

Value Based - Part 1

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

Value Function

traject. $\rightarrow \underline{s}_0, \underline{a}_0, \underline{s}_1, R(\underline{s}_0, \underline{a}_0, \underline{s}_1), \underline{a}_1, \underline{s}_2, \dots =$

♦ $V^*(s)$ = expected utility **starting in s** , and **acting optimally** in all subsequent actions.

Obj. $\downarrow$
 $\max_{a_0, a_1, \dots} \mathbb{E} \left\{ \sum_{t=0}^{\infty} \gamma^t R(\underline{s}_t, \underline{a}_t, \underline{s}_{t+1}) \right\}$
 $\downarrow$
 state trans. stoch. $\uparrow$

Return

$$V^*(\underline{s}) = \max_{\underline{a}_i | \underline{s}} \mathbb{E} \left(\sum_{t=0}^{\infty} \gamma^t R(\underline{s}_t, \underline{a}_t, \underline{s}_{t+1}) \mid \underline{s}_0 = \underline{s} \right)$$

discount factor ≤ 1
 Return

Goal: How to approx. val. function for all $\underline{s}$.
 $\pi : S \rightarrow A$

Value Iteration

Horizon $\left\{ \begin{array}{l} P(s_{t+1} | s_t, a_t) \\ R(s_t, a_t, s_{t+1}) \end{array} \right\}$ known

♦ $V_0^*(s)$ = optimal value for state s when $H=0$

$$\checkmark V_0^*(s) = 0 \forall s$$

♦ $V_1^*(s)$ = optimal value for state s when $H=1$

$$\checkmark V_1^*(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_0^*(s'))$$

♦ $V_2^*(s)$ = optimal value for state s when $H=2$

$$\checkmark V_2^*(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_1^*(s'))$$

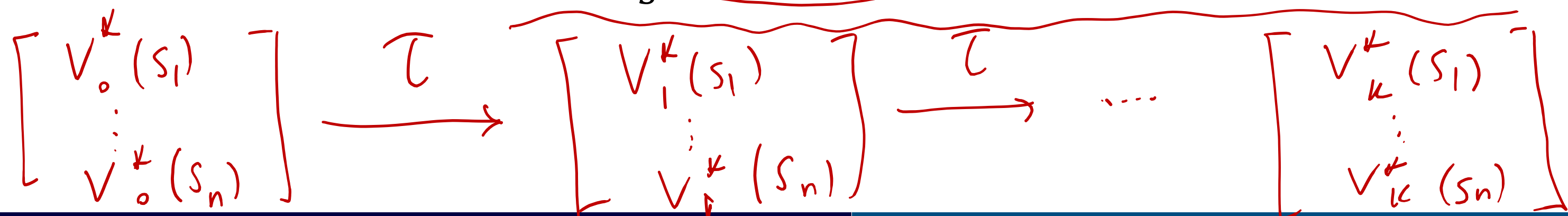
♦ $V_k^*(s)$ = optimal value for state s when $H = k$

$$\checkmark V_k^*(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_{k-1}^*(s'))$$

$V_k^*(s) :=$ max possible return in at most ∞ k steps starting @ state s .

$$V^k(s) = V_\infty^*(s)$$

* For the $V_k(s)$ to be optimal, the reward to go $(V_{k-1}(s'))$ should also be optimal.



Value Iteration

Algorithm:

Start with $V_0^*(s) = 0$ for all s .

For $k = 1, \dots, H$:

For all states s in S :

$$V_k^*(s) \leftarrow \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_{k-1}^*(s'))$$

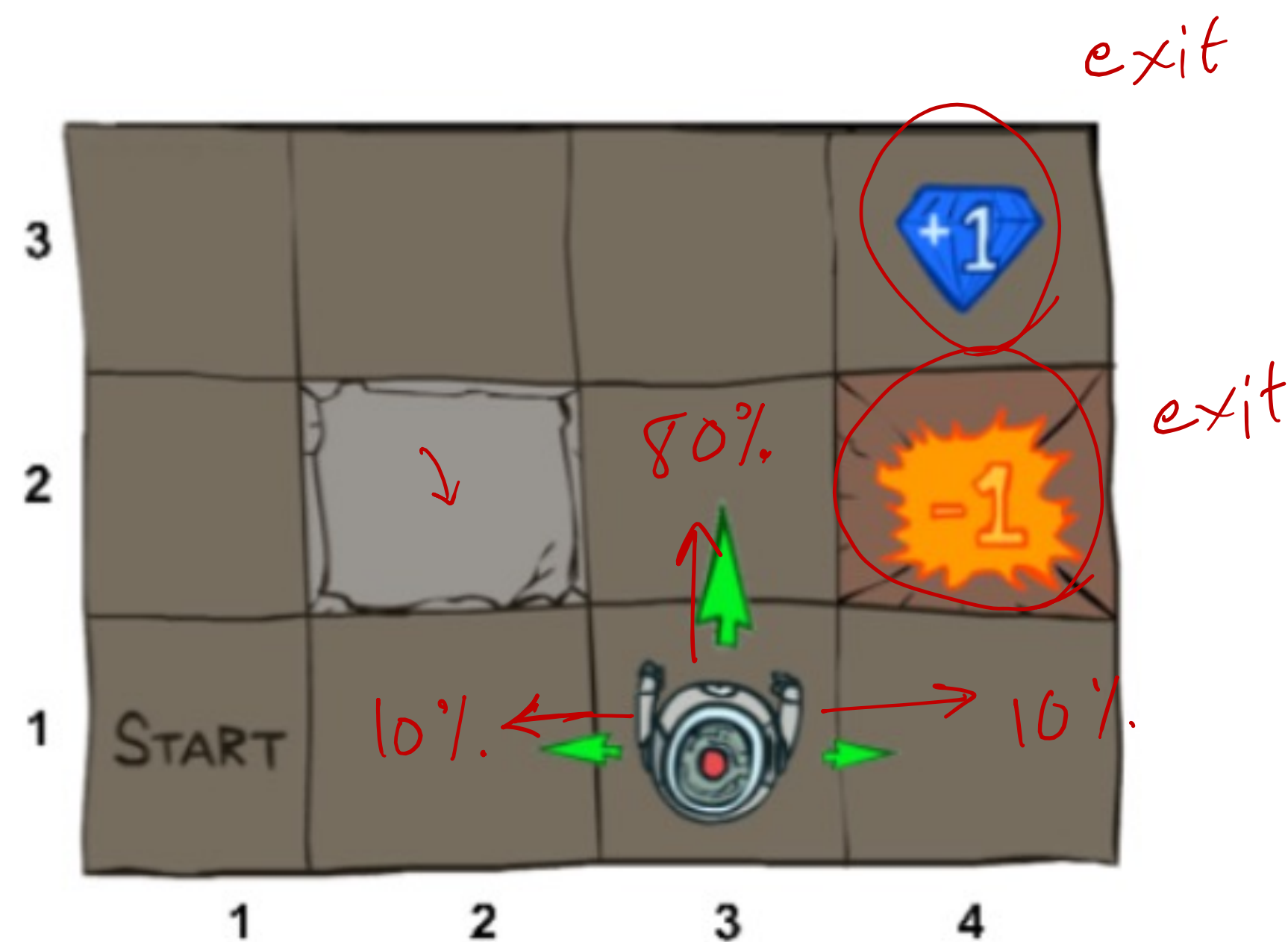
$$\pi_k^*(s) \leftarrow \operatorname{argmax}_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_{k-1}^*(s'))$$

✓ This is called a value update or Bellman update/back-up

⌊

Value Iteration

$$V_0(s) \leftarrow 0$$



Noise = 0.2

$\gamma =$ Discount = 0.9

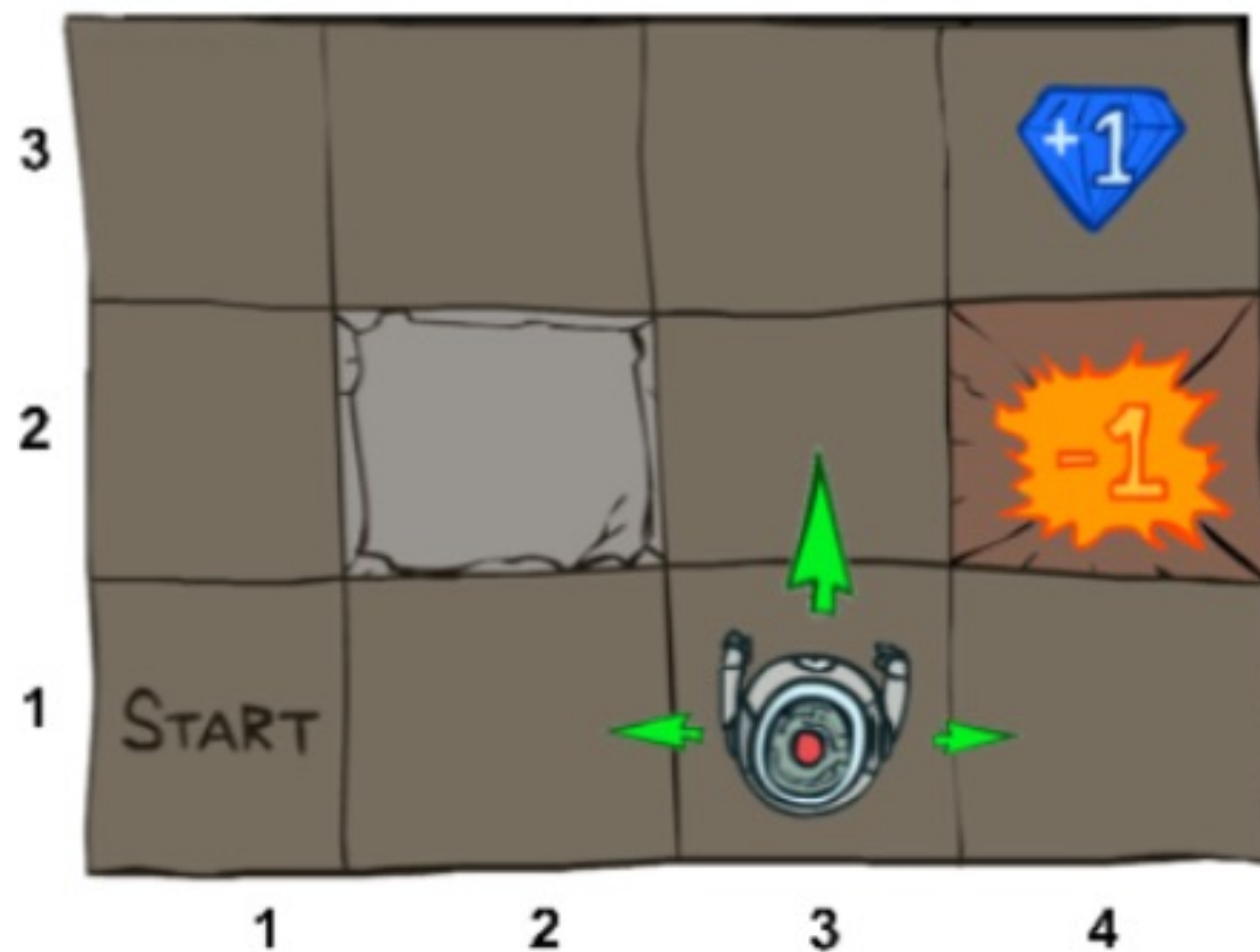
k = 0

0.00	0.00	0.00	0.00
0.00		0.00	0.00
0.00	0.00	0.00	0.00

VALUES AFTER 0 ITERATIONS

Value Iteration

$$V_1(s) = \max_a \sum_{s'} P(s'|s, a) (\underbrace{R(s, a, s')}_{k=0} + \gamma V_0(s'))$$



Noise = 0.2
Discount = 0.9

0.00	0.00	0.00	0.00
0.00		0.00	0.00
0.00	0.00	0.00	0.00

VALUES AFTER 0 ITERATIONS

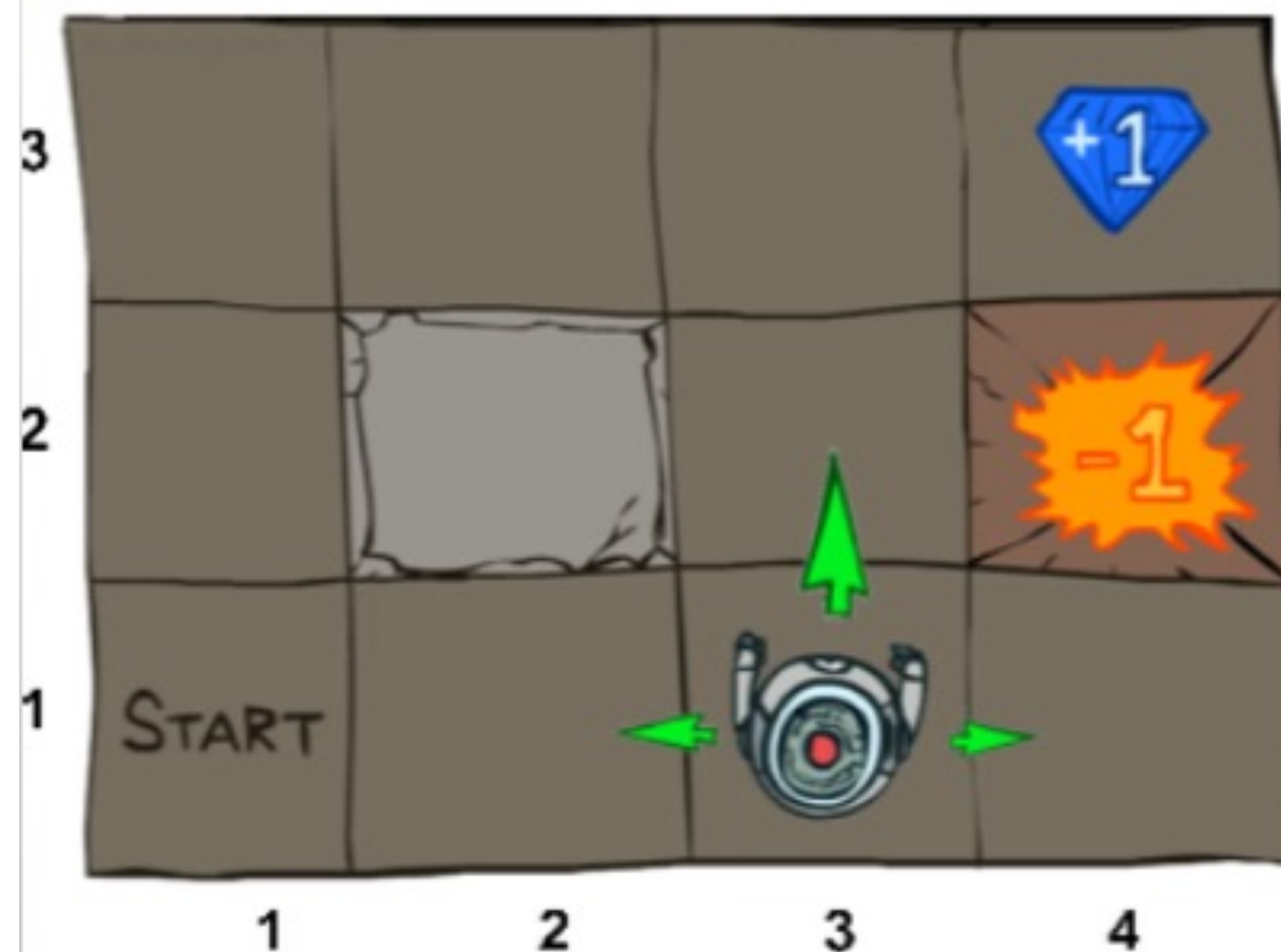
Value Iteration

$$V_2(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_1(s'))$$

$$0.80 [0 + 0.9 \times 1] + 0.1 [0 + 0.9 \times 0] + 0.1 [0 + 0.9 \times 0] = 0.72$$

$V_1^*(4, 1)$

k = 1



Noise = 0.2
Discount = 0.9

0.00	0.00	0.00	1.00
0.00		0.00	-1.00
0.00	0.00	0.00	0.00

VALUES AFTER 1 ITERATIONS

Value Iteration

$$V_2(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_0(s'))$$

k = 2

propagation



0.00	0.00	0.72	1.00
0.00		0.00	-1.00
0.00	0.00	0.00	0.00

VALUES AFTER 2 ITERATIONS

Noise = 0.2
Discount = 0.9

Value Iteration

$$V_{k+1}(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_0(s'))$$

$k = 3$



Noise = 0.2
Discount = 0.9

Value Iteration

$$V_{k+1}(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_0(s'))$$

k = 4



Noise = 0.2
Discount = 0.9

Value Iteration

$$V_{k+1}(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_0(s'))$$

$k = 5$



Noise = 0.2
Discount = 0.9

Value Iteration

$$V_{k+1}(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_0(s'))$$

k = 6



Noise = 0.2
Discount = 0.9

Value Iteration

$$V_{k+1}(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_0(s'))$$

$k = 7$



Noise = 0.2
Discount = 0.9

Value Iteration

$$V_{k+1}(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_0(s'))$$

k = 8



Noise = 0.2
Discount = 0.9

Value Iteration

$$V_{k+1}(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_0(s'))$$

$k = 9$



Noise = 0.2
Discount = 0.9

Value Iteration

$$V_{k+1}(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_0(s'))$$

$k = 10$



Noise = 0.2
Discount = 0.9

Value Iteration

$$V_{k+1}(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_0(s'))$$

k = 11



Noise = 0.2
Discount = 0.9

Value Iteration

$$V_{k+1}(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_0(s'))$$

k = 12

Almost
Converged



Noise = 0.2
Discount = 0.9

Value Iteration

$$V_{k+1}(s) = \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_0(s'))$$

k = 100



Noise = 0.2
Discount = 0.9

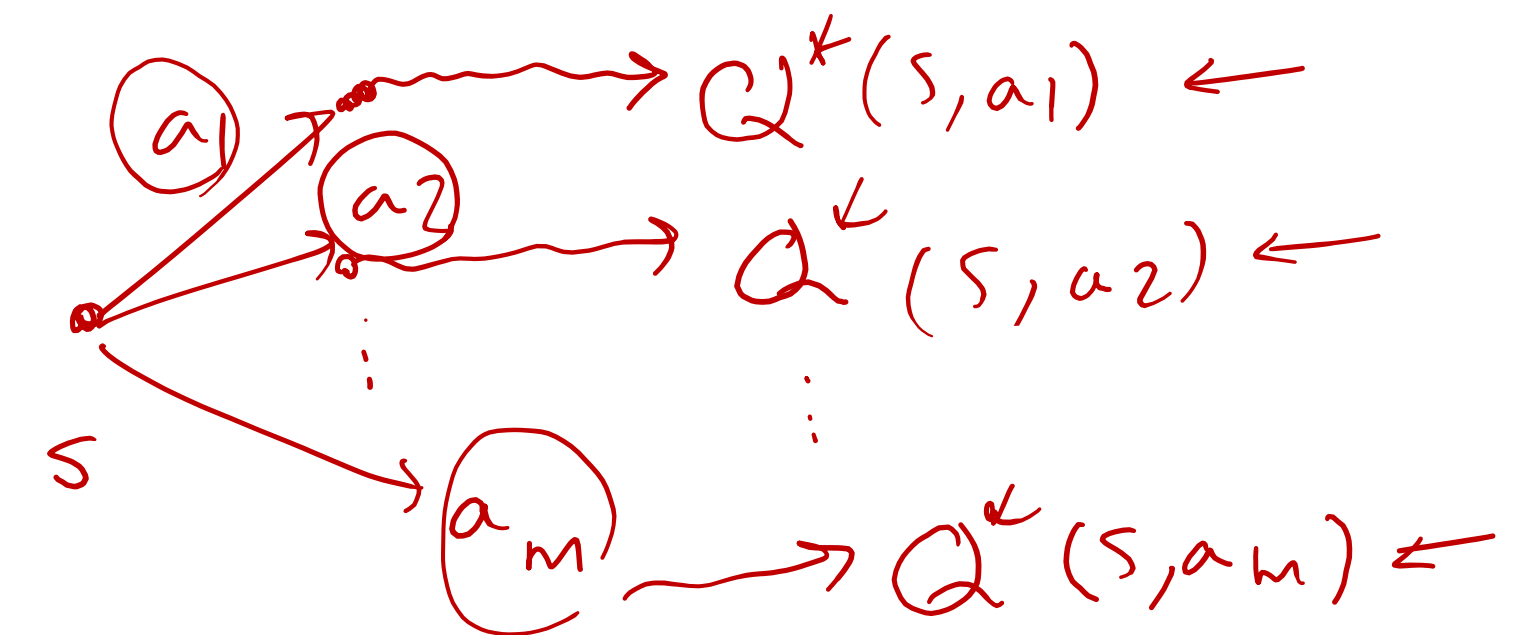
Q-Values

$$\underline{\underline{\pi^k(s)}} = \arg \max_a Q^k(s, a)$$

may be non-optimal

♦ $Q^*(s, a)$ = expected utility starting in s, taking action a, and (thereafter) acting optimally

$$V^*(s) = \max_{a'} Q^*(s, a')$$



♦ Bellman Equation:

$$Q^*(s, a) = \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma \max_{a'} Q^*(s', a'))$$

♦ Q-value Iteration:

$$Q_{k+1}^*([5, 1]) \leftarrow \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma \max_{a'} Q_k^*(s', a'))$$

$Q_{k+1}([5, 1])$
 $Q_{k+1}([6, 2])$
 $k \rightarrow \infty$
 $Q_k([5, 1])$ ✓ $Q_{k+1}([5, 1])$ ✗ Q

Q-Value Iteration

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

noise = 20% → 40%

↘ |state| < ∞

$R \leq C$

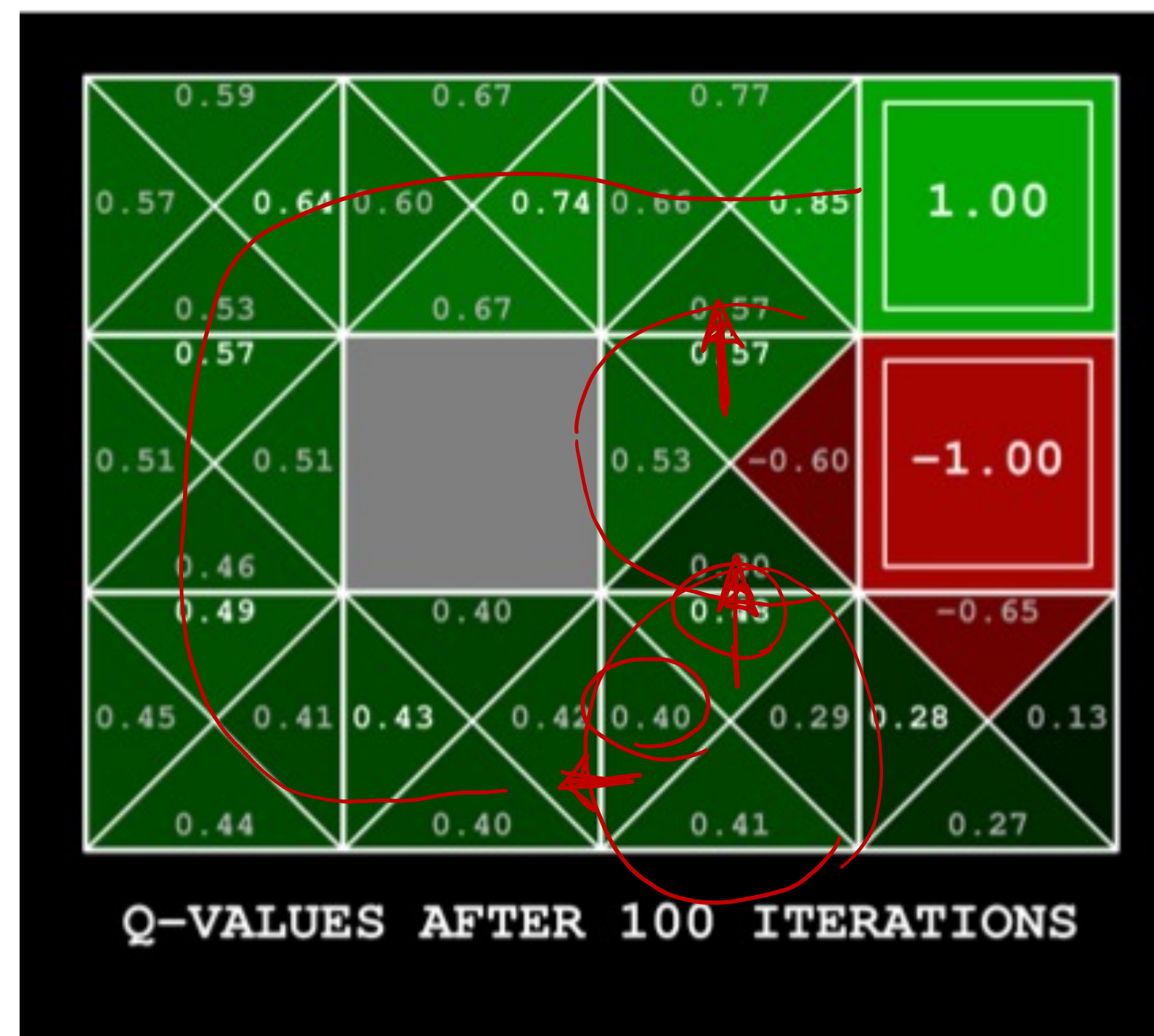
$\gamma < 1$

Alg. Converges to

V^* or Q^*

for all states

k = 100



$x = f(x)$

Noise = 0.2
Discount = 0.9

Policy Evaluation

- ♦ Recall value iteration:

$$\boxed{\checkmark} V_k^*(s) \leftarrow \max_a \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma V_{k-1}^*(s'))$$

- ♦ Policy evaluation for a given $\pi(s)$:

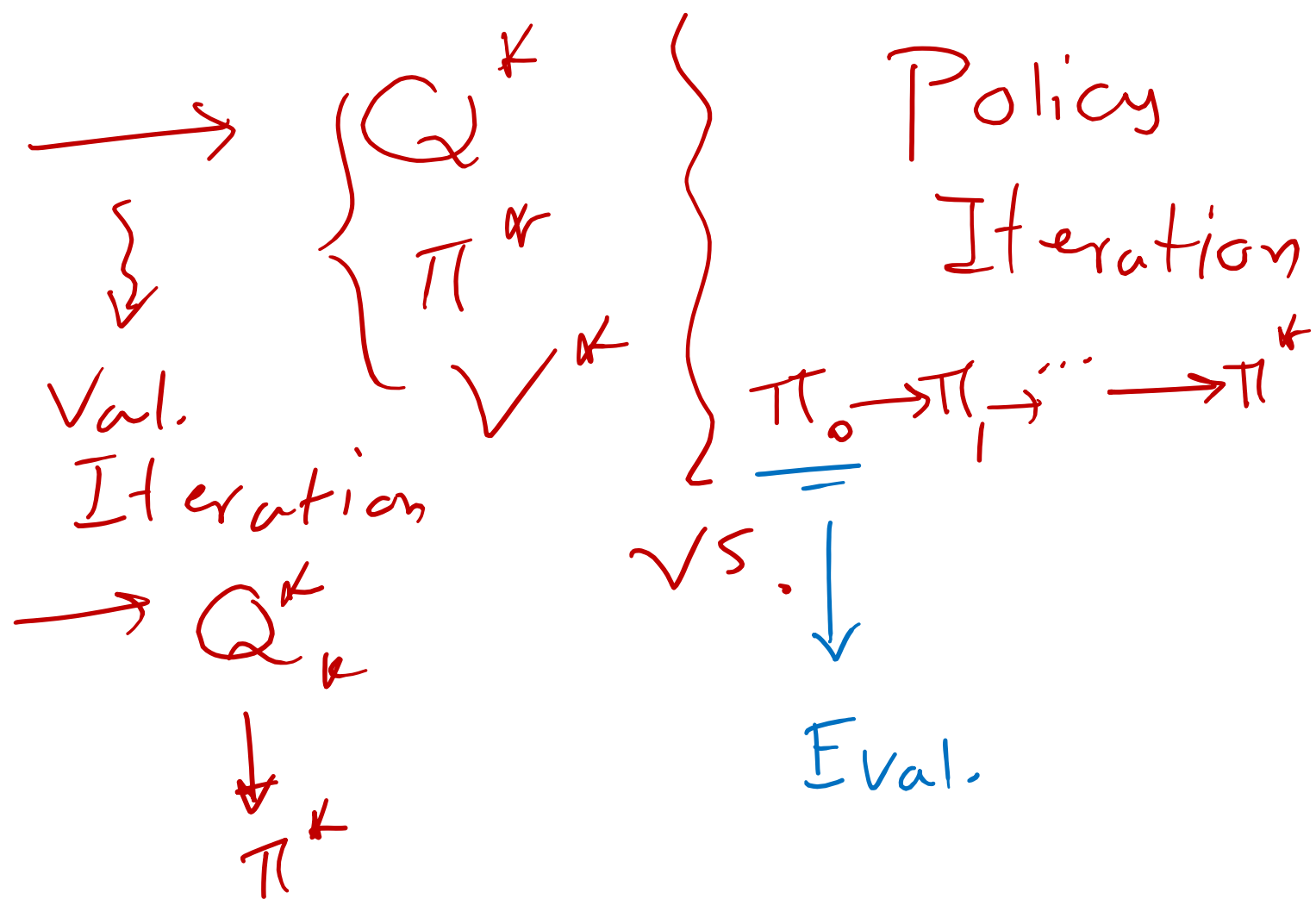
$$\boxed{\checkmark} V_k^\pi(s) \leftarrow \sum_{s'} P(s'|s, \pi(s)) (R(s, \pi(s), s') + \gamma V_{k-1}^\pi(s'))$$

- ♦ At convergence:

$$\boxed{\checkmark} \forall s V^\pi(s) \leftarrow \sum_{s'} P(s'|s, \pi(s)) (R(s, \pi(s), s') + \gamma V^\pi(s))$$



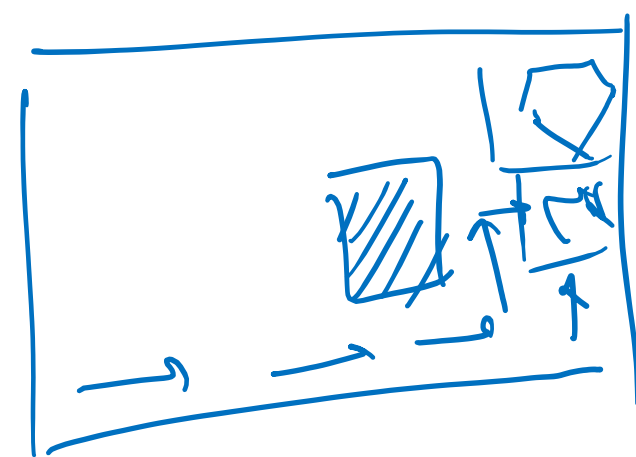
$$\left\{ \begin{array}{l} P(s_{t+1} | \underline{s_t}, a_t) \\ \underline{R(s_t, a_t, s_{t+1})} \end{array} \right.$$



What does eval. mean

for a given policy π_0 ?

$V^{\pi_0} Q^{\pi_0}(s, a)$ = expected ret. when starting @ s , take a , then follow π_0 to Choose actions!



eval.

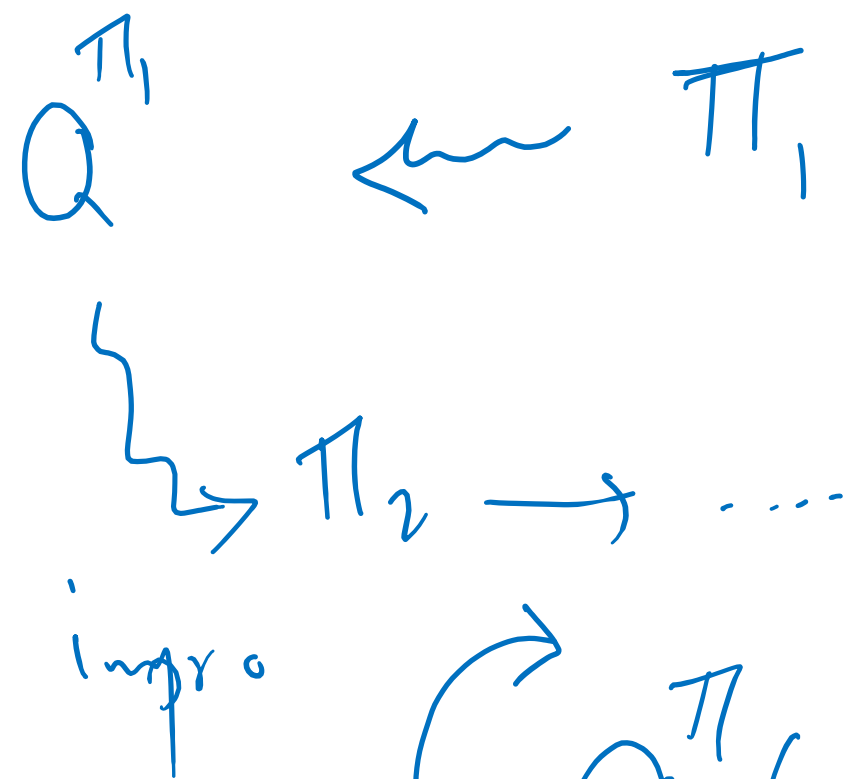
Policy

Improv. Criticism of π_0

$$\leftarrow \underline{Q^{\pi_0}(s, a)}$$

② Policy Improvement

$$\pi_1(s) = \arg \max_a \underline{Q^{\pi_0}(s, a)}$$



$$Q_{k+1}^{\pi}(s, a) \leftarrow \textcircled{\sum} P(s' | s, a) [R(s, a, s') + \gamma \overline{Q_k^{\pi}(s', \pi(s'))}]$$



$$\gamma \sum_a Q_k^{\pi}(s', a) \pi(a | s')$$

Policy Iteration

- ♦ One iteration of policy iteration

- ☑ Policy evaluation for current policy π_k . Iterate until convergence

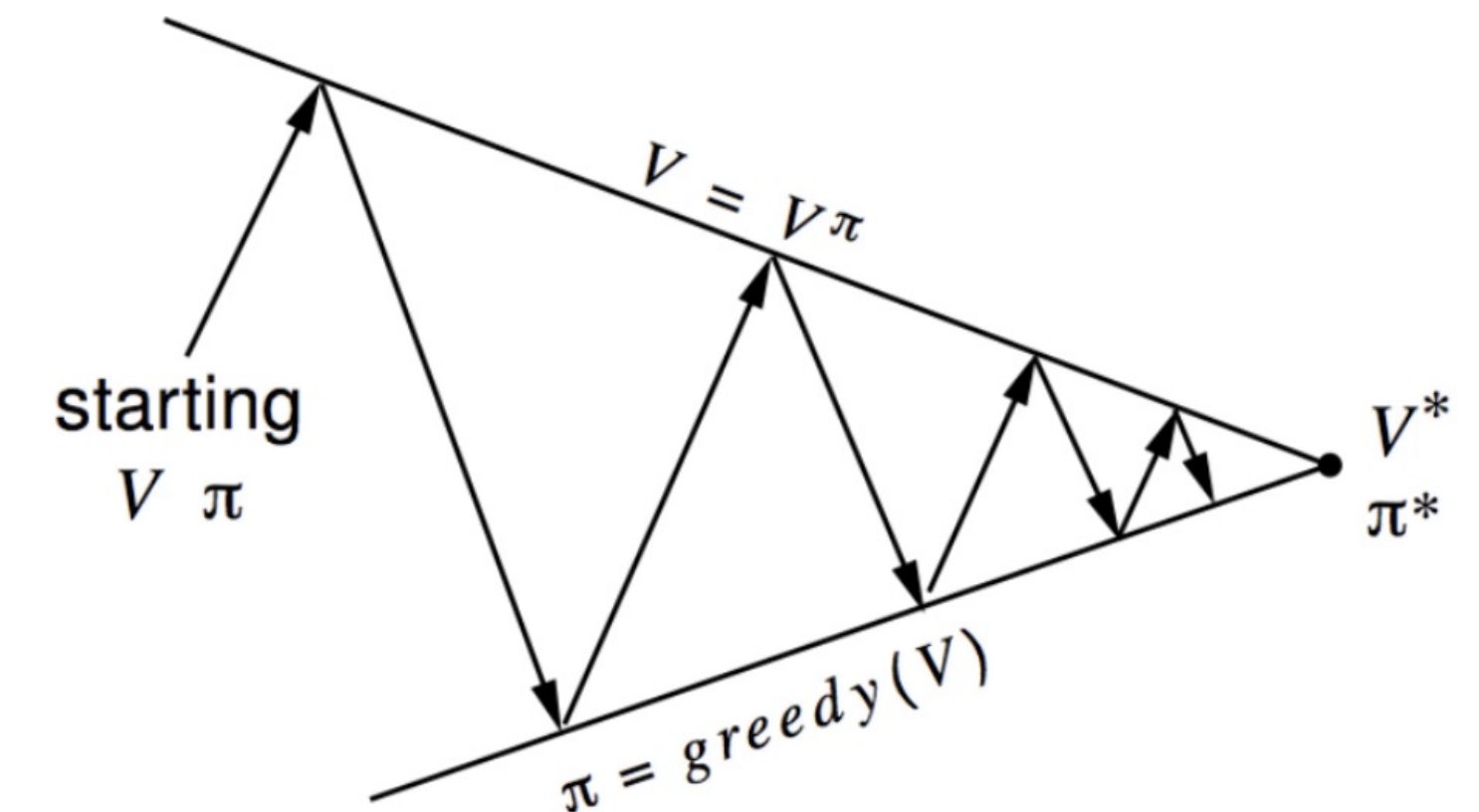
$$V_{i+1}^{\pi_k}(s) \leftarrow \sum_{s'} P(s'|s, \pi_k(s)) [R(s, \pi_k(s), s') + \gamma V_i^{\pi_k}(s')]$$

- ☑ Policy improvement: find the best action according to one-step look-ahead

$$\pi_{k+1}(s) \leftarrow \operatorname{argmax}_a \sum_{s'} P(s'|s, a) [R(s, a, s') + \gamma V^{\pi_k}(s')]$$

- ♦ Repeat until policy converges

- ☑ At convergence: optimal policy; and converges faster than value iteration under some conditions



One-step look ahead improves the policy $\rightarrow \pi_{k+1}^{(s)} = \max_a Q_k(s, a)$

$$\pi_{k+1}^{(\infty)} \geq \dots \geq \pi_{k+1}^{(2)} \geq \pi_{k+1}^{(1)} \geq \pi_k$$

- Consider an alternative policy $\pi_{(k+1)}^{(1)}(t, s)$ that takes the prescribed actions in $\pi_{k+1}(s)$ only at time $t = 0$; and stays the same as $\pi_k(s)$ in later times.
- The value function $V(s)$ for this new time-dependent policy is larger than or equal to $V(s)$ for the original policy $\pi_k(s)$ for all s . Why?
- Now let $\pi_{(k+1)}^{(2)}(t, s)$, which takes the prescribed action in $\pi_{k+1}(s)$ only at times $t = 0$ and $t = 1$, and stays the same as $\pi_k(s)$ in later times.
- Similarly, $V(s)$ gets improved for $\pi_{(k+1)}^{(2)}(t, s)$ compared to $\pi_{(k+1)}^{(1)}(t, s)$ for all s .
- Repeating this argument $\pi_{(k+1)}^{(\infty)}(t, s)$ becomes the same as $\pi_{k+1}(s)$.

An Example

noise = 0
 $\gamma = \text{disc.} = 1$

Policy Eval.

- Let this be the initial policy π_0 , show how policy improvement, makes this policy better.

π_0

+1	→	+1
-1	(↑)	-1
-1	→	-1
-1	↑	-1
-1	←	-1
-1	↑	-1
-1	→	-1

$\{ P(s_{t+1}/s_t, a_t) \}$
 $R \rightarrow \pi^*$

Policy Itr.
 Val. Itr.

Q^{π_0}

↑ ⇒ [1 | 1 | 1]
 ↑ ⇒ [-1 | 1 | -1]
 (↑) ⇒ [-1 | 1 | -1]
 Policy improvement
 [-1 | 1 | -1]

Planning vs. Learning

↪ World Model X

- Assumed to have access to the dynamics $P(s'|s, a)$.
- We **don't have access** to this in the real world.
- We need to **estimate (or learn)** the value functions.

$n \rightarrow \infty$
empirical mean

without access to this
sampling ← learning

$$Q^*(s, a) = \sum_{s'} P(s'|s, a) (R(s, a, s') + \gamma \max_{a'} Q^*(s', a'))$$

?

$\frac{1}{n} \sum_{i=1}^n X_i$

$X_1, \dots, X_n \stackrel{iid}{\sim} f_X$

$$\approx \mathbb{E}(X) = \int x \cdot f_X(x) dx$$
$$= \sum_x x \cdot P(X=x)$$

Monte-Carlo Prediction

Control

Monte Carlo Methods - Introduction

- ♦ Experience samples to learn without a model

Data

MC Pred: Estimate the val.
function based on

- ♦ MC methods require only experience— sample sequences of states, actions, and rewards from actual or simulated interaction with an environment.

the experienced
episodes.

- ♦ We can learn with samples: episodes!

World Model $\checkmark$ no access to $P(s'|s, a)$



Monte-Carlo prediction

- Suppose we wish to estimate $V_\pi(s)$, the value of a state s under policy π .
- The first-visit mc method estimates $V_\pi(s)$ as the average of the returns following first visits to s .

First-visit MC prediction, for estimating $V \approx v_\pi$

Input: a policy π to be evaluated

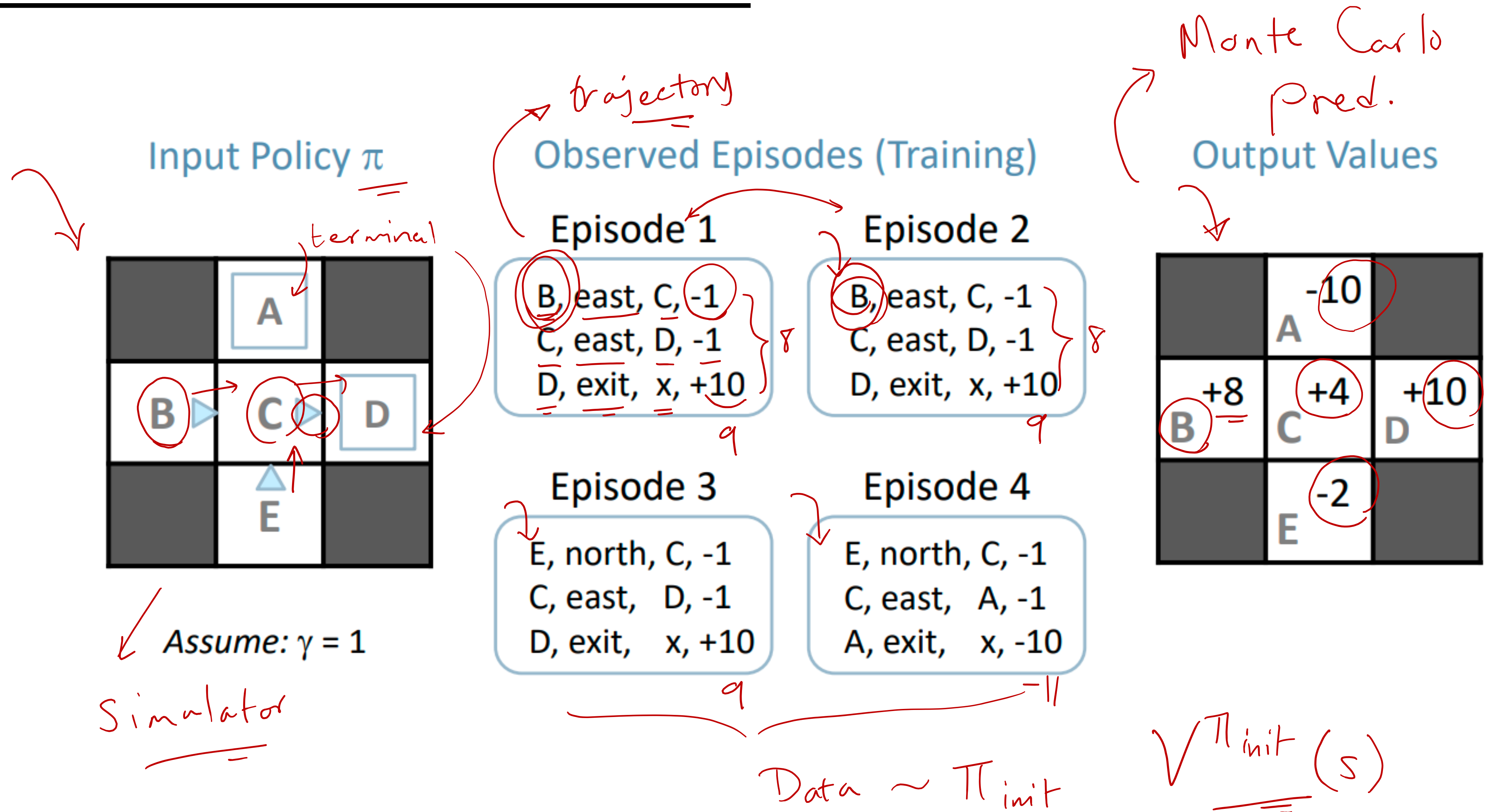
Initialize:

- $V(s) \in \mathbb{R}$, arbitrarily, for all $s \in \mathcal{S}$ $\forall s \quad V(s) \leftarrow 0$
- $Returns(s) \leftarrow$ an empty list, for all $s \in \mathcal{S}$
- Loop forever (for each episode):
 - Generate an episode following π : $S_0, A_0, R_1, S_1, A_1, R_2, \dots, S_{T-1}, A_{T-1}, R_T$
 - return $G \leftarrow 0$
 - Loop for each step of episode, $t = T-1, T-2, \dots, 0$:
 - $G \leftarrow \gamma G + R_{t+1}$
 - Unless S_t appears in $S_0, S_1, \dots, S_{t-1}$:
 - Append G to $Returns(S_t)$
 - $V(S_t) \leftarrow \text{average}(Returns(S_t))$

Handwritten notes:

- G_1, G_2, G_3, G_4 are not independent
- $S_{t-k} \neq S_t$
- $S_t = S_t$

Episodes: another example



Every Visit Monte-Carlo Policy

Every Visit Monte-Carlo Policy

Initialize $N(s) = 0, G(s) = 0 \forall s \in S$

Loop

Sample episode $i = s_{i,1}, a_{i,1}, r_{i,1}, s_{i,2}, a_{i,2}, r_{i,2}, \dots, s_{i,T_i}$

Define $G_{i,t} = r_{i,t} + \gamma r_{i,t+1} + \gamma^2 r_{i,t+2} + \dots + \gamma^{T_i-t} r_{i,T_i}$ as return from time step t onwards in i 'th episode.

For each time step t until T_i (the end of the episode i):

state s is the state visited at time step t in episodes i

Increment counter of total visits: $N(s) = N(s) + 1$

Increment total return: $G(s) = G(s) + G_{i,t}$

Update estimate $V^\pi(s) = G(s)/N(s)$

without
if

$S \dots S \dots S \dots S \dots S$

Incremental Monte-Carlo Policy

Policy Eval. $\rightarrow$ Policy Imp.

Incremental Monte-Carlo Policy

$s : G_1, G_2, \dots, G_k$

$$\frac{G_1 + \dots + G_k}{k} = V^\pi(s)$$

After each episode $i = s_{i,1}, a_{i,1}, r_{i,1}, s_{i,2}, a_{i,2}, r_{i,2}, \dots$

- Define $G_{i,t} = r_{i,t} + \gamma r_{i,t+1} + \gamma^2 r_{i,t+2} + \dots$ as return from time step t onwards in i 'th episode

- For state s visited at time step t in episode i

- Increment counter of total visits: $N(s) = N(s) + 1$

- Update estimate

$$V^\pi(s) = V^\pi(s) \frac{N(s)-1}{N(s)} + \frac{G_{i,t}}{N(s)} = V^\pi(s) + \frac{1}{N(s)} (G_{i,t} - V^\pi(s))$$

$$V^\pi(s) \leftarrow \frac{V^\pi(s)k + G_{k+1}}{k+1}$$

$$N(s) = k$$

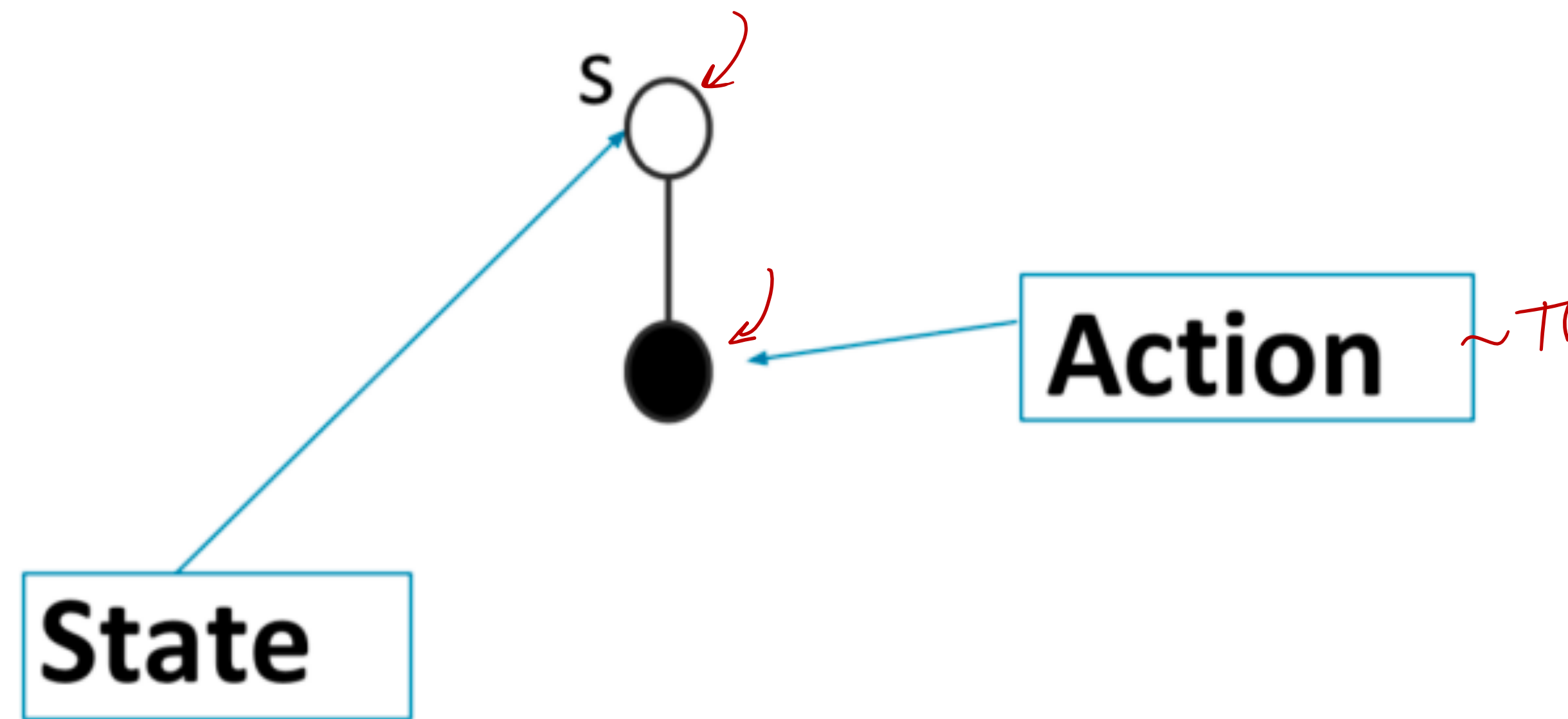
estimation error

$$\alpha = \frac{1}{n(s)}$$

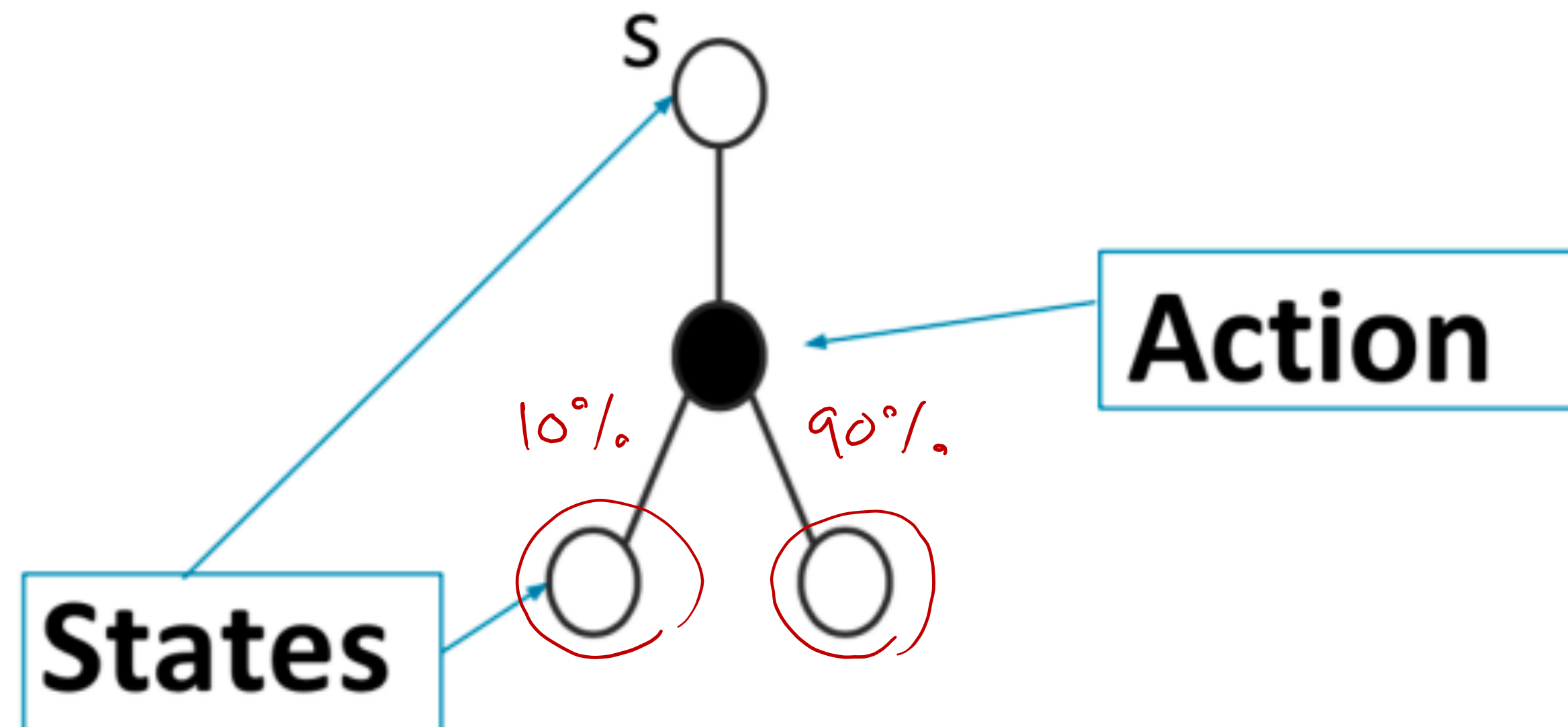
$$\alpha = \alpha_0$$

exponentially weighted average

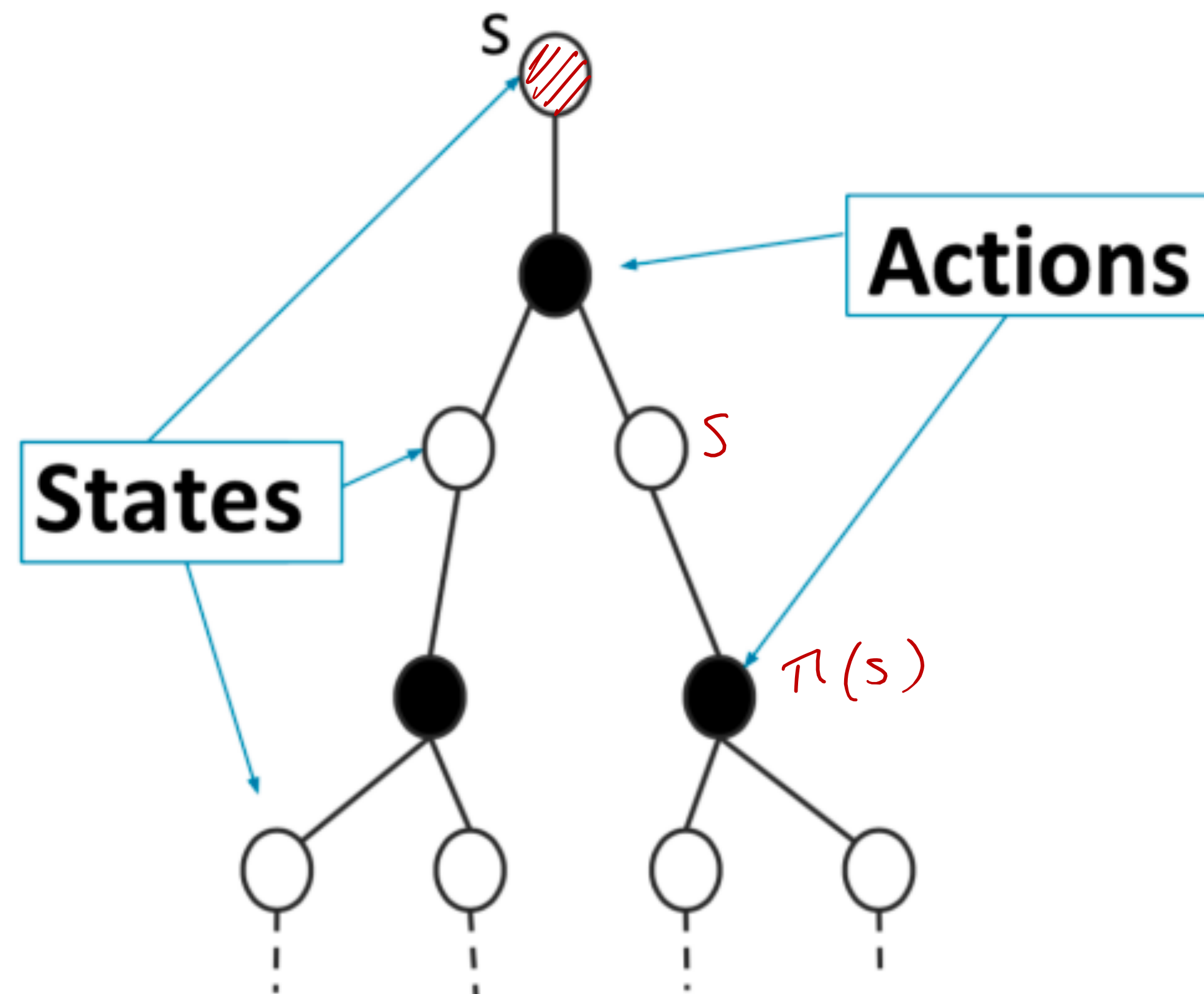
Policy Evaluation Diagram



Policy Evaluation Diagram

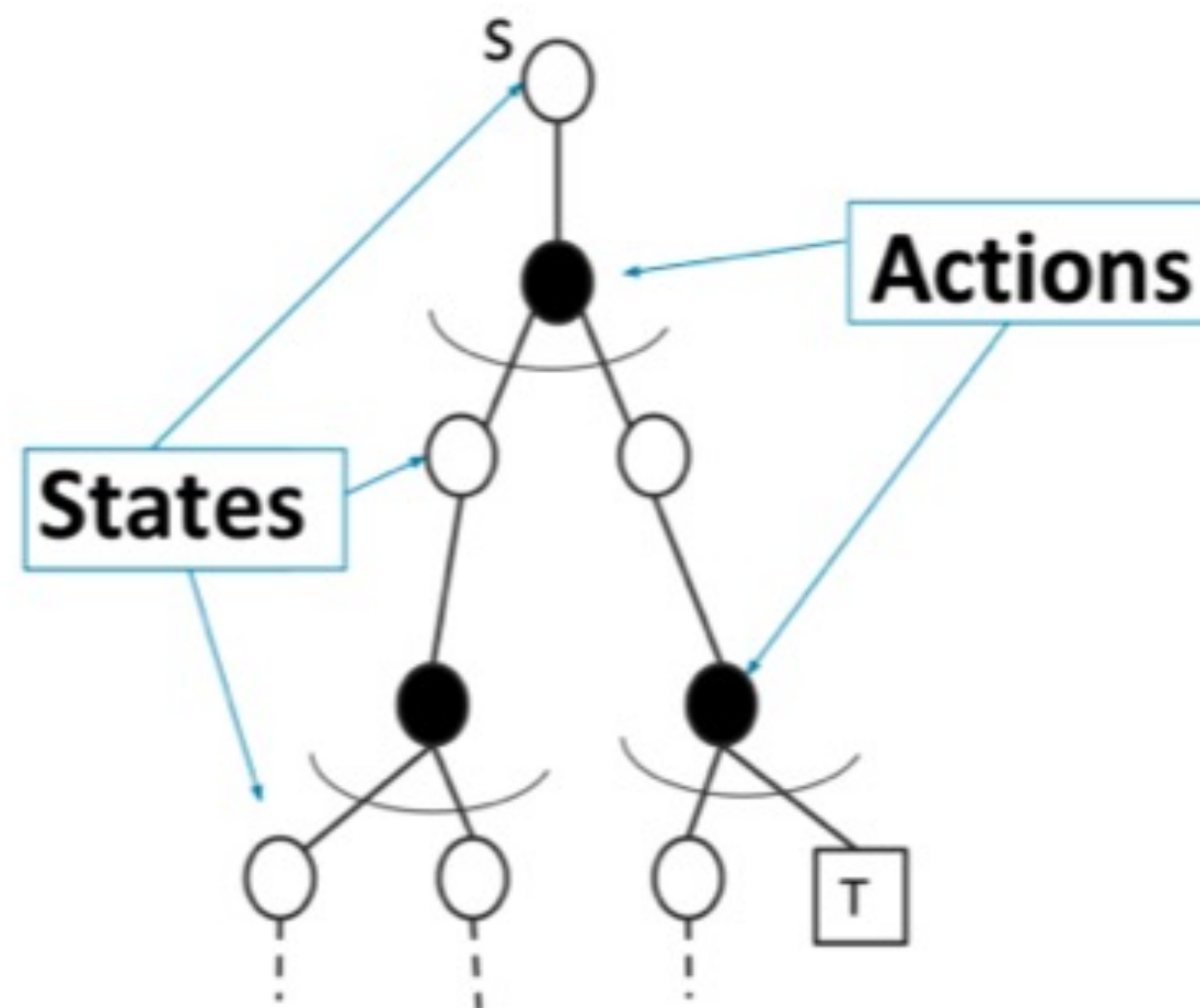


Policy Evaluation Diagram



Policy Evaluation Diagram

$$V^\pi(s) = V^\pi(s) + \underline{\alpha}(G_{i,t} - V^\pi(s))$$

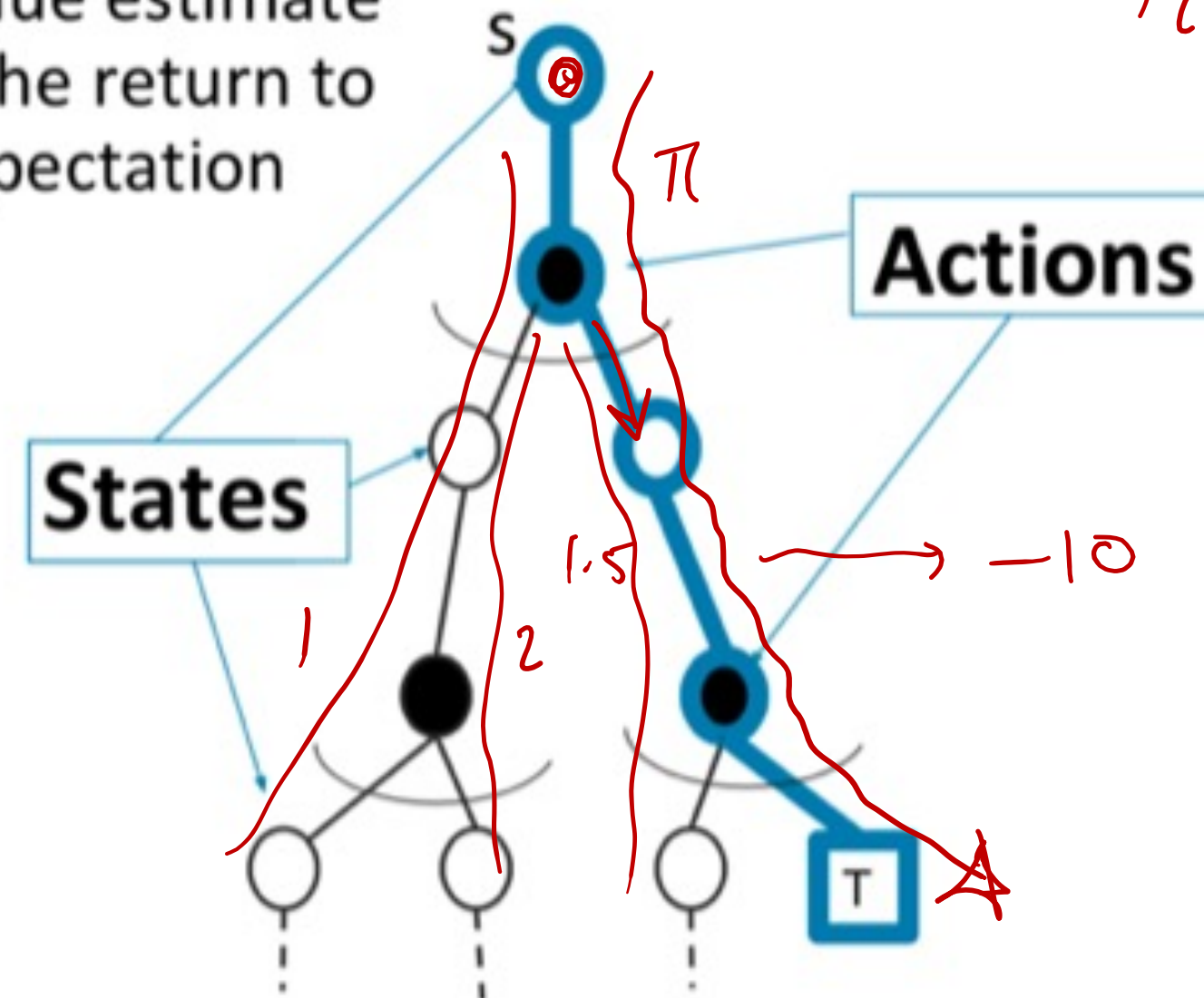


$\frown$ = Expectation
 $\boxed{T}$ = Terminal state

Policy Evaluation Diagram

$$\underline{V^\pi(s)} = V^\pi(s) + \alpha(\underline{G_{i,t}} - \underline{V^\pi(s)})$$

MC updates the value estimate using a **sample** of the return to approximate an expectation



$\underbrace{\hspace{1cm}}$ = Expectation
 $\boxed{T}$ = Terminal state

$$\left\{ \begin{array}{l} V^\pi(s) = -5 \\ V^\pi(s) = 8 \end{array} \right.$$

Val. Iter.

$$\begin{aligned} & \rightarrow V_{k+1}^\pi(s) \\ &= \sum_{s'} P(s'|s, \pi(s)) \left[R(s, \pi(s), s') + \gamma V_k^\pi(s') \right] \end{aligned}$$