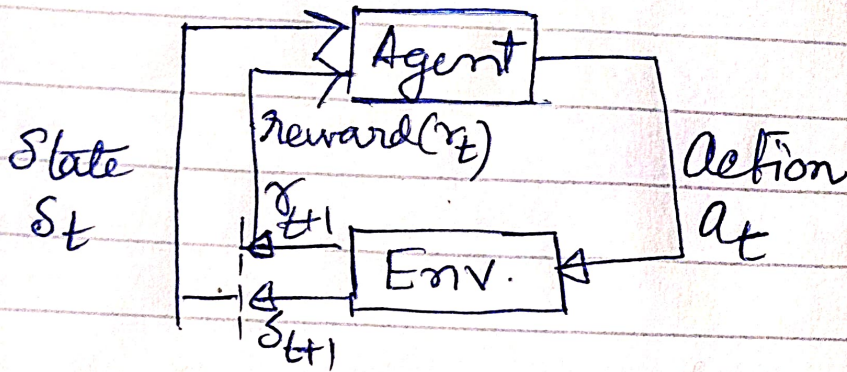


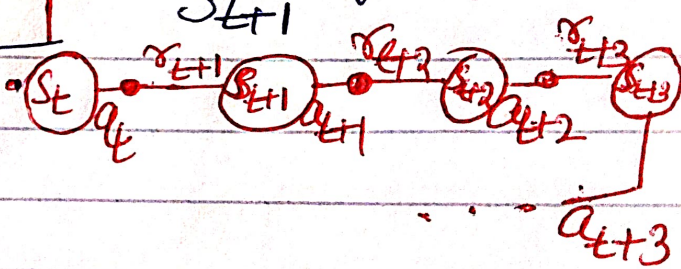
# MDP - Markov Decision Process

## The agent-env. Interface



Agent & environment interact at discrete time steps:  $t = 0, 1, 2, \dots$

- Agent observes state at step  $t$ :  $s_t \in S$
- Produces action at step  $t$ :  $a_t \in A(s_t)$
- Gets resulting reward:  $r_{t+1} \in \mathbb{R}$
- & resulting next state:  $s_{t+1}$



- Close loop
- Few assumptions:
  - discrete time steps:  $t = 0, 1, 2$  (decision steps, not in fraction)
  - decision steps  $\equiv$  opportunity  $\equiv$  not exactly how long

- 1 decision  $\sim$  one time step.

"S" - set of step states (discrete)

$\approx 1$  to  $N$ . [ $s_0 = 3 \rightarrow$  not exactly always zero<sup>th</sup> step.  
 $s_1 = 7$ ]

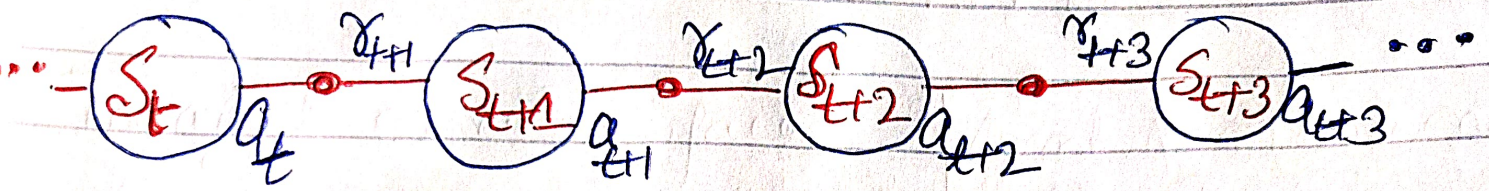
- $a_t$  is applied to Environment, causes the

environment state to change to  $s_{t+1}$ ; also produces the reward  $r_{t+1}$  [which is a result of taking  $a_t$  in state  $s_t$ , & moving to  $s_{t+1}$ ]

- It can be joint distribution

$$p(\text{next state}, r | s, a) \rightarrow p(s', r | s, a)$$

∴ What is happening bet<sup>n</sup> env & agent.



At, state  $S_t$  take action  $a_t$ , get reward  $r_{t+1}$  & move to state  $S_{t+1}$ . Then again take action  $a_{t+1}$  & get reward  $r_{t+2}$  & move to  $S_{t+1}$   
... So on ...

This is called as "trajectory" to the set of states.

Set of states comes from " $S$ " & set of discrete actions come from set of actions  $A$ .

∴ Remember ' $S$ ' is set of state space (discrete)  
Set  $A$  is a set of discrete actions;

The reward is a  $r_{t+1}$  which is a real valued scalar quantity.

∴ "The first assumption we made is Markov Property" to make things easier here.

"What is Markov Property"?

① Here "the state" at step  $t$ , means whatever information is available to the agent at step  $t$  about its environment.

② The state can include immediate "sensations" highly processed sensations & structures built up over time from sequences of sensations.

③ Ideally, a state should summarize past sensations so as to retain all "essential" information, i.e. it should have the "Markov Property".

$$\text{Pr}\{S_{t+1}=S', r_{t+1}=r \mid S_1, a_1, r_1, S_{t-1}, a_{t-1}, \dots, r_t, S_t, a_t\}$$

$$\text{Pr}\{S_{t+1}=S', r_{t+1}=r \mid S_1, a_1\}$$

\*  $S', r$ , and histories

$$S_1, a_1, r_1, S_{t-1}, a_{t-1}, \dots, r_t, S_t, a_t$$

② State could be Agent's position (robot's coordinate velocity & direction; current speed & rate in grid(x) & movement direction; sensor readings; distance to obstacles, object proximity; Energy level; battery percentage for a drone or robots etc]

"Essential" information

$$Pr \{ S_{t+1} = s', r_{t+1} = r \mid S_1, a_1, r_1, S_2, a_2, r_2, \dots, S_t, a_t, r_t, S_{t-1}, a_{t-1}, r_{t-1}, S_0, a_0, r_0 \}$$

So, the assumption is that the state should have all essential informations to predict the future (behaviors)

$$Pr \{ S_{t+1} = s', r_{t+1} = r \mid S_1, a_1, \cancel{r_1}, \cancel{S_2}, \cancel{a_2}, \dots, S_t, a_t, r_t, S_{t-1}, a_{t-1}, r_{t-1} \}$$

If current state  $S_t$  & action  $a_t$  of the prev. <sup>comp</sup> are not needed.

\* Not dependent on history, current state & actions are enough to predict the system behaviour. — Markov Property.

ex:  $S_t = \emptyset, S_t = \{ \}$  [no need of knowing history]

∴  $P(s', r \mid s, a) \rightarrow$  Markov property.

## Markov Decision Process

- MDP,  $M$  is the tuple:  $M = \langle S, A, p, r \rangle$ 
  - $S$ : set of states (finite discrete)
  - $A$ : set of actions (finite discrete)
  - $p: S \times A \times S \rightarrow [0,1]$ : probability of transition (1)
  - $r: S \times A \times S \rightarrow \mathbb{R}$ : Expected reward
    - $[p(r|s, a, s')]$ : instead saying we going to develop this distribution, we are finding expected value (2)
- Policy:  $\pi: S \times A \rightarrow [0,1]$  (can be deterministic)
  - $[ \pi \text{ is a mapping from state to action } ]$  (3)

Maximize total expected reward.

- Learn an optimal policy.

lets say we are at current state  $S_1$ , we do action  $a_1$ ; what is the probability that we will have another state  $S'$ . So for every possible  $S'$ , we will be having a value. For every combination of  $S_1$  &  $a_1$ , we will have a value (probability). For no. of states you have on values.

$p \approx Pr(S' | S_1, a_1)$

$r(S, A, S') = \int r \cdot p(r | S, A, S') dr$

Some time joint distribution also  
 $p(s', r | s, a)$

Policy  $\pi$ : The role of the agent is to learn the policy  $\pi$ , which is a mapping from state to action.

$$\pi: S \times A \rightarrow [0, 1]$$

$$\pi(a|s) \text{ or } \pi(s, a) = \text{pr}(a_t = a | s_t = s)$$

For every state you could have 'm' actions; some time transition is probabilistic/stochastic; the policy  $\pi$  can be deterministic

$$\pi(a|s) = 1 \text{ for } a = a_1 \\ = 0 \text{ otherwise}$$

Some times,  
 $\pi(s) = a_1$  (deterministic).

The goal of the agent is to maximize the total expected reward: is called optimal policy.