

7.1 N-STEP TO PREDICTION

BASIC IDEA: TD METHODS LOOK ONE STEP AHEAD,
MC LOOKS ALL THE WAY TO THE END,
LET'S BRIDGE THE GAP.

IN TERMS OF BACKUP DIAGRAMS:

2-STEP TD



2-STEP TD



N-STEP TD



∞-STEP TD,
MONTE CARLO



IN TERMS OF THE CORRECTIONS WE USE
(CALLED TARGET)

MONTE CARLO: THE TARGET IS THE RETURN

$$G_t \triangleq R_{t+1} + \gamma R_{t+2} + \gamma^2 R_{t+3} + \dots + \gamma^{T-t-1} R_T$$

ONE STEP UPDATES:

$$G_{t:t+1} \triangleq R_{t+1} + \gamma V_t(S_{t+1})$$

NEW
NOTATION

↑

REWARD

↑

ESTIMATE OF VALUE

$V_t: S \rightarrow R$ IS THE
ESTIMATE OF V_{π} AT
TIME t

TWO-STEP UPDATES:

$$G_{t:t+2} \triangleq R_{t+1} + \gamma R_{t+2} + \gamma^2 V_{t+1}(S_{t+2})$$

n-STEP UPDATE:

$$G_{t:t+n} \triangleq \begin{cases} R_{t+1} + \gamma R_{t+2} + \dots + \gamma^{n-1} R_{t+n} + \gamma^n V_{t+n-1}(S_{t+n}), \\ 0 \leq t < T-n \end{cases}$$

$$G_t, \quad \text{IF } t+n \geq T$$

OBSERVING THAT TO UPDATE THE VALUE OF A STATE, WE NEED TO WAIT UNTIL TIME $t+n$. THEREFORE, THE UPDATE IS:

$$V_{t+n}(S_t) \triangleq V_{t+n-1}(S_t) + \alpha [G_{t:t+n} - V_{t+n-1}(S_t)],$$

ALSO, $V_{t+n}(s) = V_{t+n-1}(s) \quad \forall s \neq S_t$.

n-STEP TD FOR ESTIMATING $V \approx V_{\pi}$

INPUT: POLICY π ALGORITHM PARAMETERS: $\begin{cases} \text{STEP SIZE } \alpha \in (0, 1] \\ n \in \mathbb{N}^* \end{cases}$

INITIALIZE $V(s)$ ARBITRARILY, $\forall s \in S$

ALL STATES AND ACCESS OPERATIONS (FOR S_t, R_t) CAN TAKE INDEX $\text{mod}(nH)$

LOOP FOR EACH EPISODE:

INITIALIZE AND START $S_0 \neq \text{TERMINAL}$

$T \leftarrow \infty$

LOOP FOR $t=0, 1, \dots$

IF $t < T$, THEN:

TAKE ACTION ACCORDING TO $\pi(\cdot | S_t)$

OBSERVE AND STORE NEXT REWARD R_{t+1} AND STATE S_{t+1}

IF S_{t+1} TERMINAL, $T \leftarrow t+1$

$\tau \leftarrow t-n+1$ (WE WILL UPDATE ESTIMATE FOR STATE S_t)

IF $\tau \geq 0$:

$$G \leftarrow \sum_{i=\tau}^{t+n-1} \gamma^{i-\tau-1} R_i$$

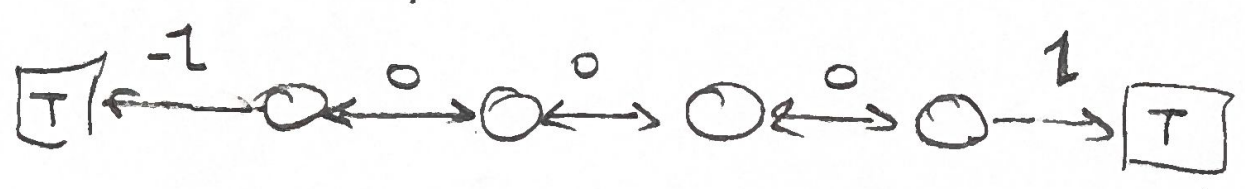
IF $\tau+n < T$, THEN $G \leftarrow G + \gamma^n V(S_{\tau+n})$

$$V(S_t) \leftarrow V(S_t) + \alpha [G - V(S_t)]$$

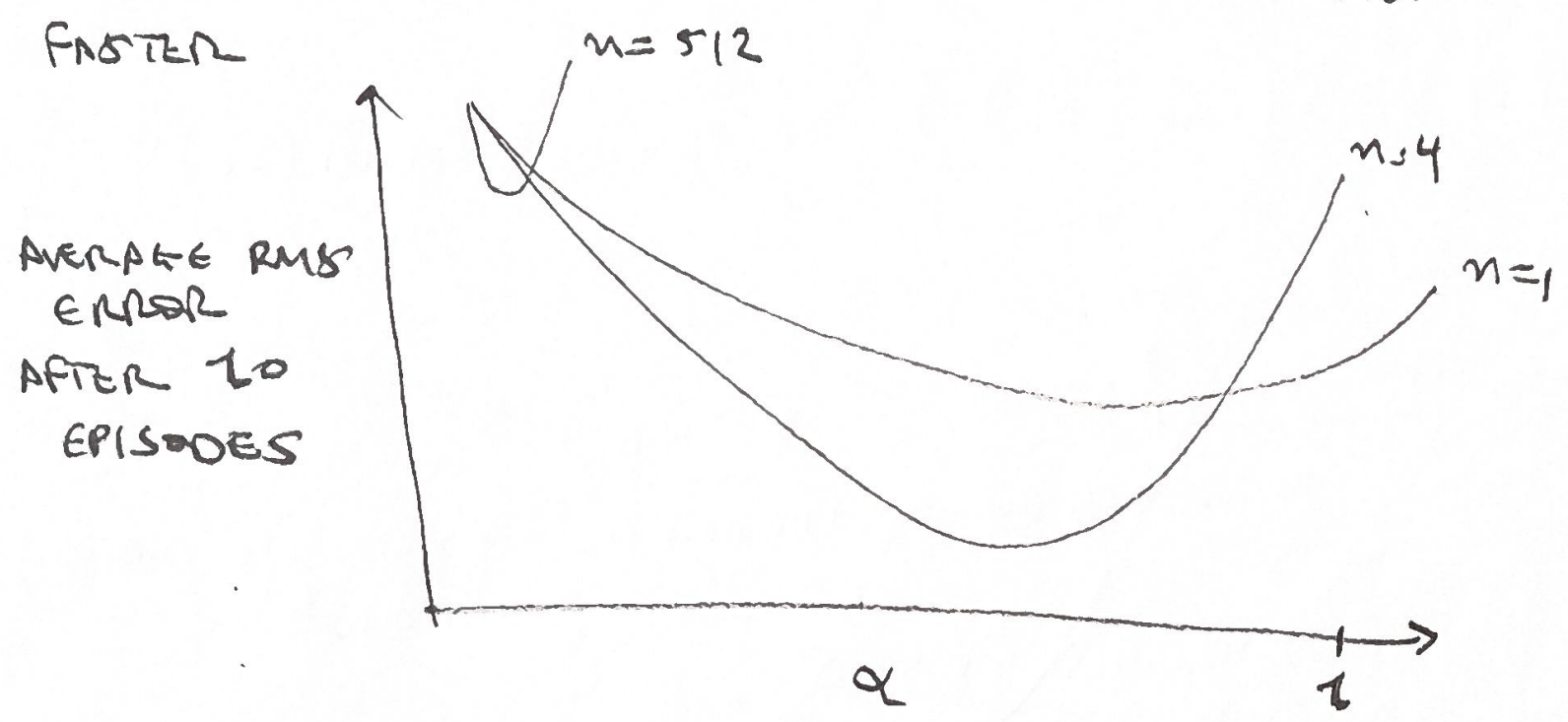
UNTIL $\tau = T-1$

THIS ALGORITHM CAN BE SHOWN TO CONVERGE.

SO, WHY IS THIS BETTER? CONSIDER EXAMPLE 6.2 (RANDOM WALK) WITH 19-STATES. THIS IS NOW EXAMPLE 7.1



WITH 1-STEP TD, AFTER THE FIRST EPISODE, ONLY ONE STATE CHANGES VALUE ESTIMATE. WITH n-STEP, n STATES WILL CHANGE THEIR VALUE ESTIMATE. SO THIS IS FASTER



7.2 n-STEP SARGA

NOW WE MOVE FROM PREDICTION TO CONTROL. WE NEED TO FOLLOW AN ACTUAL VALUES. WE DEFINE: (FOR $n \geq 1$)

$$G_{t:t+n} \triangleq \begin{cases} R_{t+1} + \gamma R_{t+2} + \dots + \gamma^{n-1} R_{t+n} + \gamma^n Q_{t+n-1}(S_{t+n}, A_{t+n}) & 0 \leq t < T-n \\ G_t = R_{t+1} + \gamma R_{t+2} + \dots + \gamma^{T-t-1} R_T, & t+n \geq T \end{cases}$$

SO THE UPDATE BECOMES

$$Q_{t+n}(S_t, A_t) \triangleq Q_{t+n-1}(S_t, A_t) + \alpha [G_{t:t+n} - Q_{t+n-1}(S_t, A_t)]$$

AND

$$Q_{t+n}(s, a) = Q_{t+n-1}(s, a) \quad \forall (s, a) \neq (S_t, A_t)$$

THIS IS n-STEP SARGA

N-STEP SARSA FOR ESTIMATING $Q \approx q^*$ (or q_π) (67)

INITIALIZE $Q(s, a)$ ARBITRARILY, FOR ALL $s \in S, a \in A$

INITIALIZE π TO BE ϵ -GREEDY WITH RESPECT TO Q , (or

ALGORITHM PARAMETERS: $\alpha \in (0, 1]$, SMALL $\epsilon > 0$, $n \in \mathbb{N}^*$ ^{FIXED POLICY π}

(ALL STATE AND ACTION OPERATIONS (FOR S_t, R_t, A_t) AND TAKE THEIR INDEX MOD $n+1$)

LOOP FOR EACH EPISODE

INITIALIZE AND STORE $S_0 \neq \text{TERMINAL}$

SELECT AND STORE AN ACTION $A_0 \sim \pi(\cdot | S_0)$

$T \leftarrow \infty$

LOOP FOR $t=0, 1, 2, \dots$

IF $t < T$, THEN:

TAKE ACTION A_t

OBSERVE AND STORE THE NEXT REWARD AS R_{t+1}
AND THE NEXT STATE AS S_{t+1}

IF S_{t+1} TERMINAL,

$T \leftarrow t+1$

ELSE:

SELECT AND STORE AN ACTION $A_{t+1} \sim \pi(\cdot | S_{t+1})$

$Z \leftarrow t - n + 1$ (# Z IS THE TIME WHERE ESTIMATE IS UPDATED)

IF $Z \geq 0$:

$G \leftarrow \sum_{i=Z+1}^{\min(Z+n, T)} \gamma^{i-Z-1} R_i$

IF $Z+n < T$, THEN $G \leftarrow G + \gamma^n Q(S_{Z+n}, A_{Z+n})$

$Q(S_t, A_t) \leftarrow Q(S_t, A_t) + \alpha [G - Q(S_t, A_t)]$

IF π IS BEING LEARNED, THEN ENSURE IT IS ϵ -GREEDY WRT. Q

UNTIL $Z = T-1$

BACKUP DIAGRAMS

1-STEP SARSA
(SARSA(0))



2-STEP SARSA



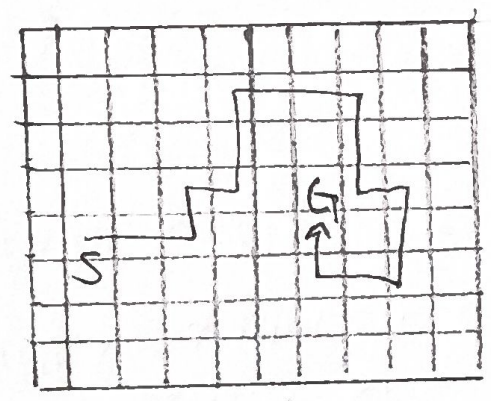
n-STEP SARSA



∞-STEP SARSA
(MONTE CARLO)



SO WHY IS IT BETTER? FIGURE 7.4



CONSIDER NEW EPISODE IN GRID WORLD. ASSUME ALL $Q=0$ INITIALLY.

- 1) 1-STEP SARSA STRENGTHENS THE LAST ACTION
- 2) 10-STEP SARSA STRENGTHENS THE LAST 10 ACTIONS

MODIFICATION FOR EXPECTED SARSA

$$G_{t:t+m} \triangleq \begin{cases} R_{t+1} + \dots + \gamma^{m-1} R_{t+m} + \gamma^m \bar{V}_{t+m-1}(S_{t+m}), & t+m < T \\ G_t, & t+m \geq T \end{cases}$$

where

$$\bar{V}_t(s) \triangleq \sum_a \pi(a|s) Q_t(s, a) \leftarrow \text{EXPECTED APPROXIMATE VALUE}$$

7.3 n-STEP OFF-POLICY LEARNING

(69)

RECALL OFF-POLICY METHODS: TWO POLICIES $\rightarrow$ TARGET π FOR WHICH WE LEARN (I.E. GREEDY WITH RESPECT TO CURRENT Q)

$\downarrow$

BEHAVIOR POLICY ϕ

TO MODIFY THE EVALUATION OF THE STATE VALUE ESTIMATE, WE USE, IN THE CASE OF n-STEP TD:

$$V_{t+n}(S_t) \triangleq V_{t+n-1}(S_t) + \alpha p_{t:t+n-1} [G_{t:t+n} - V_{t+n-1}(S_t)]$$

(THIS LEARNS π)

$$p_{t:t+n-1} \triangleq \prod_{k=t}^{t+n-1} \frac{\pi(A_k | S_k)}{\phi(A_k | S_k)}$$

(WHERE π IS THE TARGET POLICY)

FOR EXAMPLE:

- 1) IF ANY OF A_k CANNOT BE TAKEN BY π , THEN $p_{t:t+n-1} = 0$ AND WE DO NOT LEARN
- 2) IF A COMBINATION OF ACTIONS IS MUCH MORE COMMON IN π , THEN $p \gg 1$ AND WE LEARN A LOT ABOUT SOMETHING WE WOULD SEE WITH π MUCH MORE OFTEN
- 3) IF $\pi = \phi$, THEN THE COEFFICIENT $p = 1$ ALWAYS

IN THE CASE OF n-STEP SARSA, WE HAVE

$$Q_{t+n}(S_t, A_t) \triangleq Q_{t+n-1}(S_t, A_t) + \alpha p_{t+1:t+n} [G_{t:t+n} - Q_{t+n-1}(S_t, A_t)]$$

$\uparrow$
WINDOW IS NOW DIFFERENT BECAUSE WE HAVE TAKEN ACTION.

off-policy n-step SARSA FOR ESTIMATING $Q \approx q_\pi$ OR q_{π^*} (70)

INPUT: ARBITRARY BEHAVIOR POLICY b SUCH THAT $b(a|s) > 0$, $\forall a \in A, s \in S$

INITIALIZE $Q(s, a)$ ARBITRARILY $\forall s \in S, a \in A$

INITIALIZE π TO BE GREEDY WITH RESPECT TO Q , OR A FIXED GIVEN POLICY

ALGORITHM PARAMETERS: STEP SIZE $\alpha \in (0, 1]$ $n \in \mathbb{N}^*$
ALL STATE/ACTIONS OPERATIONS (FOR s_t, a_t, A_t) CAN TAKE THEIR INDEX MOD n

LOOP FOR EACH EPISODE:

INITIALIZE AND START $s_0 \neq \text{TERMINAL}$

SELECT AND START AN ACTION $a_0 \sim b(\cdot | s_0)$

$T \leftarrow \infty$

LOOP FOR $t = 0, 1, 2, \dots$

IF $t < T$ THEN:

TAKE ACTION a_t

OBSERVE AND STORE THE NEXT REWARD AS r_{t+1} AND THE NEXT STATE AS s_{t+1}

IF s_{t+1} IS TERMINAL, THEN

$T \leftarrow t + 1$

ELSE

SELECT AND START AN ACTION $a_{t+1} \sim b(\cdot | s_{t+1})$

$\tau \leftarrow t - n + 1$ (τ IS THE TIME WHOSE ESTIMATE IS BEING UPDATED)

IF $\tau \geq 0$:

$$p \leftarrow \frac{\min(n, \tau + n, T - 1)}{\prod_{i=\tau+1}^{\tau+n} \pi(a_i | s_i)} \frac{\pi(a_\tau | s_\tau)}{b(a_\tau | s_\tau)}$$

$$G \leftarrow \sum_{i=\tau+1}^{\min(\tau+n, T)} \gamma^{i-\tau-1} R_i$$

IF $\tau + n < T$ THEN: $G \leftarrow G + \gamma^n Q(s_{\tau+n}, a_{\tau+n})$

$$Q(s_t, a_t) \leftarrow Q(s_t, a_t) + \alpha p [G - Q(s_t, a_t)]$$

IF π IS BEING LEARNED, THEN ENSURE THAT $\pi(\cdot | s_t)$ IS GREEDY WRT. Q

UNTIL $\tau = T - 1$