

TODAY'S CLASS

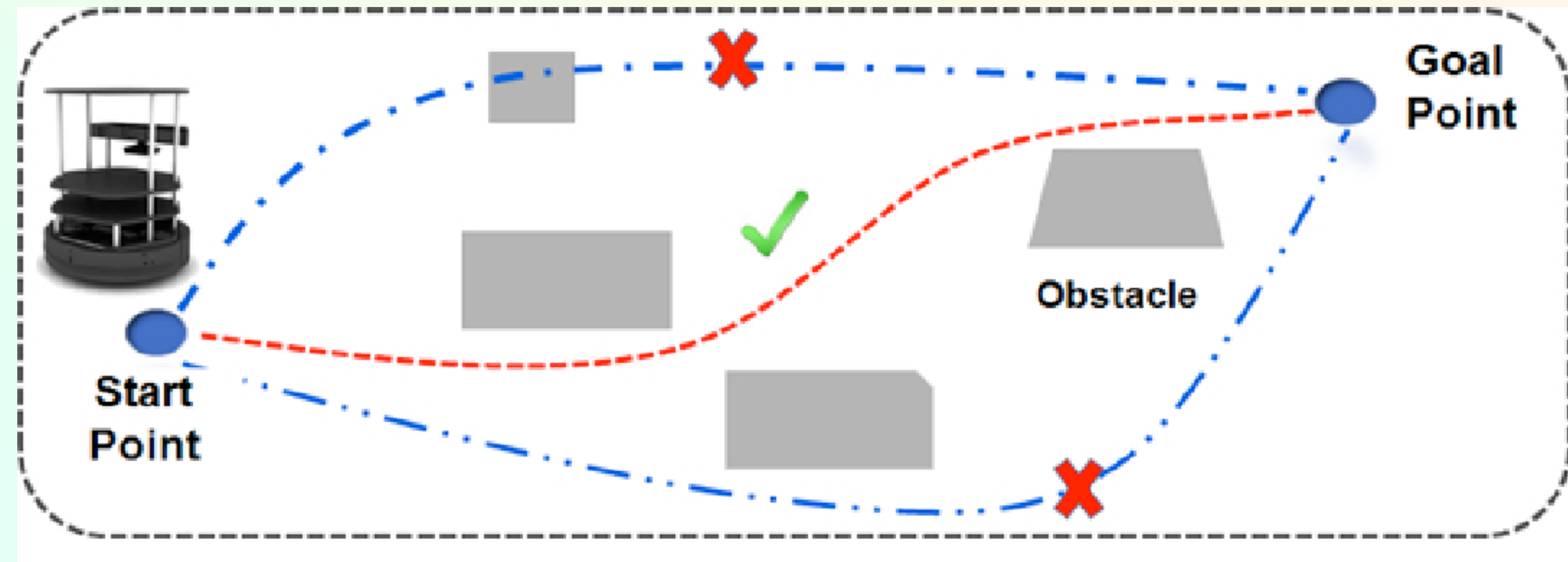
GOALS

1. Define Sequential Decision Making Problems
2. Imitation Learning (copying an expert)
3. Reinforcement Learning
4. Policy Gradient Methods

DECISION MAKING PROBLEMS

EXAMPLES

Robot navigation

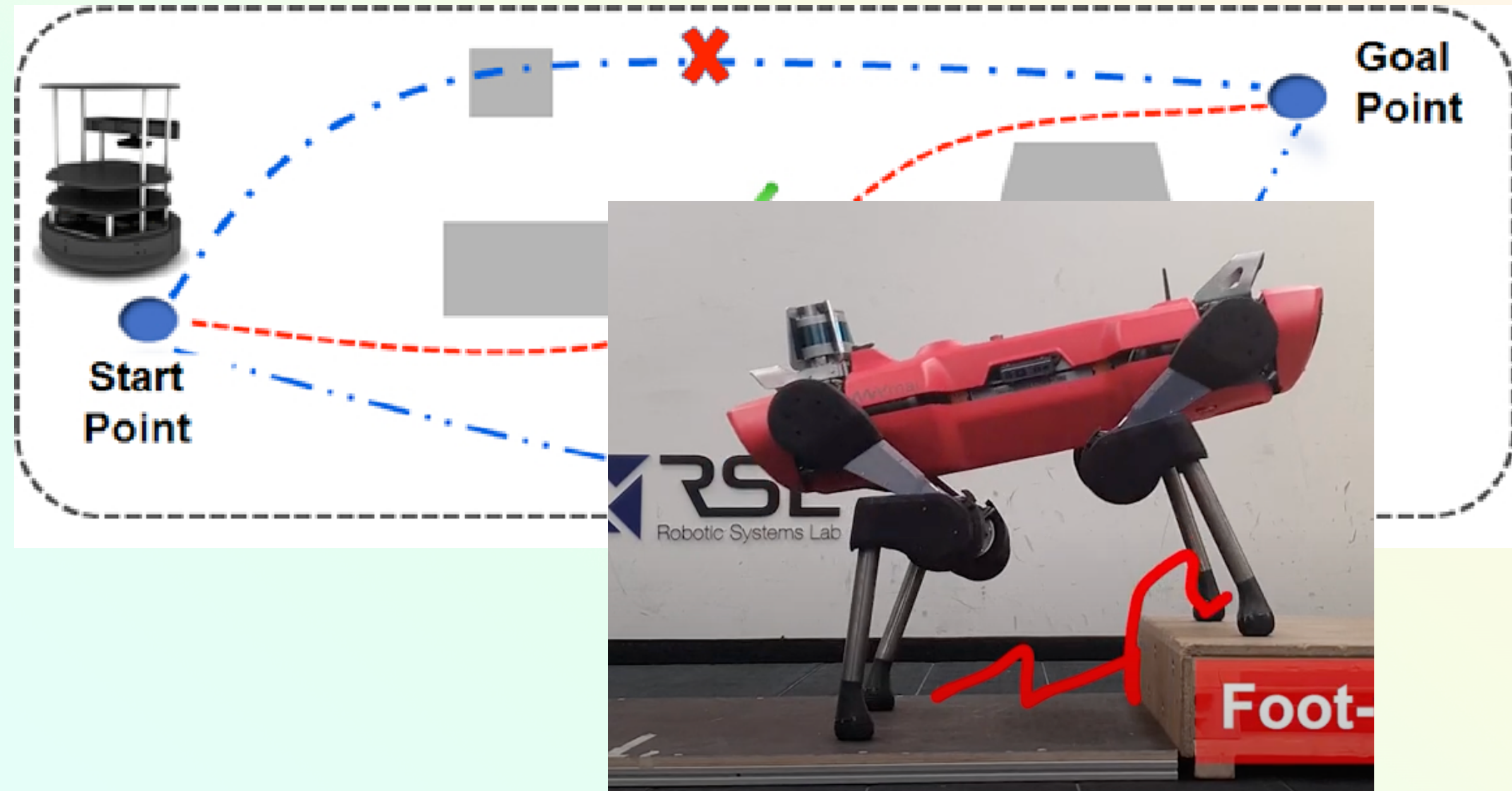


DECISION MAKING PROBLEMS

EXAMPLES

Robot navigation

Robot locomotion



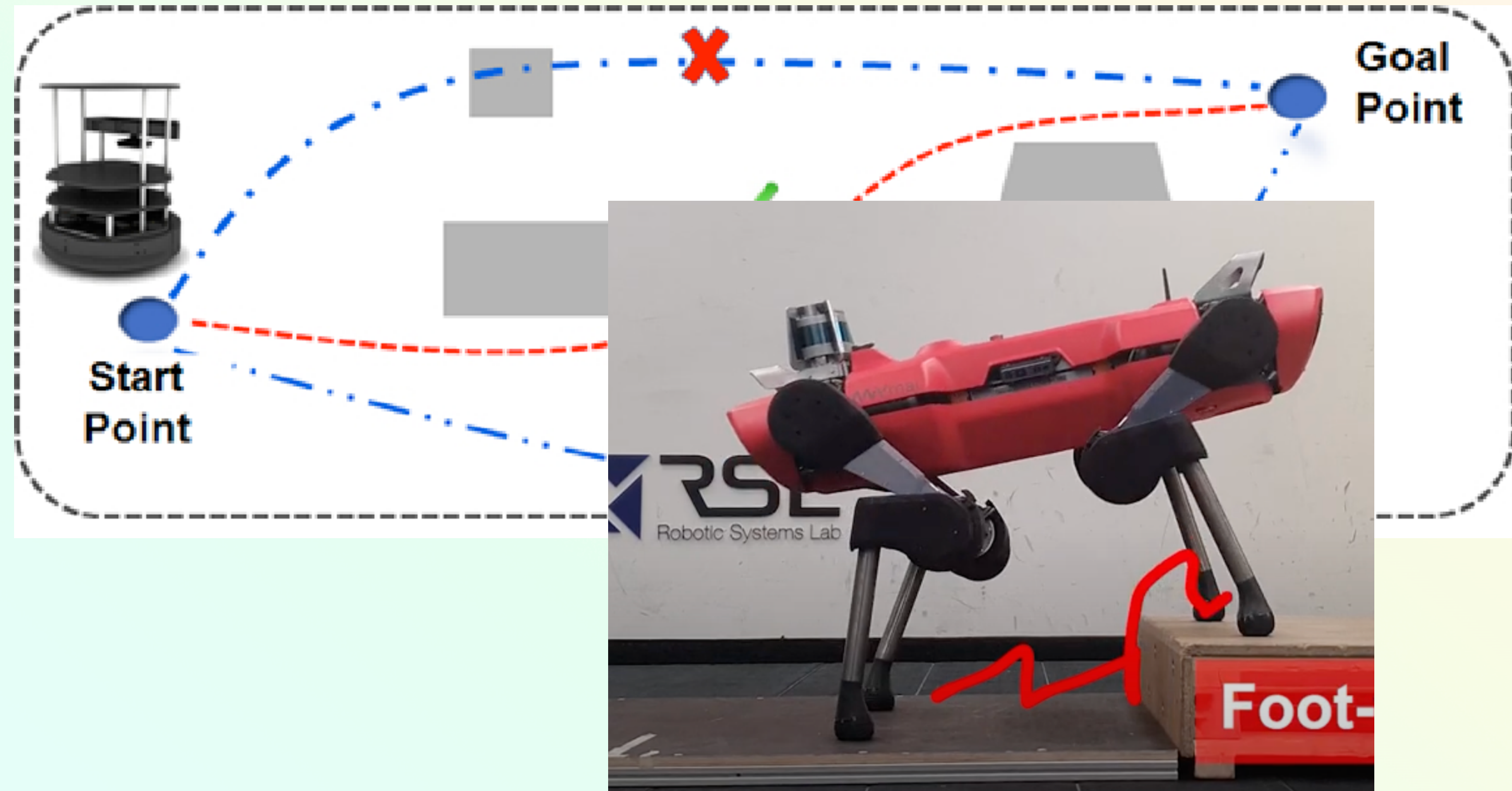
DECISION MAKING PROBLEMS

EXAMPLES

Robot navigation

Robot locomotion

Recommendation (Ads, YouTube, etc)



DECISION MAKING PROBLEMS

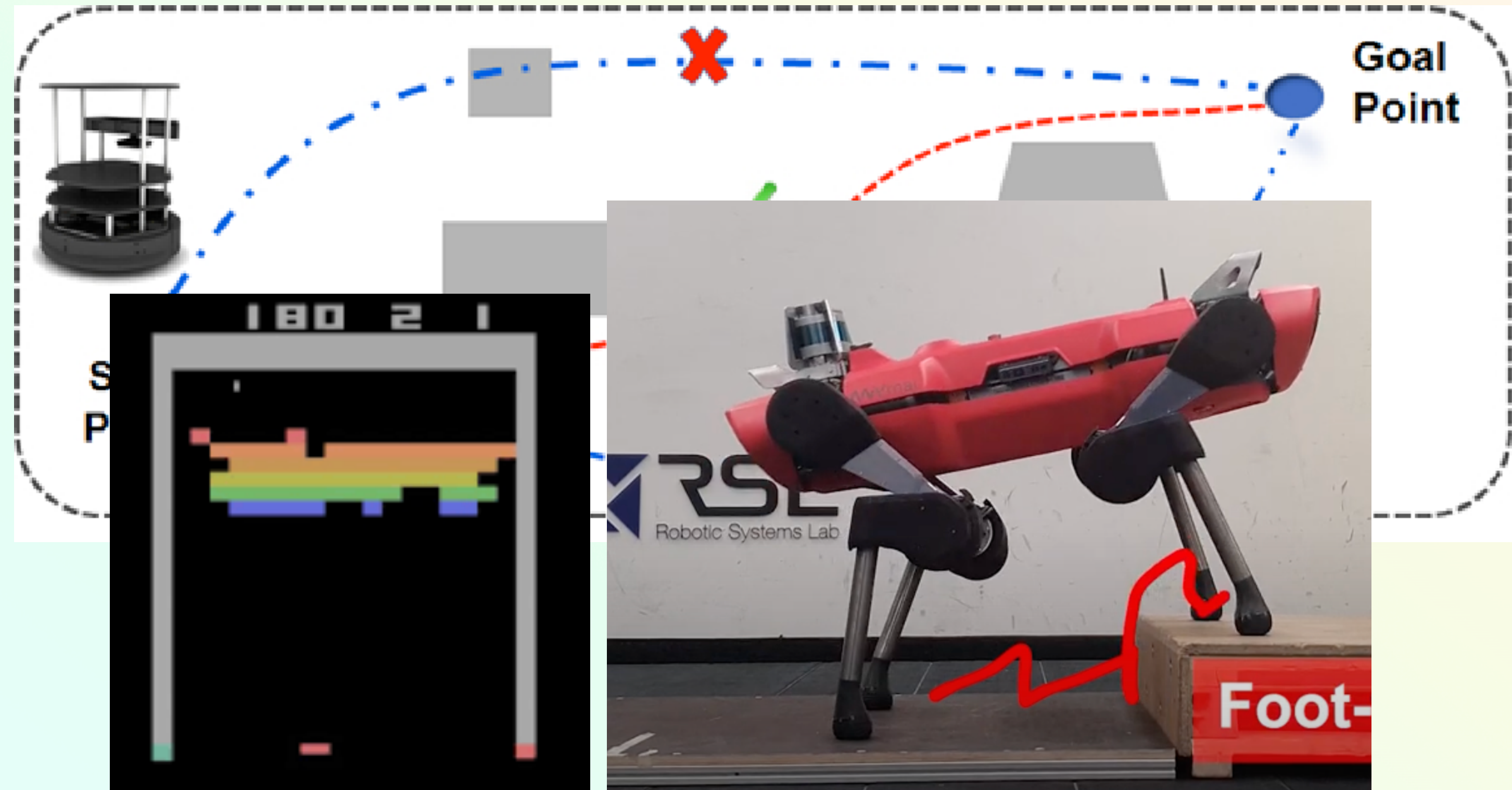
EXAMPLES

Robot navigation

Robot locomotion

Recommendation (Ads, YouTube, etc)

Game Playing



DECISION MAKING PROBLEMS

EXAMPLES

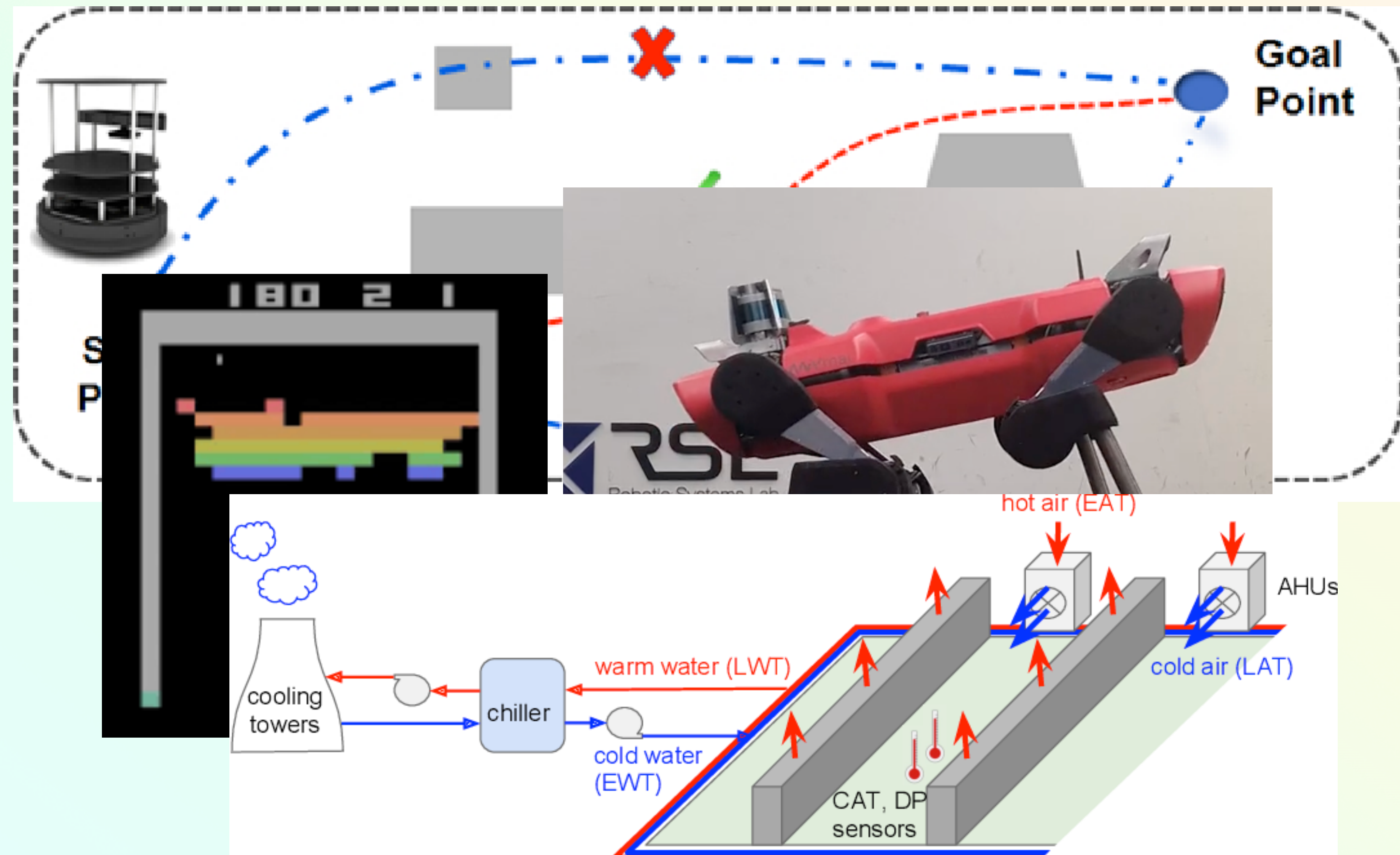
Robot navigation

Robot locomotion

Recommendation (Ads, YouTube, etc)

Game Playing

Data Center Cooling



DECISION MAKING PROBLEMS

EXAMPLES

Robot navigation

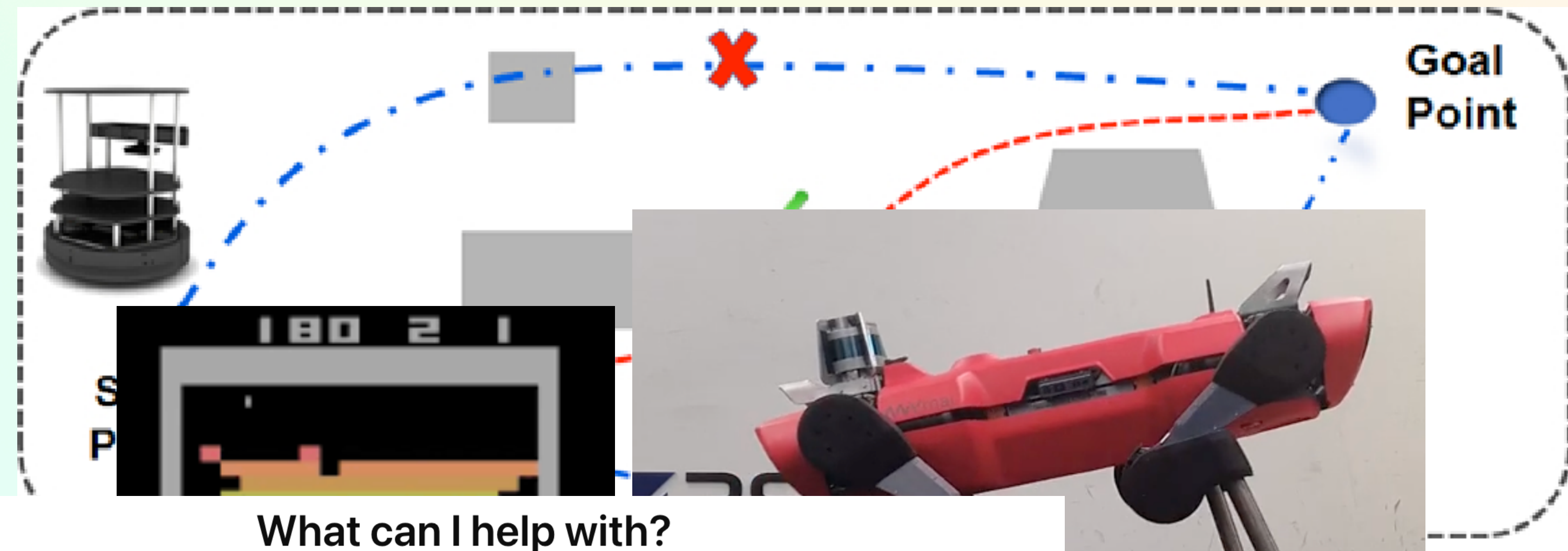
Robot locomotion

Recommendation (Ads, YouTube, etc)

Game Playing

Data Center Cooling

Chatbots/LLMs

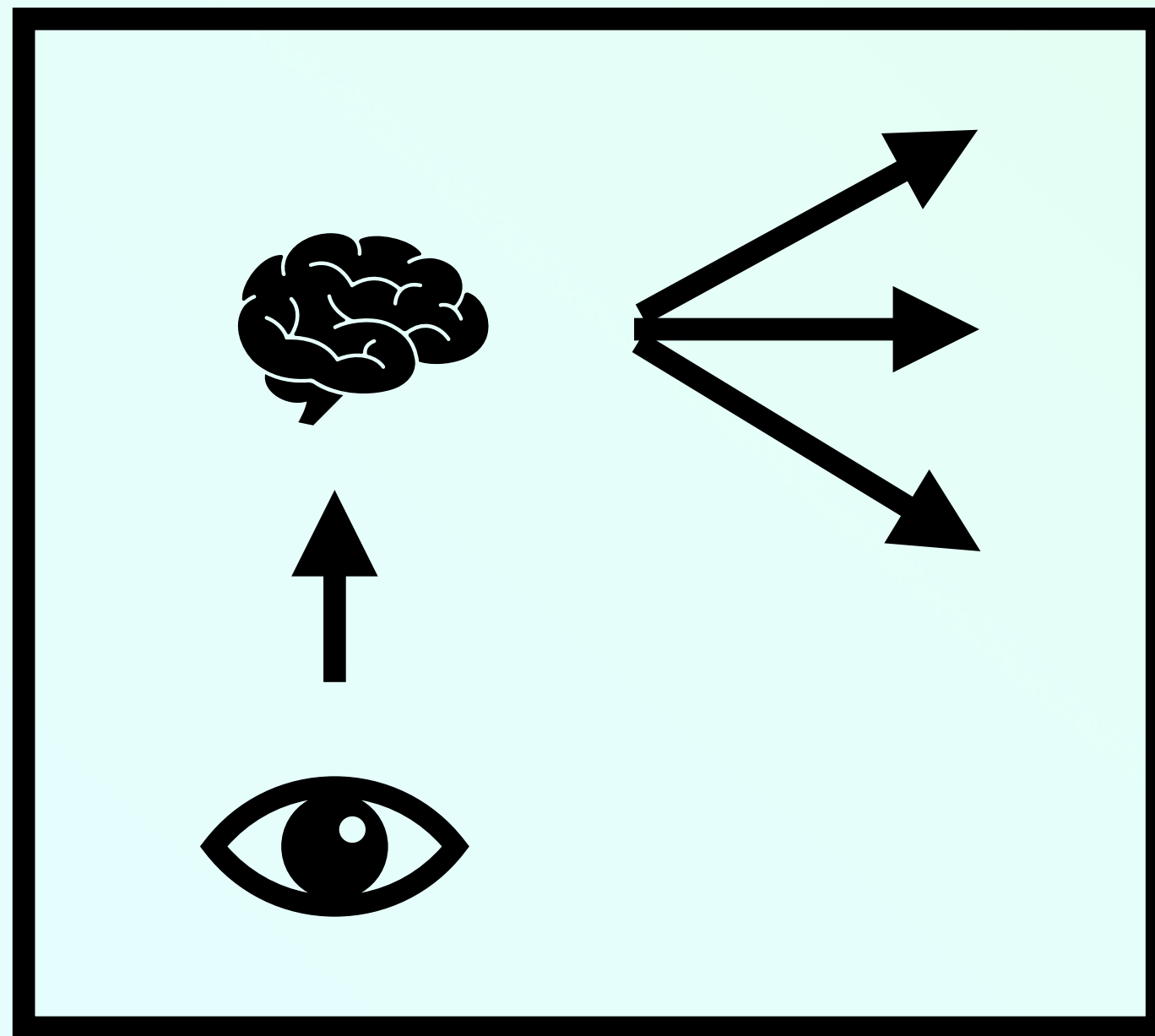


What can I help with?

Message ChatGPT

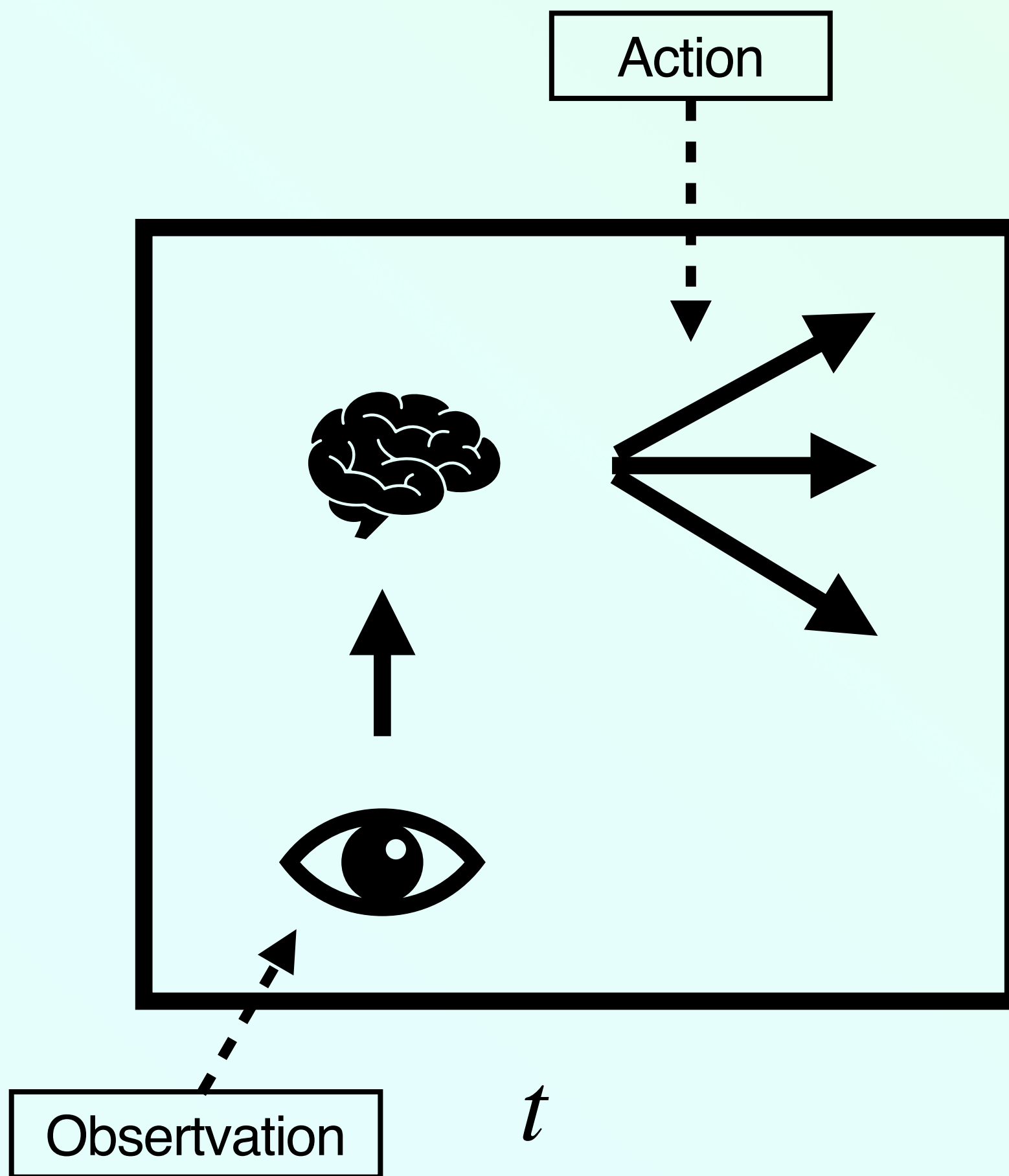


SEQUENTIAL DECISION MAKING

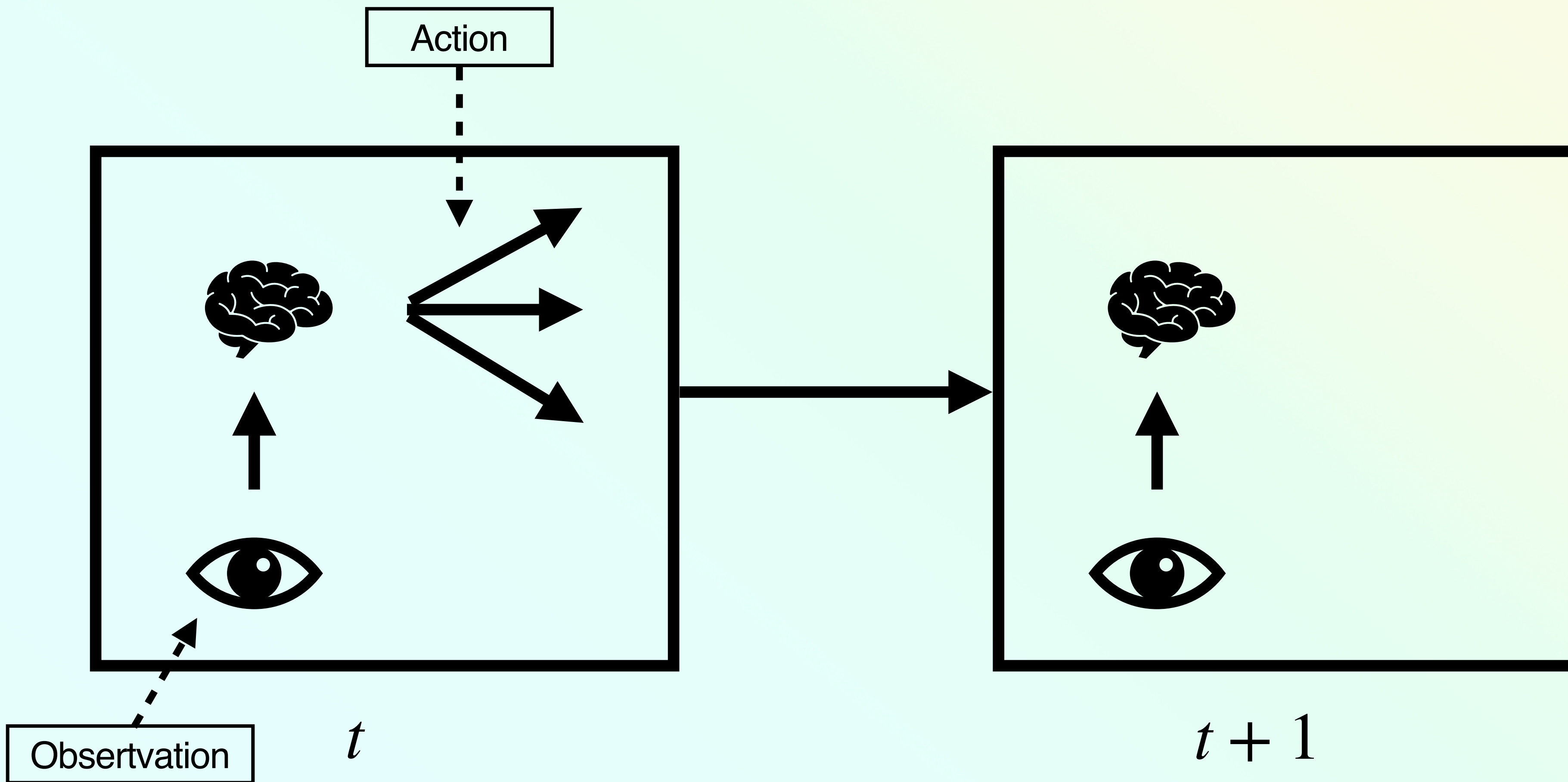


t

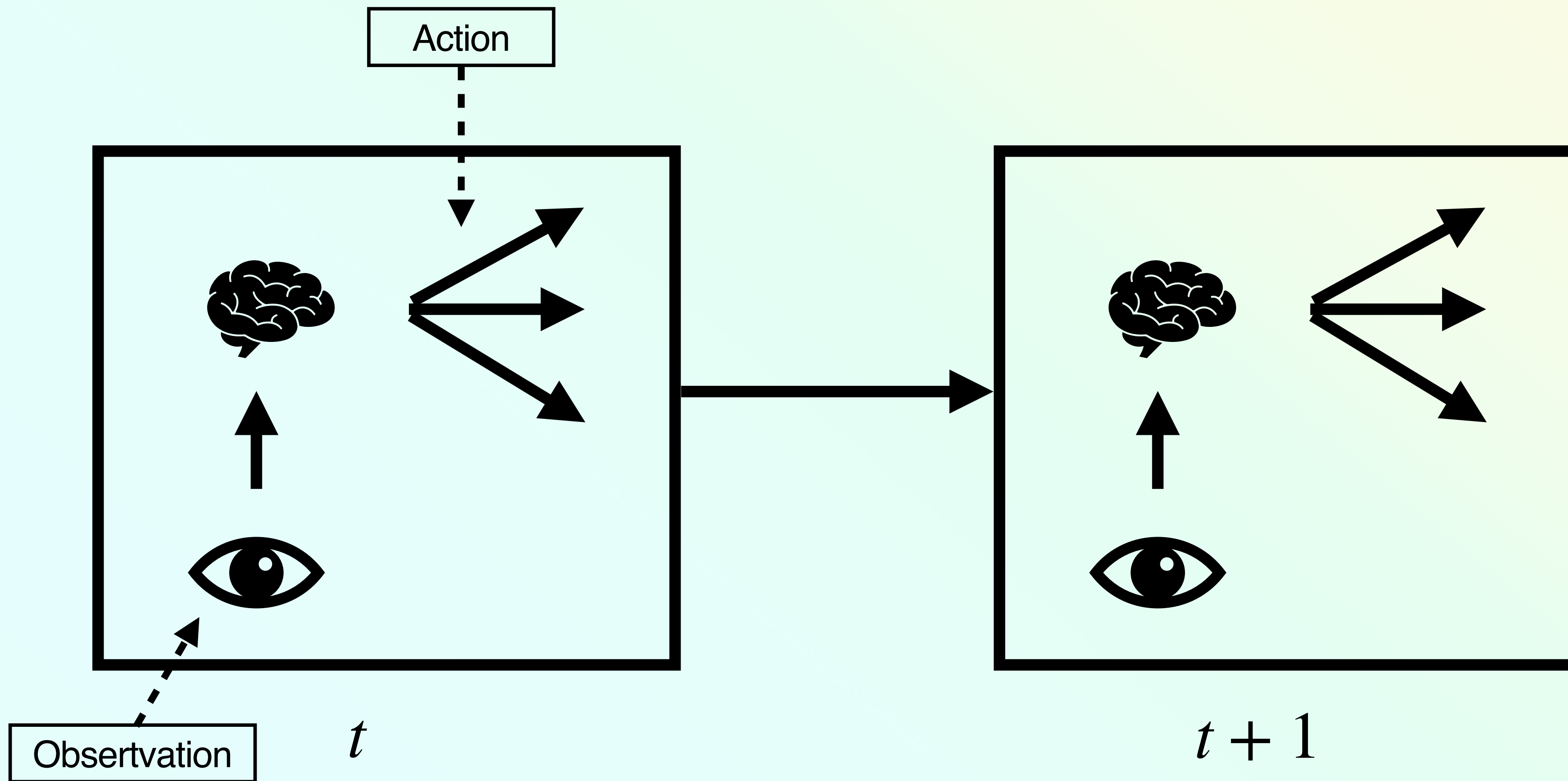
SEQUENTIAL DECISION MAKING



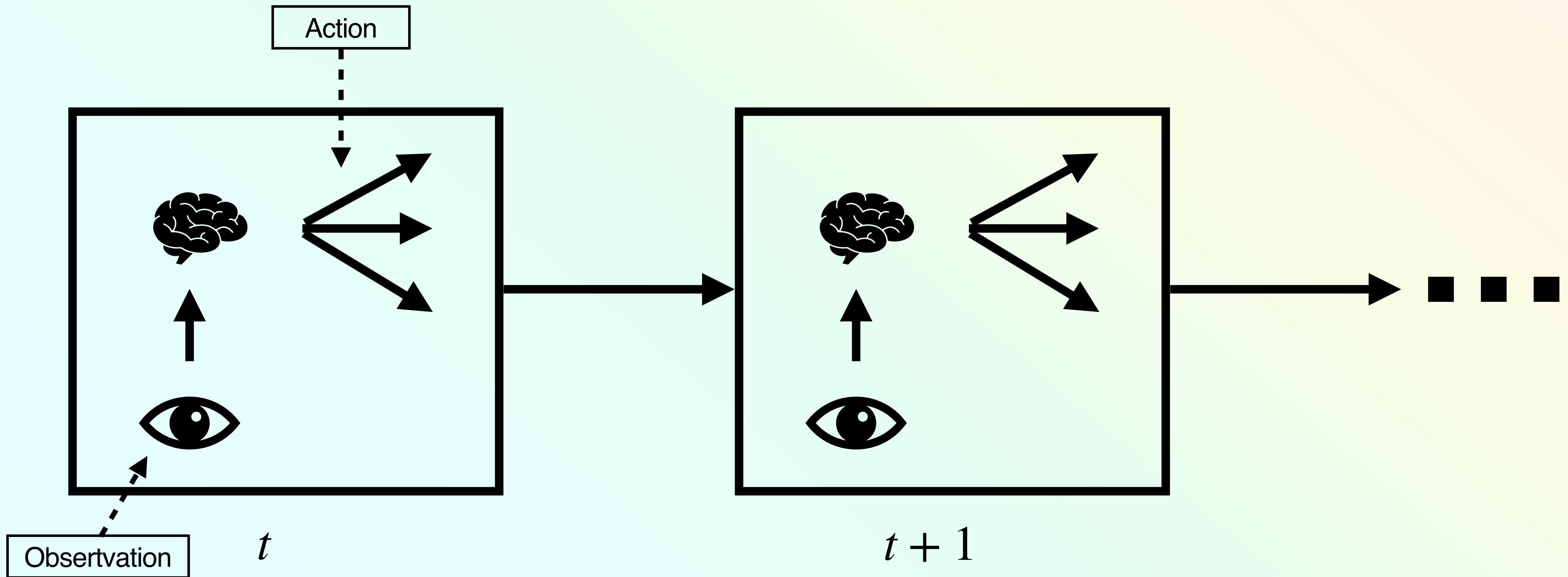
SEQUENTIAL DECISION MAKING



SEQUENTIAL DECISION MAKING



SEQUENTIAL DECISION MAKING



SEQUENTIAL DECISION MAKING

PROBLEM FORMULATION (IMITATION LEARNING)

Can we frame this problem as a function approximation problem?

X_t — observation at time t

A_t — action (agent's decision) at time t

$f_*(X_t)$ — expert's probability distribution over actions given X_t

$f(X_t)$ — agent's probability distribution over actions given X_t

$$\arg \min_f \mathbf{E} \left[\sum_{t=1}^{\infty} d(f(X_t), f_*(X_t)) \right]$$

SEQUENTIAL DECISION MAKING

PROBLEM FORMULATION (IMITATION LEARNING)

Collect data from an expert:

$$D = \{x_1, x_2, \dots, x_m\}, \{a_1, a_2, \dots, a_m\}$$

Classification (or regression) to fit f to D

$$l_D(\theta) = - \sum_{i=1}^m \ln \Pr(A_t = a_i | X_t = x_i) = - \sum_{i=1}^m \ln f(x_i, \theta)_{a_i}$$

IMITATION LEARNING

PROPERTIES

- Approximation has errors that does not match the expert
 - Errors lead to new observations, not in the data set

IMITATION LEARNING

PROPERTIES

- Approximation has errors that does not match the expert
 - Errors lead to new observations, not in the data set
- Assume perfect approximation on D
 - The environment may have noise (cannot exactly replicate the same sequences in D)
 - Tiny changes in observation may lead to new states

IMITATION LEARNING

PROPERTIES

- Imitation needs:
 - Lots of expert data
 - Collect expert decisions from observation that an expert won't see, but the agent might

IMITATION LEARNING

PROPERTIES

- Imitation needs:
 - Lots of expert data
 - Collect expert decisions from observation that an expert won't see, but the agent might

An agent can only be as good as an expert!

It is not always possible to have a good enough expert!

REINFORCEMENT LEARNING

DESIRED PROPERTIES FOR DECISION-MAKING AGENT

The agent should learn from its interactions with the environment

- collect its own data (not always necessary)

The agent needs to determine what is the best action

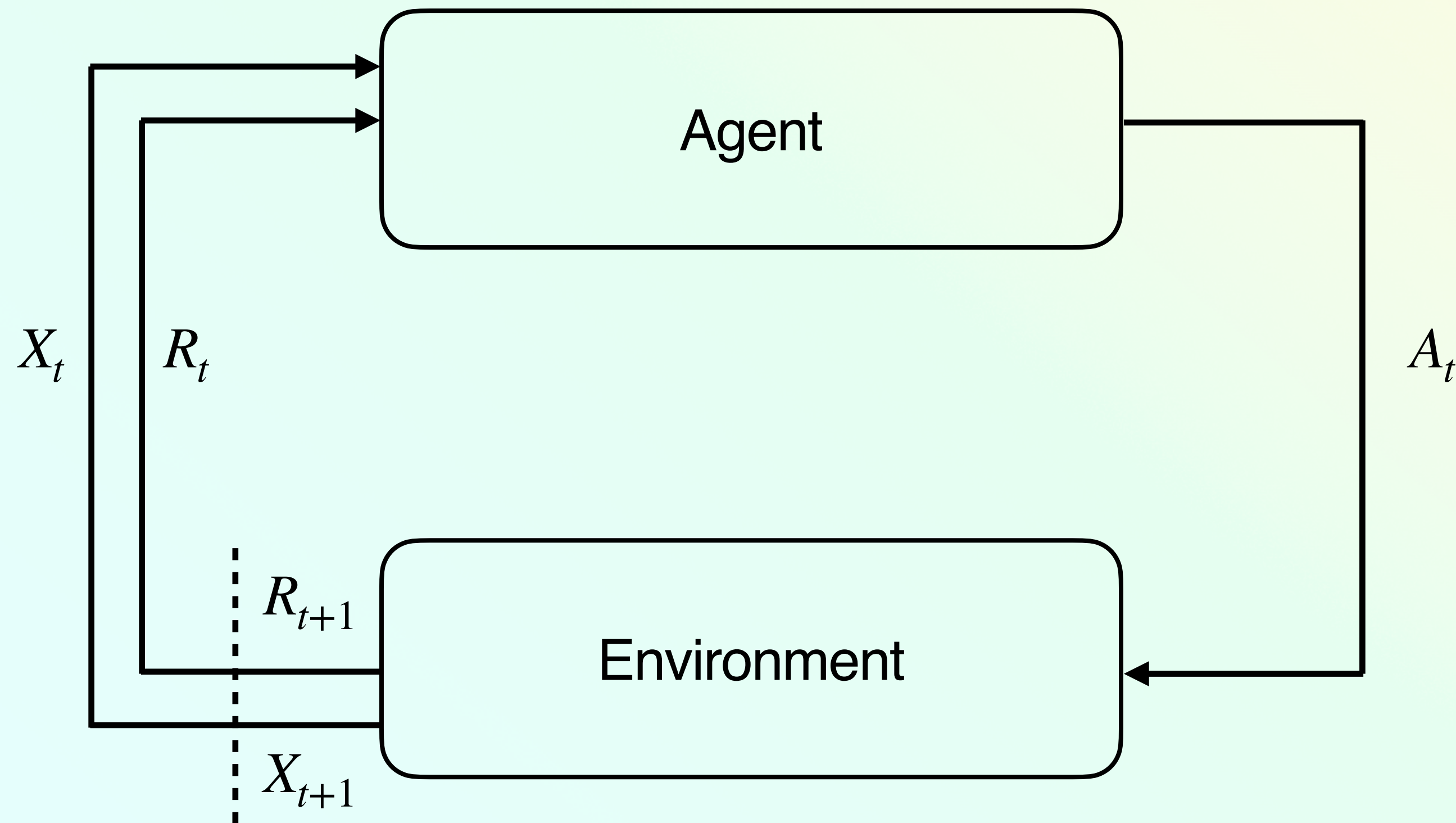
- Agent has to search for the best action (we cannot tell it what to do)

The agent needs to know how good its decision was.

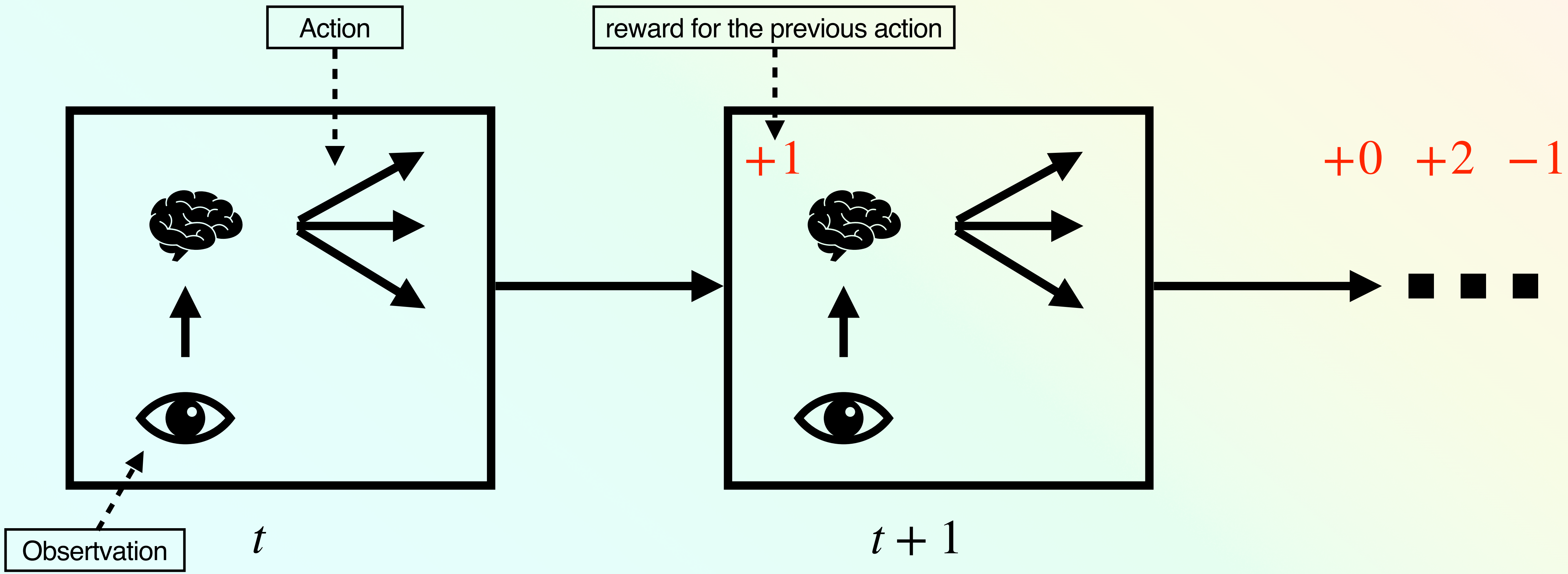
- Provide a score (reward) indicating the quality of decisions

REINFORCEMENT LEARNING

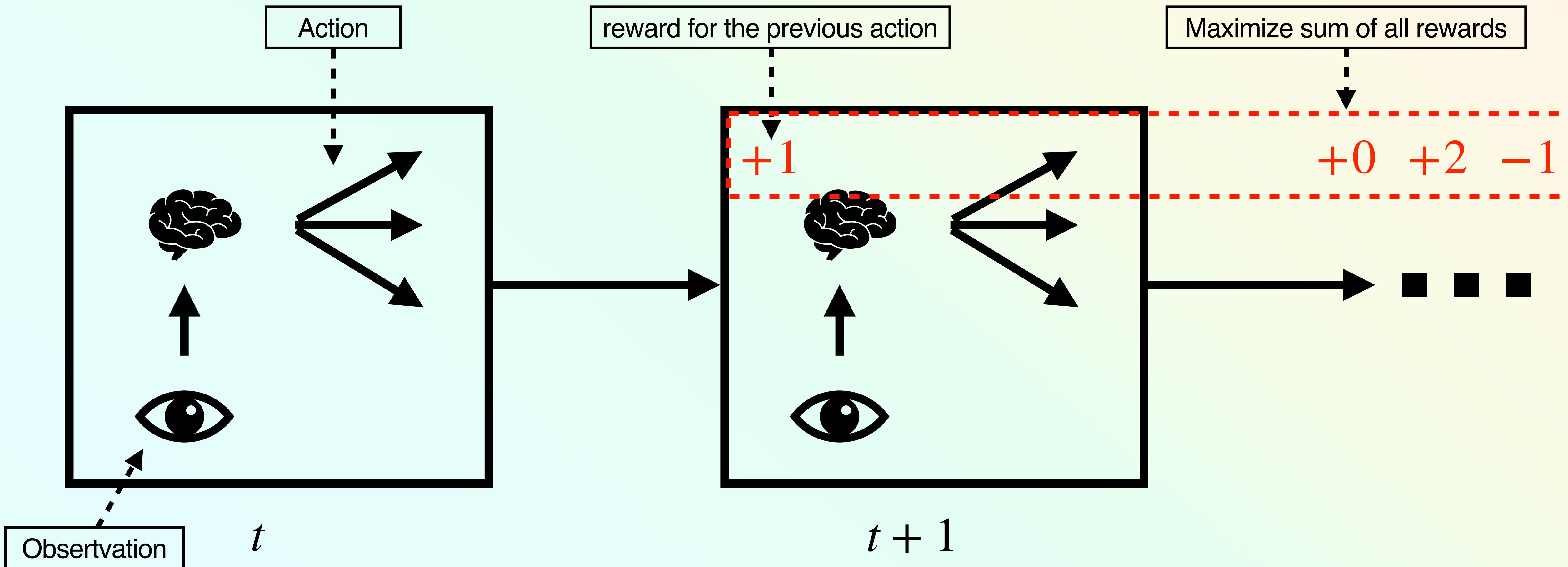
AGENT-ENVIRONMENT INTERACTION



SEQUENTIAL DECISION MAKING



SEQUENTIAL DECISION MAKING



RL PROBLEMS

EXAMPLES



Maximize score

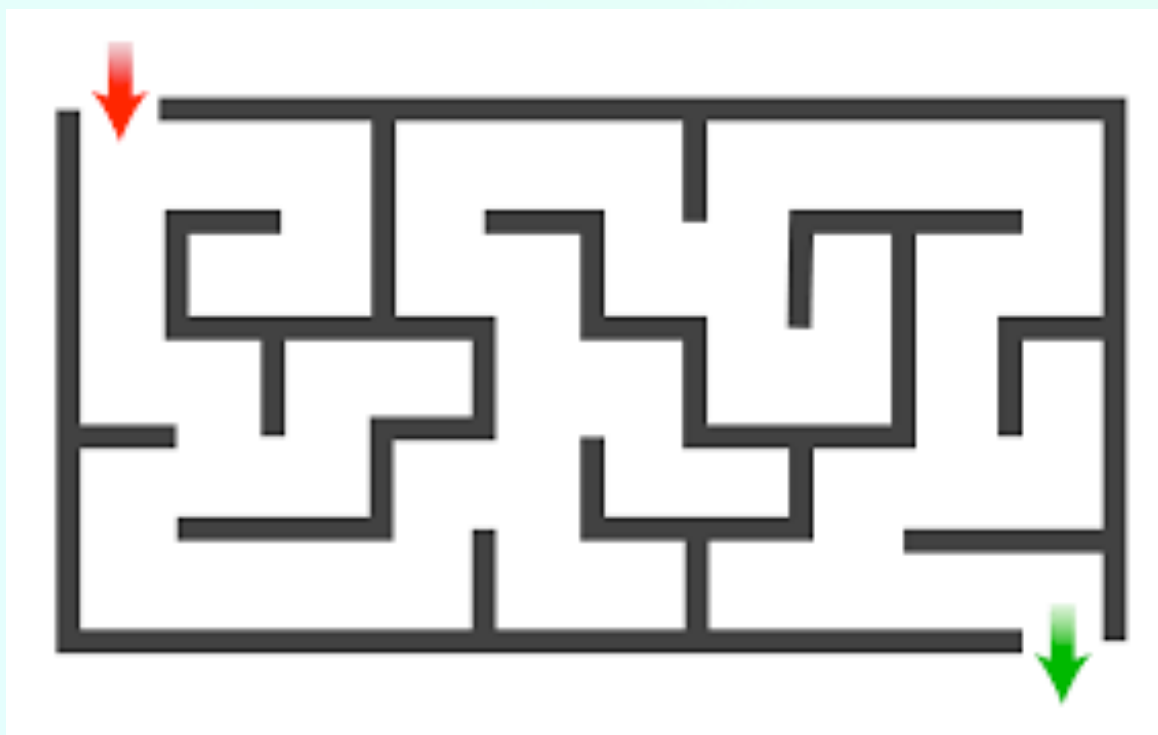
X_t = image of the game

A_t = controller movement

R_{t+1} = points for breaking bricks

RL PROBLEMS

EXAMPLES



Get out the maze as fast
as possible

X_t = Position in the maze

A_t = direction to move

$R_{t+1} = -1$

Agent penalized for being in the maze every time step

QUIZ

REINFORCEMENT LEARNING

OPTIMIZATION

$$\rho(\theta) \doteq \mathbf{E} \left[\sum_{t=1}^T R_{t+1} \right]$$

Gradient ascent:

$$\theta \leftarrow \theta + \eta \nabla \rho(\theta)$$

POLICY GRADIENT

EXPRESSION

$$\nabla \rho(\theta) = \frac{\partial}{\partial \theta} \mathbf{E} \left[\sum_{t=1}^T R_{t+1} \right]$$

POLICY GRADIENT

EXPRESSION

$$\rho(\theta) = \mathbf{E} \left[\sum_{t=1}^T R_{t+1} \right]$$

$$G_t = \sum_{k=0}^{T-t} R_{t+k+1}$$

$$H_{1:T} = \{X_1, A_1, X_2, R_2, A_2, \dots, X_T, A_T, R_{T+1}\}$$

$$\tau_{1:T} = \{x_1, a_1, x_2, r_2, \dots, x_T, a_T, r_{T+1}\}$$

POLICY GRADIENT

EXPRESSION

$$\rho(\theta) = \mathbf{E} \left[\sum_{t=1}^T R_{t+1} \right] = \mathbf{E} [G_1] = \sum_{\tau} \Pr(H_{1:T} = \tau_{1:T}) G_1$$

$$G_t = \sum_{k=0}^{T-t} R_{t+k+1}$$

$$H_{1:T} = \{X_1, A_1, X_2, R_2, A_2, \dots, X_T, A_T, R_{T+1}\}$$

$$\tau_{1:T} = \{x_1, a_1, x_2, r_2, \dots, x_T, a_T, r_{T+1}\}$$

