

Research Open House

Institute for Information Sciences

Tuesday, September 22, 2026 · 3:00–5:00 PM · Nichols Hall, University of Kansas

Machine Learning-Based Multi-Objective Optimization for HPC Workload Scheduling

Presented by Kyrian Adimora

This dissertation develops a framework that uses graph neural networks and reinforcement learning to help supercomputers schedule jobs more efficiently, balancing performance, energy use, and reliability all at once instead of optimizing for just one at a time.

Scheduling Input-Sensitive Workloads with Runtime Prediction

Presented by Sai Rithvik Gundla

This study examines whether improving runtime prediction accuracy actually improves HPC scheduling performance, using MRI-based brain segmentation jobs as a case study. The results show that how an estimate is calibrated matters more than how accurate it is: a moderately accurate model, once calibrated, can match or beat near-ideal scheduling performance.

When Memory Does the Math: Hardware-Algorithm Co-Design for Accelerating AI Workloads Using Analog In-Memory Computing

Presented by Ahsan Mojahidul

Modern AI uses an enormous amount of electricity, largely because computers waste energy constantly moving data back and forth between memory and processors. This project explores hardware-algorithm co-design of “analog in-memory computing,” a technique that lets memory chips do the math directly where data is stored to slash energy use. The project has developed frameworks to accurately simulate analog AI hardware non-idealities and make these ultra-efficient chips practical for real-world AI deployment.

AI for Networks and Networks for AI: From Learning-Enabled Edge Systems to Real-World Testbeds

Presented by Sravan Chintareddy

This work sits at the intersection of AI and networked systems, addressing two key questions: how AI can make wireless and edge networks more intelligent, and how networks can efficiently support distributed AI tasks. It develops learning-driven solutions spanning reinforcement learning, federated and distributed learning, edge inference, spectrum management, and integrated sensing and communications, complemented by real-world experimental systems including NextG testbeds and networked UAV platforms.

The Effects of Quantization on the J-Space

Presented by Daniel Neugent

This work examines quantized models to assess how quantization affects their J-Space and the fitting of the Jacobian Lens.

How Do We Know When AI Is Doing a Good Job? AI Evaluation Across Indigenous Language and Classroom Contexts

Presented by Chamisa Edmo

AI systems are traditionally evaluated using benchmarking and metrics that necessarily encode assumptions about what counts as correct, relevant, or meaningful. Drawing on case studies in Indigenous language and classroom contexts, this research examines where conventional AI evaluation approaches capture performance, where they lose important context, and how evaluation might better account for the environments in which AI systems are ultimately implemented.

Application of AI for QA/QC in Geologic Datasets

Presented by Bailey S.

StraboSpot is home to many different tools that facilitate the collection, analysis, and sharing of geologic data. To improve confidence in the quality of the data and promote collaboration, the project implements Quality Assurance and Quality Control (QAQC). The QAQC system is implemented as an AI model that parses the structure of StraboField projects to identify areas where improvements can be made.

Can One Photo Be Enough? Improving Face Recognition for Biometric Watchlists

Presented by Mashfiq Rizvee

This research introduces a learning-guided probabilistic framework for open-set biometric watchlist screening that combines ArcFace embeddings, PCA-ITQ binary coding, and multiresolution hierarchical Bloom filters to support few-shot, cross-pose, and dynamically updateable authentication.

AI-Enabled Privacy-Preserving Electrocardiogram Analysis

Presented by Fairuz Shadmani Shishir

This work explores AI-based methods for analyzing electrocardiograms (ECGs) while preserving patient privacy. It develops privacy-preserving deep learning approaches that reduce the ability to infer sensitive patient information while maintaining the clinical utility of ECG-based predictions.

AI-Enabled Atomic Resolution Modeling of Whole Cells at Long Timescales

Presented by Kamila Yunas

This work combines physics-based docking with AI-predicted protein binding conformations (AlphaFold3) to simulate whole cells at atomic resolution over biologically relevant timescales, extending simulation trajectories by orders of magnitude beyond what traditional atomic-resolution methods can reach. The project includes a case study of crystallin proteins and a whole-cell Mycoplasma model to demonstrate these methods. It also trains a neural network to predict protein motion over time as a potentially faster surrogate, with autoregressive rollouts generally yielding diffusion rates similar to those from Monte Carlo simulations.

Close Personal Relationships with AI

Presented by Oluwaseun Sanwoolu

This paper explores aspects of close personal relationships that AI cannot replicate, such as agent-relative obligations, and argues that people who believe they are in a close personal relationship with AI are simply mistaken.

JIT Compilation Policy for ML Compilers

Presented by Shriraj Vaidya

This work explores compilation policies for machine learning compilers, with a focus on improving the tradeoff between compilation overhead and runtime performance.

A Tool-Agnostic Framework for Recovering Array Bounds in Stripped Binaries

Presented by Ashish Adhikari

Recovering stack-allocated array bounds from stripped binaries remains a fundamental challenge in reverse engineering. While explicit metadata is stripped, compiled code retains sparse but complementary evidence across pointer arithmetic, loop bounds, memory accesses, and interprocedural patterns. This paper proposes a tool-agnostic framework that aggregates this normalized static analysis evidence and uses a fine-tuned large language model (LLM) to predict stack array presence, element sizes, and total byte capacities. Using a Ghidra-based reference frontend and a hybrid LLM-and-heuristics backend, the evaluation on a deduplicated dataset of 20,000 C binaries achieves 92% and 84% exact total-size accuracy at the -O0 and -O2 optimization levels, respectively.

Socially Engaging roBOTics (SEBO) Lab

Presented by Sarah Sebo

The SEBO lab explores how a robot's social behavior can be designed to best help people in the places they live, learn, and work — tutoring children at school, assisting around the home, and supporting teams in making decisions.

DeepPicar: A Low-Cost Autonomous Car Platform for ML Education

Presented by Tyler Oswald

DeepPicar provides a low-cost, vision-based autonomous driving platform for student ML education. This platform uses an end-to-end convolutional neural network to autonomously drive 1/16th-scale RC cars.

TinyLidarNet: 2D LiDAR-Based End-to-End Deep Learning Model for F1TENTH Autonomous Racing

Presented by Ethan Gao & Farrell Joswara

A 2D LiDAR-based, end-to-end deep neural network system that directly processes LiDAR input and steers the car for F1TENTH competitions. It won third place in an international F1TENTH racing competition.

When Alignment Breaks: On Inference Integrity in Text-to-Image Models

Presented by Becky Lin

This work studies whether text-to-image models preserve inference integrity — whether their output distributions remain stable under benign prompt changes and consistent with the model’s underlying knowledge and empirical reality. It audits five major text-to-image systems using more than 25,000 generated images and introduces group-based prompting as a black-box stress test for revealing distributional behavior that single-image evaluation can hide. Using occupational gender representation as a case study, the work finds that seemingly benign changes — switching between individual and group prompts, reordering demographic instructions, or adding demographic attributes — can substantially change the resulting distributions. The results show that current bias-mitigation and alignment mechanisms can be brittle, motivating integrity-focused auditing of generative AI systems.

Seeing Is No Longer Believing: A Security and Forensic Evaluation of Generative Video Platforms

Presented by Becky Lin

This evaluation examines six state-of-the-art generative video platforms using 1,100 videos generated across text-to-video and image-to-video settings, covering user-facing safeguards, identity preservation, voice similarity, lip-sync quality, watermarking, forensic detectability, and human perception. The results show substantial differences in how platforms handle identity, consent, copyright, and provenance, while current generators can preserve facial identity and produce highly realistic synchronized speech. Existing forensic tools also show important blind spots across generators and settings, and a user study finds that people still struggle to reliably distinguish generated videos from authentic content. The findings highlight a growing gap between the realism of generative video and the governance and forensic mechanisms meant to control its misuse.

PhantomSeal: Proactive Deepfake Defense with Identity/Context Protection and Forensic Tracing

Presented by Liangqin Ren

PhantomSeal is a proactive defense designed to protect users’ facial images from face-swapping deepfake attacks. It simultaneously protects both identity and contextual information while enabling forensic tracing of manipulated content through a novel identity-cloaking mechanism.

An LLM-Enabled End-to-End Attack Pipeline against Speech-Based Machine Translation

Presented by Junyi Zhao

Speech translation apps and devices let two people talk when neither speaks the other's language — which also means neither can verify what the translation says. This work shows that an attacker can alter a single word inside the audio of a voice message, causing the speech recognizer to transcribe something else and quietly change the meaning of the translation the listener receives. In a user study, listeners distinguished attacked clips from clean ones only 53% of the time — no better than chance.

Benchmarking the Benchmarks: Evaluating the Automated Safety Benchmarks for Small Language Models

Presented by Nyam Shaik

This is a large-scale study of the effectiveness and robustness of automated LLM safety benchmarks, evaluating five widely used benchmark suites across 26 open-source SLMs under a unified judging rubric that scores each response as harmful, safe, or ambiguous/irrelevant. Across the benchmarks, ambiguous judgments dominate and correlate with prompt complexity and model architecture, indicating that LLM-centric safety benchmarks are insufficient as standalone evidence for SLM safety assessment.

The Adversarial AI-Art: Understanding, Generation, Detection, and Benchmarking

Presented by Yuying Li

This work studies the security risks of adversarial AI-generated images and introduces the ARIA dataset, containing over 140,000 real and AI-generated images across multiple real-world scenarios. It also evaluates how well human users and existing AI-image detectors can distinguish AI-generated images from real ones, finding that both still face significant challenges.