

<https://introml.mit.edu/>

6.390 Intro to Machine Learning

Lecture 12: Reinforcement Learning

Shen Shen

Nov 20, 2025

11am, Room 10-250

[Interactive Slides and Lecture Recording](#)

Recall

MDP: Definition and terminologies

- $\mathcal{S}$: state space, contains all possible states s .
- $\mathcal{A}$: action space, contains all possible actions a .
- $T(s, a, s')$: the probability of transition from state s to s' when action a is taken.
- $R(s, a)$: reward, takes in a (state, action) pair and returns a reward.
- $\gamma \in [0, 1]$: discount factor, a scalar.

- $\pi(s)$: policy, takes in a state and returns an action.

The goal of an MDP is to find a "good" policy.

In 6.390,

- $\mathcal{S}$ and $\mathcal{A}$ are small discrete sets, unless otherwise specified.
- s' and a' are short-hand for the next-timestep state and action.
- $R(s, a)$ is deterministic and bounded.
- $\pi(s)$ is deterministic.

Recall



Use the definition and sum up expected rewards:

1
$$V_h^\pi(s) := \mathbb{E} \left[\sum_{t=0}^{h-1} \gamma^t R(s_t, \pi(s_t)) \mid s_0 = s, \pi \right]$$

Or, leverage the recursive structure:

2 **finite**-horizon Bellman recursions

$$V_h^\pi(s) = R(s, \pi(s)) + \gamma \sum_{s'} T(s, \pi(s), s') V_{h-1}^\pi(s')$$

3 **infinite**-horizon Bellman equations

$$V_\infty^\pi(s) = R(s, \pi(s)) + \gamma \sum_{s'} T(s, \pi(s), s') V_\infty^\pi(s')$$

Recall

horizon- h value in state s : the expected sum of discounted rewards, starting in state s and following policy π for h steps.

$$2 \quad V_h^\pi(s) = R(s, \pi(s)) + \gamma \sum_{s'} T(s, \pi(s), s') V_{h-1}^\pi(s')$$

the immediate reward for taking the policy-prescribed action $\pi(s)$ in state s .

$(h - 1)$ horizon future value at a next state s'

sum of future values weighted by the probability of reaching that next state s'

discounted by γ

Recall

the optimal state-action value functions $Q_h^*(s, a)$:

the expected sum of discounted rewards, obtained by

- starting in state s
- take action a , for one step
- act **optimally** thereafter for the remaining $(h - 1)$ steps

$$\boxed{4} \quad V_h^*(s) = \max_a [R(s, a) + \gamma \sum_{s'} T(s, a, s') V_{h-1}^*(s')] = \max_a [Q_h^*(s, a)]$$

Q^* satisfies the Bellman recursion:

$$\boxed{5} \quad Q_h^*(s, a) = R(s, a) + \gamma \sum_{s'} T(s, a, s') \max_{a'} Q_{h-1}^*(s', a')$$

Recall

Value iteration: iteratively invoke

5

$$Q_h^*(s, a) = R(s, a) + \gamma \sum_{s'} T(s, a, s') \max_{a'} Q_{h-1}^*(s', a')$$

Value Iteration

1. **for** $s \in \mathcal{S}, a \in \mathcal{A}$:

2. $Q_{\text{old}}(s, a) = 0$

3. **while** True:

4. **for** $s \in \mathcal{S}, a \in \mathcal{A}$:

5. $Q_{\text{new}}(s, a) \leftarrow R(s, a) + \gamma \sum_{s'} T(s, a, s') \max_{a'} Q_{\text{old}}(s', a')$

6. **if** $\max_{s,a} |Q_{\text{old}}(s, a) - Q_{\text{new}}(s, a)| < \epsilon$:

7. **return** Q_{new}

8. $Q_{\text{old}} \leftarrow Q_{\text{new}}$

if run this block h
times and *break*, then
the returns are Q_h^*

returns are Q_∞^*

{

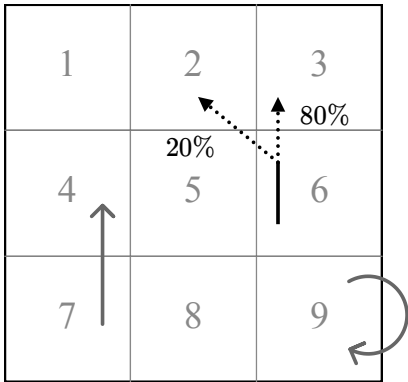
{

Outline

- Reinforcement learning setup
- Tabular Q-learning
 - exploration vs. exploitation
 - ϵ -greedy action selection
- Fitted Q-learning
- (Reinforcement learning setup again)



Mario in a grid-world *v1.0* (Markov-decision-process version)



0	0	1
0	0	1
0	0	1
0	0	-10
0	0	-10
0	0	-10
0	0	0
0	0	0
0	0	0

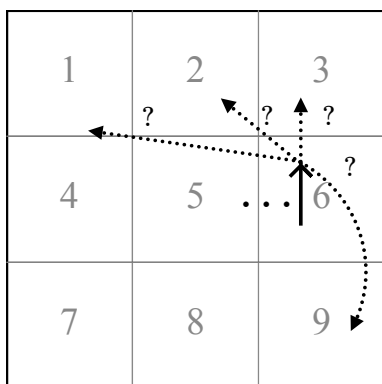
- 9 possible *states*
- 4 possible *actions*: {Up ↑, Down ↓, Left ←, Right →}
- *transition* probabilities are **known**

$$\text{e.g., } T(7, \uparrow, 4) = 1 \quad T(9, \rightarrow, 9) = 1 \quad T(6, \uparrow, 3) = 0.8 \quad T(6, \uparrow, 2) = 0.2$$

- *rewards* **known**
- *discount factor* $\gamma = 0.9$

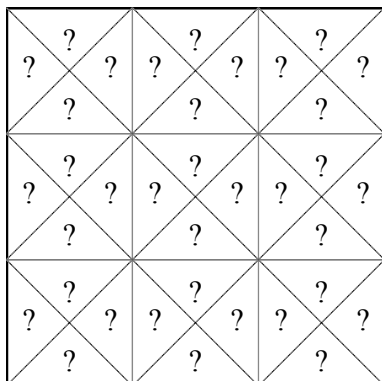


Mario in a grid-world *v2.0* (reinforcement learning version)



- 9 possible *states*
- 4 possible *actions*: {Up $\uparrow$, Down $\downarrow$, Left $\leftarrow$, Right $\rightarrow$ }
- *transition* probabilities are **unknown**

e.g., $T(7, \uparrow, 4) = ?$ $T(9, \rightarrow, 9) = ?$ $T(6, \uparrow, 3) = ?$ $T(6, \uparrow, 2) = ?$



- *rewards* **unknown**
- *discount factor* $\gamma = 0.9$

~~Markov Decision Processes~~ - Definition and terminologies

Reinforcement Learning

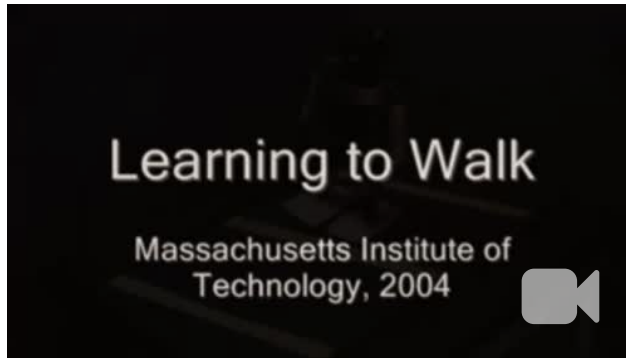
- $\mathcal{S}$: state space, contains all possible states s .
- $\mathcal{A}$: action space, contains all possible actions a .
- $T(s, a, s')$: the probability of transition from state s to s' when action a is taken
- $R(s, a)$: reward, takes in a (state, action) pair and returns a reward.
- $\gamma \in [0, 1]$: discount factor, a scalar.
- $\pi(s)$: policy, takes in a state and returns an action.

The goal of an ~~MDP~~ problem is to find a "good" policy.

RL

Reinforcement learning is very *general*:

robotics

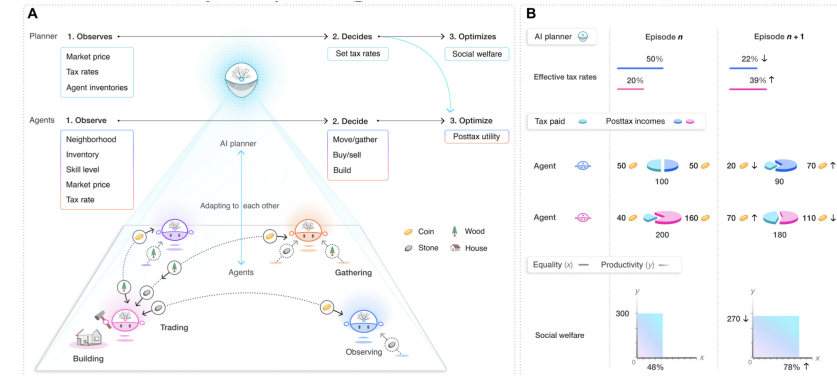


[Mastering the game of Go with deep neural networks and tree search. Silver et al. Nature 2017]

games



social sciences



[The AI Economist: Taxation policy design via two-level deep multiagent RL. Zheng et al. Science 2022]

chatbot (RLHF)



[Aligning language models to follow instructions. Ouyang et al. 2022]

health care

nature medicine

Explore content ▾ About the journal ▾ Publish with us ▾

[nature](#) > [nature medicine](#) > [articles](#) > article

Article | [Published: 13 January 2022](#)

Optimizing risk-based breast cancer treatment with reinforcement learning

[Adam Yala](#) , [Peter G. Mikhael](#), [Constance Lehman](#), [Gigin Lin](#), [Kevin Hughes](#), [Siddharth Satuluru](#), [Thomas Kim](#), [Imon Banerjee](#), [Barzilay](#)

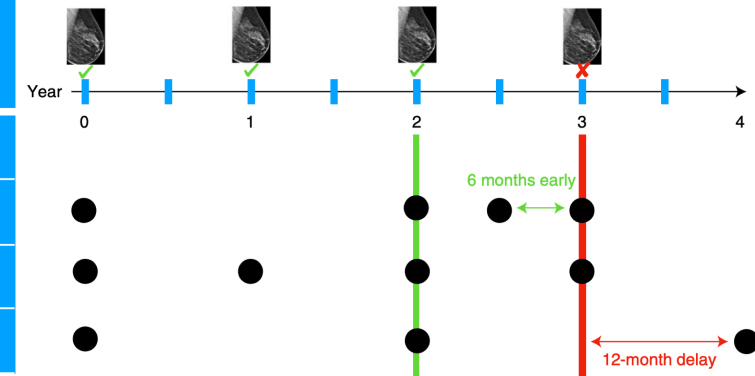
Retrospective patient trajectory

Recommended trajectories

Tempo-mirai

Annual

Biennial



Model-based RL: Learn the MDP tuple

1. Collect set of experiences (s, a, r, s')

2. Estimate $\hat{T}, \hat{R}$

e.g. $\hat{T}(6, \uparrow, 2) \approx$

$\frac{\text{observed } (6, \uparrow, 2) \text{ count}}{\text{total gameplay count}}$

e.g. $\hat{R}(6, \uparrow) =$

$\frac{\text{observed reward received from } (6, \uparrow)}{\text{total gameplay count}}$

3. Solve $\langle \mathcal{S}, \mathcal{A}, \hat{T}, \hat{R}, \gamma \rangle$ via e.g. *Value Iteration*



$\gamma = 0.9$

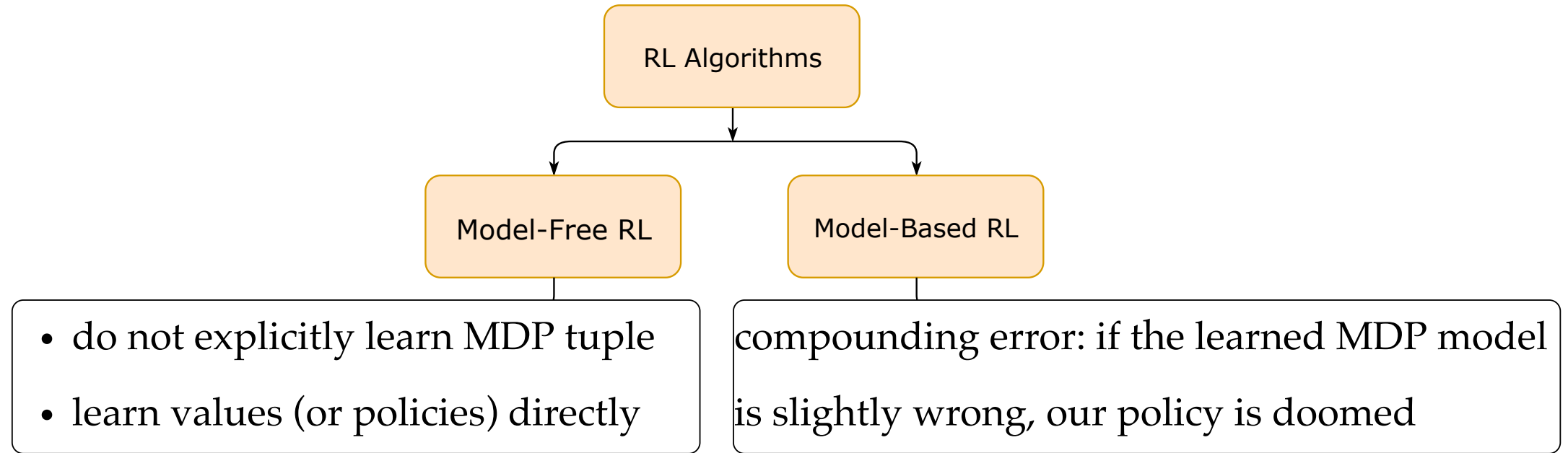
Unknown transition:

1	2	3
?	?	?
4	5	6
7	8	9

Unknown rewards

?	?	?
?	?	?
?	?	?
?	?	?

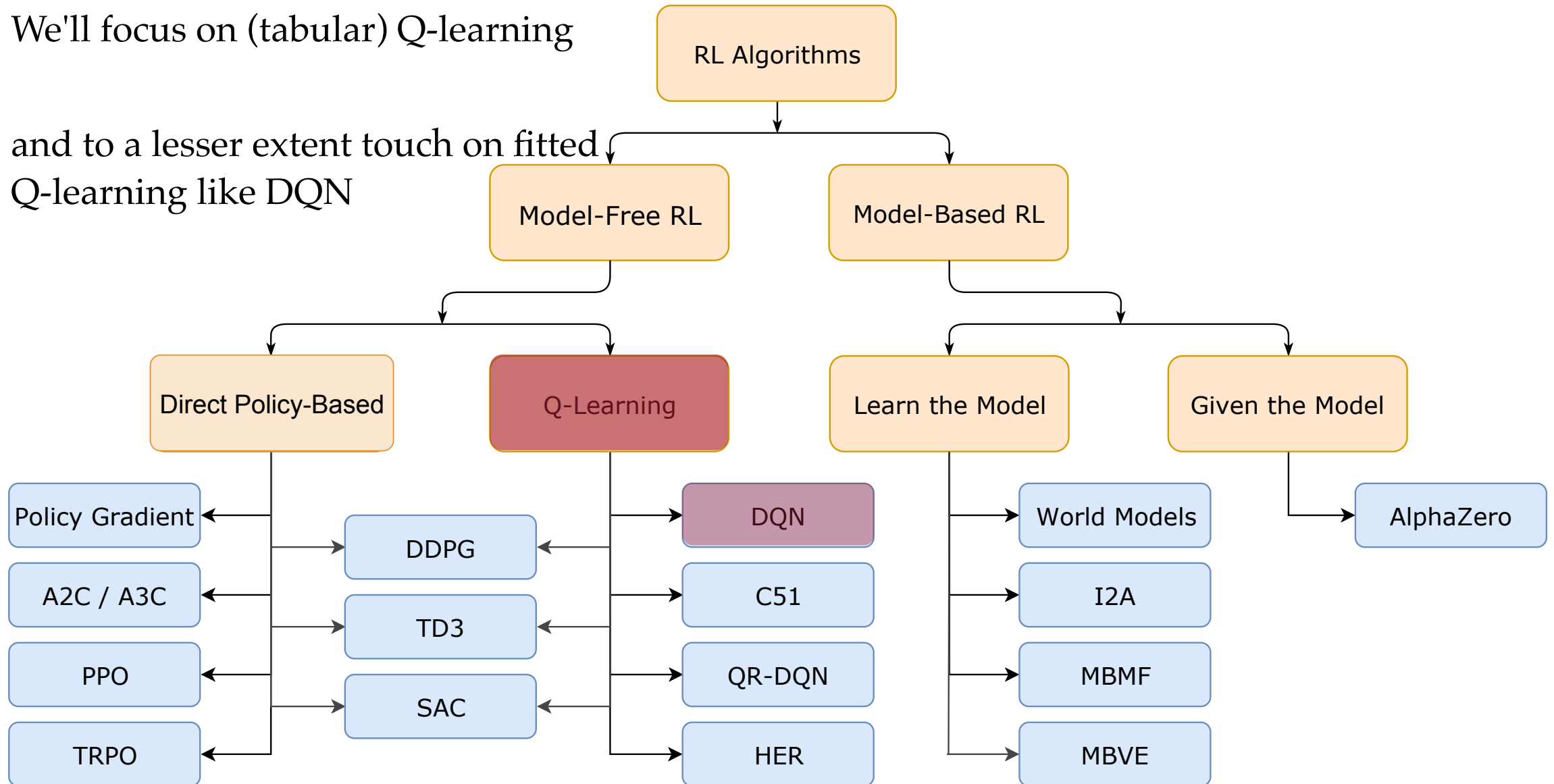
(s, a)	r	s'
$(1, \uparrow)$	0	1
$(1, \downarrow)$	0	4
$(3, \uparrow)$	1	3
$(3, \downarrow)$	1	6
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2



- We rarely say *hypothesis* in RL; instead, we refer to what's learned as a *model*, *value*, or *policy*
- *Model* in MDP/RL typically means the MDP tuple $\langle \mathcal{S}, \mathcal{A}, T, R, \gamma \rangle$

We'll focus on (tabular) Q-learning

and to a lesser extent touch on fitted
Q-learning like DQN



[A non-exhaustive, but useful taxonomy of algorithms in modern RL. [Source](#)]

Outline

- Reinforcement learning setup
- Tabular Q-learning
 - exploration vs. exploitation
 - ϵ -greedy action selection
- Fitted Q-learning
- (Reinforcement learning setup again)

Is it possible to get an optimal policy **without** learning transition or rewards explicitly?

Yes! We know one way already:

We switch to an infinite-horizon scenario. For our stochastic machine, here is the infinite-horizon Q_{∞}^* function (computed via value iteration) for γ near 1.

	wash	paint	eject
dirty	[2.32274541	-0.70048204	0.]
clean	[2.32274541	5.71581775	0.]
painted	[2.32274541	6.9	10.]
ejected	[0.	0.	0.]]

What is the optimal thing to do with a **clean** object?

What will you do if it becomes **dirty**?

Does this optimal policy make intuitive sense?

(Recall, from the MDP lab)

Optimal policy π^* easily extracted from Q^* :

$$\boxed{6} \quad \pi_h^*(s) = \arg \max_a Q_h^*(s, a)$$

But... didn't we arrive at Q^* by value iteration,
and doesn't value iteration rely on transition and rewards explicitly?

Value Iteration

1. **for** $s \in \mathcal{S}, a \in \mathcal{A}$:
2. $Q_{\text{old}}(s, a) = 0$
3. **while** True:
4. **for** $s \in \mathcal{S}, a \in \mathcal{A}$:
5. $Q_{\text{new}}(s, a) \leftarrow R(s, a) + \gamma \sum_{s'} T(s, a, s') \max_{a'} Q_{\text{old}}(s', a')$
6. **if** $\max_{s,a} |Q_{\text{old}}(s, a) - Q_{\text{new}}(s, a)| < \epsilon$:
7. **return** Q_{new}
8. $Q_{\text{old}} \leftarrow Q_{\text{new}}$

- Indeed, *value iteration* relies on having full access to R and T

$$Q_{\text{new}}(s, a) \leftarrow \underbrace{R(s, a)}_{\text{immediate reward}} + \gamma \sum_{s'} \underbrace{T(s, a, s')}_{\text{expected future value, weighted by the chance of landing in that particular future state } s'} \underbrace{\max_{a'} Q_{\text{old}}(s', a')}_{\text{future value, starting in state } s' \text{ and acting } \textit{optimally} \text{ for } (h-1) \text{ steps}}$$

immediate
reward

future value, starting in state s' and
acting *optimally* for $(h-1)$ steps

expected future value, weighted by the chance
of landing in that particular future state s'

- Without R and T , how about: execute (s, a) , observe r and s' , and update:

$$Q_{\text{new}}(s, a) \leftarrow \underbrace{r}_{\text{target}} + \gamma \max_{a'} Q_{\text{old}}(s', a')$$

target

(we will see this idea has issues)

$$Q_{\text{new}}(s, a) \leftarrow r + \gamma \max_{a'} Q_{\text{old}}(s', a')$$



$\gamma = 0.9$

$Q_{\text{old}}(s, a)$

0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0

$Q_{\text{new}}(s, a)$

0	0	0	0	1	1
0	0	0	0	1	1
0	0	0	0	-10	-10
0	0	0	0	-10	-10
0	0	0	0	0	0
0	0	0	0	0	0

(s, a)	r	s'
$(1, \uparrow)$	0	1
$(1, \downarrow)$	0	1
...		
$(3, \uparrow)$	1	3
$(3, \downarrow)$	1	6
...		
$(6, \uparrow)$	-10	3
$(6, \downarrow)$	-10	2
...		

$$Q_{\text{new}}(s, a) \leftarrow r + \gamma \max_{a'} Q_{\text{old}}(s', a')$$



$$\gamma = 0.9$$

 $Q_{\text{old}}(s, a)$

0	0	0	0	1	1
0	0	0	0	1	1
0	0	0	0	1	1
0	0	0	0	-10	-10
0	0	0	0	-10	-10
0	0	0	0	-10	-10
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0

 $Q_{\text{new}}(s, a)$

0	0	0	0	1	1
0	0	0	0	1	1
0	0	0	0	1	1
0	0	0	0	-10	-10
0	0	0	0	-10	-10
0	0	0	0	-10	-10
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0

$$Q_{\text{new}}(6, \uparrow) \leftarrow -10 + \gamma \max_{a'} Q_{\text{old}}(3, a')$$

$$= -10 + 0.9 * 1 = -9.1$$

(s, a)	r	s'
...	...	...
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
...	...	...

$$Q_{\text{new}}(s, a) \leftarrow r + \gamma \max_{a'} Q_{\text{old}}(s', a')$$

$Q_{\text{old}}(s, a)$

0	0	0	0	1	1
0	0	0	0	1	1
0	0	0	0	-9.1	-10
0	0	0	0	-10	-10
0	0	0	0	0	0
0	0	0	0	0	0

$Q_{\text{new}}(s, a)$

0	0	0	0	1	1
0	0	0	0	1	1
0	0	0	0	-10	-10
0	0	0	0	-10	-10
0	0	0	0	0	0
0	0	0	0	0	0

$$Q_{\text{new}}(6, \uparrow) \leftarrow -10 + \gamma \max_{a'} Q_{\text{old}}(2, a')$$

$$= -10 + 0.9 * 0 = -10$$



$\gamma = 0.9$

(s, a)	r	s'
...	...	...
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
...	...	...

$$Q_{\text{new}}(s, a) \leftarrow r + \gamma \max_{a'} Q_{\text{old}}(s', a')$$



$$\gamma = 0.9$$

$Q_{\text{old}}(s, a)$

0	0	0	0	1	1
0	0	0	0	1	1
0	0	0	0	-10	-10
0	0	0	0	-10	-10
0	0	0	0	0	0
0	0	0	0	0	0

$Q_{\text{new}}(s, a)$

0	0	0	0	1	1
0	0	0	0	1	1
0	0	0	0	-10	-10
0	0	0	0	-10	-10
0	0	0	0	0	0
0	0	0	0	0	0

$$Q_{\text{new}}(6, \uparrow) \leftarrow -10 + \gamma \max_{a'} Q_{\text{old}}(3, a')$$

$$= -10 + 0.9 * 1 = -9.1$$

(s, a)	r	s'
...	...	...
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
...	...	...

$$Q_{\text{new}}(s, a) \leftarrow r + \gamma \max_{a'} Q_{\text{old}}(s', a')$$



$$\gamma = 0.9$$

 $Q_{\text{old}}(s, a)$

	0	0	0	0	1	1
0	0	0	0	0	1	1
0	0	0	0	0	1	1
0	0	0	0	0	1	1
0	0	0	0	0	1	1
0	0	0	0	0	1	1
0	0	0	0	0	1	1

 $Q_{\text{new}}(s, a)$

	0	0	0	0	1	1
0	0	0	0	0	1	1
0	0	0	0	0	1	1
0	0	0	0	0	1	1
0	0	0	0	0	1	1
0	0	0	0	0	1	1
0	0	0	0	0	1	1

$$Q_{\text{new}}(6, \uparrow) \leftarrow -10 + \gamma \max_{a'} Q_{\text{old}}(2, a')$$

$$= -10 + 0.9 * 0 = -10$$

(s, a)	r	s'
...	...	...
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
...	...	...

$$Q_{\text{new}}(s, a) \leftarrow r + \gamma \max_{a'} Q_{\text{old}}(s', a')$$

target

Whenever observe $(6, \uparrow), -10, 3$:

$$\begin{aligned} Q_{\text{new}}(6, \uparrow) &\leftarrow -10 + \gamma \max_{a'} Q_{\text{old}}(3, a') \\ &= -10 + 0.9 * 1 = -9.1 \end{aligned}$$

Whenever observe $(6, \uparrow), -10, 2$:

$$\begin{aligned} Q_{\text{new}}(6, \uparrow) &\leftarrow -10 + \gamma \max_{a'} Q_{\text{old}}(2, a') \\ &= -10 + 0.9 * 0 = -10 \end{aligned}$$



$\gamma = 0.9$

(s, a)	r	s'
...	...	...
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
...	...	...

😞 Simply commit to the new target keeps "washing away" the old belief

💖 merge *old belief* and *target*

$$\begin{aligned} & (1 - \text{learning rate}) \text{ old belief} + \text{learning rate} \text{ target} \\ Q_{\text{new}}(s, a) \leftarrow & (1 - \alpha) Q_{\text{old}}(s, a) + \alpha \left(r + \gamma \max_{a'} Q_{\text{old}}(s', a') \right) \end{aligned}$$

- Core update rule of *Q-learning*
- $\alpha \in [0, 1]$, hyperparameter, trades off how much we trust new evidence versus old belief
- What we saw was $\alpha = 1$, trusting new experience entirely
- Let's see an example using $\alpha = 0.7$

$$Q_{\text{new}}(s, a) \leftarrow (1 - \alpha) Q_{\text{old}}(s, a) + \alpha (r + \gamma \max_{a'} Q_{\text{old}}(s', a'))$$



$\gamma = 0.9$

$\alpha = 0.7$

$Q_{\text{old}}(s, a)$

0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0

$Q_{\text{new}}(s, a)$

0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0

(s, a)	r	s'
$(1, \uparrow)$	0	1
$(1, \downarrow)$	0	4
$(1, \leftarrow)$	0	1
$(1, \rightarrow)$	0	2
...		
$(3, \uparrow)$	1	3
$(3, \downarrow)$	1	6
...		

$$Q_{\text{new}}(1, \uparrow) \leftarrow (1 - 0.7) Q_{\text{old}}(1, \uparrow) + 0.7 (0 + 0.9 \max_{a'} Q_{\text{old}}(1, a'))$$

$$= (1 - 0.7) * 0 + 0.7 * (0 + 0.9 * 0) = 0$$

$$Q_{\text{new}}(s, a) \leftarrow (1 - \alpha) Q_{\text{old}}(s, a) + \alpha (r + \gamma \max_{a'} Q_{\text{old}}(s', a'))$$



$\gamma = 0.9$

$\alpha = 0.7$

$Q_{\text{old}}(s, a)$

0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0

$Q_{\text{new}}(s, a)$

0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0

(s, a)	r	s'
$(1, \uparrow)$	0	1
$(1, \downarrow)$	0	4
$(1, \leftarrow)$	0	1
$(1, \rightarrow)$	0	2
...		
$(3, \uparrow)$	1	3
$(3, \downarrow)$	1	6
...		

$$Q_{\text{new}}(1, \downarrow) \leftarrow (1 - 0.7) Q_{\text{old}}(1, \downarrow) + 0.7 (0 + 0.9 \max_{a'} Q_{\text{old}}(4, a'))$$

$$= (1 - 0.7) * 0 + 0.7 * (0 + 0.9 * 0) = 0$$

$$Q_{\text{new}}(s, a) \leftarrow (1 - \alpha) Q_{\text{old}}(s, a) + \alpha (r + \gamma \max_{a'} Q_{\text{old}}(s', a'))$$



$\gamma = 0.9$

$\alpha = 0.7$

$Q_{\text{old}}(s, a)$

0	0	0	0
0	0	0	0
0	0	0	0
0	0	0	0

$Q_{\text{new}}(s, a)$

0	0	0	0
0	0	0	0
0	0	0	0
0	0	0	0

(s, a)	r	s'
$(1, \uparrow)$	0	1
$(1, \downarrow)$	0	4
$(1, \leftarrow)$	0	1
$(1, \rightarrow)$	0	2
...	...	...
$(3, \uparrow)$	1	3
$(3, \downarrow)$	1	6
...	...	...

$$Q_{\text{new}}(1, \leftarrow) \leftarrow (1 - 0.7) Q_{\text{old}}(1, \leftarrow) + 0.7 (0 + 0.9 \max_{a'} Q_{\text{old}}(1, a'))$$

$$= (1 - 0.7) * 0 + 0.7 * (0 + 0.9 * 0) = 0$$

$$Q_{\text{new}}(s, a) \leftarrow (1 - \alpha) Q_{\text{old}}(s, a) + \alpha (r + \gamma \max_{a'} Q_{\text{old}}(s', a'))$$



$\gamma = 0.9$

$\alpha = 0.7$

$Q_{\text{old}}(s, a)$

0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0

$Q_{\text{new}}(s, a)$

0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0

(s, a)	r	s'
$(1, \uparrow)$	0	1
$(1, \downarrow)$	0	4
$(1, \leftarrow)$	0	1
$(1, \rightarrow)$	0	2
...	...	...
$(3, \uparrow)$	1	3
$(3, \downarrow)$	1	6
...	...	...

$$Q_{\text{new}}(1, \rightarrow) \leftarrow (1 - 0.7) Q_{\text{old}}(1, \rightarrow) + 0.7 (0 + 0.9 \max_{a'} Q_{\text{old}}(2, a'))$$

$$= (1 - 0.7) * 0 + 0.7 * (0 + 0.9 * 0) = 0$$

$$Q_{\text{new}}(s, a) \leftarrow (1 - \alpha) Q_{\text{old}}(s, a) + \alpha (r + \gamma \max_{a'} Q_{\text{old}}(s', a'))$$



$\gamma = 0.9$

$\alpha = 0.7$

$Q_{\text{old}}(s, a)$

0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0

$Q_{\text{new}}(s, a)$

0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0

(s, a)	r	s'
$(1, \uparrow)$	0	1
$(1, \downarrow)$	0	4
$(1, \leftarrow)$	0	1
$(1, \rightarrow)$	0	2
...		
$(3, \uparrow)$	1	3
$(3, \downarrow)$	1	6
...		

$$Q_{\text{new}}(3, \uparrow) \leftarrow (1 - 0.7) Q_{\text{old}}(3, \uparrow) + 0.7 (1 + 0.9 \max_{a'} Q_{\text{old}}(3, a'))$$

$$= (1 - 0.7) * 0 + 0.7 * (1 + 0.9 * 0) = 0.7$$

$$Q_{\text{new}}(s, a) \leftarrow (1 - \alpha) Q_{\text{old}}(s, a) + \alpha (r + \gamma \max_{a'} Q_{\text{old}}(s', a'))$$



$\gamma = 0.9$

$\alpha = 0.7$

$Q_{\text{old}}(s, a)$

0	0	0.7
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0

$Q_{\text{new}}(s, a)$

0	0	0.7
0	0	0
0	0	0.7
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0

(s, a)	r	s'
$(1, \uparrow)$	0	1
$(1, \downarrow)$	0	4
$(1, \leftarrow)$	0	1
$(1, \rightarrow)$	0	2
...		
$(3, \uparrow)$	1	3
$(3, \downarrow)$	1	6
...		

$$Q_{\text{new}}(3, \downarrow) \leftarrow (1 - 0.7) Q_{\text{old}}(3, \downarrow) + 0.7 (1 + 0.9 \max_{a'} Q_{\text{old}}(6, a'))$$

$$= (1 - 0.7) * 0 + 0.7 * (1 + 0.9 * 0) = 0.7$$

$$Q_{\text{new}}(s, a) \leftarrow (1 - \alpha) Q_{\text{old}}(s, a) + \alpha (r + \gamma \max_{a'} Q_{\text{old}}(s', a'))$$



$$\gamma = 0.9$$

$$\alpha = 0.7$$

$Q_{\text{old}}(s, a)$

0	0	0.7
0	0	0
0	0	0.7
0	0	0
0	0	0
0	0	0
0	0	0
0	0	0

$Q_{\text{new}}(s, a)$

0	0	0.7
0	0	0
0	0	0.7
0	0	-6.56
0	0	0
0	0	0
0	0	0
0	0	0

(s, a)	r	s'
...	...	...
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
...	...	...

$$Q_{\text{new}}(6, \uparrow) \leftarrow (1 - 0.7) Q_{\text{old}}(6, \uparrow) + 0.7 (-10 + 0.9 \max_{a'} Q_{\text{old}}(3, a'))$$

$$= (1 - 0.7) * 0 + 0.7 * (-10 + 0.9 * 0.7) = -6.56$$

$$Q_{\text{new}}(s, a) \leftarrow (1 - \alpha) Q_{\text{old}}(s, a) + \alpha (r + \gamma \max_{a'} Q_{\text{old}}(s', a'))$$



$\gamma = 0.9$

$\alpha = 0.7$

$Q_{\text{old}}(s, a)$

				0.7
0	0	0	0	0
				0.7
0	0	0	0	-8.97
				0
0	0	0	0	0
				0
0	0	0	0	0
				0

$Q_{\text{new}}(s, a)$

				0.7
0	0	0	0	0
				0.7
0	0	0	0	-9.25
				0
0	0	0	0	0
				0
0	0	0	0	0
				0

(s, a)	r	s'
...	...	...
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
...	...	...

$$Q_{\text{new}}(6, \uparrow) \leftarrow (1 - 0.7) Q_{\text{old}}(6, \uparrow) + 0.7 (-10 + 0.9 \max_{a'} Q_{\text{old}}(3, a'))$$

$$= (1 - 0.7) * -8.97 + 0.7 * (-10 + 0.9 * 0.7) = -9.25$$



 $\alpha = 0.7$

(s, a)	r	s'
	...	
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
$(6, \uparrow)$	-10	3
$(6, \uparrow)$	-10	2
	...	

35

$$Q_{\text{new}}(s, a) \leftarrow (1 - \alpha) Q_{\text{old}}(s, a) + \alpha(r + \gamma \max_{a'} Q_{\text{old}}(s', a'))$$

- To update a particular $Q(s, a)$, we need experience from actually taking that a in that s .
- If we never try a move, its Q -value never changes.
- So if we only get one "play" quota left, which (s, a) should we try?

Depends on

- are we trying to "win" the game, or
- are we trying to get more accurate Q estimates

This is the fundamental dilemma in reinforcement learning... and in life!

Whether to *exploit* what we've already learned or *explore* to discover something better.

ϵ -greedy action selection strategy

- with probability ϵ , choose an action $a \in \mathcal{A}$ uniformly at random

During learning, especially in early stages, we'd like to explore, and observe diverse (s, a) consequences.

- with probability $(1 - \epsilon)$, choose $\arg \max_a Q_{\text{old}}(s, a)$

During later stages, can act more greedily w.r.t. the current estimated Q values

ϵ controls the trade-off between *exploration vs. exploitation*.

Value Iteration($\mathcal{S}, \mathcal{A}, T, R, \gamma, \epsilon$)

1. **for** $s \in \mathcal{S}, a \in \mathcal{A}$:
2. $Q_{\text{old}}(s, a) = 0$
3. **while** True:
4. **for** $s \in \mathcal{S}, a \in \mathcal{A}$:
5. $Q_{\text{new}}(s, a) \leftarrow R(s, a) + \gamma \sum_{s'} T(s, a, s') \max_{a'} Q_{\text{old}}(s', a')$
6. **if** $\max_{s,a} |Q_{\text{old}}(s, a) - Q_{\text{new}}(s, a)| < \epsilon$:
7. **return** Q_{new}
8. $Q_{\text{old}} \leftarrow Q_{\text{new}}$

"calculating"

Q-Learning ($\mathcal{S}, \mathcal{A}, \gamma, \alpha, s_0$ max-iter)

1. $i = 0$
2. **for** $s \in \mathcal{S}, a \in \mathcal{A}$:
3. $Q_{\text{old}}(s, a) = 0$
4. $s \leftarrow s_0$
5. **while** $i < \text{max-iter}$:
6. $a \leftarrow \text{select_action}(s, Q_{\text{old}}(s, a))$
7. $r, s' \leftarrow \text{execute}(a)$
8. $Q_{\text{new}}(s, a) \leftarrow (1 - \alpha)Q_{\text{old}}(s, a) + \alpha(r + \gamma \max_{a'} Q_{\text{old}}(s', a'))$
9. $s \leftarrow s'$
10. $i \leftarrow (i + 1)$
11. $Q_{\text{old}} \leftarrow Q_{\text{new}}$
12. **return** Q_{new}

"learning" (estimating)

Q-Learning ($\mathcal{S}, \mathcal{A}, \gamma, \alpha, s_0$ max-iter)

```
1.  $i = 0$ 
2. for  $s \in \mathcal{S}, a \in \mathcal{A}$  :
3.    $Q_{\text{old}}(s, a) = 0$ 
4.  $s \leftarrow s_0$ 
5. while  $i < \text{max-iter}$  :
6.    $a \leftarrow \text{select\_action}(s, Q_{\text{old}}(s, a))$ 
7.    $r, s' \leftarrow \text{execute}(a)$ 
8.    $Q_{\text{new}}(s, a) \leftarrow (1 - \alpha)Q_{\text{old}}(s, a) + \alpha(r + \gamma \max_{a'} Q_{\text{old}}(s', a'))$ 
9.    $s \leftarrow s'$ 
10.   $i \leftarrow (i + 1)$ 
11.   $Q_{\text{old}} \leftarrow Q_{\text{new}}$ 
12. return  $Q_{\text{new}}$ 
```

"learning"

- Remarkably,  *can* converge to the true Q_{∞}^*
- Once converged, act greedily w.r.t Q^* again.
- But convergence can be extremely slow, and often not practical:
 - We might only be interested in part of a (state, action) space
 - We might not have enough resources to visit all (state, action) pairs—even those we care about.

¹ given we visit all s, a infinitely often, and satisfy a decaying condition on the learning rate α .

Q-Learning ($\mathcal{S}, \mathcal{A}, \gamma, \alpha, s_0$ max-iter)

```
1.  $i = 0$ 
2. for  $s \in \mathcal{S}, a \in \mathcal{A}$  :
3.    $Q_{\text{old}}(s, a) = 0$ 
4.  $s \leftarrow s_0$ 
5. while  $i < \text{max-iter}$  :
6.    $a \leftarrow \text{select\_action}(s, Q_{\text{old}}(s, a))$ 
7.    $r, s' \leftarrow \text{execute}(a)$ 
8.    $Q_{\text{new}}(s, a) \leftarrow (1 - \alpha)Q_{\text{old}}(s, a) + \alpha(r + \gamma \max_{a'} Q_{\text{old}}(s', a'))$ 
9.    $s \leftarrow s'$ 
10.  $i \leftarrow (i + 1)$ 
11.  $Q_{\text{old}} \leftarrow Q_{\text{new}}$ 
12. return  $Q_{\text{new}}$ 
```

"learning"

- ϵ -greedy controls what actions to explore, versus exploiting current estimate

- three scalars in Q-learning:

- discount factor γ
- learning rate α
- greedy factor ϵ

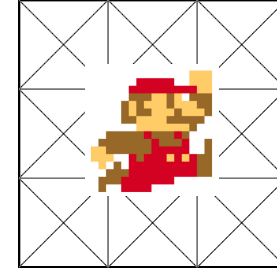
each between 0 and 1, controls some trade-off

Outline

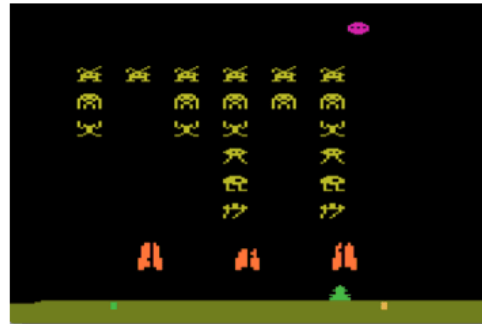
- Reinforcement learning setup
- Tabular Q-learning
 - exploration vs. exploitation
 - ϵ -greedy action selection
- Fitted Q-learning
- (Reinforcement learning setup again)

Fitted Q-learning: from table to functions

- So far, Q-learning is sensible for (small) tabular setting.
- What if $\mathcal{S}$ and/or $\mathcal{A}$ are large, or even continuous?



Game	Board size	State space
Go	19 x 19	10^{172}
Chess	8 x 8	10^{50}



10^{16992} (pixels) states



Continuous
state and
action space

- We can't keep a table — we must approximate $Q(s, a)$ with a function $Q_{\theta}(s, a)$.

- Notice that the core update rule of Q-learning:

$$Q_{\text{new}}(s, a) \leftarrow (1 - \alpha) Q_{\text{old}}(s, a) + \alpha \left(r + \gamma \max_{a'} Q_{\text{old}}(s', a') \right)$$

is equivalently:

$$Q_{\text{new}}(s, a) \leftarrow Q_{\text{old}}(s, a) + \alpha \left([r + \gamma \max_{a'} Q_{\text{old}}(s', a')] - Q_{\text{old}}(s, a) \right)$$

- Does this

$$\text{new belief} \leftarrow \text{old belief} + \text{learning rate} (\text{target} - \text{old belief})$$

remind you of something?

- Yes! Recall:

$$\theta_{\text{new}} \leftarrow \theta_{\text{old}} + \eta (\text{target} - \text{guess}_{\theta}) \nabla_{\theta} \text{guess}$$

Gradient update rule when minimizing $(\text{target} - \text{guess}_{\theta})^2$

- Same idea as before — we adjust our belief toward a target.

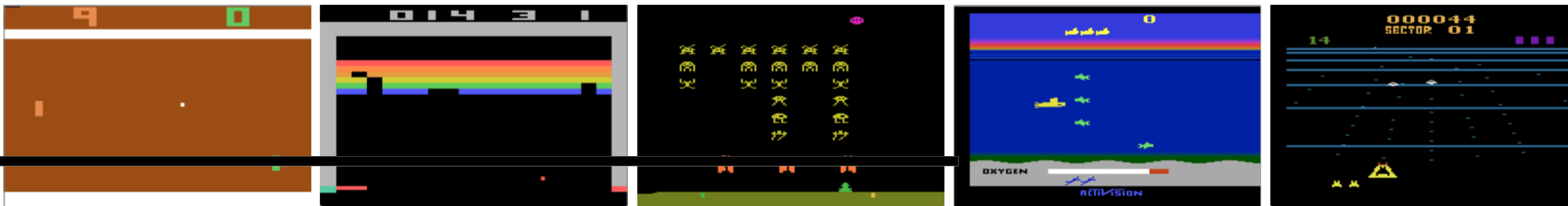
$$Q_{\text{new}}(s, a) \leftarrow Q_{\text{old}}(s, a) + \alpha ([r + \gamma \max_{a'} Q_{\text{old}}(s', a')] - Q_{\text{old}}(s, a))$$

$$\text{new belief} \leftarrow \text{old belief} + \text{learning rate} (\text{target} - \text{old belief})$$

- Now Q is a function (e.g. a neural network) and θ are its weights.
- Generalize tabular Q-learning for continuous state/action space:

1. parameterize $Q_{\theta}(s, a)$
2. execute (s, a) , observe (r, s') , construct the target $r + \gamma \max_{a'} Q_{\theta}(s', a')$
3. regress $Q_{\theta}(s, a)$ against the target, i.e. update θ via gradient-descent to minimize

$$(\text{target} - Q_{\theta}(s, a))^2$$



Playing Atari with Deep Reinforcement Learning

Volodymyr Mnih Koray Kavukcuoglu David Silver Alex Graves Ioannis Antonoglou

Daan Wierstra Martin Ried

DeepMind Technologies

{vlad,koray,david,alex.graves,ioannis,daan,mart}

Abstract

We present the first deep learning model to successfully learn to play Atari 2600 games directly from high-dimensional sensory input using reinforcement learning. The model is a convolutional neural network, trained with raw pixels as input and whose output is a value function that estimates the expected sum of discounted future rewards. We apply our method to seven Atari 2600 game environments, with no adjustment of the architecture or hyperparameters, and find that it outperforms all previous approaches on six of the seven games, achieving a human expert level on three of them.

Algorithm 1 Deep Q-learning with Experience Replay

```

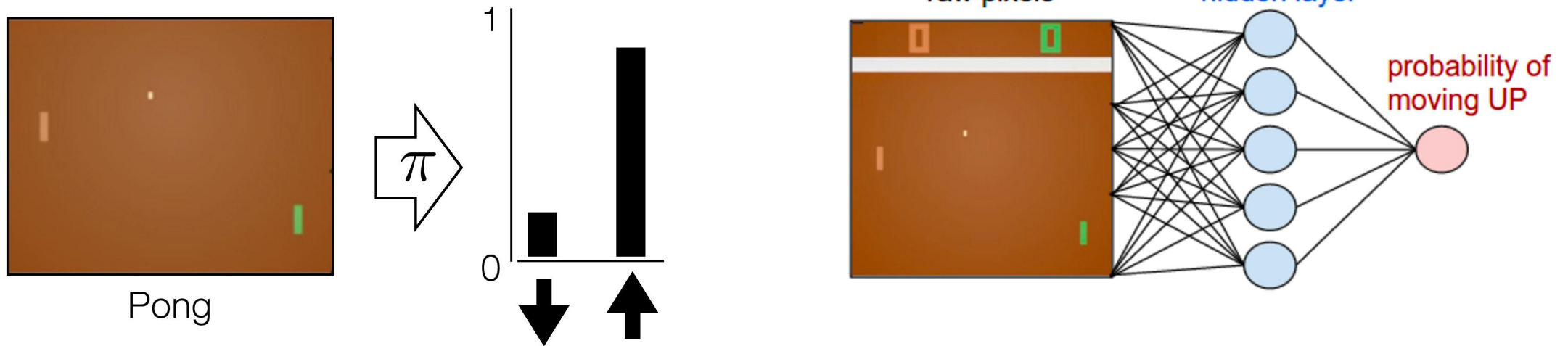
Initialize replay memory  $\mathcal{D}$  to capacity  $N$ 
Initialize action-value function  $Q$  with random weights
for episode = 1,  $M$  do
    Initialise sequence  $s_1 = \{x_1\}$  and preprocessed sequence  $\phi_1 = \phi(s_1)$ 
    for  $t = 1, T$  do
        With probability  $\epsilon$  select a random action  $a_t$ 
        otherwise select  $a_t = \max_a Q^*(\phi(s_t), a; \theta)$ 
        Execute action  $a_t$  in emulator and observe reward  $r_t$  and image  $x_{t+1}$ 
        Set  $s_{t+1} = s_t, a_t, x_{t+1}$  and preprocess  $\phi_{t+1} = \phi(s_{t+1})$ 
        Store transition  $(\phi_t, a_t, r_t, \phi_{t+1})$  in  $\mathcal{D}$ 
        Sample random minibatch of transitions  $(\phi_j, a_j, r_j, \phi_{j+1})$  from  $\mathcal{D}$ 
        Set  $y_j = \begin{cases} r_j & \text{for terminal } \phi_{j+1} \\ r_j + \gamma \max_{a'} Q(\phi_{j+1}, a'; \theta) & \text{for non-terminal } \phi_{j+1} \end{cases}$ 
        Perform a gradient descent step on  $(y_j - Q(\phi_j, a_j; \theta))^2$  according to equation 3
    end for
end for

```

Outline

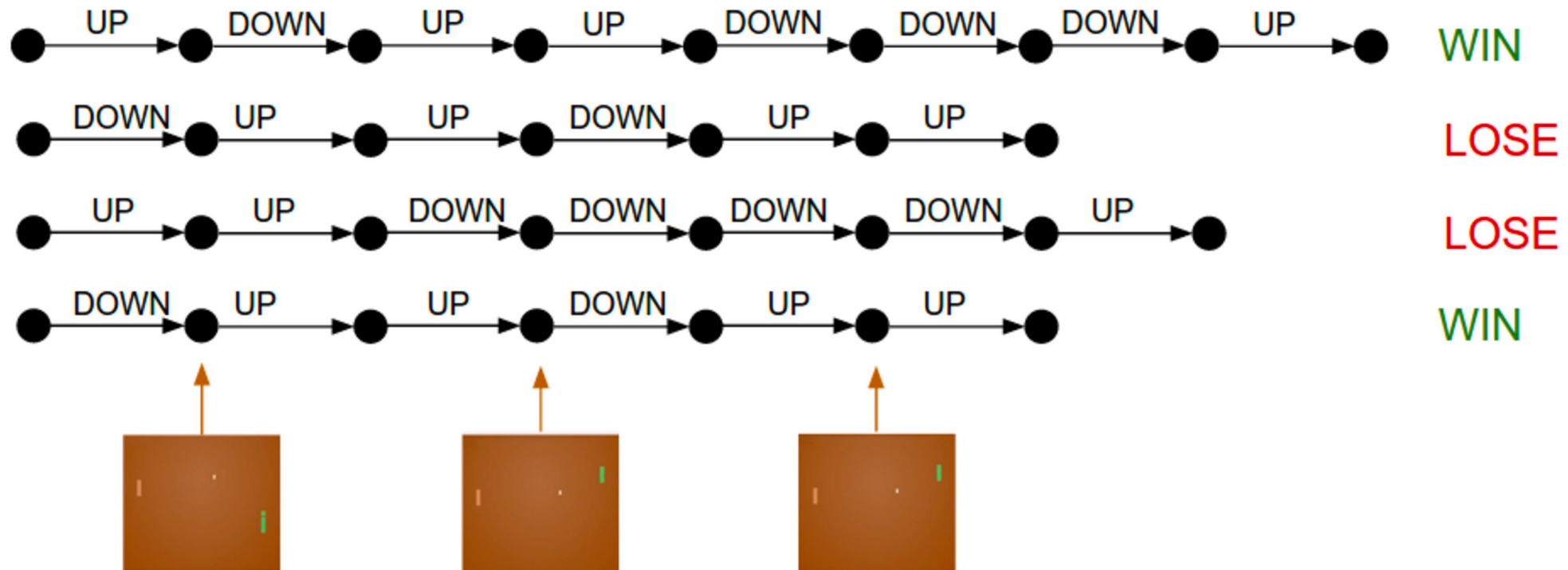
- Reinforcement learning setup
- Tabular Q-learning
 - exploration vs. exploitation
 - ϵ -greedy action selection
- Fitted Q-learning
- (Reinforcement learning setup again)

- If explicit “good” state-action pair is given, also supervised learning.
- Behavior cloning or imitation learning.
- But what if no explicit guide?



[Adapted from Andrej Karpathy: <http://karpathy.github.io/2016/05/31/rl/>]

- If no direct supervision is available?
- Strictly RL setting. Interact, observe, get data, use rewards as "coy" supervision signal.



[Adapted from Andrej Karpathy: <http://karpathy.github.io/2016/05/31/rl/>]

How Much Information is the Machine Given during Learning?

► “Pure” Reinforcement Learning (**cherry**)

- The machine predicts a scalar reward given once in a while.

► **A few bits for some samples**

► Supervised Learning (**icing**)

- The machine predicts a category or a few numbers for each input
- Predicting human-supplied data
- **10→10,000 bits per sample**

► Self-Supervised Learning (**cake génoise**)

- The machine predicts any part of its input for any observed part.
- Predicts future frames in videos
- **Millions of bits per sample**



Summary

- We saw, last week, how to find good policy in a known MDP: these are policies with high cumulative expected reward.
- In reinforcement learning, we assume we are interacting with an *unknown* MDP, but we still want to find a good policy. We will do so via estimating the Q value function.
- One problem is how to select actions to gain good reward while learning. This “exploration vs exploitation” problem is important.
- Q-learning, for discrete-state problems, will converge to the optimal value function (with enough exploration).
- “Deep Q learning” can be applied to continuous-state or large discrete-state problems by using a parameterized function to represent the Q-values.

<https://forms.gle/YLZKTx8T4h42Pf4E7>

We'd love to hear
your thoughts.

Thanks!