

Approximate Value Iteration and Variants

Chen-Yu Wei

Value Iteration

For $k = 1, 2, \dots$

$$\forall s, a, \quad Q_k(s, a) \leftarrow \underbrace{R(s, a)}_{\text{unknown}} + \gamma \sum_{s'} \underbrace{P(s'|s, a)}_{\text{unknown}} \max_{a'} Q_{k-1}(s', a')$$

Idea: In each iteration, use multiple samples to estimate the right-hand side.

Value Iteration with Samples

$$s_1, a_1, r_1, s_2, a_2, r_2, \dots, s_H, a_H, r_H, s_{H+1}$$

For $k = 1, 2, \dots$

$$(s_1, a_1, r_1, s_2) (s_2, a_2, r_2, s_3), \dots (s_H, a_H, r_H, s_{H+1}) (s' \sim P, \dots)$$

Obtain N samples $\{(s_i, a_i, r_i, s'_i)\}_{i=1}^N$ where $\mathbb{E}[r_i] = R(s_i, a_i)$, $s'_i \sim P(\cdot | s_i, a_i)$

Perform **regression** on $\{(s_i, a_i, r_i, s'_i)\}_{i=1}^N$ to find Q_k such that

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

Perform one iteration of Value Iteration

$$Q_\theta(s, a)$$

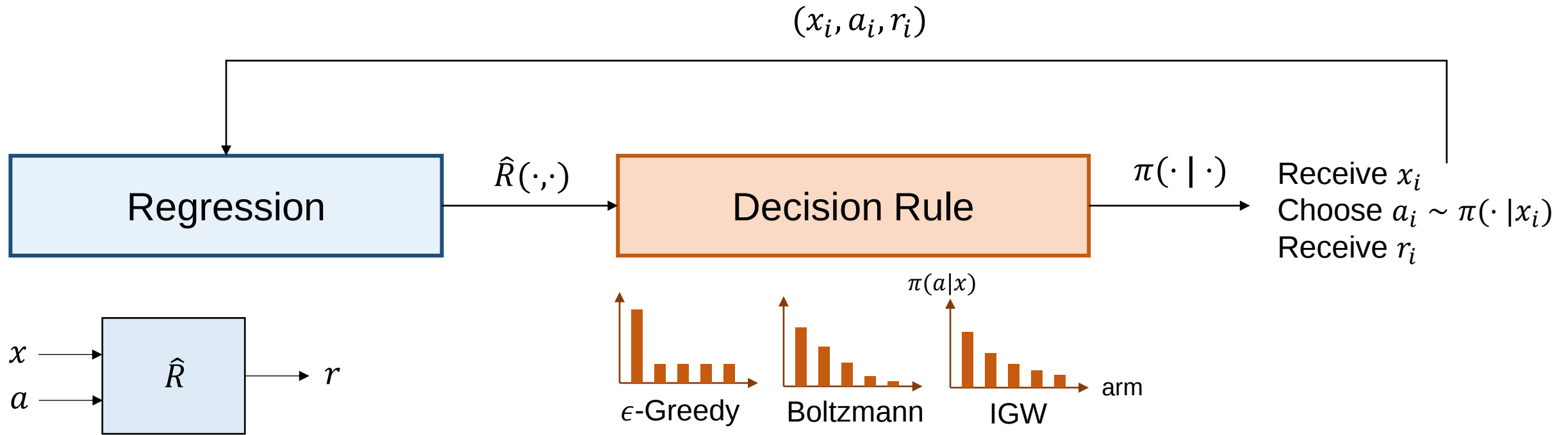
↑
parameter

Find θ_k that minimize

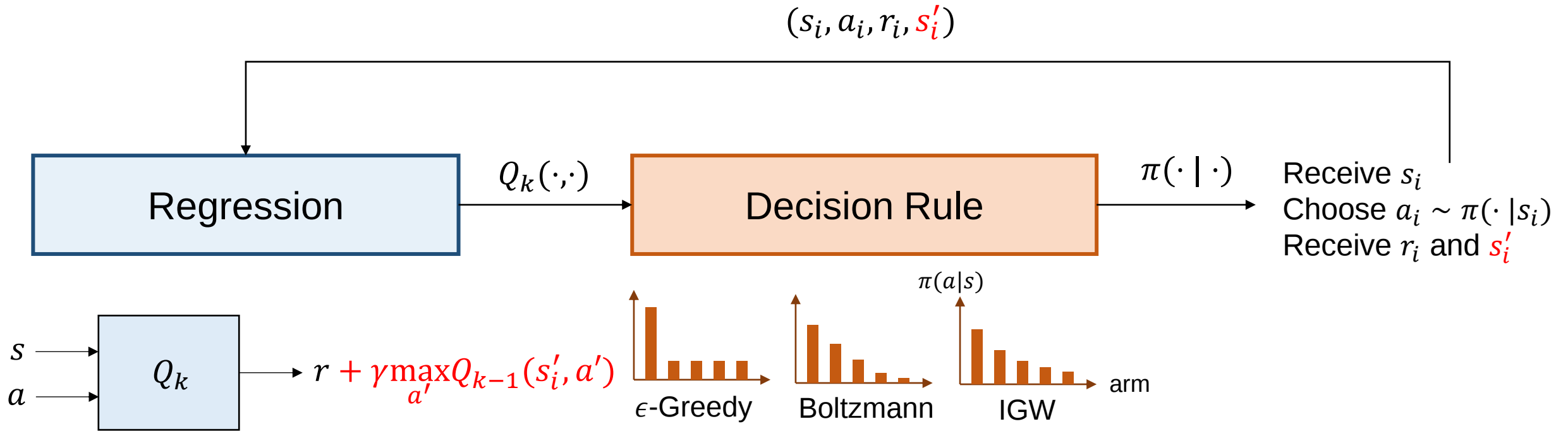
$$\theta_k = \arg \min_{\theta} \sum_{i=1}^N \left(\underbrace{Q_\theta(s_i, a_i)}_{\text{blue underline}} - \underbrace{\left(r_i + \gamma \max_{a'} Q_{\theta_{k-1}}(s'_i, a') \right)}_{\text{blue underline}} \right)^2$$

$$\mathbb{E}[\cdot | s_i, a_i] = R(s_i, a_i) + \gamma \sum_{s'} P(s' | s_i, a_i) \max_{a'} Q_{\theta_{k-1}}(s', a')$$

Recall: Contextual Bandits with Regression



Value Iteration with Regression



Train Q_k such that $Q_k(s_i, a_i) \approx r_i + \gamma \max_{a'} Q_{k-1}(s'_i, a')$

This is just one iteration of Value Iteration

Value Iteration with Samples

For $k = 1, 2, \dots$

For $i = 1, 2, \dots, N$:

Choose action $a_i \sim \text{EG}(Q_{\theta_k}(s_i, \cdot))$

Receive reward $r_i \sim R(s_i, a_i)$ and $s'_i \sim P(\cdot | s_i, a_i)$

$s_{i+1} = s'_i$ if episode continues, $s_{i+1} \sim \rho$ if episode ends

$\theta \leftarrow \theta_k$

For $m = 1, 2, \dots, M$:

Randomly pick an i (or a mini-batch) from $\{1, 2, \dots, N\}$

$\theta \leftarrow \theta - \alpha \nabla_{\theta} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\theta_k}(s'_i, a') \right)^2$

$\theta_{k+1} \leftarrow \theta$

Data collection

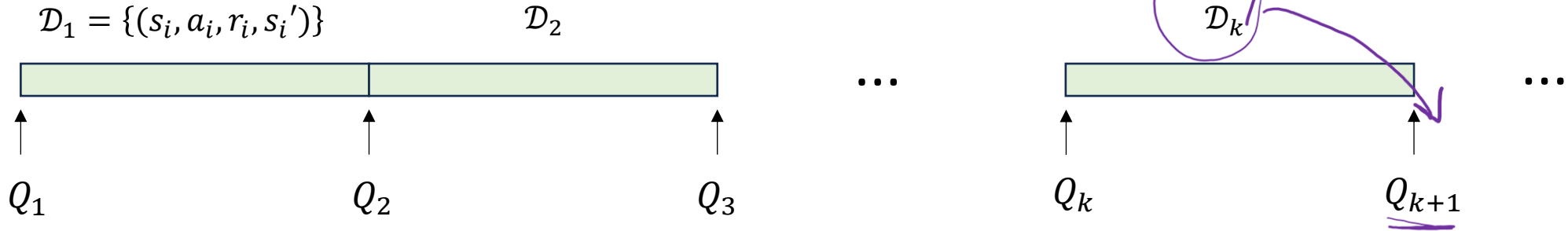
Perform one iteration
of Value Iteration

↑
Target network

2nd for-loop: trying to find $\theta_{k+1} = \underset{\theta}{\operatorname{argmin}} \sum_{i=1}^N \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\theta_k}(s'_i, a') \right)^2$

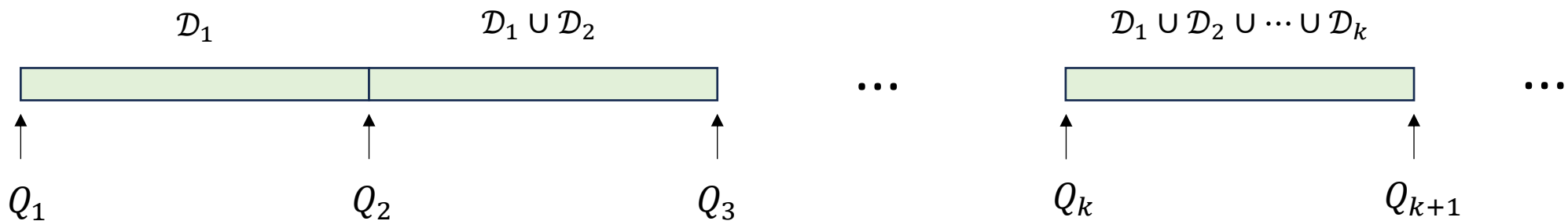
It is Valid to Reuse Samples

(e.g., using ϵ -greedy)



The algorithm in the previous slide only use $\mathcal{D}_k$ to train θ_{k+1} .

However, as the reward function R and transition P remains unchanged, it is valid (actually, even better) to reuse samples:



Benefits of Reusing Samples

- Improving data efficiency
 - Every sample is used multiple times in training – just like we usually go through multiple epochs for supervised learning tasks.
- The buffer $\mathcal{B}$ will consist of a wider range of state-actions
 - It allows better approximation of

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

Value Iteration with Reused Samples (= Deep Q-Learning or DQN)

Initialize $\mathcal{B} = \{\}$ $\leftarrow$ Replay buffer

For $k = 1, 2, \dots$

For $i = 1, 2, \dots, N$:

Choose action $a_i \sim \text{EG}(Q_{\theta_k}(s_i, \cdot))$

Receive reward $r_i \sim R(s_i, a_i)$ and $s'_i \sim P(\cdot | s_i, a_i)$

$s_{i+1} = s'_i$ if episode continues, $s_{i+1} \sim \rho$ if episode ends

Insert (s_i, a_i, r_i, s'_i) to $\mathcal{B}$

$\theta \leftarrow \theta_k$

For $m = 1, 2, \dots, M$:

Randomly pick an i (or a mini-batch) from $\mathcal{B}$

$$\theta \leftarrow \theta - \alpha \nabla_{\theta} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\theta_k}(s'_i, a') \right)^2$$

$\theta_{k+1} \leftarrow \theta$

HW4 task

Data collection

Perform one iteration
of Value Iteration

$\uparrow$
Target network

Another Popular Implementation

HW4 task

Initialize $\mathcal{B} = \{\}$ $\leftarrow$ Replay buffer

For $k = 1, 2, \dots$

For $i = 1, 2, \dots, N$:

Choose action $a_i \sim \text{EG}(Q_\theta(s_i, \cdot))$

Receive reward $r_i \sim R(s_i, a_i)$ and $s'_i \sim P(\cdot | s_i, a_i)$

$s_{i+1} = s'_i$ if episode continues, $s_{i+1} \sim \rho$ if episode ends

Insert (s_i, a_i, r_i, s'_i) to $\mathcal{B}$

For $m = 1, 2, \dots, M$:

Randomly pick an i (or a mini-batch) from $\mathcal{B}$

$$\theta \leftarrow \theta - \nabla_\theta \left(Q_\theta(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\bar{\theta}}(s'_i, a') \right)^2$$

$$\bar{\theta} \leftarrow (1 - \tau)\bar{\theta} + \tau\theta$$

$\uparrow$
Target network

When Does DQN Succeed?

DQN tries to approximate **Value Iteration** by solving

$$\theta_{k+1} = \operatorname{argmin}_{\theta} \sum_{i \in \mathcal{B}} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\theta_k}(s'_i, a') \right)^2 \quad (1)$$

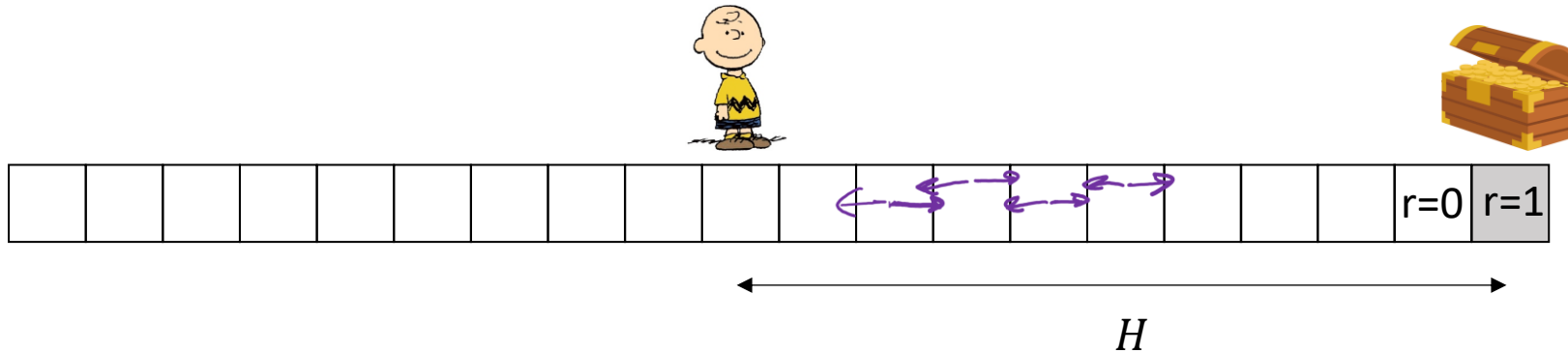
The true Value Iteration:

$$\forall \mathbf{s}, \mathbf{a}, \quad Q_{k+1}(s, a) = R(s, a) + \gamma \sum_{s'} P(s'|s, a) \max_{a'} Q_k(s', a') \quad (2)$$

Under what conditions can (1) well approximate (2)?

- $\mathcal{B}$ should contain a wide range of state-action pairs (a challenge of **exploration**)
- $Q_{\theta_{k+1}}(s, a)$ should recover $R(s, a) + \gamma \sum_{s'} P(s'|s, a) \max_{a'} Q_{\theta_k}(s', a')$ well for all state-actions (a challenge of **function approximation**, or **generalization**)

1. Exploration in MDPs (Not Easy)



Environment:

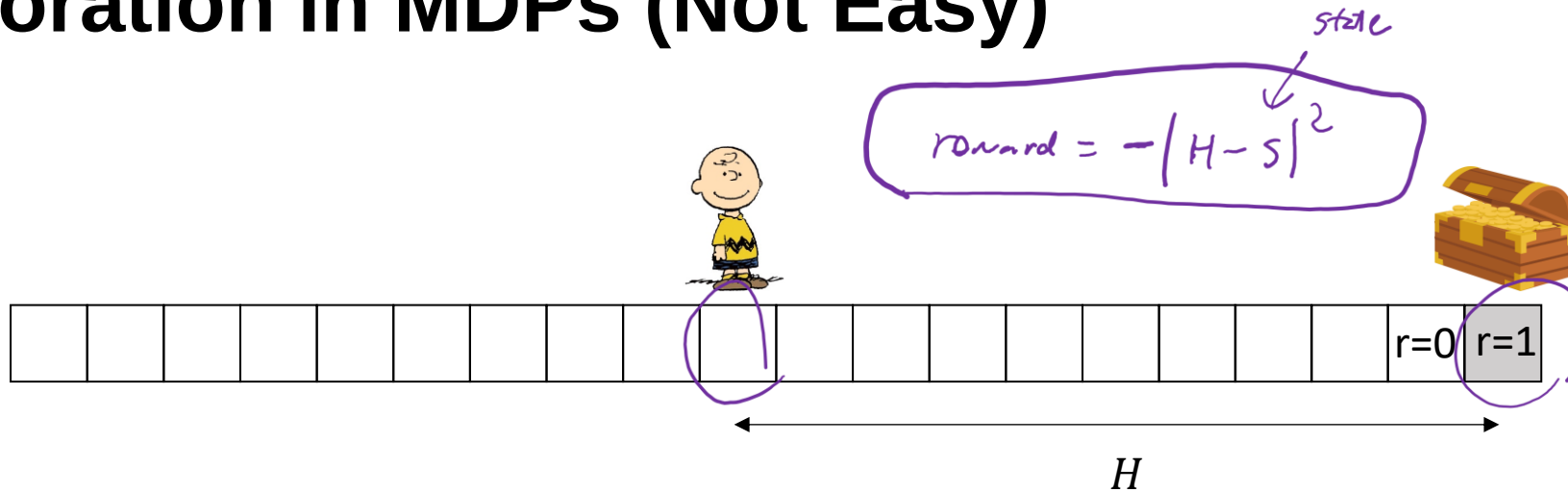
- Fixed-horizon MDP with episode length H
- Initial state at 0
- A single rewarding state at state H
- Actions: Go LEFT or RIGHT

Suppose we perform DQN with ϵ -greedy with random initialization

⇒ On average, we need 2^H episodes to see the reward

(before that, we won't make any meaningful update and will just do random walk around state 0)

1. Exploration in MDPs (Not Easy)



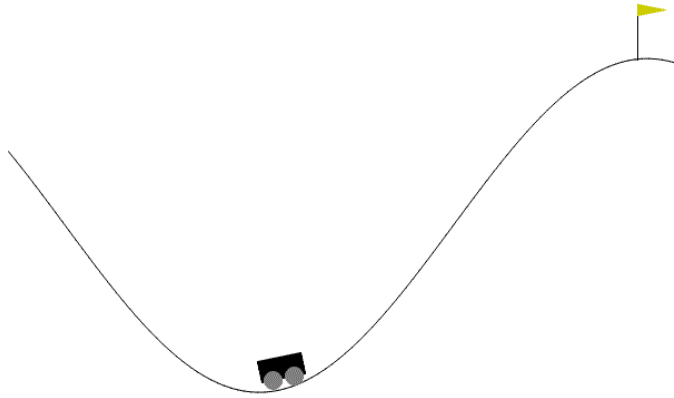
Key issue:

- The ϵ -greedy strategy (or BE, IGW) performs **action-space** exploration but not **state-space** exploration.
- This problem becomes more severe when the reward signal is **sparse**.
- To solve this, we usually require the **exploration bonus** (a form of reward shaping) technique – will be covered much later.

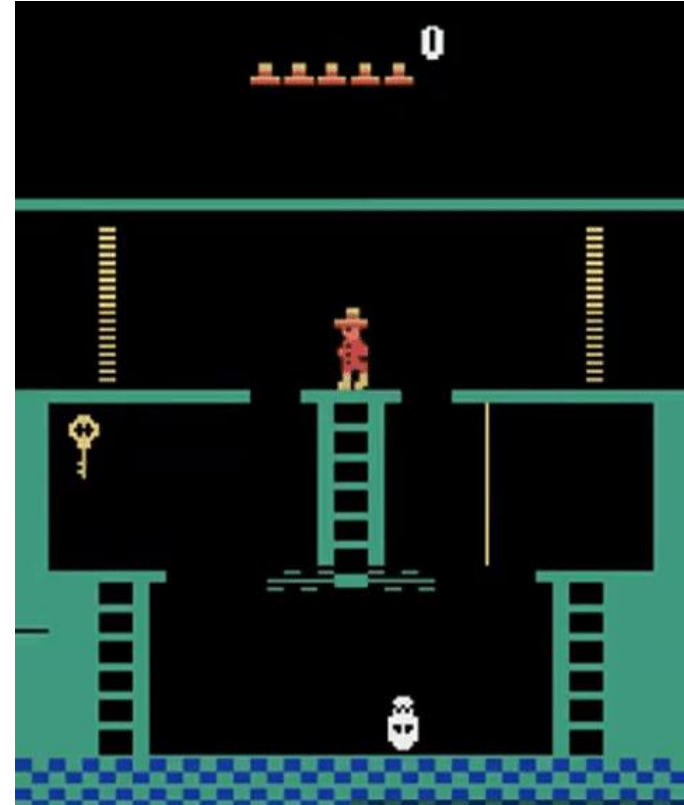
At this point (for the discussion of DQN), we pretend that EG, BE, or IGW will lead to sufficient exploration over the state space.

1. Exploration in MDPs (Not Easy)

Classic sparse-reward environments:



Mountain Car



Montezuma's Revenge

2. Function Approximation

To make DQN well approximate VI, we need

$$\forall s, a \quad Q_{\theta_{k+1}}(s, a) \approx R(s, a) + \gamma \sum_{s'} P(s'|s, a) \max_{a'} Q_{\theta_k}(s', a')$$

(ϵ -approximate) Bellman Completeness

an assumption both on the MDP and the function expressiveness

$$\forall \theta', \exists \theta \quad \forall s, a, \quad \left| Q_{\theta}(s, a) - \left(R(s, a) + \gamma \sum_{s'} P(s'|s, a) \max_{a'} Q_{\theta'}(s', a') \right) \right| \leq \epsilon$$

This allows us to quantify the regression error in each iteration.

2. Function Approximation

In HW1 you have shown

ϵ -Greedy ensures

$$\text{Regret} \lesssim \epsilon T + \sqrt{\frac{AT \cdot \text{Err}}{\epsilon}}$$

Regression error

$$\text{Err} = \sum_{t=1}^T \left(\hat{R}_t(x_t, a_t) - R(x_t, a_t) \right)^2$$

In value-based contextual bandits, the requirement / assumption for function approximation is

$$\exists \theta \quad \forall x, a \quad R_{\theta}(x, a) \approx R(x, a)$$

In value-based MDPs, the requirement / assumption for function approximation is

$$\forall \theta', \exists \theta \quad \forall s, a \quad Q_{\theta}(s, a) \approx R(s, a) + \gamma \sum_{s'} P(s'|s, a) \max_{a'} Q_{\theta'}(s', a')$$

BC assumption

Analysis of DQN assuming sufficient exploration and Bellman Completeness

Recall the analysis for the exact Value Iteration:

1. Value Iteration will terminate.

$$|Q_k(s, a) - Q_{k-1}(s, a)| \leq \epsilon \quad \forall s, a$$

2. When it terminates, it holds that

$$|Q_k(s, a) - Q^*(s, a)| \leq \frac{\epsilon}{1 - \gamma} \quad \forall s, a$$

3. When it terminates, it holds that

$$V^*(s) - V^{\hat{\pi}}(s) \leq \frac{2\epsilon}{(1 - \gamma)^2} \quad \forall s$$

$$\text{where } \hat{\pi}(s) = \operatorname{argmax}_a Q_k(s, a)$$

$$\begin{aligned} & \max_{s,a} |Q_k(s, a) - Q_{k-1}(s, a)| \\ & \leq \gamma \max_{s,a} |Q_{k-1}(s, a) - Q_{k-2}(s, a)| \end{aligned}$$

$$\text{ValueError} \leq \frac{1}{1 - \gamma} \text{BellmanError}$$

$$\text{Suboptimality} \leq \frac{1}{1 - \gamma} \text{ValueError}$$

Completing the Analysis of VI (1st Step)

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

$$\Rightarrow Q_{k-1}(s, a) = R(s, a) + \gamma \sum_{s'} P(s'|s, a) \max_{a'} Q_{k-2}(s', a')$$

$$Q_k(s, a) - Q_{k-1}(s, a) = \gamma \sum_{s'} P(s'|s, a) \left(\max_{a'} Q_{k-1}(s', a') - \max_{a'} Q_{k-2}(s', a') \right)$$

$$\begin{aligned} \max_x f(x) - \max_x g(x) &\leq \max_x |f(x) - g(x)| \end{aligned}$$

$$\begin{aligned} &\leq \gamma \sum_{s'} P(s'|s, a) \left| \max_{a'} Q_{k-1}(s', a') - \max_{a'} Q_{k-2}(s', a') \right| \\ &\leq \gamma \sum_{s'} P(s'|s, a) \max_{a'} |Q_{k-1}(s', a') - Q_{k-2}(s', a')| \\ &\leq \gamma \max_{s', a'} |Q_{k-1}(s', a') - Q_{k-2}(s', a')| \end{aligned}$$

$$Q_{k-1}(s, a) - Q_k(s, a) \leq$$

$$|Q_k(s, a) - Q_{k-1}(s, a)| \leq$$

Analysis of DQN assuming sufficient exploration and Bellman Completeness

$$\begin{aligned}
 \max_{s,a} |Q_k(s,a) - Q_{k-1}(s,a)| &\leq \gamma \max_{s,a} |Q_{k-1}(s,a) - Q_{k-2}(s,a)| + \max_{s,a} |b_k(s,a) - b_{k-1}(s,a)| \\
 \gamma = 0.95 &\leq \gamma^2 \max_{s,a} |Q_{k-2}(s,a) - Q_{k-3}(s,a)| + 2\varepsilon + \gamma \cdot 2\varepsilon' \\
 &\leq \dots \leq \gamma^{k-1} \max_{s,a} |Q_1(s,a) - Q_0(s,a)| + 2\varepsilon(1 + \gamma + \gamma^2 + \dots) \\
 Q_k(s,a) &= R(s,a) + \gamma \sum_{s'} P(s'|s,a) Q_{k-1}(s,a) + b_k(s,a) \\
 \text{where } |b_k(s,a)| &\leq \varepsilon' \\
 &\leq \gamma^{k-1} \max_{s,a} |Q_1(s,a) - Q_0(s,a)| + \frac{2\varepsilon'}{1-\gamma}
 \end{aligned}$$

1/1

DQN can be offline

Let $\mathcal{B}$ consists of (s, a, r, s') tuples collected by a mixture of **arbitrary policies**.

Data collection

For $k = 1, 2, \dots$

$\theta \leftarrow \theta_k$

For $m = 1, 2, \dots, M$:

Randomly pick an i (or a mini-batch) from $\mathcal{B}$

$$\theta \leftarrow \theta - \alpha \nabla_{\theta} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\theta_k}(s'_i, a') \right)^2$$

$\theta_{k+1} \leftarrow \theta$

Perform Value Iteration

Again, its success relies on 1) $\mathcal{B}$ contains data with sufficiently wide range of state-actions, 2) Bellman completeness.

The same theoretical analysis applies.

Handling the Non-Ideal Case

When DQN cannot well-approximate VI

In practice,

- We may not be able to collect sufficiently wide range of state-actions
- Bellman completeness may not hold

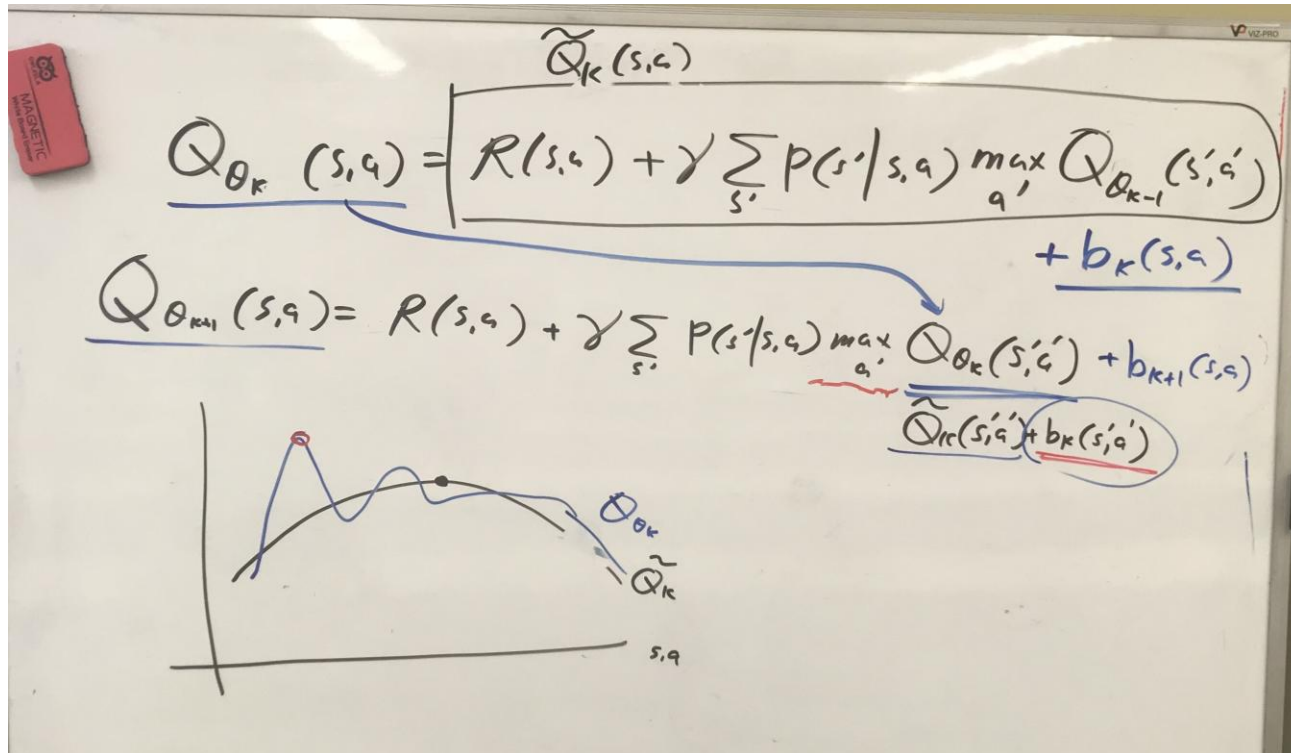
In either case, we may not have

$$\forall s, a \quad Q_{\theta_{k+1}}(s, a) \approx R(s, a) + \gamma \sum_{s'} P(s'|s, a) \max_{a'} Q_{\theta_k}(s', a')$$

This makes our previous analysis based on VI fails.

When DQN cannot well-approximate VI

In this case, $Q_{\theta_k}(s, a)$ tends to **overestimate** $Q^*(s, a)$, and the greedy policy $\hat{\pi}(s) = \operatorname{argmax}_a Q_{\theta_k}(s, a)$ could be very bad.



When DQN cannot well-approximate VI

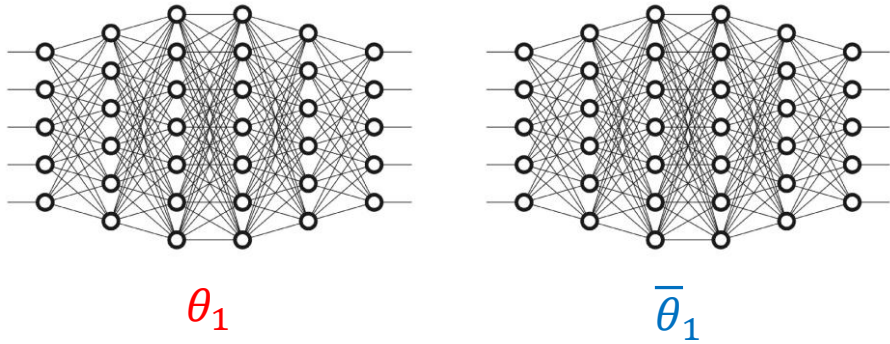
Such “seeking the error” behavior is due to “**bootstrapping**”

- An issue only in MDP but not in bandits

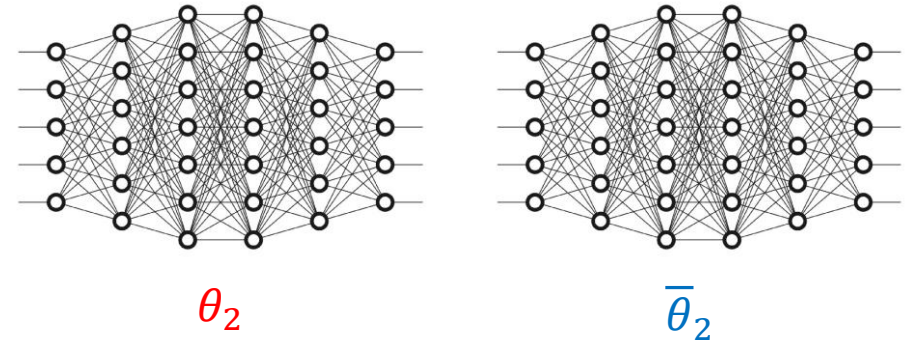
To prevent overestimation, two strategies are

- Double Q-learning: decorrelating the choice of argmax action and the error of the value function
- Conservative Q-learning: being conservative

Double DQN (v1)

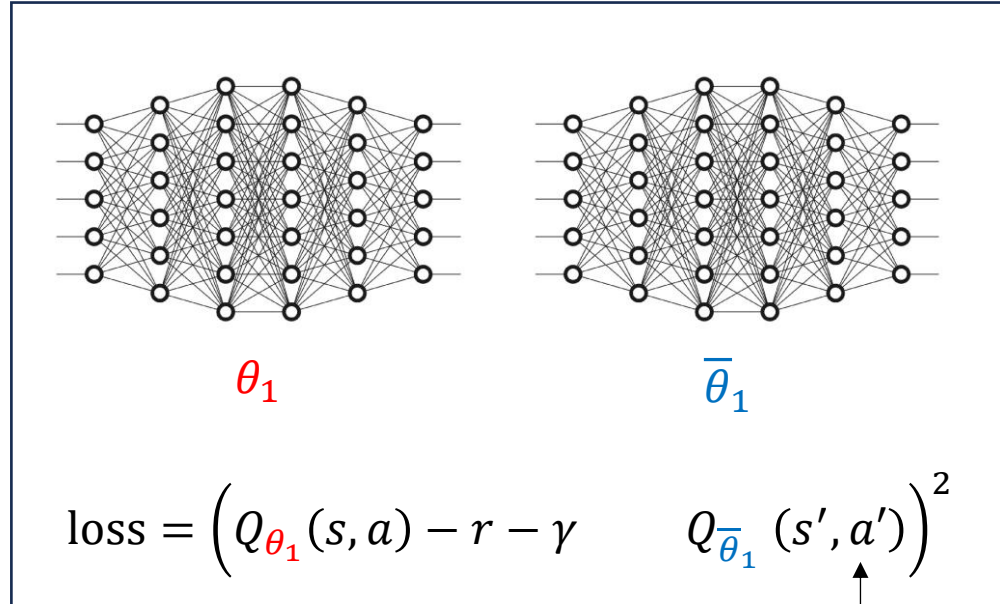


$$\text{loss} = \left(Q_{\theta_1}(s, a) - r - \gamma \max_{a'} Q_{\bar{\theta}_1}(s', a') \right)^2$$

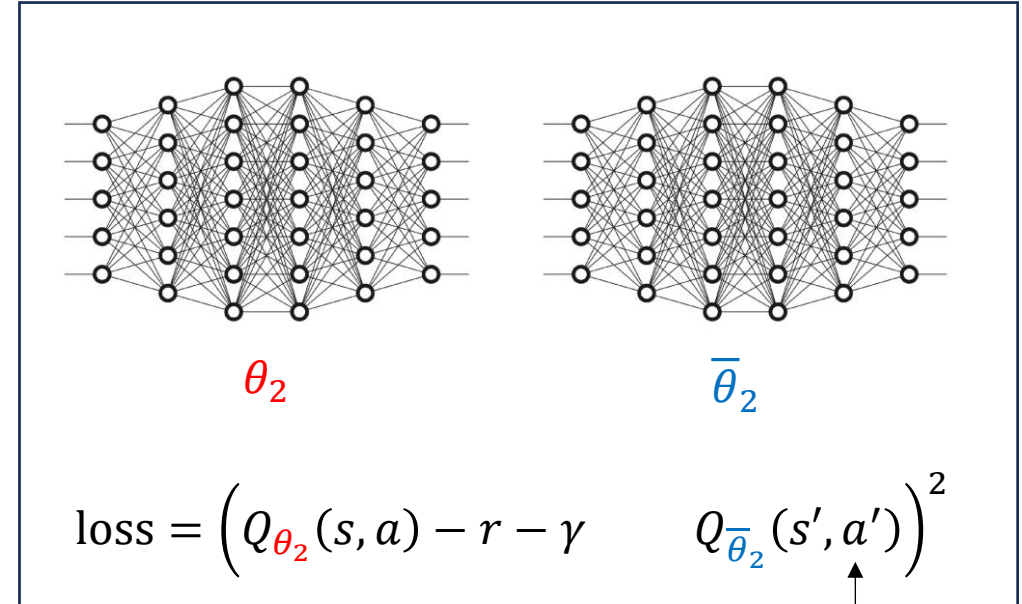


$$\text{loss} = \left(Q_{\theta_2}(s, a) - r - \gamma \max_{a'} Q_{\bar{\theta}_2}(s', a') \right)^2$$

Double DQN (v1)

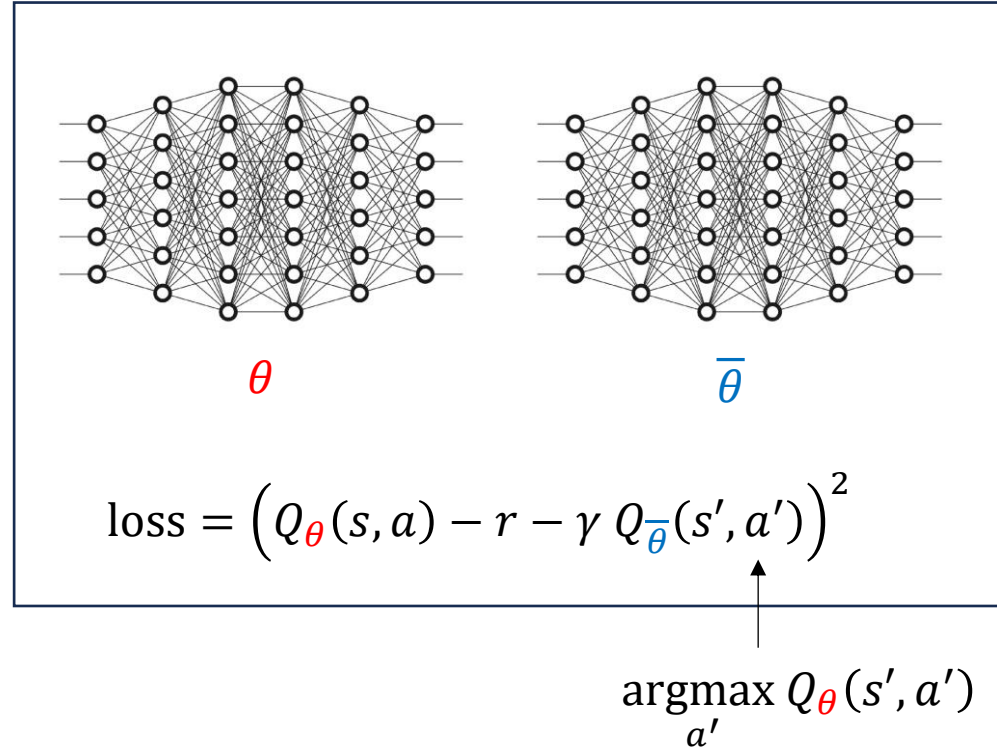


$\arg\max_{a'} Q_{\bar{\theta}_2}(s', a')$



$\arg\max_{a'} Q_{\bar{\theta}_1}(s', a')$

Double DQN (v2)



Hado van Hasselt, Arthur Guez, David Silver. Deep Reinforcement Learning with Double Q-learning. 2015.

Conservative Q-learning (CQL)

$$\begin{aligned}\theta_{k+1} &= \operatorname{argmin}_{\theta} \sum_{i \in \mathcal{B}} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\theta_k}(s'_i, a') \right)^2 \\ &\quad + \alpha \sum_{i \in \mathcal{B}} \left(\log \left(\sum_a \exp(Q_{\theta}(s_i, a)) \right) - Q_{\theta}(s_i, a_i) \right) \\ &= \operatorname{argmin}_{\theta} \sum_{i \in \mathcal{B}} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\theta_k}(s'_i, a') \right)^2 \\ &\quad + \alpha \sum_{i \in \mathcal{B}} \left(\max_{\mu} \sum_a \mu(a|s_i) Q_{\theta}(s_i, a) - Q_{\theta}(s_i, a_i) + \text{KL}(\mu(\cdot | s_i), \text{Unif}) \right)\end{aligned}$$

Comparison

- Double-Q: make the $\operatorname{argmax}_a Q_\theta(s, a)$ choice decoupled from θ
- Conservative-Q: mitigate the overestimation of $\max_a Q_\theta(s_i, a)$ over $Q_\theta(s_i, a_i)$

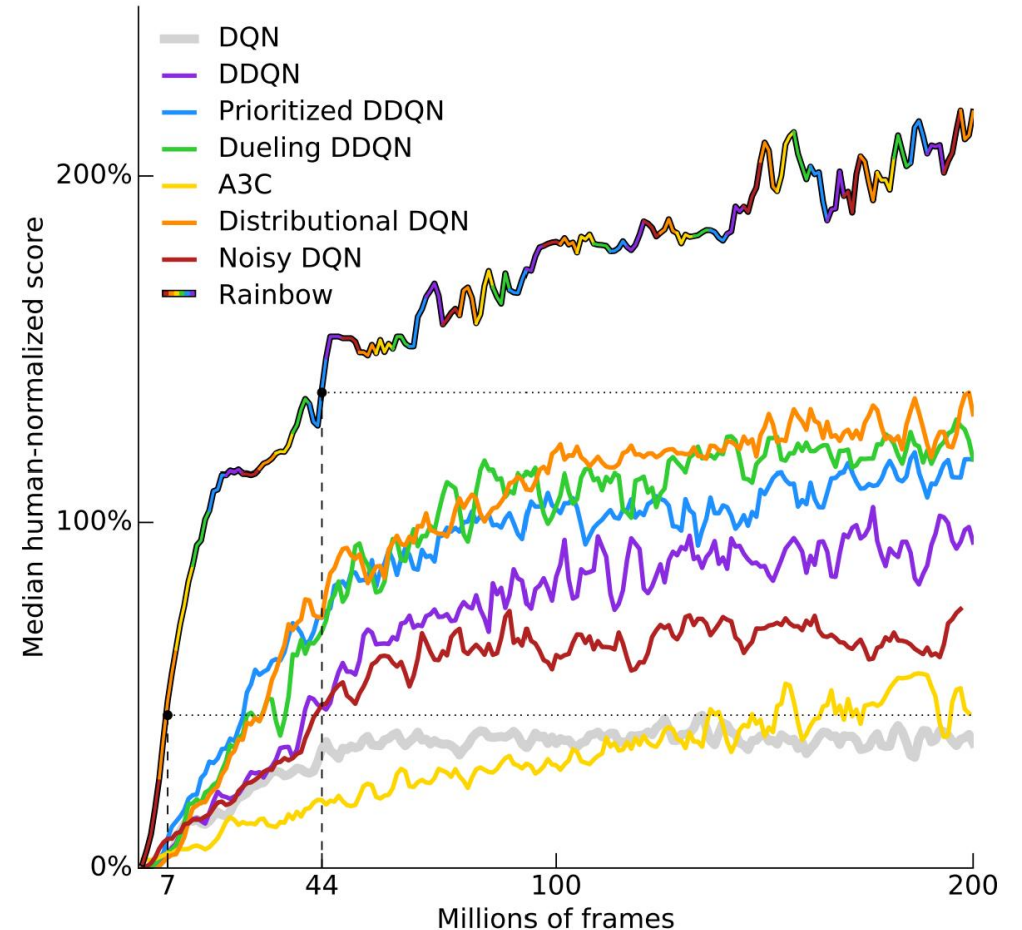
Summary for DQN

- Motivation: approximating Value Iteration using **samples** and **function approximation**
- Standard elements: target network, replay buffer
- Work as desired when both of the following conditions hold:
 - The learner is able to obtain exploratory data (online or offline)
 - Neural network is sufficiently expressive: Bellman completeness
- When the conditions above do not hold
 - Tends to overestimate Q values and suggest arbitrary actions
- Solutions
 - Double Q-learning
 - Conservative Q-learning

Improvements on DQN

- Dueling DDQN
- Prioritized replay
- Distributional DQN
- ...

Rainbow: Combining Improvements in Deep Reinforcement Learning. 2018.



Other Variants that Fail

An Unstable Variant

DQN without target network



$$\theta \leftarrow \theta - \alpha \nabla_{\theta} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \text{sg} \left[\max_{a'} Q_{\theta}(s'_i, a') \right] \right)^2$$

For $k = 1, 2, \dots$

Randomly pick an i (or a mini-batch) from $\mathcal{B}$

$$\theta \leftarrow \theta - \alpha \nabla_{\theta} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\bar{\theta}}(s'_i, a') \right)^2$$

$$\bar{\theta} \leftarrow \theta$$

stop gradient

without stop gradient

$$2 \left(\frac{d(\sim)}{d\theta} \right)$$

$$(w/o \text{ sg}) \nabla Q_{\theta}(s_i, a_i) - \gamma \nabla_{\theta} \max_{a'} Q_{\theta}(s'_i, a')$$

$$(w/ \text{ sg}) \nabla Q_{\theta}(s_i, a_i)$$

cf. DQN with target network

For $k = 1, 2, \dots$

Randomly pick an i (or a mini-batch) from $\mathcal{B}$

$$\theta \leftarrow \theta - \alpha \nabla_{\theta} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\bar{\theta}}(s'_i, a') \right)^2$$

$$\bar{\theta} \leftarrow (1 - \tau) \bar{\theta} + \tau \theta$$

For $k = 1, 2, \dots$

$$\theta \leftarrow \bar{\theta}$$

For $m = 1, \dots, M$:

Randomly pick an i (or a mini-batch) from $\mathcal{B}$

$$\theta \leftarrow \theta - \alpha \nabla_{\theta} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\bar{\theta}}(s'_i, a') \right)^2$$

$$\bar{\theta} \leftarrow \theta$$

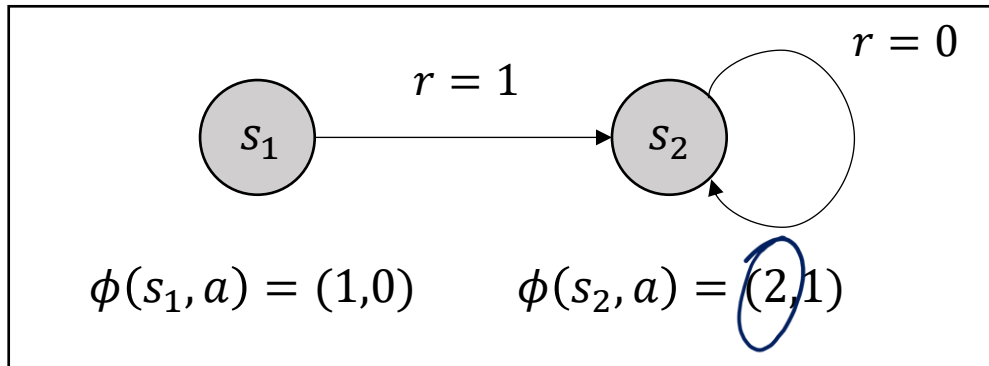
$$\nabla_{\theta} \max_a f_{\theta}(a) = \nabla_{\theta} f_{\theta}(\hat{a})$$

define $\hat{a} = \arg\max_a f_{\theta}(a)$

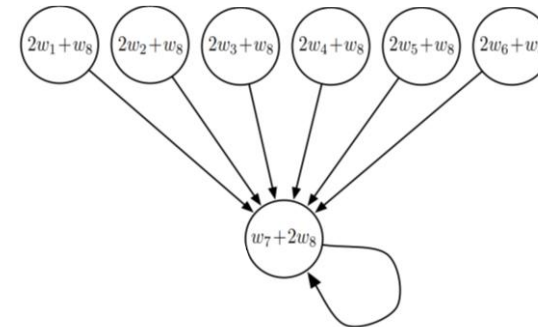
An Unstable Variant

Could diverge for **off-policy problems** (e.g. **Deep Q-learning**) even when exploration assumption and Bellman completeness hold.

Off-policy: Training Q^π using data $\left(\begin{array}{l} \text{collected from some other policy } \pi' \neq \pi \\ \text{or some distribution very different from the distribution induced by } \pi \end{array} \right)$



$$Q_\theta(s, a) = \phi(s, a)^T \theta$$



Simplified from the “Baird’s counterexample”
(see Sutton and Barto Section 11.2)

The Effect of Target Network

Let $KN = 100000$

For $k = 1, 2, \dots, K$

$$\theta_k \leftarrow \theta$$

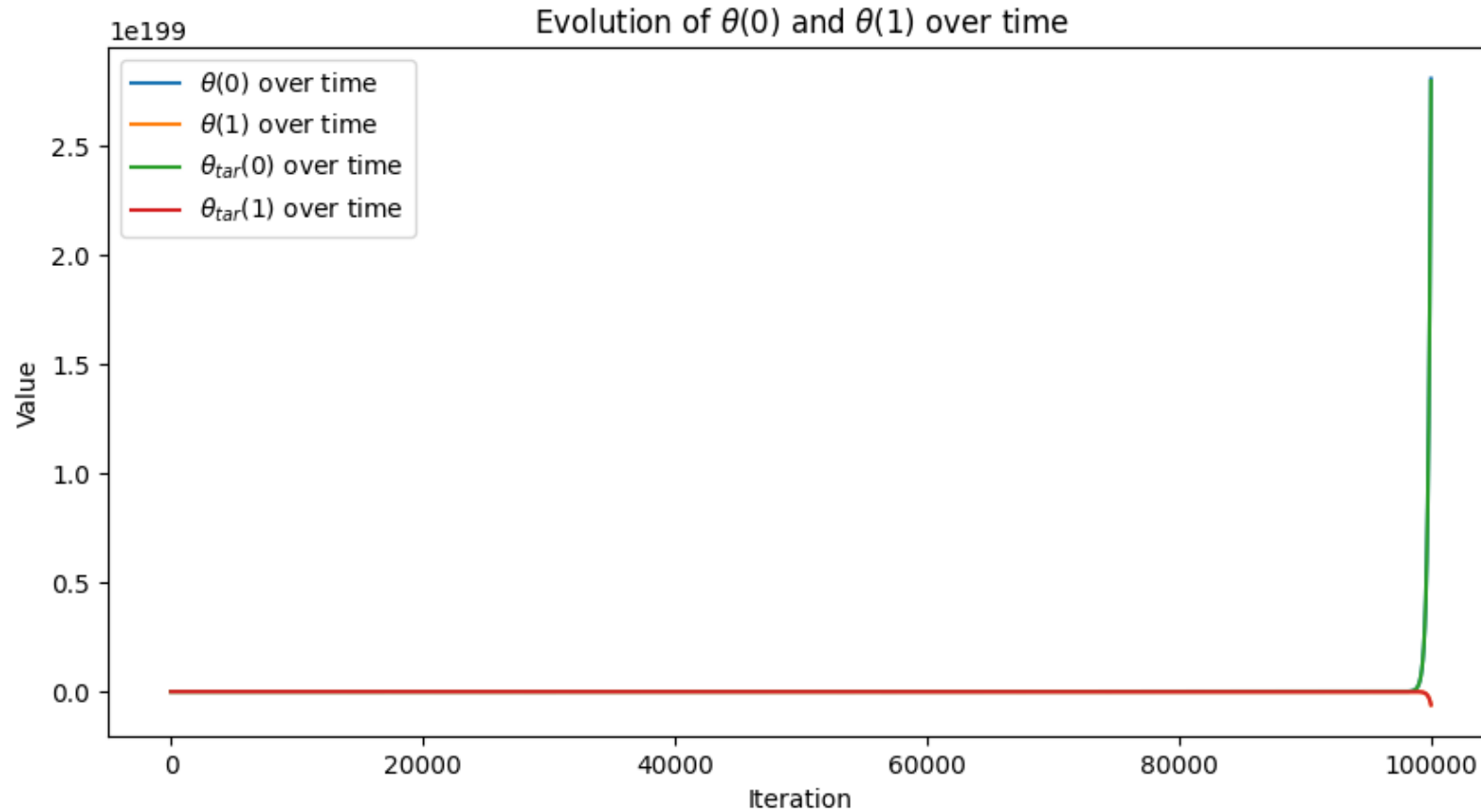
For $i = 1, \dots, N$:

Sample $(s, a, r, s') \sim \text{Uniform} \{(\underline{s_1}, a, 1, \underline{s_2}), (\underline{s_2}, a, 0, \underline{s_2})\}$

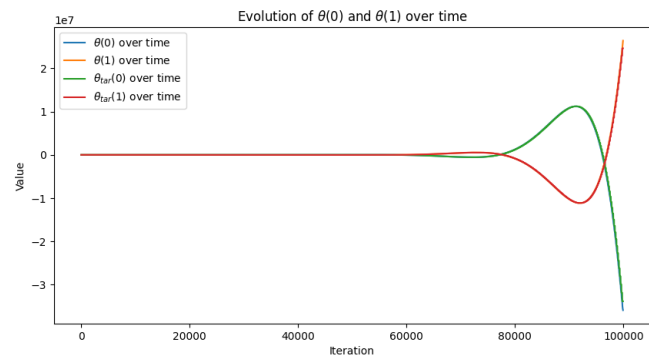
$$\theta \leftarrow \theta - \alpha \left(\phi(s, a)^\top \theta - r - \gamma \phi(s', a)^\top \theta_k \right) \phi(s, a)$$

$$\theta_{k+1} \leftarrow \theta$$

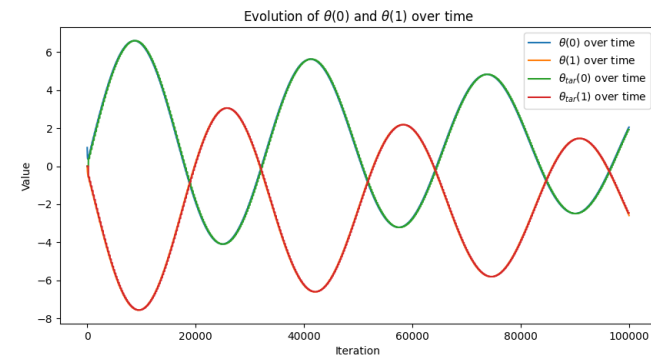
The Effect of Target Network



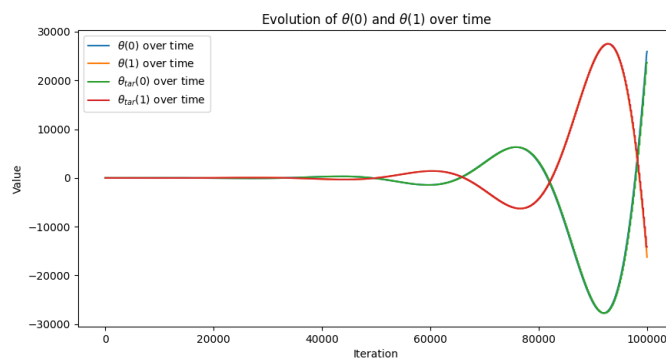
N=1



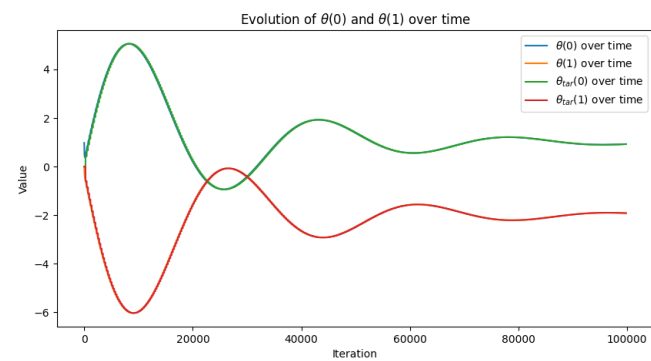
N=150



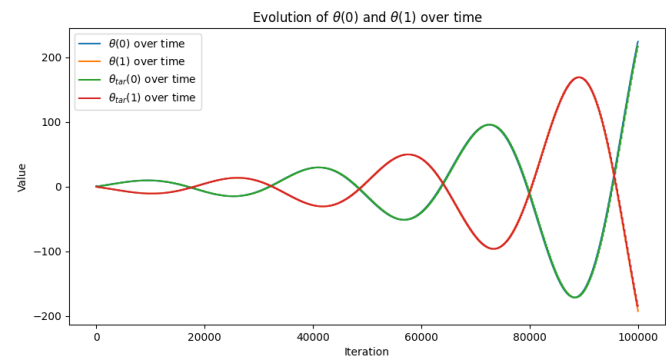
N=210



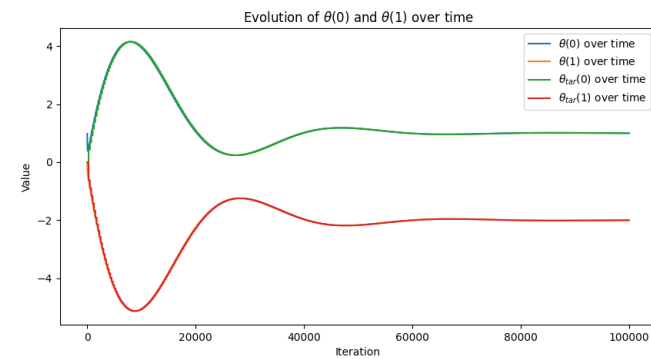
N=170



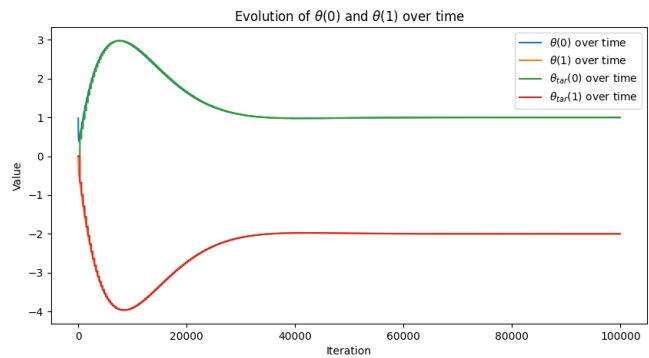
N=230



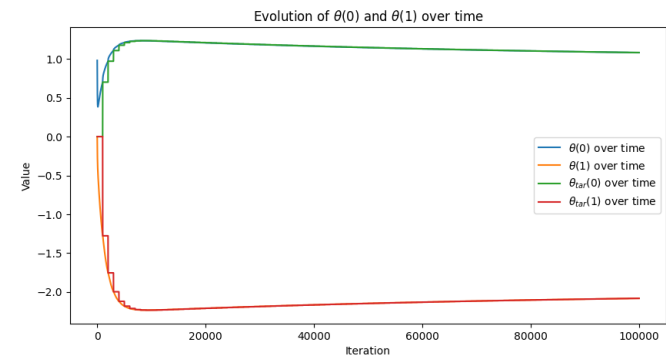
N=190



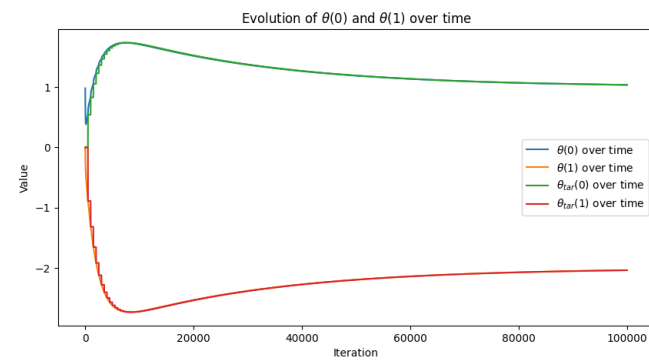
N=250



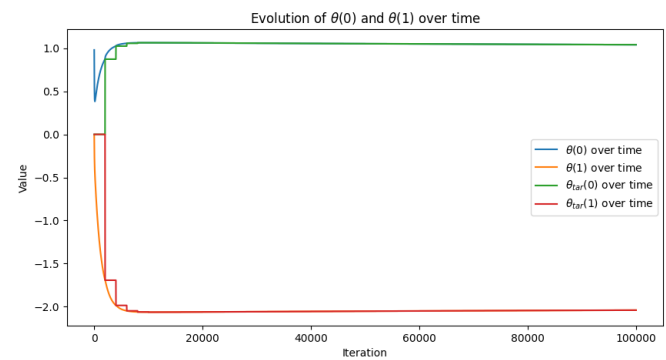
N=300



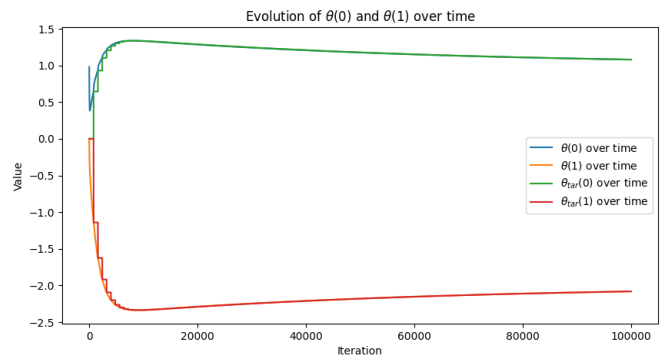
N=1000



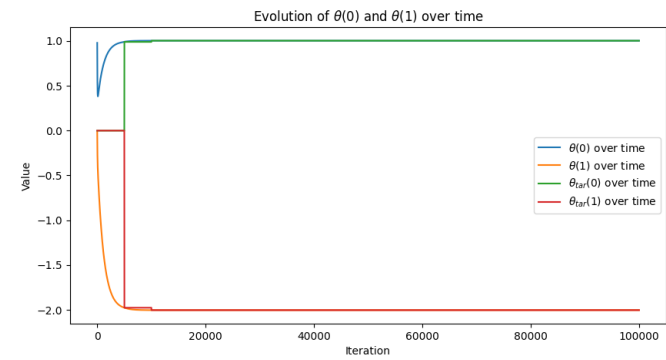
N=500



N=2000



N=800



N=5000

An Unstable Variant

- This type of algorithm (i.e., perform regression without fixing the target) is still useful in certain cases:
 - **On-policy** problems where Q^π is trained using data collected from π – will see this later
 - Off-policy Q-learning without function approximation – this is the very original “Q-learning” algorithm being proposed by Chris Watkins in 1989.
- However, in Deep Q-learning, this variant is generally less stable and less recommended.

A Biased Variant

DQN without **stop gradient**

For $k = 1, 2, \dots$

Randomly pick an i (or a mini-batch) from $\mathcal{B}$

$$\theta \leftarrow \theta - \alpha \nabla_{\theta} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\theta}(s'_i, a') \right)^2$$

Residual gradient

cf. standard DQN

For $k = 1, 2, \dots$

Randomly pick an i (or a mini-batch) from $\mathcal{B}$

$$\theta \leftarrow \theta - \alpha \nabla_{\theta} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\bar{\theta}}(s'_i, a') \right)^2$$

$$\bar{\theta} \leftarrow (1 - \tau) \bar{\theta} + \tau \theta$$

For $k = 1, 2, \dots$

$$\theta \leftarrow \bar{\theta}$$

For $m = 1, \dots, M$:

Randomly pick an i (or a mini-batch) from $\mathcal{B}$

$$\theta \leftarrow \theta - \alpha \nabla_{\theta} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\bar{\theta}}(s'_i, a') \right)^2$$

$$\bar{\theta} \leftarrow \theta$$

A Biased Variant

This variant will converge (as it is similar to standard SGD), but the solution it converges to could be undesirable.

The underlying loss function of this algorithm is

$$\text{Loss}^{\text{RG}}(\theta) = \sum_{i \in \mathcal{B}} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\theta}(s'_i, a') \right)^2$$

cf. the desired Bellman error that we truly want to minimize is

$$\text{BE}(\theta) = \sum_{i \in \mathcal{B}} \left(Q_{\theta}(s_i, a_i) - R(s_i, a_i) - \gamma \sum_{s'} P(s'|s, a) \max_{a'} Q_{\theta}(s', a') \right)^2$$

Recall the Bellman equation: $\forall s, a \quad Q(s, a) = R(s, a) + \gamma \sum_{s'} P(s'|s, a) \max_{a'} Q_{\theta}(s', a')$

Lemma 1. For random variables X, Y , let $\mu_X \triangleq \mathbb{E}_Y[Y|X]$. Then

$$\mathbb{E}[(X - Y)^2] = \mathbb{E}_X[(X - \mu_X)]^2 + \mathbb{E}_{X,Y}[(\mu_X - Y)^2].$$

Proof.

$$\begin{aligned}\mathbb{E}_{X,Y}[(X - Y)^2] &= \mathbb{E}_{X,Y}[(X - \mu_X + \mu_X - Y)^2] \\ &= \mathbb{E}_X[(X - \mu_X)]^2 + 2\mathbb{E}_{X,Y}[(X - \mu_X)(\mu_X - Y)] + \mathbb{E}_{X,Y}[(\mu_X - Y)^2].\end{aligned}$$

The second term above is zero because

$$\begin{aligned}\mathbb{E}_{X,Y}[(X - \mu_X)(\mu_X - Y)] &= \mathbb{E}_X [\mathbb{E}_Y [(X - \mu_X)(\mu_X - Y) | X]] \\ &= \mathbb{E}_X [(X - \mu_X)\mathbb{E}_Y [(\mu_X - Y) | X]] \\ &= \mathbb{E}_X [(X - \mu_X)0] \\ &= 0.\end{aligned}$$

A Biased Variant

Assuming μ is the state-action distribution in $\mathcal{B}$, we have

$$\begin{aligned}\mathbb{E}[\text{Loss}^{\text{RG}}(\theta)] &= \mathbb{E}_{(s,a) \sim \mu} \mathbb{E}_{r \sim R(s,a)} \mathbb{E}_{s' \sim P(\cdot|s,a)} \left[\left(Q_\theta(s,a) - r - \gamma \max_{a'} Q_\theta(s',a') \right)^2 \right] \\ &= \mathbb{E}_{(s,a) \sim \mu} \left[\left(Q_\theta(s,a) - R(s,a) - \gamma \mathbb{E}_{s' \sim P(\cdot|s,a)} \max_{a'} Q_\theta(s',a') \right)^2 \right]\end{aligned}\tag{1}$$

$$+ \mathbb{E}_{(s,a) \sim \mu} \mathbb{E}_{r \sim R(s,a)} \left[(r - R(s,a))^2 \right]\tag{2}$$

$$+ \mathbb{E}_{(s,a) \sim \mu} \mathbb{E}_{\tilde{s} \sim P(\cdot|s,a)} \left[\left(\max_{a'} Q_\theta(\tilde{s},a') - \mathbb{E}_{s' \sim P(\cdot|s,a)} \left[\max_{a'} Q_\theta(s',a') \right] \right)^2 \right]\tag{3}$$

where the last equality is by [Lemma 1](#). Term (1) is equal to $\text{BE}(\theta)$, which is the loss we truly want to minimize. Term (3) is an undesired term that is non-zero when the transition is non-deterministic.

A Biased Variant

- Theoretically, if the transition is deterministic, residual gradient is optimizing a correct loss function without bias.
 - Still not recommended: for other reasons, it converges slower and has worse performance when there is no sufficient 1) data coverage or 2) powerful function approximation.
 - It deviates from the idea of Value Iteration / dynamic programming.

Fujimoto et al. Why should I trust you, bellman? the bellman error is a poor replacement for value error. 2022.
Saleh and Jiang. Deterministic Bellman Residual Minimization

Summary

For $k = 1, 2, \dots$

$$\theta \leftarrow \bar{\theta}$$

For $m = 1, \dots, M$:

Randomly pick an i (or a mini-batch) from $\mathcal{B}$

$$\theta \leftarrow \theta - \alpha \nabla_{\theta} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\bar{\theta}}(s'_i, a') \right)^2$$

$$\bar{\theta} \leftarrow \theta$$

For $k = 1, 2, \dots$

Randomly pick an i (or a mini-batch) from $\mathcal{B}$

$$\theta \leftarrow \theta - \alpha \nabla_{\theta} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\bar{\theta}}(s'_i, a') \right)^2$$

$$\bar{\theta} \leftarrow (1 - \tau) \bar{\theta} + \tau \theta$$

DQN

Not taking gradient on the $Q_{\bar{\theta}}(s', a')$ term

$\bar{\theta}$ is updated slower than θ

For $k = 1, 2, \dots$

Randomly pick an i (or a mini-batch) from $\mathcal{B}$

$$\theta \leftarrow \theta - \alpha \nabla_{\theta} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\bar{\theta}}(s'_i, a') \right)^2$$

$$\bar{\theta} \leftarrow \theta$$

$\bar{\theta}$ is updated at the same speed as θ

For $k = 1, 2, \dots$

Randomly pick an i (or a mini-batch) from $\mathcal{B}$

$$\theta \leftarrow \theta - \alpha \nabla_{\theta} \left(Q_{\theta}(s_i, a_i) - r_i - \gamma \max_{a'} Q_{\theta}(s'_i, a') \right)^2$$

Taking gradient on the $Q_{\bar{\theta}}(s', a')$ term

Not Recommended