

PhD Statement of Purpose Example

Field: Machine Learning / Fairness in AI | Program: PhD in Computer Science

***What this essay does:** Opens with the research problem rather than the applicant's personal history. States the specific research question and preliminary finding by the end of the first paragraph. Names the advisor and explains the fit with precision. Quantifies contributions where possible. Distinguishes the doctoral project from existing work in a single, specific sentence. This is the structure CS programs expect.*

Hiring algorithms deployed by large employers routinely perform differently across demographic groups even when protected attributes are excluded from the feature set. The dominant technical response has been post-hoc fairness correction, applied after a model is trained, which treats disparate impact as an output to be adjusted rather than a structural property to be understood. My research takes a different approach: I want to identify the conditions under which fairness-aware regularization during training produces models that generalize their fairness guarantees to out-of-distribution data, and I want to understand when and why it fails to do so. My undergraduate research with Professor Lin Chen at Carnegie Mellon has produced a preliminary finding, currently under review at NeurIPS, that equalized odds constraints applied during training degrade out-of-distribution fairness performance by an average of 18 percentage points across four benchmark hiring datasets when the test distribution shifts by more than 0.3 in total variation distance from the training distribution.

I am applying to the PhD program in Computer Science at MIT to work with Professor Sonia Patel, whose work on distributional robustness and fairness is the most technically rigorous investigation of this problem I have encountered. Specifically, I want to extend her group's results on worst-case group performance guarantees to the sequential decision settings that characterize actual hiring pipelines, where model outputs at one stage affect the population available for evaluation at the next stage. The existing literature addresses single-stage fairness extensively but has not systematically characterized how fairness violations compound across pipeline stages, which is the gap my dissertation will address.

My technical background for this project includes proficiency in PyTorch and JAX for model development, strong foundations in statistical learning theory from Professor Chen's graduate-level reading group, and direct experience with the COMPAS, Adult Income, and LinkedIn Talent Solutions benchmark datasets used in the fairness literature. I also completed a research internship at Microsoft Research in the summer of 2025, working with Dr. Aditya Kumar's Fairness, Accountability, and Transparency group on a project examining demographic parity violations in large language model-based resume screening tools. That internship produced a second co-authored paper, currently in

preparation, and confirmed that the distributional shift problem I identified in my undergraduate work appears in production systems at scale.

The reason I am pursuing a PhD rather than continuing in industry research is the problem's timescale. The question of why fairness constraints fail under distribution shift cannot be answered in the eighteen-month cycles that govern applied research programs. It requires sustained theoretical development, empirical validation across multiple deployment domains, and the freedom to follow results that complicate the original hypothesis. I expect that freedom at MIT, in a program that treats theoretical rigor and empirical grounding as equally important, and with an advisor whose published work demonstrates the patience to pursue hard problems at the pace they require.

My goal after completing the doctorate is to build a research program at a research university that works directly with regulatory bodies and civil society organizations on the development of technically grounded standards for algorithmic fairness in high-stakes decisions. The gap between the technical fairness literature and the regulatory frameworks currently being developed in the EU AI Act and the proposed US Algorithmic Accountability Act is large and consequential. I intend to spend my career working to close it.