

ADTM 2026 Conference

Digital Twin as an additional security protection layer in Industrial Control Systems (ICS)

Dr. Guy Tel-Zur
Ben-Gurion University of the Negev

gtelzur@bgu.ac.il

15.6.2026

'The Road Not Taken'
by Robert Frost

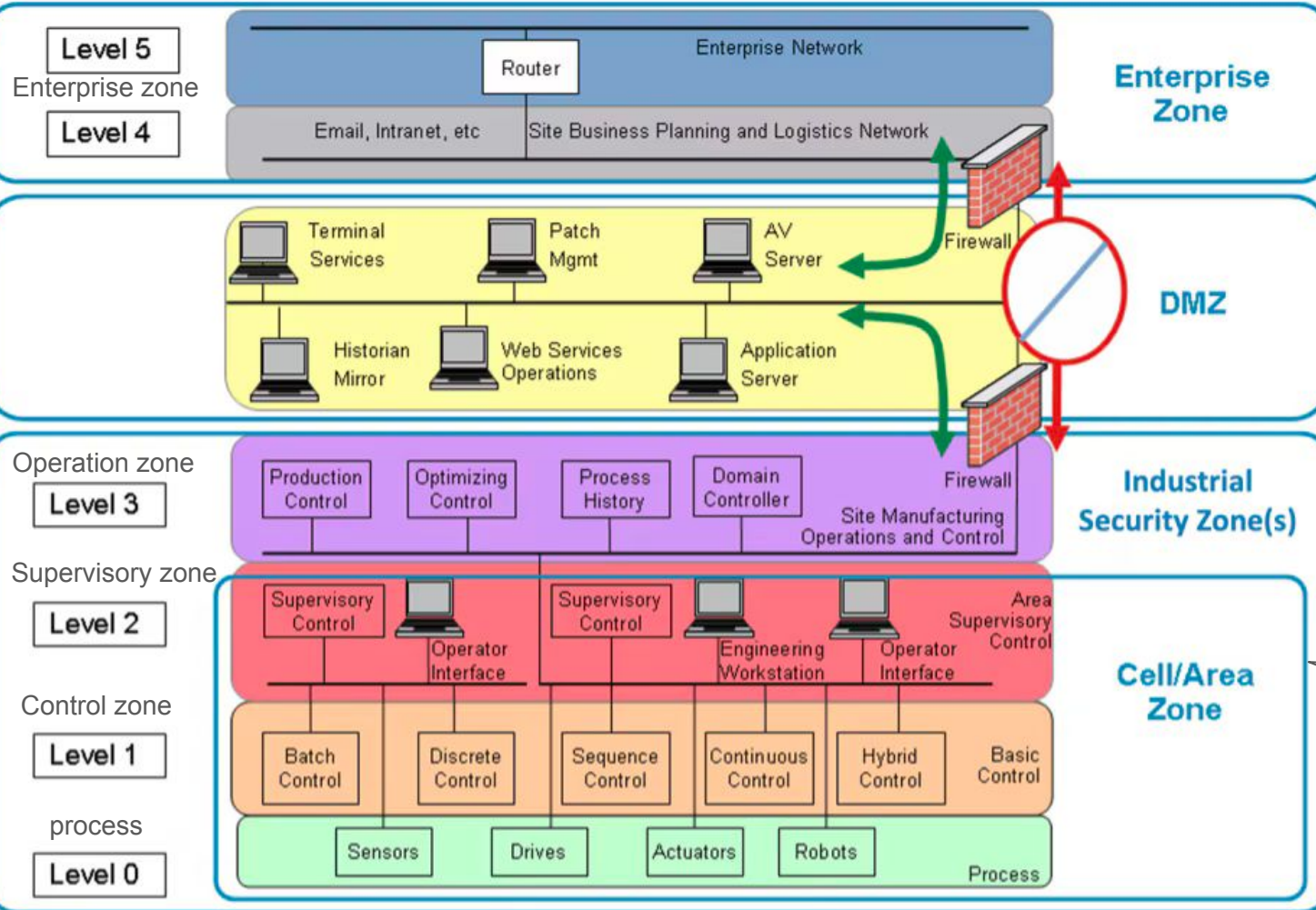
Talk agenda

- What is an ICS?
- Cybersecurity Risks
- What is a Digital Twin (DT)?
- How a Digital Twin can help protecting an ICS?
- An example: Reinforcement Learning AI-assisted DT
- Project: 4th year engineering students final project
- Outlook and Summary

What is an ICS?



Purdue Model



This work

Credit:
<https://www.cisco.com/site/us/en/learn/topics/security/what-is-ot-security.html>

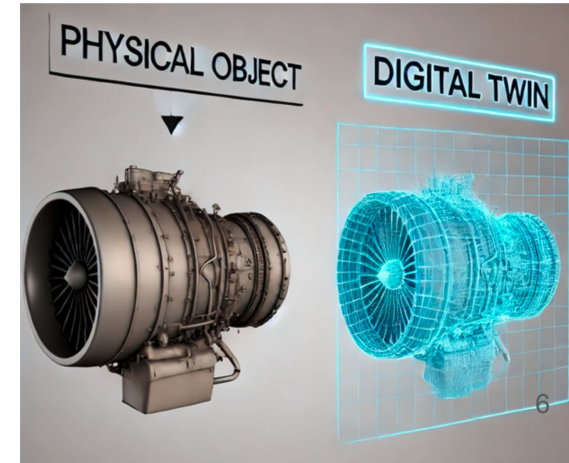
Digital Twin

A virtual replica of a physical system that updates in real-time

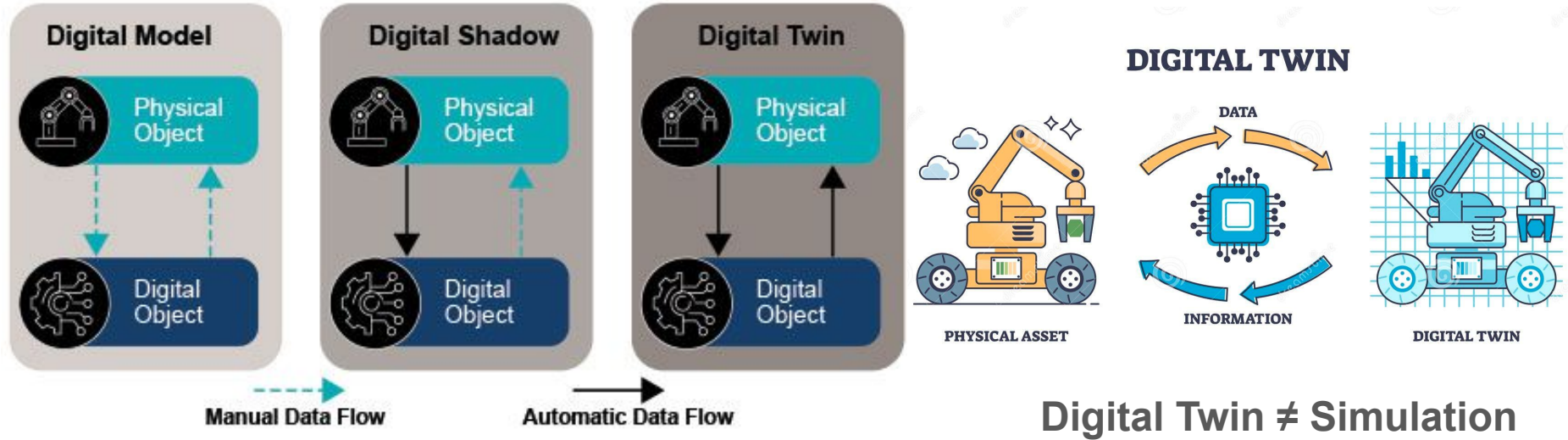
ISO/IEC 30173: *“Digital representation of a target entity with data connections that enable convergence between the physical and digital states at an appropriate rate of synchronization”*

The Digital Twin mirrors the real object's:

- Current state
- Behavior
- Performance
- Even predicts future states



Digital Twin



2 issues:

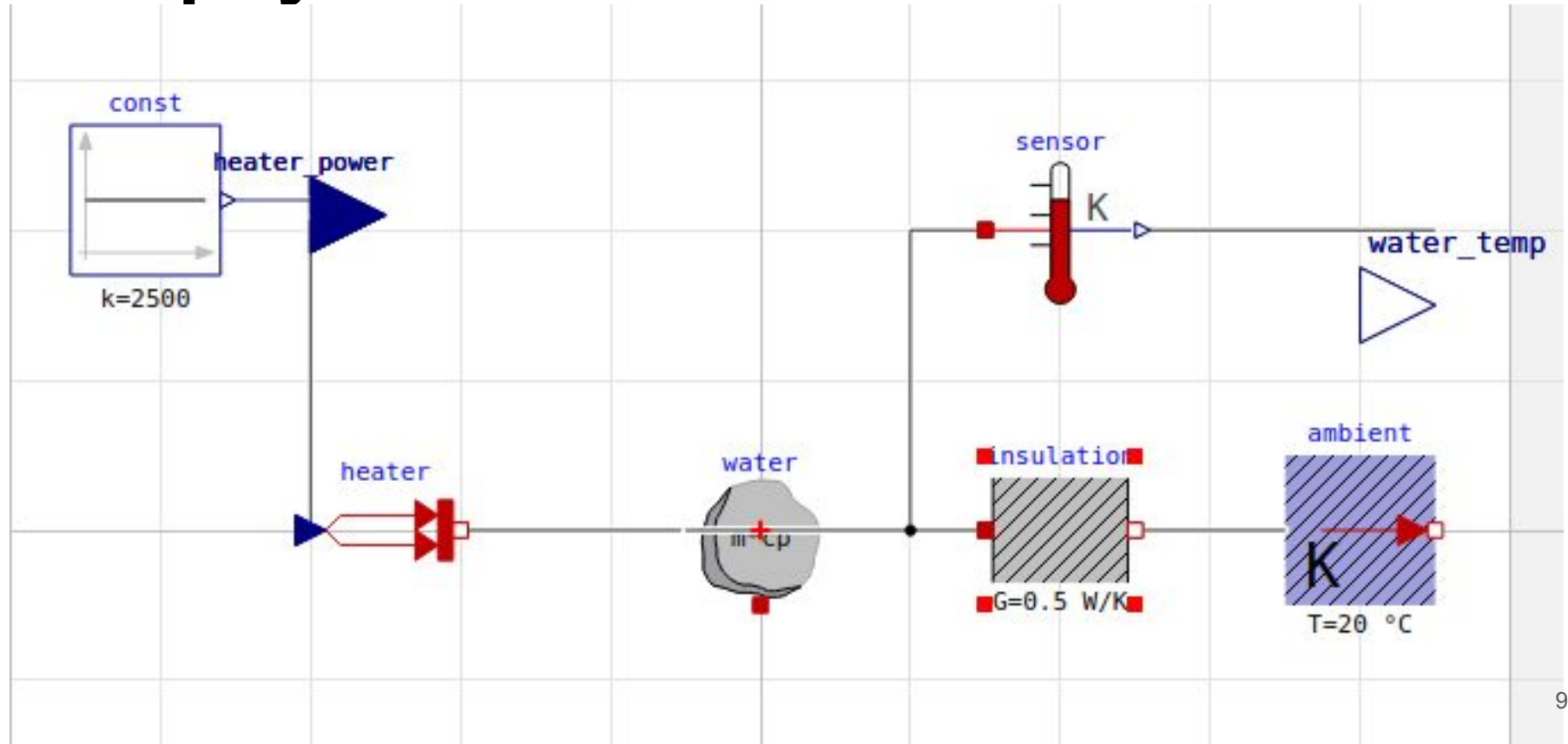
- How to use the Digital Twin for protection of physical system?
- How to protect the Digital Twin and the access to its models to make sure its main functions?

Water-Heating Control with an AI-Assisted Controller and an FMU-Based Digital Twin

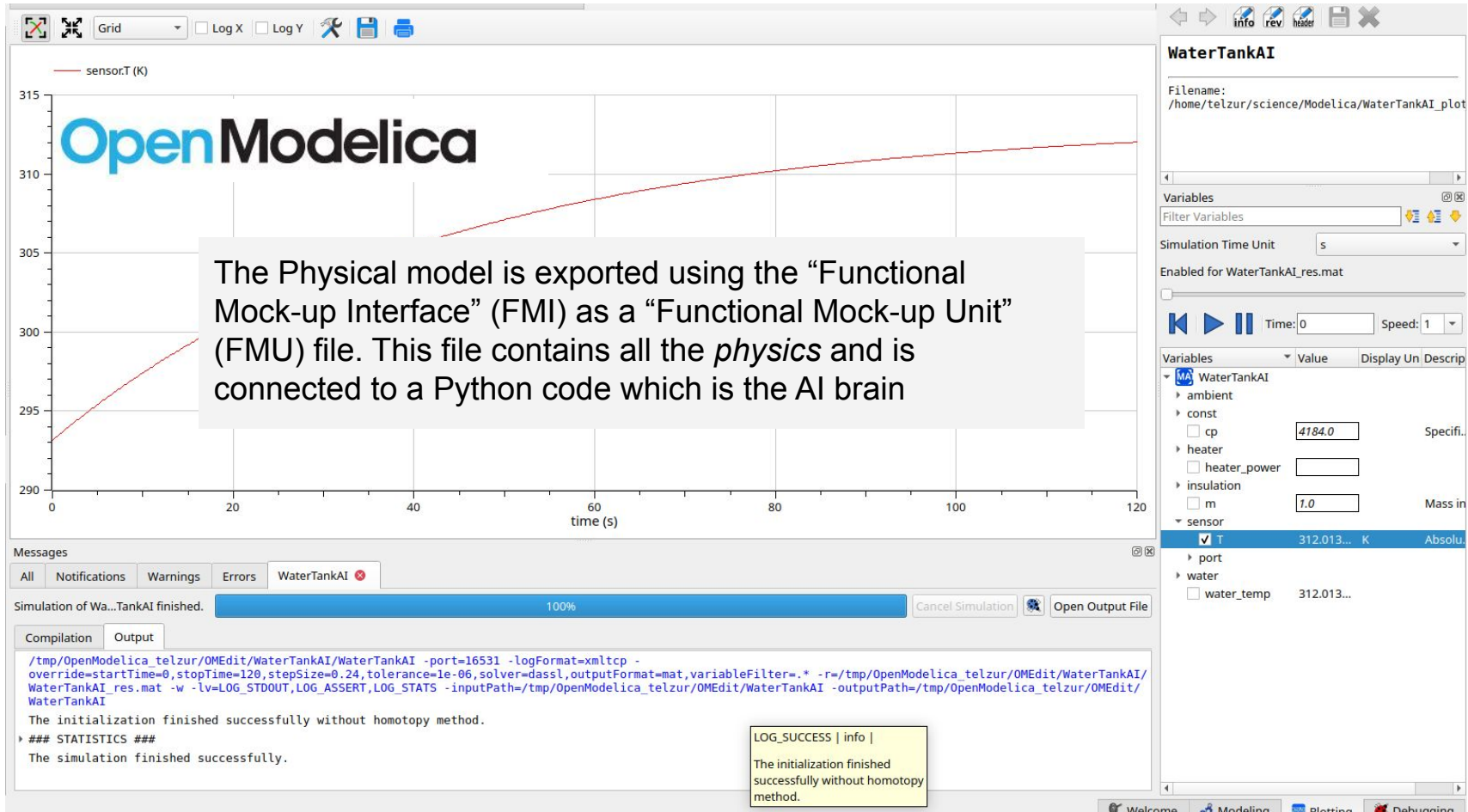
- A kettle with water
- Heater: up to 2500W
- Heating profile: Starting from room temperature (20°C), heat the water to 40°C within 120 seconds, then keep the temperature steady at 40°C for an additional 60 seconds and then let it naturally cool down.
- Water mass (1 Kg. $\rightarrow$ 1 Liter)
- Heat loss via the kettle's sides to the outside environment
- Assumptions: A simple lumped-parameter model with a uniform water/pot temperature model (i.e. no boiling, no bubbles, no turbulence, etc').



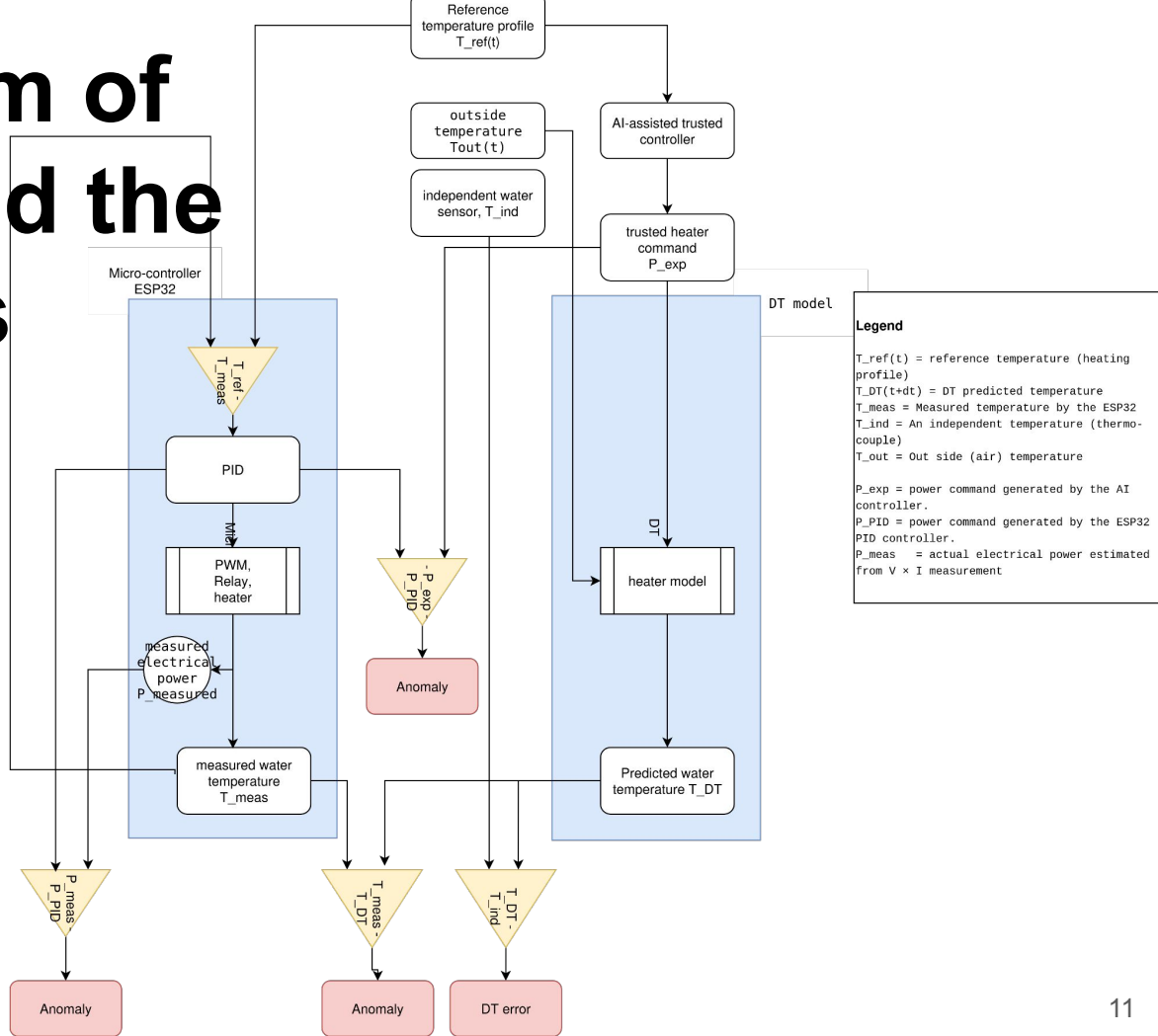
The physical model in Modelica



Model simulation in OpenModelica

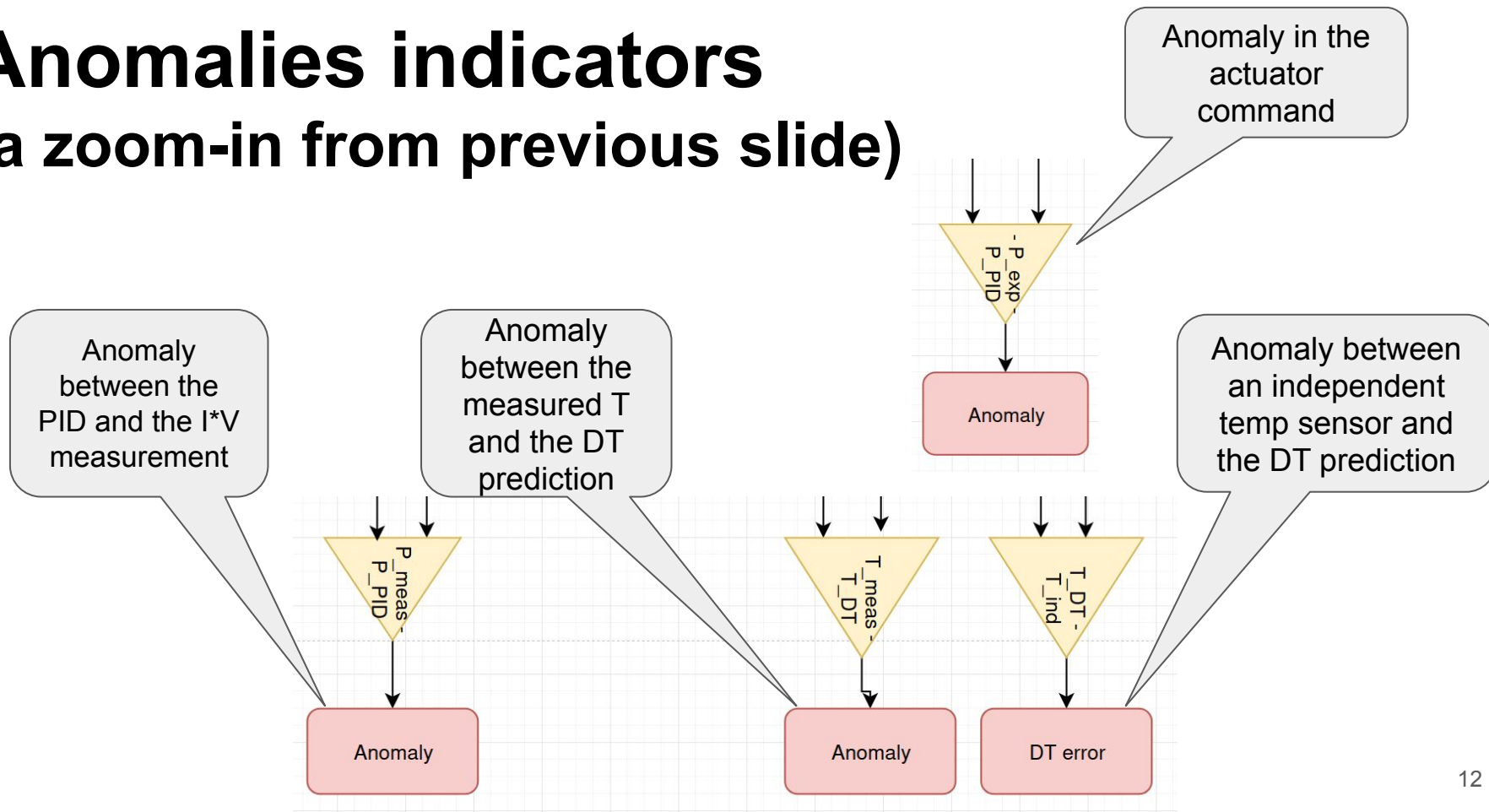


A Block diagram of the Physical and the Virtual systems



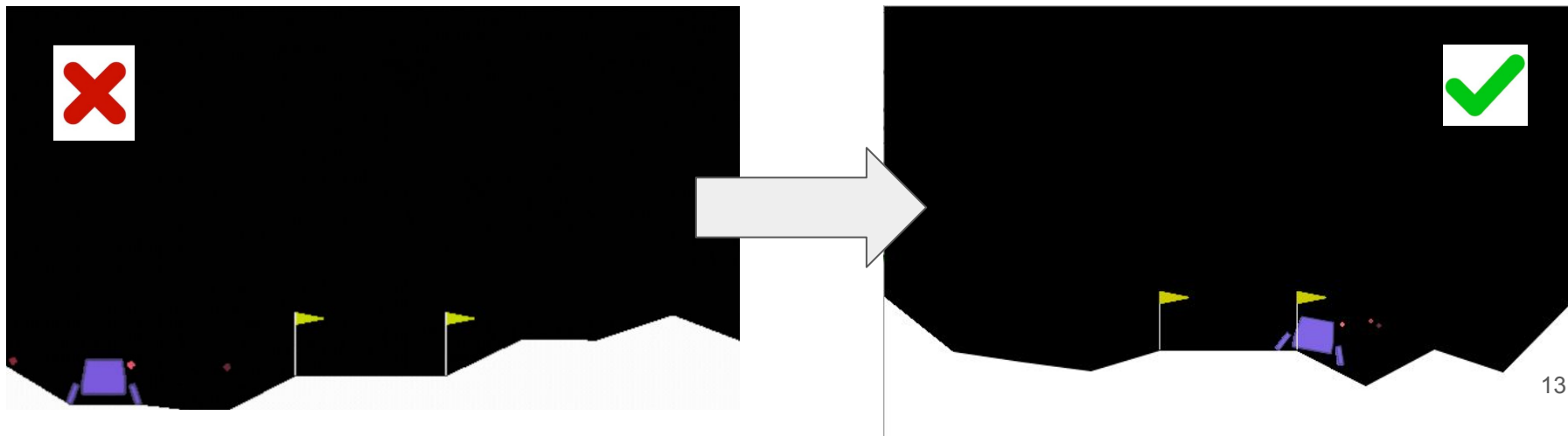
Anomalies indicators

(a zoom-in from previous slide)



Selecting an AI Method for the Trusted Controller

- MLP
- LSTM
- Other models...?
- Reinforcement Learning (RL) models:
 - **Proximal Policy Optimization (PPO)** - failed
 - **Advantage Actor-Critic (A2C)** - not evaluated; supports both discrete and continuous policies
 - **Soft Actor Critic (SAC)** - success. May be an overkill for this specific project (continuous)



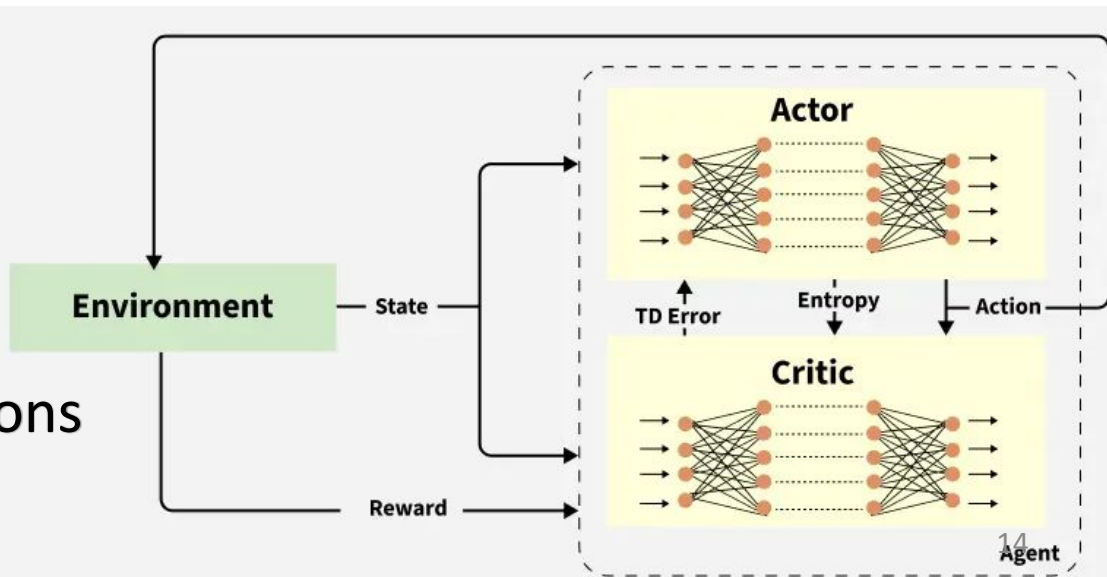
Soft Actor-Critic (SAC): Core Idea

- Actor: chooses actions probabilistically
- Critic: evaluates actions

Key idea: Maximize Reward + $\alpha \times$ Randomness (Entropy)

Why?

- Encourages exploration
- Avoids getting stuck
- Keeps multiple good options

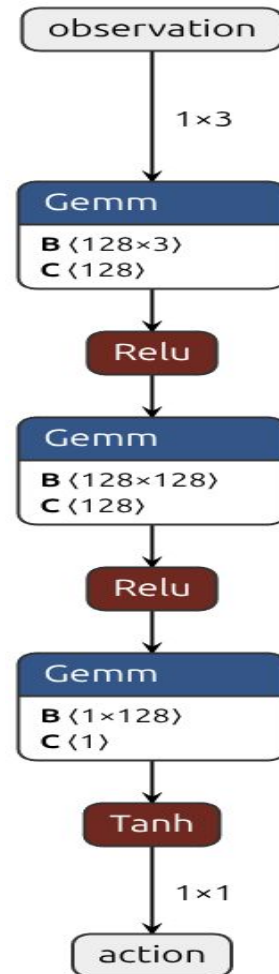


The NN structure

#Parameters = Inputs X Outputs + Biases

- Layer 1 (Gemm 1):** Inputs: 3 $\rightarrow$ Outputs: 128
 - Weights: $3 \times 128 = 384$
 - Biases: 128
 - Subtotal: 512**
- Layer 2 (Gemm 2):** Inputs: 128 $\rightarrow$ Outputs: 128
 - Weights: $128 \times 128 = 16,384$
 - Biases: 128
 - Subtotal: 16,512**
- Layer 3 (Gemm 3):** Inputs: 128 $\rightarrow$ Outputs: 1
 - Weights: $1 \times 128 = 128$
 - Biases: 1
 - Subtotal: 129**

Total number of trainable parameters: $512 + 16,512 + 129 = 17,153$



GEMM =
General
Matrix Multiply

Relu = Rectified
linear unit

$$f(x) = \max(0, x) = \begin{cases} x_i & \text{if } x_i \geq 0 \\ 0 & \text{if } x_i < 0 \end{cases}$$

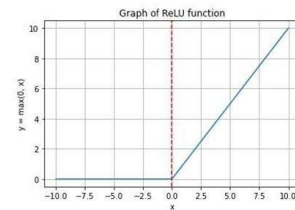
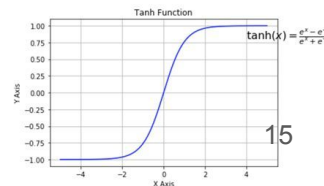


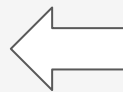
Fig 5: Plot of ReLU function.

Tanh

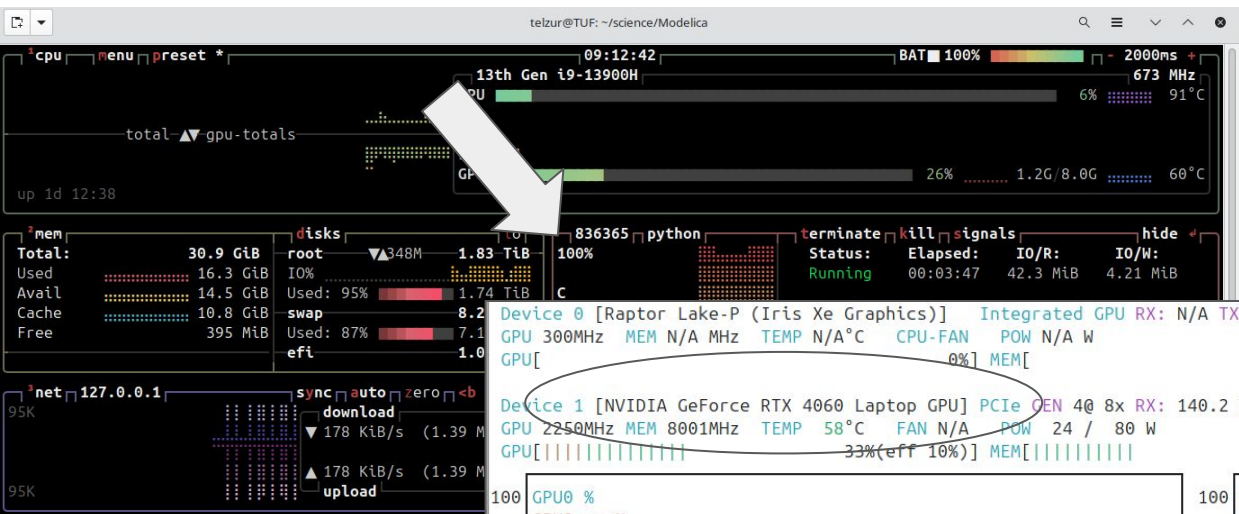


SAC training both on the CPU and the GPU

```
real    9m45.703s
user    9m22.963s
sys     0m9.045s
```

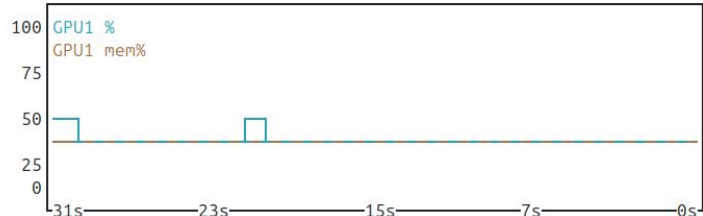
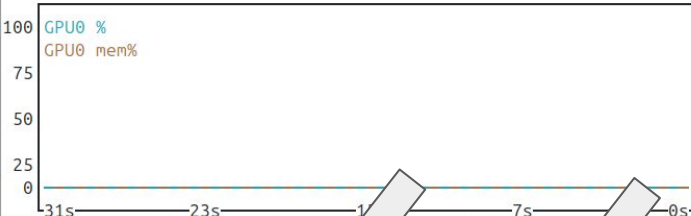


```
(mypython) telzur@TUF:~/science/Modelica$
```



Device 0 [Raptor Lake-P (Iris Xe Graphics)] Integrated GPU RX: N/A TX: N/A
GPU 300MHz MEM N/A MHz TEMP N/A°C CPU-FAN POW N/A W
GPU[|||||] MEM[|||||] 0.051Gi/30.969Gi

Device 1 [NVIDIA GeForce RTX 4060 Laptop GPU] PCIe GEN 40 8x RX: 140.2 MiB/s TX: 19.22 MiB/s
GPU 2250MHz MEM 8001MHz TEMP 58°C FAN N/A POW 24 / 80 W
GPU[|||||] MEM[|||||] 1.962Gi/7.996Gi



PID	USER	DEV	TYPE	GPU	GPU MEM	CPU	MEM	Command
6826	telzur	1	Graphic	0%	359MiB	4%	0%	/usr/lib/xorg/Xorg vt2 -displayfd 3 -auth /run/user/1000/gdm/Xauthority
1407404	telzur	1	Compute	30%	152MiB	2%	99%	/home/telzur/mypython/bin/python3 -m ipykernel_launcher -f /home/telzur
1390939	telzur	1	Compute	0%	146MiB	2%	0%	/home/telzur/mypython/bin/python3 -m ipykernel_launcher -f /home/telzur
1405066	telzur	1	Compute	0%	146MiB	2%	0%	/home/telzur/mypython/bin/python3 -m ipykernel_launcher -f /home/telzur
1405667	telzur	1	Graphic	0%	140MiB	2%	0%	/home/telzur/mypython/bin/python3 -m ipykernel_launcher -f /home/telzur

RL SAC Coding in Python



Gymnasium is a standard interface for reinforcement learning (RL) environments.

- A unified API for RL environments (like a "playground" where agents learn)
- A `reset()` → `step()` loop: observation, action, reward, done
- Popular environments (CartPole, MountainCar, Atari, etc.)
- Originally forked from OpenAI's Gym, now maintained separately

Stable Baselines 3 (SB3) is a collection of pre-implemented RL algorithms.

- Ready-to-use algorithms: PPO, A2C, DQN, SAC, TD3, etc.
- Clean, well-documented implementations
- Works on top of Gymnasium environments
- Handles the training loop, logging, and evaluation

Abstract

Model-free deep reinforcement learning (RL) algorithms have been demonstrated on a range of challenging decision making and control tasks. However, these methods typically suffer from two major challenges: very high sample complexity and brittle convergence properties, which necessitate meticulous hyperparameter tuning. Both of these challenges severely limit the applicability of such methods to complex, real-world domains. In this paper, we propose soft actor-critic, an off-policy actor-critic deep RL algorithm based on the maximum entropy reinforcement learning framework. In this framework, the actor aims to maximize expected reward while also maximizing entropy. That is, to succeed at the task while acting as randomly as possible. Prior deep RL methods based on this framework have been formulated as Q-learning methods. By combining off-policy updates with a stable stochastic actor-critic formulation, our method achieves state-of-the-art performance on a range of continuous control benchmark tasks, outperforming prior on-policy and off-policy methods. Furthermore, we demonstrate that, in contrast to other off-policy algorithms, our approach is very stable, achieving very similar performance across different random seeds.

1. Introduction

Model-free deep reinforcement learning (RL) algorithms have been applied in a range of challenging domains, from games (Mnih et al., 2013; Silver et al., 2016) to robotic control (Schulman et al., 2015). The combination of RL and high-capacity function approximators such as neural networks holds the promise of automating a wide range of decision making and control tasks, but widespread adoption

¹Berkeley Artificial Intelligence Research, University of California, Berkeley, USA. Correspondence: Tuomas Haarnoja <haarnoja@berkeley.edu>.

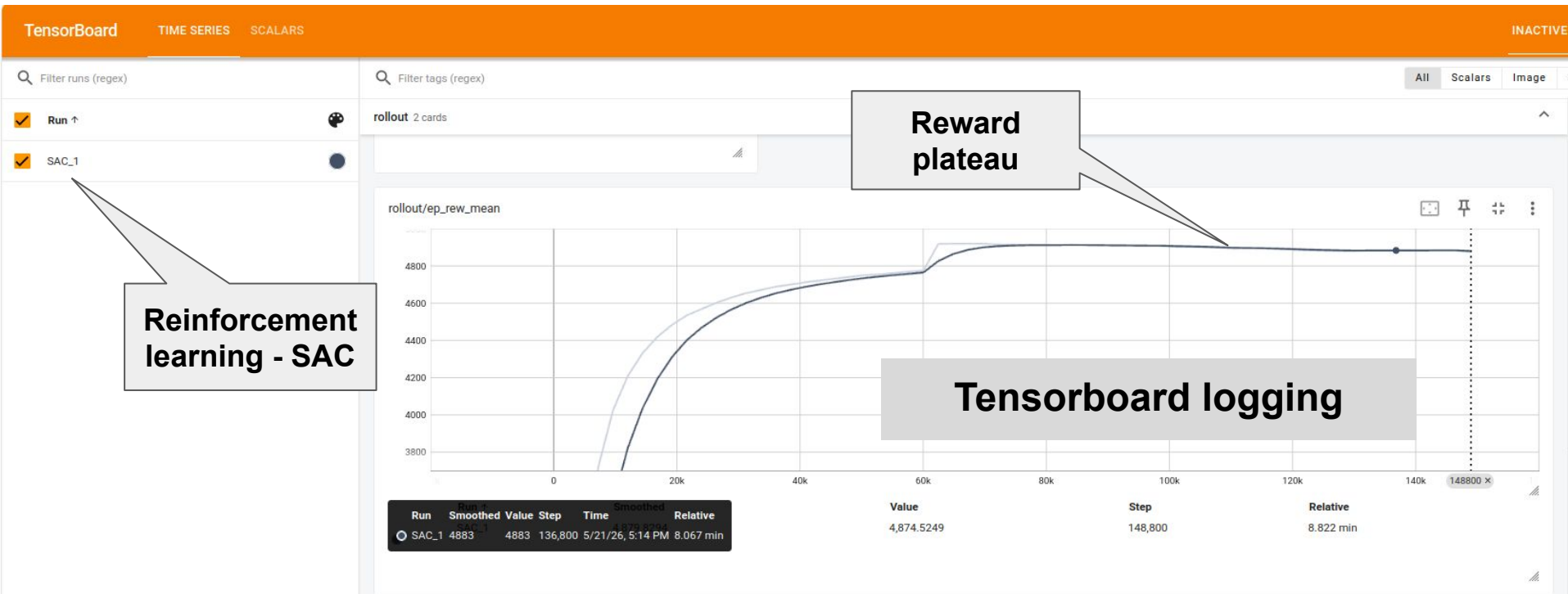
of these methods in real-world domains has been hampered by two major challenges. First, model-free deep RL methods are notoriously expensive in terms of their sample complexity. Even relatively simple tasks can require millions of steps of data collection, and complex behaviors with high-dimensional observations might need substantially more. Second, these methods are often brittle with respect to their hyperparameters: learning rates, exploration constants, and other settings must be set carefully for different problem settings to achieve good results. Both of these challenges severely limit the applicability of model-free deep RL to real-world tasks.

One cause for the poor sample efficiency of deep RL methods is on-policy learning: some of the most commonly used deep RL algorithms, such as TRPO (Schulman et al., 2015), PPO (Schulman et al., 2017b) or A3C (Mnih et al., 2016), require new samples to be collected for each gradient step. This quickly becomes extravagantly expensive, as the number of gradient steps and samples per step needed to learn an effective policy increases with task complexity. Off-policy algorithms aim to reuse past experience. This is not directly feasible with conventional policy gradient formulations, but is relatively straightforward for Q-learning based methods (Mnih et al., 2015). Unfortunately, the combination of off-policy learning and high-dimensional, nonlinear function approximation with neural networks presents a major challenge for stability and convergence (Bhatnagar et al., 2009). This challenge is further exacerbated in continuous state and action spaces, where a separate actor network is often used to perform the maximization in Q-learning. A commonly used algorithm in such settings, deep deterministic policy gradient (DDPG) (Lillicrap et al., 2015), provides for sample-efficient learning but is notoriously challenging to use due to its extreme brittleness and hyperparameter sensitivity (Duan et al., 2016; Henderson et al., 2017).

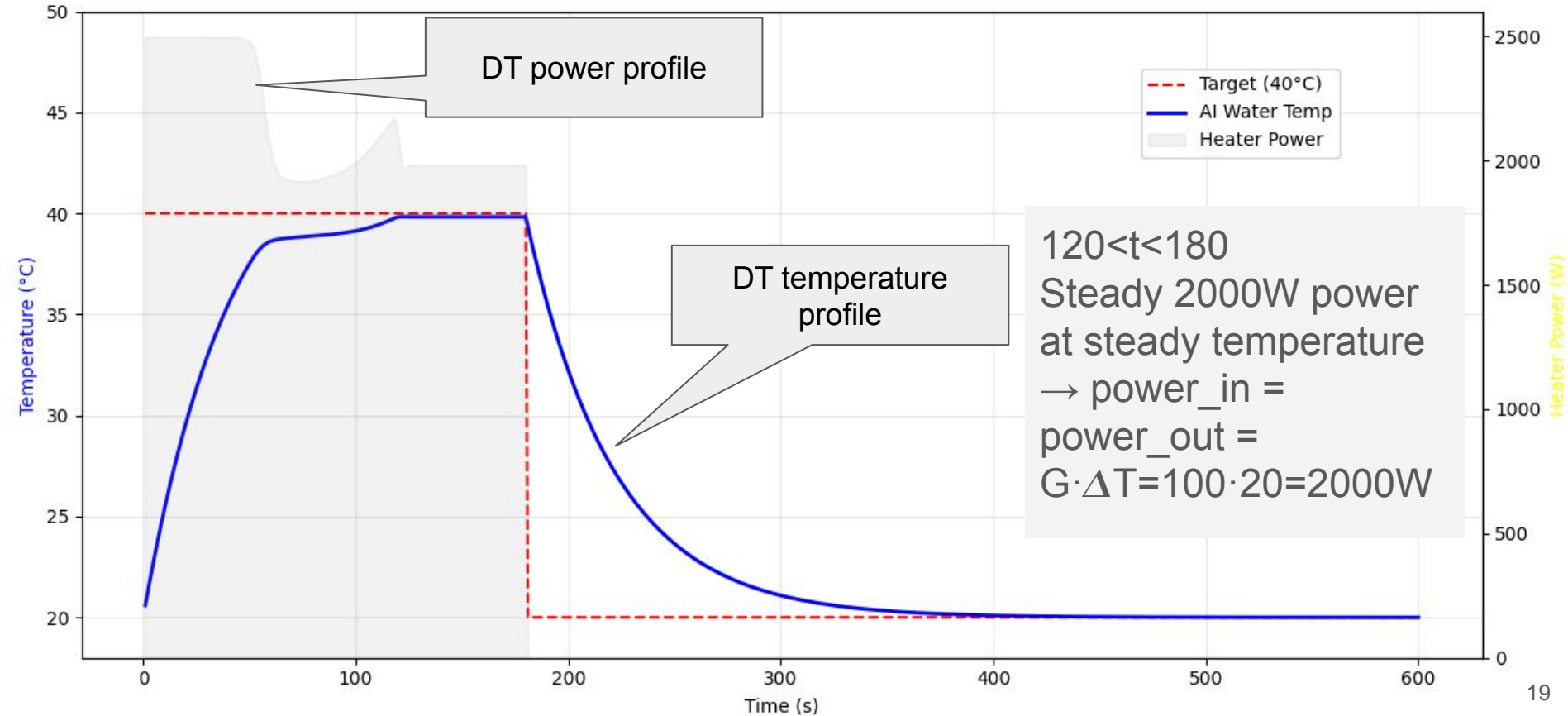
We explore how to design an efficient and stable model-free deep RL algorithm for continuous state and action spaces. To that end, we draw on the maximum entropy framework, which augments the standard maximum reward reinforcement learning objective with an entropy maximization term (Ziebart et al., 2008; Toussaint, 2009; Rawlik et al.,



Training the SAC-Based Trusted Controller



SAC Controller and Digital-Twin Response

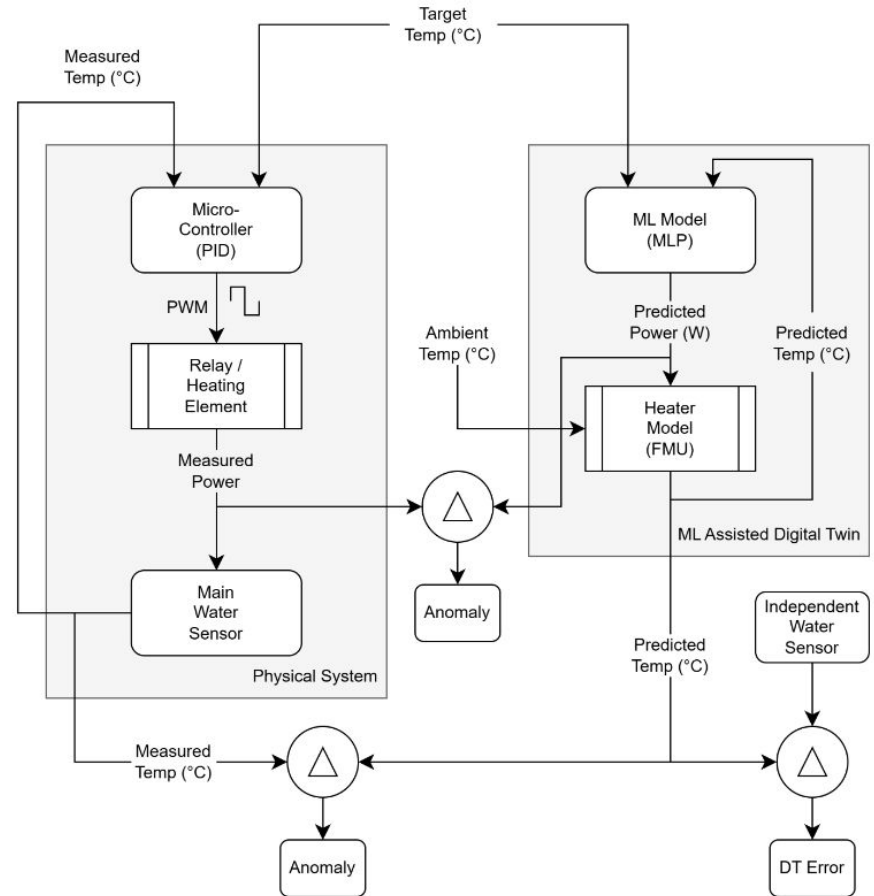


Final project. BGU 4th year engineering students project

- Ori Dagan
- Elad Elmakias

Key features

- Open Source tools
- OpenPLC
- ThingsBoard
- Real physical model. ESP32
- MLP
- Data logging



The HMI



OpenPLC

The screenshot displays the OpenPLC IDE interface. On the left, a file explorer shows the project structure for 'iot_project', including a 'Function Blocks' directory with various modules like 'clock_ladder', 'dht22_measure', and 'mqtt_data_concat_func'. Below this is a 'Library' section with categories such as 'Additional Function Blocks', 'Arduino Function Blocks', and 'Communication Blocks'.

The main workspace shows a ladder logic program for a heater control system. The program includes several function blocks and logic elements:

- RELAY_CTRL_1**: A relay control block with IN and OUT terminals.
- HEATER_1.RELAY_1.STATE**: A state variable for the heater relay.
- HEATER_1.MEASUR_ED_TEMP_1**: A measured temperature block.
- HEATER_1.TARGET_TEMP**: A target temperature block.
- HEATER_1.MEASUR_ED_POWER_1**: A measured power block.
- HEATER_1.PROFILE_D1** and **HEATER_1.PROFILE_D2**: Profile blocks for temperature and power.
- PWM_FUNC0**: A PWM function block with parameters like CYCLE, PWM_OUT, and DMS.
- BTN_CTRL_1** and **BTN_CTRL_2**: Button control blocks.
- RELAY_FORCE_ON_MODBUS**: A Modbus force on relay block.
- OR** and **AND** logic blocks for combining conditions.
- NOT** logic block for negation.

The program is connected to various I/O modules on the right, including MEASURED_TEMP, TARGET_TEMP, MEASURED_POWER, T_LOW, and T_HIGH.

A text overlay in the bottom right corner states: "OpenPLC is an open-source PLC IDE that follows the IEC 61131-3 standard, which defines the architecture and programming languages for PLCs."

At the bottom of the interface, there is a 'Console' tab and a status bar indicating 'No file selected'.

Future plans

Continuation projects at BGU and at SCE

Suggestions:2027

Advisers:

Dr. Guy Tel-zur

Project name (english)

Protecting IoT, IIoT, and ICS networks using a Digital Twin

Abstract(english)

This project focuses on the development of a Digital Twin (DT)-based monitoring system that adds an independent security layer to a conventional industrial control system. The DT models the expected physical behavior of the controlled process and continuously compares it with measurements from the real system. Significant discrepancies between the real system and the DT may indicate equipment malfunction, process anomalies, or cyberattacks. In such cases, the system will trigger an alert and display it at the Human-Machine Interface (HMI) station.

Project name (hebrew)

הגנת מערכות בקרה תעשייתית באמצעות תאום דיגיטלי

Abstract(hebrew)

מדמה את ההתנהגות הפיזית הצפויה של התהליך המבוקר ומשווה אותה באופן DT-המוסיפה שכבת אבטחה עצמאית למערכת בקרה תעשייתית קונבנציונלית. ה (DT) פרויקט זה מתמקד בפיתוח מערכת ניטור מבוססת תאום דיגיטלי עשויים להצביע על תקלה בציד, אנומליות בתהליך או מתקפות סייבר. במקרים כאלה, המערכת תפעיל התראה ותציג אותה בתחנת ממשק DT-רציף למדידות מהמערכת האמיתית. פערים משמעותיים בין המערכת האמיתית ל (HMI) אדם-מכונה.

Anomaly Detection metrics

		Predicted Values	
		Anomaly	Normal
Actual Values	Anomaly	TP	FN
	Normal	FP	TN

Confusion matrix

$$Accuracy = \frac{TN + TP}{FN + FP + TN + TP}$$

$$Precision = \frac{TP}{TP + FN}$$

$$Recall = \frac{TP}{TP + FP}$$

$$F1score = \frac{2 * Precision * Recall}{Precision + Recall}$$

Summary

Innovation

- **100% open source** tools (ThingsBoard, OpenPLC, OpenModelica, Python, local AI model)
- **AI-assisted Digital Twin** - an improved robustness and independent security layer
- Reinforcement learning implementations

Next steps

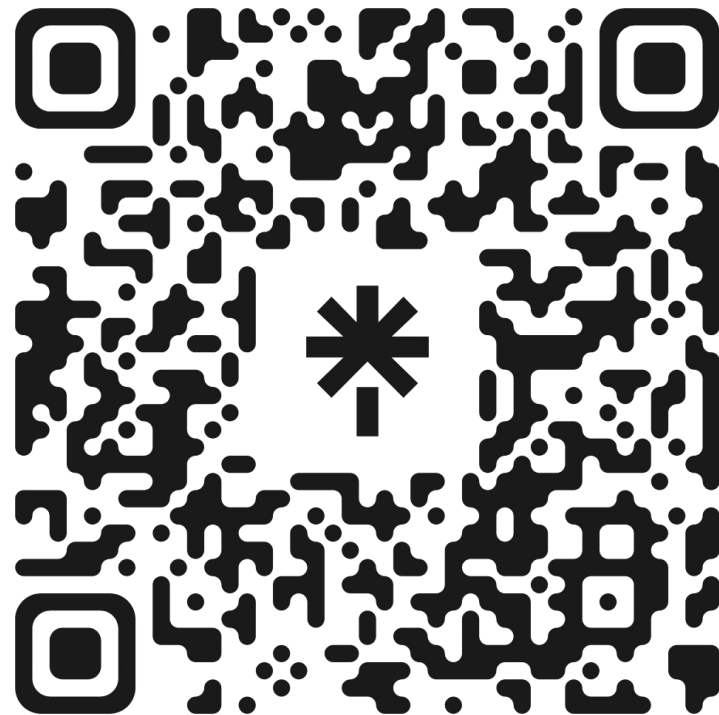
- BGU 2027 4th year ECE continuation project
- SCE 2027 CS final projects (including physical HW to be tested in the new Cyber OT laboratory).

Questions?

E-mail: g telzur@bgu.ac.il

Linkedin: [linkedin.com/in/telzur](https://www.linkedin.com/in/telzur)

That's it!



Extra Slides

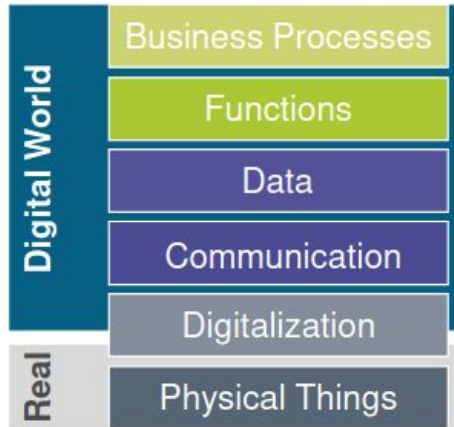
Regulation

NIST

Information technology

DIGITAL TWINS

Definitions and State of the Art
Essential Elements
Validating and Advancement
Digital Twin Standardization
Case Studies and Economics
Projects and Programs



RAMI 4.0 - Reference Architectural Model for Industrie 4.0



European Cyber Security Organisation
ECISO Technical Paper on Cybersecurity scenarios and Digital Twins

May 2023 – v1.0



ISO/IEC 30173

Edition 1.0 2023-11

INTERNATIONAL
STANDARD



Digital twin – Concepts and terminology

INTERNATIONAL
ELECTROTECHNICAL
COMMISSION

ICS 35.020

ISBN 978-2-0322-7600-0

Warning! Make sure that you obtained this publication from an authorized distributor.

INTERNATIONAL
STANDARD

ISO
23247-1

First edition
2023-10

Automation systems and
integration — Digital twin framework
for manufacturing —

Part 1:
Overview and general principles

*Systèmes d'automatisation industrielle et intégration — Cadre
technique de jumeau numérique dans un contexte de fabrication —
Partie 1: Vue d'ensemble et principes généraux*



Reference number
ISO 23247-1:2023(E)

© ISO 2023

SAC: Pros & Cons



Advantages:

- Better exploration
- More stable learning
- Robust policies
- Balanced behavior



Disadvantages:

- More complex
- Higher computation cost
- Sensitive hyperparameters
- Harder to interpret

Big Data Infrastructure for further processing

