

Chapter 1: Markov processes (MP)

• State: $S_t \in \mathcal{S}$ $t = 0, 1, 2, \dots$
state space

• Markov evolution:

$$P(\underset{\text{future}}{S_{t+1} = s'} \mid \underset{\text{present}}{S_t = s}, \underset{\text{past}}{S_{t-1}, S_{t-2}, \dots, S_0})$$

$$= P(\underset{\text{future}}{S_{t+1} = s'} \mid \underset{\text{present}}{S_t = s})$$

- Future independent of past given present
- Markov chain: $S_0 \rightarrow S_1 \rightarrow S_2 \rightarrow \dots$

• Transition probability:

$$p(s' | s) = P(S_{t+1} = s' | S_t = s)$$

• Assumed homogeneous (stationary, time-independent)

• Normalization: $\sum_{s'} p(s' | s) = 1 \quad \forall s$

• Transition matrix: $P_{ij} = \underset{\text{column}}{p(j|i)} = \underset{\text{row}}{P(i \rightarrow j)}$

• $|\mathcal{S}| \times |\mathcal{S}|$ matrix

• $\sum_j P_{ij} = 1 \quad \forall i$ (row sums = 1)

• $0 \leq P_{ij} \leq 1$ (stochastic matrix)

• Probability distribution: $P_t(s) = P(S_t = s)$

• Probability vector: P_t , $P_{t,i} = P_t(i)$

• $|\mathcal{S}|$ -dim vector

• Row vector

· Evolution equation :

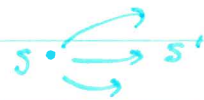
$$p_{t+1}(s') = \sum_s p(s'|s) p_t(s)$$

· Matrix notation: $p_{t+1} = p_t P$ *right product*

$$(- p_{t+1} -) = (- p_t -) \begin{pmatrix} P \end{pmatrix}$$

$$(p_{t+1})_j = \sum_i (p_t)_i P_{ij}$$

· Expectations :



$$E[f(S_{t+1}) | S_t = s] = \sum_{s'} f(s') p(s'|s)$$

· Matrix notation :

$$F_i = E[f(S_{t+1}) | S_t = i] = \sum_j f(j) p(j|i)$$

$$\begin{pmatrix} F \\ 1 \end{pmatrix} = \begin{pmatrix} P \end{pmatrix} \begin{pmatrix} f \\ 1 \end{pmatrix} = \sum_j P_{ij} f(j)$$

$$\Rightarrow F = P f \quad \text{i left product}$$

· Marginalization: $p(x) = \sum_y p(x,y)$

$$p(x|z) = \sum_y p(x,y|z)$$

· Chain rule: $p(x,y) = p(x|y) p(y) = p(y|x) p(x)$

$$p(x,y|z) = p(x|y,z) p(y|z)$$

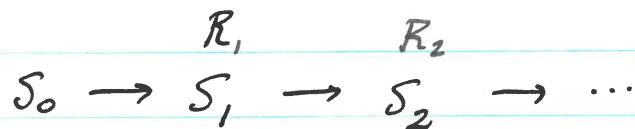
· Deterministic transition: $p(s'|s) = \begin{cases} 1 & \text{for same } s^* \\ 0 & \text{otherwise} \end{cases} = \delta_{s', s^*(s)}$
 $s \rightarrow s'$

Chapter 2: Markov reward processes (MRP)

- State: S_t $t = 0, 1, \dots$
- Reward: $R_t \in \mathbb{R}$ $t = 0, 1, \dots$
 - Cost / utility at time t
 - RV correlated with transitions $S_t \rightarrow S_{t+1}$
 - Sometimes: fct of s, s' $r(s, s')$

Transition probability (dynamics):

$$p(s', r | s) = P(S_{t+1} = s', R_{t+1} = r | S_t = s)$$

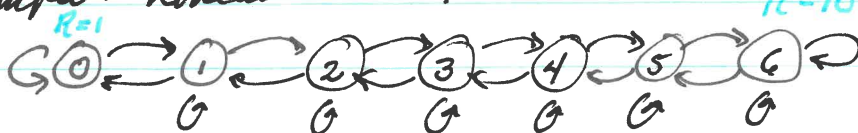


- Reward depends on s, s'
- Assumed homogeneous (time-independent, stationary)
- No actions yet: dynamics + rewards
- Reward probability: $p(r | s) = \sum_{s'} p(s', r | s)$
- Transition probability: $p(s' | s) = \sum_r p(s', r | s)$ pure dynamics
- Expected state reward:

$$\rho(s) = E[R_{t+1} | S_t = s] = \sum_r r p(r | s)$$

- One-time reward from state s

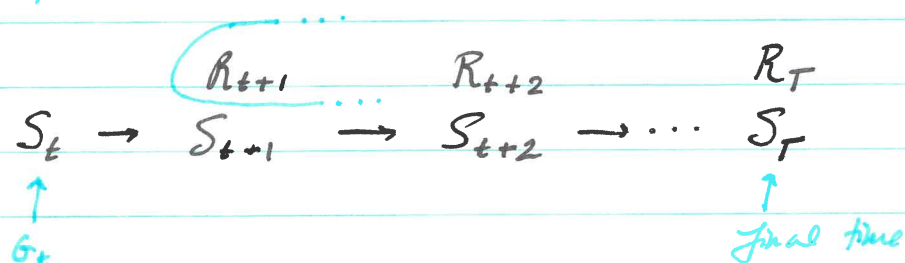
Example: linear model / Markov reward process



R deterministic
w/r S_t
 $\rho(0)=1 \quad \rho(6)=10$

Return: $G_t = R_{t+1} + R_{t+2} + \dots + R_T$

$\uparrow$ present future



• Total reward / cost to go from time t

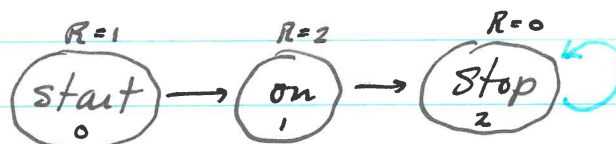
• Episodic process: T finite
Process stops / terminates

Calculate backward

$$\Rightarrow G_T = 0 \quad G_{T-1} = R_T \quad G_{T-2} = R_{T-1} + R_T \quad \dots$$

• Continuing process: $T = \infty$
Process goes on and on

• Example:



$0 \rightarrow 1 \rightarrow 2$ $1 \rightarrow 2$ 2 Episodic

$0 \rightarrow 1 \rightarrow 2 \rightarrow 2 \rightarrow 2 \rightarrow \dots$

Continuing

• Discounted return: $G_t = R_{t+1} + \gamma R_{t+2} + \gamma^2 R_{t+3} + \dots$

$$= \sum_{k=0}^{\infty} \gamma^k R_{t+1+k}$$

• Discount rate $\gamma \in [0, 1]$ Balances present / future reward

• $G_t < \infty$ for $\gamma \in [0, 1)$

• Myopic: $\gamma \approx 0$

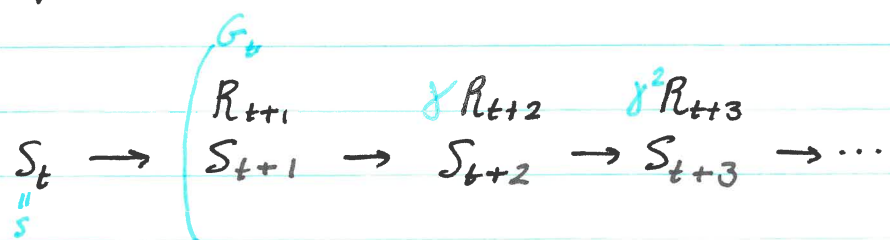
$G_t = R_{t+1}$ for $\gamma = 0$ immediate reward

• Far-sight: $\gamma \approx 1$

• Iteration:

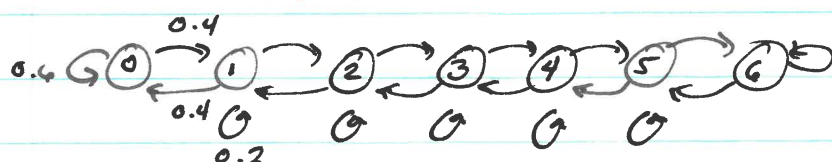
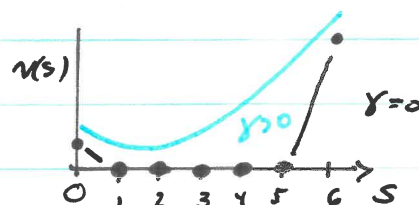
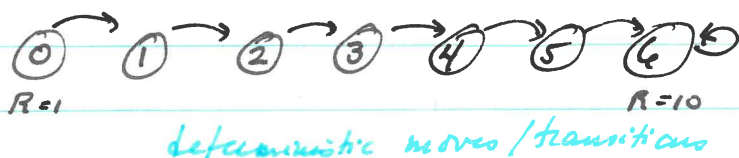
$$G_t = R_{t+1} + \gamma G_{t+1}$$

Value function: $v(s) = E[G_t | S_t = s]$



- State value fct: Expected return / cost to go from state s
- Doesn't depend on time t
 - Stationary MP
 - Continuing process / infinite return
- $\gamma = 0$: $v(s) = p(s)$ cost when leaving s

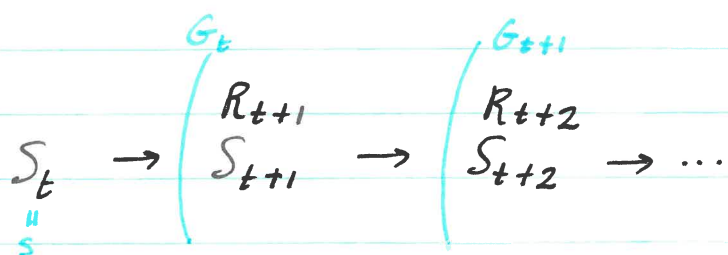
Example: linear model



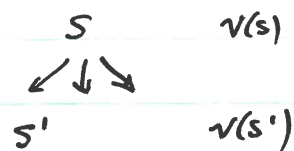
$v(s) ?$ See CW

Bellman equation:

$$\begin{aligned}
 v(s) &= E[R_{t+1} + \gamma G_{t+1} | S_t = s] \\
 &= E[R_{t+1} + \gamma v(S_{t+1}) | S_t = s]
 \end{aligned}$$



$v(s)$ $v(S_{t+1})$



· Explicit expectation:

$$\begin{aligned}
 v(s) &= \sum_r r p(r|s) + \gamma \sum_{s'} v(s') p(s'|s) \\
 &= \sum_{r, s'} r p(s', r|s) + \gamma \sum_{s', r} v(s') p(s', r|s) \\
 &= \rho(s) + \gamma (Pv)(s) \quad P_{ij} = p(j|i)
 \end{aligned}$$

· Matrix notation: $v = \rho + \gamma Pv$ $v_i = v(i)$

- $v_i = v(i)$ column vector $|S'| \times 1$
- $\rho_i = \rho(i)$ " "
- γ scalar
- $P_{ij} = p(j|i)$ $|S'| \times |S'|$ matrix

· Solution: $v = (I - \gamma P)^{-1} \rho$

- Always exists for $\gamma \in [0, 1)$
- Unique solution

· Bellman operator: $T(v) = \rho + \gamma Pv$

$$v = T(v) \quad \text{Value function} = \text{fixed pt of } T$$

· Example: Linear model.

· Write ρ , P
 7×1 7×7

· Solve numerically.

Two ways to define MDPs:

①

$$p(s', r | s)$$

$$\hookrightarrow p(r | s)$$

$$\hookrightarrow p(s) = E[R_{t+1} | S_t = s]$$

$$\hookrightarrow p(s' | s)$$

$$\hookrightarrow P$$

②

$$p(s)$$

$$P$$

$$v(s) = E[R_{t+1} + \gamma v(S_{t+1}) | S_t = s]$$

$$v = p + \gamma P v$$

Chapter 3: Markov decision processes (MDP)

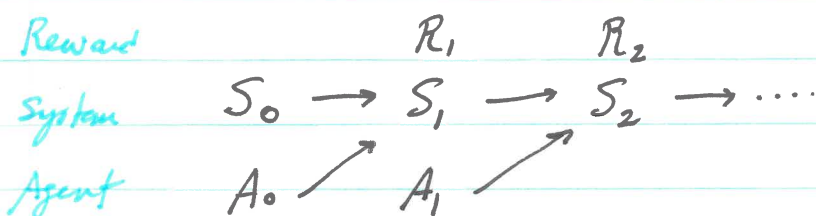
SB: Chap 3

3.1. Model

- Environment state: $S_t \in \mathcal{S}$ $t = 0, 1, \dots$
 - System to control
 - Info for making decision
- Agent state: $A_t \in \mathcal{A}(s)$ $t = 0, 1, \dots$
 - Controller
 - Action state / decision *available actions from given state*
 - State space can depend on current state
 - Simplification: $\mathcal{A}(s) = \mathcal{A} \quad \forall s$
- Reward: $R_t \in \mathbb{R}$
 - Signal from system/environment
 - Defines good/bad actions

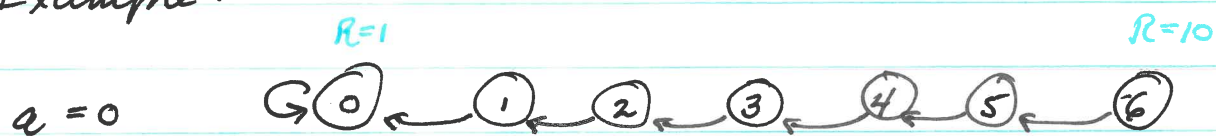
+
-
 - Zero/baseline not important
 - RV correlated with states/actions
 - Sometimes expressed as $r(s, a, s')$
- Transition probability (dynamics):

$$p(s', r | s, a) = P(S_{t+1} = s', R_{t+1} = r | S_t = s, A_t = a)$$



- Assumed homogeneous (time-independent, stationary)

Example:



- 2 deterministic actions: left, right
- Deterministic reward given state and doesn't depend on action (same for $a=0,1$)

State transition probability:

$$p(s'|s,a) = \sum_r p(s',r|s,a)$$

Transition matrix
for each a
 $p_a(j|i)$ P_a

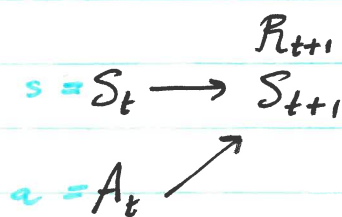
Reward probability:

$$p(r|s,a) = \sum_{s'} p(s',r|s,a)$$

↑
action

compare with
MRP

Expected reward: $\rho(s,a) = E[R_{t+1} | S_t = s, A_t = a]$



$$\rho(s,a) = \sum_r r p(r|s,a)$$

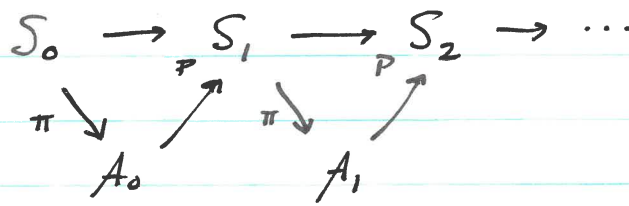
$$= \sum_{r,s'} r p(s',r|s,a)$$

- One-step reward from state s and action a
- Doesn't depend on time (stationary)

3.2. Policy

- How to choose action/control/decision at given state?
- Policy transition probability:

$$\pi(a|s) = P(\underset{\substack{\uparrow \\ \text{action}}}{A_t = a} | \underset{\substack{\uparrow \\ \text{state}}}{S_t = s})$$



} noise in environment
} explanation

- Mapping state $\rightarrow$ action
- Probabilistic policy: actions ^{can be} random
- Assumed homogeneous (time-independent, stationary)
- Deterministic policy: $\pi(s) = a$
one state one action
- Joint probability: $p(s', r, a | s) = p(s', r | s, a) \pi(a | s)$
- System transition probability:

$$\begin{aligned} P_{\pi}(s' | s) &= \sum_a p(s' | s, a) \pi(a | s) && \text{Average over actions} \\ &= \sum_{a, r} p(s', r | s, a) \pi(a | s) \end{aligned}$$

- Effective dynamics under policy π
- Transition matrix: P_{π} $(P_{\pi})_{ij} = P_{\pi}(j | i)$

$$S_0 \xrightarrow{P_{\pi}} S_1 \xrightarrow{P_{\pi}} S_2 \xrightarrow{P_{\pi}} \dots$$

Comparison:

· No control: $S_0 \xrightarrow{P} S_1 \xrightarrow{P} S_2 \rightarrow \dots$ $p(s'|s)$

· Open-loop control: $S_0 \xrightarrow{P} S_1 \xrightarrow{P} S_2 \rightarrow \dots$ $p(s'|s, a)$
 $A_0 \nearrow \quad A_1 \nearrow$

· Closed-loop control:
 Feedback $S_0 \xrightarrow{P} S_1 \xrightarrow{P} S_2 \rightarrow \dots$ $p(s'|s, a)$
 $\pi \downarrow \quad \pi \downarrow \quad \pi \downarrow$ $\pi(a|s)$
 $A_0 \nearrow \quad A_1 \nearrow$

· Reduced dynamics: $S_0 \xrightarrow{P_\pi} S_1 \xrightarrow{P_\pi} S_2 \xrightarrow{P_\pi} \dots$ $P_\pi(s'|s)$

· Expected reward under policy π :

$$\begin{aligned} \rho_\pi(s) &= E_\pi[R_{t+1} | S_t = s] \\ &= E_\pi[E[R_{t+1} | S_t = s, A_t]] \\ &= E_\pi[\rho(s, A_t)] \end{aligned}$$

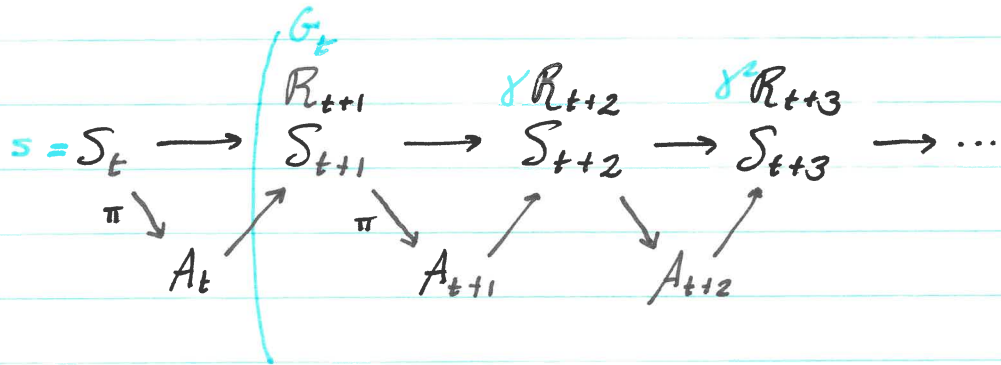
$$\begin{aligned} \rho_\pi(s) &= \sum_a \rho(s, a) \pi(a|s) \quad \text{average over actions} \\ &= \sum_{a, r} r \, p(r|s, a) \pi(a|s) \\ &= \sum_{a, r, s'} r \, p(s', r|s, a) \pi(a|s) \end{aligned}$$

- One-time / step reward from state s under policy π
- Doesn't depend on time (stationary)
- $\rho(s, a)$ doesn't depend on π
- Deterministic: $\rho_\pi(s) = \rho(s, \pi(s))$

3.3. Value Functions

- State value function:

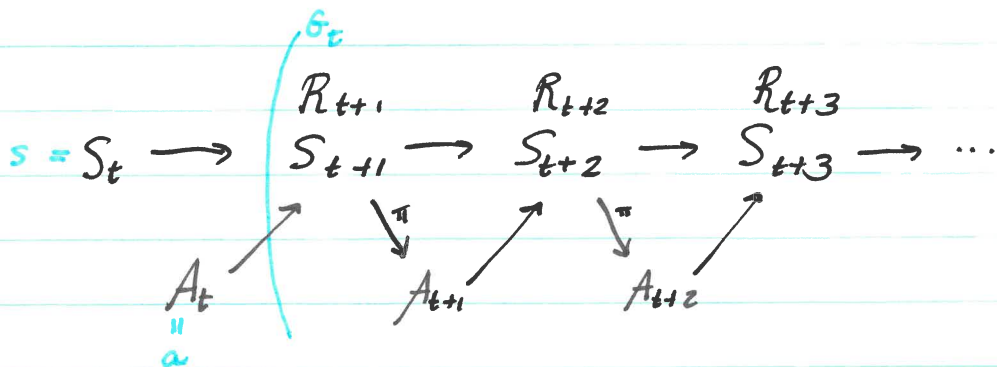
$$v_{\pi}(s) = E_{\pi}[G_t | S_t = s]$$



- Future return / cost to go from state s + policy π
- Long term value of s under π
- Function of $\pi(\cdot | s)$
- Doesn't depend on time t (stationary dynamics)
infinite return

- State-action value function:

$$q_{\pi}(s, a) = E_{\pi}[G_t | S_t = s, A_t = a]$$



- Return / cost to go if action a taken from state s
- Value of action a when in s if policy π followed
after
- Doesn't depend on time t

· Connection :

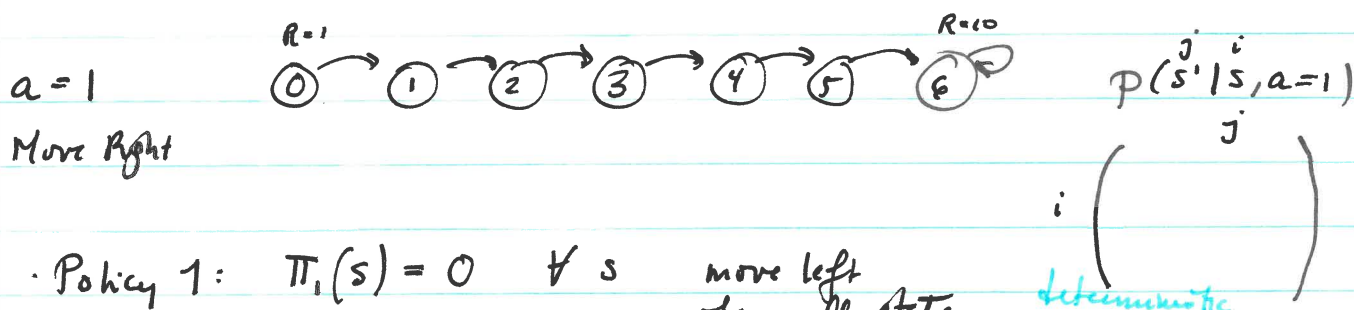
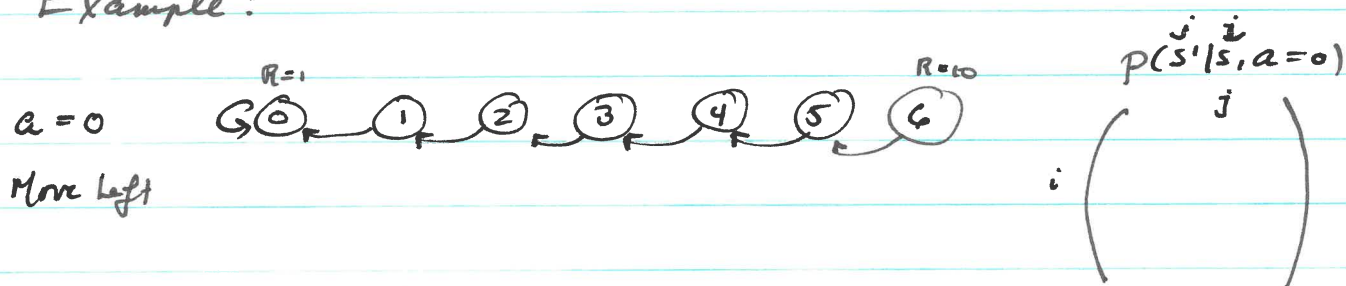
$$V_{\pi}(s) = E_{\pi} [q_{\pi}(s, A_t) | S_t = s] \quad \text{average over actions}$$

$$= \sum_a q_{\pi}(s, a) \pi(a|s)$$

$$V_{\pi}(s) = q_{\pi}(s, \pi(s)) \quad \text{for deterministic policy}$$

$a = \pi(s)$

· Example :



· Policy 1: $\pi_1(s) = 0 \quad \forall s$ move left from all state

· Policy 2: $\pi_2(s) = 1 \quad \forall s$ move right

· Policy 3: $\pi_3 = \frac{1}{2} \pi_1 + \frac{1}{2} \pi_2$ random policy/control

CW: Calculate $V_{\pi_1}, V_{\pi_2}, V_{\pi_3}$

· Note : $p_{\pi}(s) = 1$ or 10 $s=0, s=6 \quad \forall \pi$ reward only depends on s

$p(s, a) = 1$ or 10 $s=0, s=6 \quad \forall a$

· Note : 2 actions per state $\Rightarrow 2^7$ possible policies

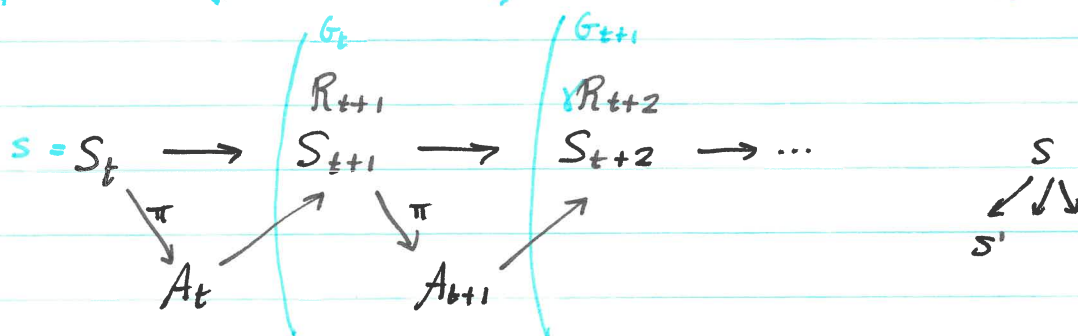
3.4. Bellman equations (BE)

• Return : $G_t = R_{t+1} + \gamma R_{t+2} + \gamma^2 R_{t+3} + \dots = R_{t+1} + \gamma G_{t+1}$

• Bellman equation for value function :

$$v_{\pi}(s) = E_{\pi} [R_{t+1} + \gamma v_{\pi}(S_{t+1}) \mid S_t = s]$$

Compare with
MRP



$$v_{\pi}(s) = \sum_{s', a} p(s', a) \pi(a|s) + \gamma \sum_{s', a} v_{\pi}(s') p(s'|s, a) \pi(a|s)$$

$$= \sum_{s', r, a} r p(s', r|s, a) \pi(a|s)$$

$$+ \gamma \sum_{s', r, a} v_{\pi}(s') p(s', r|s, a) \pi(a|s)$$

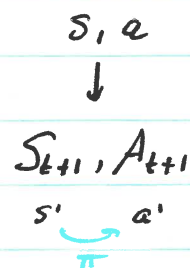
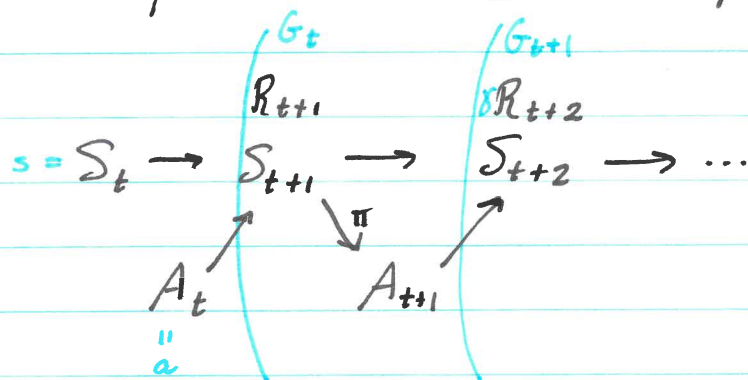
• Value = immediate reward + discounted value of
 s s s'
 success state

• Linear equation : $v_{\pi} = p_{\pi} + \gamma P_{\pi} v_{\pi}$

• Solution : $v_{\pi} = (I - \gamma P_{\pi})^{-1} p_{\pi}$

• BE for action-value function:

$$q_{\pi}(s, a) = E_{\pi}[R_{t+1} + \gamma q_{\pi}(S_{t+1}, A_{t+1}) \mid S_t = s, A_t = a]$$



$$q_{\pi}(s, a)$$

$$q_{\pi}(S_{t+1}, A_{t+1})$$

$$\begin{aligned} q_{\pi}(s, a) &= p(s, a) + \gamma \sum_{a', s'} q_{\pi}(s', a') p(s' | s, a) \pi(a' | s') \\ &= \sum_{s', r} r p(s', r | s, a) \\ &\quad + \gamma \sum_{a', s', r} q_{\pi}(s', a') p(s', r | s, a) \pi(a' | s') \end{aligned}$$

• Value = immediate reward + discounted value of
 (s, a) (s, a) success state, action
 (s', a')

• Linear equation for "matrix" q_{π}

Two ways of defining MDPs

①

②

$$p(s', r | s, a)$$

$$\hookrightarrow p(r | s, a)$$

$$\hookrightarrow p(s, a) = E[R_{t+1} | S_t = s, A_t = a]$$

 ρ

$$\hookrightarrow p(s' | s, a)$$

$$\hookrightarrow P_a$$

 P_a

$$\pi(a | s)$$

$$\hookrightarrow \rho_\pi(s) = E_\pi[R_{t+1} | S_t = s]$$

 ρ_π

$$\hookrightarrow p_\pi(s' | s)$$

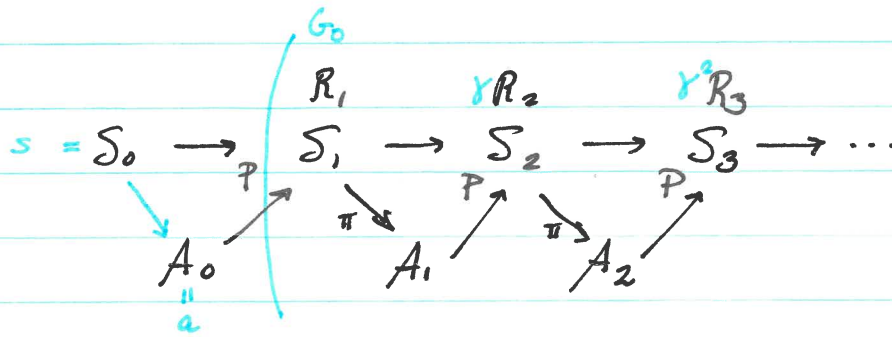
$$\hookrightarrow P_\pi$$

 P_π

$$v_\pi = \rho_\pi + \gamma P_\pi v_\pi$$

$$q_\pi = \rho + \gamma (P_a \pi q_\pi)$$

3.5 Optimal policies



Dynamics
 $p(s' | s, a)$

Policy / control
 $\pi(a | s)$

Reward
 $p(s', r | s, a)$

$$V_{\pi}(s) = E_{\pi}[G_t | S_t = s] = E_{\pi}[G_0 | S_0 = s]$$

$$q_{\pi}(s, a) = E_{\pi}[G_t | S_t = s, A_t = a] = E_{\pi}[G_0 | S_0 = s, A_0 = a]$$

infinite horizon return

- Goal: Find π with max expected return
- Optimal value function:

$$V_*(s) = \max_{\pi} V_{\pi}(s)$$

- Best actions / policies from s
- At least one solution

- Optimal action-value function:

$$q_*(s, a) = \max_{\pi} q_{\pi}(s, a)$$

- Optimal policy for success states from s, a
- Same solution π_* as V_*

- Optimal policy: $\pi_* = \arg \max_{\pi} V_{\pi}(s)$

- General results (Silver: L2): For any finite MDPs
 - There exists at least one optimal policy π_*
 - All optimal π_* achieve V_*
 - " " " " q_*
 - Optimal policy is deterministic (one action per state)
 - " " " stationary

$$\Rightarrow \pi_*(s) = a$$

- Optimal policy:

$$\pi_*(s) = \arg \max_a q_*(s, a) = \text{best action from } s$$

- Optimal value function:

$$V_*(s) = V_{\pi_*}(s) = q_*(s, \pi_*(s)) \quad \downarrow \text{best action}$$

or

$$V_*(s) = \max_a q_*(s, a) \quad \text{best action}$$

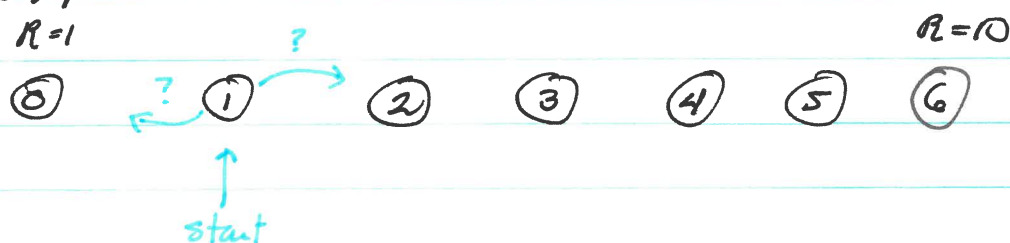
$$\pi_* \Leftrightarrow q_* \Leftrightarrow V_*$$

$$\text{Note: } V_{\pi}(s) = \sum_a q_{\pi}(s, a) \pi(a|s)$$

See p. 3-6

$$= q_{\pi}(s, \pi(s)) \quad \text{deterministic action}$$

Example:



$\gamma \approx 0$: Better to move left to get $R=1$ reward

$\gamma \approx 1$: " " " right " " $R=10$ "

Optimal policy: for each s , either left or right
 $p(s'|s, a=0)$ $p(s'|s, a=1)$

Π_* among 2^7 possible policies

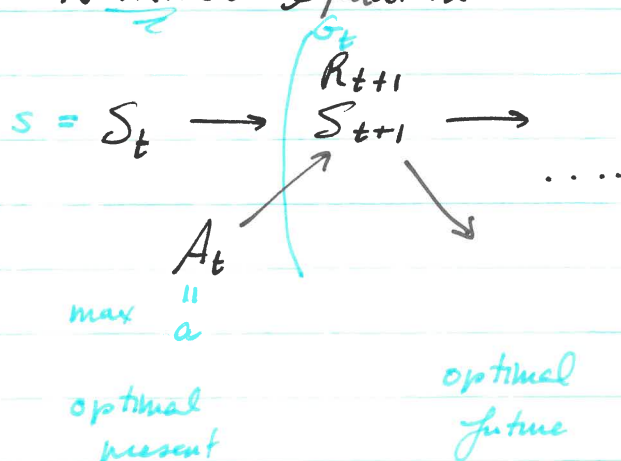
3.6 Bellman optimality equations

See 3-12b

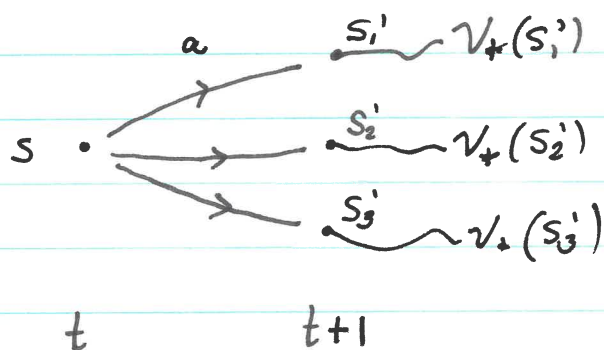
Value function:

$$V_*(s) = \max_a E[R_{t+1} + \gamma V_*(S_{t+1}) | S_t = s, A_t = a]$$

- Optimal from s = optimal action from s and then optimal policy after
- Expectation in one step - no E_π
- max outside expectation - influences reward
- arg max defines $\Pi_*(s)$
- Nonlinear equation because of max



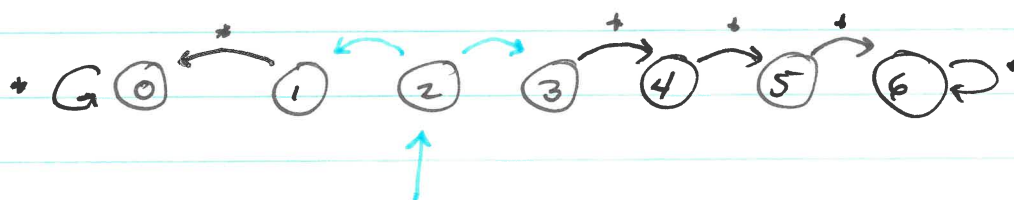
$$V_*(s) = \max_a \rho(s, a) + \gamma \sum_{s'} V_*(s') p(s'|s, a)$$



- Whatever action chosen at time t , rest of dynamics is optimal
- To be optimal at $S_t = s$, choose optimal action and then follow optimal policy from S_{t+1}

$$\Rightarrow V_*(s) = \max_a E[R_{t+1} + \gamma V_*(S_{t+1}) \mid S_t = s, A_t = a]$$

Example:



Action Value Function:

$$q_*(s, a) = E[R_{t+1} + \gamma \max_{a'} q_*(S_{t+1}, a') \mid S_t = s, A_t = a]$$

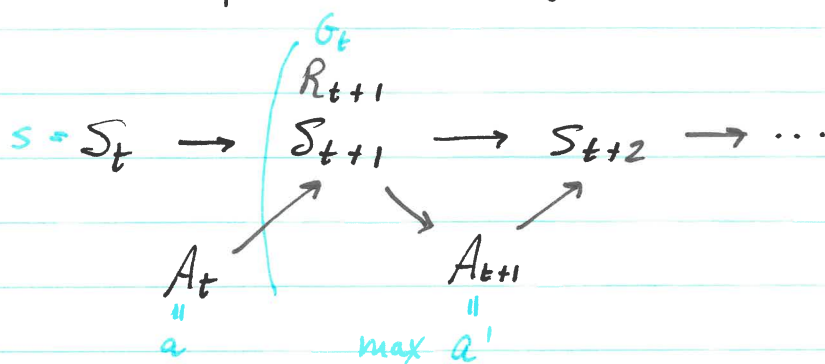
$\uparrow$ a' $\uparrow$ action after

$$= E[R_{t+1} + \gamma v_*(S_{t+1}) \mid S_t = s, A_t = a]$$

i.e.

$$q_*(s, a) = p(s, a) + \gamma \max_{a'} \sum_{s'} q(s', a') p(s' \mid s, a)$$

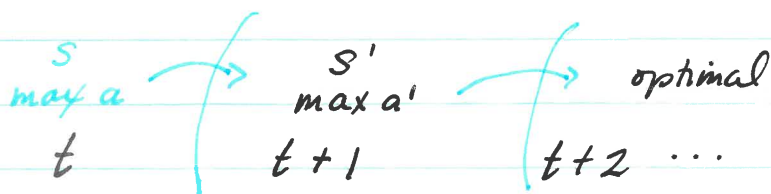
$$= p(s, a) + \gamma \sum_{s'} v_*(s') p(s' \mid s, a)$$



- Optimal from s, a = optimal after reaching S_{t+1} by taking optimal action after
- Max inside - doesn't influence reward
- arg max defines $\pi_*(s')$
- Nonlinear equation because of max

• Bellman's optimality principle:

An optimal policy has the property that whatever the initial state / decision are, the remaining decisions constitute an optimal policy with regard to the state resulting from the first decision.



Chapter 4: Dynamic programming

SB: Chap 4

4-1

4.1 Policy evaluation

- Goal: Compute V_π for given policy π
- Direct method: Solve linear equation

$$V_\pi = \rho_\pi + \gamma P_\pi V_\pi$$

$$\begin{pmatrix} V_\pi \end{pmatrix} = \begin{pmatrix} \rho_\pi \end{pmatrix} + \gamma \begin{pmatrix} P_\pi \end{pmatrix} \begin{pmatrix} V_\pi \end{pmatrix}$$

$$\Rightarrow V_\pi = (I - \gamma P_\pi)^{-1} \rho_\pi$$

- Complexity: $O(|S|^3)$

- Iterative method: $V_0 \rightarrow V_1 \rightarrow \dots \rightarrow V_\pi$

$$V_{k+1}(s) = E_\pi [R_{t+1} + \gamma V_k(S_{t+1}) | S_t = s]$$

$$V_{k+1} = \rho_\pi + \gamma P_\pi V_k$$

Bellman backups

matrix product

$$= T_\pi(V_k)$$

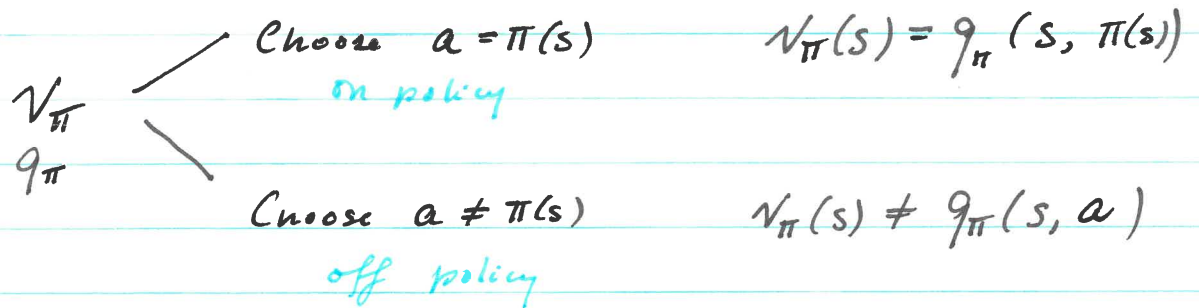
$$T_\pi(v) = \rho_\pi + \gamma P_\pi v$$

Bellman operator

- Initial value V_0 arbitrary, except $V_0(\text{terminal state}) = 0$
- V_π fixed point of T_π
- Convergence: $V_k \rightarrow V_\pi$ as $k \rightarrow \infty$
- T_π is a contraction
- Complexity: $|S|^2 \times M$
matrix product # iterations
- Not exact V_π for finite # iterations

4.2. Policy iteration

- Consider deterministic policies $\pi(s)$



- Policy improvement theorem: Let π and π' such that

$$q_\pi(s, \pi'(s)) \geq q_\pi(s, \pi(s))$$

off policy
on policy

for all $s \in \mathcal{S}$. Then π' is as good or better than π in the sense that

$$V_{\pi'}(s) \geq V_\pi(s)$$

better value under π'

for all $s \in \mathcal{S}$.

- Greedy policy selection:

$$\pi'(s) = \arg \max_a q_\pi(s, a)$$

- Equation for updating π
- Fixed point: π_* for q_* (optimal policy)
- Strict improvement unless optimal

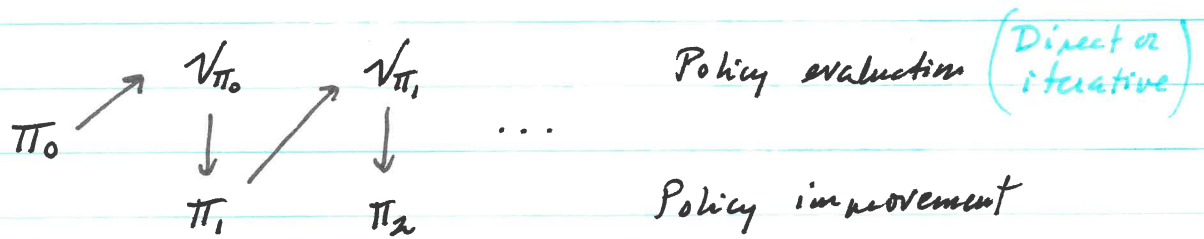
- ϵ -greedy policy selection:

$$\pi'(s) = \begin{cases} \arg \max_a q_\pi(s, a) & \text{Prob } 1-\epsilon \\ \mathcal{U}(\mathcal{S}') & \text{Prob } \epsilon \end{cases}$$

explanation

- No strict improvement

Policy iteration algorithm



Convergence:

$$\begin{aligned} \pi_k &\rightarrow \pi_* \\ V_k &\rightarrow V_* \\ q_k &\rightarrow q_* \end{aligned} \quad \text{as } k \rightarrow \infty$$

Note: Policy evaluation for q_π

$$\begin{aligned} q_\pi(s, a) &= r(s, a) + \gamma \sum_{a', s'} q_\pi(s', a') \pi(a'/s') p(s'/s, a) \\ &= r(s, a) + \gamma \sum_{s'} q_\pi(s', \pi(s')) p(s'/s, a) \\ &= r(s, a) + \gamma \sum_{s'} V_\pi(s') p(s'/s, a) \end{aligned}$$

on policy

$$\rightarrow q_\pi = r + \gamma P_a V_\pi$$

Note:

Planning

- Model known
- Solve to get V_π
- Solve to get V_*, π_*

$$V_\pi, V_*$$

Learning

- MDP model unknown
- Estimate V_π by exploration/sampling
- Converge to V_*, π_*

$$q_\pi, q_*$$

4.4. Action value iteration

$$q_{k+1}(s, a) = E[R_{t+1} + \underbrace{\delta \max_{a'} q_k(s_{t+1}, a')}_{\substack{\text{previous} \\ \text{estimate } v_k(s_{t+1})}} \mid S_t = s, A_t = a]$$

time $t+1$
↓ ↓

update

$$q_{k+1}(s, a) = p(s, a) + \delta \max_{a'} \sum_{s'} q_k(s', a') p(s' | s, a)$$

- Max defines policy improvement at stage $k+1$
- Fixed point: q_* Bellman optimality equation
- Convergence: $q_k \rightarrow q_*$ as $k \rightarrow \infty$

• Policy improvement: $\pi_{k+1}(s) = \arg \max_a q_k(s, a)$

• Variant:

$$\underbrace{q_k(s, a)}_{q_k} = E[R_{t+1} + \underbrace{\delta v_k(s_{t+1})}_{v_k} \mid S_t = s, A_t = a]$$

 $q_k \leftarrow$ v_k

$$q_k(s, a) = p(s, a) + \sum_{s'} v_k(s') p(s' | s, a)$$

$$= p(s, a) + \delta (P_a v_k)(s)$$

action a

- Max included in v_k

Chapter 5: Temporal difference methods (TD)

Planning

- MDP model known/given
- Solve to get V_* , q_* , π_*
 - Bellman opt. eq.
 - Policy iteration
 - Value iteration

Learning

- MDP model unknown
- Estimate V_* , q_* , π_*
 - Sampling
 - Exploration
 - RL

Prediction

- Predict effect of policy
- Policy evaluation V_π

Control

- Find optimal policy
- Find V_* , q_*

needed for learning

- Idea of TD:

- Use observed/visited states to estimate V_* iteratively
- Sampling / exploration informs estimation
- New states / actions depend on estimated V_*

feedback

$$V_0 \rightarrow V_1 \rightarrow V_2 \rightarrow \dots \rightarrow V_*$$

$$S_0 \xrightarrow{R_1} S_1 \xrightarrow{R_2} S_2 \rightarrow \dots$$

$$V_0 \searrow V_1 \nearrow V_2 \nearrow$$

TD

See HC course

- Stochastic approximation:

- Iteration, map: $X_{n+1} = T(X_n)$

$$X_n \rightarrow X^*$$

$X^* = T(X^*)$ attractive fixed point

- Stochastic iteration: $X_{n+1} = T(X_n)$

$$\begin{aligned} X_{n+1} &= (1 - a_n) X_n + a_n T(X_n) \\ &= X_n + a_n (T(X_n) - X_n) \end{aligned}$$

difference
TD error $\rightarrow 0$

$$a_n \geq 0$$

$$\sum a_n = \infty$$

$$\sum a_n^2 < \infty$$

5.1 TD(0)

· Policy iteration (prediction, not control)

· Bellman equation:

$$V_{\pi}(s) = E_{\pi}[R_{t+1} + \gamma V_{\pi}(S_{t+1}) | S_t = s]$$

$$V_{\pi} = \rho_{\pi} + \gamma P_{\pi} V_{\pi} = T(V_{\pi})$$

Bellman operator

· TD(0) iteration:

TD error

$$\begin{aligned} V_{k+1}(S_t) &= V_k(S_t) + \alpha_k [R_{t+1} + \gamma V_k(S_{t+1}) - V_k(S_t)] \\ &= (1 - \alpha_k) V_k(S_t) + \alpha_k [R_{t+1} + \gamma V_k(S_{t+1})] \end{aligned}$$

Bellman iteration update

· SA version of value iteration

· SA:

expectation

$$E_{\pi}[R_{t+1} + \gamma V_{\pi}(S_{t+1}) | S_t = s]$$

do nothing

RV

$$R_{t+1} + \gamma V_{\pi}(S_{t+1})$$

In transition $S_t \rightarrow S_{t+1}$

· α_k : annealing sequence

$\alpha_k = \alpha$ constant : Convergence in mean with episode loop

$$\alpha_k \sim \frac{1}{k}$$

: Convergence in probability

• TD(0) algorithm SB p.120

Input: policy π

Parameters: $\alpha \in (0,1]$, # episodes, # time steps

Initialize $V(s)$

Loop for each episode:

Initialize S

Loop for each time steps:

$A \leftarrow$ action from π for S

Get S', R from action A

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

$S \leftarrow S'$

5.2 Sarsa

• On-policy control iteration for estimating q_*

• Bellman optimality equation:

$$q_*(s, a) = E[R_{t+1} + \gamma \max_{a'} q_*(S_{t+1}, a') \mid S_t = s, A_t = a]$$

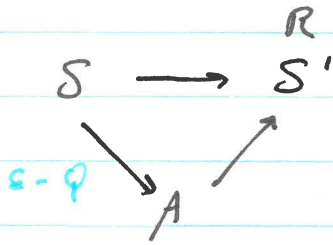
• Sarsa iteration:

$$\begin{aligned} Q_{k+1}(S_t, A_t) &= Q_k(S_t, A_t) + \alpha_k [R_{t+1} + \gamma Q_k(S_{t+1}, A_{t+1}) - Q_k(S_t, A_t)] \\ &= (1 - \alpha_k) Q_k(S_t, A_t) + \alpha_k [R_{t+1} + \gamma Q_k(S_{t+1}, A_{t+1})] \end{aligned}$$

A_t chosen ϵ -greedy from $Q_k(S_t, \cdot)$ before update
 A_{t+1} " " " " $Q_k(S_{t+1}, \cdot)$

↳ define policy π at each stage

- SA of Bellman opt. eq.
- Optimal actions A_t, A_{t+1} from current Q_n
- Use all states: s, a, r, s', a'
- On-policy: policy used to update Q
" policy used to generate new actions
- Convergence: Same as TD(0)



$$\max_a Q(S', a)$$

for update

· Maximization bias:

max estimated $q \neq$ max in expectation

$$\max \hat{f}(\cdot) \neq \max E[f(\cdot)]$$

· Choosing largest value necessarily above expectation

$$\Rightarrow Q \text{ i.e. } \hat{q} \text{ generally } > \text{ true } q_*$$

$$V \text{ i.e. } \hat{v} \quad " \quad " \quad " \quad v_*$$

Conclusions - to go further

- Represent $V(s)$ $Q(s, a)$
 $|S|$ $|S| \times |A|$ SB
Chaps
9-12

- Parameterizations / : $V(s; \theta)$
Projections $Q(s, a; \theta)$

- Neural networks : $s \rightarrow \boxed{}^{\theta} \rightarrow V(s; \theta)$
deep RL $s, a \Rightarrow \boxed{} \rightarrow Q(s, a; \theta)$

- Adapt TD(0), Sarsa, Q-learning
to update parameters θ

- Estimation, sampling : $Q_n \rightarrow Q_{n+1}$
 $V_n \rightarrow V_{n+1}$ SB
Chaps
5-7

- Use longer history
- TD(λ)
- Eligibility traces

- Policy approximation: $\pi(s) \rightarrow \pi(s; \theta)$ SB
Chap 13

- Gradient ascent : $\theta_{n+1} = \theta_n + \beta \nabla J(\theta)$
?

- Policy gradient
- Actn-critic methods



- Other topics
 - Multi-agent
 - Partially observed MDPs : POMDPs
 - etc.