

UNIVERSITY OF CALIFORNIA,
IRVINE

Stable, Robust, and Fair Solutions in Combinatorial Games

DISSERTATION

submitted in partial satisfaction of the requirements
for the degree of

DOCTOR OF PHILOSOPHY

in Computer Science

by

Rohith Reddy Gangam

Dissertation Committee:
Distinguished Professor Vijay V. Vazirani, Chair
Distinguished Professor David Eppstein
Assistant Professor Ioannis Panageas

PREVIEW

Parts of Chapter 2 and Chapter 4 © 2026 International Foundation for Autonomous Agents and Multiagent Systems (IFAAMAS).

Chapter 3 © 2024 International Foundation for Autonomous Agents and Multiagent Systems (IFAAMAS).

All other materials © 2026 Rohith Reddy Gangam.

DEDICATION

To my mom, dad, and brother.

PREVIEW

TABLE OF CONTENTS

	Page
LIST OF FIGURES	vi
LIST OF ALGORITHMS	vii
ACKNOWLEDGMENTS	viii
VITA	x
ABSTRACT OF THE DISSERTATION	xii
1 Introduction	1
2 Robust Stable Matchings	6
2.1 Introduction	6
2.1.1 Our contributions	8
2.2 Related Work	11
2.3 Model and Preliminaries	13
2.3.1 The stable matching problem and the robust stable matching problem	13
2.3.2 The lattice of stable matchings	14
2.3.3 Birkhoff's Theorem and rotations	15
2.3.4 The rotation poset	16
2.3.5 Sublattice and semi-sublattice	17
2.3.6 Linear programming formulation	17
2.4 Results	19
2.4.1 Birkhoff's Theorem on Sublattices	21
2.4.2 Structure of Robust Stable Matchings	39
2.4.3 Algorithms to Compute Worker- and Firm-Optimal Robust Stable Matchings	54
2.4.4 Integrality of the Polytope	62
2.4.5 A Robust Stable Matching Algorithm for the General Case	66
Appendix 2.A Proof of Birkhoff's Theorem using Stable Matching Lattices . . .	72
2.A.1 $L(\Pi')$ is isomorphic to a sublattice of $L(\Pi)$	72
2.A.2 Sublattice $\mathcal{L}'$ is generated by a compression Π' of Π	73
Appendix 2.B Algorithms when One Side Changes Preferences	78

2.B.1	Worker and firm optimal matchings	80
2.B.2	Proof of Correctness of Algorithm 2.9	83
2.B.3	Proof of Correctness of Algorithm 2.10	84
3	Robust Popular Matchings	88
3.1	Introduction	88
3.2	Related Work	90
3.3	Preliminaries	93
3.4	Results	97
3.4.1	Perturbations of One Agent	97
3.4.2	Perturbation by Four Downshifts	103
3.4.3	Unpopular Agents	110
3.4.4	Reduced Availability	111
3.4.5	Robust Popular Matchings in Related Models	114
3.5	Conclusion	115
	Appendix 3.A Hybrid Instances for Two Agents	117
	Appendix 3.B Limitations of LP-based Approaches	120
4	Equitable Core Imputations for Max-Flow, MST and b-Matching Games	124
4.1	Introduction	124
4.1.1	Our contributions	126
4.2	Related Work	127
4.3	Preliminaries	132
4.3.1	The games and their unfair core imputations	133
4.4	Leximin Optimal Solution of a Linear Program	136
4.5	Max-flow Game	141
4.5.1	LP formulation of the max-flow problem	141
4.5.2	Owen set for the max-flow game	143
4.5.3	Computing the leximin Owen set imputation	147
4.6	MST Game and Min-cost Branching Game	156
4.6.1	Owen set for the min-cost branching game	157
4.6.2	Finding leximin Owen set imputation using linear programming . .	161
4.7	b -Matching Game	166
4.7.1	Owen set imputations for the b -matching game	167
4.7.2	Combinatorial algorithm to compute the leximin Owen set imputation	168
	Appendix 4.A Deferred Proofs	171
	4.A.1 Proofs from Section 4.5	171
	4.A.2 Proofs from Section 4.6	179
	4.A.3 Proofs from Section 4.7	184
	Appendix 4.B Finding leximin Owen set imputation of the max-flow game using linear programming	188
	Appendix 4.C Min-cost branching	190

4.C.1	Going from maximin Owen set imputation to leximin Owen set imputation	190
4.C.2	Finding leximax Owen set imputation	192
4.C.3	More concise series of LPs for leximin Owen set imputation	193
Appendix 4.D	Omitted details from b -matching game	197
4.D.1	Finding leximin Owen set imputation using linear programming . .	197
4.D.2	Omitted definitions	198
4.D.3	Combinatorial algorithms to compute leximin Owen set imputation	200

Bibliography	207
---------------------	------------

PREVIEW

LIST OF FIGURES

	Page
2.1 Summary of our results on robust stable matchings.	10
2.2 Finding an integral stable matching from a fractional solution.	17
2.3 Examples of compressions.	22
2.4 Examples of compressions and their Hasse diagrams.	24
2.5 Examples of: (a) canonical path, and (b) bouquet.	28
2.6 An example in which $\mathcal{M}_A \setminus \mathcal{M}_B$ is not a sublattice of $\mathcal{L}_A$	43
2.7 Example for multiple agents changing preferences.	46
2.8 Lattices $\mathcal{L}_A$ and $\mathcal{L}_B$ for example in Figure 2.7	47
2.9 Example where $\mathcal{M}_A \cap \mathcal{M}_B$ is not a sublattice of $\mathcal{L}_A$	48
2.10 $\mathcal{M}_A \setminus \mathcal{M}_B$ is not a semi-sublattice of $\mathcal{L}_A$	50
2.11 $\mathcal{M}_A \setminus \mathcal{M}_B$ is not a semi-sublattice of $\mathcal{L}_A$, even when (A, B) is of type $(1, 1)$	55
2.12 Example where the robust stable matching polytope is not integral.	65
3.1 Instance $\mathcal{I}_A$ in Example 3.1.	96
3.2 Key gadget of the reduced instances in the proof of Theorem 3.3.	105
4.1 The graph for Example 4.1.	144
4.2 Illustration of Claim 4.7	152
4.3 Example to show that Owen set is a strict subset of the core.	158
4.4 Example for Theorem 4.12	164
4.5 The graph for Example 4.2.	168

LIST OF ALGORITHMS

	Page
2.1 Algorithm for finding a bouquet.	35
2.2 Subroutine for finding the next tail.	36
2.3 Subroutine for finding a flower.	37
2.4 Workers proposing Deferred Acceptance Algorithm, by [49] (from [37]) . .	56
2.5 Algorithm to find the worker-optimal robust stable matching.	57
2.6 Algorithm to find the firm-optimal robust stable matching.	58
2.7 An XP-time algorithm to find a robust stable matching for small (p, q) . . .	67
3.1 ROBUSTPOPULARMATCHING for changes of one agent	101
4.1 Algorithm to decide if an imputation is in the Owen set of a max-flow game	146
4.2 Algorithm to find the leximin Owen set imputation in a max-flow game. .	151

ACKNOWLEDGMENTS

I would like to express my deepest gratitude to my advisor and defense committee chair, Professor Vijay Vazirani. His guidance, insight, and mentorship have been central to my development as a researcher. His enthusiasm for research and constant encouragement helped me navigate both the challenges and rewards of the Ph.D. journey, and I will always cherish the time I spent working under his supervision.

I am also grateful to my other defense committee members, Professors David Eppstein and Ioannis Panageas. I especially enjoyed their courses on graph algorithms and algorithmic game theory, respectively, which piqued my interest in theoretical computer science.

Along with Vijay, I would like to thank my co-authors, including those whose papers did not appear in this thesis: Martin Bullinger, Naveen Garg, Tung Mai, Nitya Raju, Parnian Shahkar, and Shayan Taherijam, for the wonderful collaborations. Working with them has been an enriching experience, and I learned a great deal from our discussions and joint problem-solving.

I would like to thank my mother, Vijitha, my father, Prasad, and my brother, Ritesh, for their unwavering support, love, sacrifices, and constant encouragement throughout this journey. Their support and belief in me has been a source of strength at every stage of my Ph.D.

I am thankful to the CS theory group at UCI for providing a stimulating and welcoming research environment. In particular, I would like to thank the regular tea-time attendees — Hadi Khodabandeh, Ofek Gila, Shayan Taherijam, and Abraham Mathew Illickan — for many enjoyable and insightful discussions.

I am also grateful for the friends I made at UCI, far too many to name individually. Special thanks go to Sukriti Kapur, Sourav Majumdar, Praveen Ranganath Prabhakar, Rahul Bhope, Prakhar Srivastava, Aishwarya Borate, Sridipta Ghatak, and Nitish Nagesh, who made this journey both memorable and fulfilling.

Finally, I would like to thank my longtime friends, whose presence has extended well beyond this Ph.D. journey. In particular, I am grateful to Sai Chetan Chinthakindi, Jagadeesh Boddeda, Mohan Sai Teja Dantam, Vishwanadh Rapolu, and Aditya Vardhan Varre for their enduring friendship, constant availability, and the many wonderful memories we share.

This thesis is based on a collection of previously published papers:

- Chapter 2 is based on two papers co-authored with Tung Mai, Nitya Raju, and Vijay V. Vazirani: “A Structural and Algorithmic Study of Stable Matching Lattices of “Nearby” Instances, with Applications” [51], which appeared in FSTTCS 2022, for which copyright is retained by the authors; and “Stable Matching: Dealing

with Changes in Preferences” [54], which will appear in AAMAS 2026 and will be published by the International Foundation for Autonomous Agents and Multiagent Systems (IFAAMAS). IFAAMAS holds the copyright for the latter paper and has granted permission for its reuse within this dissertation.

- Chapter 3 is based on “Robust Popular Matchings” [24], co-authored with Martin Bullinger and Parnian Shahkar, which appeared in AAMAS 2024. IFAAMAS holds the copyright for this paper and has granted permission for its reuse within this dissertation. An extended version of this work [25], co-authored with Martin Bullinger, Gergely Csáji, and Parnian Shahkar, is currently under journal review.
- Chapter 4 is based on “Equitable Core Imputations for Max-Flow, MST, and b -Matching Games” [53], co-authored with Naveen Garg, Parnian Shahkar, and Vijay V. Vazirani. The paper will appear in AAMAS 2026. IFAAMAS holds the copyright for this paper and has granted permission for its reuse within this dissertation.

This work was supported in part by NSF grant CCF-2230414.

Vijay V. Vazirani directed and supervised research which forms the basis for the dissertation.

VITA

Rohith Reddy Gangam

EDUCATION

Doctor of Philosophy in Computer Science University of California, Irvine	2026 <i>Irvine, California, US</i>
Master of Science in Computer Science University of California, Irvine	2025 <i>Irvine, California, US</i>
Bachelor of Technology in Computer Science and Engineering Indian Institute of Technology Bombay	2018 <i>Mumbai, India</i>

WORK EXPERIENCE

Graduate Research Assistant University of California, Irvine	2020–2025 <i>Irvine, California, US</i>
Teaching Assistant University of California, Irvine	2020–2025 <i>Irvine, California, US</i>
Research Assistant Intern IIIT Hyderabad	2019–2020 <i>Hyderabad, India</i>
Analyst Goldman Sachs Pvt. Ltd.	2018–2019 <i>Bengaluru, India</i>
Summer Intern Microsoft	2017 <i>Hyderabad, India</i>
Summer Intern Indus OS	2016 <i>Mumbai, India</i>

REFEREED CONFERENCE PUBLICATIONS

- Stable Matching: Dealing with Changes in Preferences** 2026
Rohith Reddy Gangam, Tung Mai, Nitya Raju, and Vijay V. Vazirani
International Conference on Autonomous Agents and Multi-Agent Systems (AAMAS)
- Equitable Core Imputations for Max-Flow, MST and b -Matching Games** 2026
Rohith Reddy Gangam, Naveen Garg, Parnian Shahkar, and Vijay V. Vazirani
International Conference on Autonomous Agents and Multi-Agent Systems (AAMAS)
- Fair Rent Division: New Budget and Rent Constraints** 2025
Rohith Reddy Gangam, Shayan Taherijam, and Vijay V. Vazirani
IARCS Annual Conference on Foundations of Software Technology and Theoretical Computer Science (FSTTCS)
- Robust Popular Matchings** 2024
Martin Bullinger, Rohith Reddy Gangam, and Parnian Shahkar.
International Conference on Autonomous Agents and Multi-Agent Systems (AAMAS)
- A Structural and Algorithmic Study of Stable Matching Lattices of “Nearby” Instances, with Applications.** 2022
Rohith Reddy Gangam, Tung Mai, Nitya Raju, and Vijay V. Vazirani
IARCS Annual Conference on Foundations of Software Technology and Theoretical Computer Science (FSTTCS)

ABSTRACT OF THE DISSERTATION

Stable, Robust, and Fair Solutions in Combinatorial Games

By

Rohith Reddy Gangam

Doctor of Philosophy in Computer Science

University of California, Irvine, 2026

Distinguished Professor Vijay V. Vazirani, Chair

In this dissertation, we study *stability*, *robustness*, and *fairness* in two foundational frameworks of algorithmic game theory: matchings under preferences and cooperative games. Classical results in both settings focus on identifying stable outcomes and establishing their structural and algorithmic properties. However, stability alone does not address robustness to changes in the input or fairness at the level of individual agents. Motivated by these gaps, we investigate how stability-based solution concepts behave under perturbations and how fairness can be incorporated while retaining stability.

We first study robustness in two-sided matching markets. We introduce *robust stable matchings*, which are matchings that remain stable across multiple instances defined on the same set of agents, capturing stability under changes to preferences. We characterize when the structural, algorithmic, and geometric properties of stable matchings persist, parametrized by the number of agents changing preferences. We then extend this robustness analysis to *popular matchings*, a solution concept that generalizes stability.

We then focus on fairness in cooperative games derived from fundamental combinatorial optimization problems: max-flow, minimum spanning tree, and *b*-matching games. Although stable outcomes in these games can be computed efficiently, they can be highly unfair to individual agents. We study *lexicographically fair core imputations* and show that

computing such outcomes is computationally intractable in general. To address this, we extend the notion of the *Owen set* to these games, which leads to tractable alternatives and polynomial-time algorithms for computing leximin and leximax fair core imputations that preserve stability while providing strong fairness guarantees.

PREVIEW

Chapter 1

Introduction

Matching under preferences and cooperative games have long been central topics in algorithmic game theory and economics. In two-sided matching markets, the seminal 1962 work of Gale and Shapley [49] introduced the Deferred Acceptance algorithm for computing *stable matchings*, initiating a rich line of work on their structural, geometric, and game-theoretic properties. In this framework, *stability* forbids *blocking pairs* - pairs of agents who mutually prefer each other and deviate from a given matching. In cooperative games, an analogous notion of stability is captured by the *core*, which ensures that no coalition of agents has an incentive to secede by guaranteeing each coalition at least as much profit, or no more cost, than it can achieve on its own. In both settings, *stability* serves as a fundamental concept and has given rise to deep structural insights as well as efficient algorithmic solutions.

Despite its central role, *stability* alone does not address two additional concerns that arise naturally in practice: *robustness* and *fairness*. In matching markets, preferences are often subject to change due to evolving information, strategic behavior, or external uncertainty, and an outcome that is stable for one instance may fail to remain stable under

even small perturbations of the input. At the same time, stability places no explicit restrictions on how outcomes are distributed among individual agents. In cooperative games, *core allocations* may be highly unfair while still preventing all coalitional deviations. These considerations motivate the study of solution concepts that extend classical stability by explicitly incorporating robustness to input changes and fairness at the level of individual agents. Addressing these concerns forms a central theme of this thesis.

We begin in [Chapter 2](#) by studying *robustness* in the classical two-sided stable matching model. A stable matching instance consists of two sets of agents, workers and firms, each endowed with strict preferences over all agents on the opposite side. A matching is said to be *stable* if it admits no *blocking pair*, that is, no worker–firm pair that would both strictly prefer to be matched to each other rather than to their assigned partners. To model robustness, we consider changes to agents’ preferences and ask how stability behaves under such perturbations. Formally, we study pairs of instances defined on the same set of agents and focus on matchings that are stable in both instances simultaneously. We refer to such matchings as *robust stable matchings*. This framework captures a broad range of perturbations, from localized changes in the preference list of a single agent to arbitrary preference changes by multiple agents on either side of the market.

The primary goal of [Chapter 2](#) is to understand how the structural, algorithmic, and geometric theory of stable matchings extends to this robust setting, and to identify precise conditions under which these properties are preserved. Structurally, we study the set of robust stable matchings as the intersection of the sets of stable matchings of two instances and ask when this intersection retains the finite distributive lattice structure characteristic of the classical model. We show that the preservation of this structure depends critically on the number and location of agents whose preferences differ across the two instances. When a lattice structure is present, we further investigate whether it admits a compact Birkhoff partial order and whether this representation can be computed

efficiently, thereby enabling the enumeration of all robust stable matchings with polynomial delay. Algorithmically, we address the existence of robust stable matchings and the computation of extremal outcomes, including worker-optimal and firm-optimal robust stable matchings. From a geometric perspective, we analyze the *robust stable matching polytope* and characterize exactly when it is integral. Taken together, these results delineate sharp boundaries, parameterized by the number of agents whose preferences change, between regimes in which robustness preserves the classical structure of stable matchings and regimes in which this structure breaks down.

In [Chapter 3](#), we move beyond stable matchings and study robustness for the broader notion of *popular matchings*. The underlying model remains the same, namely two-sided markets with strict preferences over the opposite side. Rather than ruling out blocking pairs, popularity compares matchings via a majority vote among agents. Given two matchings, each agent votes for the matching that yields a more preferred outcome, and a matching is called *popular* if it does not lose this election against any other matching. Introduced by Gärdenfors [55], popularity admits a natural social-choice interpretation and strictly generalizes stability. In particular, every stable matching is popular, while the converse need not hold. Moreover, stable matchings are known to be the minimum-cardinality popular matchings, highlighting popularity as a relaxation of stability that may admit matchings of larger cardinality.

[Chapter 3](#) extends the analysis of robustness from stable matchings to this more general solution concept. We study matchings that are popular across multiple instances and investigate the computational complexity of deciding the existence of such *robust popular matchings*. Unlike stable matchings, popular matchings do not admit a known lattice structure or rotation-based characterization, and consequently the techniques developed for the stable matching setting do not directly apply. We identify sharp tractability boundaries: robustness can be handled efficiently when preference changes are confined

to a single agent, whereas allowing even limited structured changes by a small number of agents leads to computational intractability. We further study robustness under changes in availability, where edges may become unavailable across instances, and obtain a complexity dichotomy based on preference completeness. Together, these results show that robustness for popularity exhibits fundamentally different behavior from robustness for stability and requires new algorithmic techniques beyond the classical lattice-based framework.

In [Chapter 4](#), we turn to *fairness* considerations in cooperative games. We study three cooperative games derived from fundamental combinatorial optimization problems: the *max-flow game*, the *minimum spanning tree (MST) game*, and the *b-matching game*. In the max-flow game, agents correspond to the edges of a directed network, each with a given capacity, and the value generated by a coalition is the maximum amount of flow that can be routed from a designated source to a sink using only the edges in the coalition. In the MST game, agents correspond to vertices of a weighted undirected graph, and the cost incurred by a coalition is defined as the minimum cost of a network that connects its members to a designated root. The *b-matching game* generalizes the assignment game and arises from weighted bipartite matching, where edges may be selected multiple times subject to capacity constraints on the incident vertices. In this setting, agents on the two sides form mutually beneficial partnerships, and the value of a coalition is the maximum total weight achievable by matchings restricted to the coalition and the vertex capacities. In all three games, stability is captured by the *core*, which consists of profit or cost allocations under which no coalition can improve upon its outcome by seceding from the grand coalition.

Although core imputations in all three games can be found in polynomial time and guarantee stability, they do not, in general, ensure fairness at the level of individual agents. Core allocations may be highly unfair, assigning a disproportionate share of

profit or cost to a small subset of agents while remaining stable. This motivates the search for outcomes that are both stable and equitable. We focus on *lexicographic fairness* criteria, specifically *leximin* and *leximax*, which seek to equalize outcomes by, respectively, improving the welfare of the worst-off agents or limiting the advantage of the best-off agents as much as possible. We show that computing such equitable core imputations is computationally intractable for the max-flow and MST games and encounters similar barriers for b -matching games. To overcome this difficulty, we extend the notion of the *Owen set*, a well-structured and tractable subset of the core derived from optimal dual solutions of the underlying linear programs, to these settings. Restricting attention to the Owen set allows us to circumvent these hardness results and design polynomial-time algorithms for computing leximin and leximax fair imputations that preserve stability and provide fairness guarantees.

Chapter 2

Robust Stable Matchings

2.1 Introduction

Matchings under preferences is a central topic in market design, originating with the seminal 1962 paper by Gale and Shapley [49], which introduced the stable matching problem and launched the broader study of matching-based market design. The classical stable matching model assumes that agents report fixed preference rankings, yet real markets rarely satisfy such invariance: preferences may shift as new information arrives, external conditions change, or strategic behavior—including misreports or collusion—arises. These factors naturally lead to the question of how matchings behave under perturbations of the input. In this work, we investigate matchings that remain stable across multiple preference profiles—specifically, matchings that are stable in both a base instance A and a perturbed instance B , which we refer to as *robust stable matchings*. Our study provides a systematic analysis of their structural, algorithmic, and geometric properties under increasingly general models of perturbations.

A central structural feature of the stable matching problem is that the set of stable matchings of an instance forms a finite distributive lattice, ordered by the preferences of one side [74, 64]. This lattice admits a compact representation via rotations, which correspond to minimal transformations between stable matchings and underpin many fundamental algorithmic and structural results. We investigate how this lattice structure behaves under preference changes by comparing the lattices associated with two instances defined on the same set of agents. In particular, we study whether the intersection of their stable matchings retains any of the lattice structure that characterizes the classical setting.

We begin with the simplest nontrivial perturbation, namely a single shift in the preference list of one agent. Mai and Vazirani [77] establish a complete structural characterization of the resulting space of stable matchings. We vastly generalize this and consider arbitrary permutations by a single agent and subsequently extend the analysis to instances in which multiple workers and firms simultaneously modify their preferences. These models progressively weaken the lattice structure of the robust stable matchings and expose the precise conditions under which it is preserved or lost.

Alongside these structural results, we investigate the algorithmic and geometric properties of robust stable matchings. We address the complexity of deciding whether the intersection of the stable matchings of two instances is nonempty, the efficient computation of worker-optimal and firm-optimal matchings within this intersection, and the existence of succinct Birkhoff representations supporting efficient enumeration of all robust stable matchings. From a geometric perspective, we analyze the fractional robust stable matching polytope and determine exactly when it is integral. We show that the breakdown of integrality coincides with the loss of lattice structure, yielding a sharp and unifying threshold for both structural and geometric properties.

Preliminary versions of the results in this paper appeared in three earlier works [77, 51, 54], each addressing robustness under progressively more general models of preference perturbations. The work in [77] studied robustness under a single upward shift in one agent’s preference list, establishing lattice structure and polynomial-time algorithms. This was extended in [51] to arbitrary permutations by a single agent, together with efficient computation of the associated Birkhoff partial order. The most general setting, in which multiple agents on both sides may arbitrarily change their preferences, was considered in [54], which identified sharp thresholds for lattice structure and polytope integrality and developed parameterized algorithms for the general case.

2.1.1 Our contributions

Throughout, we consider stable matching instances on n workers and n firms, each with strict and complete preference lists over the opposite side. For an instance A , let $\mathcal{M}_A$ denote the set of its stable matchings and $\mathcal{L}_A$ the finite distributive lattice of these matchings under the usual dominance order (see [Section 2.3.2](#) for details). We begin with the simplest change, in which a new instance B is obtained from A by a single upward shift in the preference list of one agent. In this setting, a complete description of how the shift operation splits the lattice of stable matchings: both $\mathcal{M}_A \cap \mathcal{M}_B$ and $\mathcal{M}_A \setminus \mathcal{M}_B$ form sublattices of $\mathcal{L}_A$, is already given in [77].

We extend this to a more general model in which a single agent, whether a worker or a firm, may arbitrarily permute its preference list between instances A and B . In this setting, while $\mathcal{M}_A \cap \mathcal{M}_B$ remains a sublattice, the set $\mathcal{M}_A \setminus \mathcal{M}_B$ is no longer guaranteed to form a sublattice of $\mathcal{L}_A$. Nevertheless, some structure persists: it forms a semi-sublattice of $\mathcal{L}_A$. This key property allows us to identify and succinctly characterize all matchings that remain stable across both instances by efficiently characterizing the Birkhoff partial order on rotations (see [Section 2.3.3](#) for definitions).

A central theme in these results is the structural relationship between the lattices of two instances. This leads naturally to the following question, whose importance is underscored by the fact that stable matching lattices are *universal* among all finite distributive lattices: for every finite distributive lattice $\mathcal{L}$, there exists a stable matching instance A whose lattice $\mathcal{L}_A$ is isomorphic to $\mathcal{L}$. Consequently, an affirmative answer to the question below would yield a fundamental new property for the entire class.

Question 2.1. *Is $\mathcal{M}_A \cap \mathcal{M}_B$ always a sublattice of $\mathcal{L}_A$ and $\mathcal{L}_B$ under arbitrary preference changes?*

We show that, in general, the answer is “no.” Consider a robust stable matching instance (A, B) in which p of the n workers and q of the n firms change their preferences from A to B . In this regime, we identify a sharp structural frontier: $\mathcal{M}_A \cap \mathcal{M}_B$ forms a sublattice of both $\mathcal{L}_A$ and $\mathcal{L}_B$ precisely when $p \leq 1$ or $q \leq 1$, and we provide an example showing that this property fails beginning at $(p, q) = (2, 2)$.

In addition to these structural results, we address several fundamental computational and geometric questions concerning robust stable matchings:

1. Can we efficiently decide whether $\mathcal{M}_A \cap \mathcal{M}_B \neq \emptyset$?
2. If $\mathcal{M}_A \cap \mathcal{M}_B$ forms a lattice, can the worker-optimal and firm-optimal robust stable matchings (see [Section 2.3.2](#) for definitions) be efficiently computed?
3. Does $\mathcal{M}_A \cap \mathcal{M}_B$ admit a succinct Birkhoff partial order and support efficient enumeration?
4. Is the fractional robust stable matching polytope integral, enabling linear programming (LP) based algorithms?

For all settings in which the sublattice property holds, we also provide a Deferred Acceptance-based algorithm to compute the worker-optimal and firm-optimal matchings in the robust lattice. We show that the fractional robust stable matching polytope remains integral when $p \leq 1$ or $q \leq 1$. For $(p, q) = (2, 2)$, we provide an example in which the polytope contains fractional vertices. This yields a sharp threshold that mirrors exactly the boundary at which lattice structure breaks down, and rules out LP-based approaches in the fully general case. To complement this, we give an XP algorithm running in $O(n^{p+q+2})$ time that decides whether a robust stable matching exists, constructs one when it does, and enumerates all such matchings with the same delay bound. A summary of our results is provided in [Figure 2.1](#).

(p, q)	Computation (P?)	Structure (Lattice?)	Geometry (Integral?)
(0,0)	P [49]	Yes [69]	Yes [103]
(0,1)	P [Thm. 2.23]	Yes [Thm. 2.4]	Yes [Thm. 2.24]
(0,n)	P [Thm. 2.23]	Yes [Thm. 2.4]	Yes [Thm. 2.24]
(1,1)	P [Thm. 2.23]	Yes [Thm. 2.19]	Yes [Thm. 2.24]
(1,n)	P [Thm. 2.23]	Yes [Thm. 2.19]	Yes [Thm. 2.24]
(2,2)	P [Thm. 2.27]	No [Thm. 2.18]	No [Thm. 2.25]
(p, q)	XP: $O(n^{p+q+2})$ [Thm. 2.27]	-	-
(n, n)	NP-Complete [82]	No [Thm. 2.18]	No [Thm. 2.25]

Figure 2.1: Summary of our results on robust stable matchings. p workers and q firms permute their preference lists from instance A to B .

The “Computation” column deals with the decision problem $\mathcal{M}_A \cap \mathcal{M}_B \neq \emptyset$.

The “Structure” column answers if $\mathcal{M}_A \cap \mathcal{M}_B$ is a sublattice of $\mathcal{L}_A$ and $\mathcal{L}_B$.

The “Geometry” column answers if the robust fractional stable matching polytope is integral.

Finding the Birkhoff partial order differs between the $(0, 1)$ and $(0, n)$ cases, since $\mathcal{M}_A \setminus \mathcal{M}_B$ forms a semi-sublattice in the former but not in the latter. The problem of explicitly constructing the Birkhoff partial order for the $(1, 1)$ and $(1, n)$ cases remains open; however, we show that polynomial-time solvability of one case implies polynomial-time solvability of the other.