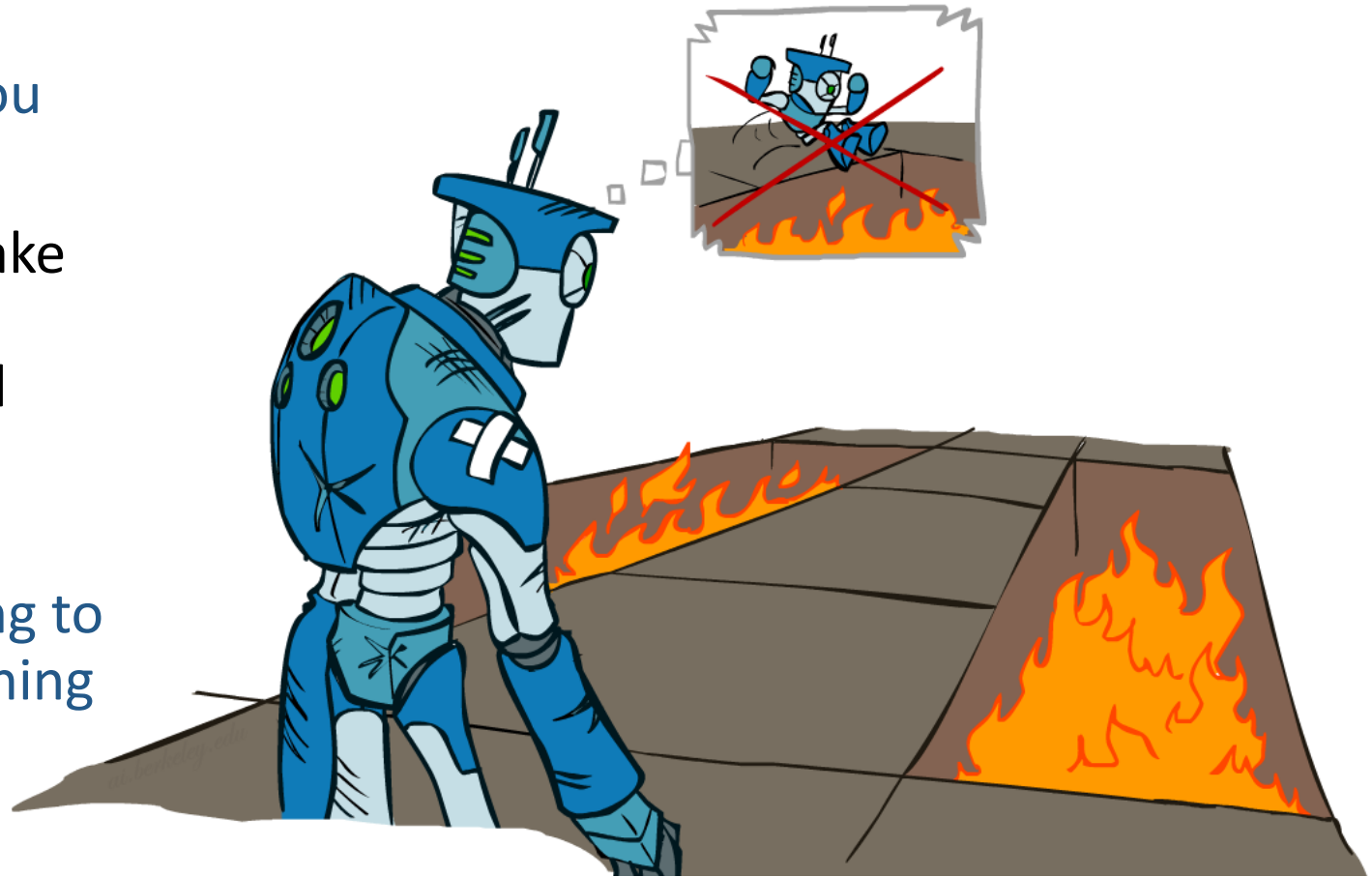
Regret

Even if you learn the optimal policy, you still make mistakes along the way!

Regret is a measure of your total mistake cost: the difference between your (expected) rewards, including youthful suboptimality, and optimal (expected) rewards

Minimizing regret goes beyond learning to be optimal – it requires optimally learning to be optimal

Example: random exploration and exploration functions both end up optimal, but random exploration has higher regret



RL Summary

Name	Monte Carlo	TD(0)	Q-learning	Approx. Q-learning
Model	$V(s)$ table	$V(s)$ table	$Q(s, a)$ table	θ for $Q_\theta(s, a)$
Action selection	$a = \pi(s)$	$a = \pi(s)$	$a' = \operatorname{argmax} Q(s', a')$ (and/or exploration)	$a' = \operatorname{argmax} Q(s', a')$ (and/or exploration)
Sampling technique	MC	Bootstrap	Bootstrap	Bootstrap
Loss	$\frac{1}{2}(\hat{y} - y)^2$	$\frac{1}{2}(\hat{y} - y)^2$	$\frac{1}{2}(\hat{y} - y)^2$	$\frac{1}{2}(\hat{y} - y)^2$
Prediction, $\hat{y}$	$V^\pi(s)$	$V^\pi(s)$	$Q(s, a)$	$Q_\theta(s, a)$
Sample, y	$ \begin{aligned} &r + \\ &\gamma r' + \\ &\gamma^2 r'' + \\ &\gamma^3 r''' + \\ &\gamma^4 r'''' + \\ &\dots \end{aligned} $	$ \begin{aligned} &r + \\ &\gamma V^\pi(s') \end{aligned} $	$ \begin{aligned} &r + \\ &\gamma \max_{a'} Q(s', a') \end{aligned} $	$ \begin{aligned} &r + \\ &\gamma \max_{a'} Q_\theta(s', a') \end{aligned} $