

When I sit at the piano, my music flows naturally as my imagination leads my hands, carried by years of technical practice and the pursuit of a Bachelor's in Piano Performance. Music also inspired my approach to CS — starting from an intuitive picture, working through technical grinds, and arriving at a precisely formulated argument.

I have broad interests in theoretical computer science and have approached it through the **foundations of modern Machine Learning (ML)** and **Algorithmic Game Theory (AGT)**. In both, I aim to understand the limits and capabilities of complex computational systems.

- **(Foundations of Modern ML.)** *For algorithmic tasks, when and why do Transformers prioritize learning statistical heuristics over generalizable algorithms, and how can we systematically induce robust reasoning?*
- **(Algorithmic Game Theory.)** *In collective decision-making, simple tournament-based voting rules are theoretically incapable of achieving optimal social welfare on their own. How can we bridge this optimality gap with minimally added overhead?*

Foundations of Modern ML. While Transformer models demonstrate remarkable capabilities, they often prioritize learning statistical shortcuts over generalizable algorithms. Advised by Professors Robin Jia and Vatsal Sharan, I investigated this failure on the graph connectivity problem and proved that the model's failure to generalize is not inherent to the architecture, but rather a deterministic consequence of the training distribution relative to the model's capacity. Crucially, this implies that by **aligning the training distribution with the model's theoretical capacity**, we can steer models toward more **robust out-of-distribution generalization**.

I developed the theoretical foundations of this project. I first proved a tight, non-asymptotic capacity bound showing that an L -layer model has the exact circuit complexity to solve connectivity for graphs with diameter up to 3^L . This theoretical threshold allowed us to model the learning process as a competition between two latent channels: a robust "algorithmic" channel that implements matrix powering and a shortcut "heuristic" channel that relies on a degree-counting shortcut. I then identified a sharp phase transition that depends solely on the data composition: "within-capacity" graphs drive the gradient descent trajectory toward the algorithmic solution, while "beyond-capacity" graphs promote the heuristic. This finding offers a prescriptive strategy for out-of-distribution generalization: restricting training data to graphs within the model's capacity paradoxically encourages the model to learn the generalizable algorithm by suppressing the heuristic channel, a prediction my collaborators empirically validated.

Algorithmic Game Theory. Under the metric distortion framework, the design of voting systems involves a fundamental tension between simplicity and efficiency. Tournament-based voting rules are practically appealing because they minimize cognitive load, requiring voters to compare only two options at a time. However, this simplicity comes at a theoretical cost: standard tournament rules are provably suboptimal compared to the general deterministic optimum.

Advised by Professor Kamesh Munagala, I sought to bridge this gap while preserving the simplicity of tournament rules. We proposed “Deliberation via Matching,” a protocol where voters with opposing views may engage in pairwise discussions. We proved that by augmenting tournament rules with minimal pairwise deliberations, we can break the aforementioned barrier. Conceptually, this establishes that the cognitive simplicity of **tournament rules** does not require sacrificing social welfare: with slight overhead, they are just as **powerful as general deterministic social choice rules**.

In more technical detail: I led the theoretical analysis to resolve the analytical intractability that had limited prior work. The distortion objective for deliberative processes is inherently non-linear and non-convex, forcing previous studies to essentially rely on black-box numerical optimizers. I identified that the objective is *supermodular* and *convex* with respect to the metric variables. These were the crucial keys to tractability. Supermodularity implied that worst-case instances must follow a specific monotonic “structure,” while convexity allowed me to apply Jensen’s inequality locally to collapse a continuum of voters into a discrete set. This reformulation converted the intractable program into a bilinear relaxation, from which the optimal solutions easily followed via vertex enumerations and linear programming. While the proof is terse in its algebra, each step closely follows an intuitive story, and this precisely echoes my belief that exciting research can start from and be driven by clean intuitions.

At Princeton, I hope to contribute to the TCS Group and work with Professors **Matt Weinberg** and **Mark Braverman**. While I previously worked on social choice, **I am now eager to explore the vast landscape of AGT**. In reading recent work at Princeton TCS, I often revisit one theme:

When people are strategic and information is limited, the hard (and practical) part is making the goal itself stable and meaningful, not just optimizing it.

I see two complementary perspectives to pursue this goal at Princeton.

First, Professor Weinberg’s work starts from a concrete way a rule can be exploited *in practice*, then builds a model that treats that exploit as *part of the mechanism design problem*. This includes, for example, treating miners’ off-chain influence as a robustness issue and showing that even widely deployed transaction-fee rules can fail it (EC’24), as well as proving separations for what truthful auction protocols can achieve under communication limits once we move beyond the simplest bidder settings (FOCS’24). I find this approach interesting because it turns real strategic threats into definitions one can test, and results that actually guide what we should design.

Professor Braverman’s work supplies a complementary approach: when hard constraints make the usual objective unrealistic, he first defines a comparison standard that respects the constraint, then *proves guarantees against it*. This is reflected in his work describing how stability shapes welfare in large two-sided matching markets (EC’23), and in his recent preprint that proposes a constraint-aware benchmark for online learning with budget-balancing constraints. This perspective is also appealing because it highlights and clarifies what it even means to “do well” when constraints are the whole point.