

Alvi Ishmam

alvi@vt.edu | 5404492513 | LinkedIn | alvi.me | Google Scholar | Research Gate

EDUCATION

VIRGINIA TECH

PHD IN COMPUTER SCIENCE

January 2021 - April 2026

M.Sc. IN COMPUTER SCIENCE

January 2021 - November 2023

GPA: 3.92 / 4.0

SKILLS

LANGUAGES

Python • C# • Java • C++

• SQL • Assembly • $\text{\LaTeX}$

TOOLS AND FRAMEWORKS

HTML • CSS • Angular • React •

Weblogic • Oracle • Apache Tomcat

• RESTful API • Spring MVC • Spring

Boot • Android • Oracle ADF • .Net

Core • Scikitlearn • pytorch

MACHINE LEARNING FRAMEWORKS

Pytorch • Keras • Tensorflow

PUBLICATION

(check out google scholar for complete list)

- Semantic Shield: Defending Vision-Language Models Against Backdooring and Poisoning via Fine-grained Knowledge Alignment [CVPR'24] (Main)
- Learning Universal Adversarial Perturbations for Multi Image Tasks via Pretrained Models [AAAI'26] (Main)
- M3D: MultiModal MultiDocument Fine-Grained Inconsistency Detection [EMNLP'24](Main)
- JourneyBench: A Challenging One-Stop Vision-Language Understanding Benchmark of Generated Images [Neurips'24] (Main)
- Model Immunization by Trapping Harmful Finetuning CVPR 2026 Main Highlight

EXPERIENCE

ARGONNE NATIONAL LAB | POSTDOC RESEARCH SCIENTIST

May 2026 – present | Lemont, IL, USA

- Develop multimodal foundation model for critical infrastructure.
- Design a ML system for National Integrated Drought Information System

FUTUREWEI TECHNOLOGIES | RESEARCH INTERN

May 2025 – Aug 2025 | Framingham, MA, USA

- Develop self evolving RAG system with a focus on advanced storage system/data abstraction process.
- Design reference free evaluation method for hallucination mitigation for enterprise. **Under review, COLM workshop 2026**

GE HEALTHCARE | AI/ML PHD INTERN

May 2024 – Aug 2024 | San Ramon, CA, USA

- Generate the largest medical image-text pairs **2M** from modalities CT, MR, US.
- Design the largest medical **CLIP** style model achieving SOTA performance.

PROJECTS

DEFENDING VISION-LANGUAGE MODEL AGAINST IMAGE/TEXT POISONING ATTACK | [CVPR2024]

- The goal is to design a robust multimodal model to defend against adversarial attack. (Multimodal foundation models)
- Proposed a robust finetuning strategy that can defend vision-language model (CLIP).

LEARNING UNIVERSAL ADVERSARIAL PERTURBATIONS FOR MULTI IMAGE TASKS VIA PRETRAINED MODELS | [AAAI2026]

- We propose the first adversarial attack for multi-image multi-model models by learning universal adversarial perturbations.
- Our method increases the attack success rate by 20% in various multi-image tasks across MLLMs compared to SOTA.

MULTIMODAL DISINFORMATION DETECTOR USING FINE GRAINED VISUAL ENTAILMENT | [EMNLP2024]

- Visual Entailment is a reasoning task where the logical relationship between text, image, and video is predicted.
- The goal of the project is to determine disinformation in multimodal news articles by representing the semantic inconsistency in knowledge elements in text, images/videos.

ROBUSTNESS ANALYSIS OF BROWSER-BASED WEB AGENTS Under review, Neurips 2026

- Proposed adaptive vulnerabilities in multimodal web agents.
- Threat models (white-box and black-box) incorporating stealth, controllability, and delayed triggers to systematically probe agent decision-making under adversarial conditions.

AUDIO ATTACKS AGAINST VIDEO AUDIO MODELS [ACL2026]

- Imperceptible audio attack to video audio models.
- Optimized audio signal to attack the video audio model.

MULTIMODAL EVENT EXTRACTION AND COREFERENCE RESOLUTION [EMNLP2026 (Main)]