

Thirulok Sundar Mohan Rasu

Ann Arbor, MI 48105 | +1 (248) 704-2911 | thirulok@umich.edu | thiruloksundar.github.io | github.com/Thiruloksundar |

Education

linkedin.com/in/thirulok-sundar-mohanrasu

University of Michigan – Ann Arbor

Aug 2025 – Dec 2026

M.S., Electrical and Computer Engineering

GPA: **4.0/4.0**

Relevant courses: Algorithmic Robotics, Mobile Robotics, Matrix Methods for ML & Signal Processing

Indian Institute of Technology (ISM) Dhanbad

Dec 2021 – May 2025

B.Tech., Electrical Engineering. Member: CyberLabs (Machine Learning)

GPA: **8.30/10**

Skills

Languages: Python, C++

ML & Robotics: Machine Learning, Computer Vision, Pose Estimation, Foundation Models, Imitation Learning, SLAM, State Estimation, Robot Control, Object Detection and Tracking, Optimization

Tools: PyTorch, OpenCV, ROS2, MoveIt, Nav2, Isaac Sim, MuJoCo, Gazebo, GitHub, Linux

Experience

Mowito

May 2026 – Aug 2026

Robotics Software Intern

- Enhanced the full software stack for a visual-servoing-based auto-fastening system that derives target poses from a single demonstration, enabling CAD-free robotic bolt fastening on parts moving on a conveyor; demonstrated the system to investors, customers, and integrators.
- Added a detector-free, semi-dense matching method and a real-time tracker so pose estimation stays robust through occlusion and low-texture parts, raising vision-stack throughput from **14.5 Hz** to **21 Hz** via compressed image transport.
- Built an Isaac Sim camera-mount pose optimization pipeline that renders and scores candidate camera mounts across multiple CAD parts of different sizes (**90–565 mm** max dimension) and with multi-plane bolts, using staged search over **1,000** candidates to converge on one robust mount pose given part classification (size, area to be accessed, bolt spacing).
- Tuned the MPC controller and implemented automatic weight tuning via preferential Bayesian optimization across **30** bolt positions, cutting mean settle time from **4.17s** to **2.67s** (**31% reduction**) on FANUC arm
- Extended working conveyor speed from **60** to **100 mm/s** on UR10e arm, cutting cycle time up to **5.9×** (**21.3s** to **3.6s**) while holding **98%** success across **150** runs.
- Ported the fastening stack to a FANUC CRX10iA-L arm, resolving driver-level jerk instability to sustain a continuous **159**-cycle run, and authored a state machine for multi-plane fastening (bolts on multiple faces) reaching **95%** success across **200** runs.

University of Michigan - Ann Arbor

Jan 2026 – Present

Graduate Research Assistant

- Evaluated LeWM (JEPA-based world model) with free-energy and coding-rate regularizers as replacements for the SIGReg anti-collapse regularizer, improving hard-protocol (100 step planning) accuracy on a two-room navigation dataset from **16%** to **75%** with the coding-rate variant.
- Showed free-energy loss (a Gaussianizing regularizer) improves reconstruction and classification across CNNs, MLPs, and autoencoders; built a joint gaussianizing classifier and published freegaussianizer as an open-source package.

Carnegie Mellon University

Jan 2025 – Sep 2025

Research Intern

- Engineered an end-to-end data generation and labeling pipeline to curate a **1M+** sample dataset of real/synthetic audio pairs stored in structured JSON format; fine-tuned the CLAP audio-language model achieving **97%** accuracy on deepfake speech detection; paper submitted to Interspeech 2026 (under review) (*preprint*)

Selected Projects

SO-101 Robot Arm – Pick-and-Place with Imitation Learning

2026

Collected teleoperated demonstrations via HuggingFace LeRobot and trained an ACT (Action Chunking with Transformers) policy for screwdriver pick-and-place from a single front-facing camera; deployed training on an AWS EC2 GPU instance.

Semantic LiDAR-SLAM with Language-Guided Navigation (GitHub)

2026

Integrated Cartographer SLAM with Grounding DINO in ROS2, Gazebo; DINO populates a persistent semantic map enabling natural-language commands to trigger Nav2 navigation goals – zero-shot object-goal navigation without environment-specific training.

VLA Model Fine-tuning for Robot Manipulation (GitHub)

2025

Fine-tuned OpenVLA (**7B** vision-language-action model) on LIBERO manipulation benchmarks via LoRA; deployed on UMich HPC (NVIDIA A40). Training loss: **19.08** → **1.16** over **3** epochs.