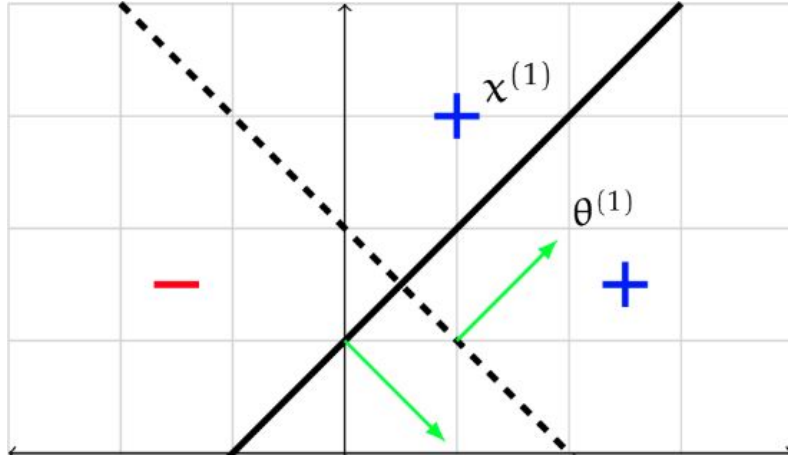
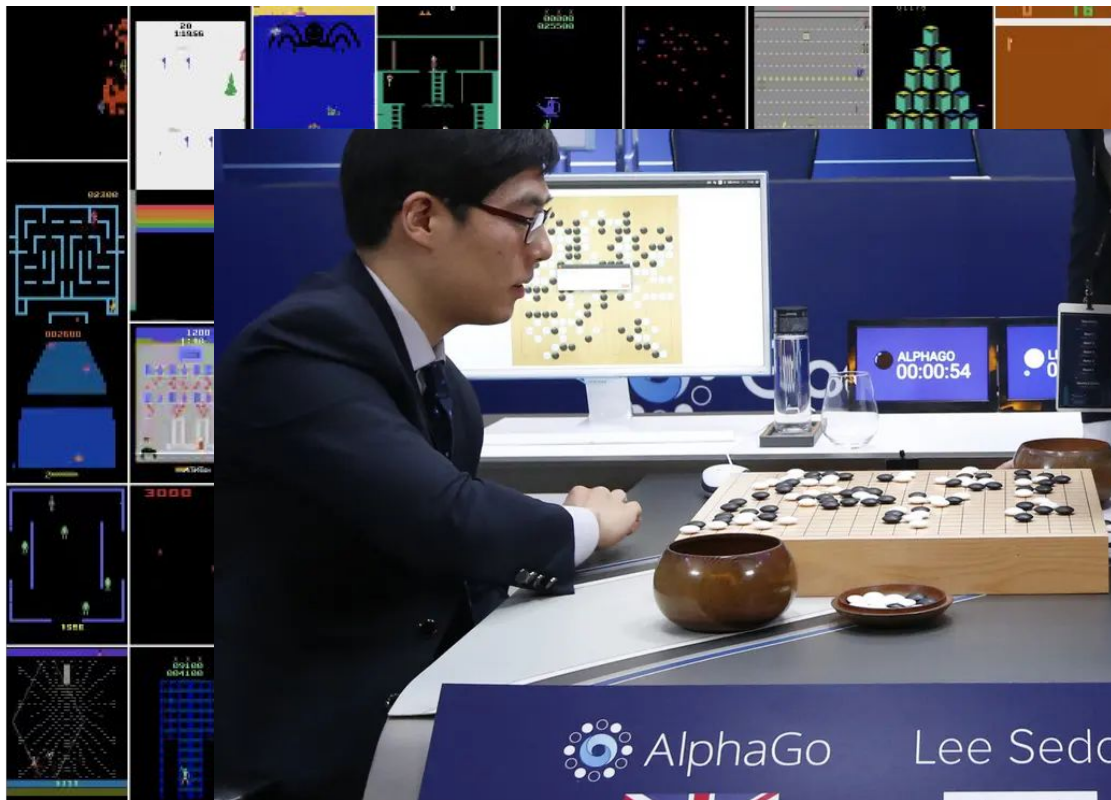


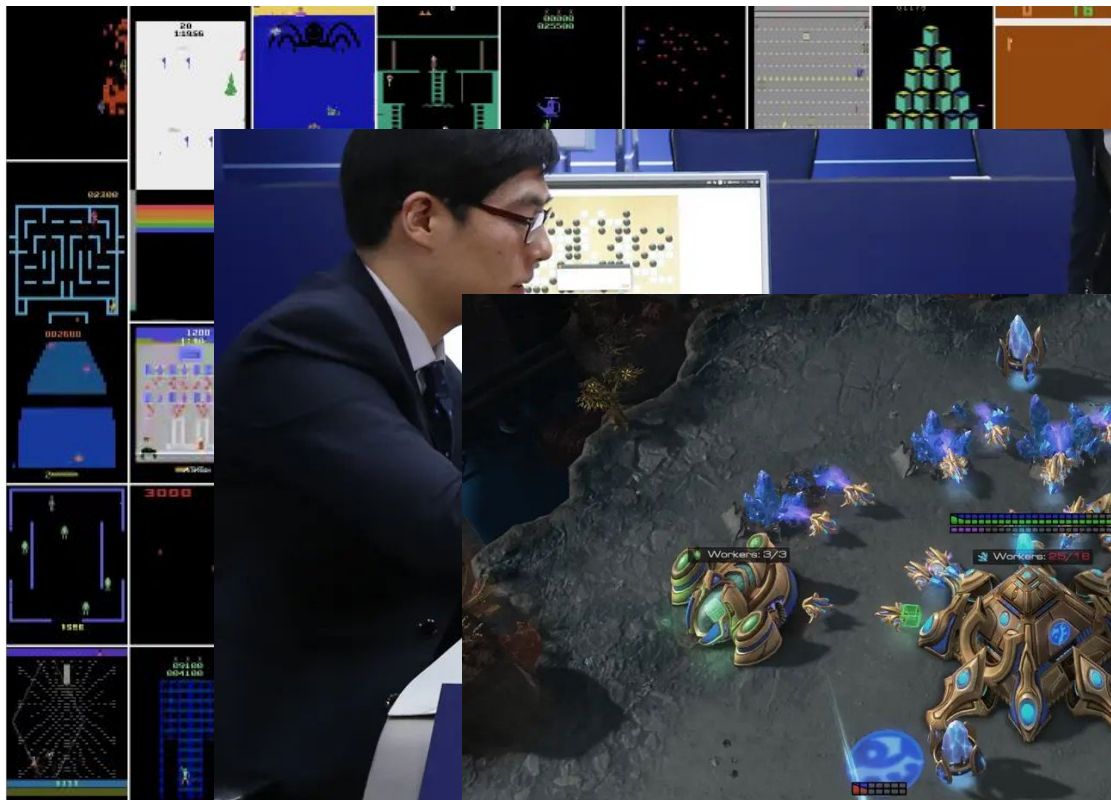
Introduction to Machine Learning



Reinforcement Learning







14:32
Catalyst LE

	AlphaStar	LiquidTLO
Score	177 / 200	147 / 172
Supply	945 +2015	335 +1595
Minerals	758 +873	442 +1030
Gas	64	61
Workers	113	86
Army	940	1377
APM	12	12
Production	12	12

Reinforcement Learning

Markov Decision Process

$$(\mathcal{S}, \mathcal{A}, \mathcal{T}, \mathcal{R}, \gamma, s_0)$$

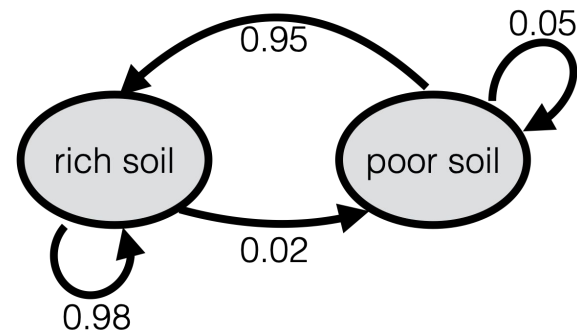
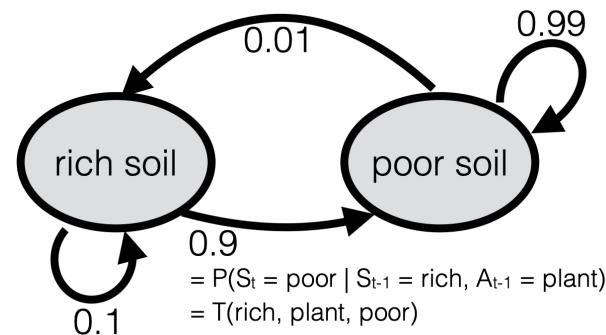
- $\mathcal{S}$ = set of possible states
- $\mathcal{A}$ = set of possible actions
- $T : \mathcal{S} \times \mathcal{A} \times \mathcal{S} \rightarrow \mathbb{R}$: transition model
- $R : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$: reward function
- γ = discount factor

Reinforcement Learning

Markov Decision Process

$$(\mathcal{S}, \mathcal{A}, \mathcal{T}, \mathcal{R}, \gamma, s_0)$$

- $\mathcal{S}$ = set of possible states
- $\mathcal{A}$ = set of possible actions
- $T : \mathcal{S} \times \mathcal{A} \times \mathcal{S} \rightarrow \mathbb{R}$: transition model
- $R : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$: reward function
- γ = discount factor

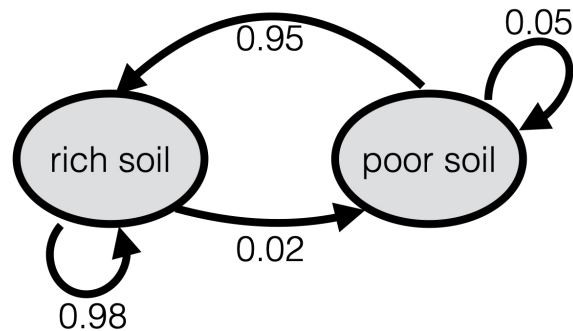
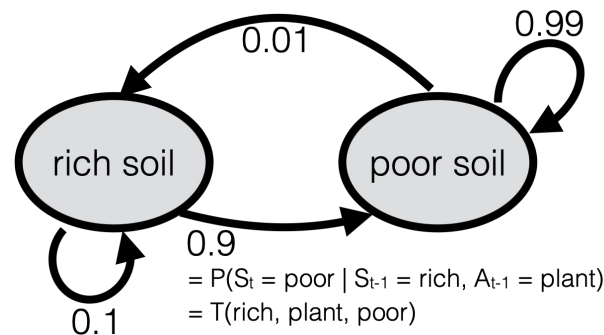


Reinforcement Learning

Markov Decision Process

$$(\mathcal{S}, \mathcal{A}, \mathcal{T}, \mathcal{R}, \gamma, s_0)$$

- $\mathcal{S}$ = set of possible states
- $\mathcal{A}$ = set of possible actions
- $T : \mathcal{S} \times \mathcal{A} \times \mathcal{S} \rightarrow \mathbb{R}$: transition model
- $R : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$: reward function
- γ = discount factor



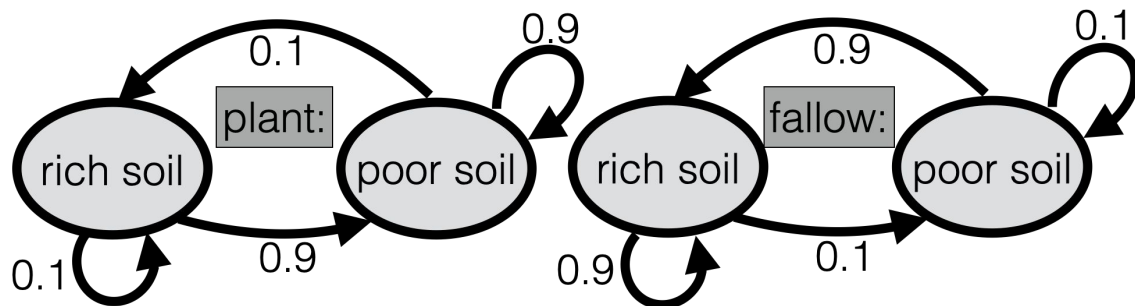
Goal in RL: find a “policy” $\pi : \mathcal{S} \rightarrow \mathcal{A}$ that maximizes reward

Review: Value of a policy

- Given an MDP and a policy $\pi : \mathcal{S} \rightarrow \mathcal{A}$ we can find the value of a policy by solving a system of linear equations.

Review: Value of a policy

- Given an MDP and a policy $\pi : \mathcal{S} \rightarrow \mathcal{A}$ we can find the value of a policy by solving a system of linear equations.



$R(\text{rich, plant})=100$
 $R(\text{poor, plant})=10$
 $R(\text{rich, fallow})=0$
 $R(\text{poor, fallow})=0$

$$V_{\pi}^0(s) = 0; V_{\pi}^h(s) = \underbrace{R(s, \pi(s))}_{\text{value of the policy on this time step}} + \underbrace{\sum_{s'} T(s, \pi(s), s') \cdot V_{\pi}^{h-1}(s')}_{\text{(expected) value of the policy across all future time steps}}$$

value of the
policy with h
steps left

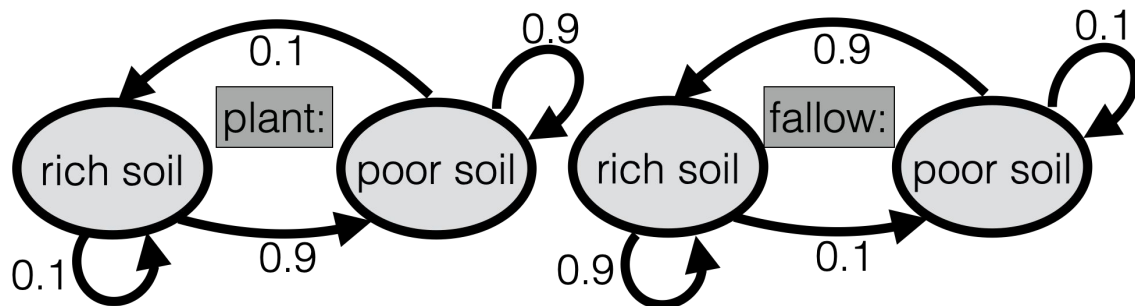
value of the
policy on this
time step

(expected) value of
the policy across
all future time steps

- h : horizon (e.g. how many growing seasons left)
- $V_{\pi}^h(s)$: value (expected reward) with policy π starting at s

Review: Value of a policy

- Given an MDP and a policy $\pi : \mathcal{S} \rightarrow \mathcal{A}$ we can find the value of a policy by solving a system of linear equations.



$R(\text{rich, plant})=100$
 $R(\text{poor, plant})=10$
 $R(\text{rich, fallow})=0$
 $R(\text{poor, fallow})=0$

$$V_{\pi}^0(s) = 0; V_{\pi}^h(s) = R(s, \pi(s)) + \sum_{s'} T(s, \pi(s), s') \cdot V_{\pi}^{h-1}(s')$$

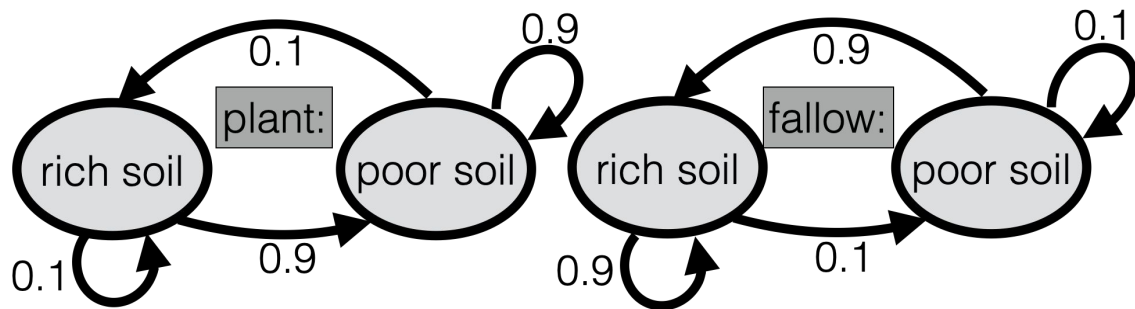
$$V_{\pi}(s) = R(s, \pi(s)) + \gamma \sum_{s'} T(s, \pi(s), s') V_{\pi}(s')$$

- h : horizon (e.g. how many growing seasons left)
- $V_{\pi}^h(s)$: value (expected reward) with policy π starting at s

Can use to evaluate which policy is better.

How to compute best policy?

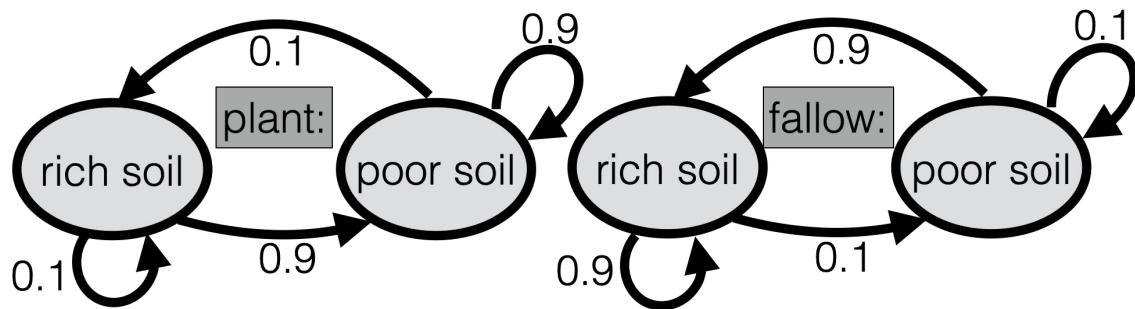
Review: Optimal policy in a known MDP



$R(\text{rich, plant})=100$
 $R(\text{poor, plant})=10$
 $R(\text{rich, fallow})=0$
 $R(\text{poor, fallow})=0$

- h : horizon (e.g. how many planting seasons)
- $Q^h(s, a)$: expected reward of starting at s , making action a , and then making the “best” action for the $h-1$ steps left
- With Q , can find **an optimal policy**: $\pi_h^*(s) = \arg \max_a Q^h(s, a)$

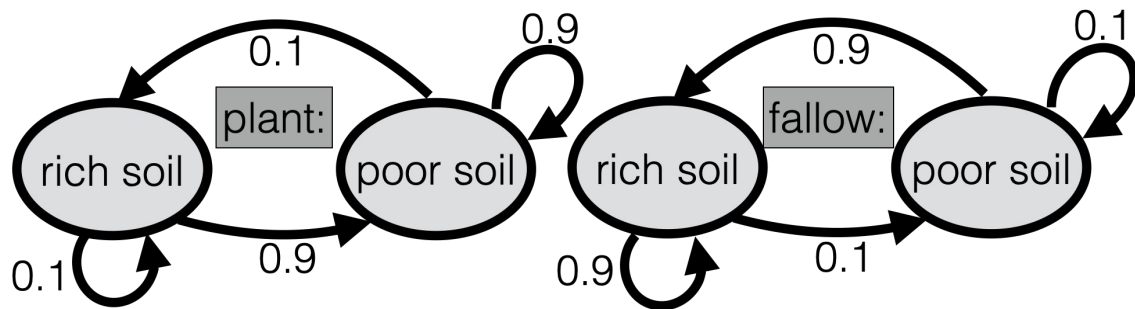
Review: Optimal policy in a known MDP



$R(\text{rich, plant})=100$
 $R(\text{poor, plant})=10$
 $R(\text{rich, fallow})=0$
 $R(\text{poor, fallow})=0$

- h : horizon (e.g. how many planting seasons)
 - $Q^h(s, a)$: expected reward of starting at s , making action a , and then making the “best” action for the $h-1$ steps left
 - With Q , can find **an optimal policy**: $\pi_h^*(s) = \arg \max_a Q^h(s, a)$
- $$Q^0(s, a) = 0; Q^h(s, a) = R(s, a) + \sum_{s'} T(s, a, s') \max_{a'} Q^{h-1}(s', a')$$

Review: Optimal policy in a known MDP



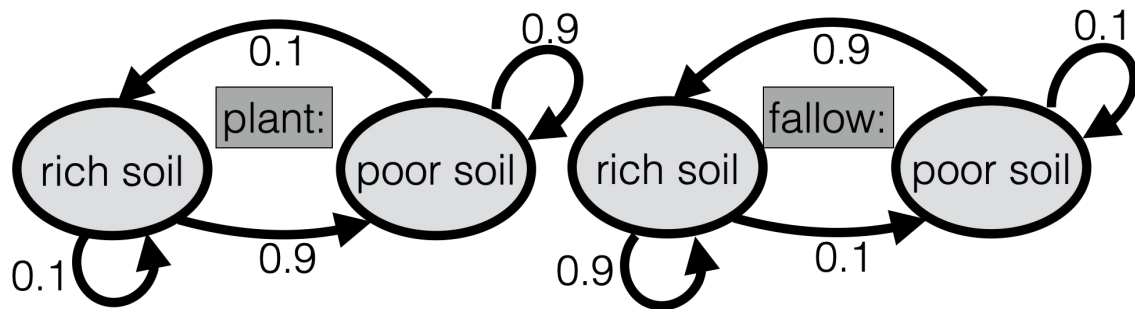
$R(\text{rich, plant})=100$
 $R(\text{poor, plant})=10$
 $R(\text{rich, fallow})=0$
 $R(\text{poor, fallow})=0$

- h : horizon (e.g. how many planting seasons)
- $Q^h(s, a)$: expected reward of starting at s , making action a , and then making the “best” action for the $h-1$ steps left
- With Q , can find **an optimal policy**: $\pi_h^*(s) = \arg \max_a Q^h(s, a)$

$$Q^0(s, a) = 0; Q^h(s, a) = R(s, a) + \sum_{s'} T(s, a, s') \max_{a'} Q^{h-1}(s', a')$$

$$Q^1(\text{rich, plant}) = 100; Q^1(\text{rich, fallow}) = 0; Q^1(\text{poor, plant}) = 10; Q^1(\text{poor, fallow}) = 0$$

Review: Optimal policy in a known MDP



$R(\text{rich, plant})=100$
 $R(\text{poor, plant})=10$
 $R(\text{rich, fallow})=0$
 $R(\text{poor, fallow})=0$

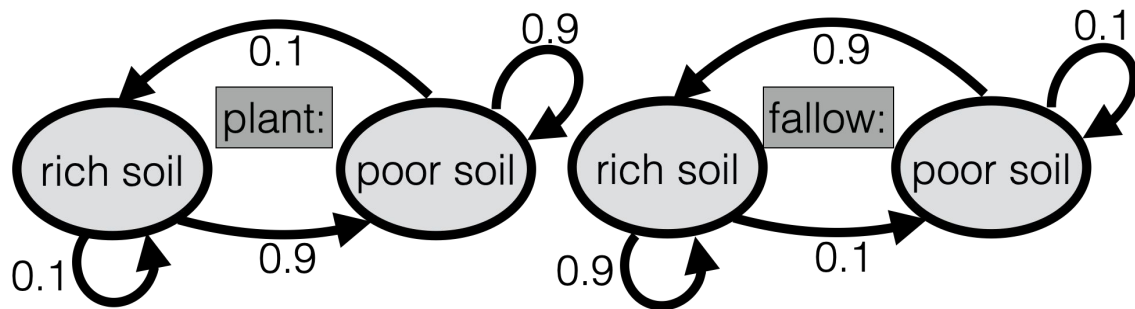
- h : horizon (e.g. how many planting seasons)
- $Q^h(s, a)$: expected reward of starting at s , making action a , and then making the “best” action for the $h-1$ steps left
- With Q , can find **an optimal policy**: $\pi_h^*(s) = \arg \max_a Q^h(s, a)$

$$Q^0(s, a) = 0; Q^h(s, a) = R(s, a) + \sum_{s'} T(s, a, s') \max_{a'} Q^{h-1}(s', a')$$

$$Q^1(\text{rich, plant}) = 100; Q^1(\text{rich, fallow}) = 0; Q^1(\text{poor, plant}) = 10; Q^1(\text{poor, fallow}) = 0$$

$$Q^2(\text{rich, plant}) = R(\text{rich, plant}) + T(\text{rich, plant, rich}) \max_{a'} Q^1(\text{rich, } a') + T(\text{rich, plant, poor}) \max_{a'} Q^1(\text{poor, } a')$$

Review: Optimal policy in a known MDP



$R(\text{rich, plant})=100$
 $R(\text{poor, plant})=10$
 $R(\text{rich, fallow})=0$
 $R(\text{poor, fallow})=0$

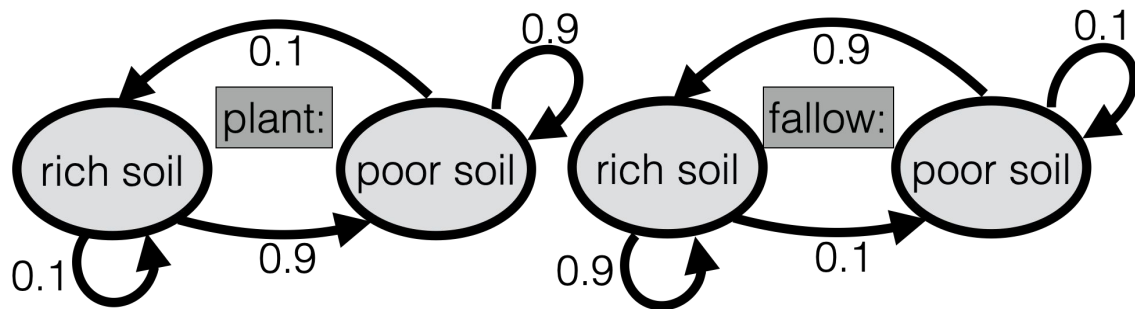
- h : horizon (e.g. how many planting seasons)
- $Q^h(s, a)$: expected reward of starting at s , making action a , and then making the “best” action for the $h-1$ steps left
- With Q , can find **an optimal policy**: $\pi_h^*(s) = \arg \max_a Q^h(s, a)$

$$Q^0(s, a) = 0; Q^h(s, a) = R(s, a) + \sum_{s'} T(s, a, s') \max_{a'} Q^{h-1}(s', a')$$

$$Q^1(\text{rich, plant}) = 100; Q^1(\text{rich, fallow}) = 0; Q^1(\text{poor, plant}) = 10; Q^1(\text{poor, fallow}) = 0$$

$$Q^2(\text{rich, plant}) = 100 + (0.1)(100) + (0.9)(10) = 119$$

Review: Optimal policy in a known MDP



$R(\text{rich, plant})=100$
 $R(\text{poor, plant})=10$
 $R(\text{rich, fallow})=0$
 $R(\text{poor, fallow})=0$

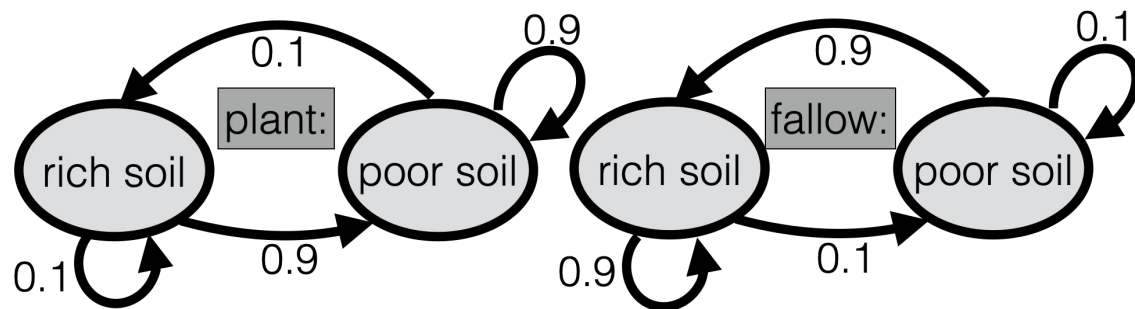
- h : horizon (e.g. how many planting seasons)
- $Q^h(s, a)$: expected reward of starting at s , making action a , and then making the “best” action for the $h-1$ steps left
- With Q , can find **an optimal policy**: $\pi_h^*(s) = \arg \max_a Q^h(s, a)$

$$Q^0(s, a) = 0; Q^h(s, a) = R(s, a) + \sum_{s'} T(s, a, s') \max_{a'} Q^{h-1}(s', a')$$

$$Q^1(\text{rich, plant}) = 100; Q^1(\text{rich, fallow}) = 0; Q^1(\text{poor, plant}) = 10; Q^1(\text{poor, fallow}) = 0$$

$$Q^2(\text{rich, plant}) = 119; Q^2(\text{rich, fallow}) = 91; Q^2(\text{poor, plant}) = 29; Q^2(\text{poor, fallow}) = 91$$

Review: (Finite-Horizon) Value Iteration



$R(\text{rich, plant})=100$
 $R(\text{poor, plant})=10$
 $R(\text{rich, fallow})=0$
 $R(\text{poor, fallow})=0$

- h : horizon (e.g. how many planting seasons)
- $Q^h(s, a)$: expected reward of starting at s , making action a , and then making the “best” action for the $h-1$ steps left
- With Q , can find **an optimal policy**: $\pi_h^*(s) = \arg \max_a Q^h(s, a)$

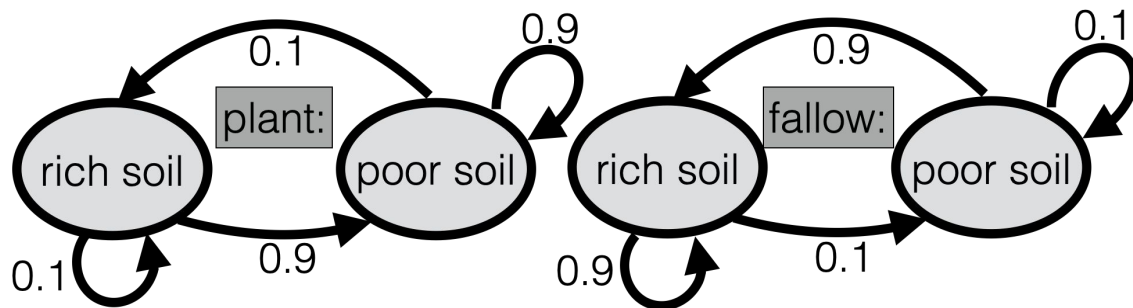
$$Q^0(s, a) = 0; Q^h(s, a) = R(s, a) + \sum_{s'} T(s, a, s') \max_{a'} Q^{h-1}(s', a')$$

$$Q^1(\text{rich, plant}) = 100; Q^1(\text{rich, fallow}) = 0; Q^1(\text{poor, plant}) = 10; Q^1(\text{poor, fallow}) = 0$$

$$Q^2(\text{rich, plant}) = 119; Q^2(\text{rich, fallow}) = 91; Q^2(\text{poor, plant}) = 29; Q^2(\text{poor, fallow}) = 91$$

What's best? Any s , $\pi_1^*(s) = \text{plant}$; $\pi_2^*(\text{rich}) = \text{plant}$, $\pi_2^*(\text{poor}) = \text{fallow}$

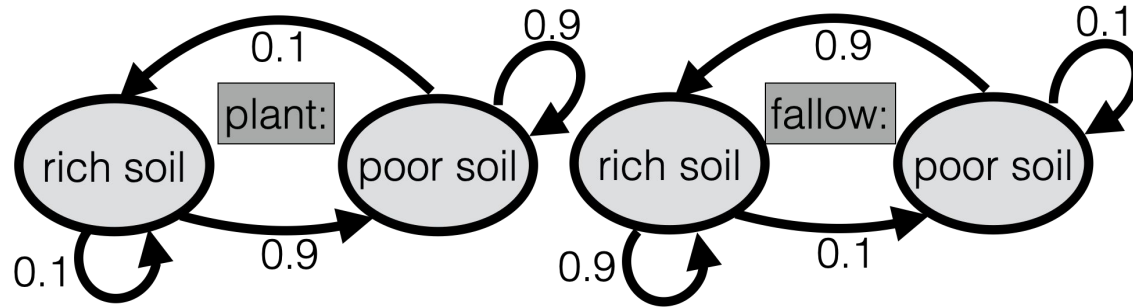
Review: (Infinite-Horizon) Value Iteration



$R(\text{rich, plant})=100$
 $R(\text{poor, plant})=10$
 $R(\text{rich, fallow})=0$
 $R(\text{poor, fallow})=0$

- What if I don't stop farming? Is there any optimal policy?
- **Theorem.** There exists a (stationary) optimal policy π^* . I.e., for every policy π and for every state $s \in \mathcal{S}$, $V_{\pi^*}(s) \geq V_{\pi}(s)$

Review: (Infinite-Horizon) Value Iteration



$R(\text{rich, plant})=100$
 $R(\text{poor, plant})=10$
 $R(\text{rich, fallow})=0$
 $R(\text{poor, fallow})=0$

- What if I don't stop farming? Is there any optimal policy?
- **Theorem.** There exists a (stationary) optimal policy π^* . I.e., for every policy π and for every state $s \in \mathcal{S}$, $V_{\pi^*}(s) \geq V_{\pi}(s)$
- $Q^*(s, a)$: expected reward if we make best actions in future
 - If we knew $Q^*(s, a)$, then: $\pi^*(s) = \arg \max_a Q^*(s, a)$
- Note: $Q^*(s, a) = R(s, a) + \gamma \sum_{s'} T(s, a, s') \max_{a'} Q^*(s', a')$
 - Not linear in $Q^*(s, a)$, so not as easy to solve as $V_{\pi}(s)$

Review: (Infinite-Horizon) Value Iteration

Finite-horizon value iteration:

$$Q^0(s, a) = 0 \quad Q^1(s, a) = R(s, a)$$

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

Review: (Infinite-Horizon) Value Iteration

Finite-horizon value iteration:

$$Q^0(s, a) = 0 \quad Q^1(s, a) = R(s, a)$$

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

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

for each state $s \in \mathcal{S}$ and each action $a \in \mathcal{A}$

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

Review: (Infinite-Horizon) Value Iteration

Finite-horizon value iteration:

$$Q^0(s, a) = 0 \quad Q^1(s, a) = R(s, a)$$

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

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

for each state $s \in \mathcal{S}$ and each action $a \in \mathcal{A}$

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

while True

for each state $s \in \mathcal{S}$ and each action $a \in \mathcal{A}$

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

Review: (Infinite-Horizon) Value Iteration

Finite-horizon value iteration:

$$Q^0(s, a) = 0 \quad Q^1(s, a) = R(s, a)$$

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

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

for each state $s \in \mathcal{S}$ and each action $a \in \mathcal{A}$

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

while True

for each state $s \in \mathcal{S}$ and each action $a \in \mathcal{A}$

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

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

return Q_{new}

$$Q_{\text{old}} = Q_{\text{new}}$$

Review: (Infinite-Horizon) Value Iteration

Finite-horizon value iteration:

$$Q^0(s, a) = 0 \quad Q^1(s, a) = R(s, a)$$

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

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

for each state $s \in \mathcal{S}$ and each action $a \in \mathcal{A}$

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

while True

for each state $s \in \mathcal{S}$ and each action $a \in \mathcal{A}$

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

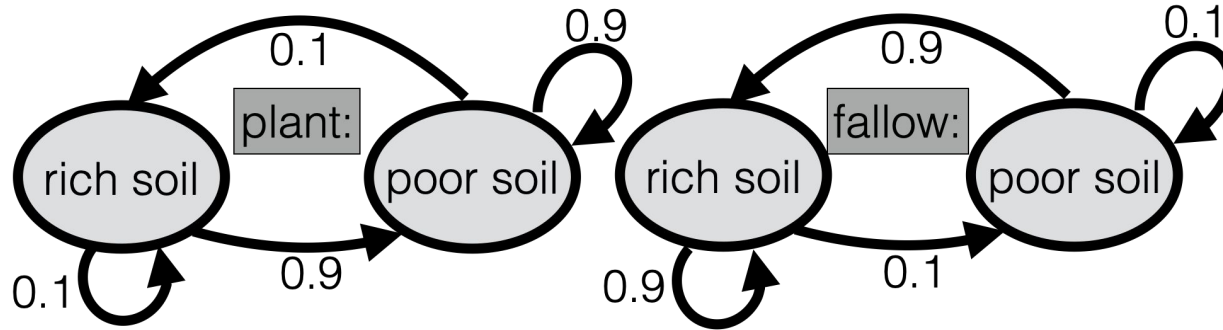
if $\max_{s,a} |Q_{\text{old}}(s, a) - Q_{\text{new}}(s, a)| < \epsilon$

return Q_{new}

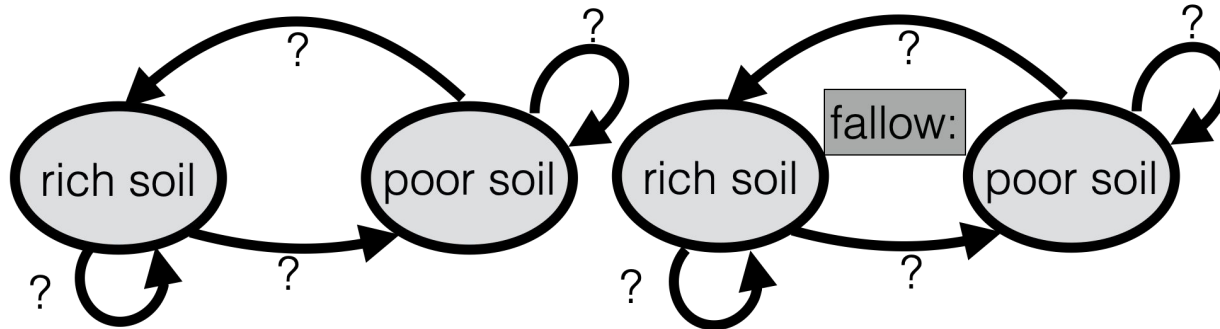
$Q_{\text{old}} = Q_{\text{new}}$

Issue: need to know reward and transition functions!

A more realistic scenario



$R(\text{rich, plant})=100$
 $R(\text{poor, plant})=10$
 $R(\text{rich, fallow})=0$
 $R(\text{poor, fallow})=0$



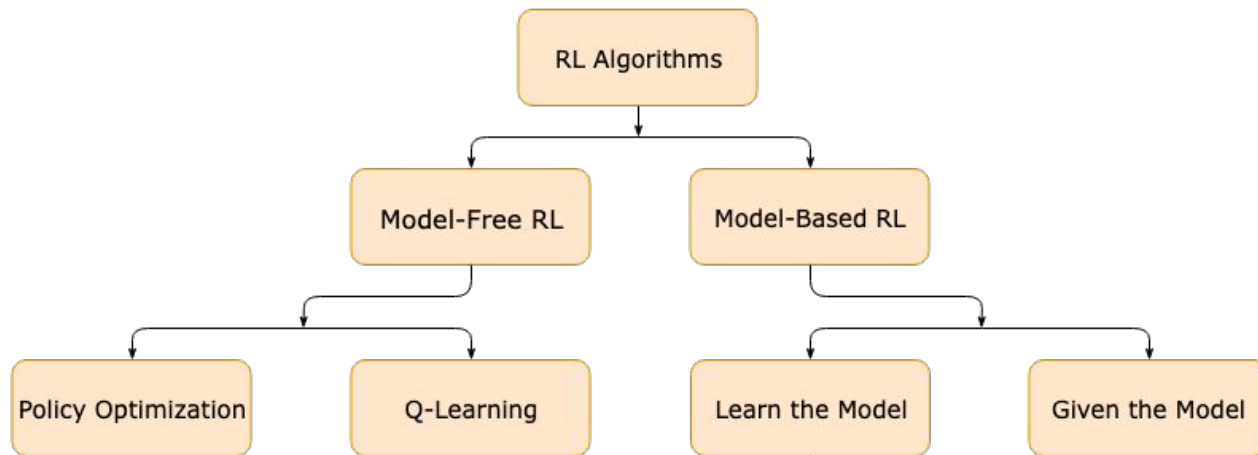
$R(\text{rich, plant})=?$
 $R(\text{poor, plant})=?$
 $R(\text{rich, fallow})=?$
 $R(\text{poor, fallow})=?$

Reinforcement Learning Overview

Goal in RL: find a “policy” $\pi : \mathcal{S} \rightarrow \mathcal{A}$ that maximizes reward in an **unknown** MDP.

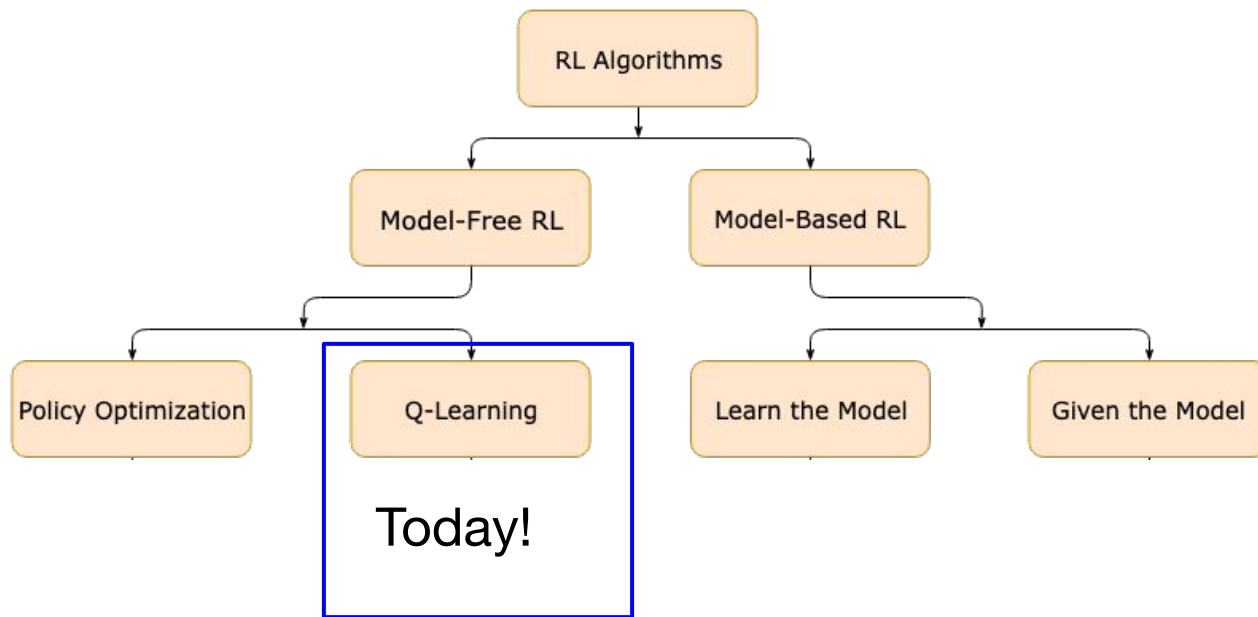
Reinforcement Learning Overview

Goal in RL: find a “policy” $\pi : \mathcal{S} \rightarrow \mathcal{A}$ that maximizes reward in an **unknown** MDP.



Reinforcement Learning Overview

Goal in RL: find a “policy” $\pi : \mathcal{S} \rightarrow \mathcal{A}$ that maximizes reward in an **unknown** MDP.



Q-Learning

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

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

$Q[s, a] = 0$

Initialize Q function to 0

Q-Learning

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

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

$Q[s, a] = 0$

Initialize Q function to 0

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

while True:

$a = \text{select_action}(s, Q)$

“Act” in the environment

Q-Learning

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

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

$Q[s, a] = 0$

Initialize Q function to 0

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

while True:

$a = \text{select_action}(s, Q)$

“Act” in the environment

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

Receive reward, transition to s'

Intuition: sampling to estimate the transition model

$$T : \mathcal{S} \times \mathcal{A} \times \mathcal{S} \rightarrow \mathbb{R}$$

Q-Learning

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

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

$Q[s, a] = 0$

Initialize Q function to 0

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

while True:

$a = \text{select_action}(s, Q)$

“Act” in the environment

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

Receive reward, transition to s'

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

Update Q function for the
sampled state and action

Q-Learning

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

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

$Q[s, a] = 0$

Initialize Q function to 0

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

while True:

$a = \text{select_action}(s, Q)$

“Act” in the environment

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

Receive reward, transition to s'

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

$s = s'$

Move onto next state

Update Q function for the
sampled state and action

What changed?

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

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

$Q[s, a] = 0$

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

while True:

$a = \text{select_action}(s, Q)$

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

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

$s = s'$

while True

for each state $s \in \mathcal{S}$ and each action $a \in \mathcal{A}$

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

What changed?

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

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

$Q[s, a] = 0$

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

while True:

$a = \text{select_action}(s, Q)$

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

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

$s = s'$

while True

for each state $s \in \mathcal{S}$ and each action $a \in \mathcal{A}$

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

Since we don't know R and T , we estimate it by “sampling” in the environment

What changed?

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

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

$Q[s, a] = 0$

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

while True:

$a = \text{select_action}(s, Q)$

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

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

$s = s'$

Moving average

What changed?

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

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

$Q[s, a] = 0$

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

while True:

$a = \text{select_action}(s, Q)$

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

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

$s = s'$

Moving average

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

Looks like gradient descent!

Q-Learning

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

Linear regression with just the bias term

$$y = \theta_0$$

$$L(\theta_0) = \frac{1}{2}(y - \theta_0)^2$$

$$\theta_0 = \theta_0 - \eta(\theta_0 - y)$$

Q-Learning

$$Q[s, a] = Q[s, a] - \alpha \left(\text{Current guess} - \text{Target} \right)$$

Linear regression with just the bias term

$$y = \theta_0$$

$$L(\theta_0) = \frac{1}{2}(y - \theta_0)^2$$

$$\theta_0 = \theta_0 - \eta(\theta_0 - y)$$

Current guess

Target

Q-Learning

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

Linear regression with just the bias term

$$y = \theta_0$$

$$L(\theta_0) = \frac{1}{2}(y - \theta_0)^2$$

$$\theta_0 = \theta_0 - \eta(\theta_0 - y)$$

Learning rate

Current guess

Target

Q-Learning

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

Linear regression with just the bias term

$$y = \theta_0$$

$$L(\theta_0) = \frac{1}{2}(y - \theta_0)^2$$

$$\theta_0 = \theta_0 - \eta(\theta_0 - y)$$

Learning rate

Current guess

Target

Key difference: in Q-learning, the target itself is a function of what is being learned (and the reward)!

Q-Learning

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

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

$Q[s, a] = 0$

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

while True:

$a = \text{select_action}(s, Q)$

How to select an action?

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

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

$s = s'$

Q-Learning

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

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

$Q[s, a] = 0$

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

while True:

$a = \text{select_action}(s, Q)$

How to select an action?

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

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

$s = s'$

Argmax action from current policy?

$\arg \max_{a \in \mathcal{A}} Q(s, a)$

Initially Q function is bad $\Rightarrow$ argmax not a good idea.

Q-Learning

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

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

$Q[s, a] = 0$

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

while True:

$a = \text{select_action}(s, Q)$

How to select an action?

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

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

$s = s'$

ϵ -greedy strategy:

- with probability $1 - \epsilon$, choose $\arg \max_{a \in \mathcal{A}} Q(s, a)$
- with probability ϵ , choose the action $a \in \mathcal{A}$ uniformly at random

Exploitation

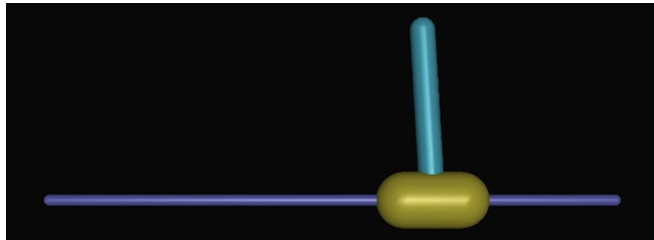
Exploration

Q-Learning

- $Q[s,a]$ is a scalar for each possible state and action. (“Tabular” Q-learning)
- What if states are high dimensional or continuous?



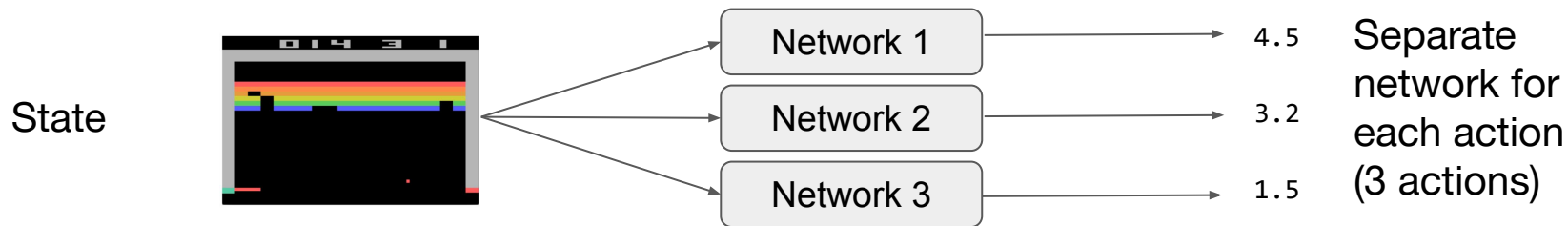
- What if actions are not discrete?



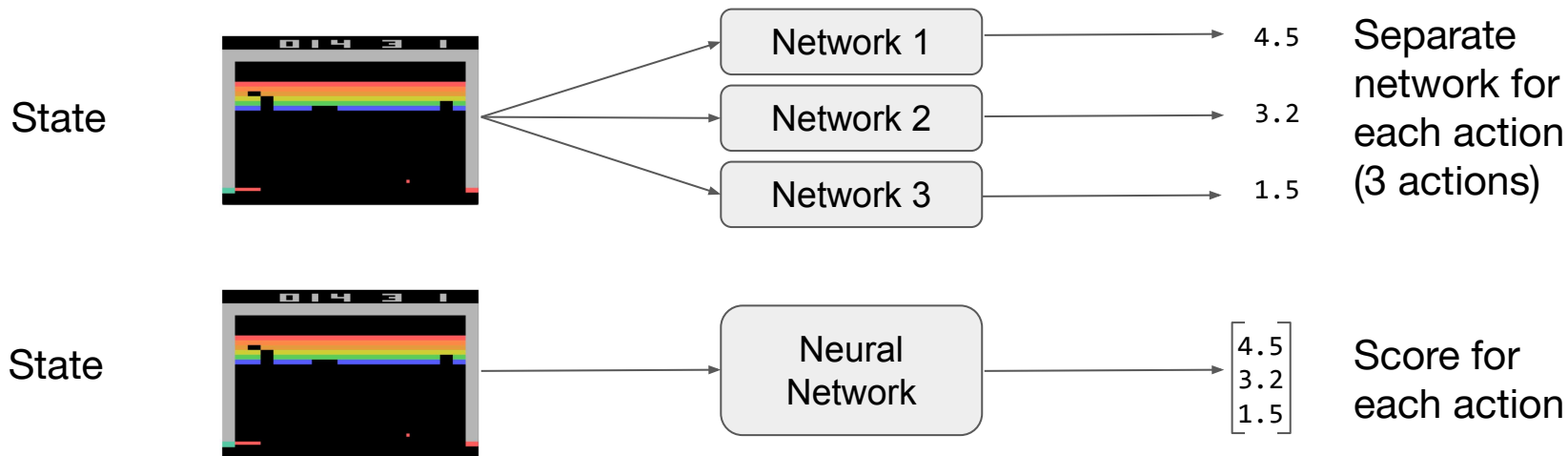
DQN: Deep Q-Networks

- Parameterize Q function to be an output from a neural network!

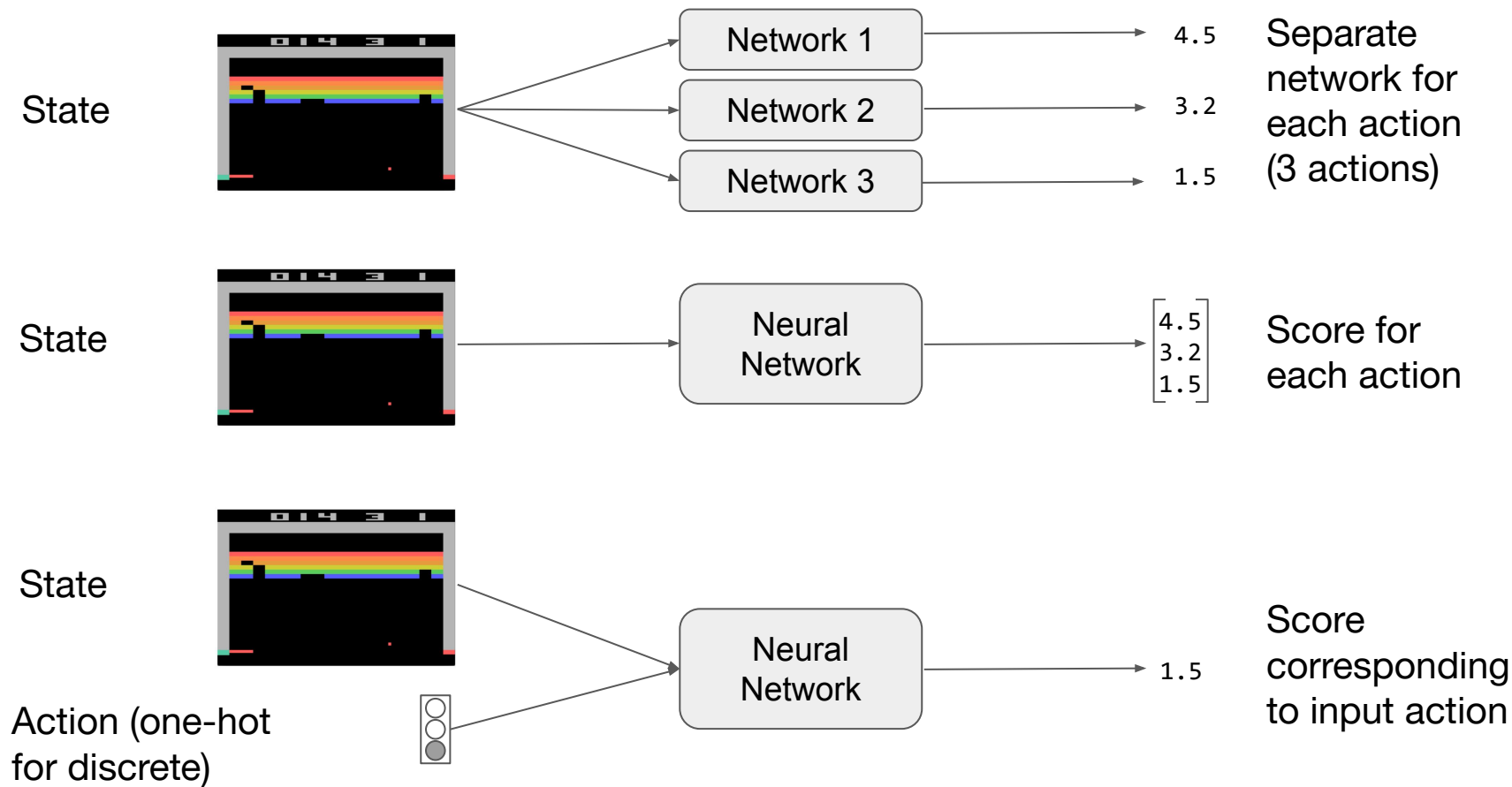
DQN: Deep Q-Networks



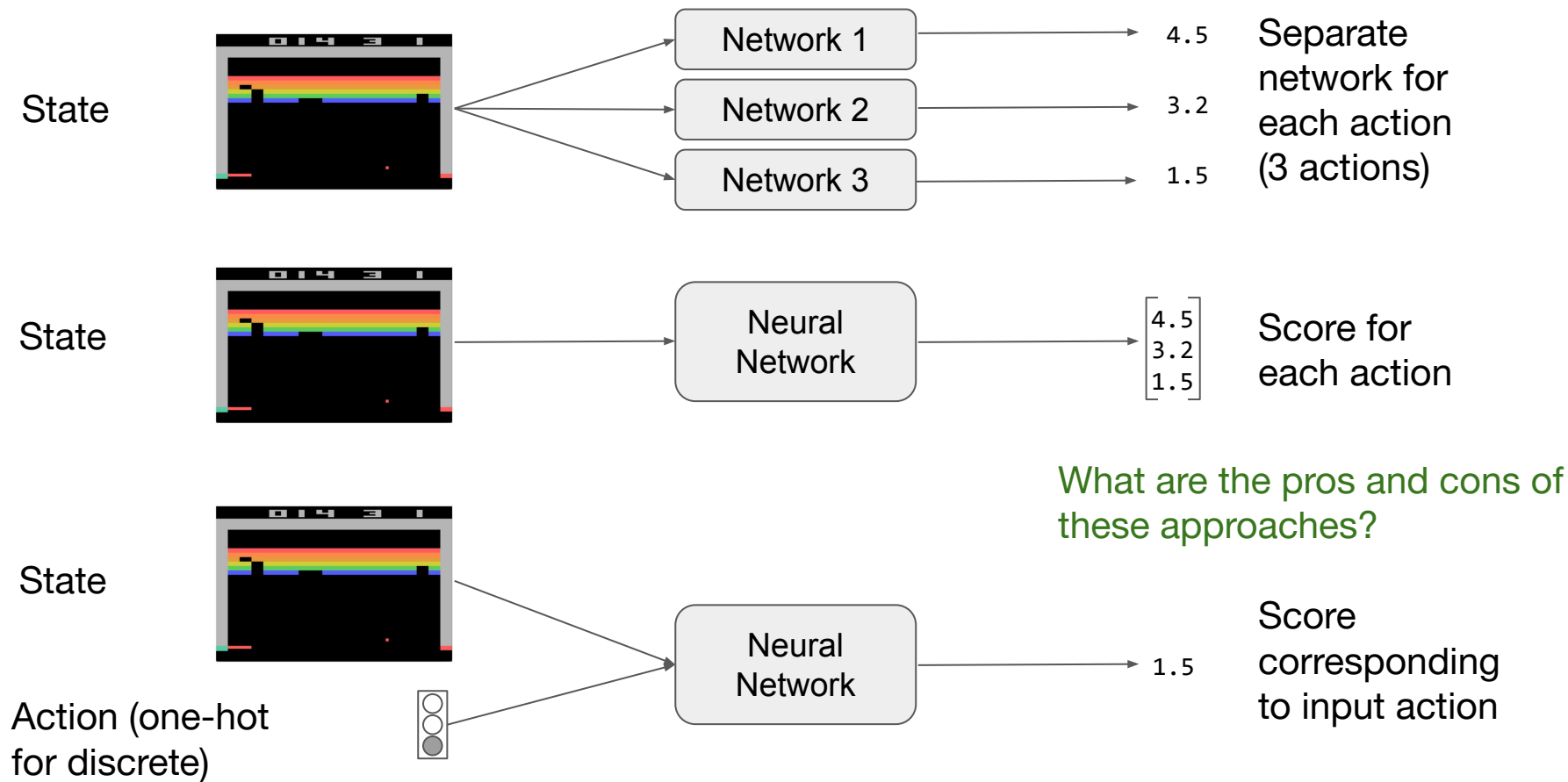
DQN: Deep Q-Networks



DQN: Deep Q-Networks



DQN: Deep Q-Networks



DQN: Deep Q-Networks

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

```
1  for  $s \in \mathcal{S}, a \in \mathcal{A}$  :  
2       $Q[s, a] = 0$   
3   $s = s_0$  // Or draw an  $s$  randomly from  $\mathcal{S}$   
4  while True:  
5       $a = \text{select\_action}(s, Q)$   
6       $r, s' = \text{execute}(a)$   
7       $Q[s, a] = (1 - \alpha)Q[s, a] + \alpha(r + \gamma \max_{a'} Q[s', a'])$   
8       $s = s'$ 
```

DQN: Deep Q-Networks

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

```
1  for  $s \in \mathcal{S}, a \in \mathcal{A}$  :  
2       $Q[s, a] = 0$   $\longrightarrow$  Randomly initialize  $Q_\theta(s, a)$   
3   $s = s_0$  // Or draw an  $s$  randomly from  $\mathcal{S}$   
4  while True:  
5       $a = \text{select\_action}(s, Q)$   
6       $r, s' = \text{execute}(a)$   
7       $Q[s, a] = (1 - \alpha)Q[s, a] + \alpha(r + \gamma \max_{a'} Q[s', a'])$   
8       $s = s'$ 
```

DQN: Deep Q-Networks

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

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

2 $Q[s, a] = 0$

→ Randomly initialize $Q_\theta(s, a)$

3 $s = s_0$ // Or draw an s randomly from $\mathcal{S}$

4 **while** True:

5 $a = \text{select_action}(s, Q)$

6 $r, s' = \text{execute}(a)$

7 $Q[s, a] = (1 - \alpha)Q[s, a] + \alpha(r + \gamma \max_{a'} Q[s', a'])$

8 $s = s'$

→ Minimize regression loss

$$(Q_\theta(s, a) - (r + \gamma \max_{a'} Q_\theta(s', a')))^2$$

DQN: Deep Q-Networks

Q-LEARNING($\mathcal{S}, \mathcal{A}, s_0, \gamma, \alpha$)

```
1  for  $s \in \mathcal{S}, a \in \mathcal{A}$  :  
2       $Q[s, a] = 0$  → Randomly initialize  $Q_\theta(s, a)$   
3   $s = s_0$  // Or draw an  $s$  randomly from  $\mathcal{S}$   
4  while True:  
5       $a = \text{select\_action}(s, Q)$   
6       $r, s' = \text{execute}(a)$   
7       $Q[s, a] = (1 - \alpha)Q[s, a] + \alpha(r + \gamma \max_{a'} Q[s', a'])$   
8       $s = s'$  → Minimize regression loss
```

Issue: instability arising from both guess and target being from a learned network.

Minimize regression loss

$$(Q_\theta(s, a) - (r + \gamma \max_{a'} Q_\theta(s', a')))^2$$

DQN: Fitted Q-Learning

FITTED-Q-LEARNING($\mathcal{A}, s_0, \gamma, \alpha, \epsilon, m$)

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

$\mathcal{D} = \{ \}$

initialize neural-network representation of Q

while True:

$\mathcal{D}_{\text{new}}$ = experience from executing ϵ -greedy policy based on Q for m steps

$\mathcal{D} = \mathcal{D} \cup \mathcal{D}_{\text{new}}$ represented as (s, a, r, s') tuples

$\mathcal{D}_{\text{sup}} = \{(\mathbf{x}^{(i)}, \mathbf{y}^{(i)})\}$ where $\mathbf{x}^{(i)} = (s, a)$ and $\mathbf{y}^{(i)} = r + \gamma \max_{a' \in \mathcal{A}} Q(s', a')$

for each tuple $(s, a, r, s')^{(i)} \in \mathcal{D}$

re-initialize neural-network representation of Q

$Q = \text{supervised_NN_regression}(\mathcal{D}_{\text{sup}})$

DQN: Fitted Q-Learning

FITTED-Q-LEARNING($\mathcal{A}, s_0, \gamma, \alpha, \epsilon, m$)

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

$\mathcal{D} = \{ \}$

initialize neural-network representation of Q

while True:

Collect data from current policy

$\mathcal{D}_{\text{new}}$ = experience from executing ϵ -greedy policy based on Q for m steps

$\mathcal{D} = \mathcal{D} \cup \mathcal{D}_{\text{new}}$ represented as (s, a, r, s') tuples

$\mathcal{D}_{\text{sup}} = \{(\mathbf{x}^{(i)}, y^{(i)})\}$ where $\mathbf{x}^{(i)} = (s, a)$ and $y^{(i)} = r + \gamma \max_{a' \in \mathcal{A}} Q(s', a')$

for each tuple $(s, a, r, s')^{(i)} \in \mathcal{D}$

re-initialize neural-network representation of Q

$Q = \text{supervised_NN_regression}(\mathcal{D}_{\text{sup}})$

DQN: Fitted Q-Learning

FITTED-Q-LEARNING($\mathcal{A}, s_0, \gamma, \alpha, \epsilon, m$)

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

$\mathcal{D} = \{ \}$

initialize neural-network representation of Q

while True:

Collect data from current policy

$\mathcal{D}_{\text{new}}$ = experience from executing ϵ -greedy policy based on Q for m steps

$\mathcal{D} = \mathcal{D} \cup \mathcal{D}_{\text{new}}$ represented as (s, a, r, s') tuples

$\mathcal{D}_{\text{sup}} = \{(\mathbf{x}^{(i)}, \mathbf{y}^{(i)})\}$ where $\mathbf{x}^{(i)} = (s, a)$ and $\mathbf{y}^{(i)} = r + \gamma \max_{a' \in \mathcal{A}} Q(s', a')$
for each tuple $(s, a, r, s')^{(i)} \in \mathcal{D}$

re-initialize neural-network representation of Q

Create supervised
learning dataset

$Q = \text{supervised_NN_regression}(\mathcal{D}_{\text{sup}})$

DQN: Fitted Q-Learning

FITTED-Q-LEARNING($\mathcal{A}, s_0, \gamma, \alpha, \epsilon, m$)

$s = s_0$ // Or draw an s randomly from $\mathcal{S}$

$\mathcal{D} = \{ \}$

initialize neural-network representation of Q

while True:

Collect data from current policy

$\mathcal{D}_{\text{new}}$ = experience from executing ϵ -greedy policy based on Q for m steps

$\mathcal{D} = \mathcal{D} \cup \mathcal{D}_{\text{new}}$ represented as (s, a, r, s') tuples

$\mathcal{D}_{\text{sup}} = \{(\mathbf{x}^{(i)}, y^{(i)})\}$ where $\mathbf{x}^{(i)} = (s, a)$ and $y^{(i)} = r + \gamma \max_{a' \in \mathcal{A}} Q(s', a')$
for each tuple $(s, a, r, s')^{(i)} \in \mathcal{D}$

re-initialize neural-network representation of Q

$Q = \text{supervised_NN_regression}(\mathcal{D}_{\text{sup}})$

Create supervised
learning dataset

Train a neural network with regression loss!