

Jihyung Park

jihyung803@gmail.com | U.S. Permanent Resident | [Website](#) | [Google Scholar](#)

Experience

IntelliChoice

Pro Bono AI Engineer

Dallas, TX

June 2026 – Present

- Developed full-stack AI applications by building backend services, multimodal workflows, real-time features, and automated end-to-end tests.

Urban Information Lab, University of Texas at Austin

Austin, TX

Graduate Research Assistant | Advisor: Dr. Junfeng Jiao

Jan 2025 – May 2026

- Led AI safety and emergency-response research by developing evaluation workflows and guardrail systems, while deploying and maintaining AI services across Docker, AWS EC2, and Linux environments.

Turing

Palo Alto (Remote), CAA

LLM Trainer

Sep 2025 – Dec 2025

- Designed and evaluated multi-turn LLM dialogue workflows, focusing on tool-use reasoning, function-calling accuracy, and natural conversational flow.

Human Language Technology Research Institute, University of Texas at Dallas

Richardson, TX

Undergraduate Research Assistant | Advisor: Dr. Vincent Ng

Dec 2023 – Jan 2025

- Conducted multimodal language-understanding research by building datasets and evaluation pipelines, and analyzing model behavior on complex visual-textual tasks.

Project

AI Tutoring Assistant

Intellichoice

- Built a multimodal tutoring platform with **Next.js/TypeScript** frontend, **FastAPI** services that analyze uploaded student work and generate personalized instructional support for tutors.
- Implemented **PostgreSQL/pgvector** retrieval, **SSE** real-time updates, and **Playwright** end-to-end tests.

Self-Reinforcing Counterfactual Reasoning for Pragmatic Language Understanding | [EMNLP](#)

UT NLP

- Developed a self-supervised training method to overcome LLMs' reliance on literal interpretations, enabling them to generate, evaluate, and learn from counterfactual reasoning traces.
- Fine-tuned Qwen3-8B/14B using **SFT** and **GRPO** with pre-rollout sampling and self-judge reward.
- Improved pragmatic reasoning accuracy by up to **9.46%** across four pragmatic benchmarks.

Anticipatory Hidden-State Monitoring for Implicit Harmful Dialogue | [arXiv](#)

UIL

- Developed a same-pass safety monitor using hidden-state probing to detect implicit harmful continuations before unsafe content becomes visible.
- Outperformed external guard models while reducing latency overhead from **79.40% to 2.34%**.

SafeMate (w/ City of Austin)

UIL

- Developed an AI agent that transforms public-safety documents into actionable emergency guidance.
- Built a **hierarchical RAG** using **SBERT embedding**, **UMAP**, **GMM clustering**, and recursive summarization.
- Built a **FastAPI** backend with **LangGraph** and **MCP**, then deployed it to **AWS EC2** using **Docker**.

MemeInterpret: Towards an All-in-One Dataset for Meme Understanding | [EMNLP](#)

HLTRI

- Addressed the difficulty of detecting implicit hate conveyed through image-text interactions and cultural context by building a 6,810-example multimodal corpus.
- Fine-tuned LLaVA and HateCLIPper, then combined their predictions through weighted joint inference, achieving **0.890 AUROC** and **state-of-the-art 0.800 accuracy**.

Education

University of Texas at Austin

GPA: 3.89

Master of Science in Computer Science

Aug 2024 - May 2026

University of Texas at Dallas

GPA: 3.94

Bachelor of Science in Computer Science

Aug 2020 - May 2024

Skill

Programming & Frameworks

Python, Java, TypeScript/JavaScript, React, Next.js, C++, Bash/Shell, SQL

AI & Machine Learning

PyTorch, Transformers, RAG, LangChain, LangGraph, LLM

Cloud & DevOps

AWS EC2, Docker, Linux, Git, GitHub Actions