

# Introduction to Reinforcement Learning based Control and Deep Reinforcement Learning based control



## Research

Centre de Recherche en Automatique de Nancy  
(CRAN) UMR 7039,  
Faculté des Sciences et Technologies  
Boulevard des Aiguillettes - BP 70239 - Bât.  
1er cycle 54506 Vandoeuvre-lès-Nancy  
Cedex France

Associate Professor  
(Maitre de Conférences)

Automatic Control, Reliability of Systems  
[mayank-shekhar.jha@univ-lorraine.fr](mailto:mayank-shekhar.jha@univ-lorraine.fr)



## Teaching

Polytech Nancy (ESSTIN),  
2 Rue Jean Lamour 54509  
Vandoeuvre-lès-Nancy,  
Cedex France



Email: [mayank-shekhar.jha@univ-lorraine.fr](mailto:mayank-shekhar.jha@univ-lorraine.fr)

# Introduction

## Markov Decision process

# Markov Decision Process (MDP)

Markov process

Markov Reward Process

Markov Decision Process

# MDP

- Markov Decision Process (MDP)
  - describes fully an environment for RL.
  - current state characterises the process
  - all essential information is contained by MDP
- Almost all RL problems can be formulated by MDPs.

# Markov Property

Definition:  $S_t$  is Markov if and only if:

$$P(S_{t+1} | S_t) = P(S_{t+1} | S_1, \dots, S_t)$$

- state describes all the relevant information.
- state is known  $\rightarrow$  history can be thrown away.

# Markov Property

Consider  $s$  and  $s'$

State transition probability

$$P(s, s') = \mathbb{P} [S_{t+1} = s' \mid S_t = s]$$

State transition Matrix  $P$

$$P = \begin{matrix} & \text{to} \\ \text{from} & \begin{bmatrix} p_{11} & \dots & p_{1n} \\ \vdots & & \vdots \\ p_{n1} & & p_{nn} \end{bmatrix} \end{matrix}$$

# Markov Process

- is a tuple  $(S, P)$
- $S \rightarrow$  finite set of states
- $P \rightarrow$  state transition matrix.

## Markov Reward Process.

- is a tuple  $(S, P, R, \gamma)$
- $R \rightarrow$  reward function  $R_S = \mathbb{E} [R_{t+1} | S_t = S]$
- $\gamma \rightarrow$  is discount factor.  $\gamma \in [0, 1]$

# Return

Return  $G_t \rightarrow$  total discounted reward from time-step  $t$ .

$$G_t = R_{t+1} + \gamma R_{t+2} - \dots = \sum_{k=0}^{\infty} \gamma^k \cdot R_{t+k+1}$$

- $\gamma \in [0, 1]$

- Value of reward ' $R$ ' after ' $k+1$ ' time-steps is  $\gamma^k \cdot R$ .

- If

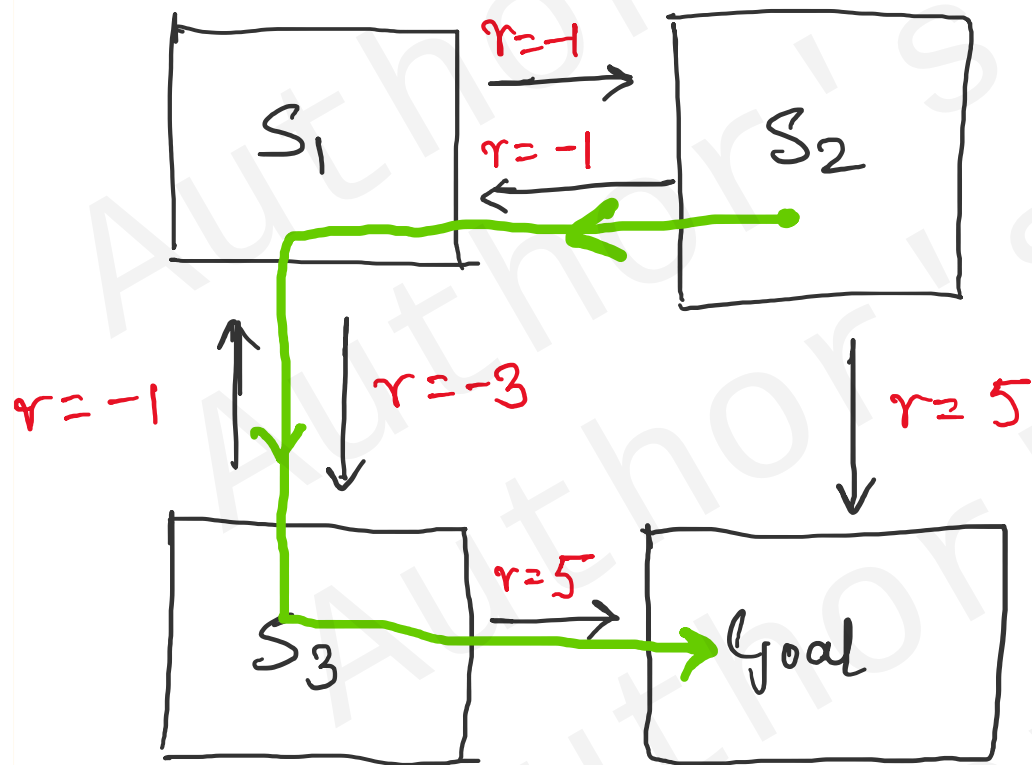
- $\gamma = 0 \rightarrow$

- near sightedness

- $\gamma = 1 \rightarrow$

- far-sighted evaluation.

# Grid World Example



$$V_{\pi}(s) = \mathbb{E}_{\pi} [ \underbrace{R_{t+1} + R_{t+2} + R_{t+3} + \dots}_{\text{future rewards}} | S_t = s ]$$

Assume: A movement from  $S_2$  as:

$$\begin{aligned} V_{\pi}(S_{2,t}) &= \mathbb{E}_{\pi} ( -1 + -3 + 5 \mid S_t = S_2 ) \\ &= (-1 + -3 + 5) = 1 \end{aligned}$$

→ we consider only one  $\pi$   
 if,  $\pi = \{ \pi_1, \pi_2, \pi_3, \dots \}$   
 then average over all  $\pi$ .

## Value Function

# Value Function

$V(s) \rightarrow$  long term value of state ' $s$ '.

State value function is expected return starting from state ' $s$ '.

$$V(s) = \mathbb{E} [G_t \mid s_t = s]$$

# Value Function

- Value Function can be divided into two parts
  - Immediate reward  $R_{t+1}$
  - discounted value of successor state  $\gamma V(S_{t+1})$

$$\begin{aligned} V(s) &= \mathbb{E} [G_t \mid S_t = s] \\ &= \mathbb{E} [R_{t+1} + \gamma R_{t+2} + \gamma^2 R_{t+3} \dots \mid S_t = s] \\ &= \mathbb{E} [R_{t+1} + \gamma (R_{t+2} + \gamma R_{t+3} + \gamma^2 R_{t+4} \dots) \mid S_t = s] \\ &= \mathbb{E} [R_{t+1} + \gamma G_{t+1} \mid S_t = s] = \mathbb{E} [R_{t+1} + \gamma V(S_{t+1}) \mid S_t = s] \end{aligned}$$

# Value Function

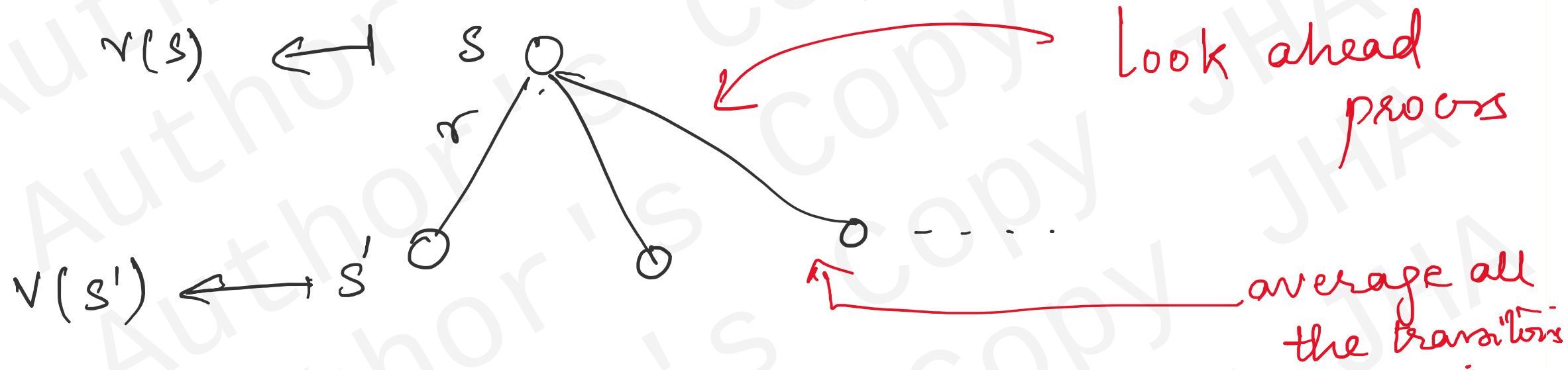
$$v(s) = \mathbb{E} [R_{t+1} + \gamma \cdot v(s_{t+1}) \mid s_t = s]$$

↑  
Reward  
at  $t+1$ .

for any  $s \in S$   
↓  
Value for any  $s \in S$  at time  $t$

(This is the Iterative Equation)

But passage from  $s \rightarrow s'$  is a probabilistic process.  
 (time  $t$ ) (time  $t+1$ )



$$V(s) = R_s + \gamma \sum_{s' \in S} P(s, s') \cdot V(s')$$

# Bellman Equation

$$V = R + \gamma P V \quad (\text{using matrices})$$

$$\begin{bmatrix} V(1) \\ \vdots \\ V(n) \end{bmatrix} = \begin{bmatrix} R_1 \\ \vdots \\ R_n \end{bmatrix} + \begin{bmatrix} P_{11} & \dots & P_{1n} \\ \vdots & & \vdots \\ P_{n1} & \dots & P_{nn} \end{bmatrix} \begin{bmatrix} V(1) \\ \vdots \\ V(n) \end{bmatrix}$$

# Markov Decision Process.

Markov process with Decisions!!

$(S, A, P, R, \gamma)$

↓  
action!

= Decision (what to do when  
in a state to obtain  
the objective)

# Markov Decision Process. $(S, A, P, R, \gamma)$

Markov process with Decisions!!

- $S \rightarrow$  set of finite states
- $A \rightarrow$  set of finite actions
- $P \rightarrow$  state transition prob. matrix

$$P_a(s, s') = \mathbb{P}[S_{t+1} = s' \mid S_t = s, A_t = a]$$

- $R$  is reward function,

$$R_a(s) = \mathbb{E}(R_{t+1} \mid S_t = s, A_t = a)$$

## Policy

# Policy

A policy  $\pi$  is a distribution over actions given states:

$$\pi(a/s) = \mathbb{P}[A_t = a \mid S_t = s]$$

→ A policy defines the behaviour of an agent.

→ Policies are stationary

$$A_t \sim \pi(\cdot \mid S_t), \forall t > 0$$

# Policies

- Given an MDP  $M = (S, A, P, R, \gamma)$  and a policy  $\pi$
- State sequence  $S_1, S_2, \dots$  Markov Process  $(S, P^\pi)$
- State and reward sequence  $(S, P^\pi, R^\pi, \gamma)$  with respect to policy  $\pi$  is a Markov Reward Process.

$$P^\pi(s, s') = \sum_{a \in A} \pi(a|s) P_a(s, s')$$

Average over all possible actions and associated probs.

$$R^\pi(s, s') = \sum_{a \in A} \pi(a|s) R_a(s)$$

Average over all possible actions.

→ Given a MDP (with respect to various policies)

↓ given a policy  $\pi$   
MRP.

Value Functions → V-functions and Q-functions

# Value Function (V-function and Q-function)

## State Value Function (V function)

V-function  $V_{\pi}(s) \rightarrow$  expected 'return' from state 's' and following policy  $\pi$

$$V_{\pi}(s) = \mathbb{E}_{\pi} [G_t \mid S_t = s]$$

(Remember:  
 $G_t = R_{t+1} + \gamma R_{t+2} + \dots$   
 $= \sum_{k=0}^{\infty} \gamma^k \cdot R_{t+k+1}$ )

## Action Value Function (Q-function)

Q-function  $\rightarrow$  expected return starting from state 's', taking action 'a', and follow policy  $\pi$ .

$$q_{\pi}(s, a) = \mathbb{E}_{\pi} [G_t \mid S_t = s, A_t = a]$$

V-function can be decomposed  $\rightarrow$  immediate reward  
+ all future states/rewards

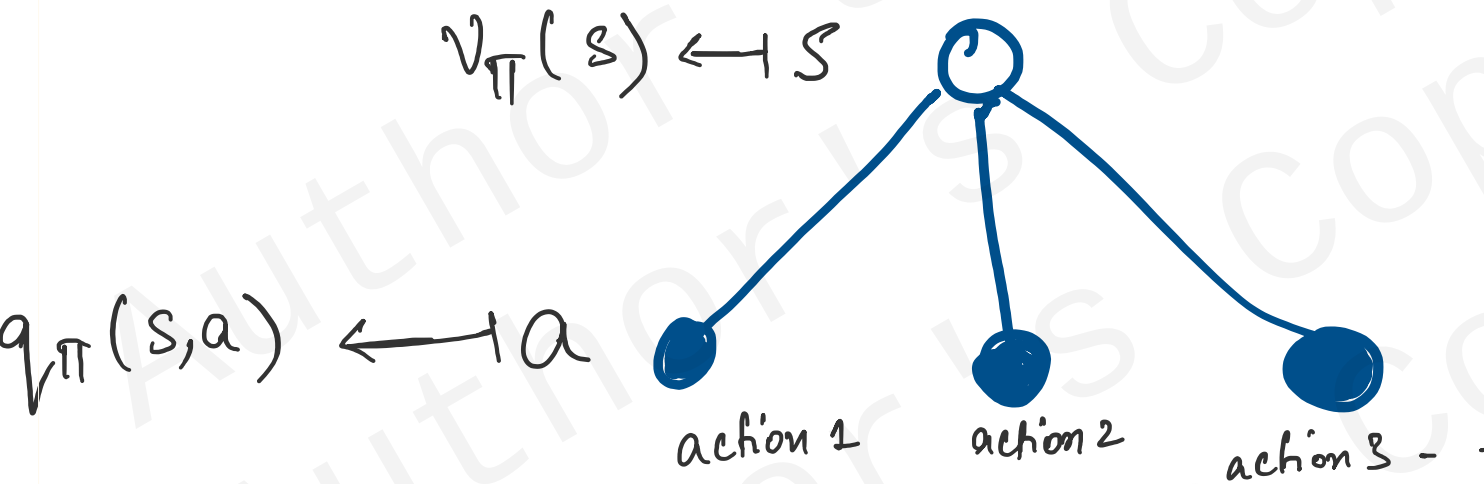
$$V_{\pi}(s) = \mathbb{E}_{\pi} \left( R_{t+1} + \gamma V_{\pi}(S_{t+1}) \mid S_t = s \right)$$

Q-function can also be decomposed.

$$Q_{\pi}(s, a) = \mathbb{E}_{\pi} \left( R_{t+1} + \gamma Q_{\pi}(S_{t+1}, A_{t+1}) \mid S_t = s, A_t = a \right)$$

- $\rightarrow$  take some action from state  $s$
- $\rightarrow$  get immediate reward for that action
- $\rightarrow$  ask Q-value of next state we end up in,  
under the next action chosen from that time.

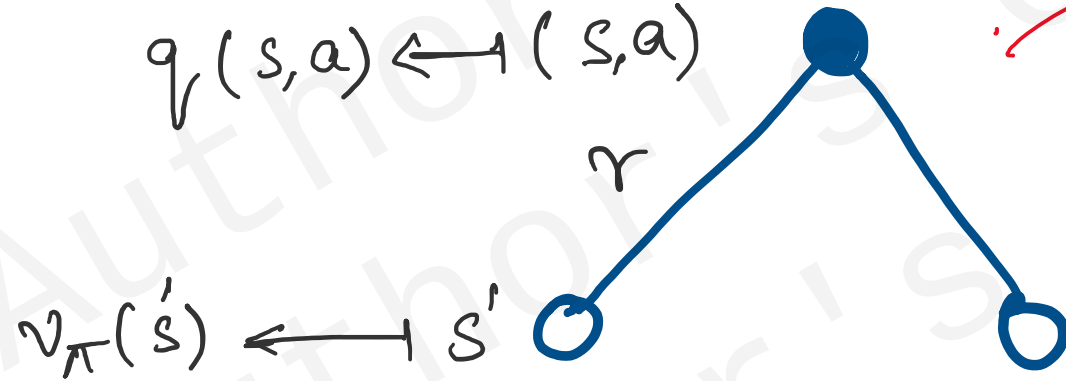
# How are V-function and Q-function related?



$$V_{\pi}(s) = \sum_{a \in A} \pi(a/s) \cdot Q_{\pi}(s,a)$$

- from state 's' under policy  $\pi$
- three actions taken.
- each action has action value.
- average over all such values → value  $V_{\pi}(s)$  of being in state 's'. (at top)

# How to interpret $Q^\pi$



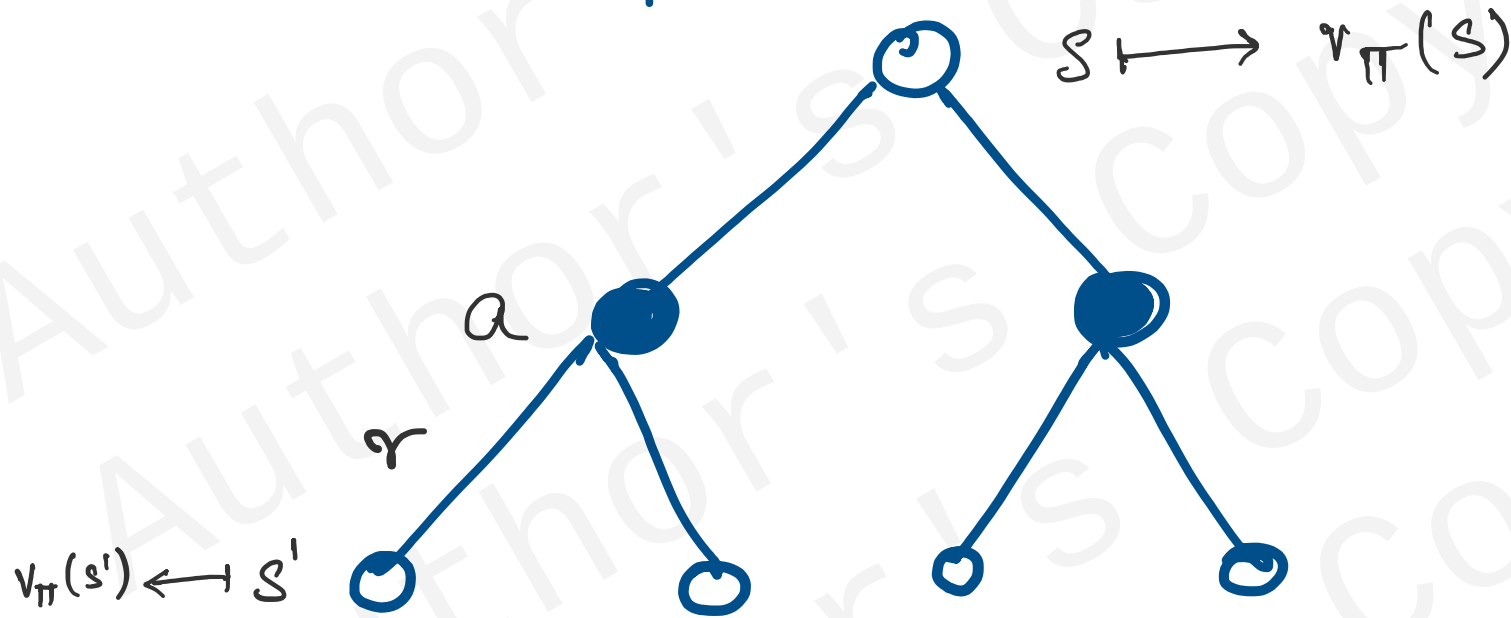
→ When in state 's'  
how good is it to take  
action 'a'.

after transition  
with reward 'r'  
how good to  
be in that state.

$$q_\pi(s,a) = R_a(s) + \gamma \sum_{s' \in S} P_a(s,s') \cdot v_\pi(s')$$

# Bellman Equation

# Bellman Equation (Expectation) for $V_\pi$



$$V_\pi(s) = \sum_{a \in A} \pi(a|s) \left( R_a(s) + \gamma \sum_{s' \in S} P_a(s, s') V_\pi(s') \right)$$

# Bellman Equation (Expectation) for $V_\pi$

$$V_\pi(s) = \sum_{a \in A} \pi(a|s) \left( R_a(s) + \gamma \sum_{s' \in S} P_a(s, s') V_\pi(s') \right)$$

Value of being  
in state 's'

probability  
of taking  
action a in  
s  
Average over  
all actions

Reward  
when  
in s

Average over  
all states

Prob of being in  
s' from s

Value of  
being in  
s'

# Bellman Equation (Expectation) for $V_\pi$

$$V_\pi(s) = \sum_{a \in A} \pi(a|s) \left( R_a(s) + \gamma \sum_{s' \in S} P_a(s, s') V_\pi(s') \right)$$

Value of being  
in state 's'

at time 't'

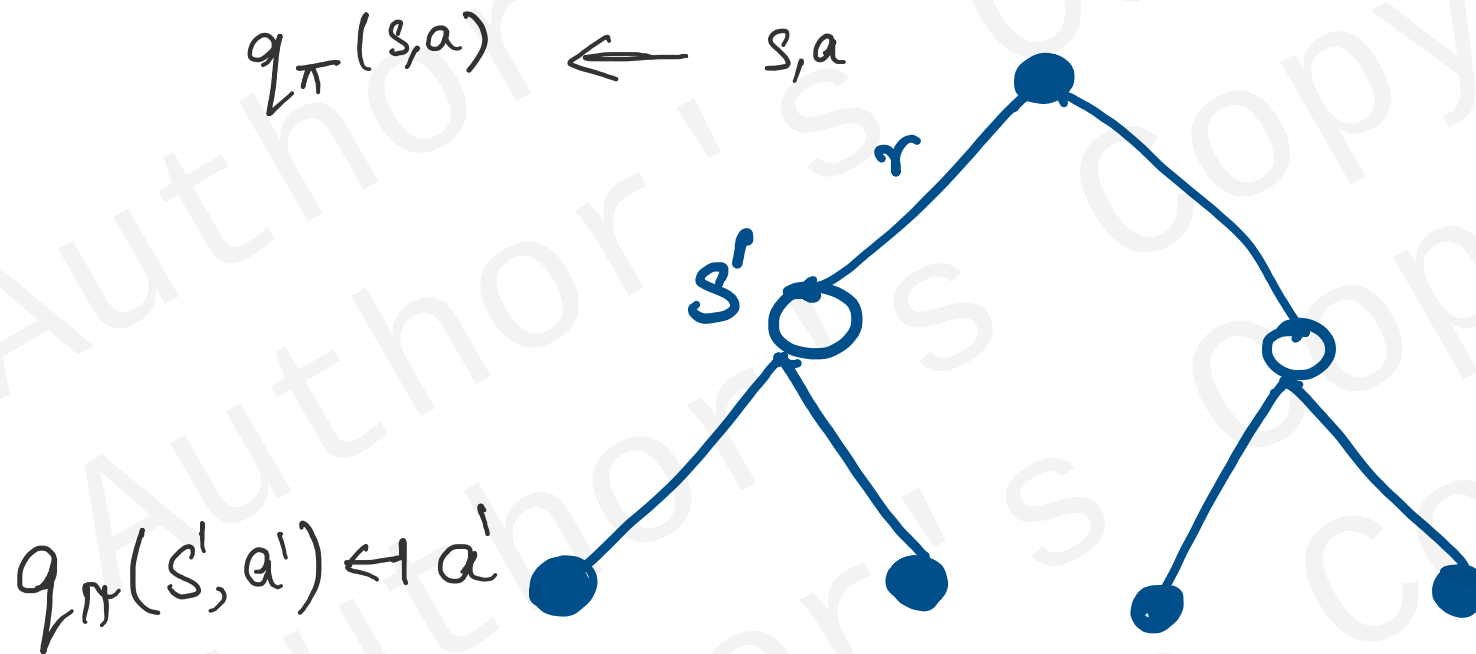
sum over  
all actions

sum over  
all states!!

Value of being  
in state

s' at  $t+1$

# Bellman Equation for Q-function.



$$q_{\pi}(s, a) = R_a(s) + \gamma \sum_{s' \in S} P_a(s, s') \sum_{a' \in A} \pi(a' | s') q_{\pi}(s', a')$$

# Bellman Equation for Q-function.

$$q_{\pi}(s, a) = R_a(s) + \gamma \sum_{s' \in S} P_a(s, s') \sum_{a' \in A} \pi(a' | s') q_{\pi}(s', a')$$

Value of being  
in state s  
with respect to  
action a

sum over all  
states

Probability of  
passing from  
s to s'

Prob of  
action a'  
when in  
state s'

Value of being  
in s' with  
respect to  
a'.

s → at time t  
s' → at time t+1.

# Bellman Eq in Matrix form.

$$V_{\pi} = R_{\pi} + \gamma P^{\pi} \cdot V_{\pi}$$

with solution

$$V_{\pi} = (I - \gamma P^{\pi})^{-1} \cdot R^{\pi}$$

$$V = R + \gamma P V \quad (\text{using matrices})$$

$$\begin{bmatrix} V(1) \\ \vdots \\ V(n) \end{bmatrix} = \begin{bmatrix} R_1 \\ \vdots \\ R_n \end{bmatrix} + \begin{bmatrix} P_{11} & \dots & P_{1n} \\ \vdots & & \vdots \\ P_{n1} & \dots & P_{nn} \end{bmatrix} \begin{bmatrix} V(1) \\ \vdots \\ V(n) \end{bmatrix}$$

If summation over all states be achieved,  
if all values are known  $\rightarrow$  this eq can be solved!

# Optimal Value Function & Optimal Policy

# Optimal Value Function

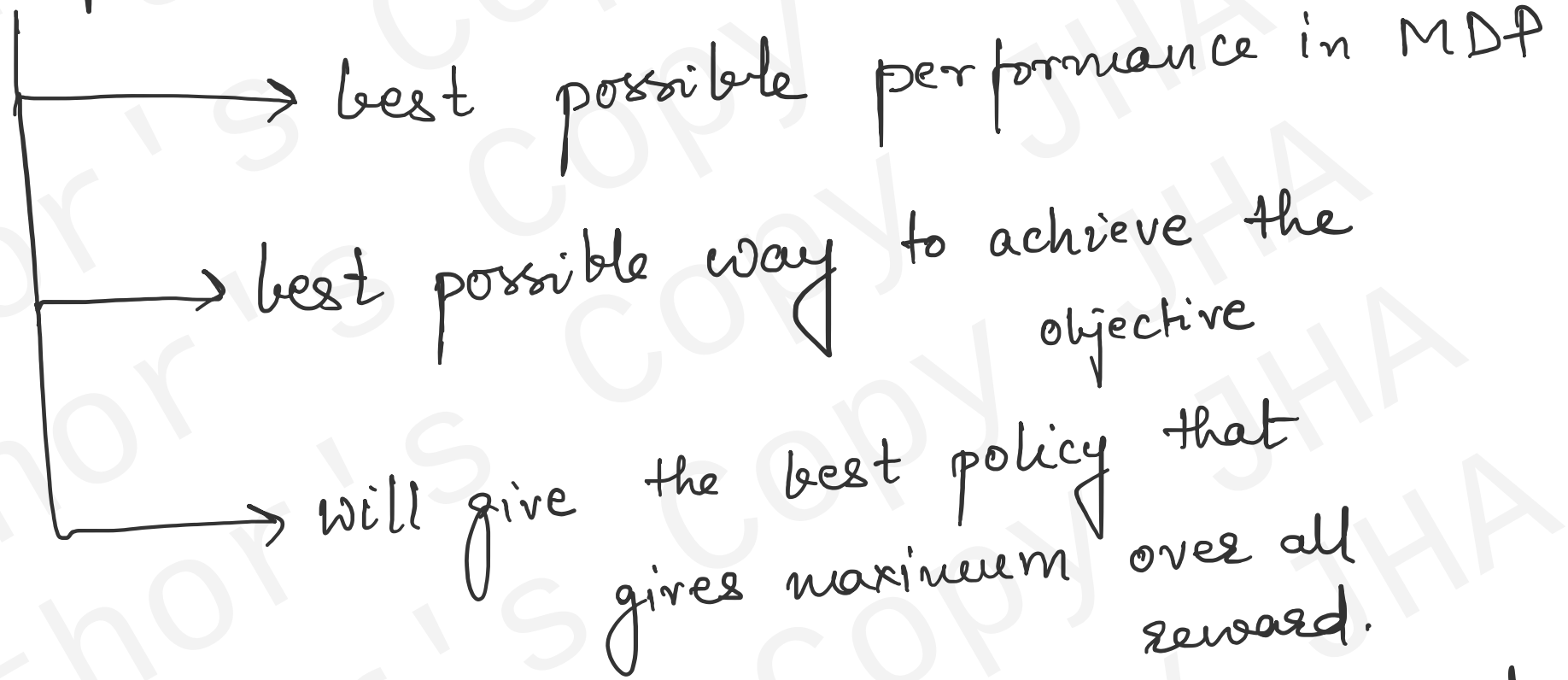
Optimal state-value function  $V_*(s) \rightarrow$  maximum V-function over all policies.

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

Optimal action-value function  $q_*(s) \rightarrow$  max Q-fn over all policies

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

# Optimal Value function



MDP is solved when optimal value function is found.

# Optimal Policy

Partial Ordering:

$$\pi \geq \pi' \text{ if } v_{\pi}(s) \geq v_{\pi'}(s); \forall s$$

Theorem: For any MDP

- ① There exists  $\pi_{\star}$  such that  $\pi_{\star} \geq \pi, \forall \pi$
- ② All optimal policy achieve optimal value function (V-fun)  
ie.  $v_{\pi_{\star}}(s) = v_{\star}(s)$
- ③ All optimal policies achieve optimal Q-function (Q-fun)  
 $q_{\pi_{\star}}(s, a) = q_{\star}(s, a)$

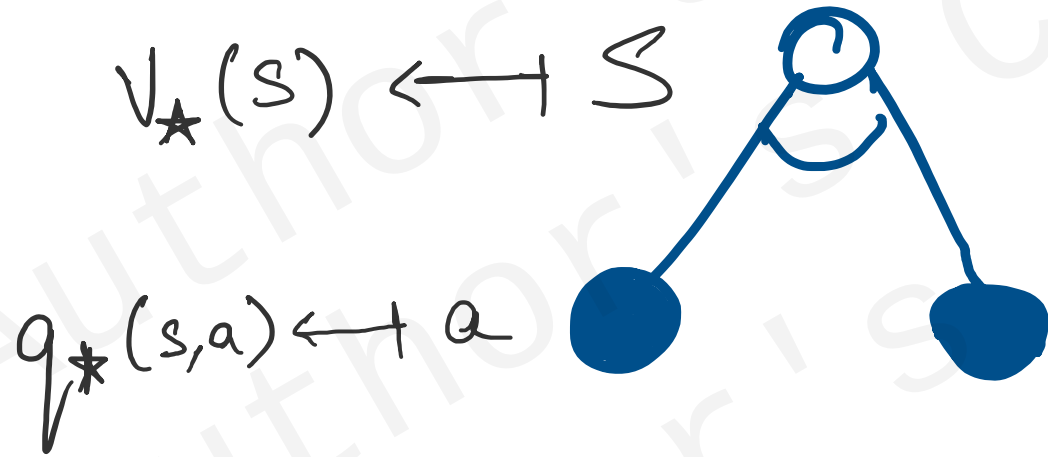
# How to find Optimal Policy?

By maximising over  $q_*(s,a)$

$$\pi_*(a|s) = \begin{cases} 1 & \text{if } a = \underset{a \in A}{\operatorname{argmax}} q_*(s,a) \\ 0 & \text{otherwise} \end{cases}$$

Note: For any MDP, there is always an optimal policy  
If  $q_*(s,a)$  is known, optimal policy can be found directly!!

# How is $V_*$ and $q_*$ related?



From  $q_*$

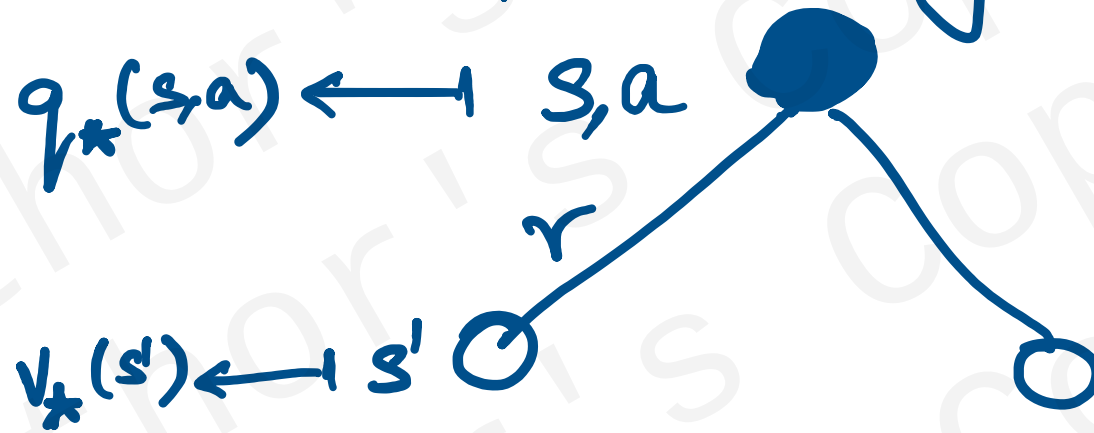
↓

$V_*$

$$V_*(s) = \max_a q_*(s,a)$$

Find max value of  $q(s,a)$  for all  $a \in A$

# Bellman Optimality for $Q_*(\cdot)$

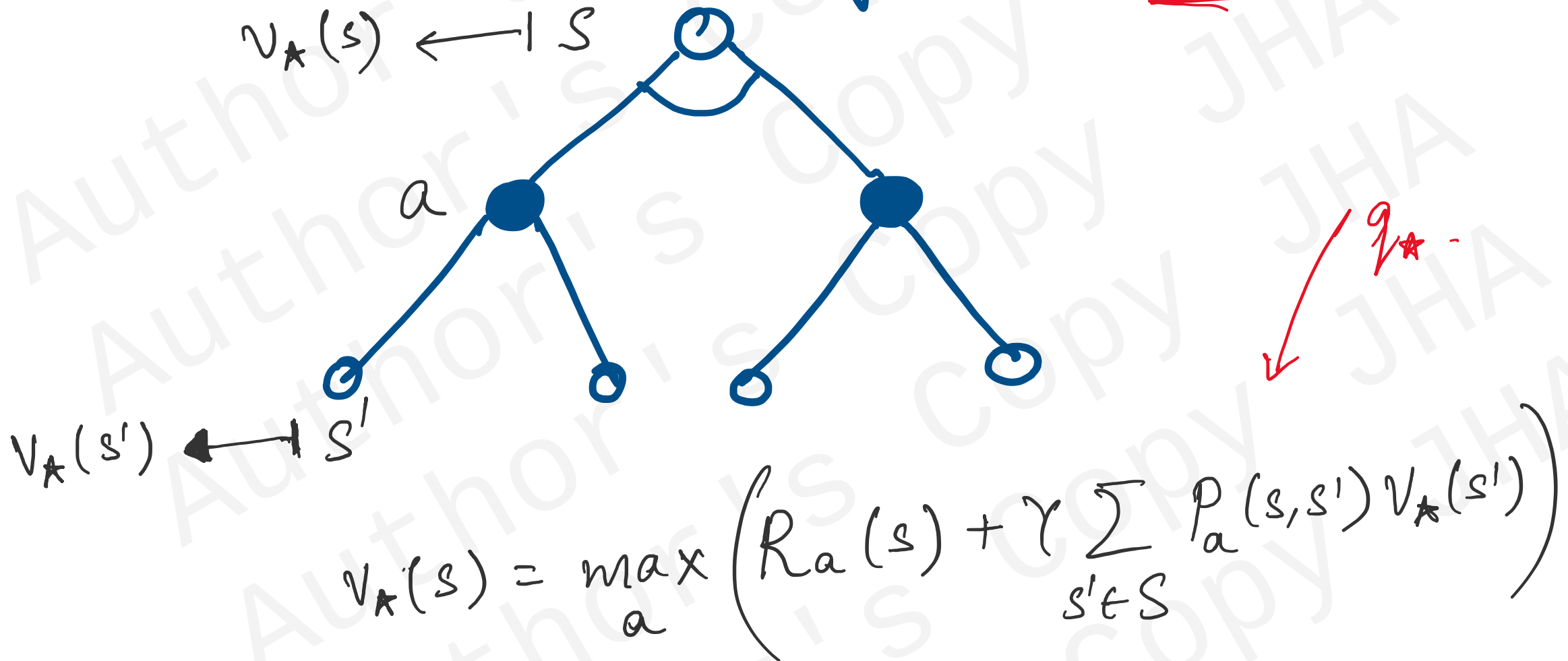


$q_*$  is given as

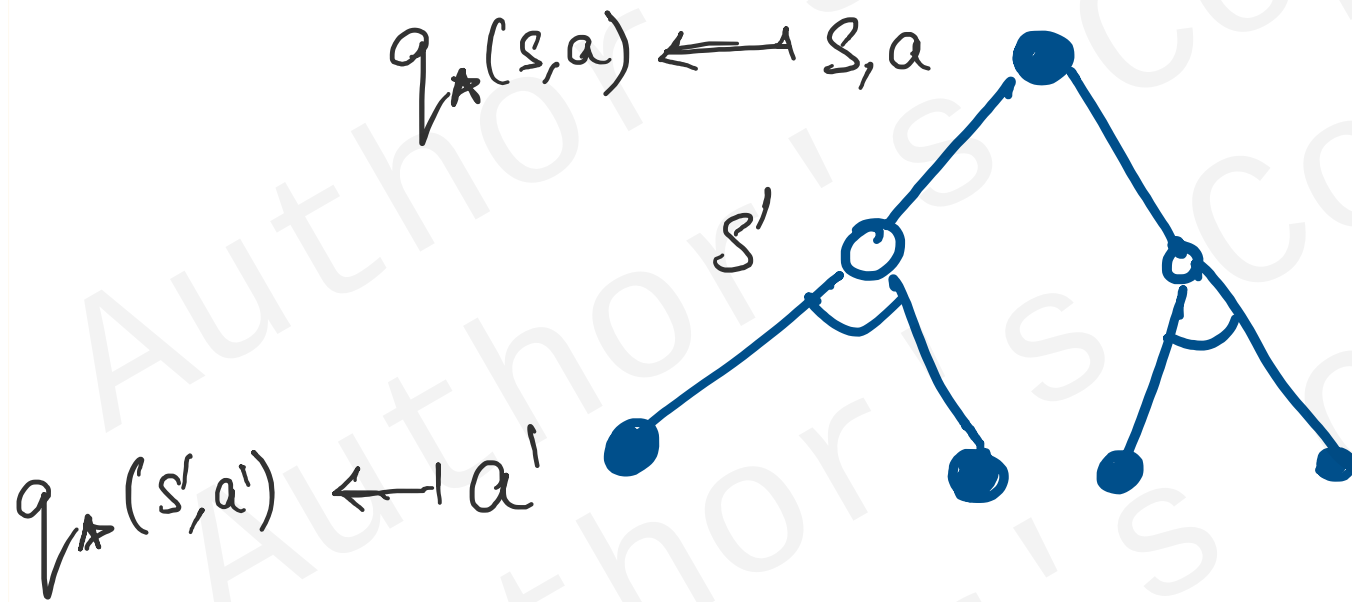
$$q_*(s,a) = R_a(s) + \gamma \sum_{s' \in S} P_a(s,s') v_*(s')$$

Average over all states

# Bellman Optimality for $v_\star$



# Bellman Optimality for $Q^*$ (2)



Another manner  
to find  $q^*$

$$q^*(s, a) = R_a(s) + \gamma \sum_{s' \in S} P_a(s, s') \max_{a'} q^*(s', a')$$

So far:

- Value functions :  $V$  funs,  $Q$ -funs
- Optimality Equations
- Optimal Value Function :  $V_*(s)$ ,  $Q_*(s)$
- From optimal value function → optimal policy

How to solve optimality equation  
to find  
optimal value → optimal policy?

# Bellman Optimality Equation

→ is non-linear

→ no closed form solution

→ Iterative approach required:

① Value Iteration

③ Q-learning

② Policy Iteration

④ SARSA.

↘ Actor-Critic Methods.