



PRESIDENCY UNIVERSITY

BENGALURU

Roll No.														
----------	--	--	--	--	--	--	--	--	--	--	--	--	--	--

End - Term Examinations – May/June 2026

Date: 14 – 05- 2026

Time: 01:00pm – 04:00pm

School: SOCSE	Program: B.Tech - CAI	
Course Code: CAI2507	Course Name: REINFORCEMENT LEARNING	
Semester: VI	Max Marks: 100	Weightage: 50%

CO - Levels	C01	C02	C03	C04
Marks	24	24	26	26

Instructions:

- (i) Read all questions carefully and answer accordingly.
- (ii) Do not write anything on the question paper other than roll number.

Part A

Answer ALL the Questions. Each question carries 2marks.

10Q x 2M=20M

1.	How do we derive the value function from the Q function?	2 Marks	L2	C01
2.	What are the steps involved in policy iteration?	2 Marks	L1	C01
3.	What is the difference between first-visit MC and every-visit MC?	2 Marks	L2	C02
4.	How does on-policy control differ from off-policy control?	2 Marks	L2	C02
5.	What is the update rule of TD learning?	2 Marks	L1	C03
6.	How does the TD prediction method work?	2 Marks	L1	C03
7.	How does Q learning differ from SARSA?	2 Marks	L2	C03
8.	How do we compute the upper confidence bound?	2 Marks	L1	C04
9.	What are the steps involved in Thompson sampling?	2 Marks	L1	C04
10.	What are contextual bandits?	2 Marks	L1	C04

Part B

Answer the Questions.

Total Marks 80M

11.	a.	Explain the concepts of value function in Reinforcement Learning. Discuss their mathematical foundations of value function with suitable example.	10 Marks	L2	C01																																			
	b.	Explain the various elements of RL with the grid world environment as an example.	10 Marks	L2	C01																																			
Or																																								
12.	a.	Discuss Bellman optimality equation and explore how it is useful for finding the optimal Bellman value and Q functions.	10 Marks	L2	C01																																			
	b.	Outline the working steps of the Value Iteration algorithm. Using the model dynamics of State A, compute the optimal value of State A after the first iteration of Value Iteration. ($\gamma = 1$) <table border="1"><thead><tr><th>State (s)</th><th>Action (a)</th><th>Next State (s')</th><th>Transition Probability $P(s' s, a)$ or $P_{ss'}^a$</th><th>Reward Function $R(s, a, s')$ or $R_{ss'}^a$</th></tr></thead><tbody><tr><td>A</td><td>0</td><td>A</td><td>0.1</td><td>0</td></tr><tr><td>A</td><td>0</td><td>B</td><td>0.8</td><td>-1</td></tr><tr><td>A</td><td>0</td><td>C</td><td>0.1</td><td>1</td></tr><tr><td>A</td><td>1</td><td>A</td><td>0.1</td><td>0</td></tr><tr><td>A</td><td>1</td><td>B</td><td>0.0</td><td>-1</td></tr><tr><td>A</td><td>1</td><td>C</td><td>0.9</td><td>1</td></tr></tbody></table>  Model dynamics of state A	State (s)	Action (a)	Next State (s')	Transition Probability $P(s' s, a)$ or $P_{ss'}^a$	Reward Function $R(s, a, s')$ or $R_{ss'}^a$	A	0	A	0.1	0	A	0	B	0.8	-1	A	0	C	0.1	1	A	1	A	0.1	0	A	1	B	0.0	-1	A	1	C	0.9	1	10 Marks	L3	C01
State (s)	Action (a)	Next State (s')	Transition Probability $P(s' s, a)$ or $P_{ss'}^a$	Reward Function $R(s, a, s')$ or $R_{ss'}^a$																																				
A	0	A	0.1	0																																				
A	0	B	0.8	-1																																				
A	0	C	0.1	1																																				
A	1	A	0.1	0																																				
A	1	B	0.0	-1																																				
A	1	C	0.9	1																																				

13.	a.	Describe the step-by-step algorithm of Monte Carlo prediction (Q function), and illustrate with a suitable example.	10 Marks	L2	C02
	b.	Explain First-Visit Monte Carlo prediction. Compare it with Every-Visit MC prediction using an illustrative example.	10 Marks	L2	C02
Or					
14.	a.	Discuss the On-policy Monte Carlo Control with Exploring Starts (MC-ES) algorithm used in reinforcement learning.	10 Marks	L2	C02
	b.	Explain the On-policy Monte Carlo Control with epsilon-greedy policy algorithm used in reinforcement learning.	10 Marks	L2	C02

15.	a.	Explain the steps involved in the off-policy TD control algorithm and compute Q-learning algorithm for the Frozen Lake environment under the given policy. Using the provided transition data, determine the updated value estimates for the state (3,2). Assume the learning rate $\alpha = 0.1$ and the discount factor $\gamma = 1$.	10 Marks	L3	C03
-----	----	--	----------	----	-----

		<table><tr><td></td><td>1</td><td>2</td><td>3</td><td>4</td></tr><tr><td>1</td><td>S</td><td>F</td><td>F</td><td>F</td></tr><tr><td>2</td><td>F</td><td>H</td><td>F</td><td>H</td></tr><tr><td>3</td><td>F</td><td>F OK</td><td>F</td><td>H</td></tr><tr><td>4</td><td>H</td><td>F</td><td>F</td><td>G</td></tr></table>		1	2	3	4	1	S	F	F	F	2	F	H	F	H	3	F	F OK	F	H	4	H	F	F	G											
	1	2	3	4																																		
1	S	F	F	F																																		
2	F	H	F	H																																		
3	F	F OK	F	H																																		
4	H	F	F	G																																		
		<table><tr><th>State</th><th>Action</th><th>Value</th></tr><tr><td>(1,1)</td><td>Up</td><td>0.5</td></tr><tr><td>(3,2)</td><td>Up</td><td>0.1</td></tr><tr><td>(3,2)</td><td>Down</td><td>0.8</td></tr><tr><td>(3,2)</td><td>Left</td><td>0.5</td></tr><tr><td>(3,2)</td><td>Right</td><td>0.6</td></tr><tr><td>(4,2)</td><td>Up</td><td>0.2</td></tr><tr><td>(4,2)</td><td>Down</td><td>0.5</td></tr><tr><td>(4,2)</td><td>Left</td><td>0.1</td></tr><tr><td>(4,2)</td><td>Right</td><td>0.8</td></tr><tr><td>(4,4)</td><td>Right</td><td>0.5</td></tr></table>	State	Action	Value	(1,1)	Up	0.5	(3,2)	Up	0.1	(3,2)	Down	0.8	(3,2)	Left	0.5	(3,2)	Right	0.6	(4,2)	Up	0.2	(4,2)	Down	0.5	(4,2)	Left	0.1	(4,2)	Right	0.8	(4,4)	Right	0.5			
State	Action	Value																																				
(1,1)	Up	0.5																																				
(3,2)	Up	0.1																																				
(3,2)	Down	0.8																																				
(3,2)	Left	0.5																																				
(3,2)	Right	0.6																																				
(4,2)	Up	0.2																																				
(4,2)	Down	0.5																																				
(4,2)	Left	0.1																																				
(4,2)	Right	0.8																																				
(4,4)	Right	0.5																																				

	b. Discuss the steps involved in the on-policy TD control algorithm and compute SARSA algorithm for the Frozen Lake environment under the given policy. Using the provided transition data, determine the updated value estimates for the state (4,2). Assume the learning rate $\alpha = 0.1$ and the discount factor $\gamma = 1$. <table><tr><td></td><td>1</td><td>2</td><td>3</td><td>4</td></tr><tr><td>1</td><td>S</td><td>F</td><td>F</td><td>F</td></tr><tr><td>2</td><td>F</td><td>H</td><td>F</td><td>H</td></tr><tr><td>3</td><td>F</td><td>F</td><td>F</td><td>H</td></tr><tr><td>4</td><td>H</td><td>F</td><td>F</td><td>G</td></tr></table><table><tr><th>State</th><th>Action</th><th>Value</th></tr><tr><td>(1,1)</td><td>Up</td><td>0.5</td></tr><tr><td>(4,2)</td><td>Up</td><td>0.3</td></tr><tr><td>(4,2)</td><td>Down</td><td>0.5</td></tr><tr><td>(4,2)</td><td>Left</td><td>0.1</td></tr><tr><td>(4,2)</td><td>Right</td><td>0.8</td></tr><tr><td>(4,3)</td><td>Up</td><td>0.1</td></tr><tr><td>(4,3)</td><td>Down</td><td>0.3</td></tr><tr><td>(4,3)</td><td>Left</td><td>1.0</td></tr><tr><td>(4,3)</td><td>Right</td><td>0.9</td></tr><tr><td>(4,4)</td><td>Right</td><td>0.5</td></tr></table>		1	2	3	4	1	S	F	F	F	2	F	H	F	H	3	F	F	F	H	4	H	F	F	G	State	Action	Value	(1,1)	Up	0.5	(4,2)	Up	0.3	(4,2)	Down	0.5	(4,2)	Left	0.1	(4,2)	Right	0.8	(4,3)	Up	0.1	(4,3)	Down	0.3	(4,3)	Left	1.0	(4,3)	Right	0.9	(4,4)	Right	0.5	10 Marks	L3	C03
	1	2	3	4																																																										
1	S	F	F	F																																																										
2	F	H	F	H																																																										
3	F	F	F	H																																																										
4	H	F	F	G																																																										
State	Action	Value																																																												
(1,1)	Up	0.5																																																												
(4,2)	Up	0.3																																																												
(4,2)	Down	0.5																																																												
(4,2)	Left	0.1																																																												
(4,2)	Right	0.8																																																												
(4,3)	Up	0.1																																																												
(4,3)	Down	0.3																																																												
(4,3)	Left	1.0																																																												
(4,3)	Right	0.9																																																												
(4,4)	Right	0.5																																																												

Or

16.	a. Describe the steps involved in the Temporal Difference (TD) prediction algorithm and compute the Temporal Difference (TD) prediction for the Frozen Lake environment under the specified policy. Using the provided transition data, determine the updated value estimates for the states (1,1), (1,2), and (1,3). Assume the learning rate $\alpha = 0.1$ and the discount factor $\gamma = 1$. <table><tr><th>State</th><th>Action</th></tr><tr><td>(1,1)</td><td>Right</td></tr><tr><td>(1,2)</td><td>Right</td></tr><tr><td>(1,3)</td><td>Left</td></tr><tr><td>⋮</td><td>⋮</td></tr><tr><td>(4,4)</td><td>Down</td></tr></table><table><tr><td></td><td>1</td><td>2</td><td>3</td><td>4</td></tr><tr><td>1</td><td>S</td><td>F</td><td>F</td><td>F</td></tr><tr><td>2</td><td>F</td><td>H</td><td>F</td><td>H</td></tr><tr><td>3</td><td>F</td><td>F</td><td>F</td><td>F</td></tr><tr><td>4</td><td>H</td><td>F</td><td>F</td><td>G</td></tr></table><table><tr><th>State</th><th>Value</th></tr><tr><td>(1,1)</td><td>0.9</td></tr><tr><td>(1,2)</td><td>0.6</td></tr><tr><td>(1,3)</td><td>0.8</td></tr><tr><td>⋮</td><td>⋮</td></tr><tr><td>(4,4)</td><td>0.7</td></tr></table>	State	Action	(1,1)	Right	(1,2)	Right	(1,3)	Left	⋮	⋮	(4,4)	Down		1	2	3	4	1	S	F	F	F	2	F	H	F	H	3	F	F	F	F	4	H	F	F	G	State	Value	(1,1)	0.9	(1,2)	0.6	(1,3)	0.8	⋮	⋮	(4,4)	0.7	10 Marks	L3	C03
State	Action																																																				
(1,1)	Right																																																				
(1,2)	Right																																																				
(1,3)	Left																																																				
⋮	⋮																																																				
(4,4)	Down																																																				
	1	2	3	4																																																	
1	S	F	F	F																																																	
2	F	H	F	H																																																	
3	F	F	F	F																																																	
4	H	F	F	G																																																	
State	Value																																																				
(1,1)	0.9																																																				
(1,2)	0.6																																																				
(1,3)	0.8																																																				
⋮	⋮																																																				
(4,4)	0.7																																																				
	b. Describe and compare DP, MC, and TD learning techniques using an illustrative example.	10 Marks	L2	C03																																																	

17.	a. Illustrate with example, how to find the best arm with the epsilon-greedy method in Multi-Armed Bandit (MAB) problem.	10 Marks	L2	C04
	b. Explain in detail Softmax exploration strategies for finding the optimal arm in Multi-Armed Bandit (MAB) problem with example.	10 Marks	L2	C04
Or				

18.	a.	Discuss the steps involved in Thompson sampling method and how Thompson sampling exploration strategy to overcome the exploration-exploitation dilemma based on a beta distribution.	10 Marks	L2	C04
	b.	Describe the Upper confidence bound for handling the exploration-exploitation dilemma in MAB problem with example.	10 Marks	L2	C04