

Michael Doyle

AI RESEARCHER & ENGINEER

An AI researcher with a passion for language, open source, and applying research to important problems.

michaeldoyle1994@gmail.com

(775) 450-6522

San Diego, California US

LinkedIn [@michaeldoyleml](#) **GitHub** [@doyled-it](#)

EXPERIENCE

The MITRE Corporation

Senior AI Research Engineer

Jun 2018 - Present

Led multiple ML projects. Researched, developed, tested, trained, and deployed machine learning models and applications.

- Led research project on LLMs, IT Modernization, and code understanding
- Led research project on training object detectors for neuromorphic cameras
- Maintained internal GPU servers and services on OpenShift and Linux servers
- Developed and open sourced fire simulator to be used in RL training for wildfire mitigation
- Developed and open sourced LLM IT modernization library
- Developed agentic LLM backend prototype for government sponsor application
- Researched adversarial object detection and classification
- Researched adversarial vision for classical depth
- Researched adversarial text for Machine Translation
- Directly deployed trained models on government sponsor systems
- Engaged in government sponsor test events in the field

Space Micro Inc.

Firmware Engineering Intern

Apr 2016 - Apr 2018

Wrote, simulated, synthesized, and implemented FPGA code in VHDL using ModelSim, Synplify Pro, Vivado, and Libero Designer. Designed PCBs in KiCAD and tested hardware.

SKILLS

MACHINE LEARNING NLP, LLMs, Computer Vision, Acoustics, RL, Simulation, Adversarial

LANGUAGES Python, Bash, SQL, LaTeX, Vue.js, JavaScript, CSS, MATLAB, VHDL

LIBRARIES NumPy, SciPy, PyTorch, TensorFlow, FastAPI, Typer, LangChain, Chroma, HuggingFace

TECHNOLOGIES Git, Docker, OpenShift, GitLab CI, GitHub Actions

EDUCATION

University of San Diego

BS/BA Dual Degrees — Electrical Engineering

Minors in Mathematics and Physics, Magna Cum Laude

Aug 2013 - Apr 2018

PROJECTS

SimFire

Aug 2019 - Present

Lead Developer, Maintainer

A wildfire simulator written in Python and meant for Reinforcement Learning research for wildfire mitigation.

Janus LLM

Jun 2023 - Present

Lead Developer, Maintainer

A library for LLM IT modernization, using LLMs, RAG, and intelligent chunking.

SimHarness

Aug 2019 - Present

Developer

A reinforcement learning harness meant to be used with SimFire for wildfire mitigation research.

PUBLICATIONS

Testing the Effect of Code Documentation on Large Language Model Code Understanding

North American Association for Computational Linguistics (NAACL)

May 2024 | [Link](#)

William Macke, **Michael Doyle**

Empirically analyzes how code documentation quality impacts the code generation and understanding capabilities of Large Language Models (LLMs). It reveals that incorrect documentation significantly hinders LLMs' code comprehension, while incomplete or missing documentation has no significant impact.

The vulnerability of UAVs: an adversarial machine learning perspective

SPIE

Mar 2021 | [Link](#)

Michael Doyle, Josh Harguess, Keith Manville, Mikel Rodriguez

Proposes a methodology to evaluate the vulnerability of unmanned aerial vehicles (UAVs) to adversarial machine learning attacks by analyzing potential attack vectors at each stage of UAV operation.

Reinforcement Learning for Wildfire Mitigation in Simulated Disaster Environments

Neural Information Processing Systems (NeurIPS)

Oct 2023 | [Link](#)

Alexander Tapley, Marissa Dotter, **Michael Doyle**, Aidan Fennelly, Dhanuj Gandikota, Savanna Smith, Michael Threet, Tim Welsh

Presents a reinforcement learning approach to wildfire mitigation in simulated disaster environments. We release two software libraries, SimFire and SimHarness, to facilitate future research in this area.

Practical Attacks on Machine Translation using Paraphrase

Association for Machine Translation in the Americas (AMTA)

Aug 2022 | [Link](#)

Elizabeth Merkhofer, John Henderson, Abigail Gertner, **Michael Doyle**, Lily Wong

Investigated the vulnerability of machine translation systems to adversarial attacks constructed with limited information. A novel attack method was proposed that generates perturbations using paraphrases and evaluates their impact on meaning preservation and translation degradation across various language pairs and systems.

LANGUAGES

English: *Native speaker*

Spanish: *Intermediate*

German: *Beginner*