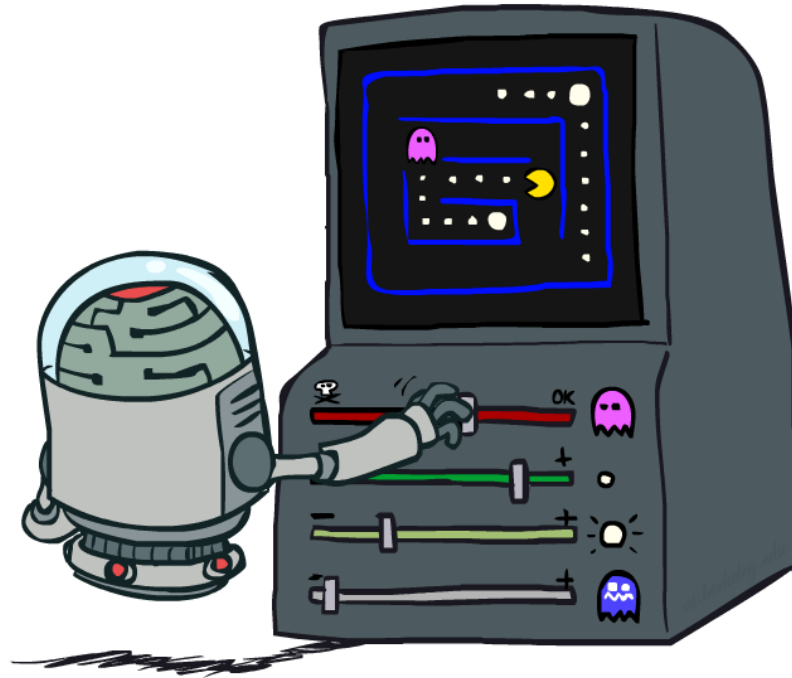


AI: Representation and Problem Solving

Reinforcement Learning II



Instructor: Pat Virtue

Slide credits: CMU AI and <http://ai.berkeley.edu>

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

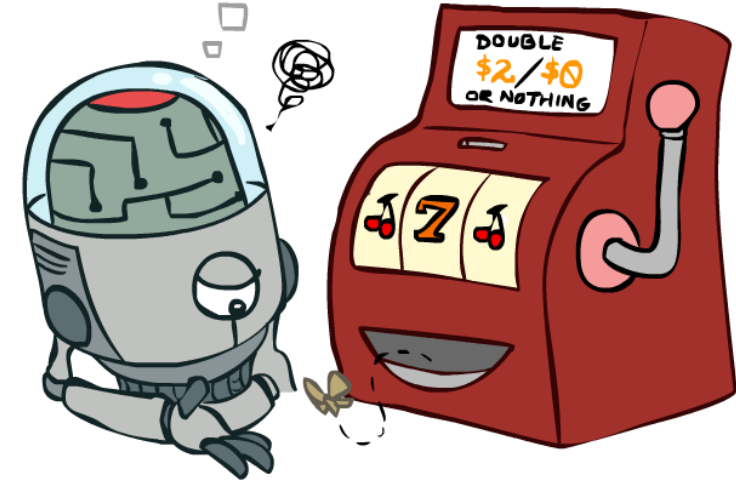
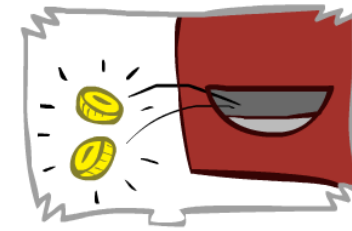
- Direct Evaluation (simple)
- TD Learning

Active Reinforcement Learning

- Q-Learning

Model
 $R(s,a,s')$
 $T(s,a,s')$
 $\emptyset$

We are given a policy
↑
to use.



Online Learning

Model-free Learning

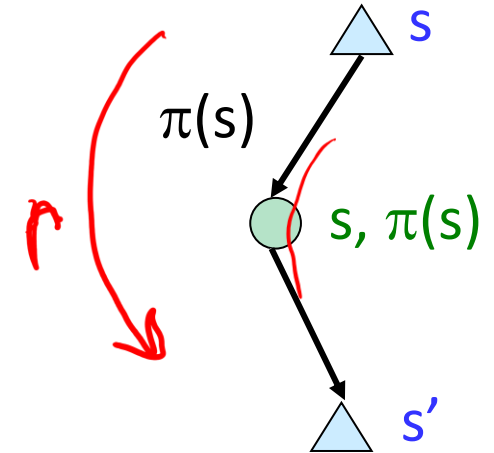
Passive Reinforcement Learning

Temporal Difference Learning

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

$$r + \gamma V^\pi(s')$$

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

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

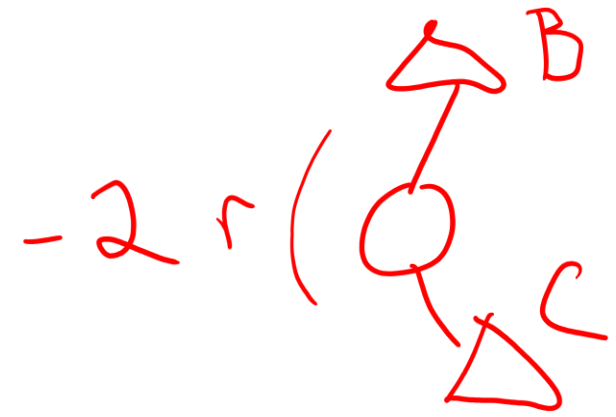
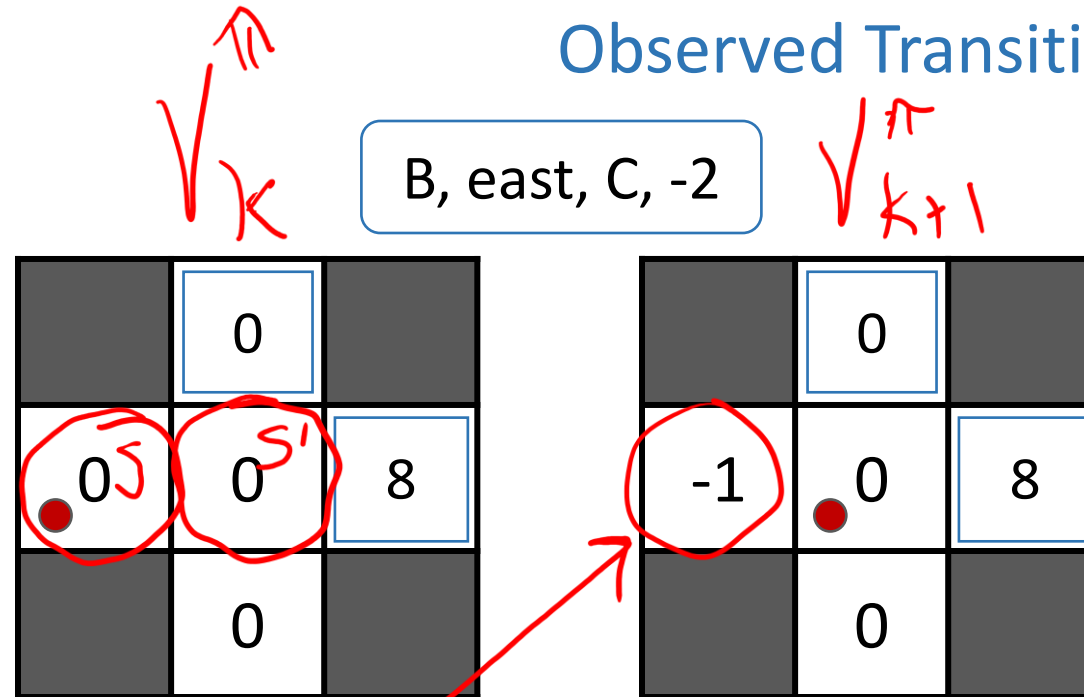
Example: Temporal Difference Learning

States

	A	
B	C	D
	E	

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

Observed Transitions



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

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

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

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	

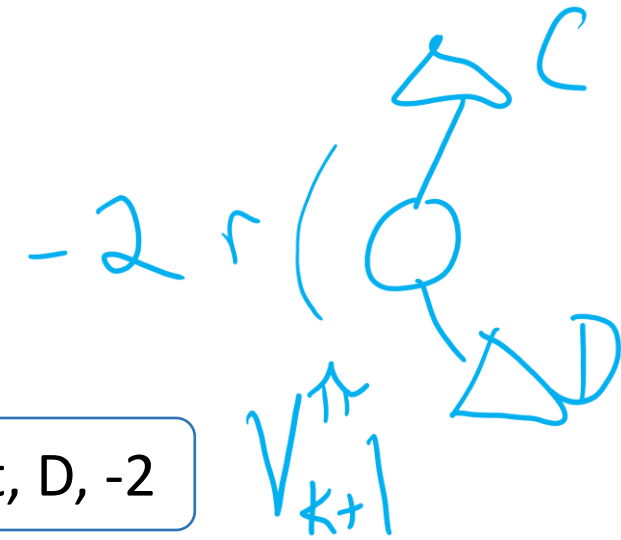
	0	
-1	0	8
	0	

	0	
-1	3	8
	0	

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

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

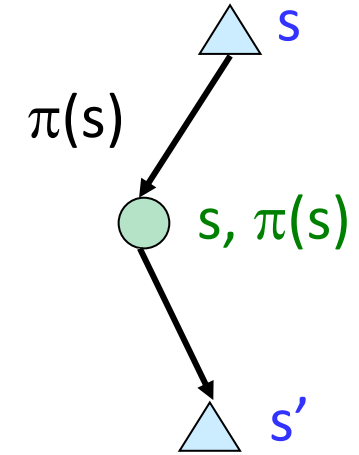
$$\underline{V^\pi(s)} \leftarrow \underline{V^\pi(s)} + \alpha(\underline{\text{sample}} - \underline{V^\pi(s)})$$



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

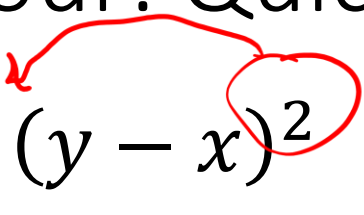
Same update: $V^\pi(s) \leftarrow V^\pi(s) + \alpha [sample - V^\pi(s)]$

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

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

Handwritten red annotations: "ML" with an arrow pointing to the error term, and "y" and "y-hat" pointing to "sample" and "V^\pi(s)" respectively.

ML detour: Quick Calculus Quiz

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


What is $\frac{df}{dx}$? $= 1 \cdot (y - x) \cdot (-1)$

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

$$V \leftarrow V - \alpha \frac{\partial \text{Error}}{\partial V} \text{Error}$$

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

$\uparrow$ Q
 $\uparrow$ Q_{func}
 $\uparrow$ Q_{net}

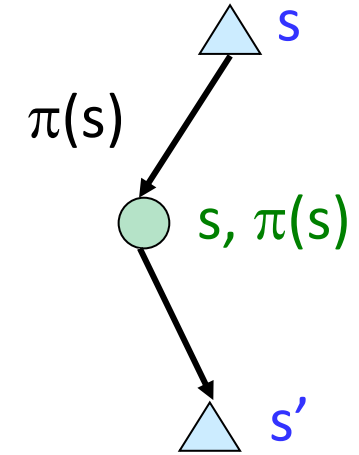
Temporal Difference Learning

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

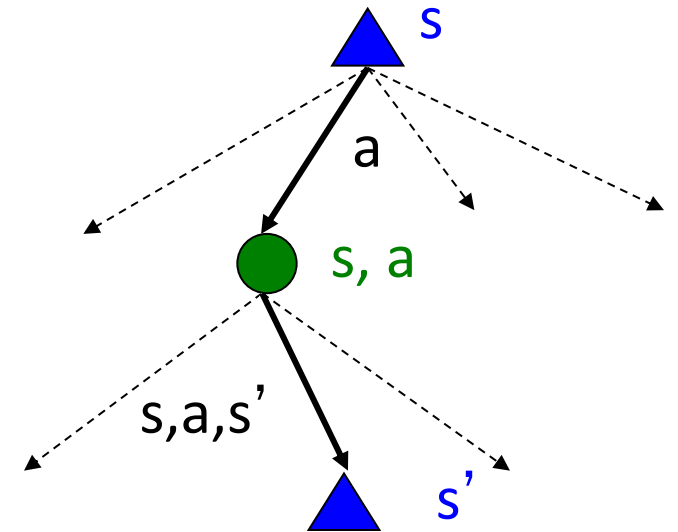
However, if we want to turn values into a (new) policy, we're sunk:

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

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

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:

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

Q-learning:

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

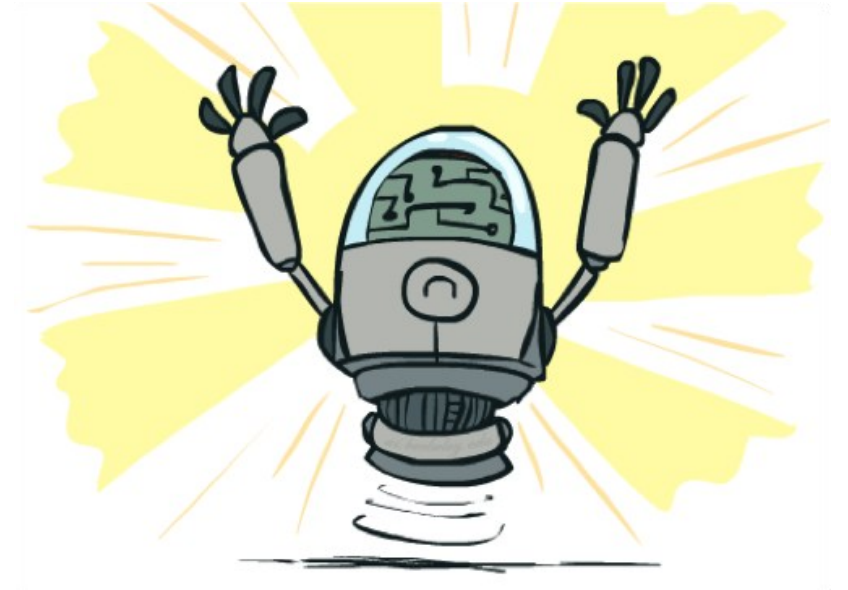
TD



Q



→ Agent gets to choose
next action



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:

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

- 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

on-policy Model-Free

Passive Reinforcement Learning

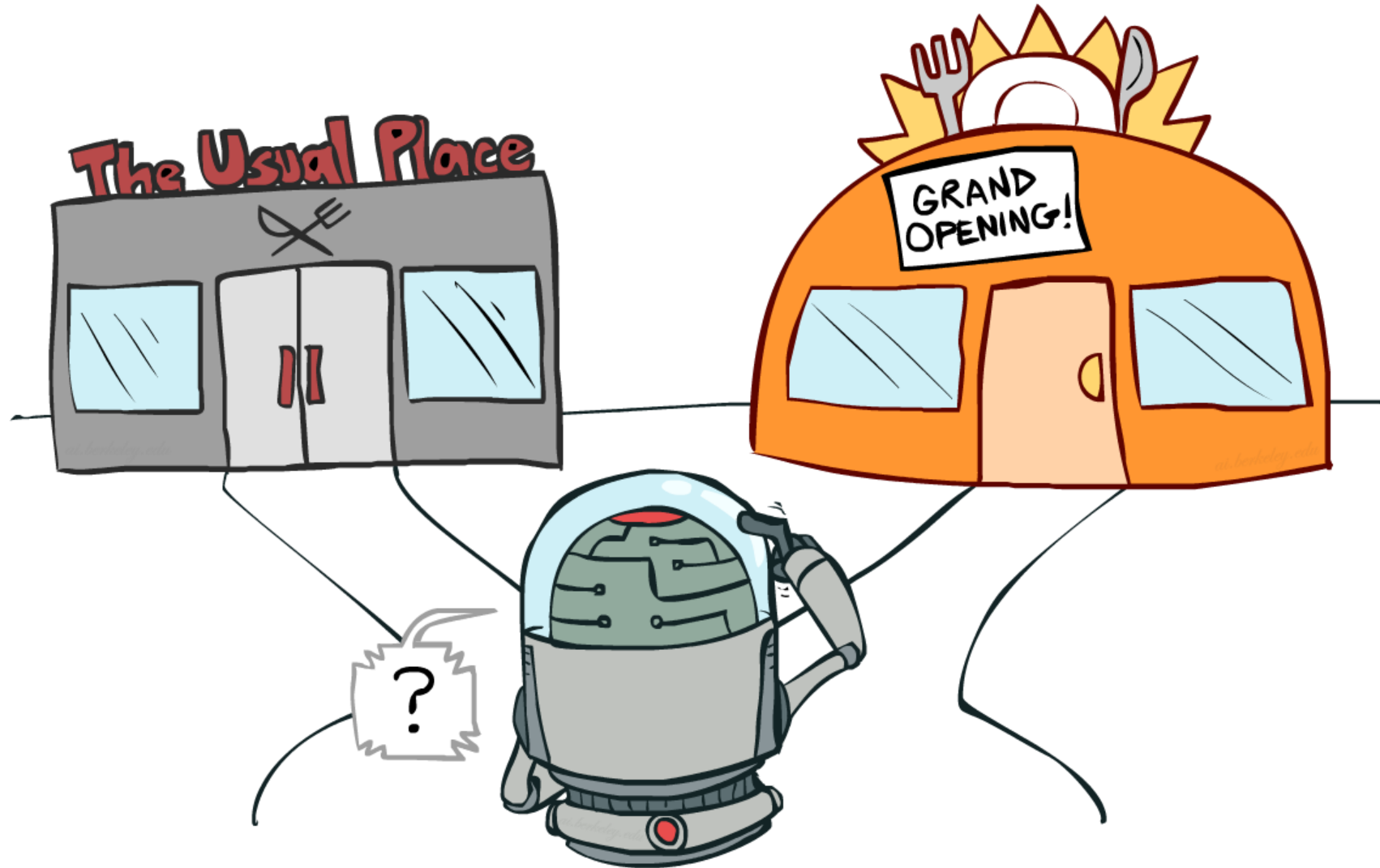
- Direct Evaluation (simple)
- TD Learning

off-policy
Active Reinforcement Learning

- Q-Learning

Demo Q-Learning Auto Cliff Grid

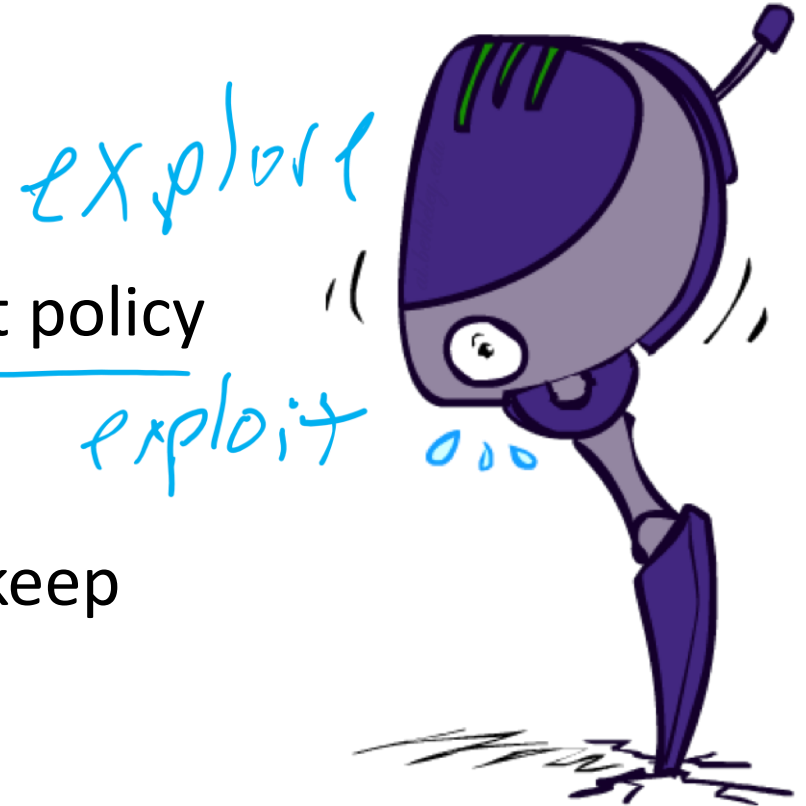
Exploration vs. Exploitation



How to Explore?

Several schemes for forcing exploration

- Simplest: random actions (ϵ -greedy)
 - Every time step, flip a coin
 - With (small) probability ϵ , act randomly
 - With (large) probability $1-\epsilon$, act on current policy
- 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 (L11D2)]

[Demo: Q-learning – epsilon-greedy -- crawler (L11D3)]

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.

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

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

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

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

$$\frac{Q}{10} \quad \frac{n}{1} \\ f(q, n) = 10 + \frac{100}{1} = 110$$

$$\frac{Q}{100} \quad \frac{n}{100} \\ f(100, 100) = 100 + 1 = 101$$

$k=100$

q



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!



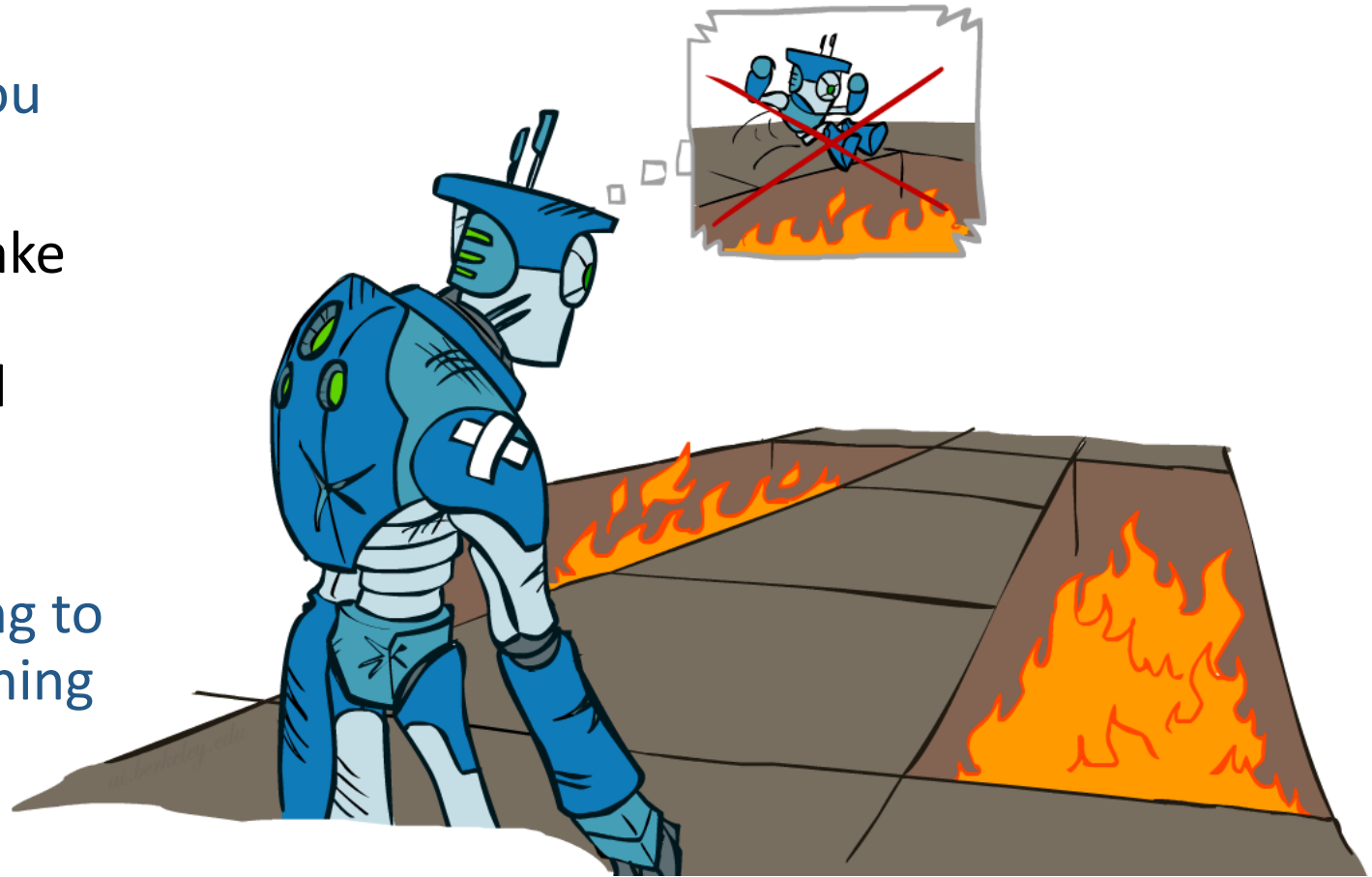
Regret

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

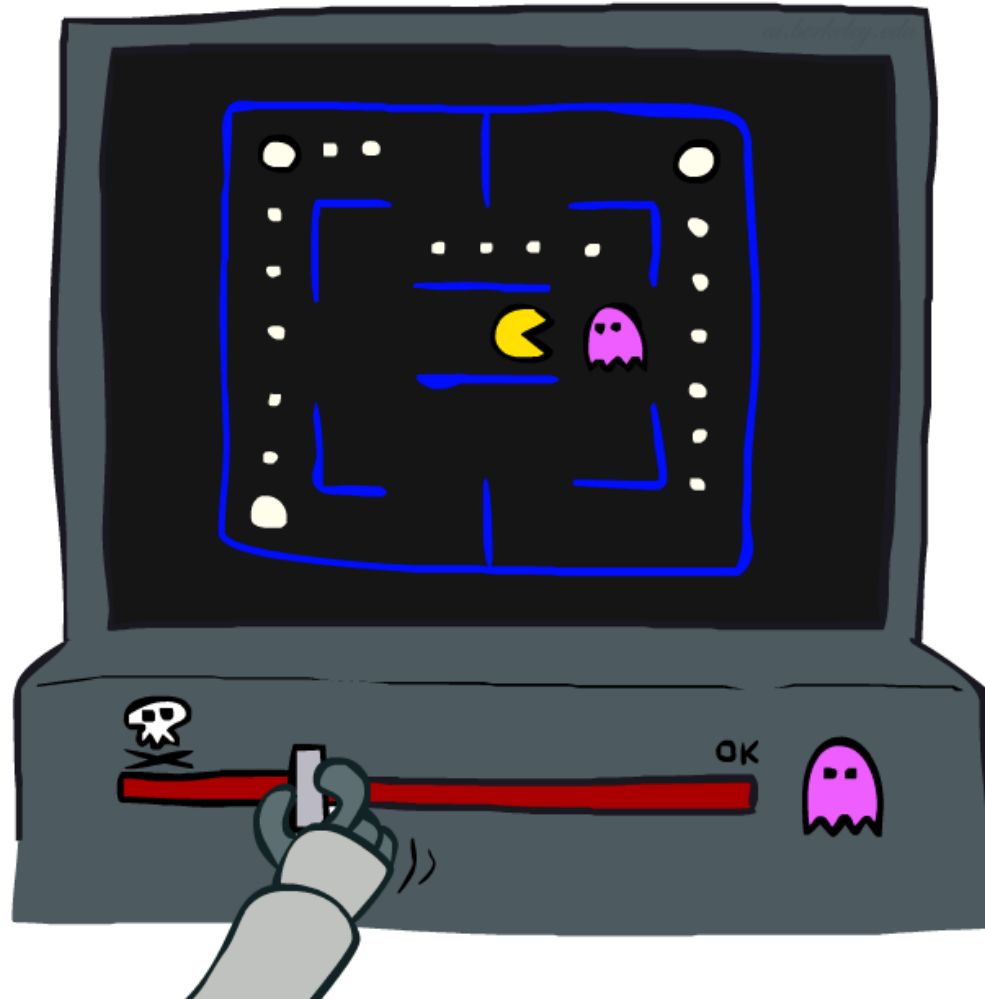
Regret is a measure of your total mistake cost: the difference between your (expected) rewards, including youthful suboptimality, and optimal (expected) rewards

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



Approximate Q-Learning



Approximate Q-Learning

What happens when we change Candy Grab to start with 1000 pieces?

Pieces Available	Take 1	Take 2
2	0%	100%
3	2%	0%
4	75%	2%
5	4%	68%
6	5%	6%
7	60%	5%

Generalizing Across States

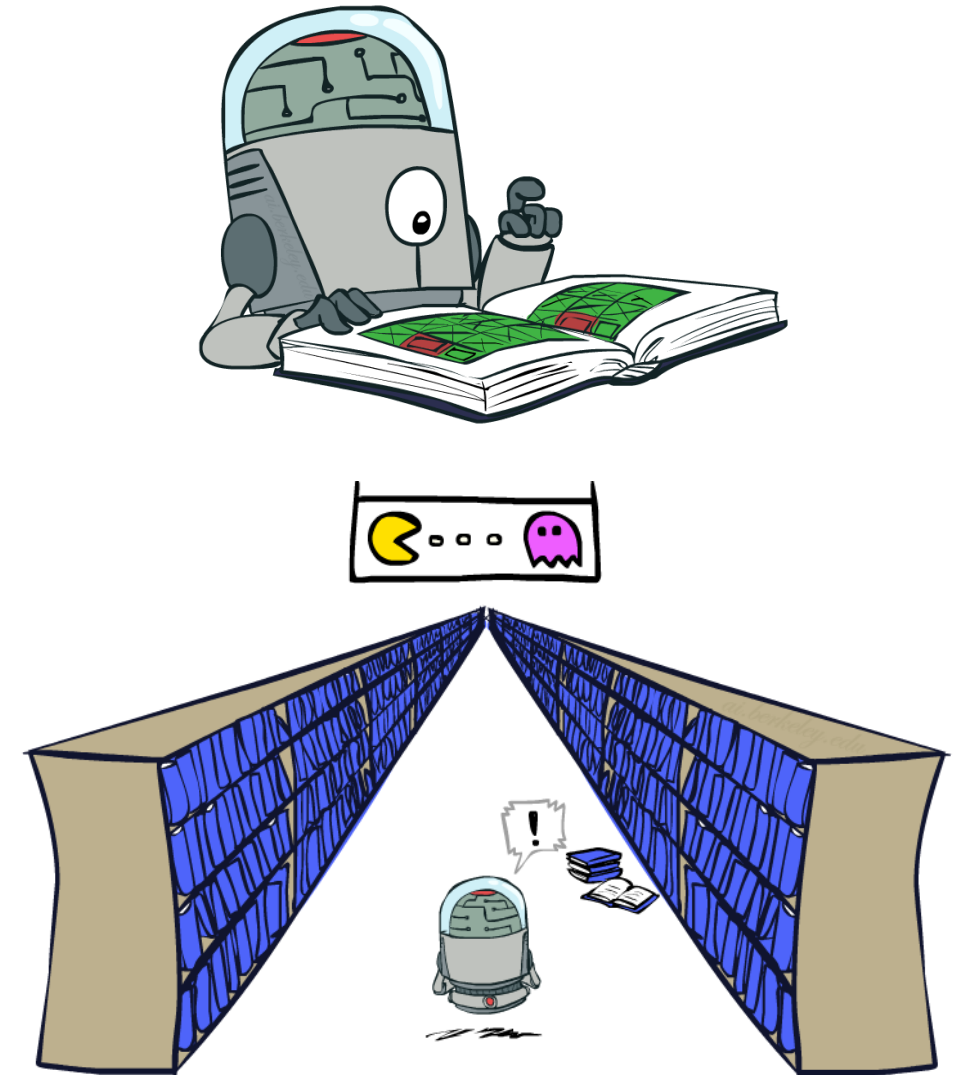
Basic Q-Learning keeps a table of all q-values

In realistic situations, we cannot possibly learn about every single state!

- Too many states to visit them all in training
- Too many states to hold the q-tables in memory

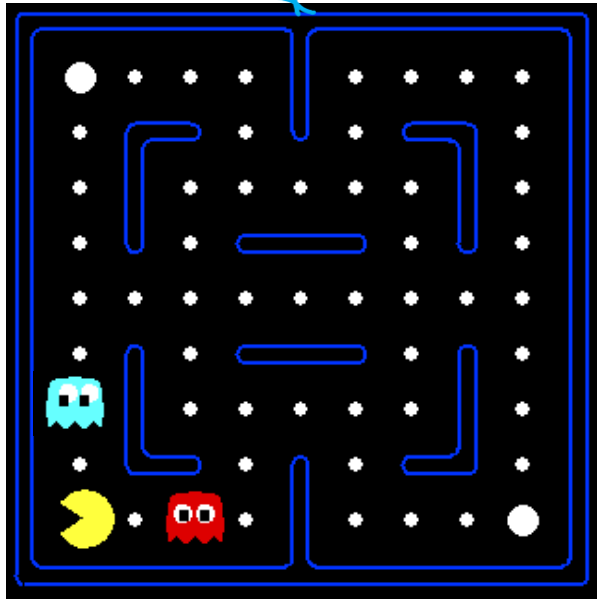
Instead, we want to generalize:

- Learn about some small number of training states from experience
- Generalize that experience to new, similar situations
- This is a fundamental idea in machine learning, and we'll see it over and over again

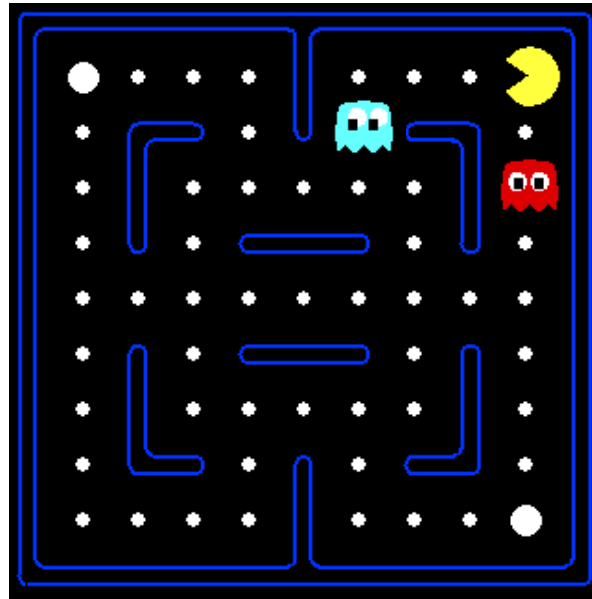


Example: Pacman

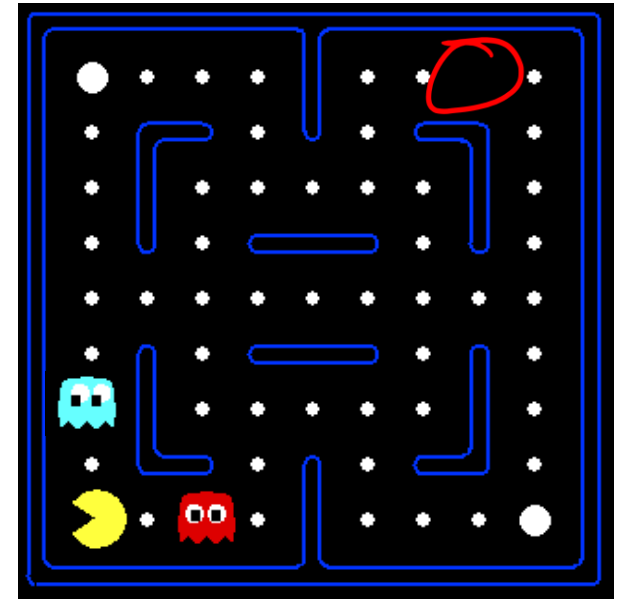
Let's say we discover through experience that this state is bad:



In naïve q-learning, we know nothing about this state:



Or even this one!

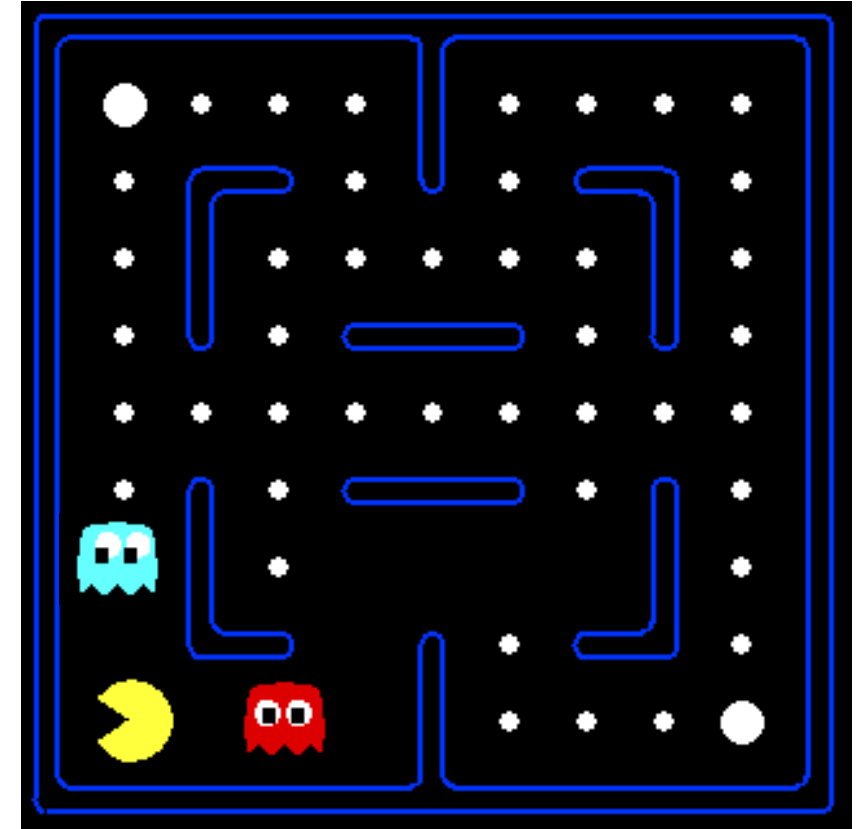


[Demo: Q-learning – pacman – tiny – watch all (L11D5)]
[Demo: Q-learning – pacman – tiny – silent train (L11D6)]
[Demo: Q-learning – pacman – tricky – watch all (L11D7)]

Feature-Based Representations

Solution: describe a state using a vector of features (properties)

- Features are functions from states to real numbers (often 0/1) that capture important properties of the state
- Example features:
 - Distance to closest ghost
 - Distance to closest dot
 - Number of ghosts
 - $1 / (\text{dist to dot})^2$
 - Is Pacman in a tunnel? (0/1)
 - etc.
 - Is it the exact state on this slide?
- Can also describe a q-state (s, a) with features (e.g. action moves closer to food)



Linear Value Functions

Using a feature representation, we can write a q function (or value function) for any state using a few weights:

- $V_w(s) = w_1 \underline{f_1(s)} + w_2 \underline{f_2(s)} + \dots + w_n \underline{f_n(s)}$

- $Q_w(s,a) = w_1 \underline{f_1(s,a)} + w_2 \underline{f_2(s,a)} + \dots + w_n \underline{f_n(s,a)}$

Advantage: our experience is summed up in a few powerful numbers

Disadvantage: states may share features but actually be very different in value!

Updating a linear value function

Original Q learning rule tries to reduce prediction error at s, a :

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

Instead, we update the weights to try to reduce the error at s, a :

$$\begin{aligned} \blacksquare \underline{w_i} &\leftarrow w_i + \alpha \cdot [r + \gamma \max_{a'} Q(s',a') - Q(s,a)] \partial Q_w(s,a) / \partial w_i \\ &= \underline{w_i} + \alpha \cdot \underbrace{[r + \gamma \max_{a'} Q(s',a') - Q(s,a)]}_{\text{error}} \underline{f_i(s,a)} \quad \leftarrow \end{aligned}$$

Quick Calculus Quiz

$$Error(w) = \frac{1}{2} (y - \underbrace{wf(x)})^2$$

Handwritten notes: "sample" with an arrow pointing to y, and "state" with an arrow pointing to f(x) in the original image.

$f_1(x)$
state

Last time

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

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

What is $\frac{dError}{dw}$? $-1 (y - \underbrace{wf(x)}) f'(x)$

Handwritten notes: "sample" with an arrow pointing to y, and "state" with an arrow pointing to f(x) in the original image.

$$\frac{d}{dw} \underbrace{wf(x)}$$

Updating a linear value function

Original Q learning rule tries to reduce prediction error at s, a :

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

Instead, we update the weights to try to reduce the error at s, a :

$$\begin{aligned} \blacksquare w_i &\leftarrow w_i + \alpha \cdot [r + \gamma \max_{a'} Q(s',a') - Q(s,a)] \frac{\partial Q_w(s,a)}{\partial w_i} \\ &= w_i + \alpha \cdot [r + \gamma \max_{a'} Q(s',a') - Q(s,a)] f_i(s,a) \end{aligned}$$


$$Q_w(s,a) = w_1 f_1(s,a) + w_2 f_2(s,a)$$

$$\frac{\partial Q}{\partial w_2} = f_2(s,a)$$

$$Error(w) = \frac{1}{2} (y - wf(x))^2$$

$$\frac{dError}{dw} = -(y - wf(x))f(x)$$

Approximate Q-Learning

$$Q(s, a) = w_1 f_1(s, a) + w_2 f_2(s, a) + \dots + w_n f_n(s, a)$$

Q-learning with linear Q-functions:

$$\text{transition} = (s, a, r, s')$$

$$\text{difference} = \left[r + \gamma \max_{a'} Q(s', a') \right] - Q(s, a)$$

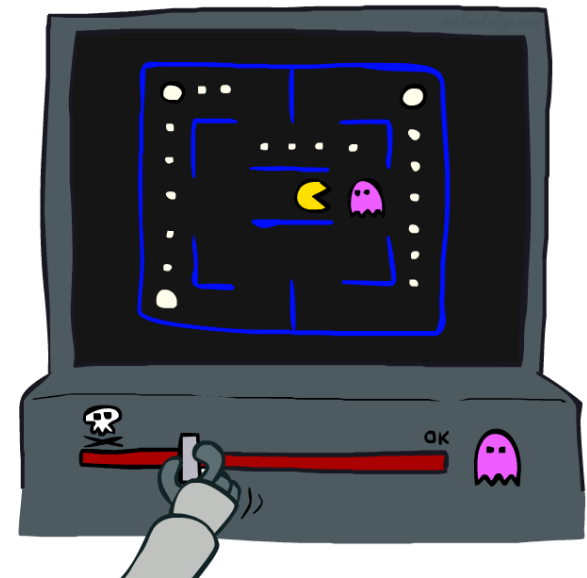
$$Q(s, a) \leftarrow Q(s, a) + \alpha [\text{difference}] \quad \text{Exact Q's}$$

$$w_i \leftarrow w_i + \alpha [\text{difference}] f_i(s, a) \quad \text{Approximate Q's}$$

Intuitive interpretation:

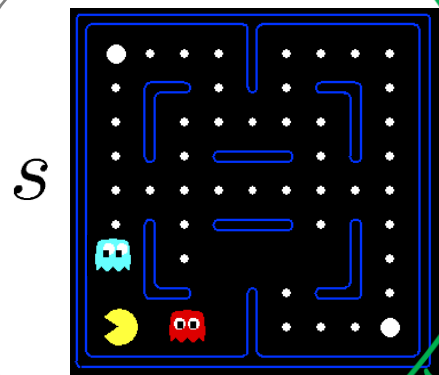
- Adjust weights of active features
- E.g., if something unexpectedly bad happens, blame the features that were on: disprefer all states with that state's features

Formal justification: online least squares



Example: Q-Pacman ✓

$$Q(s, a) = 4.0 \underline{f_{DOT}(s, a)} - 1.0 \underline{f_{GST}(s, a)}$$



s

$$f_{DOT}(s, \text{NORTH}) = 0.5$$

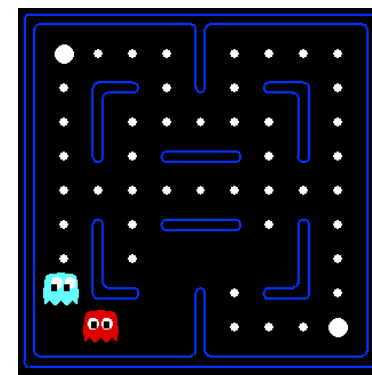
$$f_{GST}(s, \text{NORTH}) = 1.0$$

$$Q(s, \text{NORTH}) = +1$$

$$r + \gamma \max_{a'} Q(s', a') = -500 + 0$$

$$a = \text{NORTH}$$
$$r = -500$$

s'



$$Q(s', \cdot) = 0$$

difference = -501

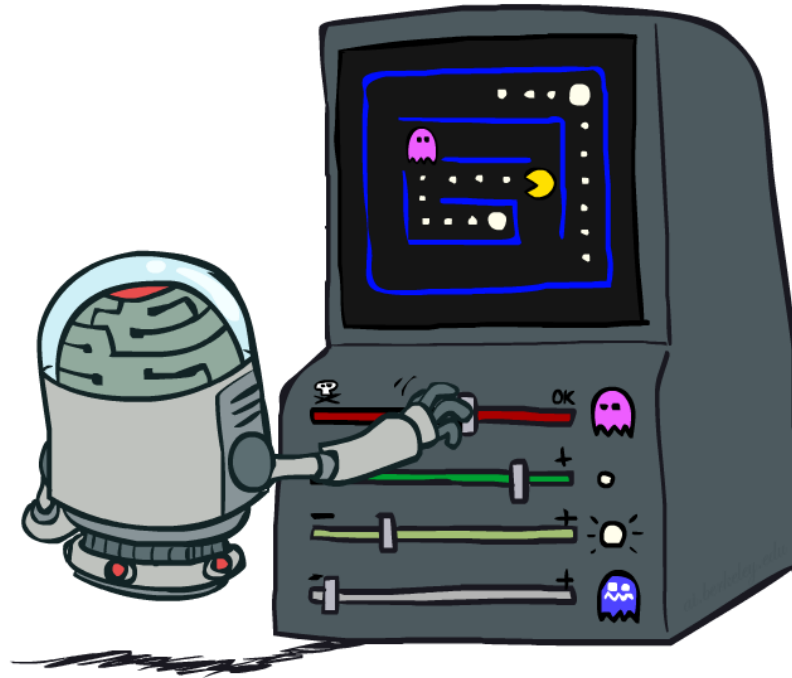
$$w_{DOT} \leftarrow 4.0 + \alpha [-501] 0.5$$

$$w_{GST} \leftarrow -1.0 + \alpha [-501] 1.0$$

$$Q(s, a) = 3.0 f_{DOT}(s, a) - 3.0 f_{GST}(s, a)$$

[Demo: approximate Q-learning pacman (L11D10)]

Reinforcement Learning Milestones



TDGammon

1992 by Gerald Tesauro, IBM

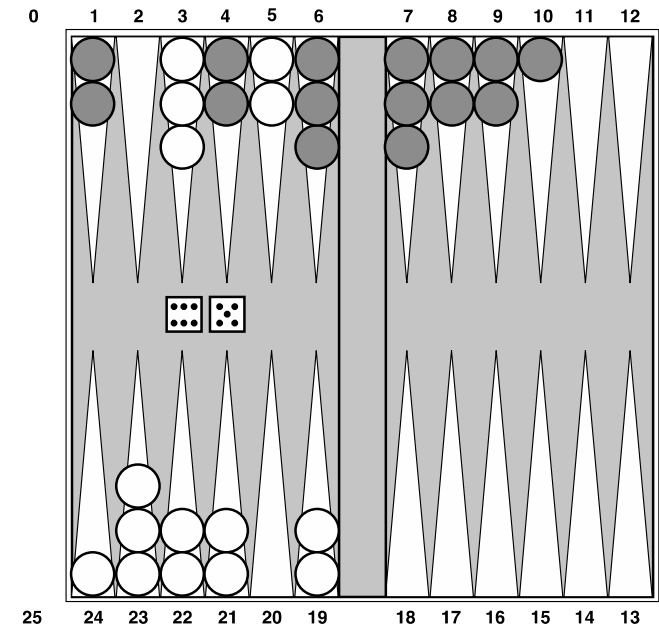
4-ply lookahead using $V(s)$ trained from 1,500,000 games of self-play

3 hidden layers, ~100 units each

Input: contents of each location plus several handcrafted features

Experimental results:

- Plays approximately at parity with world champion
- Led to radical changes in the way humans play backgammon



Deep Q-Networks

$$sample = r + \gamma \max_a Q_w(s', a')$$

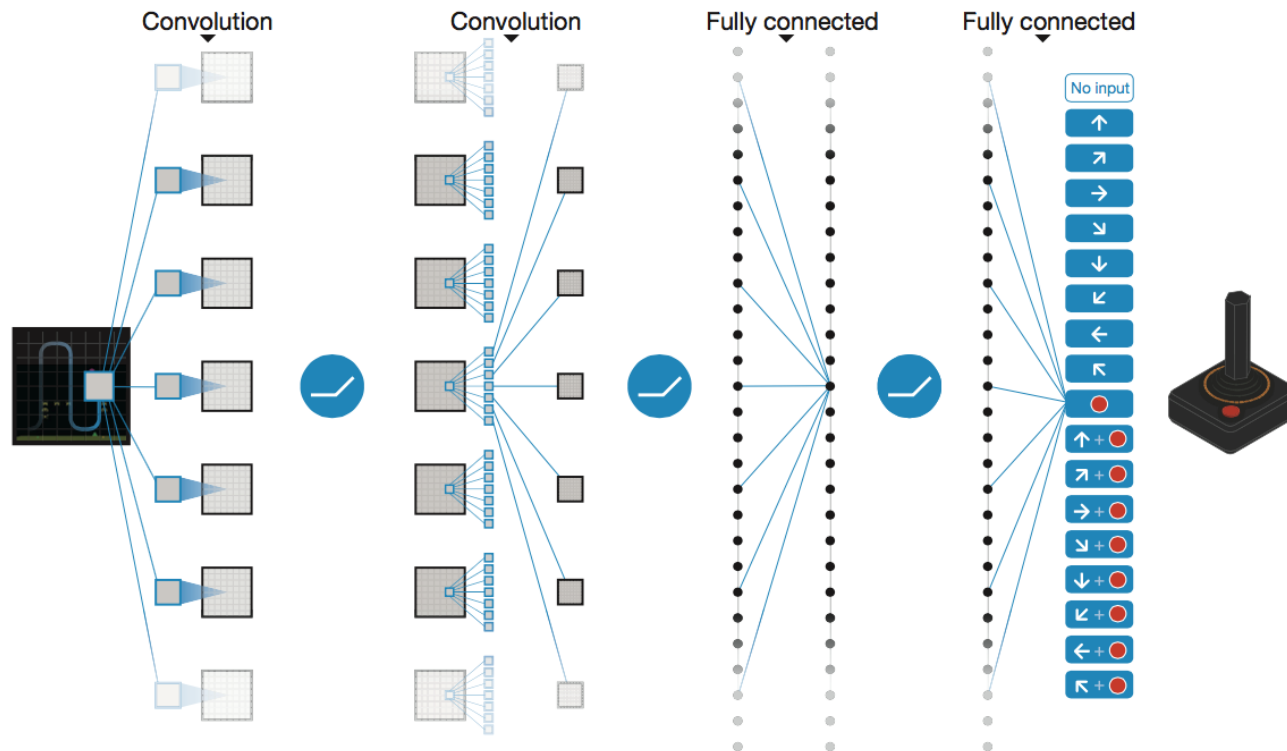
$Q_w(s, a)$: Neural network

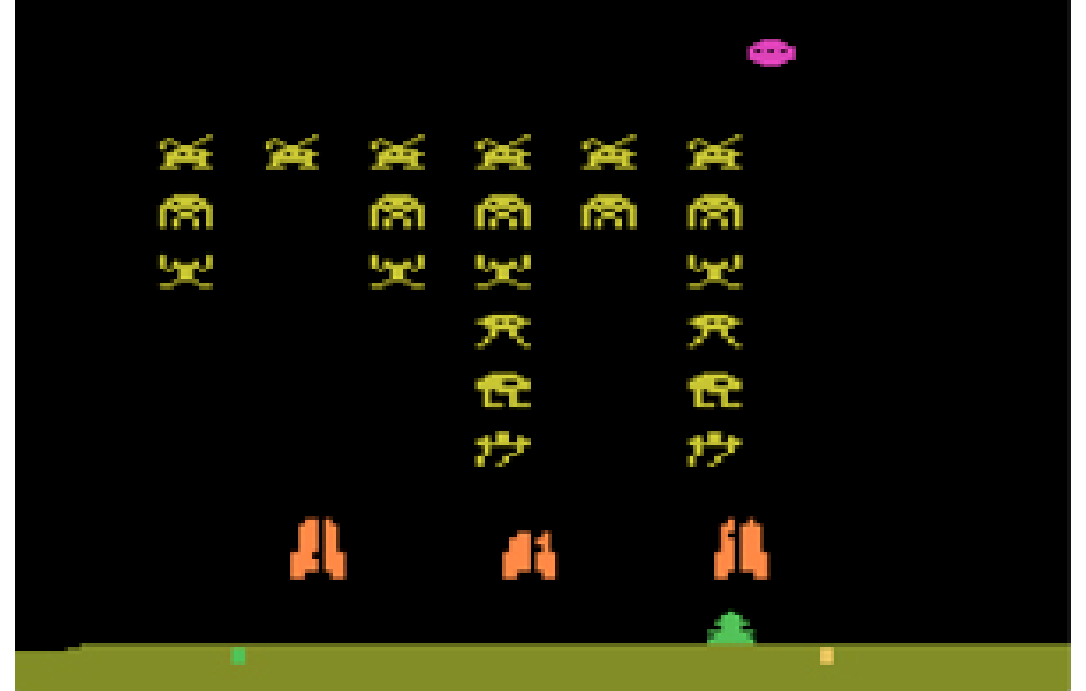
Deep Mind, 2015

Used a deep learning network to represent Q:

- Input is last 4 images (84x84 pixel values) plus score

49 Atari games, incl. Breakout, Space Invaders, Seaquest, Enduro





OpenAI Gym (now Gymnasium)

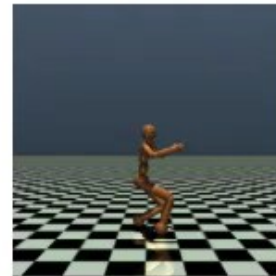
2016+

Benchmark problems for learning agents

<https://gymnasium.farama.org/>



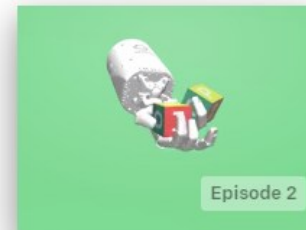
Ant



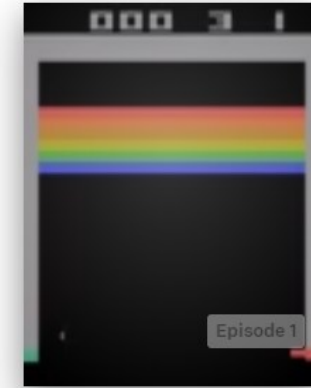
Humanoid



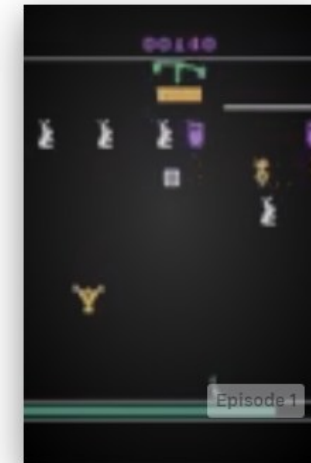
FetchPush-v0
Push a block to a goal position.



HandManipulateBlock-v0
Orient a block using a robot hand.



Breakout-ram-v0
Maximize score in the game Breakout, with RAM as input



Carnival-v0
Maximize score in the game Carnival, with screen images as input

AlphaGo, AlphaZero

Deep Mind, 2016+



Autonomous Vehicles?