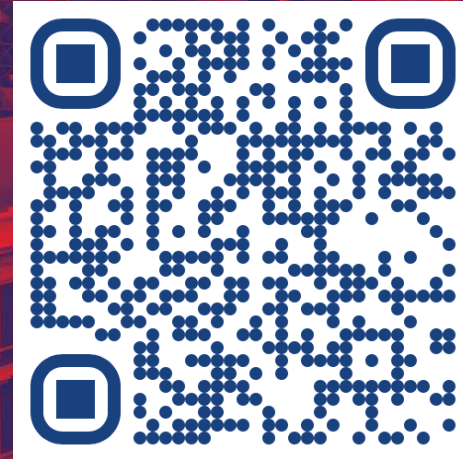


AI for Autonomous Robots: Bridging Theory and Practice

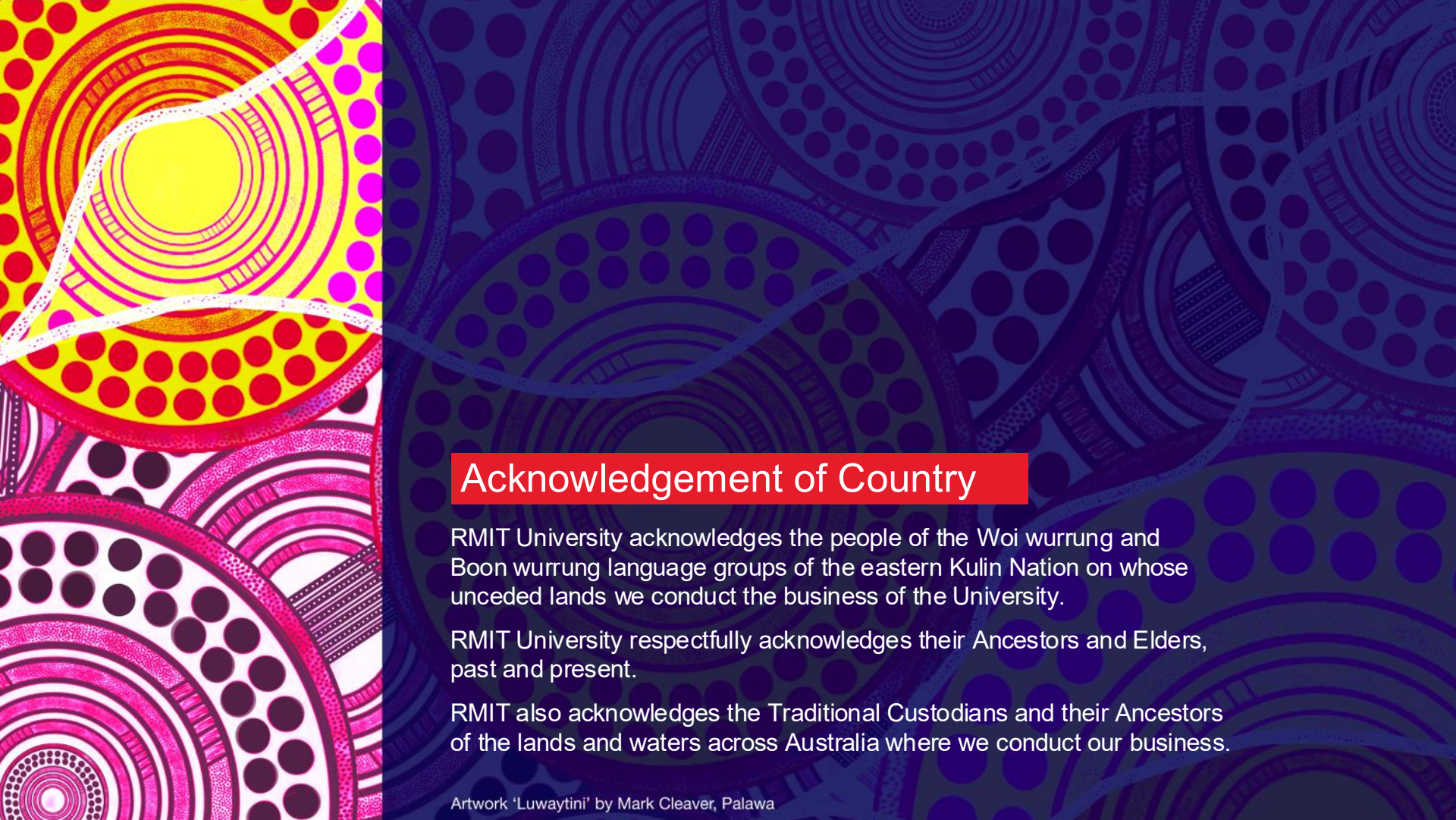
Day 1: Introduction to Robot System and Learning Robot Control.

Dr. Timothy Wiley

School of Computing Technologies
RMIT University



ESSAI July 2026



Acknowledgement of Country

RMIT University acknowledges the people of the Woi wurrung and Boon wurrung language groups of the eastern Kulin Nation on whose unceded lands we conduct the business of the University.

RMIT University respectfully acknowledges their Ancestors and Elders, past and present.

RMIT also acknowledges the Traditional Custodians and their Ancestors of the lands and waters across Australia where we conduct our business.

About me: Dr. Timothy Wiley

PhD, Online Learning of Robotic Behaviours,
UNSW Australia

Lecturer, Artificial Intelligence,
School of Computing Technologies
STEM College, RMIT University

Manager,
Artificial Intelligence Innovation Lab

Research Interests:

- Autonomous Robotics
- Reinforcement Learning
- Qualitative Reasoning
- Uncrewed Ariel Systems
- Sim2Real



— Robot Soccer



What is Autonomous Robotics?

—
ESSAI July 2026

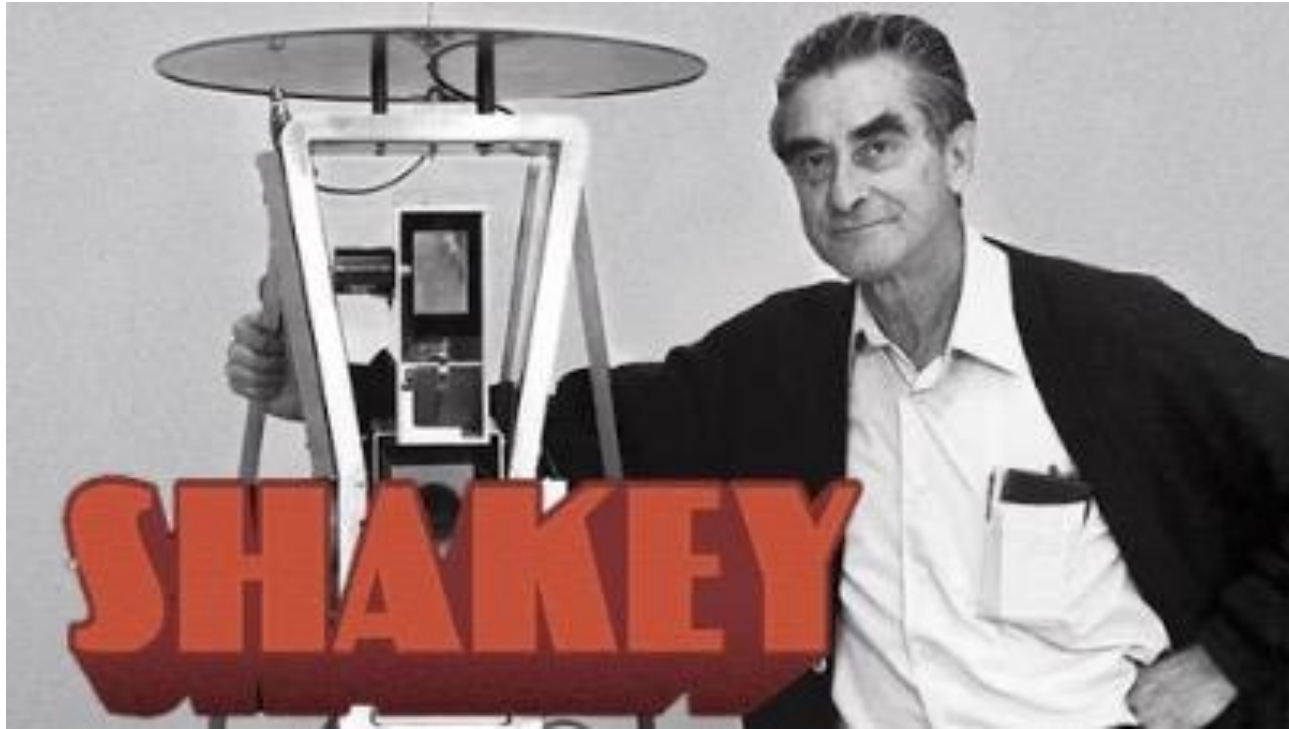


What defines an Autonomous Robot?

What defines an autonomous robot to you?



Shakey the Robot



Nilsson, N. J. (1982). *Principles of Artificial Intelligence*. Berlin Heidelberg: Springer-Verlag

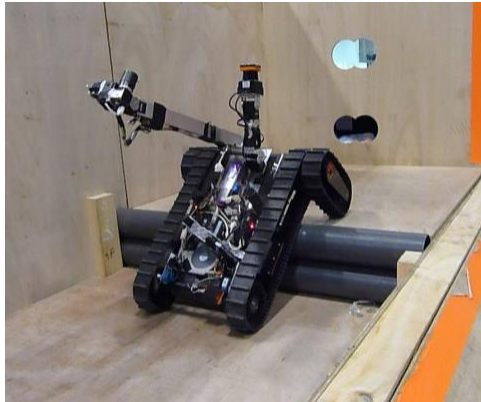


What defines an Autonomous Robot?

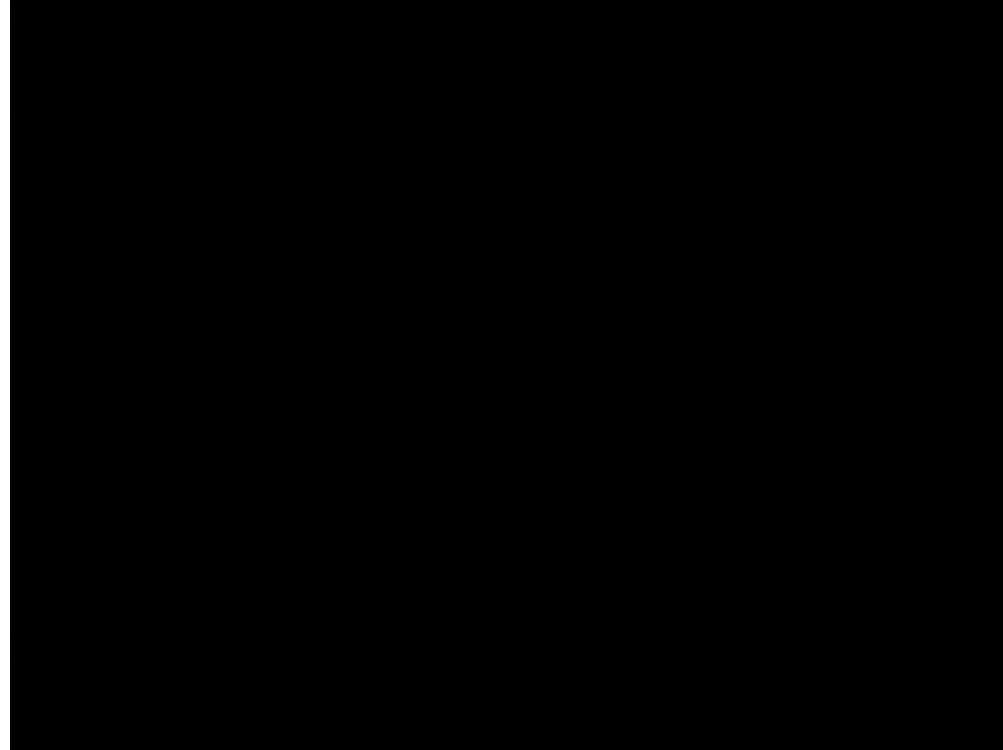
Aside from the obvious “act autonomously” that is, “think-and-act for themselves”:

- Perceive the environment
- Interact with the environment
- Internal reasoning

What separates robotics from other fields of AI is the physical input/output



RoboCup Soccer – Aibo



— RoboCup Soccer – Aibo



RoboCup Soccer – Nao



RoboCup Humanoid Soccer League



RoboCup Humanoid Soccer League



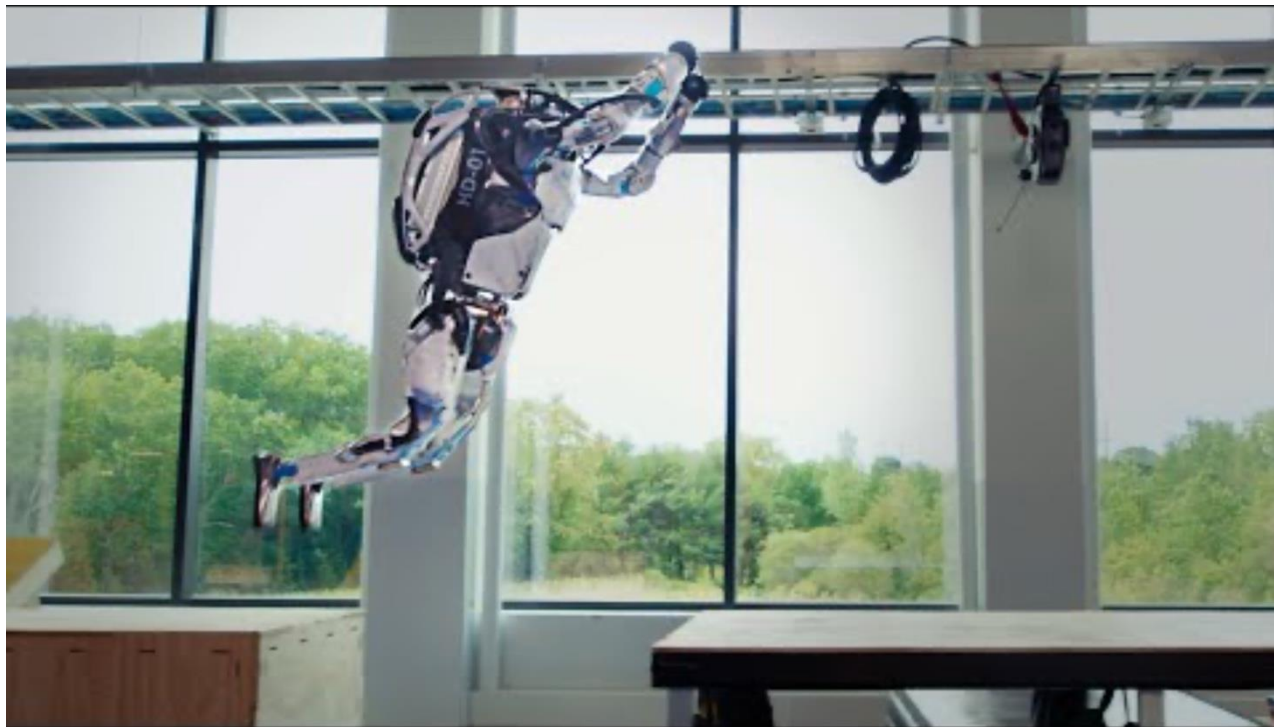
DARPA Rescue Challenge



DeepMind Learning Soccer

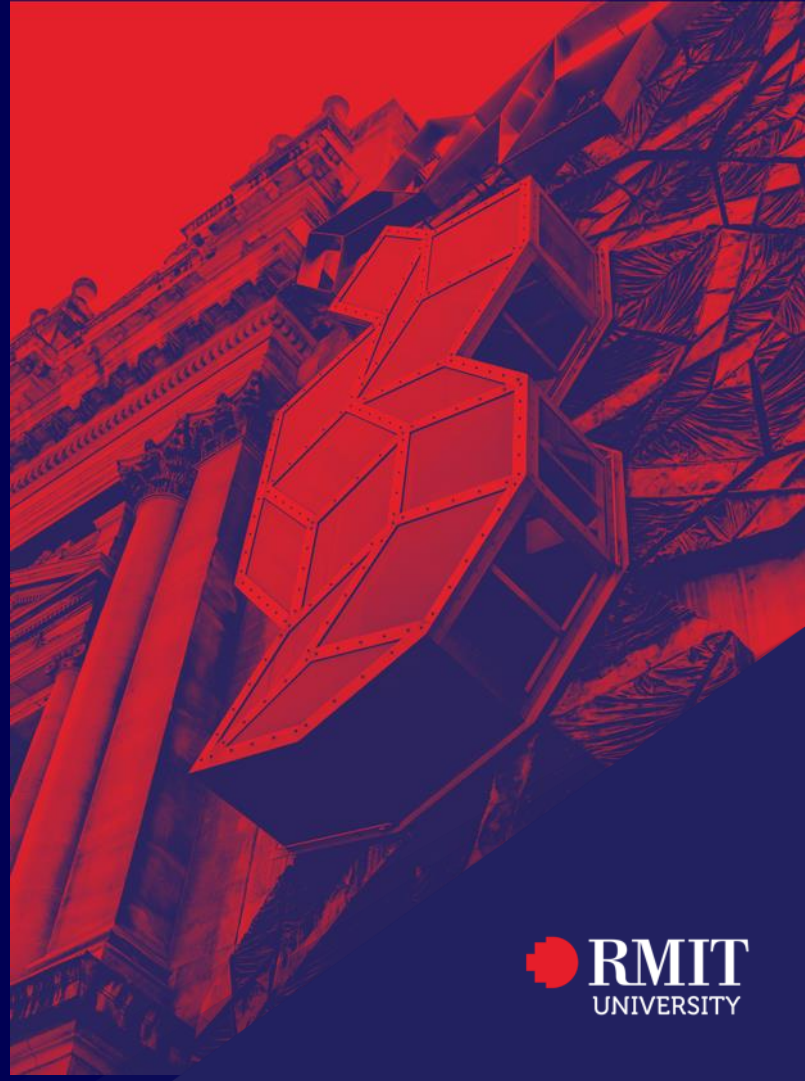


Boston Dynamics ATLAS motion



About the Course

—
ESSAI July 2026





— This Course is about...

- Overview of Autonomous Robotics
- Give a breadth of knowledge rather than a depth
- Provide operational knowledge on autonomous robotic system for later in-depth study
- Applied AI/ML
- Hopefully convince you to work in robotics 😄





— This course presumes...

- Some familiarity with common AI techniques such as:
 - Graph algorithms
 - Heuristic Search (A^*)
 - Random Sampling
 - Reinforcement Learning fundamentals
 - Machine Learning fundamentals
 - Neural Network structures





— Topics

Day 1: Introduction to Robot System and Learning Robot Control.

Day 2: Motion Planning and Navigation.

Day 3: Localisation and Mapping.

Day 4: Symbolic Task Planning and Planning with LLMs.

Not covered (unless time permits)

- Robotic Vision



Motivating Example

—
ESSAI July 2026



Problem: Maze





— Problem: Maze



Elements of Robots

—
ESSAI July 2026



ROSbot

2D LIDAR

RGB-D Camera

Wi-fi

1D
Ultrasound

Differential or
Omnidirectional
Drive

IMU
CPU



Panther

RGB-D AI
Zed2 Camera

Wi-fi x2
GPS

3D LIDAR

Cobot Arm



Differential
Drive

IMU
Pi3 (CPU)
HP-Z2 (CPU+GPU)



Nao V6

Microphone
Array

Speakers

ATOM E3845 1.91 GHz
Quad-core CPU

Lights
Interface

26 DoF Joint
System

Camera

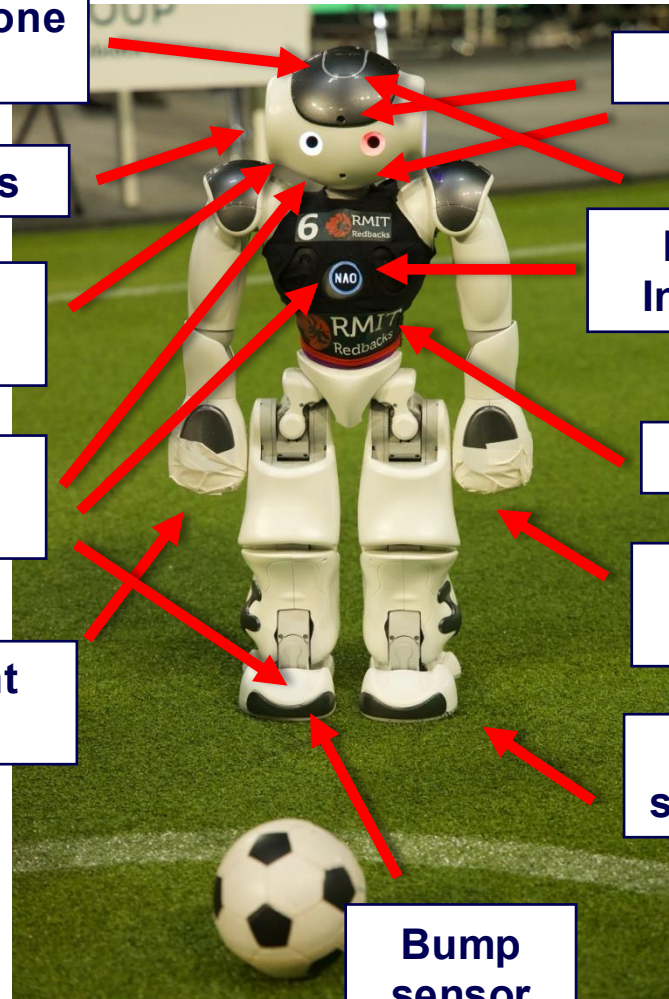
Button
Interface

IMU

Articulated
Fingers

Force
sensors

Bump
sensor



Booster K1

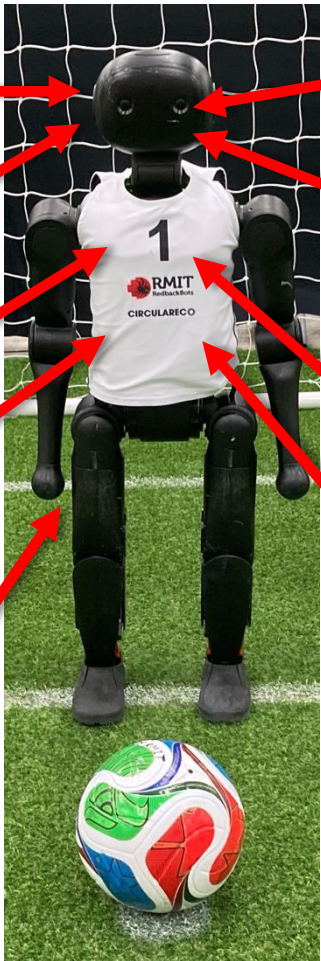
6-core Cortex-A78AE
2GHz CPU

117 TOPS GPU

6-Microphone
Array

Speakers

22 DoF Joint
System



RGB Camera

Stereo Depth
Camera

9-axis IMU

23kg
Weight



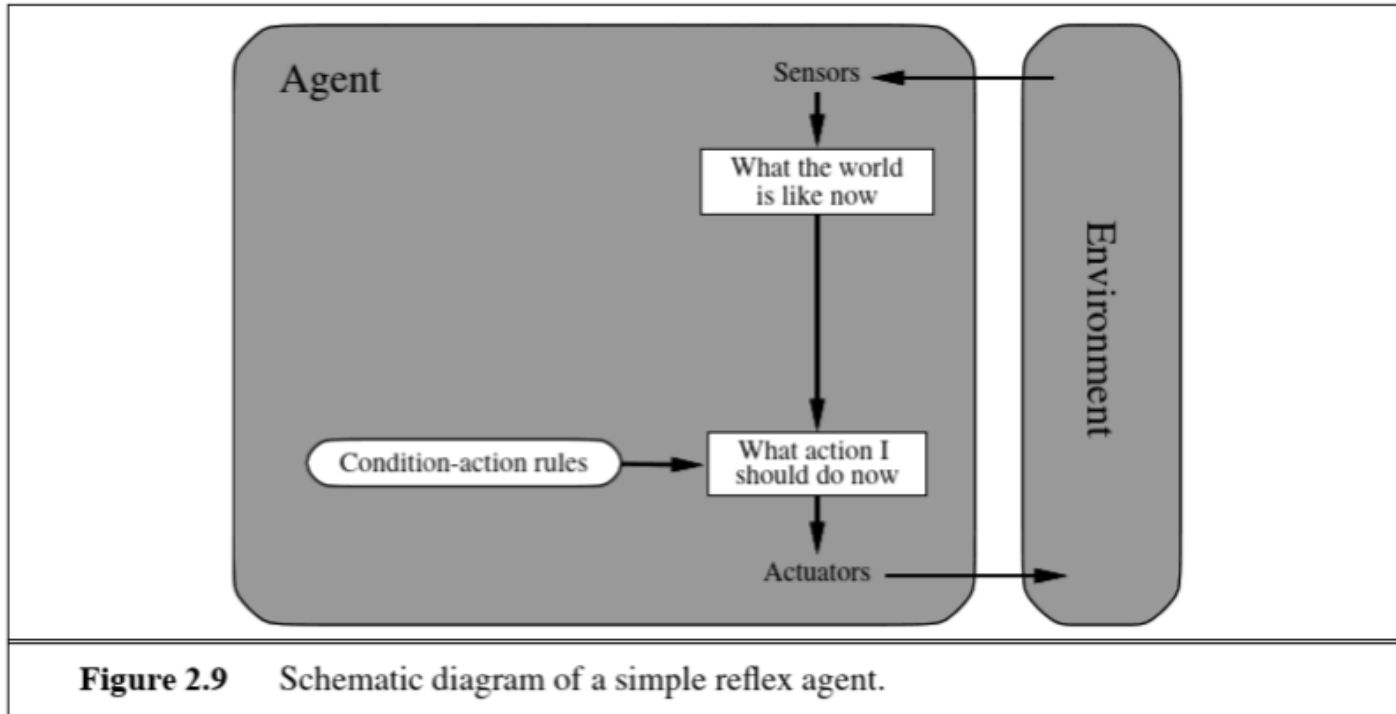
Fourier GR1 / GR3



Autonomous Agent Model

—
ESSAI July 2026

AI Agent Models



AI Agent Models

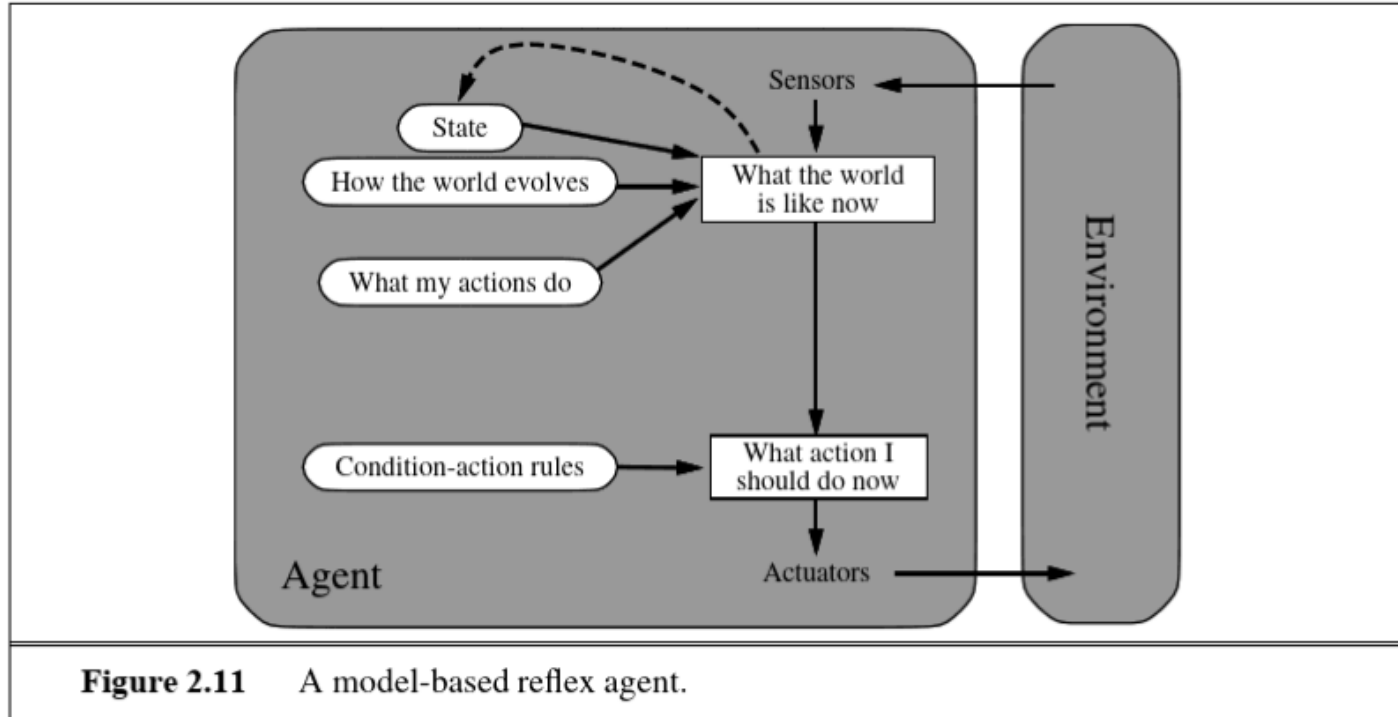


Figure 2.11 A model-based reflex agent.



AI Agent Models

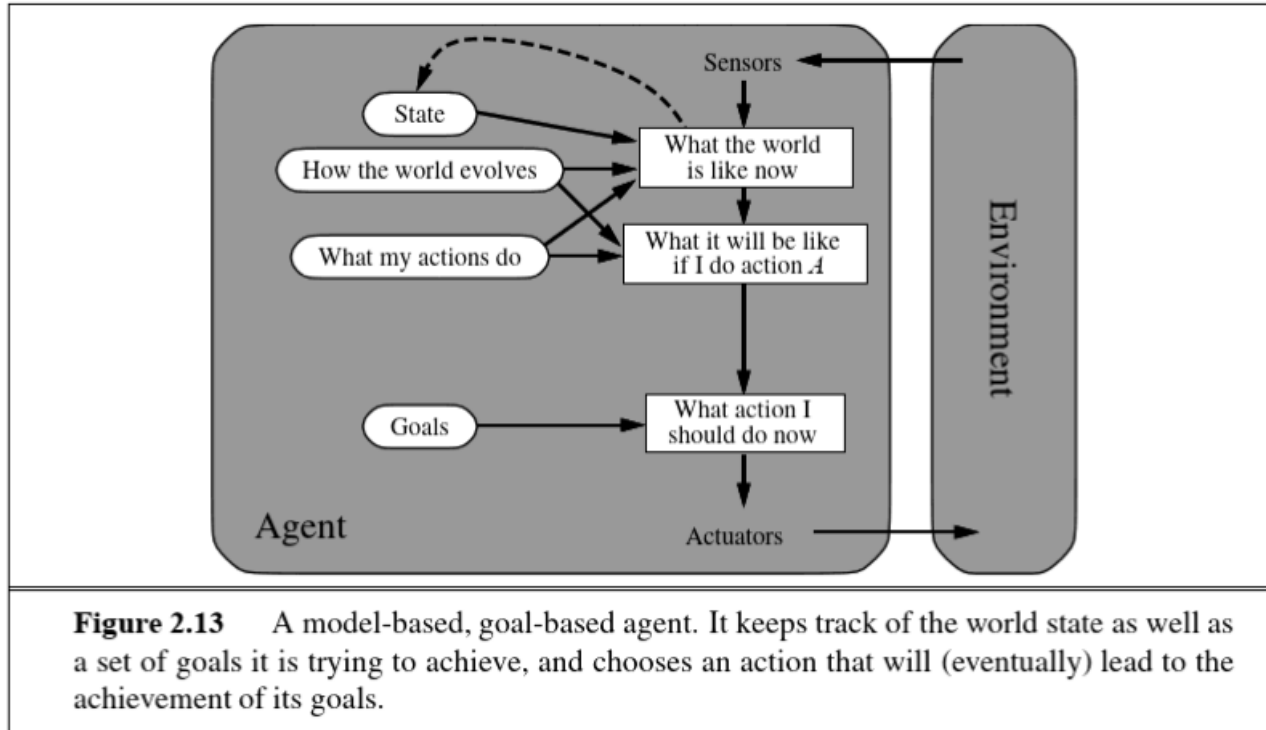
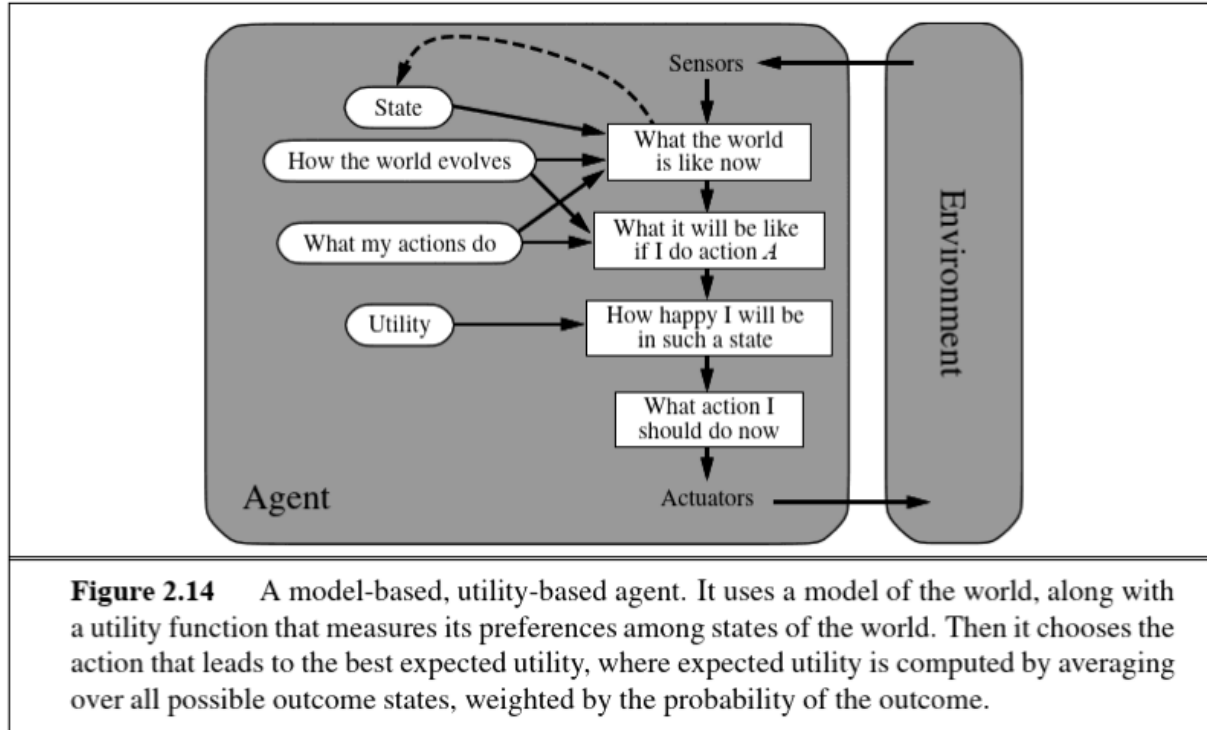


Figure 2.13 A model-based, goal-based agent. It keeps track of the world state as well as a set of goals it is trying to achieve, and chooses an action that will (eventually) lead to the achievement of its goals.



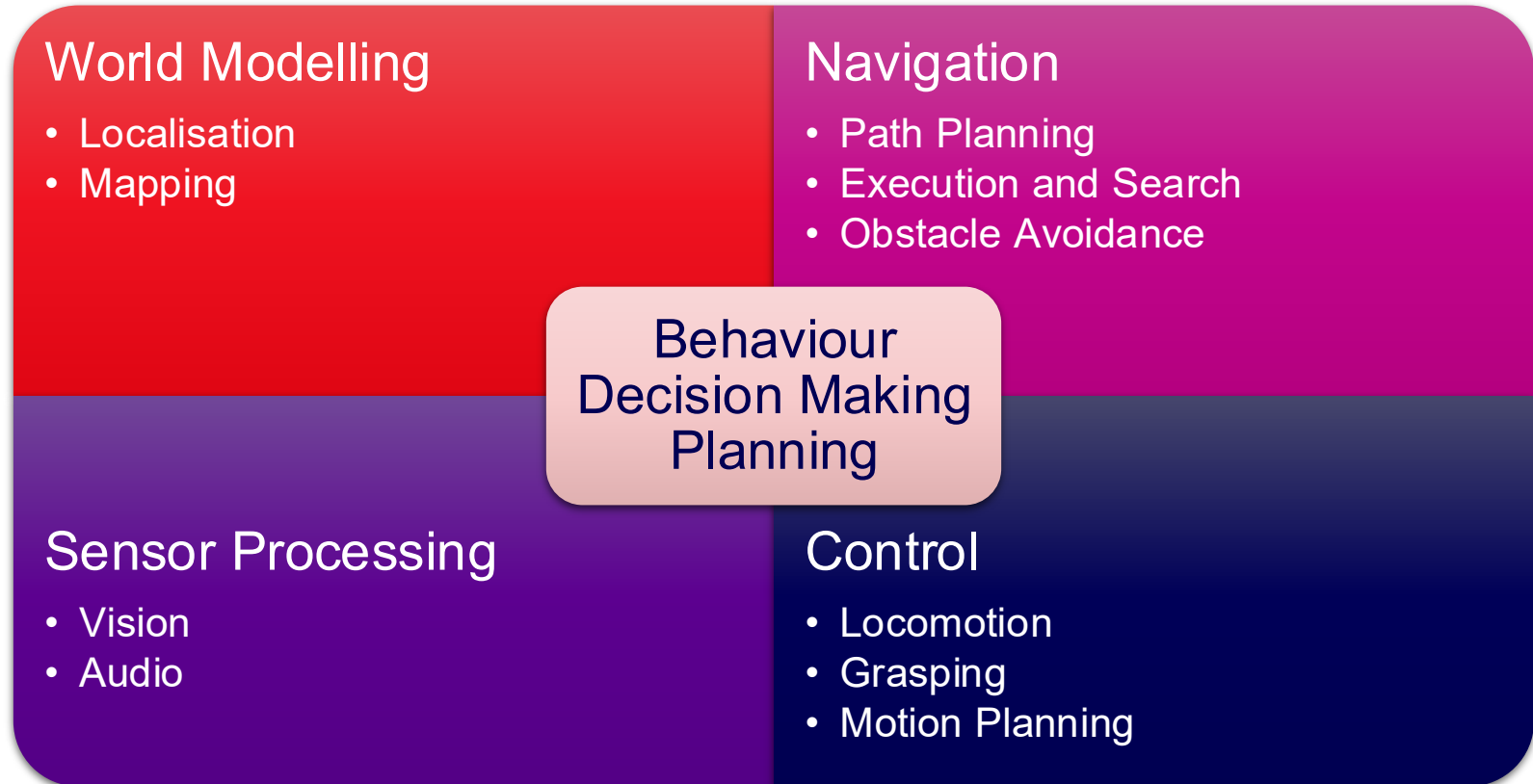
AI Agent Models



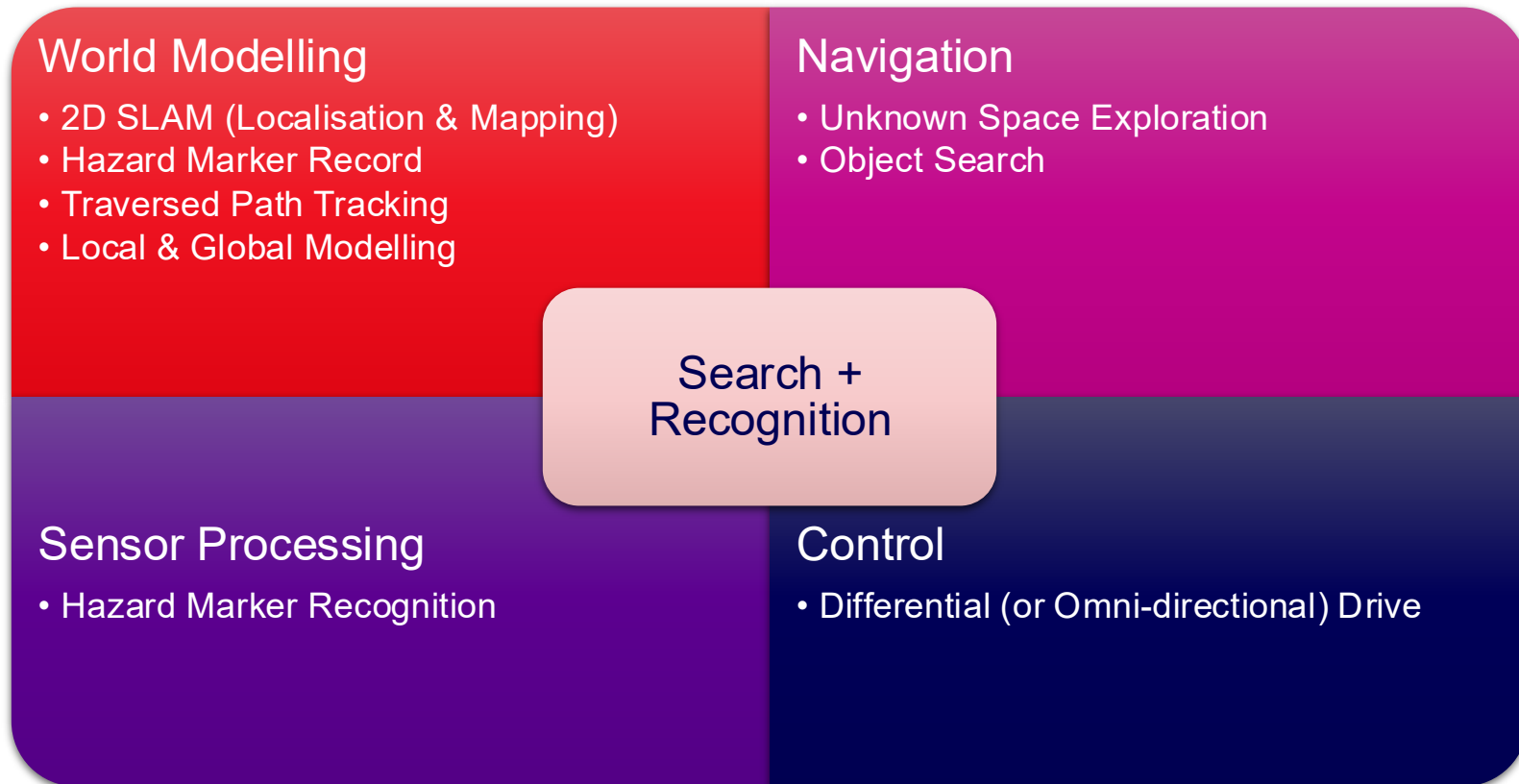
Autonomous Software Components

—
ESSAI July 2026

Autonomous Robot Software Components



ROSBot Maze



Software Architectures

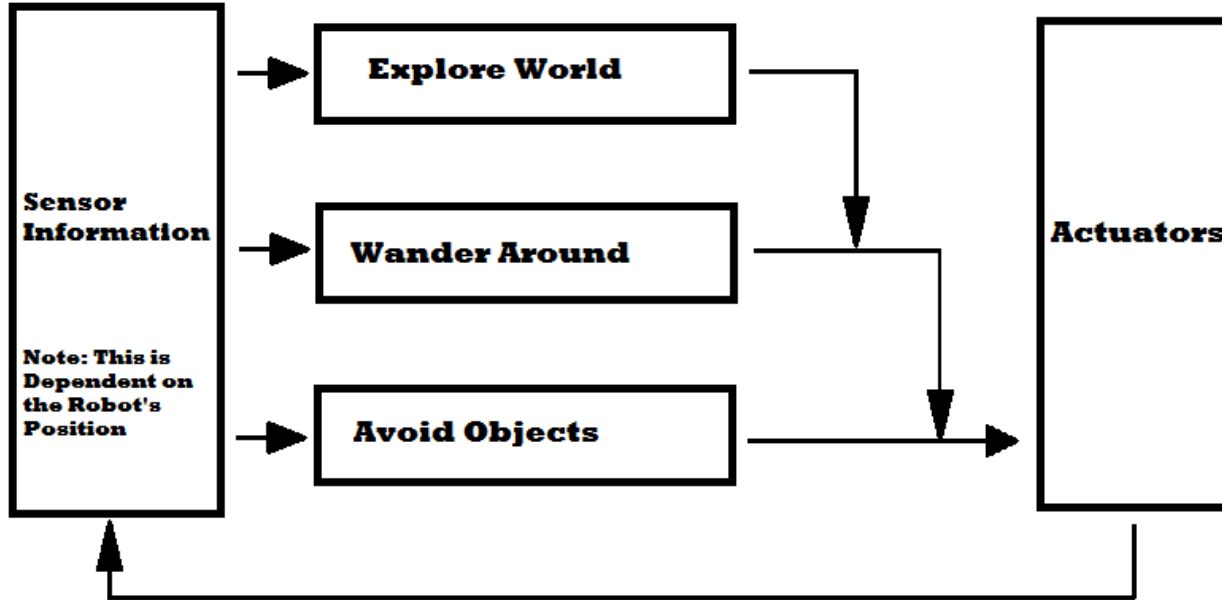
—
ESSAI July 2026



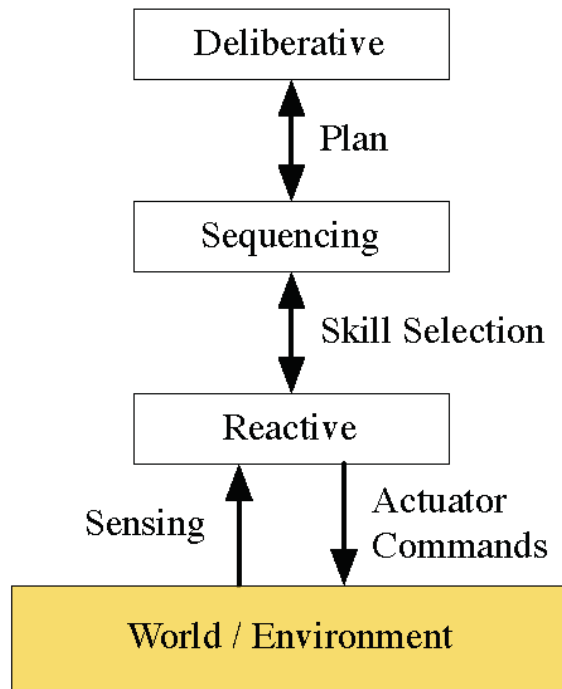
Software Architectures: Sense-Plan-Act



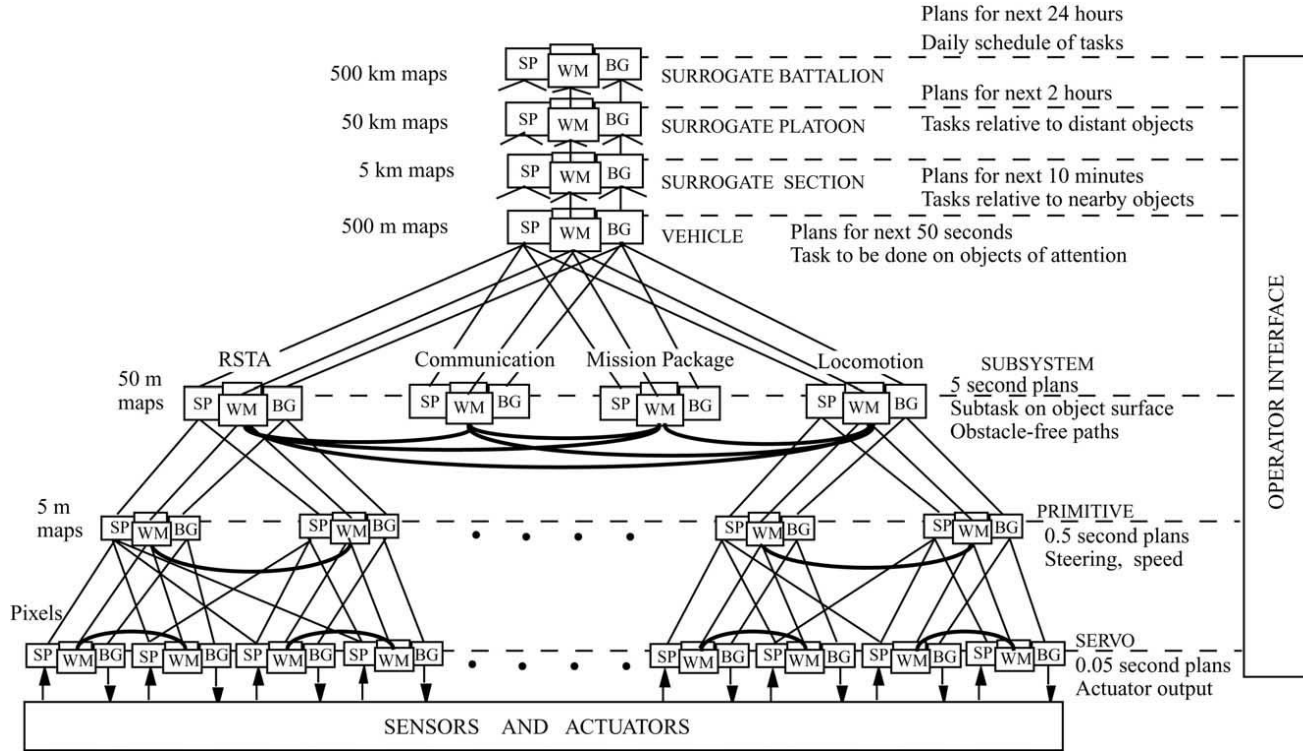
Software Architectures: Subsumption



Software Architectures: Three-Layer



Software Architectures: RCS



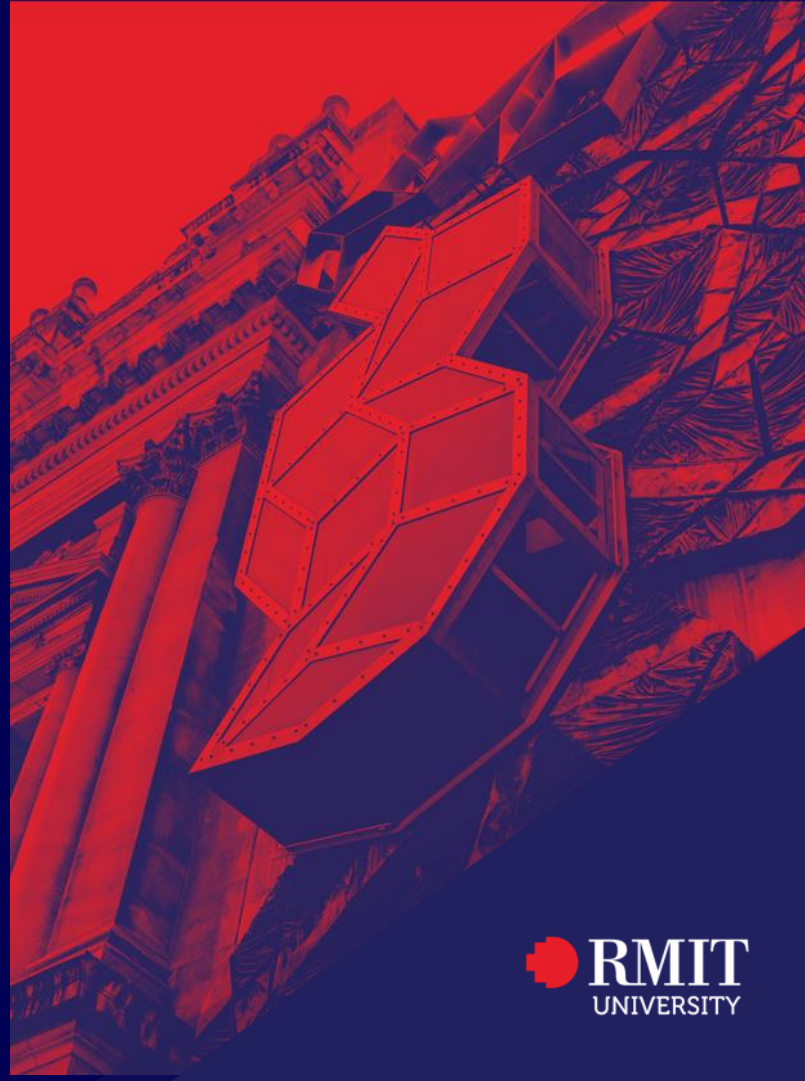
Albus, J. S. & Barbera, A. J. (2005) RCS: A cognitive architecture for intelligent multi-agent systems. *Annual Rev Control* 29, 87–99.



Learning Motion

A practical perspective for
Reinforcement Learning

—
ESSAI July 2026



Day Dreamer





Learning Motion

Algorithms for robot motion must handle a variety of issues unique to the realities of real-world control, including:

- Noise
- Slip
- Hardware failures
- Probabilistic Actions
- Non-deterministic Actions

Reinforcement Learning is a very popular approach as various RL forms can help account for these issues without ‘manual’ intervention.

We could devote a whole course to this topic. For today we’ll keep this to a high-level of how RL methods are applied for robot control, and leave a deep-dive for your own investigations.



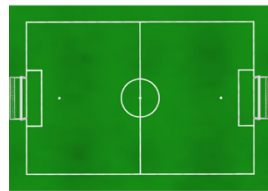
Reinforcement Learning – Overview

An RL problem can be minimally defined as:

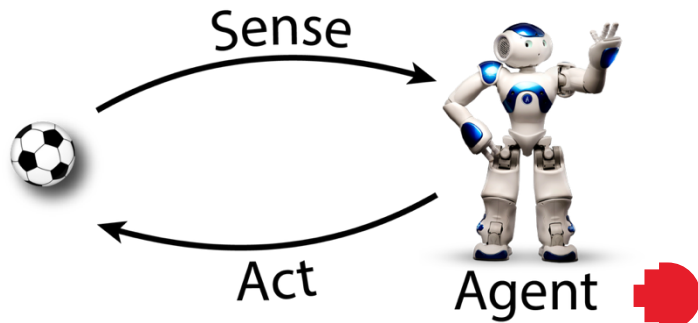
$$RL := \langle S, A, T, V, R, \gamma, \pi \rangle$$

Where:

- S – State Space (discretised, continuous)
- A – Action Set (discrete, continuous)
- T – Transition Function (model-based, or model-free)
- V – Value Function
- R – Reward Function
- γ – Discount Factor
- π – Policy Function



World

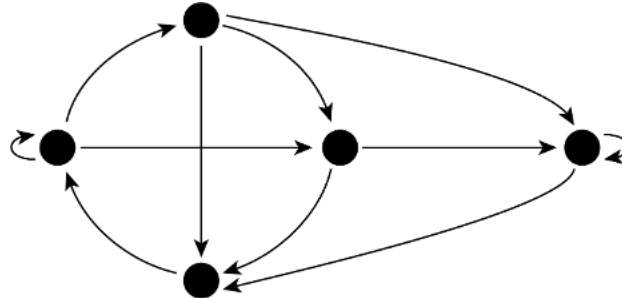


Markov Decision Process

An RL problem defines a Markov Decision Process (MDP) where:

- The decision of which action to perform in a state only requires information that is available in the current state
- The history of states and decisions that the agent took in reaching its current state has no impact on the decision

It is critical that a RL problem is properly defined as an MDP.



Value, Policy, and Q Functions

1. Value function: $V: S \rightarrow \mathbb{R}$

$$V(s) = E[r_i + r_{i+1} + \dots + r_{i+t}] = E\left[r_i \sum_{j=1}^t + r_{i+j}\right]$$
$$V(s) = (1 - \alpha)V(s) + \alpha \max_a [R(s, a) + \gamma V(T(s, a))]$$

2. Policy function: $\pi: S \rightarrow A$

$$\pi^*(s) = \operatorname{argmax}_a [R(s, a) + \gamma V(T(s, a))]$$

3. Q-Learning: $Q(S \times A) \rightarrow A$

$$Q(s, a) = (1 - \alpha)Q(s, a) + \alpha \left[R(s, a) + \gamma \max_{a'} Q(s', a') \right]$$





Learning Motion

All variates of RL forms can be used within a robotics context:

- On policy learning
 - SARS, Value/Policy iteration
- Off-policy learning
 - Classic Q-learning
- Model-free methods
 - Deep Q-Learning
 - Actor-Critic
 - PPO
- Model-based
 - Dreamer / Day Dreamer

... to name a few





Learning Motion: Issues

Reinforcement Learning methods encounter issues:

- Iterations required for convergence
- Balancing Exploration vs Exploitation
- Appropriate definition of the RL problem
- Transferring simulated behaviours to real-systems

This field is rapidly changing.



Classic Example – Cart and Pole

For simplified domains, a traditional representation with any combination of value iteration, policy iteration, q-learning, etc., are successful:

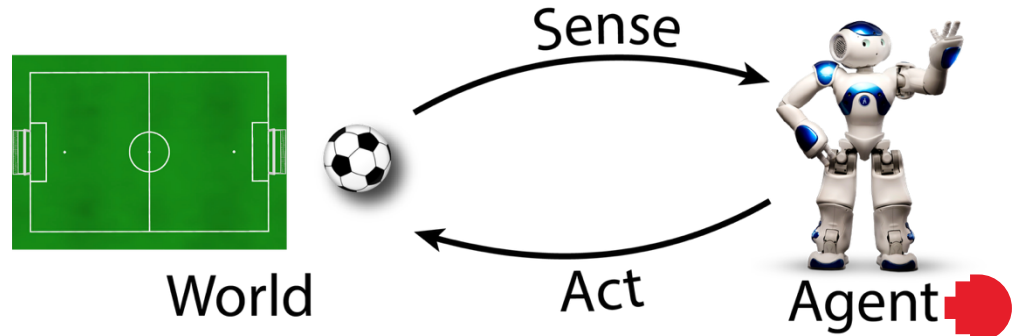
$$RL := \langle S, A, T, V, R, \gamma, \pi \rangle$$

S – Cart/pole position

A – L/R force

T – Either model-based or model-free

R – Balancing time, distance from centre, etc.



Learning Motion: Examples





Deep Reinforcement Learning

Deep Reinforcement Learning methods use the non-linear piecewise function representative power of deep neural networks to encode various parts of the RL problem definition, replacing them with Deep networks.

One of the first successful forms of DRL (2013), is the Deep Q-Network (DQN). Google DeepMind trained a

Model-Free methods: DQN, PPO, SAC, TD3

Model-Based methods: Plan2Explore, Dreamer, Day Dreamer

Mnih, Volodymyr, et al. "Playing atari with deep reinforcement learning." arXiv preprint arXiv:1312.5602 (2013).

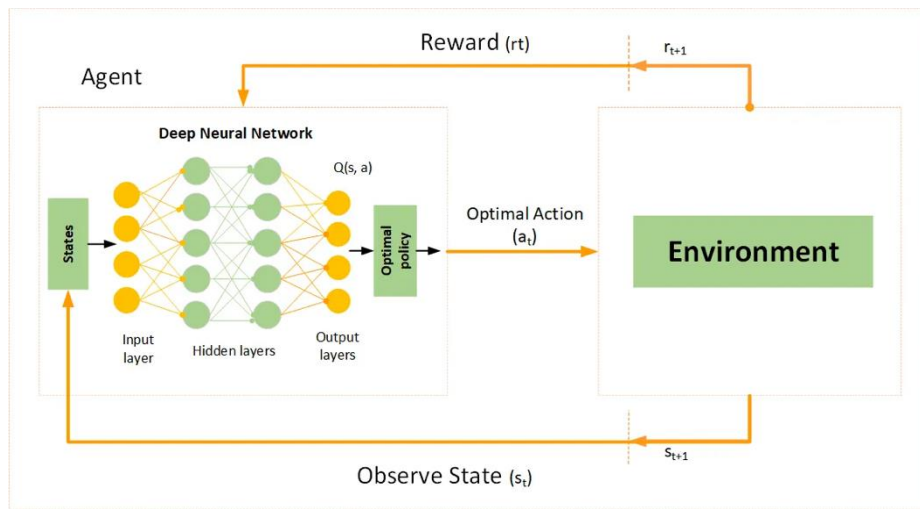


DQN (Deep Q-Network)

$$RL := \langle S, A, T, V, R, \gamma, \pi \rangle$$

This essentially represents the Q-learning Q-function $Q(s, a)$ as a Deep Network. In DQN, the Q-function gives the Loss Function:

$$L(\theta) = E[y - Q(s, a, \theta)^2]$$



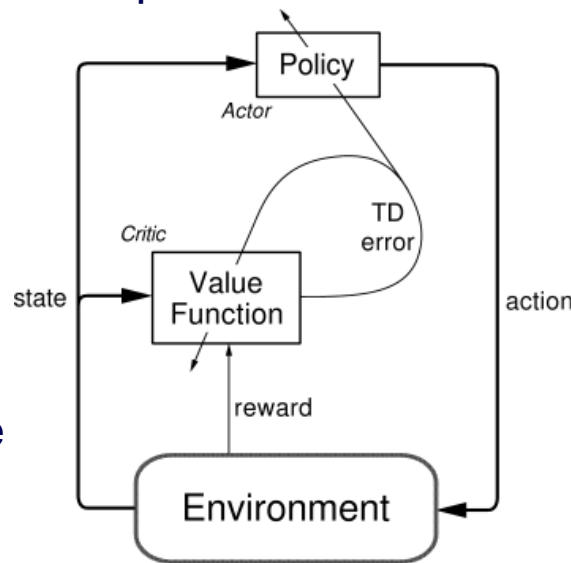
Actor-Critic RL

$$RL := \langle S, A, T, V, R, \gamma, \pi \rangle$$

Actor-Critic RL, such as TD3 and SAC, has two components:

- **Actor:** decides which action to take in a state, represented by a DNN that outputs action probabilities or deterministic actions
- **Critic:** evaluates the quality of the actions taken by the actor, estimating the value of states or actions (ie. the value function)

This essentially models classic RL, where the value and policy functions are independently computed. The error from the critic updates the actor.





DRL – PPO

PPO (2017) is an extension to SAC and DQN that address the problem of instability in the value/policy networks as well as improving the capability of learning. PPO has now become the standard method for DRL.

PPO introduced concepts of:

- The Advantage
- Clipping
- An Objective Function (to replace the Q- function)

Schulman, John; et. a. (2017), Proximal Policy Optimization Algorithms, arXiv:1707.06347





DRL – PPO

The PPO loss-function is

$$L(\theta) = \mathbb{E}_t[\min(r_t(\theta)A_t, \text{clip}(r_t(\theta), 1 - \epsilon, 1 + \epsilon)A_t))]$$

This finds the expectation of the best advantage at each timestep t

- r_t is the ratio of the probability of selecting an action under the old or new policies, and is not to be confused with the reward!
- A_t is the advantage of an action at timestep t
- ϵ is a hyper-parameter controlling how much each update is clipped



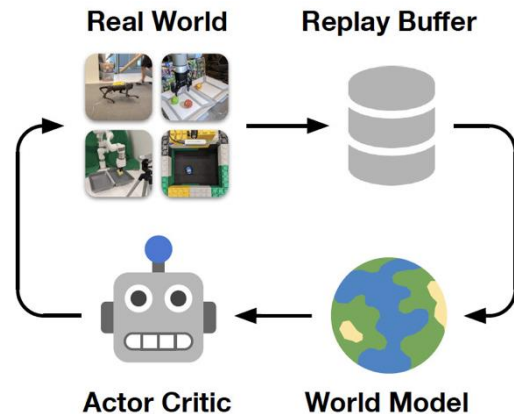
DRL – Dreamer/DayDreamer

As model-free methods don't have an explicit transition function, the the value and policy functions must also “internally encode” approximations of the transitions.

Model-based methods provide an explicit model, which can improve learning. DayDreamer (2023*) is one of the state-of-the-art methods with two deep networks:

- Latent world model models environment dynamics to predict future observations, rewards, and values in latent space.
- Behaviour model - Actor-Critic RL that incorporates world model state predictions

*Wu, P., Escontrela, A., Hafner, D., Goldberg, K. & Abbeel, P.
DayDreamer: World Models for Physical Robot Learning. in
Conference on robot learning 2226–2240 (PMLR, 2023)*



DRL – Design Challenges





DRL – Design Challenges

Compared to other applications of Deep RL, robotics poses additional challenges for efficient and accurate learning:

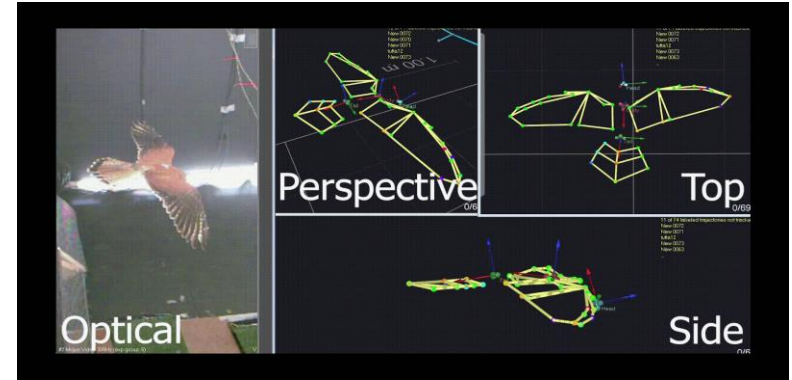
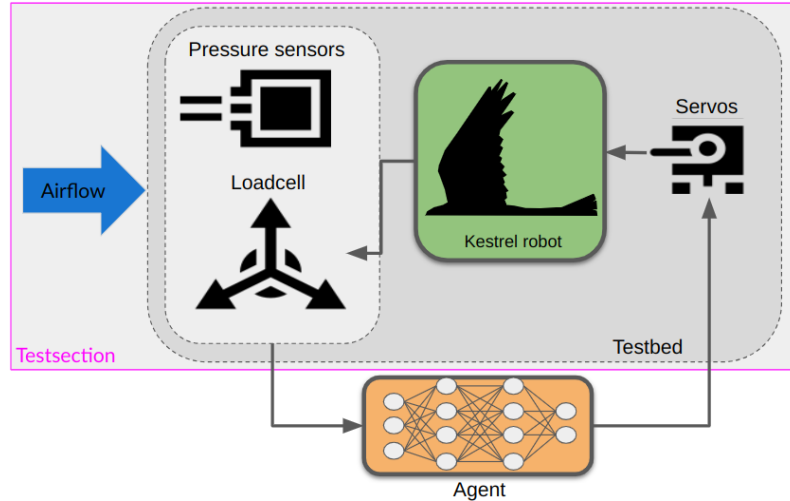
- Sensing and actuation noise
- Non-deterministic action outcome
- Semi-Markovian actions with variable time
- Simulation to real gap
- Hardware damage and safety of learning

Therefore, careful consideration of the State Space, Action Space, and Reward Function is required, including Reward Function Shaping to help guide learning.



Bio-inspired Flight

Online Deep Reinforcement Learning to control a bio-inspired model



Stiemer, L., Martinez Groves-Raines, M., Wood, L., Mohamed, A. & Wiley, T. (2025)
Online Deep Reinforcement Learning of Servo Control for a Small-Scale Bio-Inspired Wing.
AI 2024: Advances in Artificial Intelligence. Lecture Notes in Computer Science. Vol. 15443, pp. 65–76.



Bio-inspired Flight

Goal: learn target lift and trim condition

State Space:

- 3DOF Wing configuration
- error to target lift
- target pitch moment, angle of attack
- wind flow velocity
- wing force loading
- state history

Action Space:

- Continuous servo positions

(V1) Reward Function:

- Negative weighted error to target lift
- If out of tolerance, high negative

(V2) Reward Function + Shaping:

- Negative weighted error to target lift and pitch
- Penalty terms for smoothing servo deflections, minimising oscillations
- If out of tolerance, high negative



Bio-inspired Flight

Kestrel Wing Phase 2: Training

TD3 Model

Kestrel Wing Phase 3: Evaluation

TD3 Model

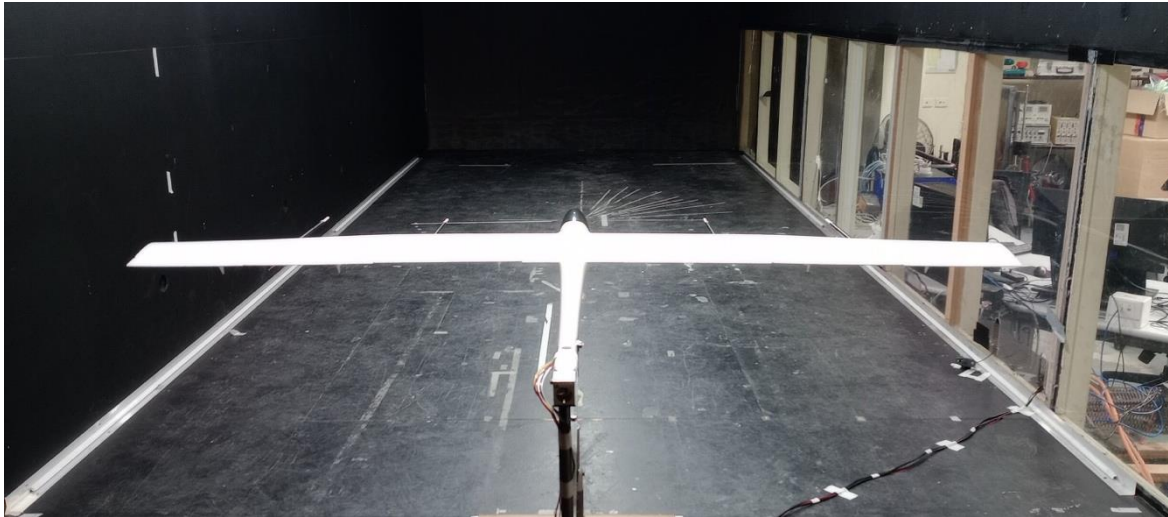


HALT Aircraft

Online Deep Reinforcement Learning to mitigate gust disturbances in HALT aircraft, with multi-segmented ailerons.

Goal: hold stable level flight in gusting conditions, minimising wing loading.

Common issues in DRL become prevalent, such as “learning something” but not learning “something good”.



HALT Aircraft

State space variations:

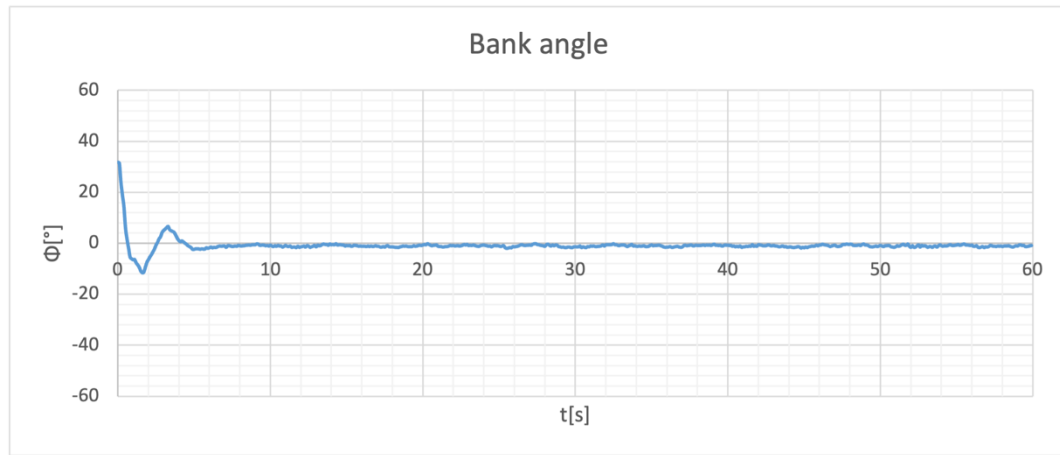
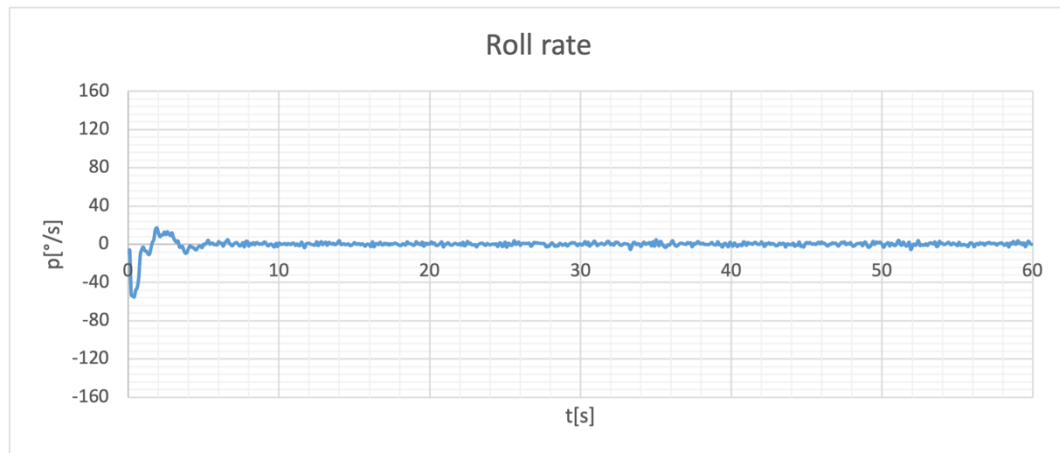
- Roll / Bank absolute values
- Roll / Bank derivatives (rates)
- Aileron deflection value
- Wind velocity

Action space:

- Aileron deflection command

Reward variations:

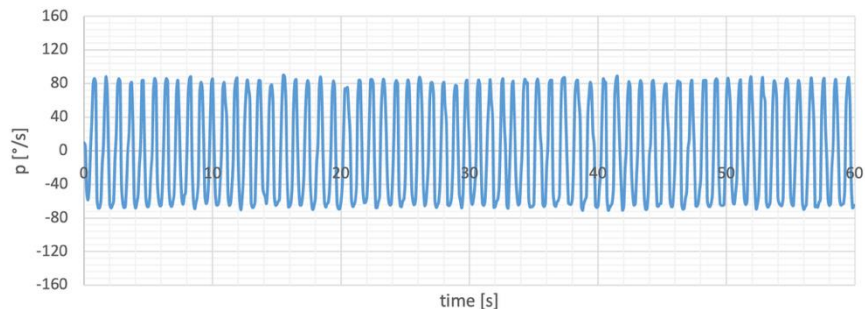
- Roll / bank values / rates
- Rate of aileron command change



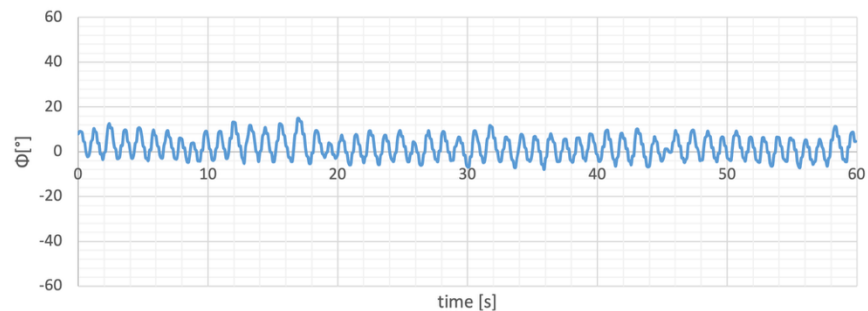
HALT Aircraft

Common issues in DRL become prevalent, such as “learning something” but not learning “something good”.

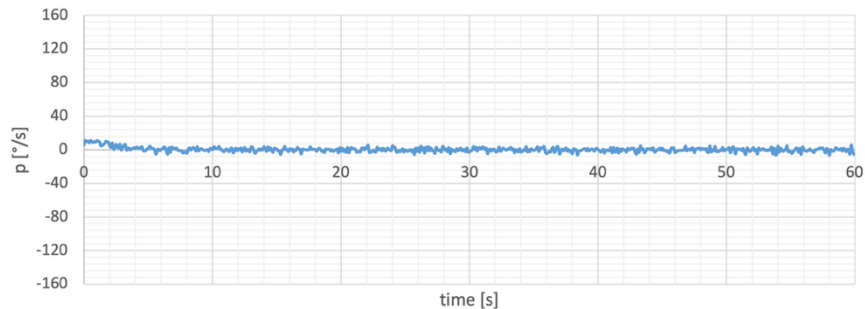
Roll rate



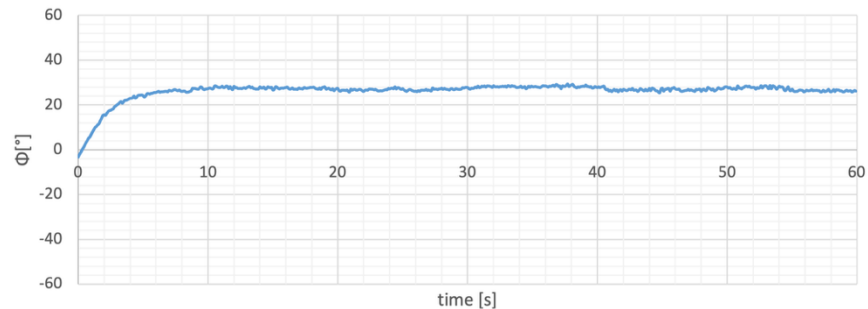
Bank angle



Roll rate



Bank angle



Learning from Demonstration: Humanoid Robots

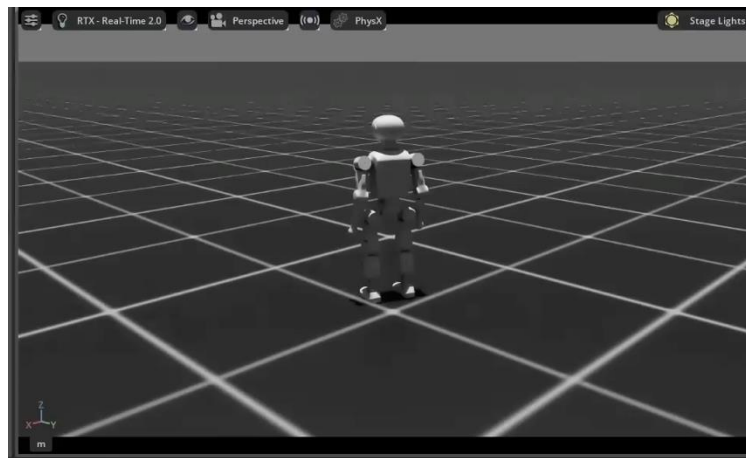
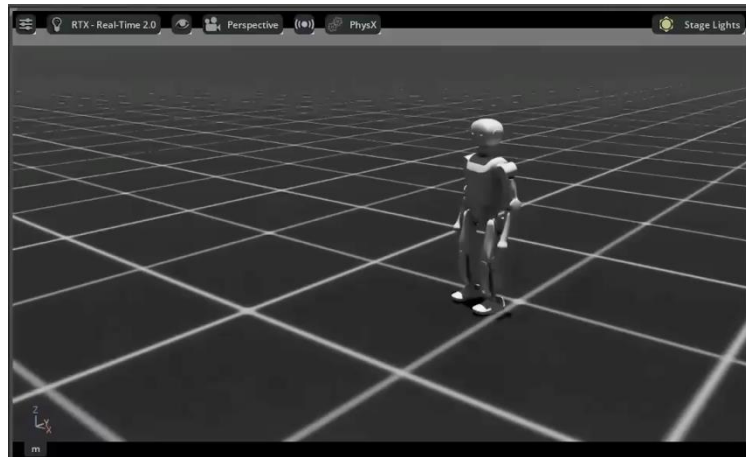
In various applications it can be challenging to define precisely the “reward” of learning.

Consider a humanoid robot walk or kick.

What is the “reward” that means success?

Alternative approach is to learn from demonstration of desired behaviour. In humanoid applications, this uses motion capture to track a human movement.

Learning “replays” the behaviour.



Learning from Demonstration: Humanoid Robots

Reward function:

- Error of joints to motion capture
- “Progression” through sequential stages of motion capture
- Penalty terms for dangerous motion

Current simulators have a very close Simulation to real gap that enable zero-shot “satisficing” deployment.

Behaviours for online still require additional learning and refinement.



Noon Gudgin

Thank you

Day 2:

Motion Planning and
Navigation

ESSAI July 2026

