

Introduction to Reinforcement Learning based Control and Deep Reinforcement Learning based control



Research

Centre de Recherche en Automatique de Nancy
(CRAN) UMR 7039,
Faculté des Sciences et Technologies
Boulevard des Aiguillettes - BP 70239 - Bât.
1er cycle 54506 Vandoeuvre-lès-Nancy
Cedex France

Associate Professor
(Maitre de Conférences)

Automatic Control, Reliability of Systems
mayank-shekhar.jha@univ-lorraine.fr



Teaching

Polytech Nancy (ESSTIN),
2 Rue Jean Lamour 54509
Vandoeuvre-lès-Nancy,
Cedex France



Email: mayank-shekhar.jha@univ-lorraine.fr

Introduction

Markov Decision process

Dynamic Programming for planning
and control

Model Free Prediction and Control
(SARSA and Q-learning)

Model-free Prediction

Temporal Difference (TD)

So far:

- Definition, setting of RL, MDP,
- Solving a known MDP

Here

- Model-free Prediction
- Estimate Value Function of unknown MDP

Next:

- Model-free control
- Optimisation of unknown MDP
- Approximate DP : Deep Learning + Reinforcement Learning.

Temporal Difference (TD)

- learns directly from 'episodes' of experiences
- TD is model-free : no knowledge of MDP transitions / rewards required
- TD: learns by bootstrapping (update a guess using a guess)
- TD: Can learn from incomplete episodes.

Goal: learn v_π online from experience under policy π

Temporal Difference TD(0) algorithm

- Update value $V(S_t)$ towards estimated return $R_{t+1} + \gamma V(S_{t+1})$

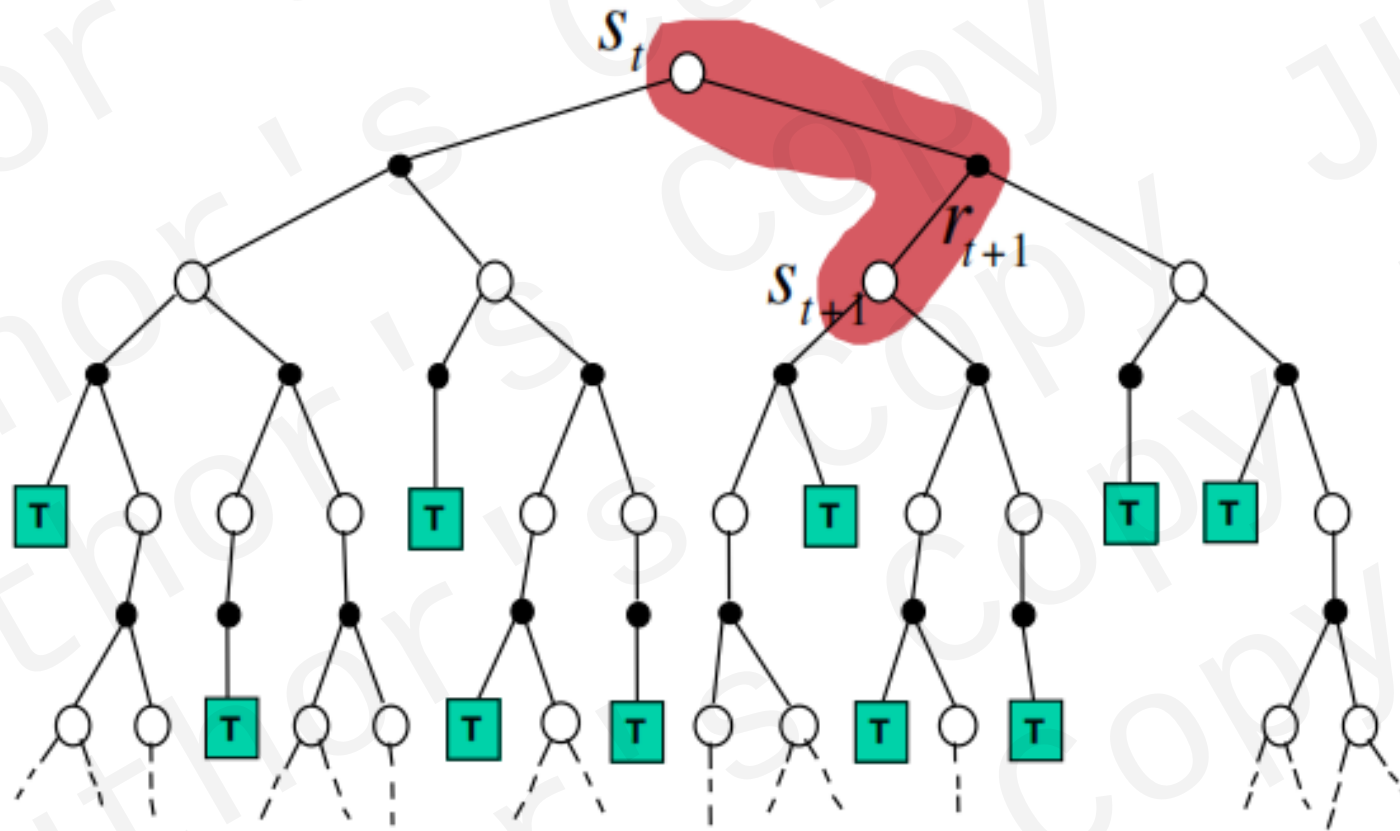
$$V(S_t) \leftarrow V(S_t) + \alpha \left(\underbrace{R_{t+1} + \gamma V(S_{t+1})}_{\text{estimated return}} - \underbrace{V(S_t)}_{\text{present value}} \right)$$

difference between some estimate of return and actual return.

TD target: $R_{t+1} + \gamma V(S_{t+1})$

TD error: $\delta_t = \text{TD target} - \underbrace{V(S_t)}_{\text{actual value}}$

TD(0) based backup



Monte Carlo (MC) learning

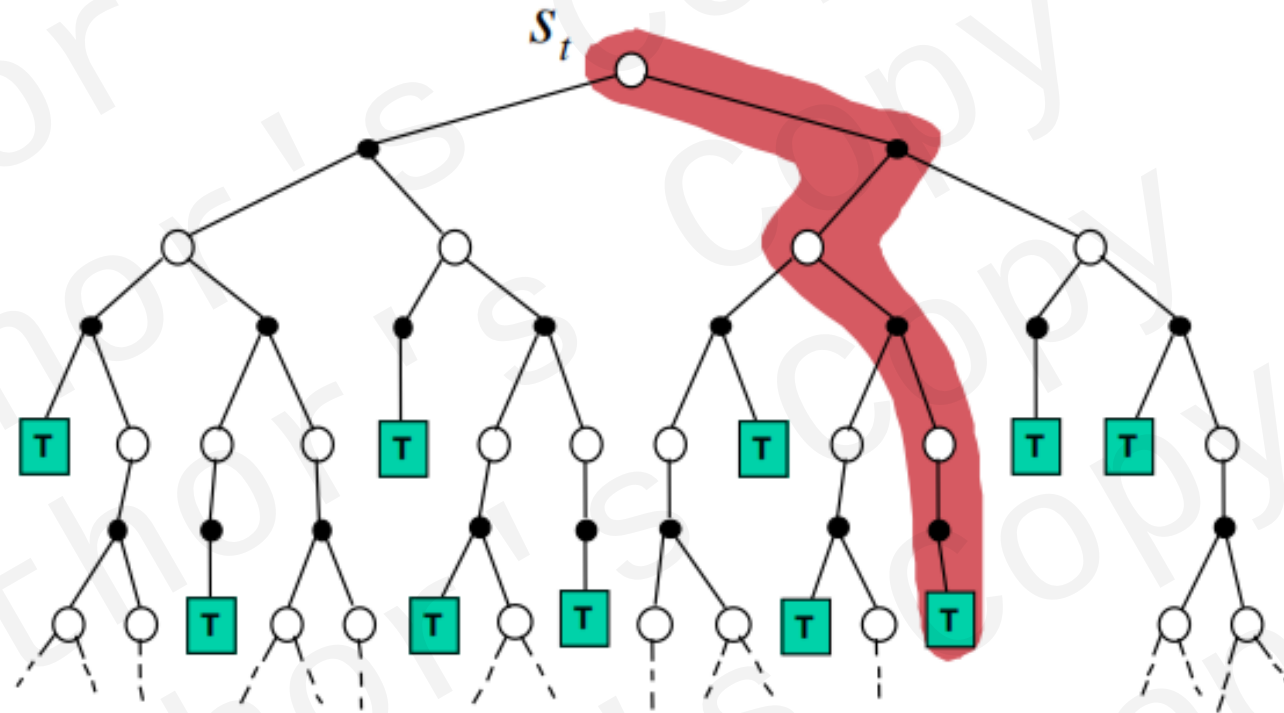
- State ' s ' is evaluated based on every state visited in an episode
- Update $V(S_t)$ towards actual (not estimated) return G_t

$$V(S_t) \leftarrow V(S_t) + \alpha (G_t - V(S_t))$$

Limitation

- Update only after episode end.
- Cannot learn on the run!!

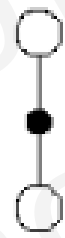
Monte Carlo based back-up.



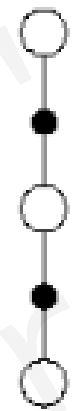
Difference TD and MC

- TD can learn before final outcome at every step.
- MC must update after completion.
- TD can learn
 - without final outcome
 - from incomplete sequences
 - in continuing environments.
- MC only works for episodic environments.
- TD exploits Markov property (efficient in Markov Environments)

TD (1-step)



2-step



3-step



n-step

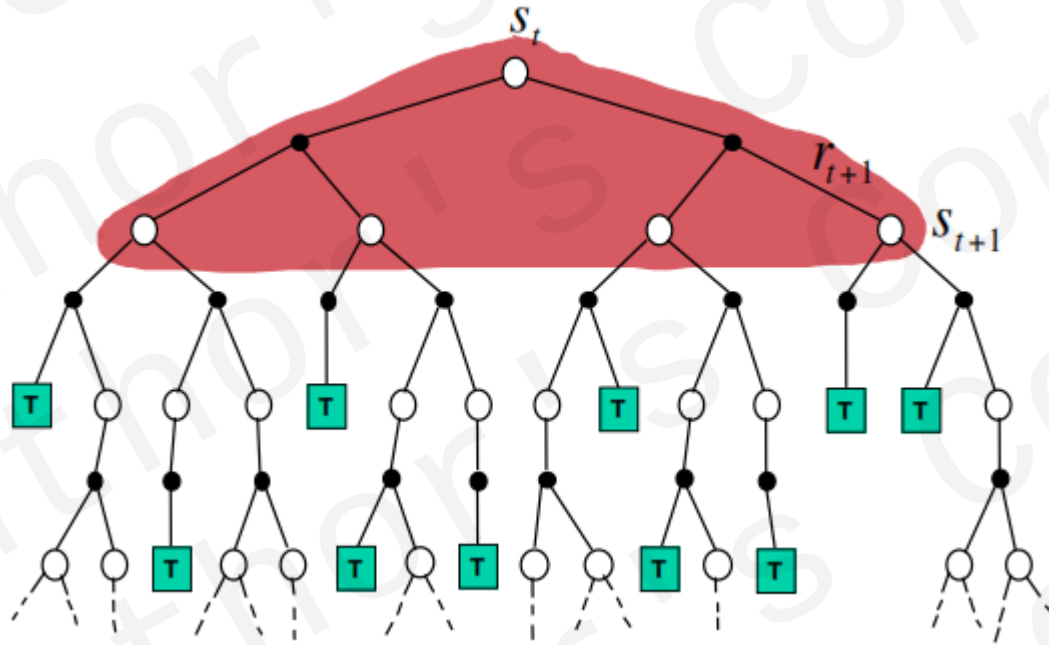


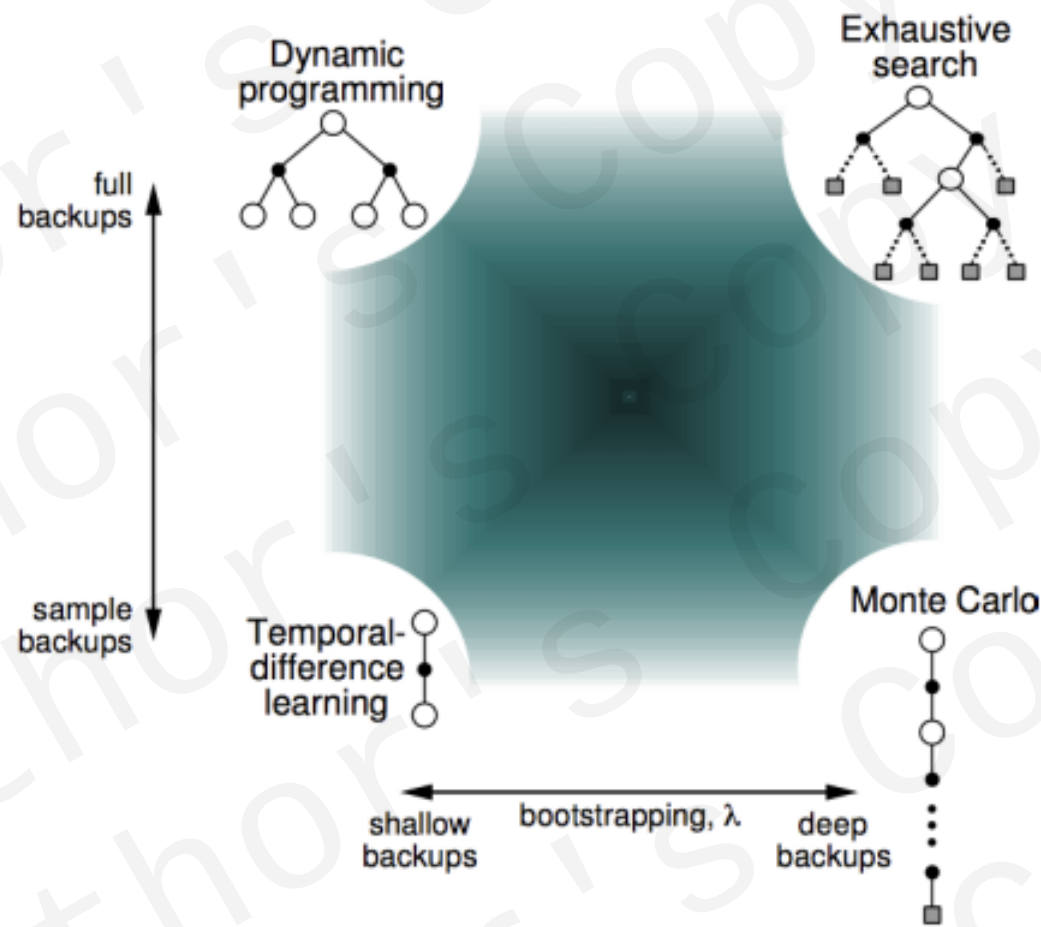
Monte Carlo



Dynamic Programming Back up.

$$V(s_t) \leftarrow \mathbb{E}_{\pi} (R_{t+1} + \gamma V(s_{t+1}))$$





Model-free Control

TD control (On-policy SARSA)

TD control (Off-Policy) Q-learning

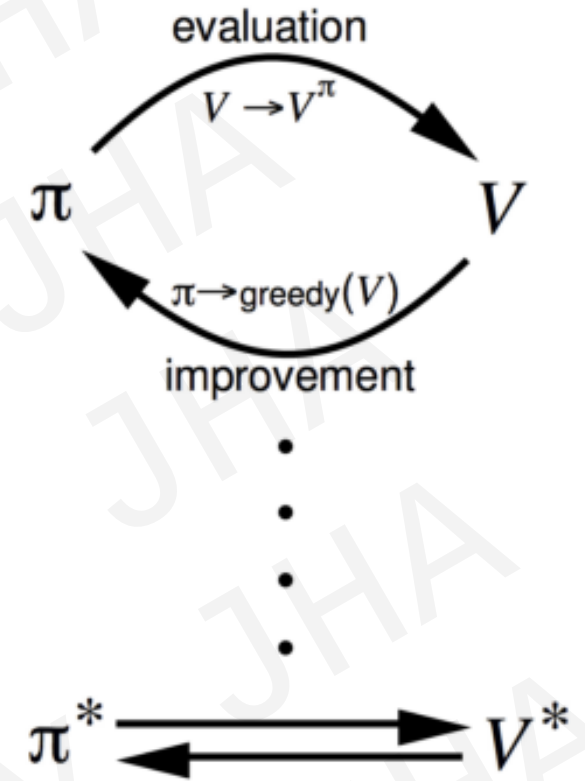
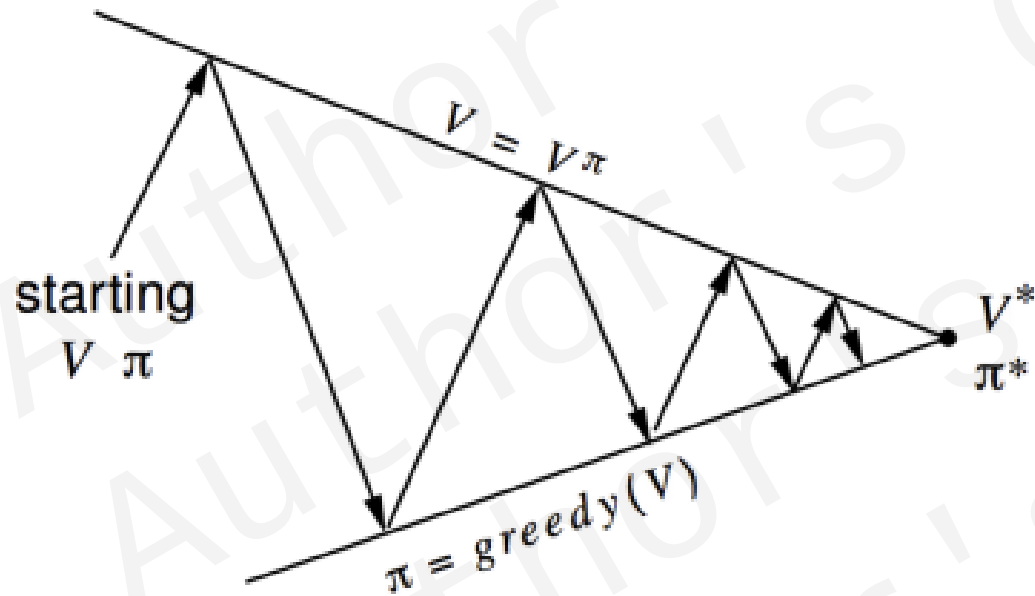
Model Free Control

→ MDP model is unknown,
experience sampling is possible!

On-Policy learning: learn while on 'RUN'
learn about policy π , from experience
sampled from π

Off-Policy learning: learn while off-line.
learn about policy π from experience
sampled from μ .

Generalised Policy Iteration



Policy Evaluation: Evaluate V_π (ϵ_n : Iterative Policy Evaluation)

Policy Improvement: Generate $\pi' \geq \pi$ (ϵ_n : Greedy Policy Improvement)

what is Greedy strategy??

Greedy Policy Improvement Example:

$$\pi'(s) = \arg \max_{a \in A} Q(s, a)$$

choose the action that maximises the Q-fun value.

the action chosen is greedy with respect to our objective

Chooses action
that gives max

maximise
the
value!

What is limitation of greedy policy → can get stuck at local minimas!
Imagine being greedy but not exploring all state space!

Exploration

ϵ -Greedy satisfies the following conditions

ϵ -Greedy Exploration.

→ Choose greedy action : with probability $(1-\epsilon)$

→ Choose random action : with probability ϵ

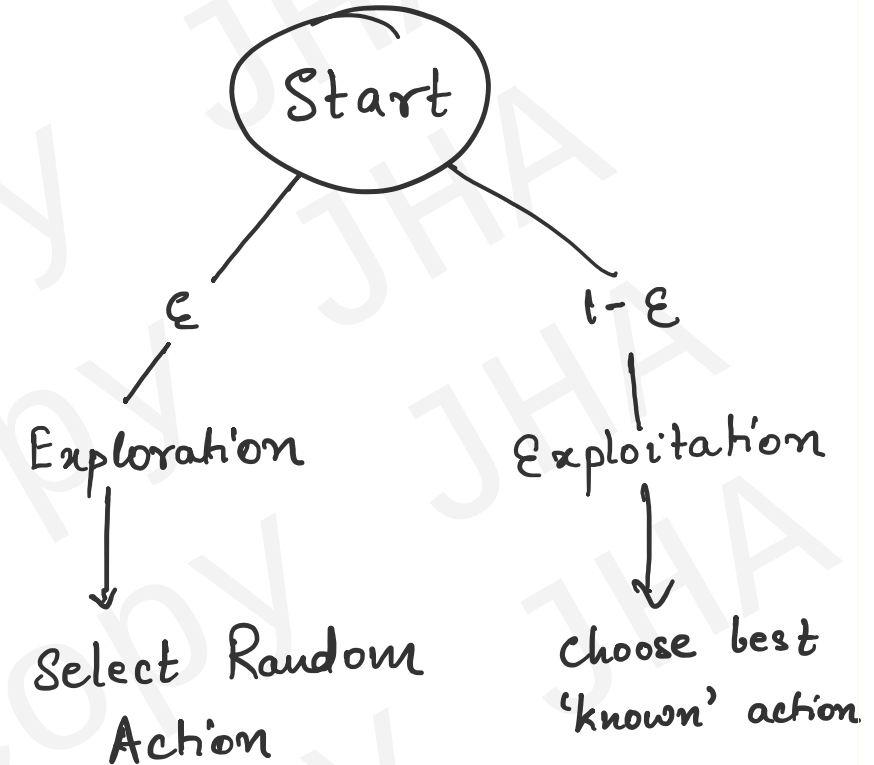
Greedy in Limit with Infinite Exploration (GLIE)

→ all state-action pairs are explored infinitely many times.

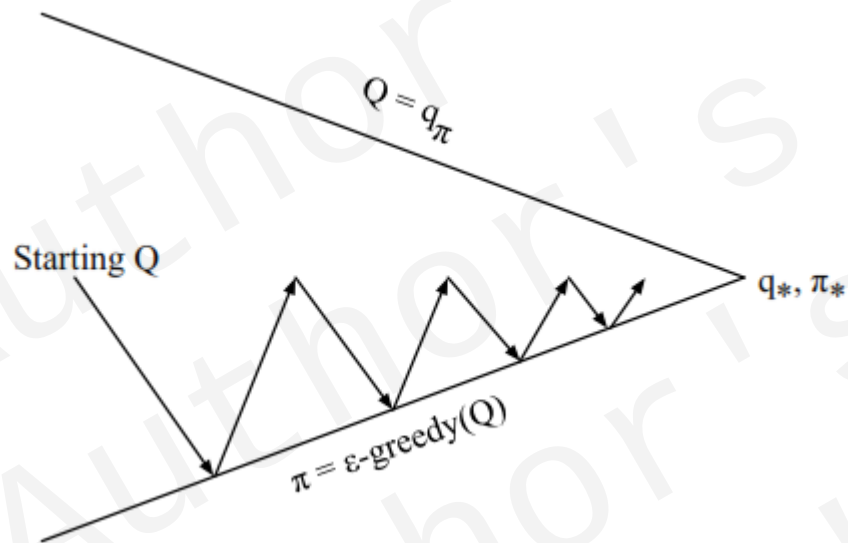
→ Policy converges on a greedy policy in asymptotic sense

$$\lim_{k \rightarrow \infty} \pi_k(a|s) = 1 \quad (a = \arg\max_{a' \in A} Q_k(s, a'))$$

Note: ϵ -greedy with $\epsilon_k = 1/k$ is GLIE.



Policy Iteration (TD control)



Idea:

- 1) Apply TD to update $Q(s,a)$
- 2) Use ϵ -greedy policy improvement
- 3) Update every time step!

Every time step:

- Policy Evaluation: Use TD to update $Q(s,a)$
$$Q(s,a) \leftarrow Q(s,a) + \alpha (R + \gamma Q(s',a') - Q(s,a))$$
- Policy Improvement: ϵ -greedy policy improvement

TD Control (Update Eqn)

$$Q(s_t, a_t) \leftarrow Q(s_t, a_t) + \alpha \left(\underbrace{R_{t+1} + \gamma Q(s_{t+1}, a_{t+1})}_{\text{TD target}} - \underbrace{Q(s_t, a_t)}_{\text{current estimate}} \right)$$

TD Control (Algorithm)

$$Q(s_t, a_t) \leftarrow Q(s_t, a_t) + \alpha \left(R_{t+1} + \gamma Q(s_{t+1}, A_{t+1}) - Q(s_t, A_t) \right)$$

Initialize : $Q(s, a) \forall s \in S, a \in A$. $\rightarrow$ initialize with all action-values to zero.

$\pi \leftarrow \epsilon\text{-greedy}(Q)$

$\rightarrow$ Choose an initial policy

s_0, A_0, R_1

$\rightarrow$ interact with environment; take action A_0
 $\rightarrow$ obtain the reward.

TD Control (Algorithm)

$$Q(s_t, a_t) \leftarrow Q(s_t, a_t) + \alpha \left(R_{t+1} + \gamma Q(s_{t+1}, A_{t+1}) - Q(s_t, A_t) \right)$$

Initialize : $Q(s, a) \forall s \in S, a \in A$. $\rightarrow$ initialize with all action values to zero.

$\pi \leftarrow \epsilon\text{-greedy}(Q)$

$\rightarrow$ Choose an initial policy

S_0, A_0, R_1, S_1, A_1

$\rightarrow$ next state is generated

$\rightarrow$ Use same π to generate next action A_1

TD Control (Algorithm)

$$Q(s_t, a_t) \leftarrow Q(s_t, a_t) + \alpha \left(\underbrace{R_{t+1} + \gamma Q(s_{t+1}, a_{t+1})}_{\text{TD target}} - \underbrace{Q(s_t, a_t)}_{\text{current estimate}} \right)$$

Initialize : $Q(s, a) \forall s \in S, a \in A$. $\rightarrow$ initialize with all action-values to zero.

$\pi \leftarrow \epsilon\text{-greedy}(Q)$

$| S_0 \ A_0 \ R_1 \ S_1 \ A_1 \ | \ R_2 \ S_2 \ A_2 \ \dots \ S_t \ A_t \ R_{t+1} \ S_{t+1} \ A_{t+1}$

Update $Q(s_0, a_0)$: $Q(s_0, a_0) \leftarrow Q(s_0, a_0) + \alpha (R_1 + \gamma Q(s_1, a_1) - Q(s_0, a_0))$
(update the value policy-evaluation)

$\pi \leftarrow \epsilon\text{-greedy}(Q)$ $\rightarrow$ choose $a = \arg\max_{a \in A} Q(s_0, a_0)$ $\rightarrow$ policy improvement.

TD Control (Algorithm)

$$Q(s_t, a_t) \leftarrow Q(s_t, a_t) + \alpha \left(\underbrace{R_{t+1} + \gamma Q(s_{t+1}, A_{t+1})}_{\text{TD target}} - \underbrace{Q(s_t, A_t)}_{\text{current estimate}} \right)$$

Initialize : $Q(s, a) \forall s \in S, a \in A$.

$\pi \leftarrow \epsilon\text{-greedy}(Q)$

$| S_0 \ A_0 \ R_1 \ S_1 \ A_1 \ | \ R_2 \ S_2 \ A_2 \ | \ \dots \ S_t \ A_t \ R_{t+1} \ S_{t+1} \ A_{t+1}$

Update $Q(S_0, A_0)$

$\pi \leftarrow \epsilon\text{-greedy}(Q)$

Update $Q(S_1, A_1)$
 $\pi \leftarrow \epsilon\text{-greedy}(Q)$

TD Control (Algorithm) SARSA (on-policy)

$$Q(s_t, a_t) \leftarrow Q(s_t, a_t) + \alpha \left(\underbrace{R_{t+1} + \gamma Q(s_{t+1}, a_{t+1})}_{\text{TD target}} - \underbrace{Q(s_t, a_t)}_{\text{current estimate}} \right)$$

Initialize : $Q(s, a) \forall s \in S, a \in A$.

$\pi \leftarrow \epsilon\text{-greedy}(Q)$

$| S_0 \ A_0 \ R_1 \ S_1 \ A_1 \ | \ R_2 \ S_2 \ A_2 \ | \ \dots \ S_t \ A_t \ R_{t+1} \ S_{t+1} \ A_{t+1} \ |$

Update $Q(S_0, A_0)$

$\pi \leftarrow \epsilon\text{-greedy}(Q)$

Update $Q(S_1, A_1)$
 $\pi \leftarrow \epsilon\text{-greedy}(Q)$

Update $Q(S_t, A_t)$
 $\pi \leftarrow \epsilon\text{-greedy}(Q)$

SARSA

off-policy Learning

Evaluate target policy $\pi(a|s)$ to compute $v_\pi(s)$ or $q_\pi(s,a)$
while following another policy $\mu(a|s)$

Advantages:

- ① Learn about optimal policy while following exploratory policy.
- ② Learn about multiple policies while following one policy.
- ③ Re-use experience generated from old policies $\pi_1, \pi_2, \pi_3, \dots$
- ④ Learn by observing other agents.

TD Control (off-policy)

Q-Learning

→ Off-policy learning of $Q(s,a)$ [Q-function]

→ Next action is chosen using a behaviour policy $A_{t+1} \sim p(\cdot | s_t)$

→ But, best successor action chosen $a' \sim \pi(\cdot | s_t)$

→ Update $Q(s_t, A_t)$ towards value given by alternative action!

$$Q(s_t, A_t) \leftarrow Q(s_t, A_t) + \alpha \left(\underbrace{R_{t+1} + \gamma Q(s_{t+1}, a')}_{\text{best alternative action}} - Q(s_t, A_t) \right)$$

best alternative
action

Q-learning (Algorithm).

→ Consider two policies $\left\{ \begin{array}{l} \text{Behaviour} \\ \text{Target} \end{array} \right.$ $\rightarrow$ used for exploration $\rightarrow$ that is learnt $\rightarrow$ improve!!

→ Behaviour policy: ϵ -greedy wrt $Q(s,a)$
 μ $\left\{ \begin{array}{l} \text{explores} \\ \text{exploits} \end{array} \right.$

→ Target policy: greedy wrt $Q(s,a)$
 π always selects 'max'

→ Q-learning target becomes:

$$\pi(S_{t+1}) = \arg \max_a Q(S_{t+1}, a)$$

$$R_{t+1} + \gamma Q(S_{t+1}, A')$$

$$= R_{t+1} + \gamma Q(S_{t+1}, \arg \max_{a' \in A} Q(S_{t+1}, a'))$$

$$= R_{t+1} + \max_{a' \in A} \gamma Q(S_{t+1}, a')$$

SARSA (on-policy) Reminder

Initialize $Q(s,a) = 0$ for all $s \in S, a \in A$

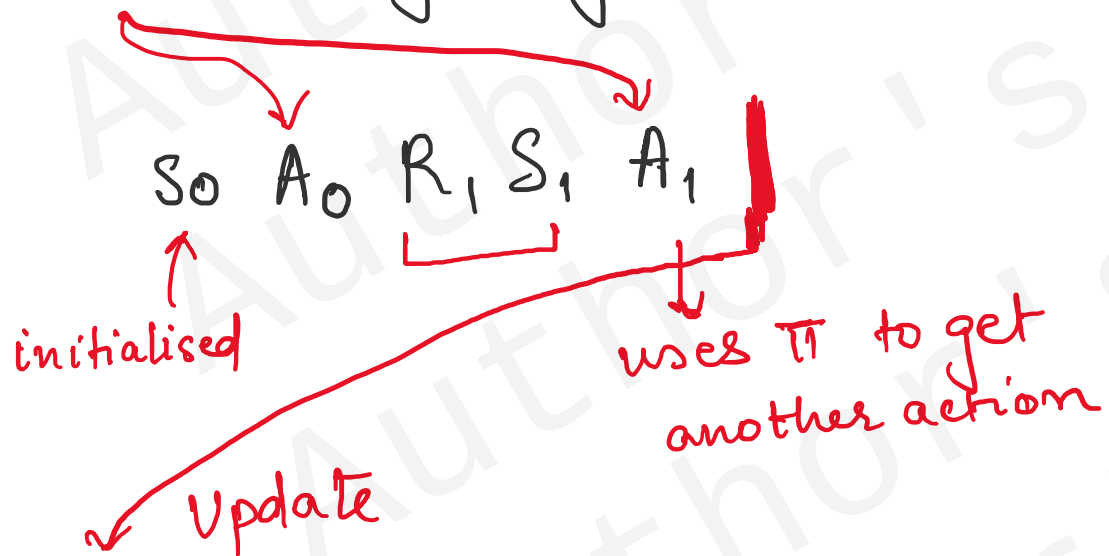
$\pi \leftarrow \epsilon\text{-greedy}(Q)$



SARSA (on-policy) Reminder

Initialize $Q(s,a) = 0$ for all $s \in S, a \in A$

$\pi \leftarrow \epsilon\text{-greedy}(Q)$



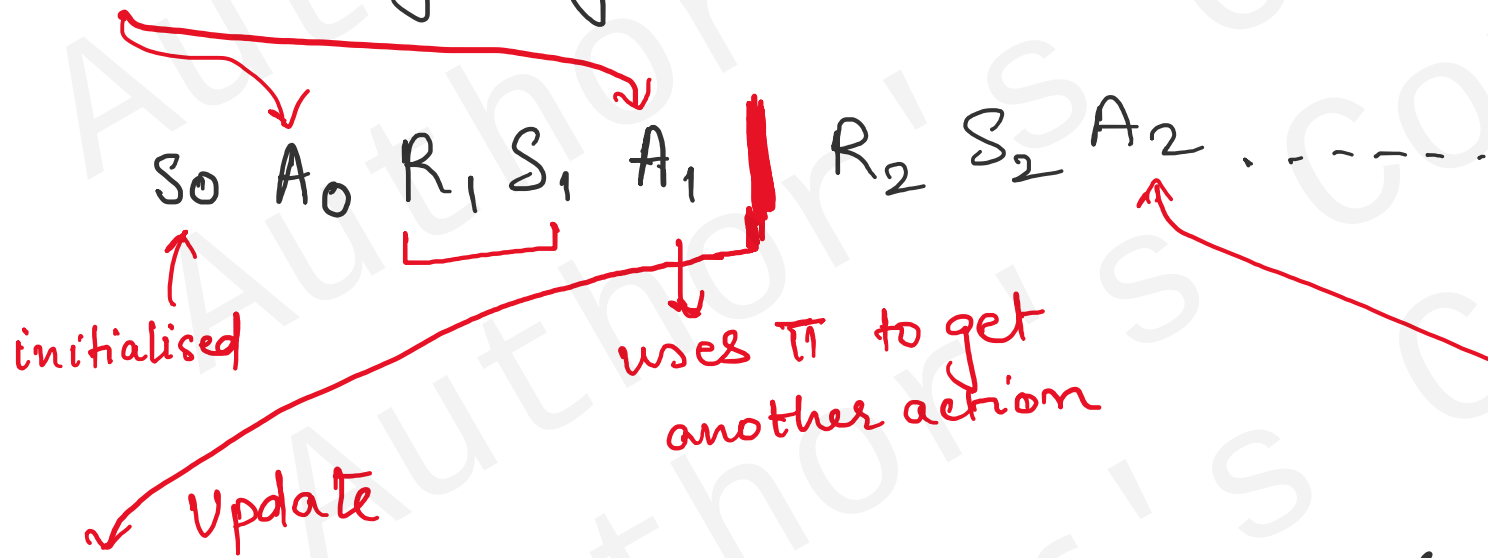
$$Q(S_0, A_0) \leftarrow Q(S_0, A_0) + \alpha (R_1 + \gamma Q(S_1, A_1) - Q(S_0, A_0)) :$$

and improve to most recent update
 $\pi \leftarrow \epsilon\text{-greedy}(Q)$

SARSA (on-policy) Reminder

Initialize $Q(s,a) = 0$ for all $s \in S, a \in A$

$\pi \leftarrow \epsilon\text{-greedy}(Q)$

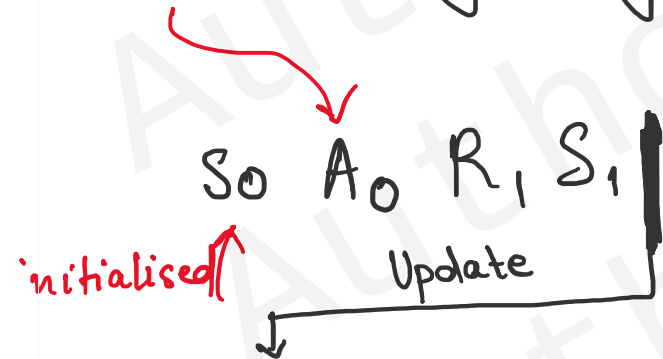


$Q(S_0, A_0) \leftarrow Q(S_0, A_0) + \alpha (R_1 + \gamma Q(S_1, A_1) - Q(S_0, A_0))$
and improve to most recent update
 $\pi \leftarrow \epsilon\text{-greedy}(Q)$

Q-learning = SARSA (max)

Initialize $Q(s, a) = 0$ for all $s \in S, a \in A$

$\pi \leftarrow \epsilon\text{-greedy}(Q)$



$$Q(S_0, A_0) \leftarrow Q(S_0, A_0) + \alpha (R_1 + \gamma Q(S_1, \underline{\quad}) - Q(S_0, A_0))$$

① What action to use here?

② SARSA uses $\epsilon\text{-greedy}$

③ Here, we could maximise directly
keep $\pi \sim \arg \max_{a \in A} (Q(s, a))$

ie. use greedy policy.

Q-learning = SARSA (max)

Initialize $Q(s, a) = 0$ for all $s \in S, a \in A$

$\pi \leftarrow \epsilon\text{-greedy}(Q)$

S_0, A_0, R, S_1

initialised
Update

$$Q(S_0, A_0) \leftarrow Q(S_0, A_0) + \alpha (R_1 + \gamma \underbrace{Q(S_1, \arg \max_{a' \in A} Q(S_1, a'))}_{\text{keep } \pi \sim \arg \max_{a' \in A} Q(S_1, a')}} - Q(S_0, A_0))$$

keep $\pi \sim \arg \max_{a' \in A} Q(S_1, a')$
ie. use greedy policy.

Q-learning = SARSA (max)

Initialize $Q(s, a) = 0$ for all $s \in S, a \in A$

$\pi \leftarrow \epsilon\text{-greedy}(Q)$

S_0, A_0, R, S_1

initialised $\uparrow$
Update $\downarrow$

$$Q(S_0, A_0) \leftarrow Q(S_0, A_0) + \alpha \left(R_1 + \gamma \max_{a \in A} Q(S_1, a) - Q(S_0, A_0) \right)$$

$\uparrow$
keep $\pi \sim \arg \max_{a \in A} (Q(s, a))$
ie. use greedy policy.

Q-learning = SARSA (max)

Initialize $Q(s, a) = 0$ for all $s \in S, a \in A$

$\pi \leftarrow \epsilon\text{-greedy}(Q)$

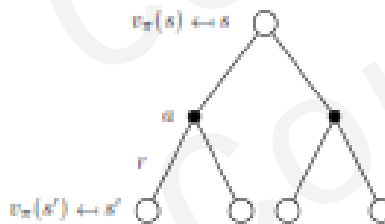
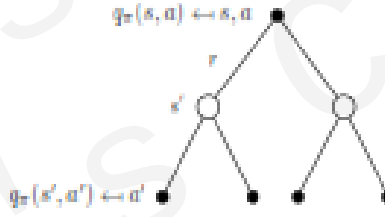
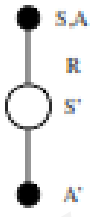
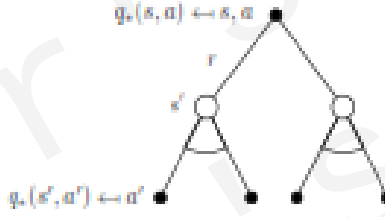
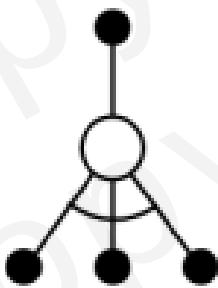
$S_0, A_0, R_1, S_1 \mid A_1, R_1, S_1$

initialised $\downarrow$ Update $\downarrow$

$$Q(S_0, A_0) \leftarrow Q(S_0, A_0) + \alpha \left(R_1 + \gamma \max_{a \in A} Q(S_1, a) - Q(S_0, A_0) \right)$$

$\pi \leftarrow \epsilon\text{-greedy}(\text{updated } Q)$

Q-learning $\rightarrow$ attempts to directly estimate the optimal value at every time step !!
(max)

	Full Backup (DP)	Sample Backup (TD)
Bellman Expectation Equation for $v_{\pi}(s)$	 Iterative Policy Evaluation	 TD Learning
Bellman Expectation Equation for $q_{\pi}(s, a)$	 Q-Policy Iteration	 Sarsa
Bellman Optimality Equation for $q_{*}(s, a)$	 Q-Value Iteration	 Q-Learning

Summary

Full Back up DP

Iterative Policy Evaluation

$$V(s) \leftarrow \mathbb{E} (R + \gamma V(s') | s)$$

Q- Policy Iteration

$$Q(s, a) \leftarrow \mathbb{E} [R + \gamma Q(s', A') | s, a]$$

Q-value Iteration

$$Q(s, a) \leftarrow \mathbb{E} [R + \gamma \max_{a' \in A} Q(s', a') | s, a]$$

Sample Back-Ups

TD learning

$$V(s) \leftarrow V(s) + \alpha (R + \gamma V(s') - V(s))$$

SARSA

$$Q(s, A) \leftarrow Q(s, A) + \alpha [R + \gamma Q(s', A') - Q(s, A)]$$

Q-learning

$$Q(s, A) \leftarrow Q(s, A) + \alpha$$

