

Value Based - Part 2

Deep Reinforcement Learning

Mohammad Hossein Rohban, Ph.D.

Courtesy: Some slides are adopted from CS 285 Berkeley, and CS 234 Stanford, and Pieter Abbeel's compact series on RL.



RIML Lab, CE Department, Sharif
University of Technology

Spring 2026

Disadvantages of Monte-Carlo Learning

- ♦ We have seen MC algorithms can be used to learn value predictions
- ♦ But when episodes are long, learning can be slow
 - ☑ we have to wait until an episode ends before we can learn...
 - ☑ return can have high variance
 - Which one is more? First-visit or every-visit
- ♦ Are there alternatives? (Spoiler: yes)

$G_{i,t}$'s



Monte Carlo Control

Policy
Eval.



Policy
Improvement

Monte-Carlo Control

Algorithm:

Repeat:

→ Sample **episode** $1, \dots, k, \dots$, using $\pi: \{S_1, A_1, R_2, \dots, S_T\} \sim \pi$

For each state S_t and action A_t in the episode:

$$q(S_t, A_t) \leftarrow q(S_t, A_t) + \alpha_t (\underline{G_t} - \underline{q(S_t, A_t)})$$

e.g.,

$$\alpha_t = \frac{1}{N(S_t, A_t)} \quad \text{of} \quad \alpha_t = 1/k$$

→ Improve **policy** based on new action-value function

$$\rightarrow \underline{\pi^{new}(s)} = \operatorname{argmax}_{a \in A} \underline{q(s, a)}$$

policy
eval

policy
imp.

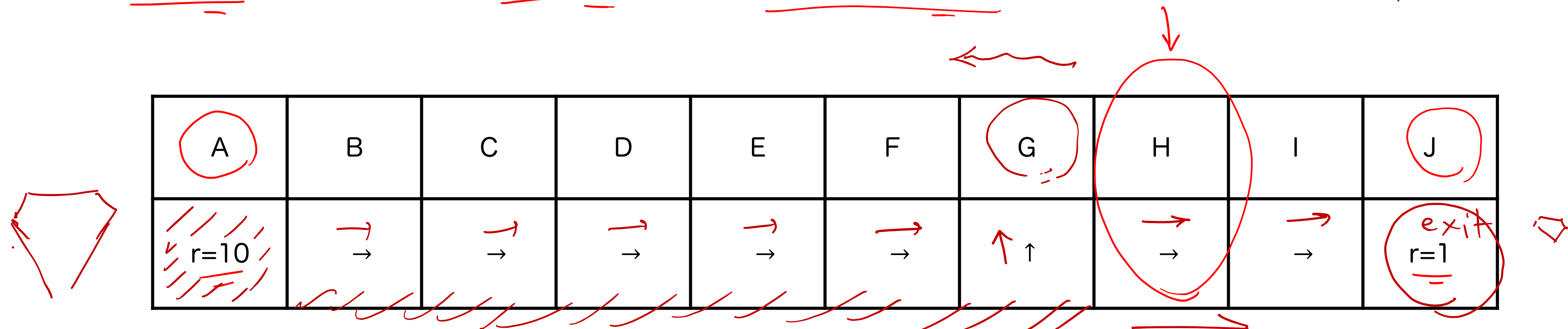
Any issues?

$$V_{(s)}^{new} \leftarrow \underbrace{V_{(s)}}_0 + \underbrace{\alpha_t}_{0.1} (\underbrace{G_t}_1 - \underbrace{V_{(s)}}_0)$$

♦ Let's consider this example:

♦ Discount = 1, learning ^{α_t} rate = 0.1, start in state H.

$$0.9 \times 0.1$$



episode 1 $\Rightarrow$ H, $\rightarrow$, I, ϕ , $\rightarrow$, J, 1

episode 2 $\Rightarrow$ "

exploration

$$\left. \begin{aligned} Q(H, \rightarrow) &= 0.1 \rightsquigarrow 0.19 \\ Q(I, \rightarrow) &= 0.1 \rightsquigarrow 0.19 \\ Q(J, \text{exit}) &= 0.1 \rightsquigarrow 0.19 \end{aligned} \right\} \begin{array}{l} \nearrow \\ \nearrow \\ \nearrow \end{array} \text{exploit}$$

Epsilon Greedy Policy

for episode generation

- ♦ Simple idea to balance exploration and achieving rewards
- ♦ Let $|A|$ be the number of actions
- ♦ Then an ϵ -greedy policy w.r.t a state action value $Q(s, a)$ is

$$\pi(a|s) =$$

☒ $\arg\max_a Q(s, a)$, w. Prob $1 - \epsilon + \frac{\epsilon}{|A|}$ Action Set Size

☒ $a' \neq \arg\max_a Q(s, a)$ w. Prob $\frac{\epsilon}{|A|}$ Risk Taking
Explorative

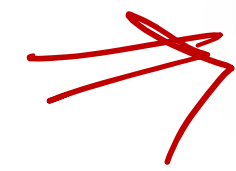
Does this hurt improvement?

$$V^{\pi^*}(s) \geq V^{\pi}(s)$$

$$N \rightarrow \infty$$

$$\forall \pi$$

Theorem



For any ϵ -greedy policy π_i , the ϵ -greedy policy w.r.t. $\underline{Q^{\pi_i}}$, π_{i+1} is a monotonic improvement
 $V^{\pi_{i+1}} \geq V^{\pi_i}$

$i = \text{itr index}$

$$\underline{\pi_{i+1}}(a|s) = \begin{cases} \arg\max_a Q^{\pi_i}(s, a) \\ \text{random act.} \end{cases}$$

$$Pr = 1 - \epsilon + \frac{\epsilon}{|A|}$$

$$Pr = \frac{\epsilon}{|A|}$$

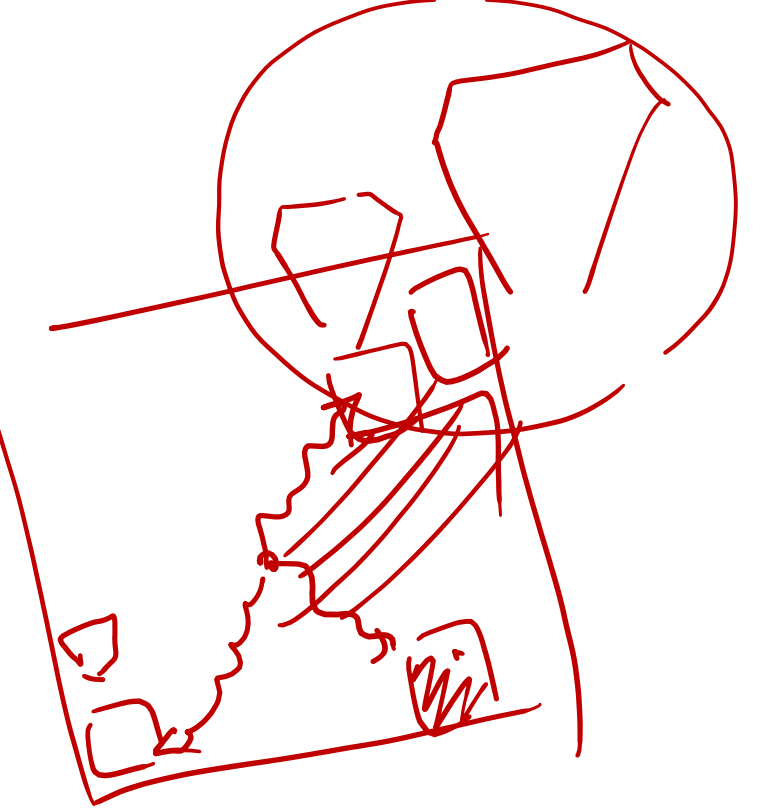
$$\forall s: V^{\pi_{i+1}}(s) \geq V^{\pi_i}(s)$$

Does this hurt improvement?

Monte-Carlo Control (done right)

$$\epsilon = 0.9$$

$$\epsilon = 0.6$$



to-the-point

efficient exploration

Algorithm:

Repeat:

Sample episode $1, \dots, k, \dots$, using $\pi: \{S_1, A_1, R_2, \dots, S_T\} \sim \pi$ stochastic

For each state S_t and action A_t in the episode:

$$q(S_t, A_t) \leftarrow q(S_t, A_t) + \alpha_t (G_t - q(S_t, A_t))$$

e.g.,

$$\alpha_t = \frac{1}{N(S_t, A_t)}$$

of $\alpha_t = 1/k$

Improve policy based on new action-value function

$$\pi(a|s) =$$

$$\operatorname{argmax}_a Q(s, a), \text{ w. Prob } 1 - \epsilon + \frac{\epsilon}{|A|}$$

$$a' \neq \operatorname{argmax}_a Q(s, a) \text{ w. Prob } \frac{\epsilon}{|A|}$$

$q(S, A)$

Disadvantages of MC Learning

- ♦ We have seen MC algorithms can be used to learn value predictions
- ♦ But when episodes are long, learning can be slow

- ☑ ...we have to wait until an episode ends before we can learn

- ☑ ...return can have high variance

→ Slower training

- ♦ Are there alternatives? (Yes)

takes longer (more episodes) for π

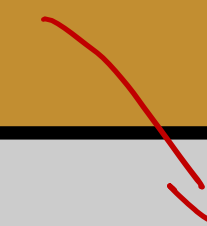
to converge to π^*

Temporal Difference Learning

Prediction & Control



V or Q



policy Imp.

TD Overview

- ✦ TD methods learn directly from episodes of experience
- ✦ TD is model-free: no knowledge of MDP transitions / rewards
- ✦ TD learns from incomplete episodes, by bootstrapping
- ✦ TD updates a guess towards a guess

TD Learning by Sampling Bellman Equations

- ♦ Bellman update equations:

$$\Rightarrow v_{k+1}(s) = \mathbb{E}[R_{t+1} + \gamma v_k(S_{t+1}) \mid S_t = s, A_t \sim \pi(S_t)]$$

- ♦ We can sample this!

↳ empirical Mean

$$v_{t+1}(S_t) = R_{t+1} + \gamma v_t(S_{t+1})$$

- ♦ Samples could be averaged, in a similar way to MC:

$$v_{t+1}(S_t) \leftarrow v_t(S_t) + \alpha \left(\underbrace{R_{t+1} + \gamma v_t(S_{t+1})}_{\text{target}} - \underbrace{v_t(S_t)}_{\text{Bootstrap}} \right)$$

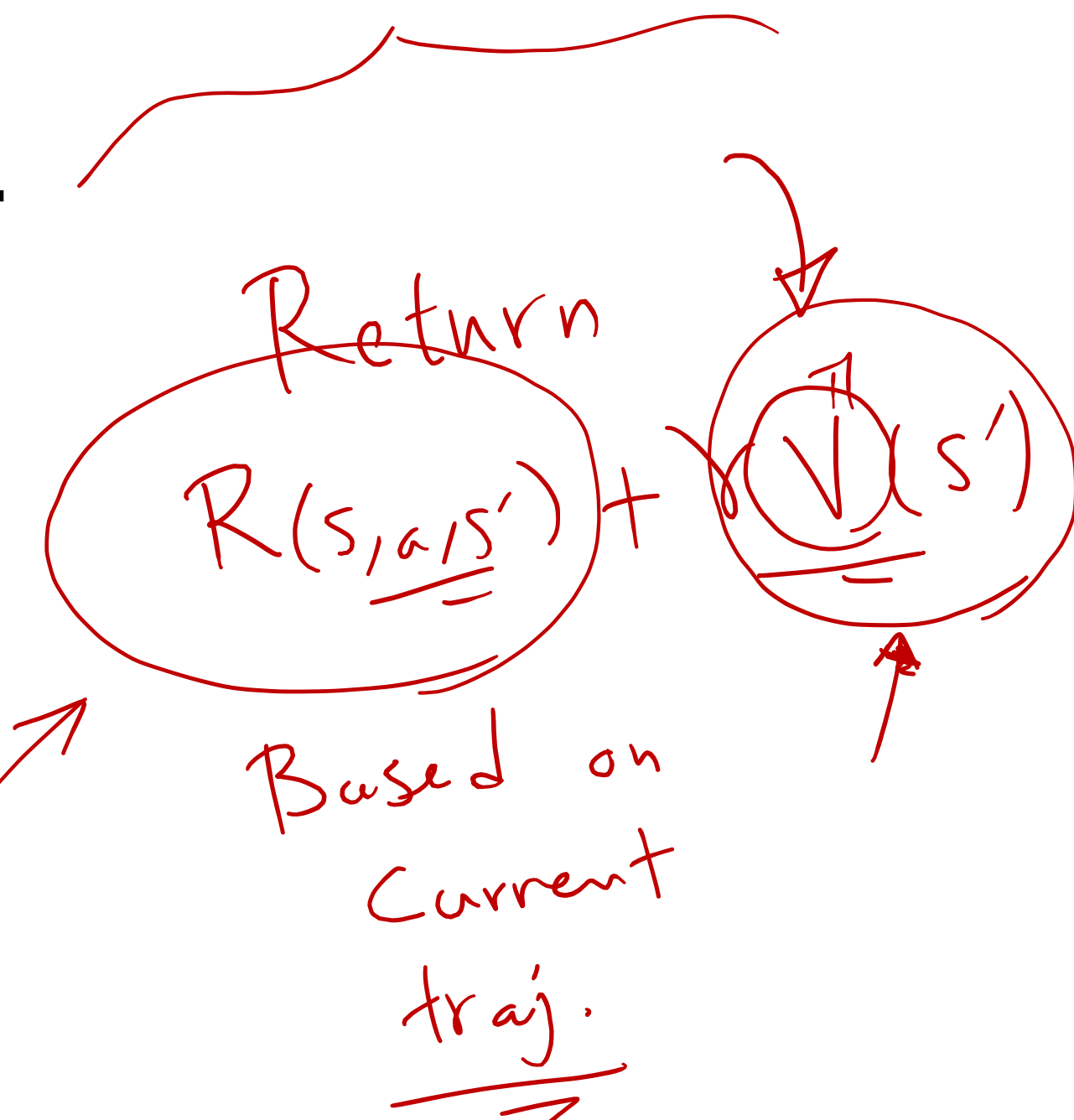
SARSA

Q-Learning

On-

off-policy
Learning

temporal difference error δ_t



- * Faster Sample Formation ✓
- * Lower Sample Variance ✓

Temporal Difference Learning

♦ Prediction setting: learn v_π online from experience under policy π

♦ Monte Carlo

☑ Update value $v_n(S_t)$ towards sampled return G_t

$$v_{n+1}(S_t) = v_n(S_t) + \alpha(G_t - v_n(S_t))$$

♦ TD Learning

☑ Update value $V_t(S_t)$ towards estimated return $R_{t+1} + \gamma v(S_{t+1})$

The diagram shows the TD learning update equation with several handwritten annotations in red:

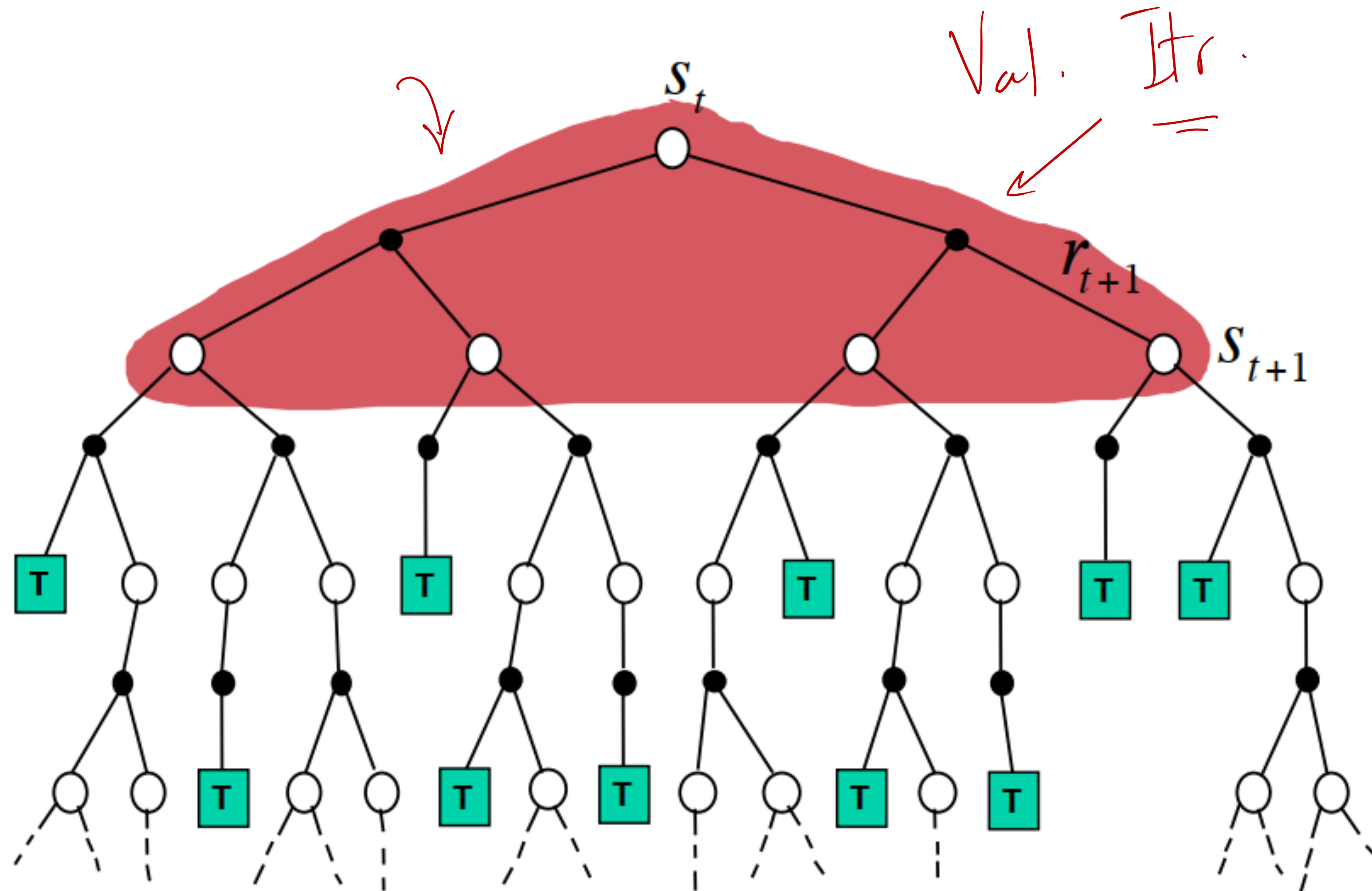
$$v_{t+1}(S_t) \leftarrow v_t(S_t) + \alpha \left(\underbrace{R_{t+1} + \gamma v_t(S_{t+1})}_{\text{target}} - \underbrace{v_t(S_t)}_{\text{Current est. of } V(S_t)} \right)$$

Annotations include:

- A red arrow pointing to the left side of the equation ($v_{t+1}(S_t)$).
- A red arrow pointing to the term R_{t+1} in the target.
- A red arrow pointing to the term $\gamma v_t(S_{t+1})$ in the target.
- A red circle around the expression $R_{t+1} + \gamma v_t(S_{t+1})$ with the label "TD error" above it.
- A red arrow pointing to the term $v_t(S_t)$ with the label "pred. of $v(S_t)$ " below it.
- A red arrow pointing to the term $v_t(S_t)$ with the label "Current est. of $V(S_t)$ " to the right.

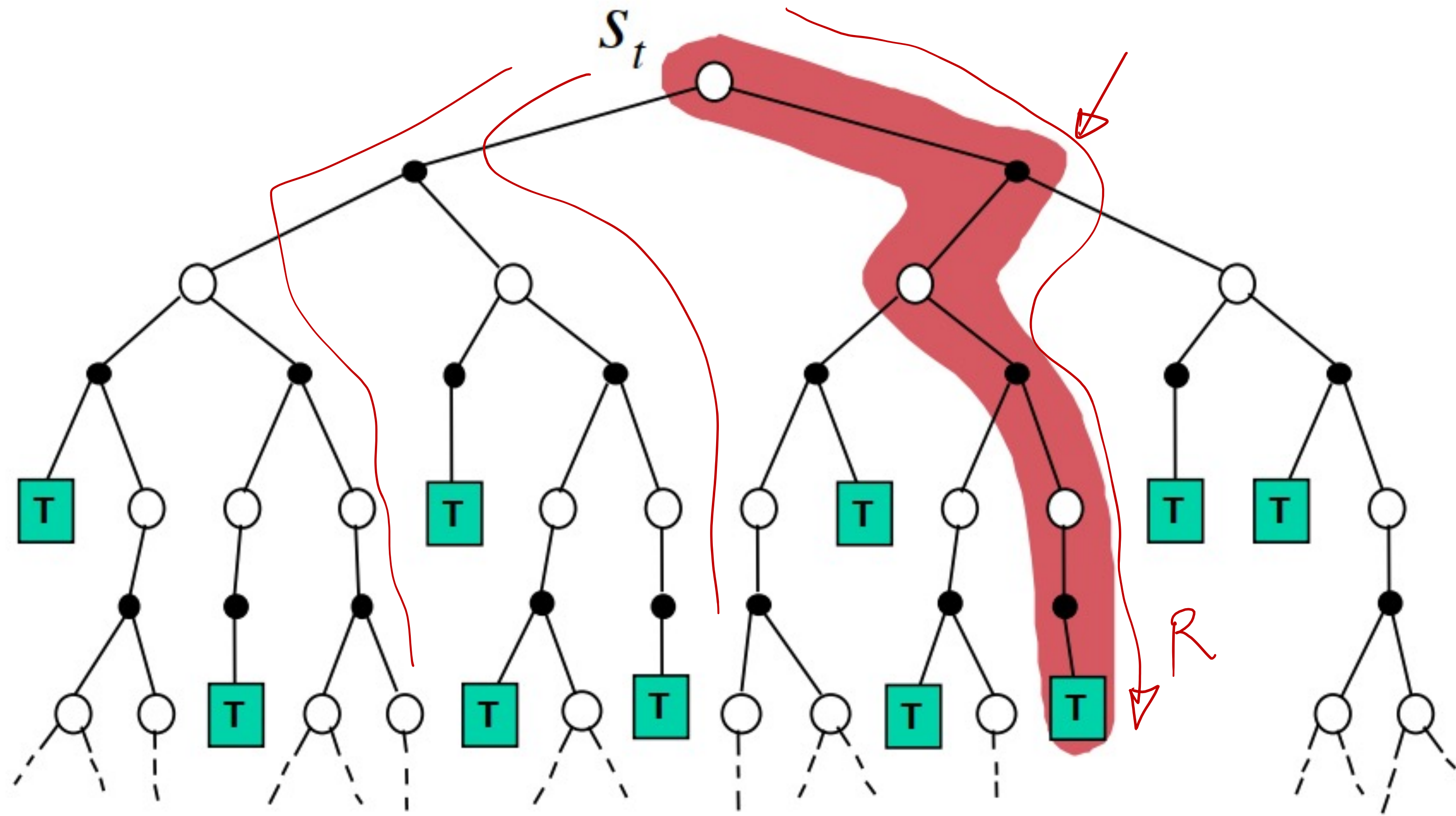
Backup (Dynamic Programming)

$$v(S_t) \leftarrow \mathbb{E}[R_{t+1} + \gamma v(S_{t+1}) \mid A_t \sim \pi(S_t)]$$



Backup (Monte Carlo)

$$\underline{v(S_t)} \leftarrow v(S_t) + \alpha(\underline{G_t} - v(S_t))$$



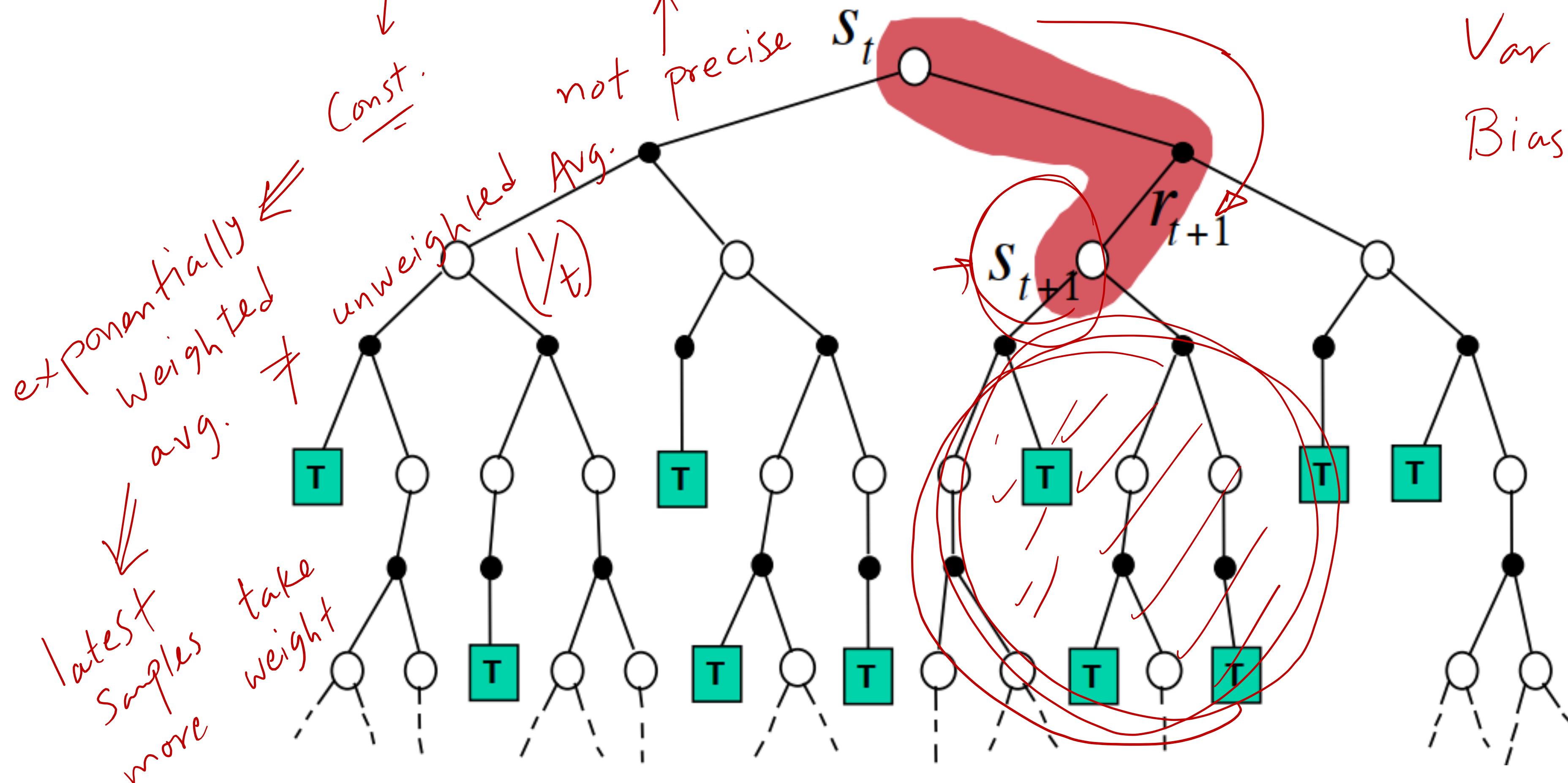
Backup (Temporal Difference)

$$v(S_t) \leftarrow v(S_t) + \alpha(R_{t+1} + \gamma v(S_{t+1}) - v(S_t))$$

$$X_{n+1} = X_n + \alpha [z_{n+1} - X_n]$$

$$X_{n+1} = \alpha z_{n+1} + \alpha^2 z_n + \alpha^3 z_{n-1} + \dots$$

Var ↓
Bias ↑



Bootstrapping and Sampling

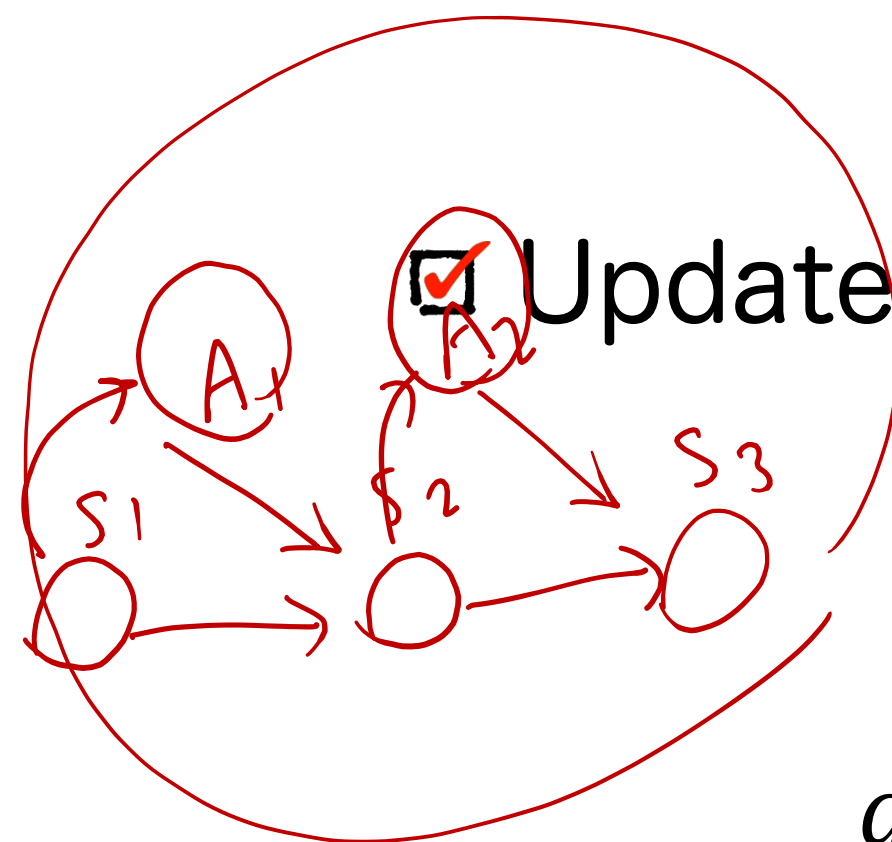
- ♦ Bootstrapping: update involves an estimate
 - ☑ MC does not bootstrap ✗
 - ☑ DP bootstraps ✓
 - ☑ TD bootstraps ✓
- ♦ Sampling: update samples an expectation
 - ☑ MC samples ✓ ✓
 - ☑ DP does not sample ✗
 - ☑ TD samples ✓ ✓

$$G_t = R_{t+1} + \gamma \underbrace{V_t(s_{t+1})}_{\text{Current est.,}}$$

TD Learning for action values

- ♦ We can apply the same idea to action values
- ♦ Temporal-difference learning for action values:

SARSA
 ==
 ↓
 S' next state next Action A'



Update value $q_t(S_t, A_t)$ towards estimated return $R_{t+1} + \gamma q(S_{t+1}, A_{t+1})$

$$q_{t+1}(S_t, A_t) \leftarrow q_t(S_t, A_t) + \alpha \left(\underbrace{R_{t+1} + \gamma q_t(S_{t+1}, A_{t+1})}_{\text{target}} - q_t(S_t, A_t) \right)$$

TD error

MDP
 ↓
 $S' \sim P(S_{t+1} | S, A)$
 $A' \sim \pi(A' | S')$

(why?)

From the current policy π

also requires A_{t+1}

← From what distribution A_{t+1} was sampled?

TD vs. MC

- ♦ TD can learn before knowing the final outcome
 - ✓ TD can learn online after every step
 - ✓ MC must wait until end of episode before return is known
- ♦ TD can learn without the final outcome
 - ✓ MC must wait until end of episode before return is known
 - ✓ MC can only learn from complete sequences
 - ✓ TD works in continuing (non-terminating) environments
 - ✓ MC only works for episodic (terminating) environments
- ♦ TD is independent of the temporal span of the prediction
 - ✓ TD can learn from single transitions
 - ✓ MC must store all predictions (or states) to update at the end of an episode
- ♦ TD needs reasonable value estimates

(S, A, S, R)

$(S_1, A_1, S_2, \dots)$

Temporal Difference Learning

Control

SARSA Algorithm for On-Policy Control

final learned Policy

$[TD(n), TD(\lambda)]$

π

$$\pi(a|s) = \underset{a}{\operatorname{argmax}} Q(s, a)$$

Algorithm:

Initialize $Q(s, a), \forall s \in \mathcal{S}, a \in \mathcal{A}(s)$, arbitrarily, and $Q(\text{terminal state}, \cdot) = 0$

Repeat (for each episode):

Initialize S

Choose A from S using policy derived from Q (e.g., ϵ -greedy)

Repeat (for each step of episode):

Take action A , observe R, S'

Choose A' from S' using policy derived from Q (e.g., ϵ -greedy)

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

$S \leftarrow S'; A \leftarrow A';$

until S is terminal

$$\epsilon = \epsilon_t$$

ϵ

$-t \leftarrow \text{itr.}$

$$R_{t+1} + \gamma R_{t+2} + \gamma^2 Q(S'', A'')$$

See the $\rightarrow$
diff with
MC.
Just need
one trans.

(promote exploration)

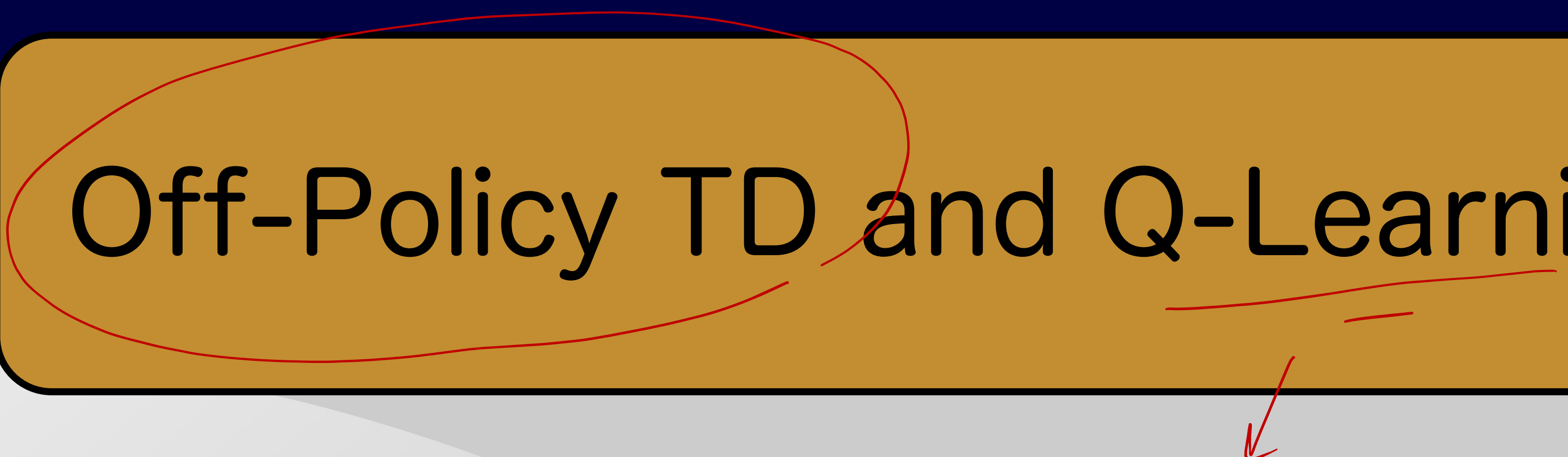
same policy

Deep RL
Approximation M.

Tabular RL

$$Q(s, a)$$

Off-Policy TD and Q-Learning



On and Off-Policy Learning

- ♦ On-policy learning
 - ☑ "Learn on the job"
 - ☑ Learn about policy π from experience sampled from π
- ♦ Off-policy learning
 - ☑ "Look over someone's shoulder"
 - ☑ Learn about policy π from experience sampled from μ

used to generate samples $\mu \neq \pi$ learned Policy

Off-Policy Learning

- ♦ Evaluate target policy $\pi(a|s)$ to compute $v_\pi(s)$ or $q_\pi(s, a)$
- ♦ While using behavior policy $\mu(a, s)$ to generate actions
- ♦ Why is this important?
 - ☑ Learn from observing humans or other agents (e.g., from logged data)
 - ☑ Re-use experience from old policies (e.g., from your own past experience)
 - ☑ Learn about multiple policies while following one policy
 - ☑ Learn about greedy policy while following exploratory policy

(Data Efficiency)

Q-Learning

→ learn the ^{action} $\sqrt{}$ values of the greedy policy

- ♦ Q-learning estimates the value of the greedy policy

$$q_{t+1}(\underline{s}, a) = q_t(S_t, A_t) + \alpha_t (R_{t+1} + \gamma \max_{a'} q_t(S_{t+1}, a') - q_t(S_t, A_t))$$

- ♦ Acting greedy all the time would not explore sufficiently

$$(\underline{s}, a, s', r_{t+1}) \sim \epsilon\text{-greedy} \leftarrow \text{data gathering}$$

Theorem

Q-learning control converges to the optimal action-value function, $q \rightarrow q^*$, as long as we take each action in each state infinitely often.

- ♦ Note: no need for greedy behavior!
- ♦ Works for any policy that **eventually selects all actions sufficiently often**

Q-Learning for Off-Policy Control

Optimal Bellman Eq.

$$q_{t+1}(s, A) = \mathbb{E} \left(R(s, A) + \gamma \max_{A'} q_t(s', A') \right)$$

Algorithm:

Initialize $Q(s, a), \forall s \in \mathcal{S}, a \in \mathcal{A}(s)$, arbitrarily, and $Q(\text{terminal-state}, \cdot) = 0$

Repeat (for each episode):

Initialize S

Repeat (for each step of episode):

Choose A from S using policy derived from Q (e.g., ϵ -greedy)

Take action A , observe R, S'

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

$S \leftarrow S'$;

until S is terminal

behavior policy

execution policy = greedy based on Q

Claim: Any other behavior

Policy should also work given that it's being random.

Q-Learning for Off-Policy Control

