

Brandon Dominique

Boston, MA

<https://bdominique.github.io/> - github.com/bdominique - dominique.b@northeastern.edu - (908) 420-0182

PhD candidate at Northeastern University researching uncertainty quantification, subgroup discovery and fairness in machine learning, with an application to medical imaging (skin lesion classification). Skilled in uncertainty quantification methods including calibration and conformal prediction, clustering, and custom loss function design.

EDUCATION

Northeastern University, Boston, MA Ongoing

Ph.D. Computer Engineering

Advisor: Jennifer Dy

Graduate Student Ambassador, Computer Engineering Department and GEM Fellowship

Northeastern University, Boston, MA May 2023

M.S. Computer Engineering

Graduate Student Ambassador, Computer Engineering Department and GEM Fellowship

New Jersey Institute of Technology, Newark, NJ May 2020

B.S. Computer Engineering, *Magna Cum Laude*

PAPERS

Dominique B., Lam P., Torop M., Kurtansky N., Weber J., Kose K., Rotemberg V., Dy J. (2026). Discovering Bias by Learning Semantically Coherent Calibration Groups. Under review.

Dominique B., Lam P., Dy J. (2026). Uncertainty-Aware Group-Conditional Conformal Prediction. Under review.

Dominique B., Lam P., Kurtansky N., Weber J., Kose K., Rotemberg V., Dy J. (2025). On the Role of Calibration in Benchmarking Algorithmic Fairness for Skin Cancer Detection. In the Journal of Machine Learning for Biomedical Imaging (MELBA).

Dominique B., Piorkowski D., Nagireddy M., Baldini I. (2024). Prompt Templates: A Methodology for Improving Manual Red Teaming Performance. In 1st HEAL Workshop at CHI Conference on Human Factors in Computing Systems.

Dominique B., Piorkowski D., El-Maghraoui K., Herger M. (2023). FactSheets for Hardware-Aware AI Models: A Case Study of Analog In Memory Computing AI Models. In IEEE International Conference on Software Services Engineering. (Best Student Paper Award)

SELECTED EXPERIENCE

IBM, Research Intern

May 2021-Aug 2023

Advisors: Lorraine Herger, Kaoutar El Maghraoui, David Piorkowski

- Conducted semi-structured interviews with participants across multiple stakeholder roles to identify transparency needs for hardware-aware AI models
- Designed and built a web-based interface to present model transparency documentation across multiple analog AI models, giving stakeholders role-specific access to training methodology, hardware performance, and model architecture details
- Designed and evaluated structured prompt guides to help non-expert users audit LLMs for biased outputs; ran a 29-participant study showing guide-assisted participants attributed 77% of successful audit prompts to the guides, versus 50% for single-guide users

FELLOWSHIPS, AWARDS, AND TALKS

ICML 2026 Gold Reviewer

April 2026

Invited Talk, "Post-Hoc Methods for Detecting and Correcting Subgroup Bias in Medical AI Systems",
University College London, Alan Turing Institute

April 2026

GEM Fellowship

August 2021

LSAMP STARS Fellowship

August 2020