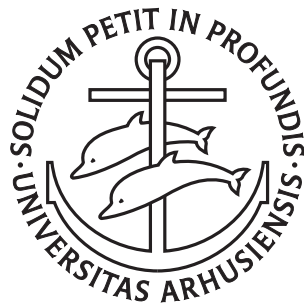

Optimization and Learning in Voting

Karl Fehrs

PhD Dissertation



Department of Computer Science
Aarhus University
Denmark

Optimization and Learning in Voting

PhD dissertation by Karl Fehrs

Published by [Aarhus University at Open Books](#)

Date of publication: July 7, 2026

Date of PhD defense: April 11, 2025

ISBN: 978-87-7507-597-3

DOI: <https://doi.org/10.7146/vsnyjm19>

© Karl Fehrs 2026

This is an Open Access publication licensed via [Creative Commons Attribution 4.0 International \(CC BY 4.0\)](#). This license enables reusers to distribute, remix, adapt, and build upon the material in any medium or format, so long as attribution is given to the creator. The license allows for commercial use.

Citation for published version (APA):

Fehrs, K. (2026). *Optimization and Learning in Voting*. Aarhus University.
<https://doi.org/10.7146/vsnyjm19>

Optimization and Learning in Voting

A Dissertation
Presented to the Faculty of Natural Sciences
of Aarhus University
in Partial Fulfillment of the Requirements
for the PhD Degree

by
Karl Fehrs
January 31, 2025

Abstract

The modern computational approach to social choice theory has greatly expanded our knowledge of the power and limitations of *voting rules*. Today, a large variety of collective decision making processes with vastly different objectives are treated as voting applications. Beyond the classical setting of political elections, voting rules are used for the allocation of public funds, for rating tasks, and for selecting representative citizen assemblies, to name only a few modern applications.

A common approach is to view voting rules as simple optimization algorithms that seek to aggregate individual preferences in a way that maximizes the collective welfare. We, as computer scientists, may then aim to leverage our insights from the study of approximation algorithms and computational complexity in the context of social choice. However, the voting setting comes with its own set of challenges which arise from an inherent lack of *information*. The voters may be unable to express their preferences exactly, for example, in numerical terms. Instead, a voter typically only provides a rough sketch of her preferences in the form of a ranking of the alternatives or by reporting those alternatives which she approves of. On the other hand, it may not always be clear to the designer of voting rules which combination of desirable properties the outcomes must satisfy in order to be considered optimal.

A central aspect of our work is to study the benefits that a limited amount of additional information can provide in the design of voting rules. This information could consist of an exact numerical quantification of the utility (or cost) that a voter has for a small subset of the alternatives. We explore whether such enhancements to the capabilities of voting rules can significantly improve their performance (measured in terms of their *distortion*). In combination with this targeted access to the voters' exact preferences, we also consider stochastic preferences which provide aggregate information to a voting rule such as a voter's average utility for the alternatives.

Another type of information that we consider is *data* about ideal outcomes of a voting application. How much data in the form of sampled preference profiles labeled by the desired outcomes is necessary and sufficient to *learn* a voting rule that is consistent with these samples and—hopefully—performs well outside of this training set? In the so-called probably approximately correct (PAC) learning model, we study this notion of sample complexity as well as the computational complexity of a related learning task for certain classes of multiwinner voting rules.

Resumé

Den moderne computerbaserede tilgang til social valgteori har i høj grad udvidet vores forståelse af fordele og ulemper ved *stemmeregler*. I dag betragtes en lang række kollektiver beslutningsprocesser med vidt forskellige mål som anvendelser af stemmesystemer. Udover det klassiske område inden for politiske valg anvendes stemmeregler til allokering af offentlige midler, vurderingsopgaver og udvælgelse af repræsentative borgerforsamlinger, for blot at nævne nogle få moderne anvendelser.

En almindelig tilgang er at betragte stemmeregler som simple optimeringsalgoritmer, der sammenfatter individuelle præferencer på en måde, der maksimerer den kollektive velfærd. Som dataloger kan vi derfor drage nytte af vores indsigt fra studiet af tilnærmelsesalgoritmer og beregningskompleksitet i konteksten af social valgteori. Dog bringer stemmeindstillingen sine egne udfordringer med sig, som opstår fra en mangel på *information*. Vælgerne kan være ude af stand til præcist at udtrykke deres præferencer, for eksempel i numeriske termer. I stedet giver en vælger typisk kun et groft billede af sine præferencer i form af en rangering af alternativerne eller ved at rapportere de alternativer, hun godkender. På den anden side kan det for designeren af stemmeregler være uklart, hvilke kombinationer af ønskelige egenskaber udfaldene skal opfylde for at blive betragtet som optimale.

Et centralt aspekt af vores arbejde er at undersøge, hvilke fordele en begrænset mængde yderligere information kan give i designet af stemmeregler. Denne information kan bestå af en præcis numerisk kvantificering af den nytte (eller omkostning), som en vælger har for en lille delmængde af alternativerne. Vi udforsker, om sådanne forbedringer af stemmeregleres kapacitet kan forbedre deres præstationer væsentligt (målt i forhold til deres *forvrængning*). I kombination med denne målrettede adgang til vælgernes præcise præferencer overvejer vi også stokastiske præferencer, som giver aggregeret information til en stemmeregler, såsom en vælgers gennemsnitlige nytte for alternativerne.

En anden type information, vi undersøger, er *data* om ideelle udfald af en stemmeanvendelse. Hvor meget data i form af præferenceprofil-stikprøver mærket med de ønskede udfald er nødvendige og tilstrækkelige for at *lære* en stemmeregler, der er konsistent med disse prøver og—forhåbentlig—klar sig godt uden for dette træningssæt? I den såkaldte “probably approximately correct” (PAC) læringsmodel undersøger vi denne forestilling om prøvekompleksitet samt beregningskompleksiteten af en relateret læringsopgave for visse klasser af flervinder-stemmeregler.

Acknowledgments

I wish to extend my deepest gratitude to my PhD advisor Ioannis Caragiannis for his commitment to supporting me on my journey to become a researcher. This work would not have been possible without his ingenuity, and I cannot thank him enough for his guidance and patience. I am also extremely grateful to my co-authors Chris Schwiegelshohn, Jakob Burkhardt, Sudarshan Shyam, and Matteo Russo for their insights and their dedication to our joint work. Special thanks go to Jakob for proofreading the Danish version of the previous abstract. Furthermore, I would like to extend my sincere thanks to Ariel Procaccia for hosting me at Harvard University and making me feel welcome there. At Aarhus University, I was always glad to meet with Claudio Orlandi and Jaco van de Pol, and our discussions helped me to successfully complete this journey. I am thankful to Ulle Endriss and Piotr Skowron for agreeing to join the PhD committee. Finally, I am immensely grateful to my mother Iris, and I would like to express my heartfelt appreciation for her support and encouragement.

*Karl Fehrs,
Aarhus, January 31, 2025.*

Contents

Abstract	i
Resumé	iii
Acknowledgments	v
Contents	vii
I Overview	1
1 Introduction	3
2 Average-Case Distortion	7
2.1 Distortion of Voting Rules	7
2.2 Beyond the Worst Case: Stochastic Preferences	9
2.3 Highlights of our Results and Techniques	11
2.4 Open Problems and Recent Developments	14
3 Distortion of Metric Multiwinner Voting Rules	17
3.1 Metric Voting	17
3.2 Clustering with Limited Cardinal Information	20
3.3 Highlights of our Results and Techniques	22
3.4 Open Problems and Recent Developments	25
4 Learnability of Approval-Based Multiwinner Voting Rules	29
4.1 Hardness of Approval-based Multiwinner Voting	29
4.2 PAC Learnability of Voting Rules	31
4.3 Highlights of our Results and Techniques	33
4.4 Open Problems and Discussion	35

II Publications	39
5 Beyond the Worst Case: Distortion in Impartial Culture Electorates	41
5.1 Introduction	41
5.2 Preliminaries	45
5.3 Average Distortion Can Be High	47
5.4 Constant Average Distortion with a Single Query per Agent	50
5.5 Randomized Mechanisms	56
5.6 Worst-Case Distortion Lower Bounds	62
5.7 Discussion and Open Problems	68
5.8 Appendix to Chapter 5	69
6 Low-Distortion Clustering with Ordinal and Limited Cardinal Information	79
6.1 Introduction	80
6.2 Preliminaries	82
6.3 Algorithms for k -Center	84
6.4 Lower Bounds for (k, z) -Clustering	90
6.5 Facility Location with Uniform Opening Costs	94
6.6 Conclusion and Open Problems	96
7 The Complexity of Learning Approval-Based Multiwinner Voting Rules	97
7.1 Introduction	97
7.2 Preliminaries	101
7.3 PAC Learning Background	103
7.4 The Learnability of ABCS Rules	105
7.5 The Complexity of TARGETABCS	110
7.6 Sequential Thiele Rules	113
7.7 Concluding Remarks	122
7.8 Appendix: Hardness of Winner Verification for the CC Rule	123
Bibliography	129

Part I

Overview

Chapter 1

Introduction

The study of voting rules is beset by the existence of strong impossibilities. The famous results by Arrow [14] as well as by Gibbard [73] and Satterthwaite [116] show that certain combinations of properties (so-called *axioms*) which appear desirable and—to some extent—natural for the outcome of an election can only be guaranteed by dictatorial rules. Beyond this axiomatic approach to voting, one could consider models of collective decision-making where there is a socially optimal outcome, for example, in terms of the electorate’s utility for different alternatives or in terms of a predefined set of ideal decisions. Is it possible to design voting rules that are guaranteed to achieve the social optimum or, at the very least, come reasonably close to achieving the optimum in most of the cases? In this dissertation, we present the progress that we made towards answering these questions by following the modern computation approach to social choice [27].

The treatment of collective decision-making as an optimization problem has a long history in social choice theory (see, for example, the discussion by Young [127]). The *utilitarian approach* to voting which can be traced back to the work of Bentham [21] in the 18th century seeks to maximize the sum of the voters’ individual *utilities* for different outcomes. The assumption that a voter’s utility for any given alternative is indeed quantifiable as an exact *cardinal* (that is, numerical) value has since found wide application among economists and social choice theorists. Still, while individual utilities may be quantifiable, it must not be the case that each voter is able to readily formulate and report her exact utility for every single alternative. The voter might lack the required insights or the computational burden of assigning a precise value to each alternative could simply be too high. Similarly, there may exist communication constraints (for example, in terms of available bandwidth) which prevent the voter from reporting her preferences in entirety and with full precision.

These challenges with *preference elicitation* are usually met with the assumption that each voter only reports a rough outline of her valuation of the different alternatives. For example, such a report could take the form of a *ranking* of all alternatives by the voter or a subset of the alternatives which the voter approves of. A voting rule which takes into account only these simplified summaries of the voters’ preferences

may choose outcomes that are far from the social optimum. In order to quantify this loss of outcome optimality, the concept of *distortion* was introduced by Procaccia and Rosenschein [111]. They adopt the utilitarian notion of *social welfare* which is defined as the sum of the voters' individual utilities for a given alternative. The distortion of a voting rule then is the worst-case ratio over all possible electorates (including their utilities and reported preferences) between the social welfare of the alternative returned by the rule and the maximum social welfare of any alternative.

Let us for now consider ranking-based voting rules. The aforementioned impossibility theorems already set hard limits to what can be achieved by these rules with respect to desirable axioms. Can we be more confident when approaching the design of ranking-based rules from an utilitarian perspective? Here, we assume that each voter's ranking is an ordering of the alternatives according to her private valuations such that alternatives with higher utility appear further up in her ranking. From simple examples (Example 1), it can be seen that the distortion of any rule which receives as input only the reported rankings without having access to the voters' underlying utilities is arbitrarily high. The situation changes if we restrict the voters' valuations, for example, by assuming that, for every voter, her utility for the alternatives sums to one (*unit-sum valuations*). While bounded distortion can now be achieved by ranking-based rules [33], it is not immediately clear in how far such restrictions are justified beyond very idealized settings.

The designer of voting rules now appears to be left in a difficult position. No single rule can satisfy all desirable axioms, and—assuming that the designer settled on a feasible combination of axiomatic properties—further concessions may be required when taking into account the utilitarian aspect of wanting to approximate the social optimum. The designer therefore might have to consider complex trade-offs between different objectives when tailoring a voting rule to the application at hand. But what if this seemingly hard task was accomplished previously, for example, by careful deliberation among the members of the electorate or by the work of a board of experts on technical, societal or ethical matters? Can we learn from such preceding experience (in the form of given sets of preferences labeled by desired outcomes) and recover a voting rule that is close to the ideal rule for the application in question? This *data-driven* approach provides yet another perspective on the challenge of designing optimal voting rules.

In this work, we explore the powers and limitations of voting rules in two utilitarian models of decision-making as well as the task of learning voting rules from data. Our investigation covers various types of preferences and classes of rules which we briefly outline in the following roadmap. However, before we proceed, we would like to make a remark to frame our work: When considering voting and decision-making in general, one might immediately think of political procedures such as presidential elections. Besides their one-shot, high-stakes nature, these settings are typically characterized by a large population of voters who need to decide between very few alternatives. While these scenarios are important, they are not our sole (or even predominant) concern. That is, we also have in mind low-stakes decision-making situations such as event participants wanting to decide between different food options as well as situations

where a small number of entities (such as a group of autonomous robots) needs to decide between a very large set of possible strategies. Another relevant application are *rating tasks* where, for example, a community of cinema aficionados wants to decide on a representative Top-20 movie selection. Therefore, from this point onward, we usually refer to the members of the electorate as *agents* (instead of *voters*). We assume that these agents want to decide between different *alternatives* (and typically do not refer to these alternatives as *candidates*).

Roadmap of the Dissertation. This dissertation combines three works which have been published as [28], [31], and [32]. We summarize the different models together with their broader social choice context and our main results in Chapters 2, 3, and 4. Part II of the dissertation then contains the complete technical presentation of two of these works which can be found in Chapter 5 [32], and Chapter 7 [31]. Our remaining work [28] is currently undergoing significant editorial changes in preparation for a journal submission. In Chapter 6, we present a selection of our results which benefit from considerable additional polish over their preliminary version.

In Chapter 2, we begin our investigation in the standard utilitarian setting where the agents report rankings which are consistent with their underlying private valuations of the alternatives. Our goal is to design low distortion rules for selecting a single winning alternative. Following a recent approach introduced by Amanatidis et al. [4], we allow our voting rules (which we now refer to as *mechanisms*) to make a limited number of *queries* to each agent about her exact cardinal values for specific alternatives. Our main contribution is a novel average-case notion of distortion which assumes that the agents draw their valuations from a common probability distribution. The distortion of any rule that does not have access to the agents' valuations can still be high in this setting. However, we show that queries to the agents can be used in very minimal ways to achieve low average distortion.

We then turn our attention to *metric voting* in Chapter 3. In particular, we consider the case of *peer selection* where the agents and alternatives are the same set of points located in a metric space. Here, our goal is to select a k -sized subset (or *committee*) of the points that minimizes the *social cost*. In this dissertation, we focus on the *egalitarian* social cost where the cost of the points for a committee is the maximum distance of any point to its closest committee member. If we had full access to the exact distance between any two points, this problem is in fact the well-known *k-center clustering* problem. However, following the usual modeling approach in metric voting, we assume that each point reports a ranking over the entire point set where points that are closer appear in positions further up in the ranking. Similar to our work in the utilitarian (non-metric) setting, our model assumes that the exact distance between any two points can be accessed by a costly *query*. Our work explores the *metric distortion* of clustering algorithms. Aside from the benefits of the additional information gathered from the distance queries, we consider the possibility of *resource augmentation*. That is, we may allow an algorithm to select more than k committee members in order to achieve—ideally, significantly—lower distortion. For the egalitarian social cost (i.e.,

the k -center objective), we obtain essentially optimal distortion/information as well as distortion/solution-size trade-offs.

Our final contribution (Chapter 7) is also concerned with *multiwinner* voting where we want to select k alternatives, i.e., a k -sized *committee*, based on the agents' preferences. However, the submitted preferences now have a different form. Instead of ranking *all* alternatives, each agent simply reports a subset of the alternatives which she approves of. A popular approach in this setting of *approval voting* is to use *scoring rules*. In our work, we investigate whether a certain class of multiwinner voting rules, so-called approval-based multiwinner scoring (ABCS) rules, can be learned from data. That is, given *samples* in the form of the agents' approval votes labeled by the desired winning committees, our goal is to identify a scoring rule which approximates the unknown target rule used for labeling the sampled ballots. We explore this *data-driven approach* to the design of voting rules in the model of *probably approximately correct* (PAC) learning. On the positive side, we show that the class of ABCS rules has low *sample complexity*: A polynomial number of samples is sufficient for learning an accurate representation of the unknown target rule with high confidence. However, on the negative side, our work demonstrates that even elementary learning tasks for ABCS rules are computationally hard. PAC learning this class of voting rules *efficiently* therefore appears to be out of reach. Our findings carry over to the class of so-called sequential Thiele rules which can be thought of as simplified greedy approximations of ABCS rules.

Chapter 2

Average-Case Distortion

2.1 Distortion of Voting Rules

In our first contribution, we consider the standard utilitarian setting of single-winner voting. In this model, there is a set N of n agents, and a set A of m alternatives. Every agent i has a private valuation function $\text{val}_i : A \rightarrow \mathbb{R}_{\geq 0}$ expressing the utility (or *value*) of the alternative to the agent. We refer to the collection of the agents' valuation functions as their *valuations*. The *social welfare* of an alternative $a \in A$ is then defined as

$$\text{SW}(a, v) = \sum_{i \in N} \text{val}_i(a).$$

Our goal is to pick a single alternative of maximum social welfare. This task is however complicated by the usual assumption that the agents do not report their exact values for every alternative but only a simplified summary of their valuations. We consider the case that every agent i reports a *ranking* $\succ_i$, that is, a strict ordering of A . Every agent's ranking $\succ_i$ is *consistent* with her valuation function in the sense that, for any two alternatives $a, b \in A$, $a \succ_i b$ only if $\text{val}_i(a) \geq \text{val}_i(b)$.

A preference *profile* $P = \{\succ_i\}_{i \in N}$ is then simply the collection of the agents' rankings. A *voting rule* takes as input a profile P and returns a single *winning* alternative. The following example shows that there are instances where any *deterministic* voting rule must select an alternative which is arbitrarily far away from the optimal choice in terms of social welfare.

Example 1. Consider a instance with two agents $N = \{1, 2\}$ and two alternatives $A = \{a, b\}$. Assume that $\text{val}_1(a) = 1$ and $\text{val}_1(b) = 0$ whereas $\text{val}_2(a) = 0$ and $\text{val}_2(b) = z$ for some large value $z \gg 1$. Then, the only rankings that are consistent with these valuations are $a \succ_1 b$ and $b \succ_2 a$. By symmetry of this profile and since we are considering only deterministic voting rules, it is without loss of generality to assume that the rule picks a as the winning alternative. The social welfare of this outcome is $\text{SW}(a) = 1$ while the social welfare maximizing alternative is b with $\text{SW}(b) = z \gg 1$.

Procaccia and Rosenschein [111] introduced the notion of *distortion* to quantify this loss of outcome optimality which results from the lack of knowledge about the

agents' exact valuations. In order to formally introduce this concept, we require some additional notation. Let us denote the agents' valuations by $v = \{\text{val}_i\}_{i \in N}$, and let $\mathcal{V}$ be the set of all possible valuations. For a given choice of valuations $v \in \mathcal{V}$, we denote by $\mathcal{P}(v)$ the set of all profiles P such that each ranking $\succ_i \in P$ is consistent with $\text{val}_i \in v$. The distortion of a deterministic voting rule $\mathcal{R}$ then is

$$\text{dist}(\mathcal{R}) = \sup_{\substack{v \in \mathcal{V} \\ P \in \mathcal{P}(v)}} \frac{\max_{a \in A} \text{SW}(a, v)}{\text{SW}(\mathcal{R}(P), v)}.$$

That is, the distortion of the rule $\mathcal{R}$ is the worst-case ratio (over all possible valuations, and profiles consistent with the respective valuations) between the maximum social welfare of any alternative and the social welfare of the alternative returned by $\mathcal{R}$. For randomized rules the denominator in the previous expression becomes the expected social welfare of the alternative returned by the randomized rule $\mathcal{R}$ where the expectation is taken over the random choices of $\mathcal{R}$ on the given profile P .

Despite the observation that the distortion of any deterministic rule is unbounded in the case of unrestricted valuations (Example 1), the concept has found wide and fruitful application in computational social choice (see the survey by Anshelevich et al. [11] for an overview). Recently, Amanatidis et al. [4] proposed an extension of the previously described voting model. Together with the profile, a voting rule receives a limited budget for making so-called *value queries* to each agent. That is, the rule is now able to uncover the exact value $\text{val}_i(a)$ of any agent $i \in N$ for any alternative $a \in A$ at the expense of one unit of its budget. This model still assumes that providing exact cardinal information is burdensome to the agents. However, in many cases, it may be feasible to ask an agent to assign exact numerical values to a small subset of the alternatives, for example, by prolonged introspection or by careful measurement. We use the term *mechanism* for any voting rule that is accompanied by such a query access to the agents' valuations. Importantly, a mechanism's choice of queried positions can be *adaptive*. This means that the mechanism is allowed to perform its queries sequentially using the information that it learned from earlier queries in its current choice.¹

In this model, Amanatidis et al. [4] showed that—for unrestricted valuations—there is a mechanism with constant distortion that makes $O(\log^2 m)$ queries per agent. Using its query access to the valuations, their mechanism performs $\log m$ rounds of binary search on each agent's ranking. In effect, the mechanism is able to approximate a subset of all values (essentially those values of significant magnitude) within a factor of two, and the analysis of Amanatidis et al. then proceeds to bound the contribution of the remaining values to the maximum social welfare of any alternative. In terms of distortion lower bounds, Amanatidis et al. proved that, for any mechanism $\mathcal{M}$ making

¹Indeed, Amanatidis et al. [4] also consider non-adaptive mechanisms that query fixed positions in the ranking of every agent. For this class of rules, they establish asymptotically tight distortion bounds of $\Theta(m/\lambda)$ where λ is the maximum number of positions that the mechanism is allowed to query in each ranking.

at most λ queries per agent, it must hold that

$$\text{dist}(\mathcal{M}) \in \Omega\left(\frac{1}{\lambda+1} \cdot m^{\frac{1}{2(\lambda+1)}}\right).$$

In particular, this bound implies that achieving constant distortion requires $\Omega\left(\frac{\log m}{\log \log m}\right)$ queries per agent. In our work (Theorem 23), we present an improved lower bound of $\Omega(\log m)$ on the number of queries per agent that are necessary to obtain constant distortion. In their follow-up work to [4], Amanatidis et al. [5] improve their previous distortion lower bound to $\Omega(m^{1/\lambda})$ for the special case that λ is constant. Furthermore, in their earlier work, Amanatidis et al. proposed the following conjecture.

Conjecture 2 (Amanatidis et al. [4]). *There is a deterministic mechanism with constant distortion that makes $O(\log m)$ queries per agent.*

Resolving this conjecture (either positively or negatively) was a major motivation for our work in this model. At the time of writing this dissertation, the conjecture remains unresolved though.

2.2 Beyond the Worst Case: Stochastic Preferences

There is a simple recipe for constructing distortion lower bound examples in the previously introduced model: For a carefully chosen profile P , it is decided in advance which values are revealed in the case that the mechanism queries any particular position of an agent's ranking. Then, for the remaining unqueried positions, a *worst-case* set of values is fixed such that the resulting valuations are consistent with the profile P and the revealed cardinal information.

This *adversarial* approach is an obvious choice from the perspective of classical algorithm design (e.g., the study of approximation algorithms for computationally hard problems). However, on the one hand, the resulting worst-case valuations necessarily have a very contrived structure which we would not expect to find among real-world electorates. Furthermore, it may often be natural to assume that the agents' preferences share common characteristics such as having similar utility for the alternatives *on average* which would prevent entirely adversarial valuations from occurring.

The main conceptional contribution of the work that we present in Chapter 5 is a novel notion of *average-case* distortion. We introduce the concept for a variant of the utilitarian voting model where every agent draws her values for the alternatives independently from a common probability distribution F . This average-case model was first considered by Boutilier et al. [25] who—among other directions—studied the average-case optimality of voting rules in this setting. By their definition, a voting rule is optimal if it maximizes the expected social welfare of the winning alternative. They find that the average-case optimal voting rule belongs to the class of positional scoring rules. A positional scoring rule is defined by a non-increasing sequence of nonnegative real-valued scores $\langle \alpha_1, \alpha_2, \dots, \alpha_m \rangle$. An alternative is assigned a score α_j for each of its occurrence in position j of an agent's ranking. The winning

alternative under the respective rule is then some alternative with maximum total score. Boutilier et al. show that the scoring rule where α_j is equal to the expected value of an alternative appearing in the j -th position of a ranking under the distribution F is indeed average-case optimal.

Recently, Gonczarowski et al. [77] introduced a notion of average-case distortion in the model of Boutilier et al. [25]. For a given profile P , they consider the ratio between the expected maximum social welfare of any alternative and the expected social welfare of the winning alternative for a random draw of valuations v conditioned on the event of observing P . That is, the profile P can still be chosen adversarially but the underlying values v follow a common distribution F for all agents. The expected distortion of a voting rule is then given by the previous ratio for a worst-case choice of profile P . The primary aim of Gonczarowski et al. is to design voting rules that are approximately optimal in terms of expected social welfare and expected distortion. Interestingly, their main results relies on the average-case optional (see discussion in previous paragraph) voting rule for the case that the valuations are drawn according to a fair Bernoulli trial. Gonczarowski et al. refer to this rule as *binomial voting*. They show that binomial voting approximates the voting rule of minimal expected distortion for any distribution supported on the interval $[0, 1]$.

Our notion of average-case distortion does not give the adversary control over the profile P . Assume that the agents' valuations are drawn independently from some common probability distribution F . In effect, this results in observing uniformly random profiles which are known in the literature as *impartial culture electorates*. These profiles are still far removed from the preferences that we would expect to observe among real-world electorates. However, the impartial culture assumption can be seen as an instructive stepping stone towards more natural models of stochastic preferences (e.g., see the variants of the model which we propose in Section 2.4).

We conclude this section with the introduction of our concept of *average distortion* and some elementary examples. From now on, we refer to the classical notion of distortion (Section 2.1) as *worst-case* distortion to distinguish the two concepts. Recall the model of Boutilier et al. [25] where every agent draws her values for the alternatives independently from a common probability distribution F which we write as $v \sim F$. We assume that F is *known* to the mechanism designer. Accounting for the possibility of ties, we write $P \sim \mathcal{P}(v)$ to indicate that each agent i chooses her ranking $\succ_i$ uniformly at random among all rankings consistent with her random draw of values. The average distortion of a mechanism $\mathcal{M}$ on a family of distributions $\mathcal{F}$ is then defined as

$$\text{avdist}(\mathcal{M}, \mathcal{F}) = \sup_{F \in \mathcal{F}} \frac{\mathbb{E}_{v \sim F} [\max_{a \in A} \text{SW}(a, v)]}{\mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\text{SW}(\mathcal{M}(P, v), v)]}.$$

With this definition in hand, let us consider the trivial voting rules which ignore the given profile and always pick a fixed alternative or an alternative uniformly at random. For *any* family for distributions, the distortion of these rules is upper-bounded by m . This essentially follows from the observation that every alternative has probability of

$1/m$ for being the social welfare maximizing alternative. Furthermore, let us consider the family of uniform distributions over an interval $[a, b]$ for $a, b \geq 0$. By elementary properties of these distributions, the expected social welfare of any alternative is $n \cdot \frac{a+b}{2}$ whereas the maximum social welfare of any alternative is trivially upper-bounded by $n \cdot b$. Hence, the distortion of the aforementioned trivial rules on this family of distributions is at most 2.

2.3 Highlights of our Results and Techniques

In our work (Chapter 5), we show that the average distortion of any voting rule that does not make queries to the agents' underlying valuations is $\Omega(m)$, see Theorem 7. We prove this result for a distribution F that returns a value of one with probability $1/(nm)$ and a value of zero otherwise. Returning to the discussion at the end of the previous section, this implies that the trivial voting rules that always pick a fixed alternative or an alternative uniformly at random are already essentially optimal.

For the proof of our result, we first observe that, for a large enough $n \in \Omega(m \log m)$, there is an alternative with positive social welfare with constant probability. The challenge then is to upper-bound the expected social welfare of the alternative picked by a voting rule. To this end, we relate the expected social welfare of *any* alternative to the expected *maximum plurality score* (that is, the maximum number of occurrences of any alternative in the first position of the agents' rankings). Using the Chernoff bound, we then upper-bound the expected maximum plurality score under F for our choice of n . We formulate our proof in the language of positional scoring rules revisiting the average-case optimal rule by Boutilier et al. [25] (see previous section) and making use of the majority rule which is defined by the scoring vector $\langle 1, 0, \dots, 0 \rangle$.

The distribution F that we used for our distortion lower bound is a member of a particular family which we call the family $\mathcal{F}^{0/1}$ of *binary distribution*. Every member $F_p \in \mathcal{F}^{0/1}$ is specified by a single value $p \in [0, 1]$ such that F_p returns a value of one with probability p and a value of zero otherwise. Our previous result led us to consider the family $\mathcal{F}^{0/1}$ also with respect to distortion upper bounds. Does there exist a mechanism that makes few queries per agent and has low—ideally, constant—average distortion on $\mathcal{F}^{0/1}$? The answer to this question is Yes. We propose a mechanism that makes only a single query per agent and achieves constant average distortion (Theorem 12). Recall that we do not intend make restrictions on the size of the electorate. In particular, we do not assume that the number n of agents is large compared to the number m of alternatives. In such cases, standard concentration inequalities imply that the social welfare of *any* alternative is very likely to be close to its expectation. Then, the previously mentioned trivial voting rules already have constant average distortion. A formal discussion of this observation together with the required conditions is given by Lemma 30 and our remarks preceding the lemma. In the absence of such simplifying assumptions, the major technical challenge lies with upper-bounding the maximum social welfare of any alternative, i.e., the enumerator in the definition of the average distortion.

Recall that our model assumes that the distribution which the agents draw their values from is known. Hence, with respect to a binary distribution $F_p \in \mathcal{F}^{0/1}$, a mechanism knows the parameter p . Our mechanism which is called MEAN queries each agent for the value of the alternative appearing in position $\tau = \max\{1, \lfloor pm \rfloor\}$ of the agent's ranking. We then employ a notion of *implied social welfare*: Consider a fixed agent that the mechanism queried at position τ of her ranking. If the mechanisms uncovered a value of zero, we assume that the agent has value zero for *all* alternatives. If the mechanism uncovers a value of one instead, we in fact know for certain that the agent has a value of one for all alternatives appearing in positions $1, \dots, \tau$ of her ranking, and we further assume that her values for the remaining alternatives are zero. The mechanism MEAN simply returns an alternative of maximum implied social welfare (summing the previously defined values over all agents).

For our proof that the mechanism MEAN has constant average distortion (Theorem 12), we separate the family $\mathcal{F}^{0/1}$ into three subfamilies:

$$\begin{aligned}\mathcal{F}_1 &= \{F_p \in \mathcal{F}^{0/1} : p \geq 1/m\} \\ \mathcal{F}_2 &= \{F_p \in \mathcal{F}^{0/1} : 1 - (1 - 1/n)^{1/m} \leq p < 1/m\} \\ \mathcal{F}_3 &= \{F_p \in \mathcal{F}^{0/1} : p < 1 - (1 - 1/n)^{1/m}\}\end{aligned}$$

Proving constant average distortion for the first subfamily $\mathcal{F}_1$ is our technically most involved contribution. Conceptionally, our proof can be thought of as a type of covering argument. Figure 2.1 provides the high-level visual intuition for our approach. Indeed, our treatment of the subfamily $\mathcal{F}_2$ relies on a simplified variant of the technique that we developed to handle the subfamily $\mathcal{F}_1$. Finally, for the proof that mechanism MEAN has constant average distortion on the subfamily $\mathcal{F}_3$, we observe that the expected maximum social welfare of any alternative is upper bounded by the expected combined social welfare of *all* alternatives, that is, the term pnm ; our distortion bound in this case then follows from elementary calculations.

Thinking back to our previous distortion lower bound, we have thus shown that, for the family $\mathcal{F}^{0/1}$ of binary distributions, using queries in a very minimal way improves the average distortion a lot. In our work, we also demonstrate that similar distortion bounds cannot be achieved in the classical worst-case setting where the agents have value either one or zero for the alternatives. Specifically, any mechanism that makes one query per agent must have worst-case distortion $\Omega(\sqrt{m})$, see Theorem 24.²

The obvious next step is then to extend our techniques to other families of distributions. Here, we propose a random threshold algorithm which we call RTMEAN. Importantly, this mechanism uses our previous 1-query mechanism MEAN for binary distributions as a subroutine. Assume that the agents draw their valuations from some general probability distribution F , and that we defined some threshold value ℓ . Let us consider the probability $p = \Pr_{z \sim F}[z \geq \ell]$. Our randomized mechanism

²We are not the first to explore the worst-case distortion of 1-query mechanisms. Amanatidis et al. [4] show that the distortion of these mechanisms is $\Omega(m)$. However, their worst-case instances require valuations that are more complex than binary.

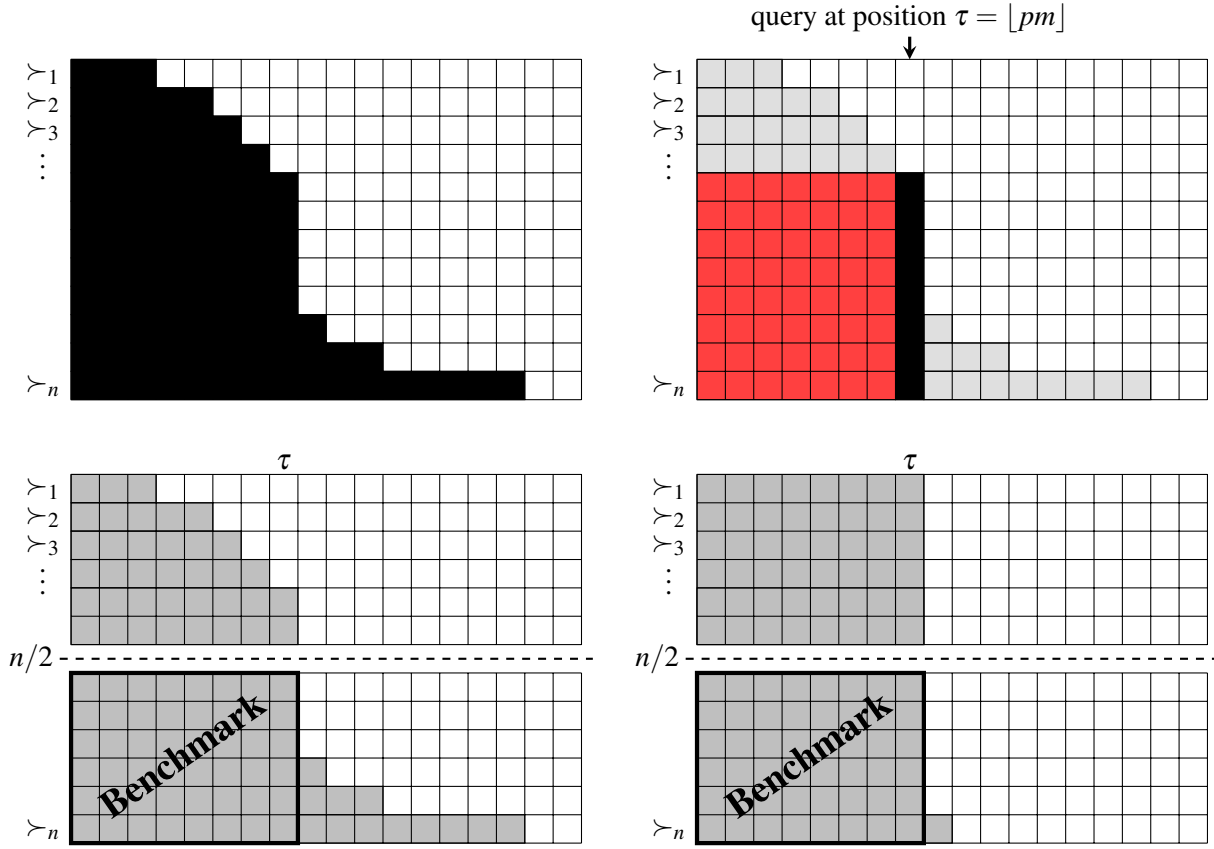


Figure 2.1: In this set of figures, we represent the agents’ rankings as horizontal bars (with the highest-ranked alternative appearing as the left-most entry of a bar). For a hypothetical draw of valuations from a binary distribution, we indicate a value of one underlying a position in an agent’s ranking as a colored box. In the top left figure, we ordered the agents $1, 2, \dots, n$ in increasing order of the number of ones that they drew for the alternatives. The top right figure illustrates how the implied social welfare is calculated when querying each agent at position τ of her ranking: If the query returned a value of one, this value is assigned to any alternative appearing in position τ (black box) or further up in the ranking of the respective agent (red boxes). The agent’s remaining values are assumed to be zero. On average, we expect to uncover at least $n/2$ ones when querying the agents at position τ . On the second row, we introduce a *benchmark value* which is the maximum expected social welfare of any alternative restricted to the box labeled “Benchmark”. The “benchmark box” is entirely filled with values of one. With probability $1/2$, the social welfare of the alternative returned by MEAN is at least this benchmark value. Our goal now is to upper-bound the expected maximum social welfare of any alternative in terms of the same benchmark value. Assume, as a thought experiment, that we re-ordered the ones underlying the profile such that the resulting arrangement is almost completely rectangular-shaped (bottom right figure). We show that such a re-arrangement only increases the expected maximum social welfare of any alternative. Very informally speaking, the constant average distortion of mechanism MEAN then results from bounding the number of “benchmark boxes” that are required on average to cover this hypothetical rectangular configuration. In our example, three benchmark boxes would be required.

RTMEAN simulates an execution of MEAN on the resulting probability distribution F_p by interpreting the answer of agent i to a query about her value for alternative a as one if $\text{val}_i(a) \geq \ell$ and zero otherwise. The winning alternative returned by RTMEAN is then exactly the alternative returned by MEAN for the simulated sequence of queries on F_p . In our work, we show how to define a suitable set of thresholds (Lemma 18), one of which is then randomly chosen for the aforementioned simulation of MEAN. This mechanism still makes one query per agent and achieves average distortion of $O(\log m)$ for the following families of distributions (Corollary 19): exponential, chi-squared, and Erlang- k . Our $O(\log m)$ distortion bound for these families follows from the average distortion guarantees of RTMEAN which in general (Theorem 17) depend on the mean and variance of the distribution F that the agents draw their values from.

2.4 Open Problems and Recent Developments

We see our contribution as laying the groundwork for future research on the distortion of mechanisms under stochastic preferences. Our primary objective was to achieve *constant average distortion* with as few queries per agent as possible. A natural follow-up to our work would be to investigate the existence of constant distortion mechanisms for other families of distributions beyond the elementary family $\mathcal{F}^{0/1}$ of binary distributions. This direction could also aim at developing techniques for constructing distortion *lower bound examples* (potentially extending to randomized mechanisms) given any family of distributions. Another important question is whether we can remove or weaken our assumption that the distribution which the agents draw their values from is known exactly to the mechanism designer. Other variants of our setting include modeling assumptions such as

- agents drawing their values from different distributions (e.g., based on location, political leaning or other indicators such as income and education),
- agents having a different distributions for each alternative, and
- stochastic valuations that satisfy the unit-sum restriction.

In the classical setting of worst-case distortion, the question whether there exists a deterministic constant distortion mechanism that makes $O(\log m)$ queries per agent remains open (Conjecture 2). Is it possible to make progress towards resolving the conjecture by relaxing certain requirements such as allowing the mechanism to make $O(n \log m)$ queries *in total* (instead of $O(\log m)$ queries on a per-agent basis)?

One of our results for the worst-case setting which we have not mentioned so far is that we propose a randomized mechanism that makes $O(\log m)$ queries per agent and achieves worst-case distortion $O(\log m)$, see Theorem 21. In very recent work, Ebadian and Shah [58] present a *deterministic* mechanism that recovers and improves upon the guarantees of our randomized mechanism. Their mechanism is inspired by a matching algorithm (also [58]) for the utilitarian setting where n agents

on one side of the market report rankings over n alternatives on the other side (one-sided matching), and value queries to the agents are permitted. Their mechanism for single-winner voting with n agents and m alternatives decides which values to query based on a sequence of so-called β -approximate stable committees. With $O(\log(\min\{n, m\}))$ queries per agent, the mechanism by Ebadian and Shah achieves distortion of $O(\log(\min\{n, m\}))$, see [58, Theorem 6] for the statement of their result for an arbitrary number λ of queries per agent. In future work, it would be interesting to explore whether there is a true separation between the power of deterministic and randomized mechanisms in the utilitarian setting of single-winner voting with value queries.

Chapter 3

Distortion of Metric Multiwinner Voting Rules

3.1 Metric Voting

For our second contribution, we turn our attention to *metric voting* which can be seen as a more constraint variant of the previous utilitarian setting (Section 2.1). In metric voting, a set N of n agents and a set A of m alternatives are located in a shared metric space $(N \cup A, d)$. Here, $d : (N \cup A)^2 \rightarrow \mathbb{R}_{\geq 0}$ is a distance function that satisfies the following properties for any $x, y, z \in N \cup A$:

- (i) $d(x, x) = 0$,
- (ii) *Symmetry*: $d(x, y) = d(y, x)$, and
- (iii) *Triangle Inequality*: $d(x, z) \leq d(x, y) + d(y, z)$.

That is, d is in fact a pseudometric where the distance between two distinct elements x, y can be zero. In the previous chapter, the agents' cardinal valuations indicated their *utility* for different alternatives. In the current model, the distance $d(i, a)$ between an agent $i \in N$ and an alternative $a \in A$ represents the *cost* that the agent has for the respective alternative. We are now in the realm of *spatial preferences*. For example, one could imagine an D -dimensional *policy space* [100] where the position of an alternative a describes a particular implementation of the D policy decisions and an agent i 's position corresponds to the—in her mind—ideal implementation. The distance (or cost) $d(i, a)$ then quantifies the degree by which agent i 's interests are *misrepresented* by the implementation of alternative a .

Whereas our goal in the previous chapter was to identify an alternative of high social welfare, the objective now is to select an alternative of minimum *social cost*. The social cost of an alternative $a \in A$ is defined as

$$\text{SC}(a, d) = \sum_{i \in N} d(i, a).$$

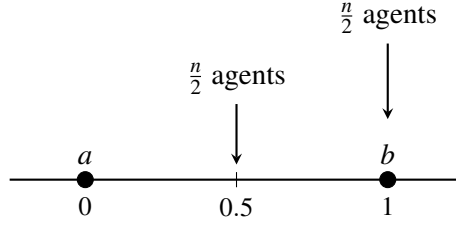


Figure 3.1: This figure illustrates a lower bound example with two alternatives a, b by Anshelevich et al. [10] which shows that all deterministic rules have distortion at least 3. Both alternatives and an even number n of agents are located on a line with Euclidean distances as depicted above. The $n/2$ agents located at position 0.5 each report $a \succ b$; the $n/2$ agents co-located at position 1 with alternative b each report $b \succ a$. By symmetry of the resulting profile and since we are only considering deterministic rules, it is without loss of generality to assume that a voting rule $\mathcal{R}$ selects alternative a as winning. Observe that the social cost of alternative a is $(3/4)n$ whereas the social cost of alternative b is $(1/4)n$. This proves that $\text{dist}(\mathcal{R}) \geq 3$.

Again, we assume that eliciting the exact numerical preferences of an agent (in this case, her distances to each alternative) is prohibitively burdensome. Instead, each agent i reports a *ranking* $\succ_i$ (i.e., a strict ordering) of all alternatives. The agent's ranking is *consistent* with the distances d in the sense that, for any $a, b \in A$, $a \succ_i b$ only if $d(i, a) \leq d(i, b)$. As before, we call the collection $P = \{\succ_i\}_{i \in N}$ a *profile*, and, for a given distance function d , we write $\mathcal{P}(d)$ to denote all profiles P where each ranking $\succ \in P$ is consistent with d . A voting rule $\mathcal{R}$ that receives as input only the ordinal information, i.e., a profile P , may then fail to identify an alternative of minimum social cost. The previous notion of *distortion* which measures this loss of outcome optimality is then also applicable in metric voting. Specifically, the distortion of a voting rule $\mathcal{R}$ is now given by

$$\text{dist}(\mathcal{R}) = \sup_{\substack{(N, A, d) \\ P \in \mathcal{P}(d)}} \frac{\text{SC}(\mathcal{R}(P), d)}{\min_{a \in A} \text{SC}(a, d)}.$$

Notice that, despite the agents now having costs instead of utilities for the alternatives, higher distortion still indicates less optimal outcomes, and that the best possible distortion of any rule is still 1.

The notion of distortion in metric voting was introduced by Anshelevich et al. [10] who, in the same work, established a distortion lower bound of 3 for all deterministic rules. We reproduce their lower bound example in Figure 3.1. The work by Anshelevich et al. inspired a remarkable series of results. Initially, Anshelevich et al. [10] accompanied their 3-distortion lower bound with the finding that the well-known *Copeland* rule has distortion at most 5. Munagala and Wang [103] then presented a distortion upper bound of $2 + \sqrt{5} \approx 4.236$ using a generalization of the approach by Anshelevich et al. [10]. Gkatzelis et al. [74] were the first to present a voting rule

	deterministic	randomized
$q \leq k/3$	∞	∞
$k/3 < q \leq k/2$	$\Theta(n)$	$\Theta(n)$
$q > k/2$	3	$3 - 2/n$

Table 3.1: Overview of distortion bounds by Caragiannis et al. [38] for deterministic and randomized multiwinner voting rules.

that achieves optimal distortion of 3. Their rule essentially inspects a sequence of so-called domination graphs (one of which is defined for every alternative) for the existence of a perfect matching which, as Gkatzelis et al. show, must be admitted by at least one of these graphs. A much simpler voting rule called *plurality veto* with optimal distortion of 3 was then proposed by Kizilkaya and Kempe [88]. The latter authors proved their distortion bound by showing that the dominance graph of the alternative returned by plurality veto always admits a perfect matching. Finally, Charikar et al. [43] demonstrated that randomization can be used to achieve distortion better than 3 by introducing a randomized voting rule with distortion 2.753.

We have so far discussed single-winner voting where the goal is to pick a single alternative of minimum social cost. However, there are many settings which call for multiple (let's say, k) alternatives to be selected. For example, we could think of a diverse research community choosing a representative committee of experts tasked with reviewing papers for a workshop or event participants selecting between multiple food options taking into account their different food preferences and dietary restrictions. These scenarios may also be modeled by spacial preferences (e.g., by rating the expertise of a potential reviewer in the D research fields relevant to the community on D numerical scales) resulting in *metric multiwinner voting*. In this context, we call any subset C of the alternatives A a *committee*.

However, how should we define the social cost of the agents for a k -sized committee C ? On the one hand, we could define the cost of an agent i for C as the distance to the alternative $a \in A$ that is *closest* to i (i.e., the alternative that represents her the most). Another option is to assume that the cost of agent i for the committee C is the sum of her distances to each alternative in C . Indeed, Caragiannis et al. [38] consider a notion of social cost for metric multiwinner voting which interpolates between the two previous extremes. In particular, they point out that our latter suggestion of defining the social cost based on the sum of the agents' distances to each alternative on the committee runs the risk of succumbing to the *tyranny of the majority*. That is, the minimum cost committee under this objective may in fact impose very low costs on a majority of the agents; however, this solution may then entirely disregard the interests of a minority which still incurs significant social cost. Instead, Caragiannis et al. define, for an integer parameter q where $1 \leq q \leq k < m$, the q -cost that agent i has for a k -sized committee C as her distance to the alternative $a \in C$ that is q -th closest to i . The definition of the q -cost reflects the desire that, for every agent, at least q committee members should be somewhat representative of the agent's interests. The

q -cost of a committee C is then the sum of the agents' individual q -costs for C . For a fixed choice of k and q , Caragiannis et al. proceed to introduce the (k, q) -distortion of a (potentially, randomized) multiwinner voting rule $\mathcal{R}$ (which takes as input a profile P , and returns a k -sized committee) in the usual way: The (k, q) -distortion of $\mathcal{R}$ is the worst-case ratio taken over all possible pseudometric spaces (N, A, d) and profiles consistent with d between the (expected) q -cost of the committee returned by $\mathcal{R}$ and the minimal q -cost of any k -sized committee. In their work, the previous authors find an interesting trichotomy of distortion bounds based on the choice of q relative to the committee size k (see Table 3.1).

3.2 Clustering with Limited Cardinal Information

Let us consider a special case of the metric multiwinner voting model where the set of agents N and the set of alternatives A are identical. This is the setting of *peer selection*. Here, a population X of n individuals wants to choose a representative committee $C \subseteq X$ of size k among themselves. Suppose that we decided that the most appropriate definition of the cost of an agent for such a committee is her distance to the individual on the committee that is *closest* to her. In the terminology of Caragiannis et al. [38], this is the 1-cost of the agent for the committee, i.e., the case that $q = 1$.

If we had full access to the distances d between the individuals, then the problem of identifying a minimum cost committee is immediately familiar from the study of algorithms. It is in fact the *k-median clustering problem*. Let us define the distance between an individual (or *point*) $x \in X$ and a set of points $S \subseteq X$ as

$$d(x, S) = \min_{y \in S} d(x, y).$$

In the k -median problem, we are given the entire pseudometric space (X, d) and are tasked with identifying a k -sized set $C \subseteq X$ that minimizes the k -median objective $\sum_{x \in X} d(x, C)$. For general metric spaces, the optimal solution under the k -median objective is known to be hard to approximate within a factor of $1 + 2/e \approx 1.74$ [81].

The k -median problem generalizes to (k, z) -clustering where, again for a k -sized set $C \subseteq X$, the previous objective is replaced by

$$\phi_z(C, d) = \sqrt[z]{\sum_{x \in X} d(x, C)^z}.$$

Besides k -median ($z = 1$), another (k, z) -clustering objective that is of particular interest to us is the case that $z \rightarrow \infty$. This is the so-called *k-center* objective where the social cost is defined as $\max_{x \in X} d(x, C)$ for a k -sized solution $C \subseteq X$. In the social choice literature, objectives such as the k -center objective that seek to minimize the cost (respectively, maximize the utility) of the *worst-off* agent are called *egalitarian*.

For our second contribution (Chapter 6), we take inspiration from the metric voting model and explore the distortion of algorithms for (k, z) -clustering. Here, we assume that each point x reports a *ranking* $\succ_x$ over X that is *consistent* with this

point's distances to the other points. The ordinal information is freely available to the algorithm in the form of a profile $P = \{\succ_x\}_{x \in X} \in \mathcal{P}(d)$. However, the exact distance between any two points $x, y \in X$ can only be accessed by a *query* which consumes one unit of the algorithm's budget for making such queries. Our work in this setting includes the study of the distortion of *randomized* algorithms. For a given (k, z) -clustering instance (X, d) , let $C^*(d)$ be the k -sized solution of minimum cost. An algorithm $\mathcal{A}$ with query access to the distances d has distortion D with constant (respectively, high) probability if

$$\sup_{\substack{(X, d) \\ P \in \mathcal{P}(d)}} \frac{\phi_z(\mathcal{A}(P, d), d)}{\phi_z(C^*(d), d)} \leq D$$

with probability at least $2/3$ (respectively, at least $1 - 1/n$). Furthermore, we define the *expected distortion* of $\mathcal{A}$ as

$$\sup_{\substack{(X, d) \\ P \in \mathcal{P}(d)}} \frac{\mathbb{E}[\phi_z(\mathcal{A}(P, d), d)]}{\phi_z(C^*(d), d)}.$$

Using standard terminology from the clustering literature, we refer to the points in a solution $C \subseteq X$ as *centers*. A set of centers C defines a partitioning of the points X into *clusters* $\{A_c\}_{c \in C}$ where, for every center $c \in C$,

$$A_c = \{x \in X : c \succ_x c' \text{ for all } c' \in C \setminus \{c\}\}.$$

We typically call the set of clusters a *clustering* of X . Notice that, given any set of points C , the clustering induced by C can be determined based on the ordinal information alone. Furthermore, this clustering is *unique* under the given ordinal profile P .

Returning to the concept of (k, q) -distortion by Caragiannis et al. [38], notice that our notions of distortion only covers the case that $q = 1$. The goal of minimizing the q -cost for $q > 1$ is known in the literature as *fault-tolerant clustering*. Let us consider the full information setting where the distances d are known exactly. For the k -median ($z = 1$) and the k -center ($z \rightarrow \infty$) objective, there exists a simple technique by Kumar and Raichel [89] to achieve a constant approximation of the optimal fault-tolerant clustering. For $m = \lfloor k/q \rfloor$, these techniques extend a non-fault tolerant clustering C of size m in an intuitive way: For every $c \in C$, the q points that are closest to c in the metric space (X, d) are included in the solution. Assuming that we used a subroutine for computing C which guarantees a constant factor approximation of the optimal non-fault tolerant clustering (which are known to exist for the k -median [79] and k -center objective [78]), the extended solution is a constant approximation of the optimal fault-tolerant clustering.

In the context of social choice, one particular application of *peer selection* in the full information setting is *sortition*. Here, a population N of n individuals wants to select a representative k -sized committee (or *panel*) $C \subseteq N$. The goal is to choose a

panel of minimum q -cost under the additional *fairness constraint* that every individual is selected for the panel with probability ideally k/n [60]. This model was introduced to capture real-world civic endeavors which aim to bring together diverse panels of citizens that are representative of the population as a whole. For example, these *citizen assemblies* [69] could then be tasked with deliberating on important societal matters such as climate change and with formulating policy recommendations based on their findings. Known clustering techniques have been leveraged in sortition by Ebadian and Micha [57] to achieve *proportional representation* (a different objective that uses the notion of the *core*) under the previous fairness constraint. Recently, Caragiannis et al. [39] weakened the typical full information assumption and explored the concept of metric distortion in a two-stage model of sortition. In particular, they re-introduce different alternatives A that lie in the same pseudometric space as the individuals N . However, only the distances between any two individuals are public while the individuals' distances to the alternatives are unknown. A k -sized panel is then chosen in a fair manner using either uniform selection among the individuals or the clustering technique by Ebadian and Micha [57]. Finally, the single alternative that minimizes the sum of distances *to the panel members* (i.e., the k -median objective) is selected as winning. In this model, Caragiannis et al. define the metric distortion based on the worst-case ratio between the social cost of the *whole* population for the winning alternative and the minimum social cost of any alternative.

3.3 Highlights of our Results and Techniques

In this section, we focus on our results for the k -center objective (i.e., the egalitarian social cost) where, for a k -sized solution $C \subseteq X$, we have that $\phi_\infty(C, d) = \max_{x \in X} d(x, C)$. In particular, we consider (α, β) -bicriteria approximations for the ordinal k -center problem with query access to the distances d . These are approximation algorithms $\mathcal{A}$ that have distortion α and choose at most $\beta \geq k$ centers. The latter parameter β can be thought of in terms of *resource augmentation*. That is, we allow the algorithm to pick an ideally small number of additional centers beyond the target size k for the purpose of attaining the desired distortion α .

With respect to lower bounds (Theorem 42), we find that any algorithm that makes no queries but only relies on the ordinal information must select $\beta \in \Omega(2^k)$ centers in order to achieve bounded distortion with constant probability. Furthermore, for any fixed α , we show that $\Omega(k)$ queries are required by any algorithm that has distortion α with constant probability. Both statements are obtained from the same set of worst-case instances which we illustrate in Figure 3.2. Here, our general approach is to design a distribution over distance functions d . The ordinal profile presented to the algorithm is then however independent of the set of distances sampled from this distribution.

More specifically, we choose the size $|X| = n$ of our hard instances such that $n = 2^{k-1}$. Let T be a complete binary tree of depth $k - 1$. The points X are the leaves of T and the distance between any two points depends on the values that we assign

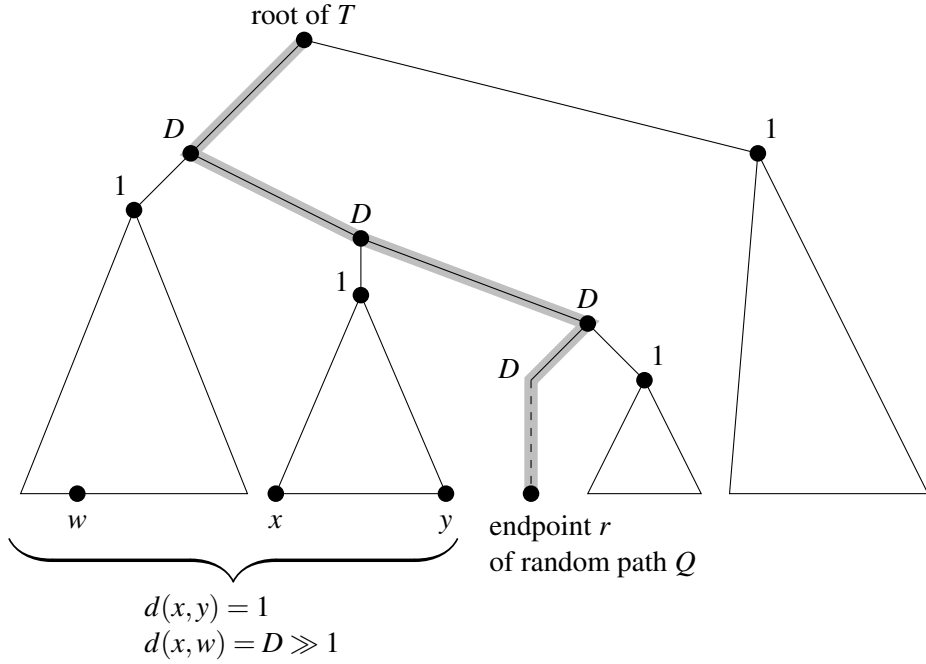


Figure 3.2: Here, we illustrate our approach of defining hard instances (X, d) for the ordinal k -center problem. T is a complete binary tree, and the point set X corresponds to the leaves of T . Q is a path from the root of T to a leaf r chosen uniformly at random. Each interior node of T is assigned a value. This value is $D \gg 1$ for every interior node that lies on the path Q , and 1 for all other interior nodes.

to the internal nodes of T . For any two $x, y \in X$, the distance $d(x, y)$ is then the value assigned to the common ancestor of x and y (i.e., the minimum-depth node on the shortest path between x and y) in T . For our distribution over distance functions d , we choose a leaf r uniformly at random among X . Let Q be the path from the root of T to r . The interior nodes lying on Q are assigned the value $D \gg 1$, and all other interior nodes are assigned the value 1. On the other hand, the ordinal preferences only depend on the relative depth of the leaves' common ancestors in T . For any two leaves x, y , we denote their common ancestor as $a(x, y)$. The preferences of any three points $w, x, y \in X$ are then defined as follows:

- If the depth of $a(x, y)$ is larger than the depth of $a(x, w)$, then $y \succ_x w$.
- Otherwise, if the depth of $a(x, y)$ and $a(x, w)$ are the same, then we choose the ordinal preference of x between y and w lexicographically.

Notice that this choice of preferences is indeed independent from the choice of the path Q (or, equivalently, from the choice of its endpoint r).

The optimal choice of centers C^* under the k -center objective (and, indeed, under any (k, z) -clustering objective) is to place two centers in the depth-1 subtree containing

r and a single center in each of the $k - 2$ subtrees rooted at a node where the value assigned to its parent is D . This solution incurs a cost of $\phi_\infty(C^*, d) = 1$. Our distortion lower bounds (Theorem 42) then essentially follow from the observation that an algorithm cannot identify the path Q (equivalently, its endpoint r) based on the ordinal information alone or will fail to do so with probability at least $1/2$ when making only $o(k)$ queries. In these cases, the cost of the algorithm's solution under the k -center objective is D which can be arbitrarily large, and thus results in unbounded distortion.

In our work, we also present deterministic constant-distortion algorithms for the ordinal k -center problem that perform essentially optimal under the previous distortion/solution-size and distortion/information trade-offs. Our algorithms are based on the *farthest-first traversal* method by Gonzalez [78]. This procedure initially includes an arbitrary point in the solution. In $k - 1$ iterations, the point with maximum distance to its closest center in the current solution is then selected as a new center. Gonzalez showed that the final solution yields a 2-approximation to the k -center problem. The farthest-first traversal method is of particular interest in the ordinal setting since, for any cluster A_c , the point $x \in A_c$ with maximum distance to the center c is identifiable from the ordinal profile alone. Recall that, for any given set $C \subseteq X$, the unique clustering $\{A_c\}_{c \in C}$ induced by C can be determined using only the ordinal information. For a center $c \in C$, the point $x \in A_c$ with maximum distance to c is simply the lowest-ranked point according to $\succ_c$ among A_c . Let us refer to this point x as the *farthest point* of the respective cluster A_c . Using our previous observation, we can execute a farthest-first traversal based on the ordinal profile alone by selecting the farthest point of *every* cluster as a center in each iteration of the procedure. This yields a deterministic $(2, 2^{k-1})$ -bicriteria approximation algorithm for the ordinal k -center problem (Theorem 37).

We then consider the limited-information setting where the algorithm can make distance queries. Here, we propose a deterministic 4-distortion algorithm for the ordinal k -center problem which always returns a solution of size k using $2k - 1$ queries (Theorem 38). Our algorithm relies on a generalization of the previous method by Gonzalez [78] which we call a γ -approximate *farthest-first traversal*. For a choice of $\gamma \in (0, 1]$ and an arbitrary initial center, this procedure includes in each iteration *any* point x such that the distance of x to its closest center is at least a γ -fraction of the maximum distance of any point to its closest center. We show that under this condition the final solution after $k - 1$ iterations is a $(2/\gamma)$ -approximation of the optimal k -center solution (Lemma 35).

Our work demonstrates that a $(1/2)$ -approximate farthest-first traversal can be performed with only $2k - 1$ distance queries. This yields the 4-distortion guarantee of our algorithm. We conclude this section by outlining the intuition for our algorithm. Again starting with an arbitrary center, the algorithm maintains a set of (center, farthest point)-pairs. However, for only some of these pairs, the actual distance between the respective points is queried by the algorithm. Using the limited number of known distances together with the ordinal information, we are able to

- bound the number of clusters affected by adding a new center to the solution,

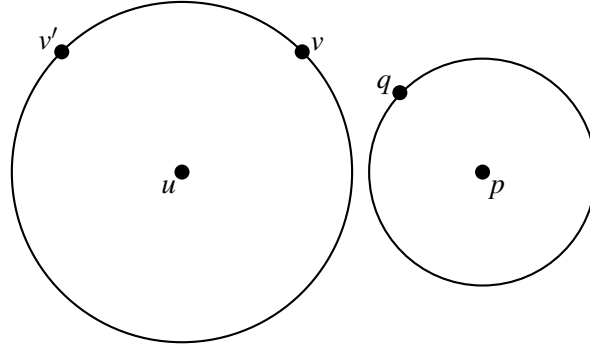


Figure 3.3: This figure gives a high-level intuition for the way that we define Algorithm 3. In this example, the points u and p are centers that have already been included in the solution. We assume that the distance $d(p, q)$ from center p to the farthest point q in its cluster has been queried at some point and is thus known to the algorithm. First, suppose that v is the farthest point in the cluster of center u . In this case, it holds that $v \succ_q p$, and we show that this implies that $d(u, v) \leq 2d(p, q)$, see Lemma 41. Hence, for the purpose of performing a $\frac{1}{2}$ -approximate farthest-first traversal, it suffices to include q in the solution instead of v . The algorithm therefore does not need to query the distance $d(u, v)$ at this point. Now, suppose that v' is the farthest point in the cluster of center u , instead. In this case, we have that $p \succ_q v'$. Thus, if the algorithm includes point v' in the solution then the point q still belongs to the cluster of center p . In particular, q is still the farthest point in the respective cluster. Hence, since we assumed that $d(p, q)$ is known already, the algorithm does not need to spend another query on the cluster of center p upon including v' in the solution.

and

- identify (center, farthest point)-pairs whose distance from each other is not more than twice a currently known distance of such a pair.

Figure 3.3 illustrates these two aspects. Thereby, we are able to reduce the number of new (center, farthest point)-pairs eligible for being queried while guaranteeing that the set of queried pairs always includes a distance that is at least half of the distance between any unqueried (center, farthest point)-pair.

3.4 Open Problems and Recent Developments

In the previous section, we presented our results for the k -center objective, that is, the egalitarian social cost. Our published work also includes bounds for the k -median objective which is more commonly encountered in social choice than the previous objective, and typically referred to as the utilitarian social cost. Indeed, the so-called *ordered k -median* objective allows to interpolate between these two notions of social

cost. Consider a non-increasing vector $w = (w_1, w_2, \dots, w_n)$ of nonnegative real-valued weights. For a given set of centers $C \subseteq X$, assume that we also labeled the points in X as $(1, 2, \dots, n)$ in non-increasing order of their distance to their closest center among C . For any given weight vector w , let us define the social cost of C as

$$\sum_{i=1}^n w_i \cdot d(i, C).$$

For example, the weight vector $(1, 0, \dots, 0)$ then yields the k -center objective whereas the weight vector $(1, 1, \dots, 1)$ corresponds to the k -median objective. Constant factor approximations to the ordered k -median problem are known in the clustering literature [29, 40]. It would be interesting to explore whether constant distortion can be achieved with a small number of queries under this generalization. How do the necessary distortion/solution-size and distortion/information trade-offs develop when interpolating between the two extremes of egalitarian cost (k -center) and utilitarian cost (k -median)?

Our work is also limited to the *peer selection* setting where we do not differentiate between agents and alternatives. It is tempting to believe that our results carry over to the latter more general setting, maybe at the expense of some constant factor increases in distortion or in the required number of queries. However, our attempts at formalizing this intuition have so far been unsuccessful.

If we instead decided to remain in the peer selection setting, it would be a natural choice to further explore the usefulness of known clustering techniques in the model of sortition under the q -cost objective (Section 3.2). This model requires as an additional fairness constraint that, for a k -sized solution (or *panel*) $C \subseteq X$, every individual $x \in X$ is selected as a panel member with probability ideally k/n . For the case that $q \leq k/2$, Ebadian et al. [60] showed that this degree of perfect fairness is essentially incompatible with approximating the q -cost of the optimal *unconstrained* solution. However, let us assume that we chose a desired level of fairness $\lambda \in (0, 1]$. That is, for a so-called λ -fair solution C returned by a selection algorithm and each $x \in X$, it shall hold that $\Pr[x \in C] \geq \lambda \cdot \frac{k}{n}$. Can we leverage existing clustering techniques in order to achieve constant factor approximations to the social cost of the optimal λ -fair solution?

Finally, we want to highlight a very recent result in metric voting and contrast it to the stochastic approach that we use to study our notion of average-case distortion in Chapter 5 (summarized in Chapter 2). The line of work on upper-bounding the metric distortion that we outlined in Section 3.1 prominently features the use of *tournament rules*. These rules rely on pairwise comparisons between alternatives by the agents. Assuming some fixed tie-breaking scheme, the outcomes of these comparisons are uniquely determined by the agents' and alternatives' positions in the metric space. However, we could also entertain the idea that there is some degree of uncertainty affecting an agent's comparison. Indeed, Goyal and Sarmasarkar [80] assume that the *probability* of an agent to prefer an alternative a over another alternative b or vice versa is proportional to her distance to the respective alternatives. Notice

that the actual position of the agent and the alternatives is still deterministic which conceptionally differs from our assumption of randomly drawn utilities in Chapter 5. A particular appeal of the approach by Goyal and Sarmasarkar is that it renders the usual lower bound constructions less prohibitive. Consider the 3-distortion lower bound by Anshelevich et al. [10] which we illustrated in Figure 3.1: The $n/2$ agents located at position 0.5 are essentially indifferent between the two alternatives. Under the probabilistic assumptions used by Goyal and Sarmasarkar, such an agent i would report $a \succ_i b$ and $b \succ_i a$ with equal probability $1/2$. Notice that the probability for observing a majority for alternative a is then exceedingly low for large enough n . In their probabilistic model, Goyal and Sarmasarkar [80] show that the Copeland rule has distortion only 2 (instead of 5 as shown by Anshelevich et al. [10] for the classic metric voting setting) in the limit for $n \rightarrow \infty$. A similar approach was previously considered by Caragiannis and Procaccia [33] in the utilitarian (non-metric) setting under the term *random embeddings into voting rules*. The results by Goyal and Sarmasarkar further demonstrate the power of this—maybe under-utilized—concept. The distortion of voting rules under the stochastic preferences that arise from these random embeddings could be further explored in the non-metric setting as an alternative to the random utility model.

Chapter 4

Learnability of Approval-Based Multiwinner Voting Rules

4.1 Hardness of Approval-based Multiwinner Voting

For our last contribution (Chapter 7), we assume that the agents report their preferences in a different form than the previously considered rankings. We again have a set $\mathcal{N}$ of n agents and a set Σ of m alternatives. Instead of asking each agent for a complete ordering of all alternatives, we simply require each agent to report those alternatives which she approves of. The formulation of these *approval ballots* is typically assumed to be less burdensome to the agents than providing complete rankings. Laslier and Sanver [94, Part I, Chapter 3] discuss this aspect of preference elicitation and other practical considerations that speak for and against the use of approval ballots in great detail. Specifically, we are interested in *multiwinner approval voting* where the goal is to pick a *committee* $C \subseteq \Sigma$ of k alternatives based on the agents' approval ballots. Multiwinner approval voting has found practical application, for example, in *scientific communities* for the purpose of electing their governing bodies [26], and for letting citizens decide the allocation of public funds to municipal projects (*participatory budgeting*) [76].

In our work, we consider the class of so-called approval-based multiwinner scoring (ABCS) rules which were introduced by Lackner and Skowron [92]. In order to formally define this family of rules, we require some additional notation. For an agent $i \in \mathcal{N}$, we denote by $\sigma_i \subseteq \Sigma$ the agent's approval ballot (or *approval vote*). A profile $P = \{\sigma_i\}_{i \in \mathcal{N}}$ is then the collection of the agents' approval votes. Under an ABCS rule, any committee $C \subseteq \Sigma$ receives a score from each agent i 's approval vote σ_i that depends on the size of the intersection between C and σ_i and on the number of approved alternatives according to σ_i . The required scoring parameters are given by a bivariate scoring function f so that the total score assigned to a committee C by a profile P is

$$\text{sc}_f(C, P) = \sum_{i \in \mathcal{N}} f(|C \cap \sigma_i|, |\sigma_i|).$$

By definition, any scoring function f needs to be monotonic non-decreasing in its first argument. The *winning committee* under an ABCS rule f on profile P is a k -sized subset $C \subseteq \Sigma$ with maximum total score.

There is a large body of work aimed at characterizing different ABCS rules and related approval-based multiwinner voting rules in terms of their *axiomatic* properties, e.g., see [65, 90, 93]. We follow the approach of Procaccia et al. [113] and assume that there is a designer that has a specific ABCS rule in mind which exhibits a satisfactory combination of axiomatic properties. The task that we set out to solve is then to recover a close approximation of this rule from data which is given in the form of approval profiles labeled by winning committees. Hence, our work is closely related to the study of computational challenges that arise when determining winners or scores under different ABCS rules. Indeed, the computational complexity of multiwinner approval voting has been investigated by the social choice community. In the following, we give an overview of relevant complexity notions, computational problems, and a few results specifically related to ABCS rules.

When considering questions such as determining all winning committees under a particular ABCS rule, an obvious combinatorial challenge is that there are $\binom{m}{k}$ different committees. Evaluating the rule on all possible committees is therefore infeasible under the typical efficiency criteria of polynomial-time computability. On the other hand, for many practical application, the committee size k is often small compared to the number n of agents and the number m of alternatives. This suggests the approach of applying notions from *parameterized complexity*. In particular, we say that a problem is *fixed parameter tractable* if, for a *parameter* k and any input instance (x, k) , there is a computable function $\phi(k)$ depending only on k such that the problem can be solved in time $\phi(k) \cdot |x|^{O(1)}$. The class FPT is then the class of fixed parameter tractable problems. The corresponding notions of hardness are given by the so-called W hierarchy (with $\text{FPT} = \text{W}[0]$). We omit the exact definition of this hierarchy and the required concept of *parameterized reductions* between problems which can, for example, be found in the book by Cygan et al. [54]. However, it is widely believed that, for any $\text{W}[i]$ -hard problem where $i \geq 1$, there is no function $\phi(k)$ such that an instance (x, k) can be solved in time $\phi(k) \cdot |x|^{O(1)}$.

At this point, we need to mention a closely related class of approval-based multiwinner voting rules that sidestep the apparent requirement to compute scores for an exponential number of committees. These are *sequential Thiele rules* which can be seen as greedy approximations for a subset of ABCS rules. A *Thiele rule* [120] is any ABCS rule defined by a scoring function f such that $f(x, y)$ depends only on x . That is, the size of any given approval vote does not affect the score. On input a profile P , a *sequential Thiele rule* builds the winning committee iteratively in k rounds starting with an empty set of alternatives. For $i \in \{0, 1, \dots, k-1\}$, let A_i be the set of alternatives that have been included in the committee until the end of iteration i . In the $(i+1)$ -th iteration, the rule selects any alternative for the committee that increases its score the most, that is, any $c \in \Sigma \setminus A_i$ such that $c \in \arg \max_{c' \in \Sigma \setminus A_i} \text{sc}_f(A_i \cup \{c'\}, P)$. The winning committee is then simply the set A_k .

We proceed to introduce some computational problems that have been considered

in the literature for ABCS rules and sequential Thiele rules.

Definition 3 (WINNERDETERMINATION (WD) problem). *Given a profile P , and a scoring function f , return any k -sized committee $C \subseteq \Sigma$ such that C is winning on P under the ABCS rule, respectively, the sequential Thiele rule specified by f .*

Notice that the WINNERDETERMINATION problem is polynomial-time solvable for any sequential Thiele rule.

Definition 4 (WINNERVERIFICATION (WV) problem). *Given a profile P , a k -sized committee $C \subseteq \Sigma$ and a scoring function f , decide whether C is a winning committee on P under the ABCS rule, respectively, the sequential Thiele rule specified by f .*

Definition 5 (CANDIDATEWINNER (CW) problem). *Given a profile P , an alternative $c \in \Sigma$ and a scoring function f , decide whether there is a winning committee C on P under the ABCS rule, respectively, the sequential Thiele rule specified by f such that $c \in C$.*

Definition 6 (WINNERTHRESHOLD (WT) problem). *Given a profile P , a scoring function f , and a threshold t , decide whether there exist a k -sized committee C such that $\text{sc}_f(C, P) \geq t$.*

In Table 4.1, we give an overview of known complexity results for the following popular ABCS rules (all of which are in fact Thiele rules): *approval voting* (AV), *proportional approval voting* (PAV), the *approval-based Chamberlin-Courant* (CC) rule, and its sequential variant CC_{seq} .

4.2 PAC Learnability of Voting Rules

In our data-driven approach to designing multiwinner voting rules, we make the assumption that data is available in the form of profiles labeled by sets of winning committees under some unknown *target rule*. The target rule however is known to be *some* ABCS rule (respectively, *some* sequential Thiele rule). Our work follows a standard approach in computational learning theory which assumes that there is a set Z of data points as well as a set Y of labels. A hypothesis class $\mathcal{H}$ consisting of functions $h : Z \rightarrow Y$ is known to contain a target function h^* which we would like to recover. To this end, a set T of points is sampled i.i.d. from Z according to some distribution D , and each point $z \in T$ is assigned the label $h^*(z)$. A *learning algorithm* then receives this *training set* $\{(z, h^*(z))\}_{z \in T}$ and returns a hypothesis $h \in \mathcal{H}$. Naturally, this function h should agree with the target function h^* as much as possible, including those data points that were not part of the training set. In other words, h should have small error under the error function

$$\text{err}(h) = \Pr_{z \sim D} [h(z) \neq h^*(z)].$$

rule	scoring function	hardness			
		WD	WV	CW	WT
AV	$f_{AV}(x, y) = x$	P	P	P	P
PAV	$f_{PAV}(x, y) = \sum_{i=1}^x \frac{1}{i}$	W[1]-h [18]	coW[1]-h [†]	Θ_2^P -h [119]	NP-c [112]
CC	$f_{CC}(x, y) = \mathbb{I}[x \geq 1]$				
CC _{seq}	$f_{CC}(x, y) = \mathbb{I}[x \geq 1]$	P	NP-h [‡]		

† Theorem 79, ‡ Theorem 84

Table 4.1: For a predicate X , we denote by $\mathbb{I}[X]$ the indicator function that returns 1 if X evaluates to true and 0 otherwise. For an alternative $a \in \Sigma$ and a profile $P = \{\sigma_i\}_{i \in \mathcal{N}}$, the *approval score* of a is $\sum_{i \in \mathcal{N}} \mathbb{I}[a \in \sigma_i]$. For the AV rule, the WINNER DETERMINATION problem can be solved efficiently by ordering the alternatives according to their approval score (breaking ties arbitrarily) and selecting the k alternatives with highest scores for the winning committee [94, Chapter 6.3]. The remaining problems can then be decided efficiently for the AV rule either by inspecting the previous ordering of the alternatives or by inspecting the score $sc_{f_{AV}}(C, P)$ of any winning committee C . With respect to the CC rule, the NP-hardness of WINNER DETERMINATION appears to be a folklore result in the computational social choice community where it is typically pointed out that this problem is equivalent to the MAXCOVER problem, e.g., see [118].

The notion of approximating the unknown target function based on this type of training data is captured by the *probably approximately correct* (PAC) learning model. We say that a hypothesis class $\mathcal{H}$ is PAC-learnable with *confidence* $\delta \in (0, 1)$, *accuracy* $\varepsilon \in (0, 1)$, and *sample complexity* $s(\delta, \varepsilon)$ if, for every distribution D over Z , and every $h^* \in \mathcal{H}$, there exists an algorithm $\mathcal{A}$ such that,

- on input a training set of size at least $s(\delta, \varepsilon)$ generated according to D and labeled by h^* ,
- the probability that $\mathcal{A}$ returns a hypothesis h such that $\text{err}(h) > \varepsilon$ is at most δ .

With respect to this definition of learnability, there are two questions which we would like to answer: First, how many samples are necessary and sufficient for learning? That is, we would like to find upper and lower bounds on the sample complexity $s(\delta, \varepsilon)$ of the hypothesis class $\mathcal{H}$ for any desired confidence and accuracy. In the literature on learning theory, there exist various measures for the *combinatorial richness* of different hypothesis classes. Intuitively, a hypothesis class that is rich in this sense has high expressive power. This expressiveness may allow many functions h to be fitted to any given training set. Thereby, accurate and correct learning becomes harder to achieve since many of these candidate hypotheses could have large error outside of this particular training set. Measures for the combinatorial richness of hypothesis classes include, for example, the VC-dimension [123] of boolean functions and the Natarajan dimension [105] of multiclass functions. Bounds on these measures imply bounds on the sample complexity of the respective hypothesis classes and

typically certify PAC learnability by so-called empirical risk minimizers (ERM). These learning algorithms seek to identify a rule $h \in \mathcal{H}$ that has minimal *empirical risk* on a given training set $\{(z, h^*(z))\}_{z \in T}$, for example, defined as the *0-1 loss* $\sum_{z \in T} \mathbb{I}[h(z) \neq h^*(z)]$ of h .

However, even if we had proven that a hypothesis class is PAC-learnable with—hopefully—low sample complexity, it is a separate question whether the class can be learned *efficiently*. Recall our assumption that the target rule h^* is known to be a member of the hypothesis class $\mathcal{H}$ which is referred to as the *realizable case* of PAC learning. In this case, we typically would like to define an efficient algorithm that always returns a hypothesis $h \in \mathcal{H}$ that assigns the correct labels to *all* points in any given training set. We say that such a hypothesis is *consistent* with the training set. Hence, another important goal is to show that there is such a learning algorithm with running time that is polynomial in n , m and the required sample size or, alternatively, to prove that the previous learning task is computationally hard.

The PAC learnability of voting rules was first considered by Procaccia et al. [113]. On a conceptional level, Procaccia et al. envision their pioneering work to enable the *automated design of voting rules*. They explore this idea in the setting of single-winner voting where the agents report *rankings* over the alternatives. One of the questions that Procaccia et al. investigate is whether the immensely popular class of *positional scoring rules* is efficiently PAC-learnable. We encountered these rules already in connection with the average-case model of Boutilier et al. [25] (Section 2.2) and used them ourselves for the proof that the average distortion of any voting rule must be high (Section 2.3). A positional scoring rule is specified by a non-increasing sequence of real-valued nonnegative numbers $\langle \alpha_1, \alpha_2, \dots, \alpha_m \rangle$. An alternative appearing in position j of an agent's ranking receives a score of α_j , and the winning alternative is an alternative with maximum total score. In their work, Procaccia et al. assume a particular way of tie-breaking between alternatives of maximum score. Hence, for the purpose of PAC learning this class of voting rules, the training set consists of preference profiles each of which is labeled by a *single* winning alternative. Notice that this means that there are m different labels. Procaccia et al. show that the class of positional scoring rules has low sample complexity. In particular, the number of samples sufficient for learning this class can be described by a low-degree polynomial depending on n , m , $1/\varepsilon$, and $\log(1/\delta)$. Their sample complexity results follow from bounds on the *generalized dimension* [106, Chapter 5] of this class of rules which the authors show to be $\Theta(m)$. Furthermore, Procaccia et al. demonstrate that *efficient* learning is possible by formulating, for any given training set, a linear program that describes feasible vectors $\alpha = \langle \alpha_1, \alpha_2, \dots, \alpha_m \rangle$ such that the corresponding scoring rules are consistent with the training set.

4.3 Highlights of our Results and Techniques

Our approach to the study of the PAC learnability of multiwinner voting rules is very much inspired by the previous work of Procaccia et al. [113]. However, we

need to point out an important conceptional difference: In our work, we do not make assumptions on tie-breaking between winning committees. As a consequence, the training sets that we consider consist of profiles (i.e., collections of approval votes) labeled by one *or more* winning committees. Hence, the number of different labels in our model is $2^{\binom{m}{k}} - 1$ whereas there are only m different labels in the model of Procaccia et al.

Nevertheless, we show that both the class of ABCS rules and the class of sequential Thiele rules have low sample complexity (Theorems 51 and 69). For ABCS rules, the number of samples sufficient for learning the unknown target rule grows polynomially in m , k , and $1/\varepsilon$ as well as logarithmically in $1/\delta$. The sample complexity of sequential Thiele can be upper-bounded in similar terms (in this case, depending only logarithmically on m), and we accompany our sample complexity upper bounds for both classes of voting rules by meaningful (although, not tight) lower bounds. We obtain our results by lower-bounding the *Natarajan dimension* and upper-bounding the so-called *graph dimension* of the respective classes which, by a theorem of Daniely et al. [56, restated by us as Theorem 50], imply our sample complexity bounds. In particular, our upper bounds on the graph dimension follow from the application of a beautiful result in algebraic combinatorics by Warren [124] and later refined by Alon [3] about the number of different sign patterns a set of linear functions may exhibit. Boutilier et al. [25] previously used the theorem by Alon for studying the PAC learnability of voting rules in a different social choice context (see our discussion in the next section).

Despite these positive results with respect to the sample complexity of ABCS and sequential Thiele rules, our work provides strong evidence that *efficient* PAC learning is unattainable. For the purpose of investigating the hardness of learning the respective classes from data, we introduce a decision problem which aims to capture the most elementary aspect of this task. Given a *single* profile P labeled by a *single* k -sized committee C , is there an ABCS rule (respectively, a sequential Thiele rule) such that C is among the winning committees on P under this rule? We call these decision problems TARGETABCS (Definition 58) and TARGETSEQTHIELE (Definition 67). Solving these problems appears to be a prerequisite to the typical ERM learning approach where the goal is to assign the correct label (which may consist of *multiple* winning committees in our setting) to *all* profiles in a given training set.

In Chapter 7, we first demonstrate that TARGETABCS is $\text{coW}[1]$ -hard (Theorem 60) when parameterized in the committee size k by a reduction from INDEPENDENTSET. Parameterized in the same way, TARGETSEQTHIELE is in fact easy in the sense that the problem is fixed parameter tractable (Theorem 68). We show that TARGETSEQTHIELE can be decided by checking a sequence of $k!$ linear programs (one program for every permutation of the alternatives in the given committee C) for feasibility. Yet, outside of the realm of parameterized complexity, TARGETSEQTHIELE is still NP-hard (Theorem 73) which we prove by a reduction from a structured variant of 3SAT, see Definition 72. The reductions that we use to obtain our two hardness results follow a similar approach: Given an instance of the respective hard input problem, we

carefully construct a profile P such that the only rule that can make certain committees of our choice winning is the (sequential) approval-based Chamberlin-Courant (CC) rule which is specified by the scoring function $f(x, y) = \mathbb{I}[x \geq 1]$. Finally, we show that a specific committee is indeed winning on P under the CC rule if and only if the input to the reduction is a YES-instance of the respective problem.

Why is it that our reductions enforce the application of the CC rule? The parts of our profiles which are more immediately related to the structure of the given input instance are in fact approval votes of size two. It is then natural to represent the alternatives as the nodes in a graph G . An approval vote $\sigma = \{a, b\}$ corresponds to an undirected edge in G between the nodes a, b . In our $\text{coW}[1]$ -hardness proof for TARGETABCS , we reduce from INDEPENDENTSET , and also with respect to the NP-hardness of TARGETSEQTHIELE we give a graph-based intuition for our reduction (see Figure 7.1). In both cases, enforcing the application of the CC rule allows us to more easily control the score that any committee may receive on a profile P . More specifically, assume that an alternative a is a member of committee A . The score of A under the CC rule is increased by one for every edge that is incident to a in G . However, also including the other endpoint b of an edge $\sigma = \{a, b\}$ in A does not further increase the score of A under the CC rule. This intuition is particularly helpful for the treatment of sequential Thiele rules: Consider the previous graph G . In every iteration, the sequential CC rule includes an alternative in the winning committee whose corresponding node currently has maximum degree in G . Selecting an alternative a for the winning committee under the sequential CC rule is equivalent to removing a together with all incident edges from G . We can then proceed to repeat the previous step of identifying a maximum degree node. In our NP-hardness proof for TARGETSEQTHIELE , we show that the only way of making a particular committee A winning requires selecting its members under the sequential CC rule in a specific order (removing nodes and incident edges in G along the way). For our reduction, we are able to relate this order to the existence of a satisfying assignment to the variables of a given structured 3-CNF formula.

4.4 Open Problems and Discussion

Our results for the sample complexity of ABCS and sequential Thiele rules demonstrate that PAC-learning these classes of voting rules is in fact possible. On the other hand, our hardness results suggest that learning related tasks cannot be performed efficiently, at least for worst-case instances. However, the latter finding does not rule out the usefulness of PAC-learning for practical purposes. In particular, it seems worthwhile to investigate the performance of empirical risk minimizing (ERM) learners on synthetic or real-world multiwinner election data [23, 61]. Regarding our case of multiclass learning where there exist exponentially many different labels for the data points, the work of Daniely and Shalev-Shwartz [55] offers learning theoretical insights on compression-based optimal learners. Could our fixed parameter tractable method for fitting a sequential Thiele rule to a single sample (Theorem 68) already

provide a starting point for practical ERM learning? Besides the *0-1 loss function* which counts the number of samples on which a hypothesis h disagrees with the unknown target function h^* , we could consider other loss functions such as the Hamming distance between committees. Designing different algorithms that fit a voting rule as best as possible (as opposed to perfectly) to the provided training data under a particular loss function appears to be an interesting technical challenge. The work of Faliszewski et al. [66] on optimizing multiwinner voting rules to resemble what they call *utopias* as close as possible and their application of local search algorithms may serve as another source of inspiration.

In our work, we consider scoring rules that are applicable to *approval* ballots. Scoring rules are however also a popular choice for aggregating ranking-based preferences in multiwinner voting. Faliszewski et al. [65] define a whole hierarchy of multiwinner scoring rules for these types of preferences. This hierarchy could also be investigated with respect to PAC learnability and the efficiency of learning-related tasks. Are the classes on different levels of the hierarchy also separable in terms of their learnability? Our techniques could serve as a starting point to explore this question.

The PAC learnability of multiwinner voting rules could also be studied in utilitarian settings. Previously, Boutilier et al. [25] investigated the possibility of PAC-learning positional scoring rules¹ in the model of single-winner voting where the agents have utilities for the alternatives. They found that this class of voting rules has low sample complexity. With respect to learning algorithms, their training data consists of sampled utilities and preference profiles (i.e., the agents' rankings) consistent with these utilities. In our notation (Section 2.1), such a t -sized training set would have the form $\{(P_i, v_i)\}_{i \in [t]}$ where v_i are valuations and $P_i \in \mathcal{P}(v_i)$. Let $\mathcal{R}_\alpha$ be the positional scoring rule defined by a scoring vector α . The ERM learning approach (or, in the authors' words, the *sample-based optimization approach*) by Boutilier et al. then seeks to compute a scoring vector α such that

$$\alpha \in \arg \max_{\alpha' \in \mathbb{R}^m} \left\{ \frac{1}{t} \cdot \sum_{i=1}^t \text{SW}(\mathcal{R}_{\alpha'}(P_i), v_i) \right\}.$$

A similar approach could be explored for ranking-based multiwinner scoring rules. Furthermore, Pierczynski and Skowron [108] introduced a model of approval preferences in metric voting. Here, the agents N and alternatives A share a metric space $(N \cup A, d)$ under a distance function d . Each agent's approval ballot is determined by an *acceptability function* $\lambda : N \rightarrow 2^A$. In particular, Pierczynski and Skowron consider functions such that, for an agent i , $\lambda(i)$ contains all alternatives that lie within a ball of a certain *acceptability radius* centered at the location of i . Let us denote the resulting profile by $P(\lambda, d)$. We could now investigate the learnability of, e.g., ABCS rules, given different forms of training data. For example, assume that both

¹We gave a definition of positional scoring rules in Section 4.2. In fact, the work of Boutilier et al. [25] is applicable to a broader class of scoring rules which allows for randomization. We omit these details in the following.

the target ABCS rule and the correct choice of acceptability radii are unknown for a given application. A t -sized training set could have the form $\{(d_i, f^*(P(\lambda, d_i)))\}_{i \in [t]}$ where $f^*(P(\lambda, d_i))$ is the set of winning committees on the profile $P(\lambda, d_i)$ under the unknown target ABCS rule specified by the scoring function f^* . Is it possible to approximate the unknown target ABCS rule together with the radii λ from few samples? Can it be done efficiently, for example, by an ERM learner under the 0-1 loss function? Other types of training data could be considered together with an ERM objective which—similar to the above approach of Boutilier et al. [25]—could seek to minimize the average social cost incurred on the training data. Finally, the assumption of utilities (respectively, costs in the metric setting) again allows for the definition of *distortion* (see Sections 2.1 and 3.1) in the respective models. Then, by taking the average over all profiles corresponding to the points in the dataset which we sampled our training data from, the unknown target voting rule exhibits a certain distortion. Assuming that we learned a voting rule from data, how does the rule differ from the target rule in terms of this notion of average-case distortion?

Part II

Publications

Chapter 5

Beyond the Worst Case: Distortion in Impartial Culture Electorates

Distortion is a well-established notion for quantifying the loss of social welfare that may occur in voting. As voting rules take as input only ordinal information, they are essentially forced to neglect the exact values the agents have for the alternatives. Thus, in worst-case electorates, voting rules may return low social welfare alternatives and have high distortion. Accompanying voting rules with a small number of cardinal queries per agent may reduce distortion considerably.

To explore distortion beyond worst-case conditions, we use a simple stochastic model according to which the values the agents have for the alternatives are drawn independently from a common probability distribution. This gives rise to so-called *impartial culture electorates*. We refine the definition of distortion so that it is suitable for this stochastic setting and show that, rather surprisingly, all voting rules have high distortion *on average*. On the positive side, for the fundamental case where the agents have random *binary* values for the alternatives, we present a mechanism that achieves approximately optimal average distortion by making a *single* cardinal query per agent. This enables us to obtain slightly suboptimal average distortion bounds for general distributions using a simple randomized mechanism that makes one query per agent. We complement these results by presenting new tradeoffs between the distortion and the number of queries per agent in the traditional worst-case setting.

5.1 Introduction

Voting has been the subject of social choice theory for centuries. Traditionally, having elections as the main application area in mind, social choice theorists have made a lot of progress in understanding the axiomatic properties of *voting rules* (e.g., see Zwicker [128] for an introduction to basic voting axioms). However, the use of voting goes far beyond elections, and today, it provides a compelling way of collective decision-making. So, in addition to the traditional axiomatic treatment of voting rules, other

approaches that have an optimization flavour have attracted a lot of attention, starting with the work of Young (see Young [127] and references therein).

Among others, the utilitarian approach [25] assumes that voters have values for the alternatives. The preferences expressed in the ballot of a voter are a very short summary of these values, e.g., a ranking of the candidates in terms of their values, a set of candidates with values exceeding a threshold, or just the alternative with the highest value for the voter. A voting rule takes as input the voters' ballots and computes an outcome (e.g., a winning alternative). As it has no access to the underlying values of the voters for the alternatives, the outcome depends only on the short summaries in the voters' ballots. Rather unsurprisingly, the outcome of the voting rule will be sub-optimal if evaluated in terms of the voters' values for the alternatives.

How far can the outcome be from the optimal alternative? To answer this question, we need to specify a measure to assess the quality of alternatives. Following a rather standard approach in the literature, we use the *social welfare*—the total value the agents have for an alternative—as our quality measure. Then, the following question arises naturally. How far from the optimum can the outcome of a voting rule be in terms of social welfare? The notion of *distortion*, introduced by Procaccia and Rosenschein [111], comes to answer this question. The distortion of a voting rule is defined as the worst-case ratio, among all voting profiles on a set of alternatives, between the maximum social welfare among all alternatives and the social welfare of the winning alternative returned by the voting rule.

The distortion of voting rules has been the subject of a long list of papers in computational social choice for more than ten years now; see the nice survey of the key results in the area by Anshelevich et al. [11]. For example, under mild assumptions about the valuations, the ubiquitous plurality rule has a distortion of $O(m^2)$, where m is the number of alternatives [33]. Nevertheless, the distortion is inherently high. Even when the values for the alternatives are restricted (e.g., the total value of each agent for all alternatives is normalized to 1), the distortion of any (possibly randomized) voting rule can be as bad as $\Theta(\sqrt{m})$ [25, 59].

A recent approach by Amanatidis et al. [4] aims to bypass such lower bounds by making very limited use of the underlying cardinal information. Here, besides the ranking submitted as a ballot by an agent, the voting rule (or, better, the *mechanism*) can pose queries to the agent regarding her value for particular alternatives. Even though optimal distortion results are now possible (by naively querying the values of all alternatives in each agent's ranking), the important problem to be solved is to design mechanisms that achieve low distortion by making a limited number of queries per agent. Among other results, Amanatidis et al. [4] present an algorithm that achieves constant distortion by making at most $O(\log^2 m)$ queries per agent. On the negative side, they show that any mechanism that achieves constant distortion must make at least $\Omega\left(\frac{\log m}{\log \log m}\right)$ per agent. Both results refer to deterministic mechanisms.

By the definition of distortion, the results discussed above are of a worst-case flavour. Voting rules and mechanisms are evaluated on a profile for which they have the worst possible performance, no matter how frequently such a profile may appear

in practice. As such, they may not be suitable to explain the success of certain voting rules in practice.

Overview of Our Contribution

We follow the utilitarian framework but—motivated by the recent trend of analyzing algorithms from non-worst-case perspectives [115]—attempt an evaluation of voting rules on an *average-case* basis. At the conceptual level, we introduce the notion of *average distortion* for such stochastic preferences. To this end, we consider a simple stochastic setting in which each agent draws a random value for each alternative according to a common probability distribution. The draws of each agent for each alternative are independent. Naturally, *impartial culture electorates* [110, 122]—i.e., profiles with uniformly random rankings of alternatives—emerge in this way. The average distortion of a mechanism is defined as the expected maximum social welfare among the alternatives over the expected social welfare of the alternative returned by the mechanism. To distinguish between the two notions, we use the term *worst-case distortion* to refer to the traditional definition. It is not hard to see that the average distortion is always upper-bounded by the worst-case distortion.

We warm up by showing (in Section 5.3) that, perhaps surprisingly, any (potentially randomized) voting rule has average distortion $\Omega(m)$, implying that the trivial voting rule that selects one of the alternatives uniformly at random (ignoring the profile) has an almost best possible average distortion. Our lower bound uses a very simple *binary* probability distribution (i.e., one that returns values of 1 and 0). In light of this seemingly discouraging result, our primary objective is to understand the additional power a limited number of value queries can give to a voting rule by seeking answers to the following questions:

Can we design mechanisms that use a few (ideally, a constant number of) queries per agent and have low (ideally, constant) average distortion?
What conditions do we need to impose on the distribution from which the agents' valuations are drawn in order to achieve this?

Notice that we do not intend to introduce other restrictions on the electorates, such as requiring the number of agents n to be large compared to the number of alternatives m . Instead, we would like our mechanisms to work for *any* choice of n and m .

Our first positive result is for the family of binary distributions. In Section 5.4, we present a deterministic mechanism, called MEAN, which makes a *single query* per agent and achieves *constant* average distortion. This is our most technically involved result and indicates that using queries (even in a minimal way) can lead to a significant improvement in average distortion.

Binary distributions are, of course, an exceedingly crude model for the agents' valuations. However, we demonstrate the usefulness of our result by employing the mechanism MEAN as a building block for a randomized mechanism, which works for a general distribution F . This mechanism, which we call RTMEAN, still requires only

one query per agent and achieves an average distortion of $O\left(\log m + \log \frac{\sigma^2}{\mu^2}\right)$, where μ and σ^2 are the mean and variance of the distribution F . For many distributions of interest, the average distortion bound obtained is only logarithmic in m . A similar idea is used to define our randomized mechanism **RTSEARCH** for the traditional model of worst-case distortion. **RTSEARCH** uses a logarithmic number of queries per agent and achieves logarithmic worst-case distortion. Such an upper bound is not known for deterministic mechanisms. These results are presented in Section 5.5. To the best of our knowledge, this is the first analysis of randomized mechanisms that make value queries in the distortion literature.

Our upper bound on the average distortion of mechanism **MEAN** is not attainable by deterministic mechanisms in the traditional worst-case model. We prove that no deterministic mechanism that makes a single query per agent can achieve distortion better than $\Omega(\sqrt{m})$, even when the valuations are binary. More importantly, for general valuations, we present a new lower bound of $\Omega(\log m)$ on the number of queries per agent that allows for constant worst-case distortion. This improves the previously best-known lower bound of Amanatidis et al. [4] by a sublogarithmic factor. These results appear in Section 5.6.

Further Related Work

Using different methodological approaches, utilities in voting have been considered in social choice theory since the work of Bentham [21] in the 18th century; for recent work on the topic see Apesteguia et al. [13], Pivato [109]. The stochastic model we follow was introduced by Boutilier et al. [25]. Among other investigations, Boutilier et al. aimed at identifying the optimal voting rule in terms of the expected social welfare of the winning alternative. Their main conclusion is that this voting rule is a positional scoring rule with parameters depending on the probability distribution the agents' values are drawn from. They do not consider notions related to average distortion and do not present related bounds. Gonczarowski et al. [77] define a version of distortion (slightly different from ours) in the model of Boutilier et al. as follows. For a particular profile, they examine the ratio between the expected maximum social welfare and the expected social welfare of the winning alternative, both terms conditioned on the random values of the alternatives being consistent with the profile. They define the expected distortion as the worst-case ratio over all profiles. It can be verified that, for every voting rule, the value of expected distortion according to Gonczarowski et al. [77] is bounded above and below by the worst-case distortion and our average distortion, respectively. The main result of Gonczarowski et al. [77] is that binomial voting (which, in the terminology of Boutilier et al. [25], is essentially the average-case optimal voting rule for a fair Bernoulli trial) approximates the rule with minimum expected distortion for values drawn from distributions supported on $[0, 1]$. Both Boutilier et al. [25] and Gonczarowski et al. [77] restrict their attention to voting rules and do not consider value queries to the agents.

In the current paper, we assume that agents have non-negative values for the

alternatives. In a series of recent papers on *metric voting*, originating from the work of Anshelevich et al. [10], agents and alternatives are assumed to be located in a metric space, and the preferences of each agent reflect her relative distance from the alternatives. In that different setting, the distortion quantifies the suboptimality of the outcome of voting rules in terms of the *social cost*. This line of research has led to challenging algorithmic problems with beautiful solutions [42, 43, 74, 88]. Cheng et al. [45, 46] and Ghodsi et al. [72] consider variants of average-case distortion that, due to the metric setting and additional modelling assumptions, are only distantly related to ours.

The study of the utilitarian framework in the recent CS literature goes beyond the study of single-winner voting. Caragiannis et al. [35, 38] and Benadè et al. [20] consider multiwinner voting and participatory budgeting settings, respectively. Filos-Ratsikas et al. [68] investigate voting in distributed settings. The utilitarian approach has also been used in settings that are more general than voting. The main idea is to explore how well algorithms that use only ordinal information about the underlying input (as opposed to cardinal values) can approximate the optimal solutions of combinatorial optimization problems. The survey of Anshelevich et al. [11] covers early work on this hot topic, as well as on the other directions discussed above. The benefits of allowing for limited access to the cardinal information were more recently explored for matching [5, 6, 95, 98] and clustering [28] problems.

We remark that stochastic models are otherwise ubiquitous in the EconCS literature, including a long line of research in mechanism design and auctions originating from the seminal paper of Myerson [104], and a rich menu of other settings like matchings [44], fair division [99], kidney-exchange [121], and many more.

5.2 Preliminaries

Throughout the paper, we denote by N and A the sets of agents and alternatives and reserve n and m for their cardinalities, respectively. Each agent $i \in N$ has a *ranking* $\succ_i$ of the alternatives, i.e., a strict ordering of the elements in A . A *profile* $P = \{\succ_i\}_{i \in N}$ is just a collection of the agents' rankings. We denote by $\mathcal{P}$ the set of all possible profiles with n agents and m alternatives. A voting rule $\mathcal{M} : \mathcal{P} \rightarrow A$ takes as input a profile and returns a single winning alternative.

We assume that the ranking of each agent results from underlying hidden non-negative *values* that the agent has for each alternative. For $i \in N$, the *valuation function* $\text{val}_i : A \rightarrow \mathbb{R}_{\geq 0}$ returns the values of agent i for the alternatives in A . Then, agent i 's ranking $\succ_i$ is *consistent* with val_i . This means that $a \succ_i a'$ only if $\text{val}_i(a) \geq \text{val}_i(a')$ for every pair of alternatives $a, a' \in A$. Given a ranking $\succ$, the function $\text{pos}_\succ : A \rightarrow [m]$ returns the position of a given alternative in the ranking. We say that alternative a is the *top-ranked* alternative of agent i if $\text{pos}_{\succ_i}(a) = 1$.

Consider a set of valuation functions $v = \{\text{val}_i\}_{i \in N}$. We typically refer to the set v as the agents' *valuations*. The *social welfare* of an alternative $a \in A$ is defined as

$$\text{SW}(a, v) = \sum_{i \in N} \text{val}_i(a).$$

Let $\mathcal{V}$ be the set of all possible valuations the agents in N can have for the alternatives in A . We denote by $\mathcal{P}(\text{val}_i)$ the set of rankings that are consistent with the valuation function val_i of agent i . We say that a profile $P = \{\succ_i\}_{i \in N}$ is consistent with the valuations v if the ranking $\succ_i$ of every agent $i \in N$ is consistent with her valuation function val_i , i.e., $\succ_i \in \mathcal{P}(\text{val}_i)$. Then, $\mathcal{P}(v)$ represents the set of profiles in $\mathcal{P}$ which are consistent with the valuations v .

Besides voting rules, we consider *mechanisms* that have, in addition to the profile, access to the valuations. Such a mechanism $\mathcal{M} : \mathcal{P} \times \mathcal{V} \rightarrow A$ takes as input the profile and the valuations and returns a winning alternative. We are particularly interested in mechanisms that use the whole profile on input but only a small part of the valuations by making a limited number of *queries* per agent. A query for the value of agent $i \in N$ for alternative $a \in A$ simply returns the value of $\text{val}_i(a)$.

The *worst-case distortion* of a mechanism $\mathcal{M}$ applied on profiles consistent with valuations from $\mathcal{V}$ is defined as

$$\text{dist}(\mathcal{M}) = \sup_{\substack{v \in \mathcal{V} \\ P \in \mathcal{P}(v)}} \frac{\max_{a \in A} \text{SW}(a, v)}{\text{SW}(\mathcal{M}(P, v), v)},$$

i.e., it is the worst-case ratio—among all valuations and consistent profiles—between the maximum social welfare and the social welfare of the alternative returned by the mechanism. For randomized mechanisms, in which the alternative returned is a random variable, we use the expectation $\mathbb{E}[\text{SW}(\mathcal{M}(P, v), v)]$ (taken over the random choices of $\mathcal{M}$) instead of $\text{SW}(\mathcal{M}(P, v), v)$ in the denominator.

We extend the notion of distortion to *stochastic environments* with random valuations and consistent profiles. We assume that the values of the agents for the alternatives are drawn from a *known* common distribution F . In particular, for each agent $i \in N$ and alternative $a \in A$, the value $\text{val}_i(a)$ is drawn independently from distribution F . Given valuations v selected in this way, a consistent profile P is selected uniformly at random among all profiles in $\mathcal{P}(v)$ (essentially, in her ranking, each agent breaks ties among alternatives in terms of value uniformly at random). This gives rise to uniformly random profiles, which are known as *impartial culture electorates* in the social choice theory literature.

For valuations v , we use $v \sim F$ to denote that, for every alternative $a \in A$ and every agent $i \in N$, the value $\text{val}_i(a)$ is drawn independently from F . We use the notation $\succ_i \sim \mathcal{P}(\text{val}_i)$ to refer to a ranking that is selected uniformly at random among all rankings that are consistent with the valuation function val_i of agent i . Similarly, we use $P \sim \mathcal{P}(v)$ for a profile that is selected uniformly at random among all profiles that are consistent with valuations v . Then, the *expected social welfare* of the winning

alternative picked by a deterministic mechanism $\mathcal{M}$ is

$$\mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\text{SW}(\mathcal{M}(P, v), v)].$$

When mechanism $\mathcal{M}$ is randomized, the expectation is taken over the randomness of $\mathcal{M}$ as well.

We now adapt the notion of distortion to this stochastic setting. The *average distortion* of a mechanism $\mathcal{M}$ on a family of distributions $\mathcal{F}$ is

$$\text{avdist}(\mathcal{M}, \mathcal{F}) = \sup_{F \in \mathcal{F}} \frac{\mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \right]}{\mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\text{SW}(\mathcal{M}(P, v), v)]}.$$

We remark that the trivial voting rules that return a fixed alternative or an alternative selected uniformly at random have average distortion at most m for every distribution F . Denoting them by $\mathcal{M}$, this is due to the simple fact that

$$\mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\text{SW}(\mathcal{M}(P), v)] = \frac{1}{m} \cdot \mathbb{E}_{v \sim F} \left[\sum_{a \in A} \text{SW}(a, v) \right] \geq \frac{1}{m} \cdot \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \right].$$

For particular distributions, the average distortion can be much lower. As an example, consider the uniform distribution over the interval $[a, b]$ with $a, b \geq 0$. Then, the expected social welfare of the alternative returned by both trivial rules is $n \cdot (a + b)/2$ while the expected maximum social welfare (over all alternatives) cannot exceed $n \cdot b$. Hence, the average distortion is at most 2 in this case.

A fundamental family that is of central importance to our work is the family $\mathcal{F}^{0/1}$ of *binary distributions*, consisting of the set of probability distributions $\{F_p\}_{p \in [0,1]}$, where F_p is such that the random variable z that is drawn according to F is equal to 1 with probability p and to 0 with probability $1 - p$.

5.3 Average Distortion Can Be High

We begin our technical exposition with a lower bound of $\Omega(m)$ on the average distortion. Essentially, Theorem 7 implies that the trivial voting rules that return a fixed alternative or an alternative selected uniformly at random on every profile have approximately optimal average distortion.

In our proof, we make use of a particular distribution F from the family $\mathcal{F}^{0/1}$. Furthermore, we exploit a class of voting rules called *positional scoring rules*. A positional scoring rule g is defined by a scoring vector $\langle \alpha_1, \alpha_2, \dots, \alpha_m \rangle$ with $\alpha_1 \geq \alpha_2 \geq \dots \geq \alpha_m \geq 0$. On input a profile P , the voting rule g assigns a score $\text{sc}_g(a, P)$ to every alternative $a \in A$ and the winning alternative is the one with the highest score (breaking ties according to some tie-breaking rule). The score of an alternative is computed as follows. For $j = 1, \dots, m$, alternative a takes α_j points each time it appears in the j -th position in an agent's ranking, i.e., $\text{sc}_g(a, P) = \sum_{i \in N} \alpha_{\text{pos}_{\prec_i}(a)}$. For

example, *plurality* is the positional scoring rule that uses the m -entry scoring vector $\langle 1, 0, \dots, 0 \rangle$. The plurality winner on profile P , denoted by $\text{plu}(P)$, is the alternative that appears most often in the top position of the agents' rankings.

Theorem 7. *For every (possibly randomized) mechanism $\mathcal{M}$, $\text{avdist}(\mathcal{M}, \mathcal{F}^{0/1}) \in \Omega(m)$.*

Proof. Consider impartial culture electorates with m alternatives, $n \geq 6m \ln m$ agents, and binary values drawn from the probability distribution $F \in \mathcal{F}^{0/1}$ with $p = \frac{1}{nm}$. Then, the probability that some alternative has positive social welfare is

$$1 - \left(1 - \frac{1}{nm}\right)^{nm} \geq 1 - e^{-1}.$$

Thus, $\mathbb{E}_{v \sim F}[\max_{a \in A} \text{SW}(a, v)] \geq 1 - e^{-1}$. Now consider any voting rule $\mathcal{M}$; we will complete the proof by showing that

$$\mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}}[\text{SW}(\mathcal{M}(P), v)] \leq \frac{4}{m}. \quad (5.1)$$

Below, we prove inequality (5.1), assuming that $\mathcal{M}$ is deterministic. For randomized $\mathcal{M}$, the expectation in the LHS of (5.1) should also be taken over the randomness of $\mathcal{M}$ as well. In that case, inequality (5.1) follows trivially by our arguments, interpreting $\mathcal{M}$ as a probability distribution over deterministic voting rules.

Let g be the positional scoring rule that uses the scoring vector $\langle \alpha_1, \alpha_2, \dots, \alpha_m \rangle$ with

$$\alpha_j = \mathbb{E}_{\substack{\text{val}_i \sim F \\ \succ_i \sim \mathcal{P}(\text{val}_i)}}[\text{val}_i(a) | \text{pos}_{\succ_i}(a) = j],$$

i.e., score α_j is equal to the expected value according to F given by each agent to the alternative ranked at position j .¹ Now, observe that for a given profile $P = \{\succ_i\}_{i \in N}$ and alternative $a \in A$, we have

$$\begin{aligned} \mathbb{E}_{v \sim F}[\text{SW}(a, v) | P \in \mathcal{P}(v)] &= \sum_{i \in N} \mathbb{E}_{\text{val}_i \sim F}[\text{val}_i(a) | \succ_i \in \mathcal{P}(\text{val}_i)] \\ &= \sum_{i \in N} \alpha_{\text{pos}_{\succ_i}(a)} \\ &= \text{sc}_g(a, P), \end{aligned}$$

where $\text{sc}_g(a, P)$ denotes the score of alternative a in profile P according to positional scoring rule g . Thus,

$$\begin{aligned} \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}}[\text{SW}(\mathcal{M}(P), v)] &= \sum_{P \in \mathcal{P}} \mathbb{E}_{v \sim F}[\text{SW}(\mathcal{M}(P), v) | P \in \mathcal{P}(v)] \cdot \Pr_{v \sim F}[P \in \mathcal{P}(v)] \\ &= \sum_{P \in \mathcal{P}} \text{sc}_g(\mathcal{M}(P), P) \cdot \Pr_{v \sim F}[P \in \mathcal{P}(v)] \\ &= \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}}[\text{sc}_g(\mathcal{M}(P), P)]. \end{aligned} \quad (5.2)$$

¹Boutilier et al. [25] refer to this positional scoring rule g as an average-case optimal social choice function.

Notice that, by the definitions of voting rules plu and g , for any profile P and alternative $a \in A$, it holds

$$\text{sc}_g(a, P) \leq \alpha_1 \cdot \text{sc}_{\text{plu}}(a, P) + n \cdot \alpha_2 \leq \alpha_1 \cdot \text{sc}_{\text{plu}}(\text{plu}(P), P) + n \cdot (mp - \alpha_1). \quad (5.3)$$

The first inequality follows since g gives α_1 points to alternative a for each of its $\text{sc}_{\text{plu}}(a, P)$ appearances in the top position in the agents' rankings and at most α_2 points for its remaining appearances. The second inequality follows since the plurality winner $\text{plu}(P)$ has the highest plurality score (and, hence, $\text{sc}_{\text{plu}}(a, P) \leq \text{sc}_{\text{plu}}(\text{plu}(P), P)$ for every alternative $a \in A$) and since $\alpha_2 \leq \sum_{j=1}^m \alpha_j - \alpha_1$ and $\sum_{j=1}^m \alpha_j = mp$ (as the expected total value given to all alternatives by an agent).

We will also need the following claim.

Claim 8. $\mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\text{sc}_{\text{plu}}(\text{plu}(P), P)] \leq \frac{3n}{m}$.

Proof. The proof will follow by a simple application of the Chernoff bound, e.g., see Motwani and Raghavan [102].

Lemma 9 (Chernoff bound, upper tail). *For every binomial random variable Q and any $\delta > 0$, we have*

$$\Pr[Q \geq (1 + \delta) \mathbb{E}[Q]] \leq \exp\left(-\frac{\delta^2 \mathbb{E}[Q]}{2 + \delta}\right).$$

Consider an alternative $a \in A$ and denote by X_i the random variable indicating whether agent i ranks a first ($X_i = 1$; this happens with probability $1/m$) or not ($X_i = 0$). Observe that $\text{sc}_{\text{plu}}(a, P) = \sum_{i \in N} X_i$, i.e., $\text{sc}_{\text{plu}}(a, P)$ is a binomial random variable with expectation $\frac{n}{m}$. By applying Lemma 9 with $\delta = 1$, we obtain that $\Pr[\text{sc}_{\text{plu}}(a, P) \geq \frac{2n}{m}] \leq \exp(-\frac{n}{3m}) \leq \frac{1}{m^2}$. The last inequality follows since $n \geq 6m \ln m$. Taking the union bound over the m alternatives then gives $\Pr[\text{sc}_{\text{plu}}(\text{plu}(P), P) \geq \frac{2n}{m}] \leq \frac{1}{m}$. Finally, since the plurality score never exceeds n , we have $\mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\text{sc}_{\text{plu}}(\text{plu}(P), P)] \leq \frac{2n}{m} \cdot (1 - \frac{1}{m}) + n \cdot \frac{1}{m} \leq \frac{3n}{m}$. $\square$

Notice that the top alternative of an agent has value 0 with probability $(1 - p)^m$ and value 1 otherwise. Thus, $\alpha_1 = 1 - (1 - p)^m$. Now, observe that $1 - (1 - p)^m = 1 - (1 - \frac{1}{nm})^m \geq 1 - \frac{1}{e^{1/n}} \geq \frac{1}{n+1}$, using the properties $(1 - r/t)^t \leq e^{-r}$ for $t > 0$ and $r \geq 0$ and $e^r \geq 1 + r$. Thus, $mp - \alpha_1 \leq \frac{1}{n} - \frac{1}{n+1} < \frac{1}{n^2}$. Also, using the property $(1 - r)^t \geq 1 - rt$ for $t \geq 1$, we get $\alpha_1 = 1 - (1 - p)^m \leq pm = \frac{1}{n}$. Combining these two last inequalities with equations (5.2) and (5.3) and Claim 8, we obtain

$$\begin{aligned} \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\text{SW}(\mathcal{M}(P), v)] &= \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\text{sc}_g(\mathcal{M}(P), P)] \\ &\leq \alpha_1 \cdot \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\text{sc}_{\text{plu}}(\text{plu}(P), P)] + n(mp - \alpha_1) \\ &\leq \frac{3}{m} + \frac{1}{n} \leq \frac{4}{m}, \end{aligned}$$

as desired by (5.1). The last inequality follows from the relation between n and m . $\square$

5.4 Constant Average Distortion with a Single Query per Agent

In the previous section, we saw that the average distortion can be high even for binary distributions if the mechanism has no access to the agents' underlying valuations. Still, binary valuations may reveal a lot of information with a single query. Consider a query for the value that an agent has for the alternative at position j of her ranking. If the query returns a value of 1, this means that the agent has a value of 1 for all alternatives in positions higher than j as well. This observation motivates the following definition of *implied social welfare*.

Definition 10 (implied social welfare). *For a distribution $F \in \mathcal{F}^{0/1}$, assume that a mechanism queried each agent for the value of the alternative in position k of her ranking. For each agent $i \in N$, let $\text{val}_{i,k}$ denote the respective value. The implied social welfare then is*

$$\text{SW}_{\text{im}}(a, P, v, k) = \sum_{i \in N} \mathbb{I}\{\text{val}_{i,k} = 1 \wedge \text{pos}_{\succ_i}(a) \leq k\}.$$

We use the notion of implied social welfare in the definition of the following mechanism.

Definition 11 (mechanism MEAN). *Given a profile P with underlying valuations drawn from a distribution $F_p \in \mathcal{F}^{0/1}$, the mechanism MEAN queries each agent for the value of the alternative at position $\tau = \max\{1, \lfloor pm \rfloor\}$. The mechanism then returns the alternative that maximizes the implied social welfare, that is,*

$$\text{MEAN}(P, v) \in \arg \max_{a \in A} \text{SW}_{\text{im}}(a, P, v, \tau),$$

breaking ties arbitrarily.

We show that mechanism MEAN has constant average distortion with a single query for the family of binary distributions.² Notably, it is impossible for deterministic 1-query mechanisms to achieve similar guarantees in the traditional setting of worst-case distortion. Indeed, even for binary valuations, the worst case distortion of these mechanisms is $\Omega(\sqrt{m})$; see Theorem 24.

Theorem 12. *Mechanism MEAN has average distortion at most 27 in impartial culture electorates with n agents and m alternatives, and underlying values drawn from any probability distribution in $\mathcal{F}^{0/1}$.*

²One may wonder whether the simplification of mechanism MEAN, which always queries the value of the top-ranked alternative in each agent and returns the alternative that maximizes the implied social welfare, achieves a constant average distortion as well. We discuss this question in Section 5.8 and show that, unless we introduce a more sophisticated tie-breaking, this variation has average distortion at least $\Omega\left(\frac{\log m}{\log \log m}\right)$.

Proof. We partition family $\mathcal{F}^{0/1}$ into the subfamilies of probability distributions $\mathcal{F}_1$, $\mathcal{F}_2$, and $\mathcal{F}_3$ defined as follows:

$$\begin{aligned}\mathcal{F}_1 &= \{F_p \in \mathcal{F}^{0/1} : p \geq 1/m\} \\ \mathcal{F}_2 &= \{F_p \in \mathcal{F}^{0/1} : 1 - (1 - 1/n)^{1/m} \leq p < 1/m\} \\ \mathcal{F}_3 &= \{F_p \in \mathcal{F}^{0/1} : p < 1 - (1 - 1/n)^{1/m}\}\end{aligned}$$

Let $F_p \in \mathcal{F}^{0/1}$; we will prove the theorem by distinguishing between the three cases $F_p \in \mathcal{F}_1$, $F_p \in \mathcal{F}_2$, and $F_p \in \mathcal{F}_3$.

We introduce some notation that we use throughout this proof. Let τ denote the position that mechanism MEAN queried in each agent's ranking, i.e., $\tau = \max\{1, \lfloor mp \rfloor\}$. For valuations v , let $X_i(v)$ be the number of alternatives for which agent i draws a value of 1 and let $X(v) = \sum_{i \in N} X_i(v)$. We will consider valuations v drawn from the probability distribution F_p . Thus, $X_i(v)$ is a random variable following the binomial distribution with m trials and success probability p . Since $X(v)$ is the sum over i.i.d. binomially distributed random variables, $X(v)$ itself follows a binomial distribution with nm trials and success probability p .

Furthermore, we denote by $Z_i(v, \tau)$ the random variable indicating whether the query at position τ of agent i 's ranking returned a value of 1 (then, $Z_i(v, \tau) = 1$) or not (then, $Z_i(v, \tau) = 0$). Clearly, $Z_i(v, \tau) = \mathbb{I}\{X_i(v) \geq \tau\}$ such that the random variable $Z(v, \tau) = \sum_{i \in N} Z_i(v, \tau)$ follows a binomial distribution with n trials and success probability $\Pr[X_i(v) \geq \tau]$. We note that the variables X_i are identically distributed for every $i \in N$.

For the first two cases below, we will need a technical lemma which we prove in Section 5.8.

Lemma 13. *For integer $s \in [nm]$, define the condition $\text{Balanced}(v, s)$ to be $X(v) = s$ with $\lfloor s/n \rfloor \leq X_i(v) \leq \lceil s/n \rceil$ for $i \in N$. Then, for any distribution $F \in \mathcal{F}^{0/1}$, it holds that*

$$\mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) | X(v) = s \right] \leq \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) | \text{Balanced}(v, s) \right].$$

Case 1. Consider impartial culture electorates with n agents and m alternatives with underlying values drawn from the distribution $F_p \in \mathcal{F}_1$. Denote by $\text{Box}(v)$ the condition $X_i(v) = \tau$ for $i \in S$ and $X_i(v) = 0$ for $i \in N \setminus S$, where S is a subset of the agents of size exactly $\lfloor n/2 \rfloor$. Now, define the quantity

$$B = \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) | \text{Box}(v) \right], \quad (5.4)$$

i.e., the expected maximum social welfare among all alternatives given that $\lfloor n/2 \rfloor$ agents draw a value of 1 for exactly τ alternatives and the remaining agents draw values of 0 for all alternatives. The quantity B will be a benchmark that will help us compare the expected social welfare of $\text{MEAN}(P, v)$ to the maximum social welfare.

As stated above, each $X_i(v)$ is a binomial random variable with m trials and success probability p . Hence, its median is either at the value $\lfloor pm \rfloor$ or $\lceil pm \rceil$, which are both at least τ since $p \geq 1/m$. Thus, $\Pr_{v \sim F}[X_i(v) \geq \tau] \geq 1/2$. The random variable $Z(v, \tau)$ then follows a binomial distribution with n trials and success probability at least $1/2$. By the same property of the median of the binomial distribution, it now holds that $\Pr_{v \sim F}[Z(v, \tau) \geq \lfloor n/2 \rfloor] \geq 1/2$. With this observation and by applying the law of total expectation, we have

$$\begin{aligned}
& \mathbb{E}_{\substack{v \sim F_p \\ P \sim \mathcal{P}(v)}} [\text{SW}(\text{MEAN}(P, v), v)] \\
&= \mathbb{E}_{\substack{v \sim F_p \\ P \sim \mathcal{P}(v)}} [\text{SW}(\text{MEAN}(P, v), v) | Z(v, \tau) \geq \lfloor n/2 \rfloor] \cdot \Pr_{v \in F_p} [Z(v, \tau) \geq \lfloor n/2 \rfloor] \\
&\quad + \mathbb{E}_{\substack{v \sim F_p \\ P \sim \mathcal{P}(v)}} [\text{SW}(\text{MEAN}(P, v), v) | Z(v, \tau) < \lfloor n/2 \rfloor] \cdot \Pr_{v \in F_p} [Z(v, \tau) < \lfloor n/2 \rfloor] \\
&\geq \frac{1}{2} \cdot \mathbb{E}_{\substack{v \sim F_p \\ P \sim \mathcal{P}(v)}} [\text{SW}(\text{MEAN}(P, v), v) | Z(v, \tau) \geq \lfloor n/2 \rfloor] \\
&\geq \frac{1}{2} \cdot \mathbb{E}_{\substack{v \sim F_p \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \text{SW}_{\text{im}}(a, P, v, \tau) | Z(v, \tau) \geq \lfloor n/2 \rfloor \right] \\
&\geq \frac{1}{2} \cdot \mathbb{E}_{v \sim F_p} \left[\max_{a \in A} \text{SW}(a, v) | \text{Box}(v) \right] = \frac{1}{2} \cdot B. \tag{5.5}
\end{aligned}$$

The first inequality above is obvious. By definition of MEAN , $\text{MEAN}(P, v)$ is the alternative that maximizes the implied social welfare when querying each agent in position τ . Hence, the implied social welfare is a lower bound for the social welfare of $\text{MEAN}(P, v)$, which yields the second inequality. Finally, under the condition that $Z(v, \tau) \geq \lfloor n/2 \rfloor$, there are at least $\lfloor n/2 \rfloor$ agents i for each of which $X_i(v) \geq \tau$. This includes the condition $\text{Box}(v)$ and the third inequality follows. The last equality follows from the definition of B in (5.4).

We proceed to bound the expected maximum social welfare from above in terms of B . For this purpose, we require another technical lemma which we prove in Section 5.8.

Lemma 14. *For every distribution $F \in \mathcal{F}^{0/1}$ and any positive integer j , it holds that*

$$\mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) | \text{Balanced}(v, jn\tau) \right] \leq 3j \cdot B.$$

For a positive integer j , define the condition $\text{Range}(v, j)$ to be $(j-1)n\tau < X(v) \leq$

$jn\tau$. By applying the law of total expectation, we obtain

$$\begin{aligned}
\mathbb{E}_{v \sim F_p} [\max_{a \in A} \text{SW}(a, v)] &= \sum_{j=1}^{\infty} \mathbb{E}_{v \sim F_p} [\max_{a \in A} \text{SW}(a, v) | \text{Range}(v, j)] \cdot \Pr_{v \sim F_p} [\text{Range}(v, j)] \\
&\leq \sum_{j=1}^{\infty} \mathbb{E}_{v \sim F_p} [\max_{a \in A} \text{SW}(a, v) | X(v) = jn\tau] \cdot \Pr_{v \sim F_p} [\text{Range}(v, j)] \\
&\leq \sum_{j=1}^{\infty} \mathbb{E}_{v \sim F_p} [\max_{a \in A} \text{SW}(a, v) | \text{Balanced}(v, jn\tau)] \cdot \Pr_{v \sim F_p} [\text{Range}(v, j)] \\
&\leq 3B \cdot \sum_{j=1}^{\infty} j \cdot \Pr_{v \sim F_p} [\text{Range}(v, j)]. \tag{5.6}
\end{aligned}$$

The first inequality follows since the condition $\text{Range}(v, j)$ includes the condition $X(v) = jn\tau$ and since the quantity $\mathbb{E}[\max_{a \in A} \text{SW}(a, v) | X(v) = t]$ is non-decreasing in terms of t . The second inequality follows from Lemma 13 while the third one follows from Lemma 14. We now bound the term $\sum_{j=1}^{\infty} j \cdot \Pr_{v \sim F_p} [\text{Range}(v, j)]$.

$$\begin{aligned}
\sum_{j=1}^{\infty} j \cdot \Pr_{v \sim F_p} [\text{Range}(v, j)] &= \sum_{j=1}^{\infty} j \sum_{k=(j-1)n\tau}^{jn\tau} \Pr_{v \sim F_p} [X(v) = k] \\
&= \frac{1}{n\tau} \cdot \sum_{j=1}^{\infty} \sum_{k=(j-1)n\tau}^{jn\tau} jn\tau \cdot \Pr_{v \sim F_p} [X(v) = k] \\
&\leq \Pr_{v \sim F_p} [X(v) \leq n\tau] + \frac{2}{n\tau} \cdot \sum_{j=2}^{\infty} \sum_{k=(j-1)n\tau}^{jn\tau} k \cdot \Pr_{v \sim F_p} [X(v) = k] \\
&\leq \frac{1}{2} + \frac{2}{n\tau} \cdot \mathbb{E}_{v \sim F_p} [X(v)] \leq 4.5.
\end{aligned}$$

The first inequality follows from the fact that $jn\tau/2 \leq (j-1)n\tau \leq k$ for $j \geq 2$. The second inequality follows from the definition of the expectation of random variable $X(v)$. Since $X(v)$ follows the binomial distribution with nm trials and success probability p , we have $\mathbb{E}[X(v)] = nmp$ which is at most $2n\lfloor mp \rfloor = 2n\tau$ since $p \geq 1/m$.

Now, inequality (5.6) implies that $\mathbb{E}_{v \sim F} [\max_{a \in A} \text{SW}(a, v)] \leq 13.5B$. Combining this bound with inequality (5.5) lets us conclude that $\text{avdist}(\text{MEAN}, \mathcal{F}_1) \leq 27$, as desired.

Case 2. Consider impartial culture electorates with n agents and m alternatives with underlying values drawn from the distribution $F_p \in \mathcal{F}_2$. Since $p < 1/m$, mechanism MEAN always queries the value of the top-ranked alternative in each agent, i.e., $\tau = 1$. Then, MEAN picks the alternative that maximizes the implied social welfare $\text{SW}_{\text{im}}(a, P, v, 1)$. It is thus immediately clear that

$$\mathbb{E}_{\substack{v \sim F_p \\ P \sim \mathcal{P}(v)}} [\text{SW}(\text{MEAN}(P, v), v)] \geq \mathbb{E}_{\substack{v \sim F_p \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \text{SW}_{\text{im}}(a, P, v, 1) \right]. \tag{5.7}$$

We intend to also relate the expected maximum social welfare to the RHS of (5.7). First, we observe that

$$\mathbb{E}_{v \sim F_p} \left[\max_{a \in A} \text{SW}(a, v) \right] = \sum_{t=1}^n \mathbb{E}_{\substack{v \sim F_p \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \text{SW}(a, v) | Z(v, 1) = t \right] \cdot \Pr_{v \sim F_p} [Z(v, 1) = t]. \quad (5.8)$$

In the following, we will upper-bound the term $\mathbb{E}[\max_{a \in A} \text{SW}(a, v) | Z(v, 1) = t]$ and show that this quantity is within a constant factor of $\mathbb{E}[\max_{a \in A} \text{SW}_{\text{im}}(a, P, v, 1) | Z(v, 1) = t]$ for every positive t . We will need another technical lemma (see Section 5.8 for the proof). Indeed, the proof of the Lemma 15 makes use of Lemma 13 and can be seen as a refined formulation of Lemma 14 for the present case where $\tau = 1$.

Lemma 15. *For every distribution $F \in \mathcal{F}^{0/1}$ and any positive integer j , it holds that*

$$\begin{aligned} & \mathbb{E}_{v \sim F_p} \left[\max_{a \in A} \text{SW}(a, v) | Z(v, 1) = t, X(v) = j \right] \\ & \leq \lceil j/t \rceil \cdot \mathbb{E}_{\substack{v \sim F_p \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \text{SW}_{\text{im}}(a, P, v, 1) | Z(v, 1) = t \right]. \end{aligned}$$

With this lemma in hand, we have that

$$\begin{aligned} & \mathbb{E}_{v \sim F_p} \left[\max_{a \in A} \text{SW}(a, v) | Z(v, 1) = t \right] \\ & = \sum_{j=t}^{\infty} \mathbb{E}_{v \sim F_p} \left[\max_{a \in A} \text{SW}(a, v) | Z(v, 1) = t, X(v) = j \right] \cdot \Pr_{v \sim F_p} [X(v) = j | Z(v, 1) = t] \\ & \leq \mathbb{E}_{\substack{v \sim F_p \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \text{SW}_{\text{im}}(a, P, v, 1) | Z(v, 1) = t \right] \sum_{j=t}^{\infty} \left\lceil \frac{j}{t} \right\rceil \Pr_{v \sim F_p} [X(v) = j | Z(v, 1) = t]. \end{aligned} \quad (5.9)$$

We proceed to bound the sum that appears in the previous inequality.

$$\begin{aligned} & \sum_{j=t}^{\infty} \left\lceil \frac{j}{t} \right\rceil \Pr_{v \sim F_p} [X(v) = j | Z(v, 1) = t] \\ & \leq \sum_{j=t}^{\infty} \left(\frac{j}{t} + 1 \right) \Pr_{v \sim F_p} [X(v) = j | Z(v, 1) = t] \\ & = \sum_{j=t}^{\infty} \Pr_{v \sim F_p} [X(v) = j | Z(v, 1) = t] + \frac{1}{t} \sum_{j=t}^{\infty} j \cdot \Pr_{v \sim F_p} [X(v) = j | Z(v, 1) = t] \\ & = 1 + \frac{1}{t} \mathbb{E}_{v \sim F_p} [X(v) | Z(v, 1) = t] \\ & = 1 + \mathbb{E}_{v \sim F_p} [X_i(v) | X_i(v) \geq 1], \end{aligned} \quad (5.10)$$

for any agent $i \in N$. The last equality is true since, under the condition $Z(v, 1) = t$, $X(v)$ is the sum $\sum_{i \in S} X_i(v)$ for a set S of t agents who are selected uniformly at random

and each satisfy $X_i(v) \geq 1$. As the random variables $X_i(v)$ are identically distributed for all agents in S , we obtain that $\mathbb{E}[X(v)|Z(v, 1) = t] = t \cdot \mathbb{E}[X_i(v)|X_i(v) \geq 1]$ for every agent i , which yields the equality.

Now, observe that

$$\mathbb{E}_{v \sim F_p}[X_i(v)|X_i(v) \geq 1] = \frac{\mathbb{E}_{v \sim F_p}[X_i(v)]}{\Pr_{v \sim F_p}[X_i(v) \geq 1]} = \frac{pm}{1 - (1 - p)^m}. \quad (5.11)$$

The derivative of the RHS of (5.11) with respect to $p \in (0, 1)$ is

$$\begin{aligned} & \frac{m}{(1 - (1 - p)^m)^2} (1 - (1 - p)^{m-1} (1 + p(m - 1))) \\ & \geq \frac{m}{(1 - (1 - p)^m)^2} (1 - ((1 - p)(1 + p))^{m-1}) \\ & = \frac{m}{(1 - (1 - p)^m)^2} (1 - (1 - p^2)^{m-1}) > 0. \end{aligned}$$

The first inequality follows from the inequality $(1 + t)^r \geq 1 + rt$ for $r \geq 1$. Hence, $\mathbb{E}[X_i(v)|X_i(v) \geq 1]$ is strictly increasing in p . Since $p < 1/m$ in the current case, it follows from (5.11) that

$$\mathbb{E}_{v \sim F_p}[X_i(v)|X_i(v) \geq 1] = \frac{pm}{1 - (1 - p)^m} \leq \frac{1}{1 - (1 - 1/m)^m} \leq \frac{1}{1 - e^{-1}},$$

using the inequality $(1 - 1/m)^m \leq e^{-1}$ for any $m \geq 1$.

Using this last observation together with inequalities (5.9) and (5.10), we obtain that

$$\mathbb{E}_{v \sim F_p} \left[\max_{a \in A} \text{SW}(a, v) \right] \leq \frac{2e - 1}{e - 1} \cdot \mathbb{E}_{\substack{v \sim F_p \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \text{SW}_{\text{im}}(a, P, v, 1) \right].$$

In combination with the lower bound on the social welfare of the mechanism MEAN in inequality (5.7), this yields an average distortion of at most 2.6.

Case 3. Again, mechanism MEAN queries the top-ranked alternative in each agent. Notice that if we have at most two agents, the average distortion of mechanism MEAN is at most 2. Indeed, mechanism MEAN always returns an alternative of positive social welfare whenever there exists one, and no alternative ever has a social welfare higher than 2. So, in the following, we assume that $n \geq 3$. We will show that the average distortion is lower than 2 in this case.

Consider impartial culture electorates with n agents and m alternatives with underlying values drawn from the distribution $F_p \in \mathcal{F}_3$. Notice that the maximum social welfare is never larger than the total social welfare of all alternatives. Hence,

$$\mathbb{E}_{v \sim F_p} \left[\max_{a \in A} \text{SW}(a, v) \right] \leq pnm.$$

With probability $1 - (1 - p)^{nm}$, there is at least one agent that gives a non-zero value to some alternative. Then, MEAN returns an alternative of social welfare at least 1. Hence,

$$\mathbb{E}_{\substack{v \sim F_p \\ P \sim \mathcal{P}(v)}} [\text{SW}(\text{MEAN}(P, v), v)] \geq 1 - (1 - p)^{nm}.$$

The average distortion of MEAN can therefore be upper-bounded by the term $\frac{pnm}{1 - (1 - p)^{nm}}$. Notice that the derivative with respect to p of this quantity is

$$\begin{aligned} & \frac{nm}{(1 - (1 - p)^{nm})^2} \cdot (1 - (1 - p)^{nm-1}(1 - p + pnm)) \\ & \geq \frac{nm}{(1 - (1 - p)^{nm})^2} \cdot (1 - (1 - p)^{nm-1}(1 + p)^{nm-1}) \\ & = \frac{nm}{(1 - (1 - p)^{nm})^2} \cdot (1 - (1 - p^2)^{nm-1}) > 0. \end{aligned}$$

The first inequality follows from the property $(1 + t)^r \geq 1 + rt$ for $r \geq 1$ and the second (strict) inequality is due to the fact that $p > 0$. Hence, the average distortion of the mechanism is strictly increasing in p . Using $p^* = 1 - (1 - 1/n)^{1/m}$ and since $p < p^*$, we have

$$\begin{aligned} \text{avdist}(\text{MEAN}, \mathcal{F}_3) &= \max_{F_p \in \mathcal{F}_3} \frac{\mathbb{E}_{v \sim F_p} [\max_{a \in A} \text{SW}(a, v)]}{\mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\text{SW}(M(P, v), v)]} \leq \frac{pnm}{1 - (1 - p)^{nm}} \\ &< \frac{p^*nm}{1 - (1 - p^*)^{nm}}. \quad (5.12) \end{aligned}$$

Substituting p^* , we get that the denominator in the RHS of (5.12) is equal to $1 - (1 - 1/n)^n \geq 1 - e^{-1}$. To bound the numerator, observe that $\left(1 - \frac{3 \ln \frac{3}{2}}{nm}\right)^{nm} < \frac{8}{27}$ and that $(1 - 1/n)^n \geq \frac{8}{27}$ for $n \geq 3$. These inequalities follow since the expression $(1 - r/t)^t$ is strictly increasing for $t \geq r$ and approaches e^{-r} from below as t goes to infinity. Thus, we have that $\left(1 - \frac{3 \ln \frac{3}{2}}{nm}\right)^m < 1 - 1/n$, which is equivalent to $\frac{3 \ln \frac{3}{2}}{nm} > 1 - (1 - 1/n)^{1/m} = p^*$, implying that the numerator of the RHS of (5.12) is at most $3 \ln \frac{3}{2}$. We conclude that $\text{avdist}(\text{MEAN}, \mathcal{F}_3) < \frac{3 \ln \frac{3}{2}}{1 - e^{-1}} < 2$, completing the proof. $\square$

5.5 Randomized Mechanisms

We now present two randomized mechanisms for impartial culture electorates with underlying valuations drawn from a *general* probability distribution and worst-case electorates, respectively. Both mechanisms are nevertheless similar in spirit. They randomly pick a single threshold from a suitably defined set of thresholds and query each agent to determine a set of alternatives that have value above the threshold. This information is then used to compute an approximation of the alternatives' respective social welfare, which is used to decide the winning alternative.

Impartial Culture Electorates

Our first “random threshold” mechanism RTMEAN uses mechanism MEAN as a building block.

Definition 16 (mechanism RTMEAN). *The mechanism RTMEAN uses k thresholds $\ell_1, \ell_2, \dots, \ell_k$ with $0 < \ell_1 < \dots < \ell_k$ as parameters. Given a profile P with underlying valuations drawn from a probability distribution F , RTMEAN selects an integer t uniformly at random from $[k]$, and sets $p = \Pr_{z \sim F}[z \geq \ell_t]$. It then simulates an execution of MEAN on the distribution $F_p \in \mathcal{F}^{0/1}$ by*

- *making the same value queries as MEAN for F_p , but*
- *interpreting the answer $\text{val}_i(a)$ to a query as 1 if $\text{val}_i(a) \geq \ell_t$ and 0 otherwise.*

RTMEAN returns as output the alternative that MEAN selects.

Notice that mechanism RTMEAN uses exactly one query per agent.

Theorem 17. *Let F be a probability distribution over non-negative, real-valued outcomes with mean μ and variance σ^2 . There is a set of thresholds $\ell_1, \ell_2, \dots, \ell_k$ such that the average distortion of mechanism RTMEAN is at most $O(\log m + \log \frac{\sigma^2}{\mu^2})$ when applied to impartial culture electorates with n agents and m alternatives, and underlying values drawn according to F .*

Proof. To prove the theorem, we use the following lemma which relates the average distortion of RTMEAN to the structure of the distribution F .

Lemma 18. *For a random variable z following F , assume that there are $L, U > 0$ such that*

$$\mathbb{E}_{z \sim F}[z \mathbb{I}\{z < L\} + (z - U) \mathbb{I}\{z \geq U\}] \leq \frac{\mu}{2m}.$$

Then, there exists a choice of thresholds $\ell_1, \ell_2, \dots, \ell_k$ such that mechanism RTMEAN yields average distortion at most $108 \lceil \log \frac{U}{L} \rceil$.

Proof. Set $k = \lceil \log \frac{U}{L} \rceil$ and define the thresholds of mechanism RTMEAN as $\ell_t = L \cdot 2^{t-1}$ for $t = 1, 2, \dots, k$ and $\ell_k = U$. We begin by observing that, for any $z \geq 0$, we have

$$\begin{aligned} z &\leq z \mathbb{I}\{z < \ell_1\} + \ell_1 \mathbb{I}\{z \geq \ell_1\} + \sum_{t=1}^{k-1} (\ell_{t+1} - \ell_t) \mathbb{I}\{z \geq \ell_t\} + (z - \ell_k) \mathbb{I}\{z \geq \ell_k\} \\ &\leq z \mathbb{I}\{z < \ell_1\} + (z - \ell_k) \mathbb{I}\{z \geq \ell_k\} + 2 \sum_{t=1}^k \ell_t \mathbb{I}\{z \geq \ell_t\}. \end{aligned} \tag{5.13}$$

The second inequality follows since the definition of the thresholds implies that $\ell_{t+1} \leq 2\ell_t$ for $t = 1, \dots, k-1$ and, hence $\ell_{t+1} - \ell_t \leq \ell_t$. We now have

$$\begin{aligned}
& \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \right] \\
&= \mathbb{E}_{v \sim F} \left[\max_{a \in A} \sum_{i=1}^n \text{val}_i(a) \right] \\
&\leq \mathbb{E}_{v \sim F} \left[\max_{a \in A} \sum_{i=1}^n \left(\text{val}_i(a) \mathbb{I}\{\text{val}_i(a) < \ell_1\} + (\text{val}_i(a) - \ell_k) \mathbb{I}\{\text{val}_i(a) \geq \ell_k\} \right. \right. \\
&\quad \left. \left. + 2 \cdot \sum_{t=1}^k \ell_t \mathbb{I}\{\text{val}_i(a) \geq \ell_t\} \right) \right] \\
&\leq \mathbb{E}_{v \sim F} \left[\max_{a \in A} \sum_{i=1}^n (\text{val}_i(a) \mathbb{I}\{\text{val}_i(a) < \ell_1\} + (\text{val}_i(a) - \ell_k) \mathbb{I}\{\text{val}_i(a) \geq \ell_k\}) \right] \\
&\quad + 2 \cdot \mathbb{E}_{v \sim F} \left[\max_{a \in A} \sum_{i=1}^n \sum_{t=1}^k \ell_t \mathbb{I}\{\text{val}_i(a) \geq \ell_t\} \right] \\
&\leq \mathbb{E}_{v \sim F} \left[\sum_{a \in A} \sum_{i=1}^n (\text{val}_i(a) \mathbb{I}\{\text{val}_i(a) < \ell_1\} + (\text{val}_i(a) - \ell_k) \mathbb{I}\{\text{val}_i(a) \geq \ell_k\}) \right] \\
&\quad + 2 \cdot \mathbb{E}_{v \sim F} \left[\max_{a \in A} \sum_{i=1}^n \sum_{t=1}^k \ell_t \mathbb{I}\{\text{val}_i(a) \geq \ell_t\} \right] \\
&\leq \frac{1}{2m} \cdot \mathbb{E}_{v \sim F} \left[\sum_{a \in A} \sum_{i=1}^n \text{val}_i(a) \right] + 2 \cdot \mathbb{E}_{v \sim F} \left[\max_{a \in A} \sum_{i=1}^n \sum_{t=1}^k \ell_t \mathbb{I}\{\text{val}_i(a) \geq \ell_t\} \right] \\
&\leq \frac{1}{2} \cdot \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \right] + 2 \cdot \mathbb{E}_{v \sim F} \left[\max_{a \in A} \sum_{i=1}^n \sum_{t=1}^k \ell_t \mathbb{I}\{\text{val}_i(a) \geq \ell_t\} \right].
\end{aligned}$$

The first inequality follows from (5.13), the second and third inequalities use the fact that the maximum among non-negative values is upper-bounded by their sum, the fourth inequality uses the assumption in the statement of the lemma for the random variable $\text{val}_i(a)$ with $\mu = \mathbb{E}_{\text{val}_i(a) \sim F}[\text{val}_i(a)]$, and the last inequality follows since the average among non-negative values is a lower-bound on the maximum value among them. The above inequality is equivalent to

$$\mathbb{E}_{v \sim F} \left[\max_{a \in A} \sum_{i=1}^n \sum_{t=1}^k \ell_t \mathbb{I}\{\text{val}_i(a) \geq \ell_t\} \right] \geq \frac{1}{4} \cdot \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \right]. \quad (5.14)$$

For valuations v and $t \in [k]$, define the valuations v^t that consists of the binary values $\mathbb{I}\{\text{val}_i(a) \geq \ell_t\}$. Notice that for every alternative $a \in A$, it is $\text{SW}(a, v) \geq \ell_t \cdot \text{SW}(a, v^t)$. Mechanism RTMEAN selects the integer t uniformly at random from $[k]$ and, for every profile P that is consistent with the valuations v , it returns the alternative

$\text{MEAN}(P, v^t)$. We thus have

$$\begin{aligned}
& \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\text{SW}(\text{RTMEAN}(P, v), v)] \\
&= \frac{1}{k} \cdot \sum_{t=1}^k \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\text{SW}(\text{MEAN}(P, v^t), v)] \geq \frac{1}{k} \cdot \sum_{t=1}^k \ell_t \cdot \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\text{SW}(\text{MEAN}(P, v^t), v^t)] \\
&\geq \frac{1}{27k} \sum_{t=1}^k \ell_t \cdot \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v^t) \right] \geq \frac{1}{27k} \mathbb{E}_{v \sim F} \left[\max_{a \in A} \sum_{t=1}^k \ell_t \cdot \text{SW}(a, v^t) \right] \\
&= \frac{1}{27k} \cdot \mathbb{E}_{v \sim F} \left[\max_{a \in A} \sum_{i=1}^n \sum_{t=1}^k \ell_t \mathbb{I}\{\text{val}_i(a) \geq \ell_t\} \right] \geq \frac{1}{108k} \cdot \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \right],
\end{aligned}$$

as desired. The second inequality follows from the average distortion guarantee for mechanism MEAN from Theorem 12, and the fifth one by inequality (5.14). $\square$

Now, let $L = \frac{\mu}{4m}$ and $U = \mu + \frac{4m\sigma^2}{\mu}$. Notice that

$$\mathbb{E}_{z \sim F} [z \mathbb{I}\{z < L\}] \leq L \cdot \Pr_{z \sim F} [z < L] \leq \frac{\mu}{4m}. \quad (5.15)$$

We slightly overload our notation and denote by F also the cumulative distribution function of a random variable z following the distribution F . For $t \geq U$, we have that

$$1 - F(t) = \Pr_{z \sim F} [z \geq t] \leq \Pr_{z \sim F} [|z - \mu| \geq t - \mu] \leq \frac{\sigma^2}{(t - \mu)^2}$$

where the last transition follows from Chebyshev's inequality. Then,

$$\begin{aligned}
\mathbb{E}_{z \sim F} [(z - U) \mathbb{I}\{z \geq U\}] &= \int_U^\infty (1 - F(t)) dt \leq \int_U^\infty \frac{\sigma^2}{(t - \mu)^2} dt \\
&= \frac{\sigma^2}{U - \mu} = \frac{\mu}{4m}. \quad (5.16)
\end{aligned}$$

By (5.15) and (5.16), the condition of Lemma 18 is satisfied with $U/L = 4m + 16m^2 \frac{\sigma^2}{\mu^2}$. The theorem then follows from the bound on the average distortion of RTMEAN provided by Lemma 18. $\square$

As an immediate consequence of Theorem 17, we obtain $O(\log m)$ bounds on the average distortion of RTMEAN for fundamental families of distributions. The relevant properties of these distributions which we use for the following corollary can be found in, e.g., Papoulis and Unnikrishna Pillai [107, table 5-2, p. 162].

Corollary 19. *The average distortion of mechanism RTMEAN is at most $O(\log m)$ when applied to impartial culture electorates with n agents and m alternatives, and underlying values drawn according to*

- the exponential distribution $E(\lambda)$ where $\mu = \frac{1}{\lambda}$, $\sigma^2 = \frac{1}{\lambda^2}$,

- the chi-squared distribution $\chi^2(k)$ where $\mu = k, \sigma^2 = 2k$, or
- the Erlang- k distribution $E_k(\lambda)$ where $\mu = \frac{1}{\lambda}, \sigma^2 = \frac{1}{k\lambda^2}$.

In the latter two cases, k denotes a positive integer.

We remark that distributions for which our analysis does not give logarithmic average distortion are the γ -distribution $G(\alpha, \beta)$ where $\mu = \alpha\beta, \sigma^2 = \alpha\beta^2$ and the β -distribution $\beta(\alpha, \beta)$ where $\mu = \frac{\alpha}{\alpha+\beta}, \sigma^2 = \frac{\alpha\beta}{(\alpha+\beta)^2(\alpha+\beta+1)}$ for respective α -parameters very close to zero.

Worst-Case Electorates

Our next random threshold mechanism RTSEARCH achieves low worst-case distortion.³

Definition 20 (mechanism RTSEARCH). *For a given profile P with underlying valuations v over m alternatives, mechanism RTSEARCH picks an integer r uniformly at random from the set $\{1, 2, \dots, \lceil \log 2m \rceil\}$. For every agent $i \in N$, mechanism RTSEARCH*

- *queries the value v_i of the agent's top-ranked alternative, and*
- *finds the set of alternatives $S_{i,r}$, for each of which the agent has value more than $v_i/2^r$ using binary search.*

The mechanism then returns an alternative $\text{RTSEARCH}(P, v) \in \arg \max_{a \in A} \sum_{i \in N} v_i \mathbb{I}\{a \in S_{i,r}\}$.

Theorem 21. *Mechanism RTSEARCH achieves worst-case distortion at most $O(\log m)$ with $O(\log m)$ queries per agent.*

Proof. Clearly, for any fixed r and any agent i , the alternatives $S_{i,r}$ can be identified by finding the lowest-ranked alternative in $\succ_i$ with value more than $v_i/2^r$. This can be accomplished using binary search using $O(\log m)$ queries per agent. For the upper bound on the distortion of RTSEARCH, we introduce the following notion of *artificial* social welfare. Given valuations v and an alternative $a \in A$, we define the artificial social welfare of a as

$$\text{SW}_{\text{art}}(a, v) = \sum_{i \in N} v_i \sum_{r=1}^{\infty} \frac{\mathbb{I}\{a \in S_{i,r}\}}{2^r}.$$

With the next lemma, we show that for any alternative, its artificial social welfare is within a factor of 2 of its social welfare.

³We remark that RTSEARCH is similar in spirit to a mechanism proposed by Benadè et al. [20] in a slightly different participatory budgeting context. There are important differences in modelling assumptions, though. First, their model allows for more powerful queries than ours, where an agent is asked to report all alternatives which she ranks above a certain threshold. More importantly, they assume unit-sum valuations; this assumption affects the design of their mechanism and simplifies its analysis.

Lemma 22. *For any set of valuations v and any alternative $a \in A$, it holds that $\text{SW}(a, v) \leq \text{SW}_{\text{art}}(a, v) \leq 2 \text{SW}(a, v)$.*

Proof. For any agent $i \in N$, let r_i^* be a non-negative integer such that $\text{val}_i(a) \in (v_i/2^{r_i^*}, v_i/2^{r_i^*-1}]$. Notice that thereby

$$\sum_{i \in N} \frac{v_i}{2^{r_i^*-1}} \geq \sum_{i \in N} \text{val}_i(a) = \text{SW}(a, v), \quad (5.17)$$

and

$$\sum_{i \in N} \frac{v_i}{2^{r_i^*-1}} < 2 \sum_{i \in N} \text{val}_i(a) = 2 \text{SW}(a, v). \quad (5.18)$$

Then, for every agent i and any alternative a , we have that

$$\sum_{r=1}^{\infty} \frac{\mathbb{I}\{a \in S_{i,r}\}}{2^r} = \sum_{r=r_i^*}^{\infty} \frac{1}{2^r} = \frac{1}{2^{r_i^*}} \sum_{r=0}^{\infty} \frac{1}{2^r} = \frac{1}{2^{r_i^*-1}},$$

and, thus,

$$\text{SW}_{\text{art}}(a, v) = \sum_{i \in N} v_i \sum_{r=1}^{\infty} \frac{\mathbb{I}\{a \in S_{i,r}\}}{2^r} = \sum_{i \in N} \frac{v_i}{2^{r_i^*-1}}.$$

The lemma follows from inequalities (5.17) and (5.18). $\square$

For a given profile P , let a_r be the alternative returned by mechanism RTSEARCH for a draw of $r = 1, 2, \dots, \lceil \log 2m \rceil$, respectively. Denote by a^* the alternative of maximum social welfare for the valuations v underlying P . Then, by definition of RTSEARCH , it holds that

$$\sum_{i \in N} \frac{v_i \mathbb{I}\{a_r \in S_{i,r}\}}{2^r} \geq \sum_{i \in N} \frac{v_i \mathbb{I}\{a^* \in S_{i,r}\}}{2^r},$$

for every $r \in \{1, 2, \dots, \lceil \log 2m \rceil\}$, and, hence,

$$\sum_{r=1}^{\lceil \log 2m \rceil} \sum_{i \in N} \frac{v_i \mathbb{I}\{a_r \in S_{i,r}\}}{2^r} \geq \sum_{r=1}^{\lceil \log 2m \rceil} \sum_{i \in N} \frac{v_i \mathbb{I}\{a^* \in S_{i,r}\}}{2^r}. \quad (5.19)$$

We now have that

$$\begin{aligned} & \sum_{r=1}^{\lceil \log 2m \rceil} \text{SW}_{\text{art}}(a_r, v) \\ &= \sum_{r=1}^{\lceil \log 2m \rceil} \sum_{i \in N} v_i \sum_{r'=1}^{\infty} \frac{\mathbb{I}\{a_r \in S_{i,r'}\}}{2^{r'}} \geq \sum_{i \in N} v_i \sum_{r=1}^{\lceil \log 2m \rceil} \frac{\mathbb{I}\{a_r \in S_{i,r}\}}{2^r} \\ &\geq \sum_{i \in N} v_i \sum_{r=1}^{\lceil \log 2m \rceil} \frac{\mathbb{I}\{a^* \in S_{i,r}\}}{2^r} = \text{SW}_{\text{art}}(a^*, v) - \sum_{i \in N} v_i \sum_{r=\lceil \log 2m \rceil+1}^{\infty} \frac{\mathbb{I}\{a^* \in S_{i,r}\}}{2^r} \\ &\geq \text{SW}_{\text{art}}(a^*, v) - \sum_{i \in N} v_i \left(\frac{1}{2^{\lceil \log 2m \rceil+1}} \sum_{r=0}^{\infty} \frac{1}{2^r} \right) \geq \text{SW}_{\text{art}}(a^*, v) - \frac{1}{2m} \sum_{i \in N} v_i \\ &\geq \frac{1}{2} \text{SW}_{\text{art}}(a^*, v). \end{aligned} \quad (5.20)$$

Here, we used (5.19) to arrive at the second inequality. The last inequality follows since the average value of the top-ranked alternatives, that is, $(1/m) \sum_{i \in N} v_i$, cannot be higher than $\text{SW}(a^*, v)$ and, by Lemma 22, $\text{SW}(a^*, v) \leq \text{SW}_{\text{art}}(a^*, v)$. Hence, for any set of valuations v and any profile $P \in \mathcal{P}(v)$, we have that

$$\begin{aligned} & \mathbb{E}_{r \sim [\lceil \log 2m \rceil]} [\text{SW}(\text{RTSEARCH}(P, v), v)] \\ &= \frac{1}{\lceil \log 2m \rceil} \sum_{r=1}^{\lceil \log 2m \rceil} \text{SW}(a_r, v) \geq \frac{1}{2\lceil \log 2m \rceil} \sum_{r=1}^{\lceil \log 2m \rceil} \text{SW}_{\text{art}}(a_r, v) \\ &\geq \frac{1}{4\lceil \log 2m \rceil} \text{SW}_{\text{art}}(a^*, v) \geq \frac{1}{4\lceil \log 2m \rceil} \text{SW}(a^*, v), \end{aligned}$$

where the first inequality follows from Lemma 22, the second inequality follows from (5.20), and the third inequality again follows from Lemma 22. This concludes the proof of the theorem. $\square$

5.6 Worst-Case Distortion Lower Bounds

We conclude our technical exposition by presenting two lower bounds on the worst-case distortion. Our basic approach in both of them is as follows. First, for every (large enough) value of m and a value of n of our choice, we decide the agents' rankings. For every position in an agent's ranking, we pre-define a value that is revealed if this particular position is queried by a mechanism. Let $\hat{a}$ be the alternative that a mechanism picked as winning on the given profile. We then show that—for every choice of $\hat{a}$ —it is possible to *fix* (i.e., to choose) the agents' remaining concealed valuations in such a way that the distortion is high. That is, for any position not queried by the mechanism, we assume an adversarial set of valuations that is consistent with the agents' rankings and with the values revealed to the mechanism.

Lower Bounding the Number of Queries for Constant Distortion

Our first lower bound on the number of queries per agent that are necessary to get constant worst-case distortion improves the previously best bound of Amanatidis et al. [4] by a sublogarithmic factor. More specifically, Amanatidis et al. [4] present a lower bound construction and show that any mechanism that uses up to λ queries per agent must have worst-case distortion at least $\Omega\left(\frac{1}{\lambda} \cdot m^{\frac{1}{2(\lambda+1)}}\right)$. Thus, in order to get constant worst-case distortion, at least $\Omega\left(\frac{\log m}{\log \log m}\right)$ queries per agent are necessary. In their follow-up work, Amanatidis et al. [5] present an improved construction which yields a higher distortion of $\Omega(m^{1/\lambda})$. Unfortunately, λ is now required to be a constant. Therefore, their new construction does not provide any lower bound on the number of queries per agent necessary to get constant worst-case distortion. We prove the next theorem using a considerably different construction which shows that mechanisms making at most λ queries per agent have worst-case distortion at least $\frac{1}{8} \cdot m^{\frac{1}{3\lambda}}$ for values of λ that are allowed to be logarithmic in m .

Theorem 23. *Any deterministic mechanism that achieves a constant worst-case distortion must make $\Omega(\log m)$ queries per agent.*

Proof. Let $m \geq 154$ and λ be an integer such that $2 \leq \lambda \leq \log m$. Consider a mechanism $\mathcal{M}$ that makes at most λ queries per agent; we will show that $\mathcal{M}$ has worst-case distortion at least $\frac{1}{8}m^{\frac{1}{3\lambda}}$.

We define the symmetric profile $P = \{\succ_i\}_{i \in N}$ with $n = m$ agents, so that agent i has the ranking

$$i \succ_i i+1 \succ_i \dots \succ_i m \succ_i 1 \succ_i \dots \succ_i i-1.$$

The ranking of every agent i is divided into $2\lambda + 1$ sets—or *buckets*— $B_1^{(i)}, \dots, B_{2\lambda+1}^{(i)}$ where

$$|B_j^{(i)}| = b_j = \left\lceil m^{\frac{j}{3\lambda}} \right\rceil,$$

for $j \in [2\lambda]$, and $|B_{2\lambda+1}^{(i)}| = b_{2\lambda+1} = m - \sum_{j=1}^{2\lambda} b_j$. Hence,

$$B_j^{(i)} = \left\{ i + \sum_{t=1}^{j-1} b_t \bmod m, \dots, i - 1 + \sum_{t=1}^j b_t \bmod m \right\}.$$

We refer to the alternatives from bucket $B_{2\lambda+1}^{(i)}$ as the *tail alternatives* of agent i .⁴

We proceed to describe the agents' valuations v . Every agent i assigns a value of 0 to each of her tail alternatives, i.e., $\text{val}_i(a) = 0$ for every $a \in B_{2\lambda+1}^{(i)}$. For $j \in [2\lambda]$, agent i assigns to all the alternatives of bucket $B_j^{(i)}$ either a *low* value of $m^{\frac{2\lambda-j}{3\lambda}}$ or a *high* value of $m^{\frac{2\lambda-j+1}{3\lambda}}$ in the following way. Whenever the mechanism $\mathcal{M}$ makes a query for the value of an alternative in bucket $B_j^{(i)}$, the concealed value of each alternative in bucket $B_j^{(i)}$ is set to the low value, i.e., $\text{val}_i(a) = m^{\frac{2\lambda-j}{3\lambda}}$ for every alternative $a \in B_j^{(i)}$; this value is also revealed as the outcome of the query. Now, consider a bucket $B_j^{(i)}$, in which mechanism $\mathcal{M}$ did not query the value of any alternative. The concealed values of all alternatives in this bucket are set to the low value $m^{\frac{2\lambda-j}{3\lambda}}$ if the winning alternative $\mathcal{M}(P, v)$ belongs to the bucket and the high value $m^{\frac{2\lambda-j+1}{3\lambda}}$ otherwise. Figure 5.1 shows an example that demonstrates this approach.

Let $\hat{a} = \mathcal{M}(P, v)$. Observe that alternative $\hat{a}$ belongs to bucket $B_j^{(i)}$ for b_j different choices of $i \in N$. Hence,

$$\text{SW}(\hat{a}, v) = \sum_{i \in N} \sum_{j \in [2\lambda]: \hat{a} \in B_j^{(i)}} m^{\frac{2\lambda-j}{3\lambda}} = \sum_{j \in [2\lambda]} b_j \cdot m^{\frac{2\lambda-j}{3\lambda}} = \sum_{j \in [2\lambda]} \left\lceil m^{\frac{j}{3\lambda}} \right\rceil \cdot m^{\frac{2\lambda-j}{3\lambda}} \leq 4\lambda m^{2/3}. \quad (5.21)$$

We now compute the sum of the social welfare over all alternatives by summing up all the values in every bucket of every agent. To do so, define the subsets H and

⁴Our assumptions $m \geq 154$ and $\lambda \leq \log m$ guarantee that $\sum_{j=1}^{2\lambda} b_j \leq m$ and, thus, buckets $B_{2\lambda+1}^{(i)}$ are well-defined.

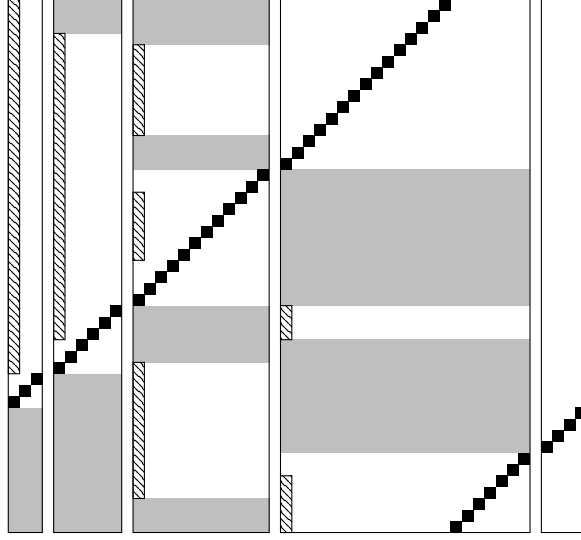


Figure 5.1: An example for our lower bound construction in Theorem 23 that illustrates the way in which the agents' valuations are defined. The alternative $\hat{a}$ picked by the mechanism $\mathcal{M}$ is marked as a black box in every agent's ranking. We assume that $\mathcal{M}$ queried the positions corresponding to the dashed boxes. In this example, the mechanism only queries the first position in a bucket which is without loss of generality. The gray areas correspond to buckets in which all alternatives have high values. White areas correspond to buckets in which alternatives have low values, either because the bucket contains the winning alternative $\hat{a}$ or because mechanism $\mathcal{M}$ queried the value of an alternative in the bucket.

L of $N \times [2\lambda]$ as follows. L consists of the pairs (i, j) such that either $\mathcal{M}(P, v) \in B_j^{(i)}$ or mechanism $\mathcal{M}$ queried the value of some alternative in bucket $B_j^{(i)}$. Let $H = N \times [2\lambda] \setminus L$. We have

$$\begin{aligned}
 \sum_{a \in A} \text{SW}(a, v) &= \sum_{i \in N} \left(\sum_{j \in [2\lambda]: (i, j) \in L} b_j \cdot m^{\frac{2\lambda-j}{3\lambda}} + \sum_{j \in [2\lambda]: (i, j) \in H} b_j \cdot m^{\frac{2\lambda-j+1}{3\lambda}} \right) \\
 &= \sum_{i \in N} \left(\sum_{j \in [2\lambda]} b_j \cdot m^{\frac{2\lambda-j}{3\lambda}} + \sum_{j \in [2\lambda]: (i, j) \in H} b_j \cdot \left(m^{\frac{2\lambda-j+1}{3\lambda}} - m^{\frac{2\lambda-j}{3\lambda}} \right) \right) \\
 &\geq \sum_{i \in N} \sum_{j \in [2\lambda]: (i, j) \in H} b_j \cdot m^{\frac{2\lambda-j+1}{3\lambda}} \\
 &\geq \sum_{i \in N} \sum_{j \in [2\lambda]: (i, j) \in H} m^{\frac{2\lambda+1}{3\lambda}} = |H| \cdot m^{\frac{2\lambda+1}{3\lambda}}. \tag{5.22}
 \end{aligned}$$

Now, observe that for every agent i , the set L contains at most $\lambda + 1$ pairs $(i, \cdot)$ for the up to λ buckets in which $\mathcal{M}$ queried the value of some alternative and (possibly) one extra bucket that contains alternative $\hat{a}$. Hence, $|L| \leq (\lambda + 1) \cdot n$ and, consequently,

$|H| \geq (\lambda - 1) \cdot n$. Recalling that $n = m$, inequality (5.22) yields

$$\max_{a \in A} \text{SW}(a, v) \geq \frac{1}{m} \cdot \sum_{a \in A} \text{SW}(a, v) \geq (\lambda - 1) \cdot m^{\frac{2\lambda+1}{3\lambda}}. \quad (5.23)$$

The desired lower bound of $\frac{\lambda-1}{4\lambda} \cdot m^{\frac{1}{3\lambda}} \geq \frac{1}{8} \cdot m^{\frac{1}{3\lambda}}$ on the distortion now follows from inequalities (5.21) and (5.23). $\square$

Lower Bounding the Distortion of 1-Query Mechanisms

In Section 5.4, we saw that a single query per agent is sufficient to guarantee constant average distortion when the agents draw their valuations according to a binary distribution. Such a guarantee is not attainable in the traditional setting of worst-case distortion. Indeed, Amanatidis et al. [4] proved that any deterministic 1-query mechanism must have distortion $\Omega(m)$. However, their lower bound construction uses valuations that are more complex than binary. Our next result states that, even in the case of binary valuations, the worst-case distortion of deterministic 1-query mechanisms must still be high.

Theorem 24. *Every deterministic 1-query mechanism has a worst-case distortion of at least $\Omega(\sqrt{m})$. This is true even if the agents have binary values for each of the alternatives.*

Proof. Let $m \geq 16$ and t be the largest even integer such that $t^2 \leq m$. Clearly, $t \in \Omega(\sqrt{m})$. We consider the profile $P = \{\succ_i\}_{i \in N}$ with $n = t^2$ agents so that for every agent $i \in N$, the ranking $\succ_i$ has the form

$$i \succ_i i+1 \succ_i \dots \succ_i t^2 \succ_i 1 \succ_i \dots \succ_i i-1 \succ_i t^2+1 \succ_i \dots \succ_i m.$$

We divide the t^2 agents into t groups, each containing t agents. We call these groups *cohorts*. For $k \in [t]$, the k -th cohort $\mathcal{C}_k$ contains the agents $(k-1)t+1, \dots, kt$. Due to the symmetry of P and the assumption that $n = t^2$, an element in $\mathcal{C}_k$ may refer to an agent $i \in \mathcal{C}_k$ as well as to an alternative j that is the top-ranked alternative of agent $j \in \mathcal{C}_k$. Figure 5.2 shows an example of our lower bound construction.

Let $\text{val}_{i,j}$ denote the value that agent i has for the alternative in the j -th position of her ranking. For every query that the mechanism makes, we reveal the following information:

- If agent i is the first agent of cohort $\mathcal{C}_k$ who is queried by the mechanism at a position $j \leq t$, we reveal $\text{val}_{i,j} = 1$.
- Otherwise, we reveal $\text{val}_{i,j} = 0$.

The latter item includes the case where the mechanism queries an agent at any position $j > t$ as well as the case where the mechanism already queried an agent of cohort $\mathcal{C}_k$ at a position $j \leq t$.

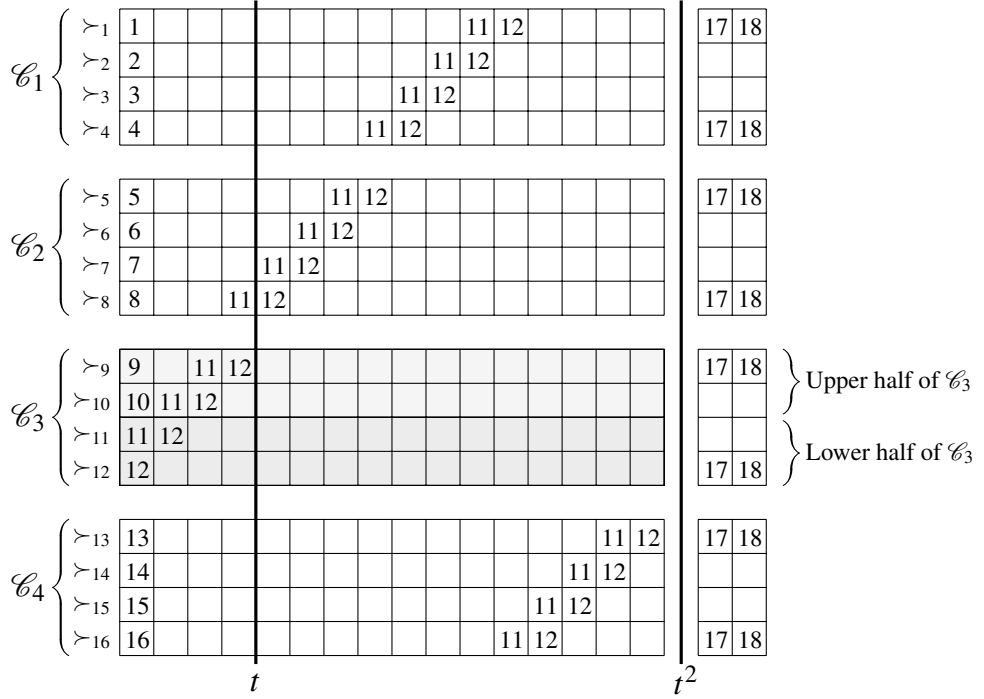


Figure 5.2: An example of our lower bound construction in Theorem 24 where $m = 18$. Hence, we set $t = 4$ and $n = t^2 = 16$ resulting in the profile P shown above. In every agent's ranking (horizontal bars), the example mentions the top-ranked alternative. Additionally, we marked alternatives 11 and 12 in every ranking in order to showcase the symmetry of P . Alternatives 17 and 18 are shown for two agents of every cohort. Notice that these alternatives appear in the exact same positions of every agent's ranking.

With the next two lemmas, we distinguish two cases that—taken together—cover all possible choices a mechanism may make for querying positions in P . In each case, we are able to fix the agents' valuations in a way that is consistent with the values revealed to the mechanism and results in a high distortion. The following piece of notation will be helpful for this purpose. Let $N_{\leq j}(a)$ be the set of agents that rank alternative a at position j or higher, that is,

$$N_{\leq j}(a) = \{i \in N : \text{pos}_{\gamma_i}(a) \leq j\}.$$

Lemma 25. *Assume that there is a cohort $\mathcal{C}_k$ such that mechanism $\mathcal{M}$ does not query any agent $i \in \mathcal{C}_k$ at a position $j \leq t$. Then, there exists a choice of valuations v that is consistent with profile P and the values revealed to $\mathcal{M}$, such that*

$$\frac{\max_{a \in A} \text{SW}(a, v)}{\text{SW}(\mathcal{M}(P, v), v)} \geq t/2.$$

Proof. First, assume that $\mathcal{M}(P, v) = a \notin \mathcal{C}_k$. Consider the alternative $a' = kt \in \mathcal{C}_k$. Note that $N_{\leq t}(a') = \mathcal{C}_k$, that is, all agents of cohort $\mathcal{C}_k$ rank alternative a' at a position $j \leq t$. Since the algorithm did not query any of these positions, we are free to fix the concealed values for every $i \in N_{\leq t}(a')$ such that

$$\text{val}_{i,j} = \begin{cases} 1 & \text{for every position } j \leq \text{pos}_{\succ_i}(a') \\ 0 & \text{otherwise.} \end{cases}$$

The remaining concealed values (outside of cohort $\mathcal{C}_k$) are set to 0 except for those positions where a value of 1 is implied by a revealed value of 1 at a position further down in the ranking. Thereby, $\text{SW}(a', v) = t$. Furthermore, notice that the set $N_{\leq t}(a)$ intersects with at most two cohorts and alternative a receives a value of 0 from every agent $i \in N_{\leq t}(a')$. Thus, $\text{SW}(a, v) \leq 2$, which shows that the lemma holds for $\mathcal{M}(P, v) = a \notin \mathcal{C}_k$.

Now, let $\mathcal{M}(P, v) = a \in \mathcal{C}_k$. We will further distinguish between the cases where $a \leq (k - 1/2)t$ (i.e., a is in the *upper half* of $\mathcal{C}_k$; see Figure 5.2) and $a > (k - 1/2)t$ (i.e., a is in the *lower half* of $\mathcal{C}_k$). For both cases, we first show that there is an alternative $a' \neq a$ such that $\text{SW}(a', v) \geq t/2$.

Case 1: $a \leq (k - 1/2)t$. Consider alternative $a' = kt$. Note that $N_{\leq t}(a') = \mathcal{C}_k$. Furthermore, since a is in the upper half of the cohort, it holds that

$$|N_{\leq t}(a') \setminus N_{\leq t}(a)| \geq t/2.$$

Since $\mathcal{M}$ did not query any agent from $\mathcal{C}_k$ at a position $j \leq t$, we can fix the concealed values for every $i \in N_{\leq t}(a') \setminus N_{\leq t}(a)$ such that

$$\text{val}_{i,j} = \begin{cases} 1 & \text{for every position } j \leq \text{pos}_{\succ_i}(a') \\ 0 & \text{otherwise.} \end{cases}$$

Thereby, $\text{SW}(a', v) \geq t/2$.

Case 2: $a > (k - 1/2)t$. Let $a' = a - 1$ and note that

$$|N_{\leq t}(a') \cap \mathcal{C}_k| \geq t/2.$$

Since $\mathcal{M}$ did not query any agent of cohort $\mathcal{C}_k$ at a position $j \leq t$, we are free to fix the concealed values for every agent $i \in N_{\leq t}(a') \cap \mathcal{C}_k$ such that

$$\text{val}_{i,j} = \begin{cases} 1 & \text{for every position } j \leq \text{pos}_{\succ_i}(a') \\ 0 & \text{otherwise.} \end{cases}$$

Thereby, $\text{SW}(a', v) \geq t/2$ in this case as well.

In both cases, the remaining concealed values are set to 0 except for those positions where a value of 1 is implied by a revealed value of 1 at a position further down in the ranking. In particular, this means that $\text{val}_i(a) = 0$ for every agent $i \in N_{\leq t}(a) \cap \mathcal{C}_k$. By

assumption, $\mathcal{M}$ did not query any of these agents at a position $j \leq t$. In case 1, among cohort $\mathcal{C}_k$, only the agents $N_{\leq t}(a') \setminus N_{\leq t}(a) = \mathcal{C}_k \setminus N_{\leq t}(a)$ assign a value of 1 to any alternative. In case 2, among cohort $\mathcal{C}_k$, only the agents $N_{\leq t}(a') \cap \mathcal{C}_k$ have a value of 1 for any alternative and exclusively for those alternatives ranked above a . Finally, by our construction, there is at most one other cohort k' such that $N_{\leq t}(a) \cap \mathcal{C}_{k'} \neq \emptyset$. Hence, $\text{SW}(a, v) \leq 1$ and the lemma follows. $\square$

Lemma 26. *Assume that mechanism $\mathcal{M}$ queries every cohort at at least one position $j \leq t$. Then, there exists a choice of valuations v that is consistent with profile P and the values revealed to $\mathcal{M}$, such that*

$$\frac{\max_{a \in A} \text{SW}(a, v)}{\text{SW}(\mathcal{M}(P, v), v)} \geq t/2 - 1.$$

Proof. Since no agent gives any value to alternatives $t^2 + 1, \dots, m$, the distortion is infinite when $\mathcal{M}(P, v) > t^2$. Hence, assume that there is a cohort $\mathcal{C}_k$ such that $\mathcal{M}(P, v) = a \in \mathcal{C}_k$. Consider the alternative $a' = a - 1 \pmod{t^2}$. The set of agents $N_{\leq t}(a)$ and $N_{\leq t}(a')$ intersect with at most two cohorts, namely, cohort k and cohort $k' = k - 1 \pmod{t^2}$. For every cohort $\mathcal{C}_\ell$ where $\ell \neq k, k'$, it holds by our assumption that there is an agent i_ℓ who is the *first agent in this cohort* to be queried by $\mathcal{M}$ at a position $j_\ell \leq t$. At this position, there is an alternative other than a or a' , and we revealed the value $\text{val}_{i_\ell, j_\ell} = 1$; see above. Since $\mathcal{M}$ can make at most one query to each agent, we can set the remaining concealed values such that

$$\text{val}_{i_\ell, j} = \begin{cases} 1 & \text{for every position } j \leq \text{pos}_{\succ_{i_\ell}}(a') \\ 0 & \text{otherwise} \end{cases}$$

for every $\ell \neq k, k'$. This implies that $\text{val}_{i_\ell}(a) = 0$ for every $\ell \neq k, k'$. The remaining concealed values are set to 0 except for those positions in $\mathcal{C}_k, \mathcal{C}_{k'}$ where a value of 1 is implied by a revealed value of 1 at a position further down in the ranking. Then, $\text{SW}(a, v) \leq 2$ since $N_{\leq t}(a)$ can intersect only with cohorts $\mathcal{C}_k$ and $\mathcal{C}_{k'}$. On the other hand, by setting the concealed values for every cohort $\ell \neq k, k'$ as described, it holds that $\text{SW}(a', v) \geq t - 2$. From this, the lemma follows. $\square$

Theorem 24 now follows by combining Lemmas 25 and 26. $\square$

5.7 Discussion and Open Problems

We have initiated the study of average distortion in a simple stochastic setting that creates impartial culture electorates. The main open problem is whether constant average distortion is possible with a small number of queries per agent for general probability distributions of valuations. Throughout the paper, we assume that the distribution is given as part of the input, and this information is crucial to make our mechanisms work. It would be interesting to explore whether this—admittedly strong—requirement can be removed. Other natural extensions of our model include different

distributions per alternative or distributions that produce random valuations satisfying the unit-sum or unit-range assumption. These latter assumptions are clearly beyond the reach of our current analysis techniques, as they necessarily imply correlations between the random values an agent has for the alternatives.

For the worst-case setting, we have improved the previously best-known lower bound on the number of queries per agent that are necessary for constant worst-case distortion by deterministic mechanisms. Still, the conjecture of Amanatidis et al. [4] that constant worst-case distortion is possible with $\Theta(\log m)$ deterministic queries per agent is wide open. Furthermore, we have also demonstrated that the use of randomization can yield worst-case bounds that the known deterministic mechanisms cannot achieve. Exploring whether there is a separation between deterministic and randomized mechanisms in terms of their worst-case distortion for a given number of queries per agent is another challenging problem that deserves investigation.

5.8 Appendix to Chapter 5

Proof of Lemma 13

For the proof of Lemma 13, we present another technical statement as Lemma 27. The latter lemma formalizes the intuition that the expectation for the maximum social welfare only increases when a given number of values of 1 is distributed more evenly between two agents. The proof of the lemma leans heavily on notation. We therefore highlight two key properties that the proof of Lemma 27 exploits. These properties depend crucially on the fact that the agents' rankings are uniformly random and independent. Let r be any non-negative integer.

- **Property 1:** Let $u = (u_1, u_2, \dots, u_n)$ be a vector with non-negative integer entries. Now, consider the question whether, for a random draw of $v \sim F$, $P \sim \mathcal{P}(v)$, there exists an alternative $a \in A$ that appears in the rankings of at least r agents such that for each of these agents i it holds that $\text{pos}_{\succ_i}(a) \leq u_i$. Whether or not such an alternative exists does not depend on the number of values of 1 underlying the agents' rankings.
- **Property 2:** Let a_1 be *any* alternative appearing in *any* position in the ranking of agent 1 and let a_2 be *any* alternative appearing in *any* position in the ranking of agent 2. The probability of having social welfare of exactly r for agents $3, \dots, n$ is exactly the same for a_1 and a_2 .

We continue with the statement and formal proof of Lemma 27.

Lemma 27. Consider two vectors $t = (t_1, t_2, \dots, t_n)$ and $t' = (t'_1, t'_2, \dots, t'_n)$ with non-negative integer entries such that $t_j \geq t_k + 2$ for two entries $j, k \in [n]$ and $t'_j = t_j - 1$, $t'_k = t_k + 1$ and $t'_i = t_i$ for all $i \in [n] \setminus \{j, k\}$. Then,

$$\mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) | X_i(v) = t_i, i \in [n] \right] \leq \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) | X_i(v) = t'_i, i \in [n] \right].$$

Proof. Without loss of generality, we assume that $j = 1$ and $k = 2$. For each choice of vector $\hat{t} \in \{t, t'\}$, we will abbreviate the event $X_i(v) = \hat{t}_i$ for $i \in [n]$ by $C(\hat{t})$. Define the vector $u = (u_1, u_2, \dots, u_n) = (t_1 - 1, t_2, t_3, \dots, t_n)$ and, for a non-negative integer r , let $M(r)$ be the event that there is an alternative $a \in A$ such that $\sum_{i \in [n]} \mathbb{I}\{\text{pos}_{\succ_i}(a) \leq u_i\} \geq r$. We denote by $\bar{M}(r)$ the complement event. By the properties of the expectation and using the law of total expectation, we have

$$\begin{aligned} & \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) | C(\hat{t}) \right] \\ &= \sum_{r=1}^n \Pr_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \geq r | C(\hat{t}) \right] \\ &= \sum_{r=1}^n \left(\Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [M(r) | C(\hat{t})] \cdot \Pr_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \geq r | M(r), C(\hat{t}) \right] \right. \\ & \quad \left. + \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\bar{M}(r) | C(\hat{t})] \cdot \Pr_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \geq r | \bar{M}(r), C(\hat{t}) \right] \right) \\ &= \sum_{r=1}^n \left(\Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [M(r)] + \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\bar{M}(r)] \cdot \Pr_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \geq r | \bar{M}(r), C(\hat{t}) \right] \right). \end{aligned}$$

On the last line, the conditions $M(r)$ and $C(\hat{t})$ combined imply that there is indeed an alternative with social welfare at least r . We also used the observation (i.e., property 1 above) that the conditions $M(r)$ and $\bar{M}(r)$ only depend on the realization of the agents' (uniformly random) rankings. Now, notice that in the previous equality only the term $\Pr[\max_{a \in A} \text{SW}(a, v) \geq r | \bar{M}(r), C(\hat{t})]$ depends on the choice of $\hat{t}$ among t, t' .

In the following, we distinguish between $\hat{t} = t$ (case 1) and $\hat{t} = t'$ (case 2). Let a_1, a_2 be the random alternatives appearing in position t_1 of agent 1's ranking and position $t_2 + 1$ of agent 2's ranking. Under the conditions $\bar{M}(r), C(\hat{t})$ combined, only the alternatives a_1, a_2 may have social welfare of at least r .

Case 1. $\hat{t} = t = (t_1, t_2, t_3, \dots, t_n)$. Alternative a_2 has social welfare of at most $r - 1$ due to $\bar{M}(r), C(t)$ and the fact that $\text{pos}_{\succ_2}(a_2) > t_2$. Hence, only alternative a_1 can have social welfare of at least r such that

$$\begin{aligned} & \Pr_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \geq r | \bar{M}(r), C(t) \right] \\ &= \Pr_{v \sim F} [\text{SW}(a_1, v) \geq r | \bar{M}(r), C(t)] \\ &= \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[1 + \mathbb{I}\{\text{pos}_{\succ_2}(a_1) \leq t_2\} + \sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_1) \leq t_i\} \geq r | \bar{M}(r) \right] \\ &= \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\mathbb{I}\{\text{pos}_{\succ_2}(a_1) \leq t_2\} + \sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_1) \leq t_i\} = r - 1 \right]. \end{aligned}$$

The last equality follows since, under the condition $\overline{M}(r)$, alternative a_1 can appear at positions t_i or above for at most $r-1$ agents among the agents $2, 3, \dots, n$. Applying the law of total probability yields

$$\begin{aligned}
& \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\mathbb{I}\{\text{pos}_{\succ_2}(a_1) \leq t_2\} + \sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_1) \leq t_i\} = r-1 \right] \\
&= \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_1) \leq t_i\} = r-1 \right] \cdot 1 \\
&\quad + \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_1) \leq t_i\} = r-2 \right] \cdot \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} [\mathbb{I}\{\text{pos}_{\succ_2}(a_1) \leq t_2\} = 1] \\
&\quad + \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_1) \leq t_i\} < r-2 \right] \cdot 0 \\
&= \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_1) \leq t_i\} = r-1 \right] \\
&\quad + \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_1) \leq t_i\} = r-2 \right] \cdot \frac{t_2}{m}.
\end{aligned}$$

Case 2. $\hat{t} = t' = (t_1 - 1, t_2 + 1, t_3, \dots, t_n)$. Similar to the previous case, we obtain

$$\begin{aligned}
& \Pr_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \geq r | \overline{M}(r), C(t') \right] \\
&= \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\mathbb{I}\{\text{pos}_{\succ_1}(a_2) \leq t_2\} + \sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_2) \leq t_i\} = r-1 \right] \\
&= \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_2) \leq t_i\} = r-1 \right] \\
&\quad + \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_2) \leq t_i\} = r-2 \right] \cdot \frac{t_1 - 1}{m}.
\end{aligned}$$

We now use the equalities obtained for the two cases to show that, for every $r \in [n]$, the probability $\Pr[\max_{a \in A} \text{SW}(a, v) \geq r | \overline{M}(r), C(\hat{t})]$ is greater for $\hat{t} = t'$ than for $\hat{t} = t$ which proves the lemma.

As highlighted above by property 2, it holds that

$$\Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_1) \leq t_i\} = r-1 \right] = \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_2) \leq t_i\} = r-1 \right],$$

and

$$\Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_1) \leq t_i\} = r-2 \right] = \Pr_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\sum_{i=3}^n \mathbb{I}\{\text{pos}_{\succ_i}(a_2) \leq t_i\} = r-2 \right]$$

due to the agents' rankings being uniformly random. Finally, by the assumption that $t_1 \geq t_2 + 2$, we have that $(t_1 - 1)/m > t_2/m$ such that, for every r ,

$$\Pr_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \geq r | \overline{M}(r), C(t') \right] > \Pr_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \geq r | \overline{M}(r), C(t) \right]$$

which concludes the proof. $\square$

We are now ready to prove Lemma 13. Starting from any vector t with non-negative integer entries and applying Lemma 27 repeatedly, we obtain that

$$\begin{aligned} & \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) | X_i(v) = t_i, i \in [n] \right] \\ & \leq \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) | X_i(v) \in \{\lfloor s/n \rfloor, \lceil s/n \rceil\}, i \in [n], X(v) = s \right] \\ & = \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) | \text{Balanced}(v, s) \right]. \end{aligned}$$

Then,

$$\begin{aligned} & \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) | X(v) = s \right] \\ & = \sum_{\substack{t=(t_1, \dots, t_n): \\ \sum_{i \in [n]} t_i = s}} \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) | X_i(v) = t_i, i \in [n] \right] \cdot \Pr_{v \sim F} [X_i(v) = t_i, i \in [n] | X(v) = s] \\ & \leq \sum_{\substack{t=(t_1, \dots, t_n): \\ \sum_{i \in [n]} t_i = s}} \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) | \text{Balanced}(v, s) \right] \cdot \Pr_{v \sim F} [X_i(v) = t_i, i \in [n] | X(v) = s] \\ & = \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) | \text{Balanced}(v, s) \right], \end{aligned}$$

implying the lemma. $\square$

Proof of Lemma 14

In the proof of Lemmas 14 and 15, we will use the following simple claim.

Claim 28. *Let k be a positive integer, S and S' sets of agents with $|S| \leq |S'|$, and T_i a set of k consecutive positions in agent $i \in S$. Then,*

$$\mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \sum_{i \in S} \mathbb{I}\{\text{pos}_{\succ_i}(a) \in T_i\} \right] \leq \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \sum_{i \in S'} \mathbb{I}\{\text{pos}_{\succ_i}(a) \leq k\} \right]$$

To see why the claim holds, observe that the expectations are defined over uniformly random profiles. Hence, the probability that an alternative appears in any position of any agents' ranking is equal to $1/m$, i.e., independent from both the agent and the position.

Returning to the proof of Lemma 14, define the sets of agents $N_1 = \{1, \dots, \lfloor n/2 \rfloor\}$, $N_2 = \{1, \dots, \lfloor n/2 \rfloor + 1, \dots, 2\lfloor n/2 \rfloor\}$, and $N_3 = N \setminus (N_1 \cup N_2)$. We have

$$\begin{aligned}
& \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \mid \text{Balanced}(v, jn\tau) \right] \\
&= \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \sum_{i \in N} \mathbb{I}\{\text{pos}_{\succ_i}(a) \leq j\tau\} \right] \\
&= \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \sum_{\ell \in [3]} \sum_{i \in N_\ell} \sum_{k=1}^j \mathbb{I}\{(k-1)\tau < \text{pos}_{\succ_i}(a) \leq k\tau\} \right] \\
&\leq \sum_{\ell \in [3]} \sum_{k=1}^j \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \sum_{i \in N_\ell} \mathbb{I}\{(k-1)\tau < \text{pos}_{\succ_i}(a) \leq k\tau\} \right] \\
&\leq \sum_{\ell \in [3]} \sum_{k=1}^j \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \sum_{i \in N_1} \mathbb{I}\{\text{pos}_{\succ_i}(a) \leq \tau\} \right] \\
&= 3j \cdot \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \mid \text{Box}(v) \right] = 3j \cdot B.
\end{aligned}$$

The first equality is due to the fact that the condition $\text{Balanced}(v, jn\tau)$ requires that the contribution to the social welfare of each alternative comes from positions from 1 to $j\tau$ only. The first inequality uses a simple property of the maximum function and linearity of expectation, while the second one follows from Claim 28. The last equality follows since the condition $\text{Box}(v)$ requires that the contribution to the social welfare of each alternative comes from positions from position 1 to τ of exactly $\lfloor n/2 \rfloor$ agents. The last equality uses the definition of B . $\square$

Proof of Lemma 15

Let $N_1 = [t]$. Define the condition $\text{PartiallyBalanced}(v, t, s)$ to be $Z(v, 1) = t$, $\lfloor s/t \rfloor \leq X_i(v) \leq \lceil s/t \rceil$ for $i \in [N_1]$ and $X_i(v) = 0$ for $i \in N \setminus N_1$. Using Lemma 27 in a similar way we used it to prove Lemma 13, we can show that

$$\begin{aligned}
& \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \mid Z(v, 1) = t, X(v) = j \right] \\
&\leq \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \mid \text{PartiallyBalanced}(v, t, j) \right].
\end{aligned}$$

So, we have

$$\begin{aligned}
& \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \mid Z(v, 1) = t, X(v) = j \right] \\
& \leq \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \mid \text{PartiallyBalanced}(v, t, j) \right] \\
& \leq \mathbb{E}_{v \sim F} \left[\max_{a \in A} \text{SW}(a, v) \mid \text{PartiallyBalanced}(v, t, \lceil j/t \rceil \cdot t) \right] \\
& = \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \sum_{i \in N_1} \mathbb{I}\{\text{pos}_{\succ_i}(a) \leq \lceil j/t \rceil\} \right] \\
& = \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \sum_{i \in N_1} \sum_{k=1}^{\lceil j/t \rceil} \mathbb{I}\{\text{pos}_{\succ_i}(a) = k\} \right] \\
& \leq \sum_{k=1}^{\lceil j/t \rceil} \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \sum_{i \in N_1} \mathbb{I}\{\text{pos}_{\succ_i}(a) = k\} \right] \\
& \leq \lceil j/t \rceil \cdot \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \sum_{i \in N_1} \mathbb{I}\{\text{pos}_{\succ_i}(a) = 1\} \right] \\
& = \lceil j/t \rceil \cdot \mathbb{E}_{\substack{v \sim F \\ P \sim \mathcal{P}(v)}} \left[\max_{a \in A} \text{SW}_{\text{im}}(a, P, v, 1) \mid Z(v, 1) = t \right].
\end{aligned}$$

The second inequality holds since the condition $\text{PartiallyBalanced}(v, t, \lceil j/t \rceil \cdot t)$ implies the condition $\text{PartiallyBalanced}(v, t, j)$ and the quantity $\mathbb{E}[\max_{a \in A} \text{SW}(a, v) \mid \text{PartiallyBalanced}(v, t, j)]$ is non-decreasing in j . The first equality is due to the fact that condition $\text{PartiallyBalanced}(v, jn\tau)$ requires that the contribution to the social welfare of each alternative comes from positions 1 to $\lceil j/t \rceil$ of t agents only. Next, the fourth inequality follows from a simple property of the maximum function and linearity of expectation, while the fifth inequality follows from Claim 28. Finally, for the last inequality, observe that under the condition $Z(v, 1) = t$, the implied social welfare of each alternative comes from the top position of a set of t agents. $\square$

Two Comments Regarding Mechanism MEAN and its Analysis

We devote this section to discussing two issues related to mechanism MEAN. First, notice that MEAN queries the value of the top-ranked alternative in each agent for $p < 2/m$ and an alternative at a lower position in the agents' rankings otherwise. Let us consider the simpler variant of MEAN which always queries the value of each top-ranked alternative and returns the one of maximum implied social welfare. Our Lemma 29 below shows that this variant of MEAN has super-constant distortion when ties in implied social welfare are resolved arbitrarily, that is, ignoring the content of the profile below the top position in each ranking.

Our proof uses a profile with m alternatives, $n = m^{1/3}$ agents, and $p = m^{-1/3}$. In this way, there are, on average, $m^{2/3}$ alternatives that have a value of 1 in each

agent, but the mechanism recovers very little information about which alternatives do get a value of 1. In particular, our choice of parameters n , m , and p guarantees that all top-ranked alternatives are different with high probability. Then, the alternative selected among them by the mechanism has only constant expected social welfare. In contrast, the parameters are such that the expected maximum social welfare is $\Omega\left(\frac{\log m}{\log \log m}\right)$.

Lemma 29. *Let $\mathcal{M}$ be the mechanism that queries the value of the top-ranked alternative in every agent and returns an alternative that maximizes the implied social welfare (breaking ties arbitrarily). It holds that $\text{avdist}(\mathcal{M}, \mathcal{F}^{0/1}) \in \Omega\left(\frac{\log m}{\log \log m}\right)$.*

Proof. Let $n \geq 2$, $m = n^3$, and consider the binary distribution F_p with $p = 1/n = m^{-1/3}$. We first lower-bound the probability that all alternatives appearing in the first position of any agent's ranking are different. For a set of binary valuations v , we refer to this condition as $\text{Distinct}(v)$. Recall that, in an impartial culture electorate, the top-ranked alternative of each agent is, in effect, uniformly random and independent from the top-ranked alternatives of the remaining agents. We thus have that

$$\begin{aligned} \Pr_{v \sim F_p} [\text{Distinct}(v)] &= \prod_{i=0}^{n-1} \frac{m-i}{m} = \prod_{i=0}^{m^{1/3}-1} \left(1 - \frac{i}{m}\right) \geq \left(1 - m^{-2/3}\right)^{m^{1/3}} \\ &\geq 1 - m^{-1/3}. \end{aligned} \quad (5.24)$$

Here, we used the fact that $(1+x)^r \geq 1+rx$ for $x \geq -1, r \in \mathbb{R} \setminus (0, 1)$.

Now, notice that the probability that an alternative gets a value of 1 from an agent, given that it is not top-ranked by this agent, is at most p . Then, the expected social welfare of the alternative returned by mechanism $\mathcal{M}$ under the condition $\text{Distinct}(v)$ is at most $1 + p(n-1) \leq 2$. Thus,

$$\begin{aligned} &\mathbb{E}_{\substack{v \sim F_p \\ P \sim \mathcal{P}(v)}} [\text{SW}(\mathcal{M}(P, v), v)] \\ &\leq \Pr_{v \sim F_p} [\text{Distinct}(v)] \cdot \mathbb{E}_{\substack{v \sim F_p \\ P \sim \mathcal{P}(v)}} [\text{SW}(\mathcal{M}(P, v), v) | \text{Distinct}(v)] \\ &\quad + (1 - \Pr_{v \sim F_p} [\text{Distinct}(v)]) \cdot n \\ &\leq 2 + m^{-1/3} \cdot n = 3. \end{aligned} \quad (5.25)$$

The first inequality follows from the law of total expectation together with the observation that the expected social welfare of any alternative is trivially upper-bounded by n under binary valuations. For the second inequality, we use inequality (5.24).

We complete the proof by showing that the expected maximum social welfare is at least $\Omega\left(\frac{\log m}{\log \log m}\right)$. We do so by proving that the probability that the maximum social

welfare is at least $\left\lfloor \frac{\log m}{\log \log m} \right\rfloor$ is lower-bounded by a constant. First, notice that

$$\begin{aligned} \Pr_{v \sim F_p} [\text{SW}(a, v) \geq k] &\geq \Pr_{v \sim F_p} [\text{SW}(a, v) = k] = \binom{n}{k} \cdot p^k \cdot (1-p)^{n-k} \\ &\geq \left(\frac{m^{1/3}}{k} \right)^k \cdot m^{-k/3} \cdot (1 - m^{-1/3})^{m^{1/3}} = \frac{(1 - m^{-1/3})^{m^{1/3}}}{k^k}, \end{aligned}$$

where the second inequality follows since $\binom{n}{k} \geq \left(\frac{n}{k}\right)^k$. We now observe that $(1 - m^{-1/3})^{m^{1/3}}$ is strictly increasing in m approaching e^{-1} from below. For $n \geq 2$, we have that $m \geq 8$ and, thus, $(1 - m^{-1/3})^{m^{1/3}} \geq \frac{1}{4}$. This yields $\Pr_{v \sim F_p} [\text{SW}(a, v) \geq k] \geq \frac{1}{4k^k}$ and, hence

$$\Pr_{v \sim F_p} \left[\max_{a \in A} \text{SW}(a, v) \geq k \right] \geq 1 - \left(1 - \frac{1}{4k^k} \right)^m \geq 1 - \exp \left(-\frac{m}{4k^k} \right), \quad (5.26)$$

where the last inequality is due to the fact that $e^x \geq 1 + x$ for any real x . By selecting $k = \left\lfloor \frac{\log m}{\log \log m} \right\rfloor$, we have $k^k \leq m$ and

$$\Pr_{v \sim F_p} \left[\max_{a \in A} \text{SW}(a, v) \geq \left\lfloor \frac{\log m}{\log \log m} \right\rfloor \right] \geq 1 - e^{-1/4} > 0.22,$$

as desired. □

One may still wonder whether concentration inequalities like Chernoff bounds could replace (parts of) our analysis of mechanism MEAN. Intuitively, if the expected social welfare of each alternative is high (e.g., $np \in \Omega(\log m)$), then the social welfare of all alternatives will be sharply concentrated around this expectation, and the expected maximum and expected minimum social welfare will only be a constant factor apart. This would imply constant average distortion for *all* mechanisms, including those that make no queries at all. We include a formal proof of this fact as Lemma 30 below. Unfortunately, the assumption that $np \in \Omega(\log m)$ required by the lemma does not subsume any of the three cases in our analysis of mechanism MEAN in Section 5.4. Furthermore, we do not see how to extend the use of concentration inequalities to a broader range of parameters (e.g., satisfying $np \in o(\log m)$) where we have proven that queries are necessary.

Lemma 30. *Any voting rule has average distortion at most 13 in impartial culture electorates with n agents and m alternatives, and underlying values drawn from a binary distribution F_p such that $n \cdot p \geq 8 \ln(2m)$.*

Proof. We will show that returning any alternative (including the one of minimum social welfare) yields an average distortion of at most 13. We will use the upper tail Chernoff bound (Lemma 9) as well as its next lower tail version.

Lemma 31 (Chernoff bound, lower tail). *For every binomial random variable Q and any $\delta \in [0, 1]$, we have*

$$\Pr[Q \leq (1 - \delta) \mathbb{E}[Q]] \leq \exp\left(-\frac{\delta^2 \mathbb{E}[Q]}{2}\right).$$

For valuations v drawn according to F_p and an alternative $a \in A$, let $Q_i(v)$ be a random variable that indicates whether agent i has value 1 (then, $Q_i(v) = 1$) or 0 (then, $Q_i(v) = 0$) for this alternative. Define $Q(v) = \sum_{i \in N} Q_i(v)$ and $\mu = \mathbb{E}_{v \sim F_p}[Q(v)]$. Thereby, Q is the social welfare of alternative a , and μ is its expected value. By our assumption, it holds that $\mu = n \cdot p \geq 8 \ln(2m)$.

Notice that for $\delta \geq 2$, we have $\frac{\delta^2}{2+\delta} \geq \frac{\delta}{2}$ and Lemma 9 now implies that

$$\Pr_{v \sim F_p}[Q(v) \geq (1 + \delta)\mu] \leq \exp\left(-\frac{\delta\mu}{2}\right). \quad (5.27)$$

For any $t \geq 3\mu$, we use the union bound and apply inequality (5.27) with $\delta = \frac{t}{\mu} - 1 \geq 2$ to obtain

$$\Pr_{v \sim F_p}\left[\max_{a \in A} \text{SW}(a, v) \geq t\right] \leq m \cdot \Pr_{v \sim F_p}[Q(v) \geq t] \leq m \cdot \exp\left(-\frac{t - \mu}{2}\right). \quad (5.28)$$

By the definition of the expectation and using inequality (5.28), it holds that

$$\begin{aligned} \mathbb{E}_{v \sim F_p}\left[\max_{a \in A} \text{SW}(a, v) \geq t\right] &= \int_0^\infty \Pr_{v \sim F_p}\left[\max_{a \in A} \text{SW}(a, v) \geq t\right] dt \\ &\leq \int_0^{3\mu} dt + \int_{3\mu}^\infty \Pr_{v \sim F_p}\left[\max_{a \in A} \text{SW}(a, v) \geq t\right] dt \\ &\leq 3\mu + m \int_{3\mu}^\infty \exp\left(-\frac{t - \mu}{2}\right) dt \\ &= 3\mu + 2m \exp(-\mu) \leq 3\mu + (2m)^{-7}, \end{aligned} \quad (5.29)$$

where the last inequality follows since $\mu \geq 8 \ln(2m)$.

Now, using again the union bound and the lower tail Chernoff bound (Lemma 31), we get

$$\Pr_{v \sim F_p}\left[\min_{a \in A} \text{SW}(a, v) \leq \frac{\mu}{2}\right] \leq m \cdot \Pr_{v \sim F_p}\left[Q(v) \leq \frac{\mu}{2}\right] \leq m \cdot \exp\left(-\frac{\mu}{8}\right) \leq \frac{1}{2}.$$

Thus,

$$\mathbb{E}_{v \sim F_p}\left[\min_{a \in A} \text{SW}(a, v)\right] \geq \frac{\mu}{2} \cdot \Pr_{v \sim F_p}\left[\min_{a \in A} \text{SW}(a, v) \geq \frac{\mu}{2}\right] > \frac{\mu}{4}. \quad (5.30)$$

By inequalities (5.29) and (5.30), we obtain that the ratio between $\mathbb{E}_{v \sim F_p}[\max_{a \in A} \text{SW}(a, v)]$ and $\mathbb{E}_{v \sim F_p}[\min_{a \in A} \text{SW}(a, v)]$ and, consequently, the distortion of any mechanism is at most 13, as desired. $\square$

Chapter 6

Low-Distortion Clustering with Ordinal and Limited Cardinal Information

Motivated by recent work in computational social choice, we extend the metric distortion framework to clustering problems. Given a set of n agents located in an underlying metric space, our goal is to partition them into k clusters, optimizing some social cost objective. The metric space is defined by a distance function d between the agent locations. Information about d is available only implicitly via n rankings, through which each agent ranks all other agents in terms of their distance from her. Still, even though no cardinal information (i.e., the exact distance values) is available, we would like to evaluate clustering algorithms in terms of social cost objectives that are defined using d . This is done using the notion of distortion, which measures how far from optimality a clustering can be, taking into account all underlying metrics that are consistent with the ordinal information available.

Unfortunately, the most important clustering objectives (e.g., those used in the well-known k -median and k -center problems) do not admit algorithms with finite distortion. To sidestep this disappointing fact, we follow two alternative approaches: We first explore whether resource augmentation can be beneficial. We consider algorithms that use more than k clusters but compare their social cost to that of the optimal k -clustering. In this dissertation, we focus on our results for the k -center objective. We show that using exponentially (in terms of k) many clusters, we can get constant distortion. Interestingly, such an exponential blowup is shown to be necessary. More importantly, we explore whether limited cardinal information can be used to obtain better results. Somewhat surprisingly, we show that a number of queries that is linear in k is enough to get constant distortion under the k -center objective.

6.1 Introduction

The typical computational social choice problem consists of optimizing a function over alternatives, each with a different associated cost or value. A classic example is given by representative election. Each voter has a different representation score for every candidate, which we assume to correspond to the distance in some underlying metric. Ideally, the representation minimizes the sum of distances of each voter to their closest representative. In the full information setting, this corresponds to solving the classic k -median problem. But this example already illustrates the difficulty of implementing any voting mechanism: Even if the representation scores are assumed to be distances, they might be unknown even to the participating voters. However, we may readily know if a voter prefers alternative a over alternative b .

Such examples have given rise to *ordinal algorithms*. An ordinal algorithm mainly allows for comparisons between distances in the underlying metric. That is, given three points a, b, c , we are freely given information whether $d(a, b) \leq d(a, c)$, but we are not given the exact numerical values of $d(a, b)$ and $d(a, c)$. The objective is to solve a given problem relying primarily on the ordinal information, while using as few (ideally zero) distance queries as possible. The goodness of such an algorithm is measured in terms of the quality of the computed solution C compared to the quality of the optimal solution OPT that is given full information, commonly known as the *metric distortion*.

Finding the median is arguably the most important problem in this field. Given a set of points X and a distance function d , the median m is defined to be the point minimizing the sum of distances. Following a long line of work [11, 12, 67, 75, 87, 103], there now exists a deterministic algorithm with optimal metric distortion 3 [74], which is also optimal [10]. Using randomization, Charikar et al. [43] recently achieved an important breakthrough, achieving a metric distortion of 2.753. The best known lower bound is at least 2.1126 [42].

Extensions to more general clustering objectives such as (k, z) -clustering and facility location are comparatively much harder, see Anshelevich and Zhu [8], Caragiannis et al. [38]. In facility location, we ask for a set of centers C such that

$$\sum_{x \in X} \min_{c \in C} d(x, c) + f \cdot |C|$$

is minimized, where f is the cost of opening a center. For (k, z) -clustering, we instead consider the objective

$$\sqrt[z]{\sum_{x \in X} \min_{c \in C} d(x, c)^z},$$

i.e., the algorithm does not incur a cost for opening the centers, but instead has a budget of at most k centers that can be placed. Special cases include k -median where $z = 1$ and k -center which corresponds to $z \rightarrow \infty$.¹

¹Sometimes the $\sqrt[z]{\cdot}$ operation is omitted, as is the case for k -means which corresponds to $(k, 2)$ -clustering. An α -approximation to $\sqrt[z]{\sum_{x \in X} \min_{c \in C} d(x, c)^z}$ implies an α^z -approximation to $\sum_{x \in X} \min_{c \in C} d(x, c)^z$.

Unfortunately, there are strong impossibility results for purely ordinal algorithms. Even for 2-median, it is not possible to obtain an algorithm with bounded metric distortion [8]. Therefore, research has begun to design algorithms that are given more power than purely ordinal information. Indeed, there has been some recent success in providing guarantees using only a constant number of queries per point, see Amanatidis et al. [5, 6]. For clustering, recent work by Pulyassary [114] has shown that using at most $\text{polylog}(n)$ distance queries per point, or $n \cdot \text{polylog}(n)$ queries overall, it is possible to achieve a constant factor approximation. The same work also showed that k queries per point, or $O(nk)$ queries overall are sufficient to achieve a constant factor approximation for k -median. Thus, we ask:

Question 32. *What is the minimum number of queries necessary for an algorithm to achieve constant metric distortion for k -median, k -center, and facility location?*

While distance queries are a natural way of lending more power to the algorithm designer, obtaining the distances may be expensive as mentioned above. This leads to the question whether other models exist that allow the algorithm designer to bound the metric distortion. A very natural way of doing so for clustering algorithms is by allowing the algorithm to return a (α, β) -bicriteria approximation. Such algorithms bound the clustering cost by at most α times the cost of an optimal k -clustering, while using β many centers. We ask:

Question 33. *What is the minimum value of β such that a bicriteria clustering algorithm using only ordinal information has constant metric distortion?*

Our Results

The work that we present in the following together with our remaining preliminary results [28] makes substantial progress towards answering both questions. In the low-query setting, we give two deterministic polynomial time algorithms for k -center that, using at most $O(k^2)$ overall distance evaluations, obtain a distortion of 2 and, using at most $O(k)$ overall distance evaluations, achieve distortion of 4. We show that the latter result is optimal in terms of the number of necessary queries. Furthermore, in our preliminary work [28], we present a randomized polynomial time algorithm for (k, z) -clustering that uses at most $\text{poly}(k, \log n)$ overall distance queries and achieves constant metric distortion. Note that all of these bounds are sub-linear in the input size, that is assuming $k \ll n$, we make $o(1)$ queries per point. Finally, for facility location, there exists a simple adaptation of the seminal Meyerson algorithm [101] that achieves a constant distortion using exactly one query per point or n queries overall, see also Section 4.1 of Pulyassary [114]. We show that no algorithm can achieve a constant factor approximation using less than $\Omega(n)$ queries, effectively closing the problem.

Related Work

Ordinal Preferences and Distortion. The first paper to consider optimization problems using ordinal information was probably Procaccia and Rosenschein [111].

Subsequently, two main directions have been established. Continuing to work with the model introduced by Procaccia and Rosenschein, one line focuses mainly on maximizing welfare subject to normalization assumptions, but without assuming any metric properties, see Amanatidis et al. [4, 5, 6], Caragiannis and Procaccia [33], Filos-Ratsikas et al. [68]. The other line of work studies problem without the normalization assumptions, but assuming that the preferences are metric, i.e., they satisfy the triangle inequality. Beyond clustering papers covered in the introduction, several other distortion problems have been studied [24, 45, 46, 108]. While rare, it is also possible to achieve some results without making either a normalization or metric assumptions, see Abramowitz and Anshelevich [1].

Clustering and Facility Location. (k, z) -clustering is APX-hard in general metrics [49], though it is possible to obtain very accurate algorithms when making assumptions on either the metric [50, 71] or the input [7, 17, 47]. For k -center, Gonzalez [78] gave an optimal 2-approximation algorithm. For k -median, k -means and facility location, following a long line of research [16, 52, 53, 82, 83, 97], the current state of the art is a 2.613 approximation for k -median [79], a 9 approximation for k -means [2], and a 1.488 approximation for facility location [96]. For general (k, z) -clustering, there are few claimed bounds, though most of the proofs for k -median and k -means go through while losing a $\exp(z)$ approximation factor. Explicit results can be found in Cohen-Addad et al. [48, 51].

6.2 Preliminaries

Let (X, d) be a metric space where X is a set of n points and $d : X \times X \rightarrow \mathbb{R}_{\geq 0}$ is a metric. The distance between any two points $x, y \in X$ can be accessed by a *query* for the value $d(x, y)$. We assume that such a query is associated with a cost. An algorithm is given a budget and each query that the algorithm makes consumes one unit of its budget. While querying the exact distance between two points is costly, our model assumes that, for every point, *ordinal* information about its relative distance to the other points is freely available. More specifically, each point $x \in X$ provides a strict total order $\succ_x$ on X which we refer to as the *preference list* (or *ranking*) of x . Each ranking $\succ_x$ is *consistent* with the distance function d in the sense that, for any two points $y, z \in X$, $y \succ_x z$ only if $d(x, y) \leq d(x, z)$. That is, in a pairwise comparison between y and z , x *prefers* the point that is closest. An ordinal *preference profile* P is then just the collection of the points' rankings, i.e., $P = \{\succ_x\}_{x \in X}$. We write $\mathcal{P}(d)$ for the set of profiles where each point's ranking is consistent with the distances d .

For a point $x \in X$, the function $\pi_x : X \rightarrow [n]$ maps any point $y \in X$ to the *position* of y in the ranking of x . That is, the position $\pi_x(y) = i$ indicates there are exactly $i - 1$ points z such that $z \succ_x y$, and exactly $n - i$ points z such that $y \succ_x z$. We say that the former set consists of the points *ranked higher* by x than y , and the latter set comprises those points *ranked lower* than y . Furthermore, it is often convenient to restrict π_x to a

certain subset of X . Let $S \subseteq X$ and $m = |S|$. The *restriction of π_x to S* is a function $\pi_{x,S} : S \rightarrow [m]$ such that, for any two $y, z \in S$, $\pi_{x,S}(y) < \pi_{x,S}(z)$ if and only if $y \succ_x z$.

The ordinal preference profile provides a very rough sketch of the underlying distance function d . However, the relative distances expressed by the profile can enable an algorithm to allocate its budget in a very economic way. Consider the following notion: For a set of points $S \subseteq X$ and a point $x \in X$, we define the *distance of x to S* to be $d(x, S) = \min_{y \in S} d(x, y)$. Given the ordinal information, there is a point $z \in \arg \min_{y \in S} d(x, y)$ which can be readily identified as x 's highest ranked point among S . Hence, an algorithm can determine the distance of x to S with a single query $d(x, z)$. Clearly, the same observation can be made about finding a point $z \in \arg \max_{y \in S} d(x, y)$ and the distance $d(x, z)$.

We intend to study the loss in outcome optimality if we restrict an algorithm $\mathcal{A}$ to the ordinal information and a fixed query budget. We consider a variety of clustering problems where the goal is to find a solution (i.e., a set $C \subseteq X$ of so-called centers) that minimizes a given cost function ϕ . We denote by $\mathcal{M}$ the set of all metric spaces. For a metric space $(X, d) \in \mathcal{M}$ and a profile $P \in \mathcal{P}(d)$, let $\mathcal{A}(P, d)$ be the solution computed by algorithm $\mathcal{A}$, and let $C^*(d)$ be a solution of minimal cost. We say that an algorithm $\mathcal{A}$ has *distortion D* with constant (respectively high) probability, if

$$\sup_{\substack{(X, d) \in \mathcal{M} \\ P \in \mathcal{P}(d)}} \frac{\phi(\mathcal{A}(P, d))}{\phi(C^*(d))} \leq D$$

with probability at least $2/3$ (respectively probability at least $1 - 1/n$). The *expected distortion* of $\mathcal{A}$ is given by the ratio

$$\sup_{\substack{(X, d) \in \mathcal{M} \\ P \in \mathcal{P}(d)}} \frac{\mathbb{E}[\phi(\mathcal{A}(P, d))]}{\phi(C^*(d))}.$$

We now state the definition of the (k, z) -clustering problem in the ordinal setting and introduce a few standard terms that are commonly used in the context of clustering problems.

Definition 34. In the ordinal (k, z) -clustering problem, we are given positive integers k, z and a set X of n points that form a metric space (X, d) under the distance function d . Each point $x \in X$ reports a ranking $\succ_x$ that is consistent with the distances d . Let $P = \{\succ_x\}_{x \in X}$. For a subset $S \subseteq X$ of the points, we denote the cost of a given solution $C \subseteq X$ by

$$\phi_z(C, S, d) = \sqrt[z]{\sum_{x \in S} d(x, C)^z}.$$

The goal is to find a set C of k points such that the cost function $\phi_z(C, S, d)$ is minimized.

Given a solution C to an ordinal (k, z) -clustering instance, we typically call the elements of C *centers*. C naturally induces a partition of X into k clusters $\{A_c\}_{c \in C}$ where, for each $c \in C$, $A_c = \{x \in X : \pi_{x,C}(c) = 1\}$. We refer to the collection of these

clusters as a *clustering* of X . Notice that, given a solution C , the clustering induced by C is *unique* since its constituting clusters are defined based on the points' rankings which are in turn strict orderings. A typical operation that we require in our work is to identify a point in a cluster A_c with maximum distance from its center c . This point is again the *unique* point that is ranked lowest by c among A_c . Together with the previous observation that A_c is unique given the ordinal information, we can thus assume that the operation $\arg \max_{x \in A_c} d(c, x)$ returns a *single* point (instead of a set of points).

6.3 Algorithms for k -Center

We present three deterministic algorithms for solving the ordinal k -center $((k, \infty)$ -clustering) problem. Our algorithms are based on a greedy procedure by Gonzalez [78], which is known to yield a 2-approximation of the k -center problem. This procedure simply chooses an arbitrary center to begin with and then, in $k - 1$ iterations, selects as a new center the point that is farthest away from the already chosen centers (*farthest-first traversal*). The next lemma gives a slight generalization of the classic result by Gonzalez. That is, we consider an approximate variant of the previous approach which we refer to as a γ -approximate *farthest-first traversal*. The following formulation includes the case that *multiple* new centers are selected in any iteration of the procedure.

Lemma 35. *Fix an $\gamma \in (0, 1]$. Let $x \in X$ be an arbitrary point, and set $C_0 = \{x\}$. Suppose that, for $i \in [k - 1]$, we define $C_i = C_{i-1} \cup S_i$ where $S_i \subseteq X \setminus C_{i-1}$ contains some point z such that $d(z, C_{i-1}) \geq \gamma \cdot \max_{x \in X} d(x, C_{i-1})$. Then, C_{k-1} is a $\frac{2}{\gamma}$ -approximation of the optimal k -center clustering, i.e.,*

$$\max_{x \in X} d(x, C_{k-1}) \leq \frac{2}{\gamma} \cdot \phi_{OPT},$$

where ϕ_{OPT} is the cost of an optimal k -center clustering.

Proof. Let $C = C_{k-1}$ be the set of centers that we picked according to the procedure described by the lemma. Denote by C^* the optimal k -center solution, and consider the corresponding clustering $\{A_c\}_{c \in C^*}$. If $C \cap A_c$ is non-empty for all $c \in C^*$, then C is a 2-approximation of ϕ_{OPT} by the triangle inequality.

Otherwise, let $i \in [k - 1]$ be the first iteration such that some point $z \in S_i$ satisfying $d(z, C_{i-1}) \geq \gamma \cdot \max_{x \in X} d(x, C_{i-1})$ belongs to the same cluster A_c of the optimal clustering as some point $y \in C_{i-1}$. By definition of the procedure, this must occur in some iteration $i \in [k - 1]$. Then, for any $u \in X$, we have that

$$d(u, C_{i-1}) \leq \frac{1}{\gamma} d(z, C_{i-1}) \leq \frac{1}{\gamma} d(z, y) \leq \frac{1}{\gamma} (d(y, c) + d(z, c)) \leq \frac{2}{\gamma} \max_{x \in A_c} d(x, c) \leq \frac{2}{\gamma} \phi_{OPT}.$$

Here, the first inequality is due to our assumption that z satisfies $d(z, C_{i-1}) \geq \gamma \cdot \max_{x \in X} d(x, C_{i-1})$. For the second inequality, we use that $y \in C_{i-1}$. The remaining

inequalities follow from the triangle inequality, the assumption that $y, z \in A_c$, and the optimality of C^* under the k -center objective.

Hence, C_{i-1} is already a $\frac{2}{\gamma}$ -approximation of ϕ_{OPT} , and adding more centers to the solution can never increase its cost. This concludes the proof of the lemma. $\square$

Warm-up: 2-Distortion Algorithms

Algorithm 1: Ordinal k -center with $O(k^2)$ queries

Input: X, d, P, k

```

1  $x \leftarrow$  arbitrary point from  $X$ 
2  $C \leftarrow \{x\}$ 
3 for  $1 \dots (k-1)$  do
4   Set  $\delta_{\max} = 0$ 
5   for  $c \in C$  do
6     //define cluster with center  $c$ 
6     Define  $A_c = \{x \in X : \pi_{x,C}(c) = 1\}$ 
6     //query distance from  $c$  to farthest point among  $A_c$ 
7     Let  $p = \arg \max_{x \in A_c} d(c, x)$ 
8     Query  $\delta = d(c, p)$ 
9     if  $\delta \geq \delta_{\max}$  then
10       $\delta_{\max} = \delta$ 
11       $z \leftarrow p$ 
12    $C \leftarrow C \cup \{z\}$ 
13 return  $C$ 
```

The farthest-first traversal method lends itself well to be adapted to the ordinal setting (Algorithm 1). Given a set of centers, the farthest point from these centers can be determined by inspecting the clustering induced by this set of centers and making one distance query per cluster.

Theorem 36. *Algorithm 1 makes $\frac{k^2-k}{2}$ distance queries and has distortion 2.*

Proof. In every iteration $i \in [k-1]$ of the outer loop, Algorithm 1 adds exactly one new center to the current solution C . Observe that, at the start of the i -th iteration of the outer loop (line 4), we have that $|C| = i$. Then, in the inner loop, the algorithm performs exactly $|C| = i$ distance queries. Hence, the algorithm makes $\sum_{i=1}^{k-1} i = \frac{k^2-k}{2}$ queries in total.

With respect to the distortion of the algorithm, let C_{i-1} be the current set of centers in some iteration $i \in [k-1]$ at the beginning of the outer loop (line 4). Assume that, at the end of the outer loop (line 12), the set of centers is $C_i = C_{i-1} \cup \{z\}$ where the point z belongs to the cluster of center $c' \in C_{i-1}$. By definition of the inner loop, we

have that

$$d(z, C_{i-1}) = \max_{x \in A_{c'}} d(c', x) = \max_{c \in C_{i-1}} \max_{x \in A_c} d(c, x) = \max_{x \in X} d(x, C_{i-1}).$$

Hence, Algorithm 1 performs a 1-approximate farthest first traversal, and the desired distortion bound follows from Lemma 35. $\square$

Algorithm 2: Ordinal k -center without queries

Input: X, d, P, k

```

1  $x \leftarrow$  arbitrary point from  $X$ 
2  $C \leftarrow \{x\}$ 
3 for  $1 \dots (k-1)$  do
4    $S \leftarrow \emptyset$  //new set of centers
5   for  $c \in C$  do
6     //define cluster with center  $c$ 
6     Define  $A_c = \{x \in X : \pi_{x,C}(c) = 1\}$ 
6     //add farthest point from  $c$  among  $A_c$  to solution
7      $z \leftarrow \arg \max_{x \in A_c} d(c, x)$ 
8      $S \leftarrow S \cup \{z\}$ 
9    $C \leftarrow C \cup S$ 
10 return  $C$ 
```

For the zero-query regime, we modify the farthest-first traversal method such that, in every iteration, the farthest point in *every* cluster is chosen (Algorithm 2).

Theorem 37. *Algorithm 2 returns a set of centers C of size $|C| = 2^{k-1}$ such that $\max_{x \in X} d(x, C) \leq 2\phi_{OPT}$, where ϕ_{OPT} is the cost of an optimal k -center clustering.*

Proof. Let C be the solution returned by Algorithm 2. We first show that C has the right size. Note that after initially having size 1, in each iteration i of the outer loop, the algorithm adds 2^{i-1} points to C . Thus,

$$|C| = 1 + \sum_{i=0}^{k-2} 2^i = 1 + 2^{k-1} - 1 = 2^{k-1}.$$

We proceed to prove that Algorithm 2 is 2-approximate with respect to ϕ_{OPT} . In any iteration $i \in [k-1]$ of the algorithm, let C_{i-1} be the set of already selected centers at the beginning of the outer loop (line 4). Assume that the solution at the end of the outer loop (line 9) in this iteration is $C_i = C_{i-1} \cup S_i$. We now argue that there must be a point $z \in S_i$ such that $d(z, C_{i-1}) = \max_{x \in X} d(x, C_{i-1})$ which implies that the algorithm performs a 1-approximate farthest-first traversal (Lemma 35).

Consider any center $c' \in \arg \max_{c \in C_{i-1}} \max_{x \in A_c} d(c, x)$, and assume that Algorithm 2 included z' in S_i when considering the cluster $A_{c'}$ on line 7 of the inner loop.

Hence,

$$d(z', C_{i-1}) = \max_{x \in A_{c'}} d(c', x) = \max_{c \in C_{i-1}} \max_{x \in A_c} d(c, x) = \max_{x \in X} d(x, C_{i-1})$$

as desired.

We have thus shown that Algorithm 2 performs a 1-approximate farthest-first traversal. By Lemma 35, the cost of the solution returned by the algorithm is therefore at most $2\phi_{\text{OPT}}$ which concludes the proof of the theorem. $\square$

4-Distortion Algorithm with $O(k)$ Queries

In this section we present an algorithm that achieves constant distortion by performing a $\frac{1}{2}$ -approximate farthest-first traversal. Surprisingly, using ordinal information, we can execute such an approximate traversal with an asymptotically optimal query bound. That is, Algorithm 3 makes at most $2k - 1$ queries, and, as we will later demonstrate (see Theorem 42), any constant distortion algorithm for the k -center problem requires $\Omega(k)$ queries.

Theorem 38. *There exists a deterministic 4-distortion algorithm for the k -center problem that makes at most $2k - 1$ queries.*

We first give a high level intuition for the way that we defined the algorithm. Algorithm 3 builds the final solution iteratively over $k - 1$ rounds (starting with a solution that consists of an arbitrary point). Conceptionally, the algorithm keeps track of (center, farthest point)-pairs for all clusters in its current solution. However, not all distances between the points of the respective pairs are queried. Instead, these distances will only be known for a subset of the pairs, and the latter pairs are chosen in a way that will allow us to bound the number of new pairs created during the algorithm's execution. At the same time, for every unqueried pair, the set of queried pairs will contain at least one distance that is at least half the respective distance for the unqueried pair.

Throughout the algorithm's execution, we denote its current solution by C . Furthermore, the algorithm keeps a so-called *query set* $Q \subseteq C$ such that, for every point $p \in Q$, the distance to the farthest point in its cluster is known. Both C and Q will change over time, and we use the notion of points *entering* and *leaving* the query set.

We begin our formal discussion of Algorithm 3 by showing that the queries that we make on lines 4 and 26 indeed suffice for its execution. That is, by the second statement of the following invariant, it holds that in every iteration of the algorithm, each distance required on line 10 is known.

Invariant 39. *At the start of each iteration, the following two statements are true.*

1. *For any two points $p, u \in Q$, let q, v be the farthest points in their respective clusters, that is, $q = \arg \max_{x \in A_p} d(x, p)$, and $v = \arg \max_{x \in A_u} d(x, u)$. It holds that $d(p, q) \leq d(v, q)$.*

Algorithm 3: Ordinal k -center with $2k - 1$ queries

Input: X, d, P, k

```

1  $y \leftarrow$  arbitrary point from  $X$ 
2  $z \leftarrow \arg \max_{x \in X} d(x, y)$ 
3  $C, Q \leftarrow \{y\}$ 
4 Query  $d(y, z)$  //query to simplify the analysis
5  $A \leftarrow$  clustering induced by  $C$ 
6 for  $1 \dots (k - 1)$  do
7    $\delta_{\max} = 0$ 
8   for  $p \in Q$  do
9      $q \leftarrow \arg \max_{x \in A_p} d(x, p)$ 
10    //distance  $d(p, q)$  is known by Invariant 39
11    if  $d(p, q) \geq \delta_{\max}$  then
12       $\delta_{\max} = d(p, q)$ 
13       $y \leftarrow p$ 
14    $z \leftarrow \arg \max_{x \in A_y} d(x, y)$ 
15    $C \leftarrow C \cup \{z\}$ 
16    $Q \leftarrow Q \setminus \{y\}$ 
17    $A \leftarrow$  clustering induced by  $C$  //update clustering according to
    new solution
    //decide whether any center (including  $y$  and  $z$ ) that is
    not currently contained in the query set should enter
    the query set
18   for  $u \in C \setminus Q$  do
19      $add \leftarrow \text{true}$ 
20      $v \leftarrow \arg \max_{x \in A_u} d(x, u)$ 
21     for  $p \in Q$  do
22        $q \leftarrow \arg \max_{x \in A_p} d(x, p)$ 
23       if  $v \succ_q p$  then
24          $add \leftarrow \text{false}$ 
25     if  $add$  then
26        $Q \leftarrow Q \cup \{u\}$ 
27       Query  $d(u, v)$ 
28 return  $C$ 

```

2. For any point $p \in Q$ and the farthest point $q = \arg \max_{x \in A_p} d(x, p)$ in its cluster, the exact distance $d(p, q)$ is known.

Proof. We prove both statements by induction over the iterations $i = 1, 2, \dots, (k-2)$ of the algorithm. At the start of iteration $i = 1$, both statements hold trivially, since Q contains only a single point y , and the algorithm queried the distance from y to the farthest point in its cluster on line 4.

Assume now that both statements are true at the beginning of iteration i . We show that both statements then also hold at the beginning of iteration $i+1$. Let z_i be the point that was included in the solution C as the farthest point in the cluster of center y_i in iteration i . For the points that remain in Q after the removal of y_i on line 15, the first statement holds by the induction hypothesis. Now, consider any point $u \in C \setminus Q$ and the farthest point v in its cluster at some iteration of the loop on line 17. The point u enters the query set only if Q does *not* contain a point p such that the farthest point q in its cluster prefers v over p . Hence, if u entered the query set, then it holds that $d(p, q) \leq d(v, q)$ for any such pair p, q where $p \in Q$. This shows that the first statement is true at the beginning of iteration $i+1$.

With respect to the second statement, we label the query sets at the beginning of iteration i and $i+1$ as Q_i and Q_{i+1} , respectively. Let Q'_i be the query set in iteration i after the removal of y_i on line 15. By the induction hypothesis, the first statement holds for $Q_i = Q'_i \cup \{y_i\}$. That is, for any $p \in Q'_i$ and the farthest point q in its cluster, we have that $d(p, q) \leq d(z_i, q)$. Hence, q is not assigned to the cluster of the new center z_i , and is therefore still the farthest point in the cluster of center p . By the induction hypothesis (second statement), the distance $d(p, q)$ is known. Finally, for any point $u \in C \setminus Q'_i$ that entered the query set in iteration i and the farthest point v in its cluster, the algorithm queries the distance $d(u, v)$. Thus, for any point in $Q_{i+1} \setminus Q'_i$ the distance required by the second statement is also known, which concludes the proof of the invariant. $\square$

We proceed to bound the number of queries that the algorithm makes.

Lemma 40. *The total number of queries that Algorithm 3 makes is at most $2k - 1$.*

Proof. For $i \in [k-1]$, let (y_i, z_i) be the pair of points such that, in iteration i , the algorithm included z_i in C as the farthest point in the cluster of center y_i . We will show that for each of these pairs, the algorithm makes at most two queries. Together with the single query that the algorithm performs on line 4, this yields the desired bound.

First, for any iteration i , notice that after the pair (y_i, z_i) was determined by the algorithm, neither of these points is included in the query set (in particular, y_i was removed from Q on line 15). In the following, we focus on the point y_i since the same line of reasoning will hold for z_i . The algorithm performs a single query for the distance from center y_i to the farthest point in its cluster, if—in some iteration— y_i enters the query set. The only way for y_i to leave the query set again is that the algorithm chooses the farthest point in its cluster as a center. Assume that this happens in iteration $i' > i$. At this point, y_i (respectively, z_i) becomes $y_{i'}$, and we can repeat the previous argument for $y_{i'}$. $\square$

Finally, we show that Algorithm 3 achieves the desired distortion which essentially follows from the observation that the algorithm performs a $\frac{1}{2}$ -approximate farthest-first traversal. The combination of Lemmas 40 and 41 then yields Theorem 38.

Lemma 41. *Algorithm 3 has distortion 4.*

Proof. In each iteration, the algorithm considers only the clusters which have as a center a point from $Q \subseteq C$. Among the points contained in these clusters, the algorithm then includes the point z with maximum distance to its most preferred center $y \in Q$ in its solution. We will show that, for any remaining center $u \in C \setminus Q$ and the farthest point v in its cluster, $d(u, v) \leq 2d(y, z)$. This implies that the sequence of centers picked by the algorithm constitutes a $\frac{1}{2}$ -approximate farthest-first traversal (see Lemma 35).

Consider any point $u \in C \setminus Q$ at the beginning of some iteration of the algorithm (line 7). Let v be the farthest point in u 's cluster. Since u did not enter the query set in the previous iteration, Q must contain a point p such that the farthest point q in its cluster prefers v over p . Hence,

$$d(v, q) \leq d(p, q). \quad (6.1)$$

We now have that

$$d(u, v) \leq d(p, v) \leq d(p, q) + d(v, q) \leq 2d(p, q) \leq 2d(y, z).$$

Here, the first inequality is due to the fact that u is v 's most preferred center among C . For the second and third inequality, we use the triangle inequality and (6.1), respectively. The last inequality follows from the fact that the pair (y, z) maximizes the distance between center and farthest point among all centers in the current query set Q which includes p .

We have thus shown that Algorithm 3 performs a $\frac{1}{2}$ -approximate farthest-first traversal. The lemma now follows from Lemma 35. $\square$

6.4 Lower Bounds for (k, z) -Clustering

We begin the discussion of our lower bounds by presenting a construction for the k -center objective. Our lower bounds for k -center are simple and optimal. Moreover, the same bounds carry over to any (k, z) -clustering objective. We then turn our attention to the k -median objective in particular. Our lower bounds for k -median are significantly more complicated, but use a similar construction as the k -center lower bound.

Theorem 42. *For any fixed α , every bicriteria algorithm $\mathcal{A}$ for k -center that has distortion at most α with at least constant probability must return a solution of size at least $\Omega(2^k)$. Moreover, any algorithm that has distortion at most α with at least constant probability must make at least $\Omega(k)$ queries.*

Proof. The hard instance is the same for the low query and zero query setting. We start with an analysis for the latter.

The hard instance: Our hard instance consists of 2^{k-1} points. We begin by describing the ordinal information and the underlying metric. Consider a complete binary tree T of depth $k - 1$. For any two nodes p, q , we say that a is the common ancestor of p and q if a is the minimum depth node in the shortest path between p and q in T .

The interpretation of this tree is that the leaves are the points and for any interior node a , the value $d(a)$ stored in a denotes the distances between all points p, q that have a as the common ancestor. Thus, we now require the following invariant to ensure that the tree encodes a metric.

Invariant 43. *If the subtree rooted at a contains the interior node b , then $d(a) \geq d(b)$.*

We now specify the ordinal preferences, which we fix *before* determining the values $d(a)$ of the interior nodes. Let p, q, o be three leaves and let $a(p, q)$, $a(p, o)$, and $a(q, o)$ be common ancestors of these pairs of nodes, respectively.

- If the depth of $a(p, q)$ is larger than the depth of $a(p, o)$ and $a(q, o)$ then the preference list of p determines q to be closer to p than to o .
- If the depth of $a(p, q)$ and $a(p, o)$ is equal then the relative ordering of q and o in the preference list of p is arbitrary (w.l.o.g., it may be chosen lexicographically).

We now describe a hard input distribution that satisfies the invariant and is consistent with the ordinal preferences. Select a random path Q between the root of T and an arbitrary node r at depth $k - 1$. All nodes a along that path receive the value $d(a) = D$. All remaining nodes receive the value $d(a) = 1$.

Analysis: Note that, for any two trees sampled from the distribution, the values assigned to the interior nodes satisfy Invariant 43 and thereby induce a metric on the set of leaf nodes. Since the ordinal preferences are independent from these values, the two trees cannot be distinguished using the ordinal information.

We now determine an optimal k -center solution C . For every interior node a in Q , the children of a form subtrees $T(a, \text{small})$ and $T(a, \text{large})$. The root b of $T(a, \text{small})$ satisfies $d(b) = 1$ and the root c of $T(a, \text{large})$ satisfies $d(c) = D$. For the largest depth interior node a in Q , we introduce the convention that $T(a, \text{large})$ contains the leaf r (i.e., the end point of Q). C now places exactly one center on an arbitrary leaf of $T(a, \text{small})$ and one center on r . The cost of C is therefore 1. Now consider any other solution C' . If C' does not place a center on r , then the cost of C' is D . Otherwise, there must exist some $a \in Q$ for which $T(a, \text{small})$ does not receive a center. Hence, the points in $T(a, \text{small})$ must be served by some center contained in $T(a, \text{large})$ or by a point not contained in the subtree rooted at a . In both cases, the cost of these points is D .

To conclude, it now suffices to analyze the performance of the best deterministic algorithm placing K centers against this hard input distribution. Since the algorithm does not make any queries and cannot determine Q based on the ordinal information, its choice of centers is fixed. There are 2^{k-1} many different nodes at depth $k - 1$. Hence, the probability that K includes the leaf node r is $K/2^{k-1}$. Conversely, if

$K \notin \Omega(2^k)$ then the probability that K does not include r is at least constant, which leads to a distortion of D .

Finally, we remark on some generalizations of this lower bound. For the low query regime, an algorithm needs to find the entire path Q or, equivalently, identify the leaf r . If it decides to not do so then with probability at least $\frac{1}{2}$ it will have unbounded distortion of D . Again, consider the performance of the best deterministic algorithm against the input distribution. Given that T is a binary tree, at least one query queries is required to reduce the search space for Q (equivalently, for r) by a factor of $\frac{1}{2}$ in expectation. Hence, if the algorithm does not make $\Omega(k)$ queries, its distortion is unbounded. $\square$

Notice that, according to Theorem 42, the distortion bound α has no influence on the number of queries or the number of centers. That is, our lower bounds hold for arbitrary values of α . This property together with the observation that the cost of all (k, z) -clustering objectives are within a $\text{poly}(n)$ factor implies that the same bounds indeed hold for any (k, z) -clustering objective.

Next, we give a different lower bound for any bicriteria algorithm for k -median. Specifically, we show that any bicriteria algorithm for k -median requires $\Omega(2^k \log n)$ centers. For 2-median, this becomes $\Omega(\log n)$, which stands in contrast with 2-center, where we can obtain a true 2-distortion using only ordinal information (see Theorem 37).

Theorem 44. *For any fixed α , every bicriteria algorithm $\mathcal{A}$ for k -median that has distortion less than α with at least constant probability must return a solution of size at least $\Omega\left(\frac{\log n}{\log \alpha} \cdot 2^k\right)$. Moreover, any algorithm achieving a constant factor approximation for k -median must make at least $\Omega(k + \log \log n)$ queries.*

Proof. As with the proof for k -center above, we first describe the hard instance for the zero-query regime and then remark on how to extend it. To simplify the calculations, we prove the lower bound for an input of size $\Theta(n)$, where we make the following two assumptions:

- There is an integer n' such that $n = 2^{k-2} \cdot n'$.
- n' and $\alpha + 1$ are powers of 2.

The claim for general n and α carries over with very minor details.

The hard instance: The first part of the instance is almost identical to that of k -center in the proof of Theorem 42. Indeed, since the distortion of k -center is unbounded, it is also unbounded for k -median as both costs are within a factor n of each other. Recall that our hard instance for k -center used a complete binary tree T . In our hard instance for k -median, we augment this tree by adding a hard instance for 2-median below each of its leaves.

We proceed to describe these 2-median instances. For a leaf node u in T , we refer to the 2-median instance below u as I_u . For every leaf u , I_u consists of $n' = n/(2^{k-2})$

points. We group the points in I_u into bundles B_i for $i \in \{0, 1, \dots, \frac{\log n'}{\log(\alpha+1)}\}$. The i -th bundle has the property that $|B_i| = (\alpha + 1)^i$.

We now introduce the ordinal preferences among the points in I_u , as well as between the points of different 2-median instances. Then, we describe the distribution over metrics consistent with said preferences.

- Consider only points from a 2-median instance I_u . For any two points $p, q \in B_i$ and any point $o \notin B_i$ we have $d(p, q) \leq d(p, o), d(q, o)$. The remaining ordinal preferences among the points in I_u may be chosen arbitrarily.
- Let $p \in I_u$, and let q, o be two points such that either $q \in I_v, v \neq u$ or $o \in I_v, v \neq u$. Then, whether p prefers q over o or not, depends on the depths of the common ancestors $a(p, q), a(p, o)$ in T . We refer the reader to the description of the ordinal preferences in our hard k -center instance (see the proof of Theorem 42).

We now specify the hard input distribution over metrics that is consistent with these preferences. We initialize T as a binary tree of depth $k - 2$ and pick a leaf node r uniformly at random. Let Q be the path in T from its root to r . We now assign values to each node in T including its leaves. These values are $d(a) = D$ for every node a that lies on the path Q (where D is some sufficiently large number) and $d(a) = \varepsilon$ otherwise (where $\varepsilon > 0$ is arbitrarily small).

The distance between any two points $p, q \in I_u, u \neq r$ is ε . For the 2-median instance I_r , we select an $\ell \in \{0, 1, \dots, \frac{\log n'}{\log(\alpha+1)} - 1\}$ uniformly at random. The distances now satisfy the following properties:

- For every pair of points p, q from a bundle B_j with $j \geq \ell$, we set $d(p, q) = \varepsilon$.
- For every pair of points p, q from a bundle B_j with $j < \ell$, we set $d(p, q) = 1$.
- For every $p \in B_j$ and $q \in B_k, j \neq k$, we set $d(p, q) = 1$ if $k \leq \ell$. If $k, j > \ell$, we set $d(p, q) = \varepsilon$.

Analysis: Since the ordinal preferences were determined before sampling ℓ , no algorithm using only ordinal information can determine any information about ℓ . As before, the distance between any two points $p \in I_u, q \in I_v, u \neq v$ is given by the value that is stored at their common ancestor node $a(p, q)$ in T (see the proof of Theorem 42).

We now consider the cost of an optimal solution C . As in our hard instance for k -center (Theorem 42), the optimal solution must place at least one center in every subtree $T(a, \text{small})$ where a is a node on the path Q . Otherwise, the solution has cost at least D which can be arbitrarily high. Consider any subtree $T(a, \text{small})$ and note that its root b satisfies $d(b) = \varepsilon$. Hence, if the solution places a center on any leaf in $T(a, \text{small})$, then the contribution of the points in $T(a, \text{small})$ to the cost of the solution is negligible. Hence, the cost of any optimal solution C depends only on the cost incurred for the points in I_r .

We claim that C places a single center c_ℓ in B_ℓ and a single center in some bundle B_j with $j > \ell$. The cost of the points in bundles B_j with $j > \ell$ is now ε . The cost of a point p served by c_ℓ is ε , if $p \in B_\ell$ and 1 if $p \in B_k$, $k < \ell$. Thus the overall cost is $\sum_{i=0}^{\ell-1} \alpha^i = \frac{(\alpha+1)^\ell - 1}{\alpha}$, ignoring negligible contributions from the ε -valued distances.

Any solution that does not intersect with a bundle B_j , $j > \ell$ costs at least $(\alpha + 1)^{\frac{\log n'}{\log(\alpha+1)}} = n'$. Finally, any solution that does not intersect with B_ℓ costs at least $(\alpha + 1)^\ell$. Both of those terms are larger than $\frac{(\alpha+1)^\ell - 1}{\alpha}$ by at least a factor α if n' is large enough so we can conclude that C is optimal.

Again, it suffices to consider the performance of a deterministic algorithm placing K centers against the hard input distribution. Since the ordinal information offers no information on either Q or ℓ , the choice of centers C' is fixed. As in the proof of Theorem 42, the probability that C' includes any point from I_r is $K/2^{k-2}$ such that if $K \notin \Omega(2^k)$ the distortion is D with at least constant probability. Furthermore, the probability that C' intersects with B_ℓ is at most $\frac{\log(\alpha+1)}{\log n'}$. Thus, C' must consist of

$$\Omega\left(\frac{\log n'}{\log(\alpha+1)} \cdot 2^k\right) = \Omega\left(\frac{\log n - k}{\log(\alpha+1)} \cdot 2^k\right)$$

centers to improve over an α distortion. The first part of the theorem now follows by choosing n' such that n is large enough compared to k .

For the low-query regime, the argument that we require at least $\Omega(k)$ queries is equivalent to that of the k -center instance, being that the instances for the first $k - 2$ levels of the tree are identical. The $\Omega(\log \log n)$ query lower bounds follows from the fact that there are $\log n'$ many choices for the bundle B_ℓ and every query can rule out half of the remaining possible choices. $\square$

6.5 Facility Location with Uniform Opening Costs

In this section, we present both lower and upper bounds for the facility location problem in the ordinal information setting. For our upper bound, we revisit the seminal algorithm for online facility location by Meyerson [101]. For worst case input orders, it is known to achieve an optimal $O(\log n / \log \log n)$ approximation [70]. For random order inputs, it is known to achieve a 4-approximation [86]. In effect, our next algorithm simulates the latter random order arrival of input points. In every iteration, the algorithm needs to determine the exact distance of a point x to the set C of already opened facilities, see line 5 in the description of Algorithm 4. This operation requires a single query given the ordinal information. Furthermore, the operation is performed for every point exactly once. The following theorem summarizes this discussion.

Theorem 45 (See also Pulyassary [114]). *Algorithm 4 achieves constant expected distortion for the ordinal facility location problem with uniform opening costs using one query per point.*

Algorithm 4: Meyerson's algorithm for facility location

Input: X, d, P, f

- 1 $R \leftarrow$ random permutation of X
- 2 $C \leftarrow \{R(1)\}$
- 3 **for** $t = 2 \dots n$ **do**
- 4 $x \leftarrow R(t)$
- 5 $p \leftarrow \min \left\{ 1, \frac{d(x, C)}{f} \right\}$
- 6 Set $C \leftarrow C \cup \{x\}$ with probability p
- 7 **return** C

Lower Bound. We now lower-bound the number of queries necessary to achieve any given distortion for the facility location problem. Using no queries, it is not possible to obtain bounds on the distortion [9] beyond the trivial $O(n)$ bound. Our lower bound essentially shows that $\Omega(n)$ queries are necessary to achieve constant distortion, making the adaptation of Meyerson's algorithm optimal.

Theorem 46. *For any fixed α , every algorithm $\mathcal{A}$ for facility location that has distortion less than α with probability at least $1/2$ must make $\Omega\left(\frac{n}{\alpha}\right)$ distance queries.*

Proof. We assume that the opening costs per facility are 1. As before, we first describe the ordinal preferences and then sample a metric from a distribution consistent with these preferences.

The hard instance: For a sufficiently large choice of n , we group the points into $\frac{n}{s}$ clusters A_i , each consisting of $s \in \Omega(\alpha)$ points. Any two points from a cluster A_i prefer each other over any point from some other cluster A_j . The preferences within the clusters, as well as across clusters are arbitrary as long as every point $p \in A_i$ prefers any point $q \in A_i$ over any point $o \in A_j$, $j \neq i$. Moreover, the distances between any two points $p \in A_i$ and $q \in A_j$, $i \neq j$ are set to be ∞ (or a sufficiently large number if finite values are required).

The hard input distribution is now defined as follows. We select a cluster A_i uniformly at random, and toss a fair coin. With probability $\frac{1}{2}$, all of the points in A_i have pairwise distance ε . With probability $\frac{1}{2}$, the pairwise distances in A_i are chosen to be a sufficiently larger number $N \gg n/s + s - 1$. For each cluster other than A_i , the distances between any two points within this cluster are ε .

Analysis: In the case that the pairwise distances in A_i are ε , the optimal solution consists of placing exactly one facility in every cluster, leading to a cost of n/s (if we ignore the arbitrarily small connection costs). In the case that the pairwise distances are N , the optimal solution consists of placing exactly one facility in every cluster A_j , $j \neq i$ and placing a facility on every point in A_i , leading to an overall cost of $n/s + s - 1$.

To distinguish between these two cases, the algorithm has to query at least one distance between two points in A_i . Suppose the algorithm makes Q queries. If the algorithm fails to determine whether A_i has pairwise distances N or not, it must place a center on every point of a cluster it has not queried, as otherwise the distortion is N and therefore unbounded.

By Yao's minimax principle, we may assume that the centers queried by the algorithm are fixed until A_i is detected. The probability that A_i is detected is $\frac{Q \cdot s}{n}$. Thus, if A_i is not detected, the algorithm must place $s - 1$ additional facilities on each unqueried cluster, that is, $(n/s - Q) \cdot (s - 1)$ additional facilities in total. Hence, in the case that the distances between the points in A_i are ε (which occurs with probability $\frac{1}{2}$), the algorithm incurs a distortion of $\frac{n/s + (n/s - Q) \cdot (s - 1)}{n/s} \in \Omega(\alpha)$ for $Q \in o(n/\alpha)$. Here, we used that we chose the cluster size s such that $s \in \Omega(\alpha)$. The claim now follows by scaling s so that the distortion becomes at least α . $\square$

6.6 Conclusion and Open Problems

We gave optimal algorithm for computing bicriteria approximations for k -center both in terms of the number of distance queries as well as number of additional centers in the purely ordinal setting. In our remaining preliminary work [28], we also give substantially improved low-query and purely ordinal bicriteria algorithms for k -median. Aside from closing the small remaining gaps left in our analysis, several interesting open problems present themselves. A popular way to interpolate between k -median and k -center is ordered clustering [30, 40]. Is it possible to achieve low distortion algorithms for this problem as well? Furthermore, there exist many other clustering objectives, such as graph clustering. Which distortion/query tradeoffs are possible for sparsest cut and metric max cut?

Chapter 7

The Complexity of Learning Approval-Based Multiwinner Voting Rules

We study the PAC learnability of multiwinner voting, focusing on the class of approval-based committee scoring (ABCS) rules. These are voting rules applied on profiles with approval ballots, where each voter approves some of the candidates. According to ABCS rules, each committee of k candidates collects from each voter a score, which depends on the size of the voter's ballot and on the size of its intersection with the committee. Then, committees of maximum score are the winning ones. Our goal is to learn a target rule (i.e., to learn the corresponding scoring function) using information about the winning committees of a small number of sampled profiles. Despite the existence of exponentially many outcomes compared to single-winner elections, we show that the sample complexity is still low: a polynomial number of samples carries enough information for learning the target rule with high confidence and accuracy. Unfortunately, even simple tasks that need to be solved for learning from these samples are intractable. We prove that deciding whether there exists some ABCS rule that makes a given committee winning in a given profile is a computationally hard problem. Our results extend to the class of sequential Thiele rules, which have received attention recently due to their simplicity.

7.1 Introduction

Voting has been used for centuries to aggregate individual preferences into a common decision. In addition to its traditional use for electing governments or for decision making in management boards, it has also been proved useful in novel applications where individual ratings need to be summarized as collective knowledge. But, is there a general recipe on how preferences should be aggregated? Fortunately, there is no “golden” voting rule and this has led social choice theory—and, in particular, its modern computational branch [27]—onto exciting research endeavours.

A popular approach has aimed, quite successfully, to evaluate voting rules in terms of desirable axioms they must satisfy. Well-known impossibilities, e.g., see Arrow [15], showcase the limitations of this approach. Deviating from this *axiomatic* treatment, recent works view voting rules as optimized decision making methods, perhaps tailored to particular applications. In this context, the *data-driven design* of voting rules is a very natural approach. The goal is to derive a voting rule from a set of known preferences which are accompanied by favored outcomes under these preferences. The hope then is for the resulting rule to be equally well-suited to more general preferences where the favored outcomes are unknown. The current paper aims to study the potentials and limitations of this approach.

We focus on *multiwinner* voting rules [64], which on input the preferences of n voters over m available candidates, return as outcome one or more committees of candidates of fixed size k . In particular, we study *approval-based* voting [93, 94], where the preference of a voter is simply the set of candidates she approves. And, more concretely, we consider the class of *approval-based committee scoring (ABCS)* rules, defined by Lackner and Skowron [92]. An ABCS rule follows a common format. It employs a scoring function, according to which each voter awards a score to each committee of k candidates. This score depends on the ballot size (the number of candidates the voter approves) and the size of its intersection with the committee. Different scoring functions can be used to define different voting rules.

A natural application area of approval-based voting are rating tasks. For example, consider a website specialized on cultural events that aims to present every week the top-20 performances in the theaters of a city. A simple way to compute this is to ask website visitors for their opinions and aggregate them to a top-20 list. Approval-based multiwinner voting can be used here as follows.

- Each participant is asked for her favorite performances (i.e., for her approval vote) among the ones available (i.e., among the alternatives).
- Then, an ABCS rule can be used to compute the top-20 (i.e., the winning committee).

Deciding on the best ABCS rule depends on the application at hand. For example, under a rule that favours *individual excellence*, each voter assigns to each committee a score that is equal to the number of candidates in the committee the voter approves. Another rule could give just one point to each committee that has a non-empty intersection with the voter's ballot; such a rule would promote *representation of voters*. In practice, situations with such a one-dimensional objective for a voting rule are extremely rare. This issue has been extensively studied in the literature, e.g., see the work of Faliszewski and Talmon [62], Faliszewski et al. [63], Jaworski and Skowron [84], Lackner and Skowron [91]. At the same time, hand-picking an ABCS rule that satisfies multiple objectives (and even accounting for potential trade-offs between objectives) might prove difficult. Instead, it may be easier to derive the characteristics of the desired rule from data when suitable data is available or when it is possible to

generate such data. In this paper, we assume the availability of data in the form of preference profiles and corresponding desired winning committees.

Consider the above example of rating cultural events. A good ABCS rule for picking the top-20 theater performances may aim at ensuring that the choice is representative of the website visitors' common interests while also taking into account fringe works that nevertheless appear highly outstanding to a few of the theater-goers. A data-driven approach to decide such a rule could be implemented by the website operators as follows.

- For the first ten weeks of operations, the website collects the input from the visitors but uses a small set of (expensive) experts every week to decide an “ideal” top-20.
- Then, these ten pairs of visitor input and expert top-20 constitute a set of profile-winning committee examples.
- An ABCS rule derived from these data is applied in the subsequent weeks to visitors' input to simulate (in an inexpensive way) the expert opinion.

This indicative scenario involves several thousands of voters, more than 100 alternatives, and winning committees of size 20. Other rating applications, e.g., for hotels, restaurants, or local businesses, or platforms for evaluating business proposals, investment opportunities, and microloan applications could benefit from a similar approach. Arguably, the best rule for the application at hand should at least agree with given profile-winning committee data, and, ideally, produce desirable outcomes for unknown preference profiles. Can such a data-driven selection of an ABCS rule be effective?

We explore this question using the *PAC* (probably approximately correct) *learning* framework. We follow a similar methodological approach with Procaccia et al. [113], who addressed the same question for single-winner voting rules. In the terminology of PAC learning, we would like to determine the *sample complexity* of the class of ABCS rules. How many samples (profiles and corresponding winning committees) are necessary and sufficient so that an ABCS rule that agrees with these data points can be learnt? However, the answer to this question addresses our challenge only partially. Indeed, low sample complexity does not necessarily imply *efficient* learning, as the computational problem of finding an ABCS rule that fits the given data can be hard.

Our contribution and techniques

Our first result states that the class of ABCS rules has only polynomial sample complexity (Section 7.4). Using a variant of the *multiclass fundamental theorem* in PAC learning (Theorem 50), we obtain our sample complexity bounds by proving upper bounds on the *graph dimension* and lower bounds on the *Natarajan dimension* of the class of ABCS rules. For our upper bound, we establish a connection between the graph dimension and the number of different sign patterns of a set of linear functions.

Then, a result in algebraic combinatorics—originally proved by Warren [124] and later refined by Alon [3]—is used to upper-bound this number of sign patterns and, consequently, the graph dimension and the sample complexity of the class of ABCS rules.

On the negative side, we give strong evidence that efficient PAC learnability of ABCS rules is not possible. We show that given a profile of approval votes and a committee, deciding whether there is an ABCS rule that makes this committee winning is a $\text{coW}[1]$ -hard¹ problem, when parameterized by the committee size k (Section 7.5). Our proof uses a quite involved reduction from INDEPENDENTSET , which, on input a graph, defines a profile consisting of several parts and a committee. Some of the parts of the profile guarantee that the only ABCS rule that can make the committee winning has a very particular form: it takes into account only votes with two candidates (ignoring the rest), and mimics the approval-based CC rule (henceforth, simply, the CC rule), a famous rule that is inspired by the work of Chamberlin and Courant [41]. Then, the main part of the profile guarantees that the committee is indeed winning under this rule if and only if the graph does not have a large independent set. Our reduction can be modified to give $\text{coW}[1]$ -hardness for the following *winner verification* problem: given a profile and a committee, is the committee winning under the CC rule? This result strengthens a recent one by Sonar et al. [119].

We also consider sequential Thiele rules (Section 7.6). These can be thought of as greedy approximations of a subclass of ABCS rules which originate from the work of Thiele [120]. However, their definition is considerably different from ABCS rules, so that our sample complexity analysis techniques need revision. Still, we are able to show polynomial sample complexity bounds for learning sequential Thiele rules. Interestingly, the problem of deciding whether there is some sequential Thiele rule that makes a given committee winning in a given profile is now fixed-parameter tractable (parameterized by the committee size). Despite this seemingly positive result, we provide evidence that efficient learning is out of reach for sequential Thiele rules as well, by showing NP-hardness. We do so by a novel reduction from a structured version of 3SAT, which equates the ordering in which several candidates are greedily included in the winning committee with a boolean assignment to the 3SAT variables. As a corollary, our reduction can be modified to yield the first NP-hardness result for the winner verification problem for the sequential CC rule.

Related work

The paper by Procaccia et al. [113] is the most related to ours. Among other results, they prove that the class of single-winner positional scoring rules is efficiently PAC-learnable. We remark that our setting is much more demanding. In particular, the number of possible outcomes is doubly exponential in our case, i.e., $2^{\binom{m}{k}} - 1$, the number of all possible non-empty sets of winning committees, while it is just m in

¹We follow standard notions from parameterized complexity theory, such as W -hierarchy hardness and fixed-parameter tractability; e.g., see Cygan et al. [54].

theirs (where fixed tie-breaking is used to produce a single winning candidate). Hence, even though we have not been able to prove efficiency of learning, the low sample complexity of ABCS rules is rather surprising.

PAC learning in voting has been considered, among other economic paradigms, by Jha and Zick [85] and, in relation to the notion of the distortion, by Boutilier et al. [25]. Actually, the use of sign patterns has been inspired by the latter paper, even though the particular way in which we employ the result of Alon [3] here is different.

More distantly related to our setting, the data-driven approach in the design of voting rules has been followed by a series of papers which focus on particular applications like rating [36], evaluation of online surveys [19], and peer grading [34, 37]. Other foundational work in this direction includes the papers by Faliszewski et al. [66] and Xia [125].

The computational complexity of multiwinner voting rules has received much attention; see the book by Lackner and Skowron [93]. The CC rule has been central in most related studies regarding ABCS rules. Procaccia et al. [112] proved that the problem of deciding whether there is a committee that exceeds a given threshold under the CC rule on a given profile is NP-hard. The problem was later proved to be W[2]-hard by Betzler et al. [22]. Sonar et al. [119] considered the question of whether a given candidate belongs to a winning committee for a given profile. They also prove that winner verification for the CC rule is coNP-hard, using a reduction from a variant of 3-HITTINGSET. We are not aware of published hardness results (of a similar spirit) for sequential Thiele rules.

Roadmap

The rest of the paper is structured as follows. We begin with preliminary definitions and notation in Section 7.2 and present briefly the necessary background on PAC learning in Section 7.3. Sections 7.4-7.6 contain our technical contributions, as outlined in Section 7.1 above. We conclude in Section 7.7, where we highlight our two byproduct results on winner verification; the formal statements of these results and their proofs, which are simplifications of our main hardness proofs, appear in Section 7.8.

7.2 Preliminaries

We consider approval-based voting with a set $\mathcal{N}$ of n voters (or *agents*), each approving a subset from a set Σ of m candidates (or *alternatives*). An approval-based multiwinner voting rule is defined for an integer k with $1 < k < m$. It takes as input a profile $P = \{\sigma_i\}_{i \in \mathcal{N}}$, where $\sigma_i \subset \Sigma$ is the non-empty set of alternatives approved by agent $i \in \mathcal{N}$ (or, her *approval vote*), and returns one or more k -sized subsets of Σ . We use the term *committee* to refer to any k -sized set of alternatives; then, the outcome of a multiwinner voting rule is one or more *winning* committees. We are interested in a specific class of multiwinner voting rules called *approval-based committee scoring* (ABCS) rules, defined by Lackner and Skowron [92]. These rules are specified by a

set of scoring parameters. Using these parameters, an agent's approval vote gives a score to each committee and the winning committees are those that receive the highest total score from all agents.

More formally, an ABCS rule is specified by a bivariate scoring function f . The parameter $f(x, y)$ denotes the non-negative score that an approval vote σ gives to the committee C when σ consists of y alternatives and has x alternatives in common with C . Notice that, under this interpretation, the function f needs only be defined over the set of pairs

$$\mathcal{X}_{m,k} = \{(x, y) : y \in [m-1], x \in \max\{0, y-m+k\}, \dots, \min\{k, y\}\}.$$

Indeed, an approval vote with y alternatives can intersect with a committee in at least $\max\{0, y-m+k\}$ and at most $\min\{k, y\}$ alternatives.

Hence, formally $f : \mathcal{X}_{m,k} \rightarrow \mathbb{R}_{\geq 0}$. By definition, f is monotone non-decreasing in its first argument. To keep the presentation concise, we slightly overload notation and use f to refer both to the scoring function f and the ABCS rule specified by f . On input a profile $P = \{\sigma_i\}_{i \in \mathcal{N}}$, the ABCS rule f assigns a score of

$$\text{sc}_f(C, P) = \sum_{i \in \mathcal{N}} f(|C \cap \sigma_i|, |\sigma_i|)$$

to each committee C ; then, any committee of maximum score is winning in profile P under rule f . We write $f(P)$ for the set of all winning committees in profile P . We denote by $\mathcal{F}^{m,k}$ the class of ABCS rules with m alternatives and committee size k . We use the term *trivial* to refer to the ABCS rule f with $f(x, y) = 0$ for every $(x, y) \in \mathcal{X}_{m,k}$; obviously, all committees are winning in any profile under this rule.²

An important subclass of ABCS rules is that of *Thiele* rules. Thiele rules use scoring functions f where the scoring parameter $f(x, y)$ does not depend on y . In this case, we can assume that f is univariate, defined over $\{0, 1, \dots, k\}$, non-negative, and monotone non-decreasing. A specific Thiele rule that we use extensively is the CC rule that uses $f(0) = 0$ and $f(x) = 1$ for $x > 0$.

To bypass the necessity of computing the scores of all committees, *sequential Thiele rules* have been introduced to approximate ABCS rules by computing a winning committee in a greedy manner. Starting from an empty subcommittee, such rules build a winning committee gradually in k steps; in each step, they include an alternative that increases the score of the current subcommittee the most. The sequential Thiele rule that uses the univariate scoring function f computes the intermediate score of a set of alternatives A of size up to k on profile $P = \{\sigma_i\}_{i \in \mathcal{N}}$ as $\text{sc}_f(A, P) = \sum_{i \in \mathcal{N}} f(|A \cap \sigma_i|)$. Then, given a profile P , a committee C is winning under the sequential Thiele rule f

²We remark that scoring functions with different parameters may correspond to the very same ABCS rule. Indeed, as Lackner and Skowron [92] observe, the ABCS rule g specified by $g(x, y) = c \cdot f(x, y) + d(y)$ for $c > 0$ and $d : [m-1] \rightarrow \mathbb{R}_{\geq 0}$ is identical to the ABCS rule f , in the sense that, for every profile, they define the same set of winning committees. Throughout the paper, we consider ABCS rules whose scoring function is normalized to have $f(\max\{0, y-m+k\}, y) = 0$ for every $y \in [m-1]$. So, the trivial ABCS rule is representative of the rules in which $f(x, y)$ depends only on y ; all these rules have all committees as winning.

in profile P if there is an ordering of the alternatives in C , e.g., as $C = \{c_1, c_2, \dots, c_k\}$, so that

$$c_i \in \arg \max_{c \in \Sigma \setminus \{c_1, \dots, c_{i-1}\}} \text{sc}_f(\{c_1, \dots, c_{i-1}\} \cup \{c\}, P),$$

for every $i \in [k]$. We denote by $\mathcal{F}_{\text{seq}}^k$ the class of sequential Thiele rules for committee size k and any number of alternatives m higher than k . Again, the term *trivial* is reserved for the sequential Thiele rule that uses a scoring function f with $f(x) = 0$ for every x .

We conclude this section by defining the two decision problems we study: TARGETABCS and TARGETSEQTHIELE. In both, we are given a profile of approval votes $P = \{\sigma_i\}_{i \in \mathcal{N}}$ over the set Σ of m alternatives and a k -sized subset C of Σ . Our goal is to decide whether there exists a non-trivial rule f from $\mathcal{F}^{m,k}$ (for TARGETABCS) or $\mathcal{F}_{\text{seq}}^k$ (for TARGETSEQTHIELE), so that C is a winning committee in profile P according to f . Separate definitions of these problems are given as Definitions 58 and 67 in the sections discussing the complexity of the respective problems (Sections 7.5 and 7.6).

7.3 PAC Learning Background

We follow a standard PAC learning model. In this model, a learning algorithm has to learn a target function from a hypothesis class $\mathcal{H}$ of functions which assign labels from the set Y to the points of a set Z . The learning algorithm is given a *training set of examples* T consisting of points from the sample space Z paired with the labels from Y that the target function assigns to the respective data point. The data points in the training set are sampled i.i.d. according to some probability distribution D over Z . We consider the *realizable case* and assume that there exists a function $h^* \in \mathcal{H}$ that is used to label the examples in the training set as $\{(z, h^*(z))\}_{z \in T}$. The learning algorithm outputs a function $h \in \mathcal{H}$. The error of function h is defined as

$$\text{err}(h) = \Pr_{z \sim D} [h(z) \neq h^*(z)].$$

Clearly, $\text{err}(h^*) = 0$. The terms “probably” and “approximately correct” refer to the existence of two parameters $\delta, \epsilon \in (0, 1)$, indicating the required *confidence* and *accuracy* of learning, respectively.

Definition 47 (PAC learnability). *A hypothesis class $\mathcal{H}$ of functions from set Z to set Y is PAC-learnable if there exist a function $s : (0, 1)^2 \rightarrow \mathbb{N}$ —the sample complexity of $\mathcal{H}$ —and a learning algorithm $\mathcal{A}$ with the following property: For every $\delta, \epsilon \in (0, 1)$, every distribution D over Z , and every function h^* from $\mathcal{H}$, on input a training set of at least $s(\delta, \epsilon)$ examples generated by D and labelled by h^* , the probability (over the choice of the training examples) that algorithm $\mathcal{A}$ returns a hypothesis h of error more than ϵ is at most δ .*

Extending the relation of the well-known VC dimension with the PAC learnability of boolean functions, Natarajan [105] relates the sample complexity of a hypothesis class $\mathcal{H}$ to the notions of graph dimension and Natarajan (or generalized) dimension, both capturing the combinatorial richness of $\mathcal{H}$. To define them, we need to define the notion of *shattering* first.

Definition 48 (shattering). *Let $\mathcal{H}$ be a class of functions from Z to Y and let $T \subseteq Z$. We say that $\mathcal{H}$ G-shatters T if there exists a function $g \in \mathcal{H}$ such that*

- *For all $S \subseteq T$, there exists $h_S \in \mathcal{H}$ such that $h_S(z) = g(z)$ for all $z \in S$, and $h_S(z) \neq g(z)$, for all $z \in T \setminus S$.*

We say that $\mathcal{H}$ N-shatters T if there exist two functions $g_1, g_2 \in \mathcal{H}$ such that

1. *For all $z \in T$, $g_1(z) \neq g_2(z)$.*
2. *For all $S \subseteq T$, there exists $h_S \in \mathcal{H}$ such that $h_S(z) = g_1(z)$ for all $z \in S$, and $h_S(z) = g_2(z)$, for all $z \in T \setminus S$.*

Definition 49 (graph and Natarajan dimension). *Let $\mathcal{H}$ be a class of functions from a set Z to a set Y . The graph dimension of $\mathcal{H}$, denoted by $D_G(\mathcal{H})$, is the greatest integer d such that there exists a set of cardinality d that is G-shattered by $\mathcal{H}$. The Natarajan dimension of $\mathcal{H}$, denoted by $D_N(\mathcal{H})$, is the greatest integer d such that there exists a set of cardinality d that is N-shattered by $\mathcal{H}$.*

We will use the relation of the graph and Natarajan dimension to the sample complexity that is given by the next statement (Theorem 50). This is a variant of the multiclass fundamental theorem which, in its standard form (e.g., see Shalev-Shwartz and Ben-David [117]), uses only the Natarajan dimension to both upper- and lower-bound the sample complexity. In the variant we use here, the graph dimension is used instead to upper-bound the sample complexity.³

Theorem 50 (multiclass fundamental theorem, Daniely et al. [56]). *There exist constants $C_1, C_2 > 0$ such that the hypothesis class $\mathcal{H}$ is PAC-learnable (assuming realizability) with sample complexity $s(\delta, \epsilon)$ that satisfies*

$$C_1 \cdot \frac{D_N(\mathcal{H}) + \ln(1/\delta)}{\epsilon} \leq s(\delta, \epsilon) \leq C_2 \cdot \frac{D_G(\mathcal{H}) \cdot \ln(1/\epsilon) + \ln(1/\delta)}{\epsilon}.$$

³In the conference version of the paper [31], we used a formulation of the multiclass fundamental theorem which provides sample complexity upper bounds in terms of the Natarajan dimension and the number of different labels a function from the hypothesis class $\mathcal{H}$ may realize. We prefer to use Theorem 50 instead in this revised version, which leads to considerably simpler proofs and better sample complexity bounds in all cases besides the extreme ones where the accuracy parameter ϵ is exponentially small in terms of $|\mathcal{X}_{m,k}|$ for ABCS rules or in terms of k for sequential Thiele rules.

7.4 The Learnability of ABCS Rules

We are ready to prove that the class $\mathcal{F}^{m,k}$ of ABCS rules is PAC-learnable with sample complexity that depends polynomially on the number of alternatives m and the committee size k .

Theorem 51. *The class $\mathcal{F}^{m,k}$ of ABCS rules with m alternatives and committee size k is PAC-learnable with sample complexity $s(\delta, \varepsilon)$ such that*

$$s(\delta, \varepsilon) \in \Omega\left(\varepsilon^{-1}(|\mathcal{X}_{m,k}| + \ln(1/\delta))\right),$$

and

$$s(\delta, \varepsilon) \in \mathcal{O}\left(\varepsilon^{-1}(|\mathcal{X}_{m,k}| \cdot k \cdot \ln m \cdot \ln(1/\varepsilon) + \ln(1/\delta))\right).$$

Notice that $|\mathcal{X}_{m,k}| \in \Theta(k(m-k))$; so the sample complexity grows only polynomially in m , k , and $1/\varepsilon$, and logarithmically in $1/\delta$. Our proof of Theorem 51 will follow by Theorem 50 after proving an upper bound of $\mathcal{O}(|\mathcal{X}_{m,k}| \cdot k \cdot \ln m)$ on the graph dimension (Lemma 53) and a lower bound of $\Omega(|\mathcal{X}_{m,k}|)$ on the Natarajan dimension of $\mathcal{F}^{m,k}$ (Lemma 54).

Upper-bounding the graph dimension

To bound the graph dimension, we will use an important result in algebraic combinatorics that bounds the number of different sign patterns a set of polynomials may have. Consider a set $\mathcal{L}$ of K polynomials $p_1, p_2, \dots, p_K$, each defined over the ℓ real variables $x_1, x_2, \dots, x_\ell$ (i.e., $p_i : \mathbb{R}^\ell \rightarrow \mathbb{R}$ for $i \in [K]$). A sign pattern $\mathbf{s}$ is just a vector of values in $\{-1, 0, +1\}$ with K entries. We say that the set of polynomials $\mathcal{L}$ realizes the sign pattern $\mathbf{s}$ if there exist values $x_1^*, x_2^*, \dots, x_\ell^*$ for the variables $x_1, x_2, \dots, x_\ell$ such that $\text{sgn}(p_i(x_1^*, x_2^*, \dots, x_\ell^*)) = s_i$, for $i = 1, 2, \dots, K$. Here, sgn is the signum function returning -1 , 0 , or $+1$, depending on whether its argument is negative, zero, or positive.

Clearly, the number of different sign patterns K polynomials may realize is at most 3^K . Usually, this is a very weak upper bound; Alon [3] provides a much better bound, extending a previous statement due to Warren [124].

Theorem 52 (Alon [3], Warren [124]). *The number of different sign patterns a set of K polynomials of degree τ over ℓ real variables may realize is at most $\left(\frac{8e\tau K}{\ell}\right)^\ell$.*

Using Theorem 52, we can prove the next upper bound on the graph dimension $D_G(\mathcal{F}^{m,k})$.

Lemma 53. $D_G(\mathcal{F}^{m,k}) \in \mathcal{O}(|\mathcal{X}_{m,k}| k \log m)$.

Proof. Assume that the graph dimension of $\mathcal{F}^{m,k}$ is N with $N \geq 4 \cdot |\mathcal{X}_{m,k}|$ since, clearly, the upper bound on $D_G(\mathcal{F}^{m,k})$ holds if $N < 4 \cdot |\mathcal{X}_{m,k}|$. According to Definition 49, we thus have a set of N different profiles $\{P_j\}_{j \in [N]}$ and a voting rule $g \in \mathcal{F}^{m,k}$ such that for any subset $S \subseteq [N]$ there exists a rule $h_S \in \mathcal{F}^{m,k}$ with the property that

- $h_S(P_j) = g(P_j)$ for all $j \in S$, and
- $h_S(P_j) \neq g(P_j)$ for all $j \in [N] \setminus S$.

Let $\mathcal{C}^k$ be the set of all k -sized committees. For every pair of committees $C, D \in \mathcal{C}^k$ and every $j \in [N]$, we define

$$L_{C,D}^j(s) = \text{sc}_s(C, P_j) - \text{sc}_s(D, P_j),$$

where s is any scoring function specifying a voting rule in $\mathcal{F}^{m,k}$. Let $P_j = \{\sigma_i^j\}_{i \in \mathcal{N}}$; then

$$L_{C,D}^j(s) = \sum_{i \in \mathcal{N}} s(|C \cap \sigma_i^j|, |\sigma_i^j|) - \sum_{i \in \mathcal{N}} s(|D \cap \sigma_i^j|, |\sigma_i^j|).$$

Hence, $L_{C,D}^j(s)$ is a linear function (a polynomial of degree 1) on the variables $s(x, y)$ for $(x, y) \in \mathcal{X}_{m,k}$. Let $\mathcal{L} = \{L_{C,D}^j(s) : j \in [N] \text{ and } C, D \in \mathcal{C}^k\}$ be the set of linear functions defined for the N different profiles and all pairs of k -sized committees. We note that $|\mathcal{L}| \leq Nm^{2k}$ since $|\mathcal{C}^k| = \binom{m}{k} \leq m^k$.

Now, consider two different subsets S and S' of $[N]$ such that $S \not\subseteq S'$ (notice that this is without loss of generality). Let j^* be such that $j^* \in S \setminus S'$. Since $\mathcal{F}^{m,k}$ G -shatters the set of profiles $\{P_j\}_{j \in [N]}$ by our assumption, there exist two rules $h_S, h_{S'} \in \mathcal{F}^{m,k}$ such that $h_S(P_{j^*}) = g(P_{j^*})$ and $h_{S'}(P_{j^*}) \neq g(P_{j^*})$. Hence, there must be a pair of committees $C, D \in \mathcal{C}^k$ such that $C \in h_S(P_{j^*}), D \in h_{S'}(P_{j^*})$ and either $C \notin h_{S'}(P_{j^*})$ or $D \notin h_S(P_{j^*})$ (not necessarily exclusively). Then,

$$\text{sgn}(L_{C,D}^{j^*}(h_S)) = \text{sgn}(\text{sc}_{h_S}(C, P_{j^*}) - \text{sc}_{h_S}(D, P_{j^*})) = \begin{cases} 0, & \text{if } D \in h_S(P_{j^*}) \\ +1, & \text{if } D \notin h_S(P_{j^*}) \end{cases}$$

and

$$\text{sgn}(L_{C,D}^{j^*}(h_{S'})) = \text{sgn}(\text{sc}_{h_{S'}}(C, P_{j^*}) - \text{sc}_{h_{S'}}(D, P_{j^*})) = \begin{cases} 0, & \text{if } C \in h_{S'}(P_{j^*}) \\ -1, & \text{if } C \notin h_{S'}(P_{j^*}) \end{cases}$$

Since it is either $C \notin h_{S'}(P_{j^*})$ or $D \notin h_S(P_{j^*})$, we get that

$$\text{sgn}(L_{C,D}^{j^*}(h_S)) \neq \text{sgn}(L_{C,D}^{j^*}(h_{S'})).$$

Hence, each of the 2^N voting rules h_S for $S \subseteq [N]$ —corresponding to a distinct assignment of values $s(x, y)$ for $(x, y) \in \mathcal{X}_{m,k}$ —yields a different sign pattern to the set of polynomials $\mathcal{L}$. We now apply Theorem 52 to $\mathcal{L}$ for $K = Nm^{2k}$, $\tau = 1$, $\ell = |\mathcal{X}_{m,k}|$.

This gives an upper bound of $\left(\frac{8eNm^{2k}}{|\mathcal{X}_{m,k}|}\right)^{|\mathcal{X}_{m,k}|}$ on the number of different sign patterns with entries in $\{-1, 0, +1\}$ for the set of polynomials $\mathcal{L}$. Hence,

$$2^N \leq \left(\frac{8eNm^{2k}}{|\mathcal{X}_{m,k}|}\right)^{|\mathcal{X}_{m,k}|}$$

and, equivalently,

$$\frac{|\mathcal{X}_{m,k}|}{N} \cdot 2^{N/|\mathcal{X}_{m,k}|} \leq 8em^{2k}.$$

We now apply the property $2^{z/2} \geq z$ for $z \geq 4$. Using $z = \frac{N}{|\mathcal{X}_{m,k}|}$, combined with our assumption that $N \geq 4|\mathcal{X}_{m,k}|$, we get

$$2^{\frac{1}{2} \cdot N/|\mathcal{X}_{m,k}|} \leq 8em^{2k}$$

and, $N \leq 2|\mathcal{X}_{m,k}| \log(8em^{2k})$, as desired. $\square$

Lower bounding the Natarajan dimension

We now prove a lower bound on $D_N(\mathcal{F}^{m,k})$. In our proof, we construct a large set of profiles that can be N -shattered by the set of ABCS rules $\mathcal{F}^{m,k}$.

Lemma 54. $D_N(\mathcal{F}^{m,k}) \in \Omega(|\mathcal{X}_{m,k}|)$.

Proof. For a given $m \geq 3$ and k such that $2 \leq k \leq m-1$, consider the set of alternatives

$$\Sigma = \{a, b_1, \dots, b_{k-1}, c, d_1, \dots, d_{m-k-1}\}.$$

Our goal is to define a set of profiles, where for each profile we are able to pick rules from $\mathcal{F}^{m,k}$ such that either committee $A = \{a, b_1, \dots, b_{k-1}\}$ or committee $C = \{b_1, \dots, b_{k-1}, c\}$ is the single winning committee under the respective rule.

Let $\mathcal{T}_{m,k}$ be the following set of pairs:

$$\mathcal{T}_{m,k} = \left\{ (x, y) : y \in \{2, \dots, m-1\}, \right. \\ \left. x \in \{1 + \max\{0, y - m + k\}, \dots, \min\{k, y\}\} \setminus \{(k, k)\} \right\}.$$

Even though some of the pairs of set $\mathcal{X}_{m,k}$ have been omitted from $\mathcal{T}_{m,k}$, they have asymptotically the same size as the next lemma indicates.

Lemma 55. $|\mathcal{T}_{m,k}| \in \Omega(|\mathcal{X}_{m,k}|)$.

Proof. The lemma clearly holds for $m < 4$ since both $\mathcal{X}_{m,k}$ and $\mathcal{T}_{m,k}$ have constant size in this case. In the following, we assume $m \geq 4$. Observe that the (x, y) pairs of $\mathcal{X}_{m,k}$ that are missing from $\mathcal{T}_{m,k}$ are the following: $(\max\{0, y - m + k\}, y)$ for $y = 2, \dots, m-1$, $(0, 1)$, $(1, 1)$, and (k, k) . Hence,

$$|\mathcal{T}_{m,k}| = |\mathcal{X}_{m,k}| - m - 1 \tag{7.1}$$

Also, observe that

$$|\mathcal{X}_{m,k}| \geq |\mathcal{X}_{m,2}| = 3m - 5 \geq \frac{7}{5}(m + 1). \tag{7.2}$$

The second inequality follows since $m \geq 4$. Now, using equations (7.1) and (7.2), we get

$$|\mathcal{T}_{m,k}| \geq |\mathcal{X}_{m,k}| - \frac{5}{7}|\mathcal{X}_{m,k}| = \frac{2}{7}|\mathcal{X}_{m,k}|,$$

as desired. $\square$

We now define the set of profiles $\{P_{xy}\}_{(x,y) \in \mathcal{T}_{m,k}}$. Each profile P_{xy} contains four approval votes:

$$\sigma_1^{xy} = \{a\}, \sigma_2^{xy} = A, \sigma_3^{xy} = C,$$

and

$$\sigma_4^{xy} = \{b_1, \dots, b_{x-1}, c, d_1, \dots, d_{y-x}\}.$$

We introduce the family of rules $F \subseteq \mathcal{F}^{m,k}$ which, for every subset $S \subseteq \mathcal{T}_{m,k}$, contains the voting rule h_S defined as:

$$\begin{aligned} h_S(1, 1) &= 1, \\ h_S(k, k) &= 4k - 1, \\ h_S(\max\{0, y - m + k\}, y) &= 0, \quad \text{for } y \in [m - 1], \\ h_S(x, y) - h_S(x - 1, y) &= \begin{cases} 0, & (x, y) \in S, \\ 2, & (x, y) \in \mathcal{T}_{m,k} \setminus S. \end{cases} \end{aligned}$$

Notice that the function h_S is monotonically non-decreasing in its first argument, as required by the definition of voting rules in $\mathcal{F}^{m,k}$.

We will show that the family F N -shatters the set of profiles $\{P_{xy}\}_{(x,y) \in \mathcal{T}_{m,k}}$. To do so, we will make use of the following two lemmas. Lemma 56 guarantees that no other committee besides A and C is ever winning in any profile of $\{P_{xy}\}_{(x,y) \in \mathcal{T}_{m,k}}$. Lemma 57 identifies the winning committee among A and C in each of these profiles for every voting rule in set F .

Lemma 56. *For every committee $X \neq A, C$, every profile $P_{xy} \in \{P_{xy}\}_{(x,y) \in \mathcal{T}_{m,k}}$ and any voting rule $h \in F$, it holds that $\text{sc}_h(A, P_{xy}) - \text{sc}_h(X, P_{xy}) \geq 1$.*

Proof. Let us first assume that at least one of the following two assumptions hold: either $y \neq k$ or $|X \cap \sigma_4^{xy}| < k$. Notice that $h(|X \cap \sigma_4^{xy}|, y) \leq 2k$ in this case, due to the definition of the family F . Then, since X is different than both A and C , we have that

$$\begin{aligned} \text{sc}_h(X, P_{xy}) &= h(|X \cap \{a\}|, 1) + h(|X \cap A|, k) + h(|X \cap C|, k) + h(|X \cap \sigma_4^{xy}|, y) \\ &\leq h(1, 1) + 2h(k - 1, k) + 2k, \end{aligned}$$

while

$$\text{sc}_h(A, P_{xy}) = h(1, 1) + h(k, k) + h(k - 1, k) + h(x - 1, y).$$

Hence,

$$\text{sc}_h(A, P_{xy}) - \text{sc}_h(X, P_{xy}) \geq h(k, k) + h(x-1, y) - h(k-1, k) - 2k \geq 1$$

since $h(x-1, y) \geq 0$, $h(k, k) = 4k-1$, and $h(k-1, k) \leq 2k-2$.

Now, assume that $y = k$ and $|X \cap \sigma_4^{xy}| = k$; then $X = \sigma_4^{xy}$. We have

$$\begin{aligned} \text{sc}_h(X, P_{xy}) &= h(0, 1) + h(x-1, k) + h(x, k) + h(k, k) \\ &\leq h(0, 1) + h(x-1, k) + h(k-1, k) + h(k, k). \end{aligned}$$

The inequality follows since $(k, k) \notin \mathcal{T}_{m,k}$ and, thus, $x \leq k-1$, and by the monotonicity of voting rule h . Hence,

$$\text{sc}_h(A, P_{xy}) - \text{sc}_h(X, P_{xy}) \geq h(1, 1) - h(0, 1) = 1. \quad \square$$

Lemma 57. *Let $S \subseteq \mathcal{T}_{m,k}$. For every profile $P_{xy} \in \{P_{xy}\}_{(x,y) \in \mathcal{T}_{m,k}}$, it holds that*

$$\text{sc}_{h_S}(A, P_{xy}) - \text{sc}_{h_S}(C, P_{xy}) = \begin{cases} 1, & (x, y) \in S \\ -1, & (x, y) \in \mathcal{T}_{m,k} \setminus S \end{cases}$$

Proof. By the definition of profile P_{xy} , we have

$$\text{sc}_{h_S}(A, P_{xy}) = h_S(1, 1) + h_S(k, k) + h_S(k-1, k) + h_S(x-1, y),$$

and

$$\text{sc}_{h_S}(C, P_{xy}) = h_S(0, 1) + h_S(k-1, k) + h_S(k, k) + h_S(x, y).$$

Recall that $h_S(0, 1) = 0$ and $h_S(1, 1) = 1$. Thereby,

$$\text{sc}_{h_S}(A, P_{xy}) - \text{sc}_{h_S}(C, P_{xy}) = 1 - (h_S(x, y) - h_S(x-1, y)),$$

and the lemma follows from the definition of voting rule h_S . $\square$

Together, Lemmas 56 and 57 imply that, when applied to profile P_{xy} , the voting rule h_S returns

- A as the unique winning committee if $(x, y) \in S$, and
- C as the unique winning committee if $(x, y) \in \mathcal{T}_{m,k} \setminus S$.

By Definition 48, this implies that the family F (and, consequently, the family $\mathcal{F}^{m,k}$) N -shatters the set of profiles $\{P_{xy}\}_{(x,y) \in \mathcal{T}_{m,k}}$. Indeed, it suffices to define functions g_1 and g_2 as $g_1 = h_{\mathcal{T}_{m,k}}$ and $g_2 = h_\emptyset$, while the set of profiles $\{P_{xy}\}_{(x,y) \in \mathcal{T}_{m,k}}$ plays the role of set T in Definition 48. By Definition 49, we conclude that the Natarajan dimension of $\mathcal{F}^{m,k}$ is at least $|\mathcal{T}_{m,k}|$. Lemma 54 now follows from Lemma 55. $\square$

7.5 The Complexity of TARGETABCS

Unfortunately, despite the low sample complexity of the class of ABCS rules, learning from samples is notoriously hard. We prove this for TARGETABCS, which captures the elementary task of learning from a single sample. We repeat the formal definition of TARGETABCS here.

Definition 58 (TARGETABCS). *Given a profile of approval votes $P = \{\sigma_i\}_{i \in \mathcal{N}}$ over the set Σ of m alternatives and a k -sized subset C of Σ , decide whether there exists a non-trivial rule f from $\mathcal{F}^{m,k}$, so that C is a winning committee in profile P according to f .*

The next statement uses a polynomial-time reduction from (the complement of) the INDEPENDENTSET problem.

Definition 59 (INDEPENDENTSET). *Given a graph G and a positive integer K , decide whether G contains a set of at least K nodes that form an independent set.*

INDEPENDENTSET is known to be $W[1]$ -hard, parameterized by the independent set size [54, Theorem 13.18].

Theorem 60. *TARGETABCS parameterized by the committee size k is $\text{co}W[1]$ -hard.*

Proof. For a given instance of INDEPENDENTSET consisting of a graph G and an integer K , we construct an instance of TARGETABCS with $k = K$ such that there is a non-trivial rule $f \in \mathcal{F}^{m,k}$ that outputs A as a winning committee in P if and only if G contains no independent set of size K .

Let Δ denote the maximum degree among the vertices of G . We can assume that $\Delta \geq 2$, since INDEPENDENTSET would be trivially solvable in polynomial time otherwise. As a first step in our construction, we modify G to another graph G' as follows. For every vertex $v \in V$, we add $\Delta - \deg(v)$ dummy vertices that are adjacent only to v . Let $G' = (V', E')$ be the resulting graph and let $|V'| = r$. We note that G' may contain more independent sets of size k than G . However, observe that each additional k -sized independent set in G' contains at least one dummy vertex of degree 1, whereas the vertices of G (including those vertices belonging to any independent set in G) have degree $\Delta \geq 2$ in G' . As we will see in the later parts of our proof, the construction of our TARGETABCS instance is not affected by the existence of additional k -sized independent sets in G' due to this difference in vertex degree.

Without loss of generality, we can assume that V' is the set of positive integers in $[r]$. The set of alternatives Σ consists of alternatives a_i and b_i for every vertex $i \in V'$, and the additional alternatives c and d . Let $A = \{a_1, a_2, \dots, a_k\}$. The profile P consists of three parts:

- Part 1 consists of a vote $\{b_i, b_j\}$ for every edge $(i, j) \in E'$.
- Part 2 consists of $k\Delta - 1$ copies of each of the following votes: vote $\{a_i, b_j\}$ for every $i, j \in [r]$, votes $\{a_i, c\}$ and $\{b_i, d\}$ for every $i \in [r]$, vote $\{a_1, d\}$, and vote $\{c, d\}$.

- Part 3 consists of a vote containing alternative d , alternatives $a_1, a_2, \dots, a_{x-1}$, and $y - x$ additional alternatives among alternatives $a_{k+1}, a_{k+2}, \dots, a_r, b_1, \dots, b_r$, for every $(x, y) \in \widehat{\mathcal{X}} = \mathcal{X}_{m,k} \setminus (\{(\max\{0, y - m + k\}, y) : y \in [m - 1]\} \cup \{(1, 2), (2, 2)\})$.

We use P_1 , P_2 , and P_3 to denote the three subprofiles of votes in part 1, 2, and 3, respectively.

Parts 2 and 3 of profile P have important properties that are given in Lemmas 61 and 62.

Lemma 61. *Let $f \in \mathcal{F}^{m,k}$ and $C = \{a_1, \dots, a_{k-1}, d\}$. Then, $\text{sc}_f(A, P_3) = \text{sc}_f(C, P_3)$ if $f(x, y) = 0$ for every $(x, y) \in \mathcal{X}_{m,k} \setminus \{(1, 2), (2, 2)\}$ and $\text{sc}_f(A, P_3) < \text{sc}_f(C, P_3)$, otherwise.*

Proof. Committees A and C get scores of $f(x - 1, y)$ and $f(x, y)$ from the vote of subprofile P_3 corresponding to the pair $(x, y) \in \widehat{\mathcal{X}}$. Hence,

$$\begin{aligned} \text{sc}_f(A, P_3) - \text{sc}_f(C, P_3) &= \sum_{(x,y) \in \widehat{\mathcal{X}}} (f(x - 1, y) - f(x, y)) \\ &= - \sum_{y \in [m-1] \setminus \{2\}} \sum_{x=1+\max\{0, y-m+k\}}^{\min\{k, y\}} (f(x, y) - f(x - 1, y)) \\ &= - \sum_{y \in [m-1] \setminus \{2\}} f(\min\{k, y\}, y). \end{aligned}$$

Thus, the difference $\text{sc}_f(A, P_3) - \text{sc}_f(C, P_3)$ is equal to zero if $f(x, y) = 0$ for every $(x, y) \in \mathcal{X}_{m,k} \setminus \{(1, 2), (2, 2)\}$ and negative otherwise. $\square$

Lemma 62. *Let $f \in \mathcal{F}^{m,k}$ and $C = \{a_1, \dots, a_{k-1}, d\}$. Then, $\text{sc}_f(A, P_2) = \text{sc}_f(C, P_2)$ if $f(1, 2) = f(2, 2)$ and $\text{sc}_f(A, P_2) < \text{sc}_f(C, P_2)$, otherwise.*

Proof. In the subprofile P_2 , committee A gets score $f(1, 2)$ from each of the $k\Delta - 1$ copies of vote $\{a_i, b_j\}$ for $i \in [k]$ and $j \in [r]$, from each of the $k\Delta - 1$ copies of vote $\{a_i, c\}$ for $i \in [k]$ and from each of the $k\Delta - 1$ copies of $\{a_1, d\}$, i.e., $(k\Delta - 1)(kr + k + 1)f(1, 2)$ in total. Committee C gets score $f(1, 2)$ from each of the $k\Delta - 1$ copies of vote $\{a_i, b_j\}$ for $i \in [k - 1]$ and $j \in [r]$, from each of the $k\Delta - 1$ copies of vote $\{a_i, c\}$ for $i \in [k - 1]$, from each of the $k\Delta - 1$ copies of $\{b_i, d\}$ for $i \in [r]$, and from each of the $k\Delta - 1$ copies of vote $\{c, d\}$, and $f(2, 2)$ from each of the $k\Delta - 1$ copies of vote $\{a_1, d\}$, i.e., $(k\Delta - 1)(kr + k)f(1, 2) + (k\Delta - 1)f(2, 2)$. Hence,

$$\text{sc}_f(A, P_2) - \text{sc}_f(C, P_2) = (k\Delta - 1)(f(1, 2) - f(2, 2)),$$

which yields the lemma. $\square$

As no vote in part 1 of profile P includes any alternatives in committees A and C , Lemmas 61 and 62 imply that a non-trivial rule $f \in \mathcal{F}^{m,k}$ can make committee A winning in P only if it satisfies $f(1, 2) = f(2, 2) > 0$ and $f(x, y) = 0$ for any pair (x, y) of $\mathcal{X}_{m,k}$ different than $(1, 2)$ and $(2, 2)$. We complete the proof assuming—without loss of generality—that f furthermore satisfies $f(1, 2) = f(2, 2) = 1$.

Claim 63. *It holds that $\text{sc}_f(A, P) = (k\Delta - 1)(kr + k + 1)$.*

Proof. Committee A gets one point from the $k\Delta - 1$ copies of vote $\{a_i, b_j\}$ for $i \in [k]$ and $j \in [r]$, the $k\Delta - 1$ copies of vote $\{a_i, c\}$ for $i \in [k]$, and the $k\Delta - 1$ copies of vote $\{a_1, d\}$. $\square$

Consider a committee B and let t be the number of its alternatives from $\{b_1, \dots, b_r\}$, and λ , μ , and ν be binary variables indicating whether alternative c , d , and a_1 belongs to B , respectively.

Lemma 64. *If $t < k$, then $\text{sc}_f(B, P) \leq (k\Delta - 1)(kr + k + 1)$.*

Proof. We will show that

$$\text{sc}_f(B, P) \leq (k\Delta - 1)(\mathbb{I}\{t > 0\} + kr + k + \mu + (1 - \mu)\nu - (t + \lambda)(k - t - \lambda)). \quad (7.3)$$

Inequality (7.3) yields the claim by distinguishing between three cases:

- If $t + \lambda = 0$, then the terms $\mathbb{I}\{t > 0\}$ and $(t + \lambda)(k - t - \lambda)$ are equal to 0 and $\mu + \nu - \mu\nu \leq 1$.
- If $k - t - \lambda = 0$, then B contains only alternatives from $\{b_1, \dots, b_r\}$ and alternative c . Hence, μ , ν , and term $(t + \lambda)(k - t - \lambda)$ are all equal to 0.
- If $t + \lambda > 0$ and $k - t - \lambda > 0$, then $(t + \lambda)(k - t - \lambda) \geq 1$, while $\mu + \nu - \mu\nu \leq 1$.

Due to (7.3), all these cases yield the desired inequality for $\text{sc}_f(B, P)$.

To prove (7.3), first observe that B gets at most Δ points for each of the t alternatives from set $\{b_1, \dots, b_r\}$ that B contains. Hence, by the assumption that $t < k$, we obtain that

$$\text{sc}_f(B, P_1) \leq t\Delta \leq (k\Delta - 1)\mathbb{I}\{t > 0\}. \quad (7.4)$$

To bound $\text{sc}_f(B, P_2)$, observe that B gets a point for each of the $k\Delta - 1$ copies of

- vote $\{a_i, b_j\}$ for $i, j \in [r]$ with either $a_i \in B$ or (not exclusively) $b_j \in B$. The total number of these votes is $(k\Delta - 1)(tr + (k - t - \lambda - \mu)r - t(k - t - \lambda - \mu))$.
- vote $\{a_i, c\}$ for $i \in [r]$ with $c \in B$ or $a_i \in B$. The total number of these votes is $(k\Delta - 1)(k - t - \lambda - \mu + r\lambda - \lambda(k - t - \lambda - \mu))$.
- vote $\{b_i, d\}$ for $i \in [r]$ with $d \in B$ or $b_i \in B$. The total number of these votes is $(k\Delta - 1)(r\mu + t - \mu t)$.
- vote $\{a_1, d\}$ if at least one of a_1 and d belongs to B . The number of these votes is $(k\Delta - 1)(\mu + \nu - \mu\nu)$.
- vote $\{c, d\}$ if at least one of c and d belongs to B . The number of these votes is $(k\Delta - 1)(\lambda + \mu - \lambda\mu)$.

Summing these numbers of votes, we get

$$\text{sc}_f(B, P_2) = (k\Delta - 1)(kr + k + \mu + \nu - \mu\nu - (t + \lambda)(k - t - \lambda)). \quad (7.5)$$

Inequality (7.3) then follows by summing equations (7.4) and (7.5) and since $\text{sc}_f(B, P_3) = 0$. $\square$

By Claim 63 and Lemma 64, if committee B has score higher than $\text{sc}_f(A, P)$, then it must be $t = k$. We conclude the proof by reasoning about $\text{sc}_f(B, P)$ in this case.

Claim 65. *Let B be a committee with $t = k$. Then, $\text{sc}_f(B, P_2) = (k\Delta - 1)(kr + k)$.*

Proof. In part 2 of the profile, committee B gets one point from the $k\Delta - 1$ copies of vote $\{a_i, b_j\}$ for $i \in [r]$ and $b_j \in B$ and the $k\Delta - 1$ copies of vote $\{b_i, d\}$ for $b_i \in B$. $\square$

Lemma 66. *Consider any committee B with $t = k$. If G has no independent set of size k , then $\text{sc}_f(B, P_1) \leq k\Delta - 1$.*

Proof. Let S be the set of vertices in G' to which the alternatives in B correspond. Then, $\text{sc}_f(B, P_1)$ is equal to the number of edges in G' that are incident to the vertices of S . These vertices have degree either 1 or Δ . If one of them has degree 1, then $\text{sc}_f(B, P_1) \leq (k - 1)\Delta + 1 \leq k\Delta - 1$. Otherwise, if all of them have degree Δ in G' , then they correspond to vertices of G . Since G has no independent set of size k , at least two vertices of S are connected by an edge in G and, consequently, in G' . Hence, the number of edges incident to the vertices of S and, consequently, $\text{sc}_f(B, P_1)$ is at most $k\Delta - 1$. $\square$

By Claim 65 and Lemma 66, we obtain that if G has no independent set of size k , then $\text{sc}_f(B, P) \leq (k\Delta - 1)(kr + k + 1)$. Thus, by Claim 63, A is a winning committee in this case.

Now, assume that G has an independent set of size k . This implies that G' has an independent set S of k vertices of degree Δ . Now, consider the committee B consisting of the alternatives that correspond to the vertices of S . As the number of edges that are incident to vertices of S is $k\Delta$, we have that $\text{sc}_f(B, P_1) = k\Delta$ as well. Then, by Claims 63 and 65, we have $\text{sc}_f(B, P) = 1 + (k\Delta - 1)(kr + k + 1) > \text{sc}_f(A, P)$ indicating that A is not winning. The proof of correctness of our reduction is now complete. $\square$

By restricting the profile to have only part 1 and a simplified variant of part 2, our reduction yields $\text{coW}[1]$ -hardness of winner verification for the CC rule as stated in Theorem 79 (see Section 7.8 for the detailed statement and proof).

7.6 Sequential Thiele Rules

We now turn our attention to the PAC learnability of sequential Thiele rules and the complexity of the related elementary learning task. We repeat the formal definition of TARGETSEQTHIELE here.

Definition 67 (TARGETSEQTHIELE). *Given a profile of approval votes $P = \{\sigma_i\}_{i \in \mathcal{N}}$ over the set Σ of m alternatives and a k -sized subset C of Σ , decide whether there exists a non-trivial rule f from $\mathcal{F}_{seq}^k$, so that C is a winning committee in profile P according to f .*

In Section 7.4, we saw how the sign of a single linear function can be used to compare the score of two committees in a profile according to an ABCS rule. Due to the different definition of sequential Thiele rules, such a direct comparison is not possible. Still, deciding whether a committee is winning can be done by examining the signs of a *block* of linear functions. This will be our main tool to show that TARGETSEQTHIELE is in FPT and that the class $\mathcal{F}_{seq}^k$ is PAC-learnable.

Assume an arbitrary ordering of the alternatives in Σ . For a committee A and integer $i \in [k]$, we denote by $A(i)$ the i -th alternative of committee A (according to the assumed ordering). For a committee A , permutation $\pi : [k] \rightarrow [k]$, and integer $i \in [k]$, the notation $A[\pi, i]$ is used to denote the set of alternatives $\cup_{j=1}^i \{A(\pi(j))\}$.

Now, assume that the sequential Thiele rule s returns committee A as winning when applied on profile P . Assume that the order in which rule s decides the alternatives in A as winning is given by permutation π : at step i , the rule includes alternative $A(\pi(i))$ in the winning committee. By the definition of the sequential Thiele rule s , this decision can be expressed by the set of inequalities

$$sc_s(A[\pi, i], P) - sc_s(A[\pi, i-1] \cup \{a\}, P) \geq 0, \quad (7.6)$$

for every alternative $a \in \Sigma \setminus A[\pi, i]$. Non-negativity is necessary and sufficient so that alternative $A(\pi(i))$ is (weakly) preferred for inclusion in the winning committee at step i over any alternative $a \in \Sigma \setminus A[\pi, i]$.

For a sequential Thiele rule s , committee A , and permutation π , we define the block $B_{A,\pi}^P(s)$ consisting of the LHS expression of equation (7.6) for every $i \in [k]$ and every alternative $a \in \Sigma \setminus A[\pi, i]$. By the discussion above, committee A is winning in profile P under rule s if and only if there is a permutation π so that all expressions in block $B_{A,\pi}^P(s)$ are non-negative. Otherwise, if the block $B_{A,\pi}^P(s)$ contains a negative expression for every permutation π , committee A is not winning.

We can use this observation to show that TARGETSEQTHIELE can be solved in time $k! \cdot \text{poly}(m, n)$ and, hence, is fixed-parameter tractable. This can be done as follows. For each of the $k!$ permutations π , consider the linear program that has parameters $s(1), \dots, s(k)$ as variables (assuming $s(0) = 0$) and its constraints require that each expression of block $B_{A,\pi}^P(s)$ —each of which is a linear function of the variables—is non-negative and, furthermore, $0 \leq s(1) \leq \dots \leq s(k)$ and $s(k) \geq 1$ to ensure non-negativity, monotonicity, and non-triviality. If the linear program is feasible for some permutation π , then the corresponding scoring function s gives a sequential Thiele rule that makes A a winning committee in profile P . Otherwise, no such rule exists. Checking feasibility can be done in polynomial time using well-known algorithms for linear programming. The next statement summarizes this discussion.

Theorem 68. TARGETSEQTHIELE parameterized by the committee size k is in FPT.

Sample complexity bounds

By adapting our analysis in Section 7.4 and using blocks of linear functions to witness winning committees as discussed above, we can prove Theorem 69. The sample complexity of sequential Thiele rules is polynomial, too. Indeed, it is asymptotically equivalent to the sample complexity of ABCS rules, after replacing $|\mathcal{X}_{m,k}|$ with $k + 1$.

Theorem 69. *The class $\mathcal{F}_{\text{seq}}^k$ of sequential Thiele rules with m alternatives and committee size k is PAC-learnable with sample complexity $s(\delta, \varepsilon)$ such that*

$$s(\delta, \varepsilon) \in \Omega(\varepsilon^{-1}(k + \ln(1/\delta))),$$

and

$$s(\delta, \varepsilon) \in \mathcal{O}(\varepsilon^{-1}(k^2 \cdot \ln m \cdot \ln(1/\varepsilon) + \ln(1/\delta))).$$

Again, the proof of Theorem 69 relies on Theorem 50 and follows after proving an upper bound on quantity $D_G(\mathcal{F}_{\text{seq}}^k)$ (Lemma 70) and a lower bound on quantity $D_N(\mathcal{F}_{\text{seq}}^k)$ (Lemma 71).

Lemma 70. $D_G(\mathcal{F}_{\text{seq}}^k) \in \mathcal{O}(k^2 \ln m)$.

Proof. We follow a similar approach to the proof of Lemma 53. Assume that the graph dimension of $\mathcal{F}_{\text{seq}}^k$ is N , with $N \geq 4k$. Thus, there are N different profiles $\{P_j\}_{j \in [N]}$ and a sequential Thiele rule $g \in \mathcal{F}_{\text{seq}}^k$ such that for any subset $S \subseteq [N]$ there exists a rule $h_S \in \mathcal{F}_{\text{seq}}^k$ with the property that

- $h_S(P_j) = g(P_j)$ for all $j \in S$, and
- $h_S(P_j) \neq g(P_j)$ for all $j \in [N] \setminus S$.

Again, let $\mathcal{C}^k$ be the set of all k -sized committees. We now argue that the set $\mathcal{L}$ of linear functions defined by the blocks $B_{C,\pi}^{P_j}(s)$ for every committee $C \in \mathcal{C}^k$, every permutation $\pi : [k] \rightarrow [k]$, and every $j \in [N]$ realize at least 2^N distinct sign patterns.

Now, consider two different subsets S and S' of $[N]$ such that $S \not\subseteq S'$. Let j^* be such that $j^* \in S \setminus S'$. Since $\mathcal{F}_{\text{seq}}^k$ G -shatters the set of profiles $\{P_j\}_{j \in [N]}$ by our assumption, there exist two rules $h_S, h_{S'} \in \mathcal{F}_{\text{seq}}^k$ such that $h_S(P_{j^*}) = g(P_{j^*})$ and $h_{S'}(P_{j^*}) \neq g(P_{j^*})$. Then, there must be a pair of committees $C, D \in \mathcal{C}^k$ such that $C \in h_S(P_{j^*}), D \in h_{S'}(P_{j^*})$ and either $C \notin h_{S'}(P_{j^*})$ or $D \notin h_S(P_{j^*})$. Thus, at least one of the following cases must happen:

- $C \notin h_{S'}(P_{j^*})$. Then, for every permutation π , one of the linear functions of block $B_{C,\pi}^{P_{j^*}}(s)$ is negative for $s = h_{S'}$. In contrast, by our assumption that $C \in h_S(P_{j^*})$, we know that there is a permutation π so that all linear functions of block $B_{C,\pi}^{P_{j^*}}(s)$ are non-negative for $s = h_S$. Hence, setting $s = h_S$ and $s = h_{S'}$ yields different sign patterns to the set of linear functions $\mathcal{L}$.

- $D \notin h_S(P_{j^*})$. Then, for every permutation π , one of the linear functions of block $B_{D,\pi}^{P_{j^*}}(s)$ is negative for $s = h_S$. In contrast, by our assumption that $D \in h_{S'}(P_{j^*})$, there is a permutation π so that all linear functions of block $B_{D,\pi}^{P_{j^*}}(s)$ are non-negative for $s = h_{S'}$. Again, setting $s = h_S$ and $s = h_{S'}$ yields different sign patterns to $\mathcal{L}$.

We now apply Theorem 52 to $\mathcal{L}$ with $\tau = 1$ and $\ell = k$. To bound the number K of linear functions, observe that $\mathcal{L}$ contains for each of the N profiles $P_1, \dots, P_N$, every committee $C \in \mathcal{C}^k$, and each permutation π a block with at most $k \cdot m$ functions. Since $|\mathcal{C}^k| = \binom{m}{k}$, we have $K \leq N \cdot \binom{m}{k} \cdot k! \cdot k \cdot m \leq N \cdot m^{k+1} \cdot k$. Theorem 52 then gives us an upper bound of $(8eNm^{k+1})^k$ different sign patterns; this must be at least as high as the number of distinct voting rules defined by the 2^N possible selections of a set $S \subseteq [N]$. We obtain that

$$2^{N/k} - 8eNm^{k+1} \leq 0,$$

which implies

$$\frac{k}{N} \cdot 2^{N/k} \leq 8em^{k+2}.$$

We now apply the inequality $2^{z/2} \geq z$ for $z \geq 4$. Using $z = N/k$, combined with our assumption that $N \geq 4k$, we get $2^{\frac{1}{2} \cdot N/k} \leq 8em^{k+2}$, which yields $N \leq 2k \log(8em^{k+2})$, as desired. $\square$

Lemma 71. $D_N(\mathcal{F}_{seq}^k) \in \Omega(k)$.

Proof. Let $m = k + 1$ and consider the set of alternatives $\Sigma = \{b_1, b_2, \dots, b_{k-1}, a, c\}$. We define $k - 1$ profiles and a set F of $k - 1$ rules from $\mathcal{F}_{seq}^k$ which N -shatters these profiles.

For $x = 2, 3, \dots, k$, profile P_x contains: three votes $\sigma_1^{x,i} = \sigma_2^{x,i} = \sigma_3^{x,i} = \{b_i\}$ for $i = 1, 2, \dots, k - 1$, a vote $\sigma_4^x = \{a\}$, and a vote $\sigma_5^x = \{b_1, b_2, \dots, b_{x-1}, c\}$. Also, define the family of voting rules $F \subseteq \mathcal{F}_{seq}^k$ which, for every subset S of $\{2, 3, \dots, k\}$, contains the rule h_S defined as:

$$\begin{aligned} h_S(0) &= 0 \\ h_S(1) &= 1 \\ h_S(x) - h_S(x-1) &= \begin{cases} 0, & \text{if } x \in S \\ 2, & \text{if } x \in \{2, 3, \dots, k\} \setminus S \end{cases} \end{aligned}$$

Let us see how rule S works for profile P_x . Initially, the winning committee is empty. Adding alternative b_i to the winning committee increases the score by 4 if $1 \leq i < x$ and by 3 if $x \leq i \leq k - 1$. Adding alternatives a or c would increase the score by 1. Hence, some alternative b_i for $i < x$ is included in the winning committee in the first step.

Next, we claim that all b -alternatives are included in the winning committee in the first $k - 1$ steps. Indeed, assume that j b -alternatives (with $1 \leq j < k - 1$) have been included in the winning committee. At step $j + 1$, adding another b -alternative increases the score by at least 3. Adding alternative c would increase the score by $h_S(j + 1) - h_S(j)$ if $j < x$ and by $h_S(x) - h_S(x - 1)$ if $j \geq x$, i.e., by no more than 2. Adding alternative a would increase the score by just 1. Hence, some b -alternative is included in the winning committee at step $j + 1$, too.

At the final k -th step, we are left with alternatives a and c . Including alternative a in the committee increases the score by 1. Including alternative c increases the score by $h_S(x) - h_S(x - 1)$, which is either 0 or 2. Hence, $A = \{b_1, \dots, b_{k-1}, a\}$ or $C = \{b_1, \dots, b_{k-1}, c\}$ is the unique winning committee returned by rule h_S when applied on profile P_x , depending on whether x belongs to S or not.

By Definition 48, this means that the family F (and, consequently, the family $\mathcal{F}_{\text{seq}}^k$) N -shatters the profiles $P_2, \dots, P_k$ (by defining g_1 and g_2 as $g_1 = h_{\{2,3,\dots,k\}}$ and $g_2 = h_\emptyset$, respectively). By Definition 49, we conclude that $D_N(\mathcal{F}_{\text{seq}}^k) \geq k - 1$. $\square$

NP-hardness of TARGETSEQTHIELE

The last statement of this section is negative and provides evidence that learning in class $\mathcal{F}_{\text{seq}}^k$ is hard as well. The proof employs a novel reduction from a structured version of 3SAT, known to be NP-hard [126].

Definition 72 (2P2N-3SAT problem). *In the 2P2N-3SAT problem, we are given a 3-CNF formula ϕ with r variables $x_1, \dots, x_r$ and t clauses $c_1, \dots, c_t$. Each variable appears in ϕ twice as a positive literal and twice as a negative literal. The goal is to decide whether there exists an assignment α of boolean values to the variables so that ϕ is satisfied.*

Theorem 73. TARGETSEQTHIELE is NP-hard.

Proof. We prove the theorem by presenting a polynomial-time reduction from the 2P2N-3SAT problem. Let ϕ be the given 3-CNF formula on r variables and t clauses. In our reduction, we have various types of alternatives. There are eight *padding alternatives* $\{p, w_1, \dots, w_7\}$. The $2r$ literals derived from the variables of ϕ form another set of *literal alternatives* $\{x_1, \bar{x}_1, x_2, \bar{x}_2, \dots, x_r, \bar{x}_r\}$. The t clauses of ϕ are represented by a set of *clause alternatives* $\{c_1, \dots, c_t\}$. There is a set of *special alternatives* $\{s_1, \dots, s_t, z\}$. In addition, the set of alternatives Σ includes a number of dummy alternatives which we denote by d followed by a subscript. Let $k = 2r + t + 8$ and

$$A = \{p, w_1, \dots, w_7, x_1, \bar{x}_1, x_2, \bar{x}_2, \dots, x_r, \bar{x}_r, c_1, \dots, c_t\}.$$

Before we define profile P , we introduce some notation. For any clause c of ϕ , let

$$\text{lit}(c) = \{l_1, l_2, l_3\} \subseteq \{x_1, \bar{x}_1, x_2, \bar{x}_2, \dots, x_r, \bar{x}_r\}$$

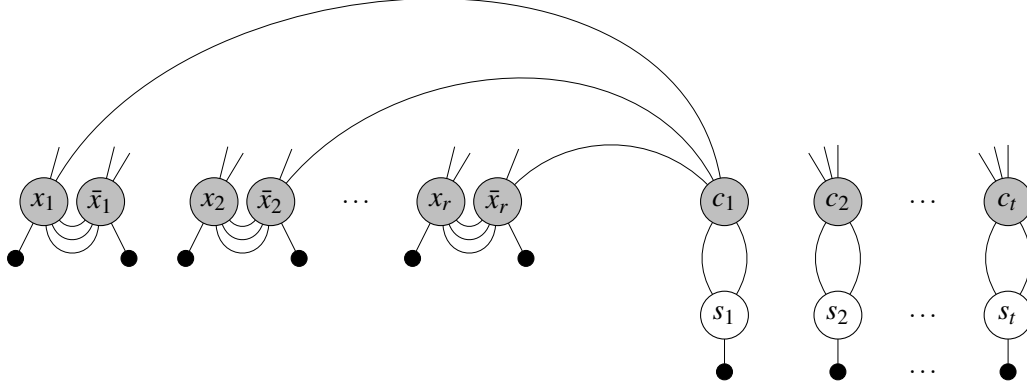


Figure 7.1: A graph representation of part 1 of profile P . The vertices correspond to alternatives and the edges to votes. Alternatives belonging to committee A are highlighted in gray. In this example, we assume that $\text{lit}(c_1) = \{x_1, \bar{x}_2, \bar{x}_r\}$. The other votes containing a literal alternative and a clause alternative are only indicated as edge stubs. White vertices are the special alternatives. Dummy alternatives are shown as small black vertices without a label.

be the set of literals appearing in c . Let S be a set of alternatives. For some ordering of the alternatives in S and any $i \in [|S|]$, we use $S(i)$ to denote the i -th alternative of the set S . The profile P then again consists of three parts.

- Part 1 is derived from the given 3-CNF formula ϕ . For every $i \in [r]$, profile P includes three copies of vote $\{x_i, \bar{x}_i\}$ and the votes $\{x_i, d_{x_i}\}, \{\bar{x}_i, d_{\bar{x}_i}\}$. For every $i \in [t]$, there are the votes $\{c_i, l_1\}, \{c_i, l_2\}, \{c_i, l_3\}$ where $\{l_1, l_2, l_3\} = \text{lit}(c_i)$, two copies of the vote $\{c_i, s_i\}$, and a vote $\{s_i, d_{s_i}\}$. Figure 7.1 shows a graph representation of part 1 of profile P .
- Part 2 consists of the following votes: votes $\{p, z\}$ and $\{p, d_{p,j}\}$ for every $j \in [9]$, and votes $\{w_i, z\}$ and $\{w_i, d_{w_i,j}\}$ for every $i \in [7], j \in [6]$.
- For part 3, we define $S = A \cup \{s_1, \dots, s_t\}$ and impose an arbitrary ordering over the alternatives in S . Notice that $|S| = k + t$. Then, part 3 consists of the votes $\{S(1), S(2), \dots, S(i-1), z, S(i+1), \dots, S(k+t)\}$ for $i \in [k+t]$.

We refer to the subprofiles of P in part 1, 2, and 3 as P_1, P_2 , and P_3 , respectively.

Throughout the proof, for the sequential Thiele rule f applied to profile P_j , we use the notation $\Delta_f(C_i, a, P_j)$ to denote the *score increase* when including alternative a in the subcommittee C_i , i.e.,

$$\Delta_f(C_i, a, P_j) = \text{sc}_f(C_i \cup \{a\}, P_j) - \text{sc}_f(C_i, P_j).$$

In Claims 74 and 75, we make two important observations about part 3 of profile P .

Claim 74. For any rule $f \in \mathcal{F}_{seq}^k$, let $C_i \subseteq A$ be the set of winning alternatives in step $i \in \{0, \dots, k-1\}$. For any two alternatives $a, a' \in S \setminus C_i$, it holds that

$$\Delta_f(C_i, a, P_3) = \Delta_f(C_i, a', P_3).$$

Proof. For any alternative $a \in S \setminus C_i$, observe that a appears in $k+t-1$ votes of P_3 . In i of these votes, a appears together with z and $i-1$ alternatives from C_i . In the remaining $k+t-1-i$ votes, a appears together with all i alternatives from C_i . Recall that $z \notin A$ and therefore $z \notin C_i \subseteq A$. We thus have that

$$\Delta_f(C_i, a, P_3) = (k+t-1-i)(f(i+1) - f(i)) + i(f(i) - f(i-1)). \quad (7.7)$$

Notice that this statement is well-defined also for $i=0$ since the term that contains $f(i-1) = f(-1)$ vanishes in this case. Thus, any two alternatives $a, a' \in S \setminus C_i$ receive the same additional score from subprofile P_3 . $\square$

Claim 75. For any rule $f \in \mathcal{F}_{seq}^k$, let $C_i \subseteq A$ be the set of winning alternatives in step $i \in \{0, \dots, k-1\}$. For any alternative $a \in S \setminus C_i$, it holds that

$$\Delta_f(C_i, z, P_3) \geq \Delta_f(C_i, a, P_3),$$

with strict inequality if $f(i+1) - f(i) > 0$.

Proof. Again, recall that $z \notin A$ such that $z \notin C_i \subseteq A$. Alternative z appears in all $k+t$ votes of P_3 . In i of these votes, z appears together with $i-1$ alternatives from C_i . In the remaining $k+t-i$ votes, z appears together with all i alternatives from C_i . The claim then follows from the observation that

$$\begin{aligned} \Delta_f(C_i, z, P_3) &= (k+t-i)(f(i+1) - f(i)) + i(f(i) - f(i-1)) \\ &= \Delta_f(C_i, a, P_3) + f(i+1) - f(i), \end{aligned}$$

where we used (7.7) to obtain the second equality. $\square$

We proceed by narrowing down the choices of parameters that may make A a winning committee in P .

Lemma 76. For any non-trivial rule $f \in \mathcal{F}_{seq}^k$, committee A is winning in P only if $f(1) = f(2) > 0$.

Proof. We distinguish between three cases where the condition $f(1) = f(2) > 0$ is not true, in each of which A is not winning. First, assume that $f(1) = f(2) = 0$. Since f is non-trivial, there is some $i \in \{2, \dots, k-1\}$ such that $f(i+1) - f(i) > 0$. Let i be minimal such that this is the case and let $C_i \subseteq A$ be the winning committee in step i . Since $i \geq 2$, subprofiles P_1, P_2 do not grant additional score to any alternative $a \in A \setminus C_i$. From Claim 75, it then follows that alternative z has strictly higher score increase than any alternative $a \in A \setminus C_i$. Thus, alternative z is included in the winning committee such that A is not winning.

Now, assume that $f(1) = 0$ and $f(2) > 0$. Note that the padding alternatives $p, w_1, \dots, w_7$ belong to committee A . Consider the first step i in which an alternative from $\{p, w_1, \dots, w_7\}$ is included in the winning subcommittee $C_i \subseteq A$. Since no two alternatives from $\{p, w_1, \dots, w_7\}$ belong together in a vote of P_2 and since $f(1) = 0$, we have

$$\Delta_f(C_i, a, P) = \Delta_f(C_i, a, P_3)$$

for every $a \in \{p, w_1, \dots, w_7\} \setminus C_i$. Consider the alternative z . We observe that z does not appear in any vote of P_1 and shares exactly one vote with a among the votes of P_2 . Hence,

$$\begin{aligned} \Delta_f(C_i, z, P) &= \Delta_f(C_i, z, P_2) + \Delta_f(C_i, z, P_3) = f(2) + \Delta_f(C_i, z, P_3) \\ &\geq f(2) + \Delta_f(C_i, a, P_3) \\ &= f(2) + \Delta_f(C_i, a, P) \\ &> \Delta_f(C_i, a, P). \end{aligned}$$

Here, the first inequality is due to Claim 75, and the second (strict) inequality follows from the assumption that $f(2) > 0$. Hence, alternative z is included in the winning committee before any alternative $a \in \{p, w_1, \dots, w_7\} \setminus C_i$.

Lastly, assume that $f(2) > f(1) > 0$. In this case, the winning committee in step 1 is $C_1 = \{p\}$, since alternative p appears in the most votes from profile P . More specifically, p appears in 10 votes of P_2 and $k+t-1$ votes of P_3 (i.e., $9+k+t$ votes in total) while all other alternatives from A appear in at most 7 votes among $P_1 \cup P_2$ and $k+t-1$ votes of P_3 (i.e., at most $6+k+t$ votes in total). Alternative z appears in 8 votes of P_2 , and $k+t$ votes of P_3 (i.e., $8+k+t$ votes in total). By our construction, the remaining alternatives clearly appear in strictly less votes than p .

Consider an alternative $a \in A \setminus \{p\}$. Again, observe that, in profiles P_1 and P_2 , alternative a appears in at most 7 votes, and none of these votes contain p . Furthermore, alternative a appears in $k+t-1$ votes of P_3 , $k+t-2$ of which also contain p . Hence,

$$\Delta_f(\{p\}, a, P) \leq 8f(1) + (k+t-2)(f(2) - f(1)). \quad (7.8)$$

On the other hand, alternative z appears in 8 votes of P_2 , and in all $k+t$ votes of P_3 . Among these votes, p belongs to one vote of P_2 , and to $k+t-1$ votes of P_3 . Hence,

$$\begin{aligned} \Delta_f(\{p\}, z, P) &= 7f(1) + f(2) - f(1) + f(1) + (k+t-1)(f(2) - f(1)) \\ &= 8f(1) + (k+t-2)(f(2) - f(1)) + 2(f(2) - f(1)) \\ &\geq \Delta_f(\{p\}, a, P) + 2(f(2) - f(1)) \\ &> \Delta_f(\{p\}, a, P), \end{aligned}$$

where we used (7.8) to obtain the first inequality, and the second (strict) inequality is due to the assumption that $f(2) > f(1) > 0$. Thus, in the second step, alternative z is included in the winning committee before any alternative $a \in A \setminus \{p\}$. $\square$

We complete the proof assuming—without loss of generality—that $f(1) = f(2) = 1$. For any such rule, we make a simple, yet important observations which follows from Claim 74.

Observation 77. *Consider any two alternatives $a, a' \in S$ and let rule f be such that $f(1) = f(2) = 1$. At any step during the evaluation of f on profile P , whether alternative a has larger score increase than alternative a' depends only on the votes of P_1 and P_2 .*

Thus, in order for A to be winning, alternatives $p, w_1, \dots, w_7$ are included in the winning committee first. Then, for every variable x_i in ϕ , either alternative x_i or $\bar{x}_i$ (exclusively) is included. After that, any clause alternative c is selected as long as there is at least one literal alternative $l \in \text{lit}(c)$ such that l is not yet included in the winning committee.⁴

Lemma 78. *If there is no assignment to the variables $x_1, \dots, x_r$ satisfying ϕ , then there is some $j^* \in [t]$ such that alternative s_{j^*} is included in the winning committee before alternative c_{j^*} .*

Proof. Let $\tilde{A}$ be the winning committee after the inclusion of alternatives $p, w_1, \dots, w_5$ and either x_i or $\bar{x}_i$ (exclusively) for every $i \in [r]$. Since there is no assignment to the variables that satisfies ϕ , there must now be a clause alternative c_{j^*} for which all literal alternatives corresponding to literals in $\text{lit}(c_{j^*})$ are already included in the winning committee. Otherwise, the ordering in which the r literal alternatives have been included in the winning committee would correspond to an assignment α that satisfies ϕ . That is, we could pick α such that

$$\alpha(l) = \begin{cases} 0 & \text{if } l \in \tilde{A}, \\ 1 & \text{otherwise,} \end{cases}$$

for every $l \in \{x_1, \bar{x}_1, \dots, x_r, \bar{x}_r\}$.

In order for A to be winning on profile P , c_{j^*} needs to be included in the winning subcommittee in some step. However, notice that for any $C_i \subseteq A$, it holds that

$$\Delta_f(C_i, s_{j^*}, P_1) = 3 > 2 = \Delta_f(C_i, c_{j^*}, P_1).$$

Outside of P_1 , the alternatives c_{j^*}, s_{j^*} appear only in votes of P_3 . But, from Claim 74, it follows that $\Delta_f(C_i, c_{j^*}, P_3) = \Delta_f(C_i, s_{j^*}, P_3)$. Hence, $\Delta_f(C_i, s_{j^*}, P) > \Delta_f(C_i, c_{j^*}, P)$ such that s_{j^*} is included in the winning committee before c_{j^*} . $\square$

⁴There is an intuitive way of keeping track of the score increase that the alternatives earn from profile P_1 using the graph representation in Figure 7.1. For $f(1) = f(2) = 1$, including an alternative in the winning committee is equivalent to removing all edges incident with this alternative's node from the graph. At any step, the score increase that an alternative receives from P_1 is the number of edges incident with this alternative's node in the respective step.

So far, we have shown that committee A is not winning on profile P under any non-trivial rule $f \in \mathcal{F}_{\text{seq}}^k$ if there exists no assignment that satisfies ϕ . Now, assume that there is an assignment α to the variables $x_1, \dots, x_r$ that satisfies ϕ . We rename the literals according to their values under the assignment α such that

$$l_i^1 = \begin{cases} x_i & \text{if } \alpha(x_i) = 1 \\ \bar{x}_i & \text{otherwise,} \end{cases}$$

and

$$l_i^0 = \begin{cases} x_i & \text{if } \alpha(x_i) = 0 \\ \bar{x}_i & \text{otherwise.} \end{cases}$$

Let f be such that $f(1) = f(2) = \dots = f(k) = 1$. We conclude the proof by showing that there is an ordering of the alternatives in A such that the sequential Thiele rule specified by f returns A as the winning committee. Using Observation 77 and, with respect to the score increase of alternative z , also Claim 75, the feasibility of the following sequence of inclusions can easily be verified.

First, the alternatives $p, w_1, \dots, w_7$ and the literal alternatives $l_1^0, \dots, l_r^0$ are added to the committee. Conversely, none of the alternatives $l_1^1, \dots, l_r^1$ have been included in the winning committee yet. Since α is a satisfying assignment, each clause alternative c_i for $i \in [t]$ is included in at least one approval vote $\{l, c_i\}$ where $l \in \{l_1^1, \dots, l_r^1\}$. Thus, every clause alternative c_i has a score increase of at least 3 which is at least as high as the score increase of the corresponding special alternative s_i and of any remaining literal alternative. This implies that there is an ordering of the clause alternatives such that each of these alternatives is included in the winning committee. At this point, all the remaining alternatives have a score increase of at most 1. All alternatives from A that have not yet been selected for the winning committee have a score increase of exactly 1. Hence, the alternatives $l_1^1, \dots, l_r^1$ can be included in the winning committee as well. Thereby, A is indeed winning under rule f on profile P . $\square$

We remark that by considering only part 1 of profile P , our reduction yields the NP-hardness of winner verification for the sequential CC rule (see Section 7.8 for a detailed statement and proof).

7.7 Concluding Remarks

We studied complexity aspects of learning ABCS and sequential Thiele rules. In a nutshell, our results suggest that learning from these classes is feasible in the PAC learning framework but—in a worst-case sense—only in computationally inefficient ways. We believe that our techniques for assessing PAC learnability can be extended to other rules. Faliszewski et al. [65] define a hierarchy of classes of ranking-based multiwinner voting rules that are specified using scoring functions. These are natural candidates for extending our analysis. We also remark that, en route to proving hardness of TARGETABCS and TARGETSEQTHIELE, parts of our reductions show

hardness of winner verification for the CC and the sequential CC rule. Formal statements and proofs of these two byproduct results appear in Section 7.8.

Acknowledgements. We thank Piotr Faliszewski for his insights on the complexity of sequential Thiele rules and for pointers to the literature. We also thank Allan Grönlund for helpful discussions on the multiclass fundamental theorem, Chris Schwiegelshohn for discussion on the graph dimension, Krzysztof Sornat for comments on an earlier version of this paper, and an anonymous reviewer for spotting a subtle error in an earlier version of the proof of Theorem 73.

7.8 Appendix: Hardness of Winner Verification for the CC Rule

We introduced the decision problems TARGETABCS and TARGETSEQTHIELE in Section 7.2. A more restricted variant of these problems is known under the term *winner verification*. In addition to the inputs defined for our problems, that is, a profile P and a k -sized committee C , we are now also given a specific rule f from $\mathcal{F}^{m,k}$ or $\mathcal{F}_{seq}^k$. The goal of the winner verification problem is to decide whether committee C is winning on profile P under the given rule f . Our hardness proofs for TARGETABCS and TARGETSEQTHIELE produce as a byproduct corresponding hardness proofs for winner verification under the CC rule, and its sequential counterpart, the sequential CC rule. Recall that we introduced the CC rule in Section 7.2 as an example of a Thiele rule, a rule whose scoring function does not depend on the y parameter and can therefore be treated as univariate. The CC rule is specified by the univariate scoring function f_{CC} where $f_{CC}(0) = 0$, and $f_{CC}(x) = 1$ for $x > 0$. The sequential CC rule is the sequential Thiele rule that uses f_{CC} as its scoring function.

In the following, we show coW[1]-hardness of winner verification for the CC rule (Theorem 79) and thereby strengthen a recent result by Sonar et al. [119]. We then proceed to show NP-hardness of winner verification for the sequential CC rule (Theorem 84). Both proofs rely on simplifications to our constructions in the hardness proofs of TARGETABCS (Theorem 60) and TARGETSEQTHIELE (Theorem 73), respectively.

Theorem 79. *Winner verification for the CC rule parameterized by the committee size k is coW[1]-hard.*

Proof. We reduce from INDEPENDENTSET (Definition 59), proceeding along the same lines as the proof of Theorem 60. We set $k = K$ and define the graph $G' = (V', E')$ in the same way as before. In particular, every node $v \in V'$ has degree $\Delta \geq 2$. Without loss of generality, we can assume that V' is the set of positive integers in $[r]$. The set of alternatives Σ consists of alternatives a_i and b_i for every vertex $i \in V'$, and an additional alternative c . Let $A = \{a_1, a_2, \dots, a_k\}$. The profile P consists of two parts:

- Part 1 consists of a vote $\{b_i, b_j\}$ for every edge $(i, j) \in E'$.

- Part 2 consists of $k\Delta - 1$ copies of each of the following votes: vote $\{a_i, b_j\}$ for every $i, j \in [r]$, and a vote $\{a_1, c\}$.

We use P_1 , and P_2 to denote the two subprofiles of votes in parts 1 and 2, respectively.

Our first claim gives the score of committee A in profile P under the CC rule.

Claim 80. *It holds that $\text{sc}_{f_{cc}}(A, P) = (k\Delta - 1)(kr + 1)$.*

Proof. Committee A gets one point from the $k\Delta - 1$ copies of vote $\{a_i, b_j\}$ for $i \in [k]$ and $j \in [r]$, and from the $k\Delta - 1$ copies of vote $\{a_1, c\}$. $\square$

Now, consider a committee B and let t be the number of its alternatives from $\{b_1, \dots, b_r\}$, and λ , and ν be binary variables indicating whether alternative c , and a_1 belong to B , respectively.

Claim 81. *If $t < k$, $\text{sc}_{f_{cc}}(B, P) \leq (k\Delta - 1)(kr + 1)$.*

Proof. First, observe that B gets at most Δ points for each of the $t < k$ alternatives from set $\{b_1, \dots, b_r\}$ that B contains. Hence,

$$\text{sc}_{f_{cc}}(B, P_1) \leq t\Delta \leq (k\Delta - 1)\mathbb{I}\{t > 0\}. \quad (7.9)$$

To bound $\text{sc}_{f_{cc}}(B, P_2)$, observe that B gets a point for each of the $k\Delta - 1$ copies of

- vote $\{a_i, b_j\}$ for $i, j \in [r]$ with either $a_i \in B$ or (not exclusively) $b_j \in B$. The total number of these votes is $(k\Delta - 1)(tr + (k - t - \lambda)r - t(k - t - \lambda))$.
- vote $\{a_1, c\}$ if at least one of a_1 and c belong to B . The number of these votes is $(k\Delta - 1)(\lambda + \nu - \lambda\nu)$.

Summing these numbers of votes, we get

$$\text{sc}_{f_{cc}}(B, P_2) \leq (k\Delta - 1)(kr - \lambda(r - 1) - t(k - t - \lambda) + (1 - \lambda)\nu). \quad (7.10)$$

Combining inequalities (7.9) and (7.10), we obtain

$$\text{sc}_{f_{cc}}(B, P) \leq (k\Delta - 1)(\mathbb{I}\{t > 0\} + kr - \lambda(r - 1) - t(k - t - \lambda) + (1 - \lambda)\nu). \quad (7.11)$$

We now distinguish between three cases:

- If $k - t - \lambda = 0$, then, since $k \geq 2$ and $t < k$, it must hold that $t = k - 1 > 0$ implying $\mathbb{I}\{t > 0\}$ equals 1, and $\lambda = 1$.
- If $k - t - \lambda > 0$ and $t = 0$, then $\mathbb{I}\{t > 0\} = 0$.
- Lastly, if $k - t - \lambda > 0$ and $t > 0$, we get $\mathbb{I}\{t > 0\} = 1$.

The claim for $\text{sc}_{f_{cc}}(B, P)$ trivially follows from inequality (7.11) for all the three cases above. $\square$

By Claims 80 and 81, if committee B has score higher than $\text{sc}_{f_{CC}}(A, P)$, then it must be that $t = k$. We conclude the proof by reasoning about $\text{sc}_{f_{CC}}(B, P)$ in this case.

Claim 82. *Let B be a committee with $t = k$. Then, $\text{sc}_{f_{CC}}(B, P_2) = (k\Delta - 1)kr$.*

Proof. In part 2 of the profile, committee B gets one point from the $k\Delta - 1$ copies of vote $\{a_i, b_j\}$ for $i \in [r]$ and $b_j \in B$. $\square$

Note that we defined subprofile P_1 in exactly the same way as we did for the proof of Theorem 60. We also have that $f_{CC}(1) = f_{CC}(2) = 1$ by definition. Thus, Lemma 66 is also applicable in the context of this proof. We state the lemma here again as a corollary.

Corollary 83. *Consider any committee B with $t = k$. If G has no independent set of size k , then $\text{sc}_{f_{CC}}(B, P_1) \leq k\Delta - 1$.*

By Claim 82 and Corollary 83, we obtain that if G has no independent set of size k , then $\text{sc}_{f_{CC}}(B, P) \leq (k\Delta - 1)(kr + 1)$. Thus, by Claim 80, A is a winning committee in this case.

Now, assume that G has an independent set of size k . This implies that G' has an independent set S of k vertices of degree Δ . Now, consider the committee B consisting of the alternatives that correspond to the vertices of S . As the number of edges that are incident to vertices of S is $k\Delta$, we have that $\text{sc}_{f_{CC}}(B, P_1) = k\Delta$ as well. Then, by Claims 80 and 82, we have $\text{sc}_{f_{CC}}(B, P) = 1 + (k\Delta - 1)(kr + 1) > \text{sc}_{f_{CC}}(A, P)$, indicating that A is not winning. The proof of correctness of our reduction is now complete. $\square$

Theorem 84. *Winner verification for the sequential CC rule is NP-hard.*

Proof. We reduce from 2P2N-3SAT (Definition 72). Let ϕ be the given 3-CNF formula on r variables and t clauses. In our reduction, we define different types of alternatives based on ϕ . The $2r$ literals derived from the variables of ϕ form a set of *literal alternatives* $\{x_1, \bar{x}_1, x_2, \bar{x}_2, \dots, x_r, \bar{x}_r\}$. The t clauses of ϕ are represented by a set of *clause alternatives* $\{c_1, \dots, c_t\}$. There is a set of *special alternatives* $\{s_1, \dots, s_t\}$. In addition, for the literal alternatives and the special candidates, there exists a corresponding set of *dummy alternatives* $\{d_{x_1}, d_{\bar{x}_1}, \dots, d_{x_r}, d_{\bar{x}_r}, d_{s_1}, \dots, d_{s_t}\}$. Let

$$A = \{x_1, \bar{x}_1, x_2, \bar{x}_2, \dots, x_r, \bar{x}_r, c_1, \dots, c_t\}.$$

Recall that, for any clause c of ϕ ,

$$\text{lit}(c) = \{l_1, l_2, l_3\} \subseteq \{x_1, \bar{x}_1, x_2, \bar{x}_2, \dots, x_r, \bar{x}_r\}$$

is the set of literals appearing in c . The profile P is then defined in the exact same way as subprofile P_1 in the proof of Theorem 73. For every $i \in [r]$, profile P includes three copies of vote $\{x_i, \bar{x}_i\}$ and the votes $\{x_i, d_{x_i}\}$ and $\{\bar{x}_i, d_{\bar{x}_i}\}$. For every $i \in [t]$, there are the votes $\{c_i, l_1\}, \{c_i, l_2\}, \{c_i, l_3\}$ where $\{l_1, l_2, l_3\} = \text{lit}(c_i)$, two copies of the vote $\{c_i, s_i\}$, and a vote $\{s_i, d_{s_i}\}$. The reader can inspect Figure 7.1 for a graph representation of the profile.

Claim 85. *If there is no assignment to the variables $x_1, \dots, x_r$ satisfying ϕ , then there is some $j^* \in [t]$ such that alternative s_{j^*} is included in the winning committee before alternative c_{j^*} .*

Proof. Observe that under the sequential CC rule, for every $i \in [r]$, either x_i or $\bar{x}_i$ (exclusively) is included in the winning committee in the first r rounds. Let C_r be the winning committee after these rounds. Since there is no assignment to the variables that satisfies ϕ , there must now be a clause alternative c_{j^*} for which all literal alternatives corresponding to literals in $\text{lit}(c_{j^*})$ are already included in the winning committee. Otherwise, the ordering in which the r literal alternatives have been included in the winning committee would correspond to an assignment α that satisfies ϕ . That is, we could pick α such that

$$\alpha(l) = \begin{cases} 0 & \text{if } l \in C_r, \\ 1 & \text{otherwise,} \end{cases}$$

for every $l \in \{x_1, \bar{x}_1, \dots, x_r, \bar{x}_r\}$.

In order for A to become the winning committee, there must be a step i where c_{j^*} is included in the winning committee. However, observe that

$$\Delta_{f_{CC}}(C_i, s_{j^*}, P) = 3 > 2 = \Delta_{f_{CC}}(C_i, c_{j^*}, P),$$

where $\Delta_{f_{CC}}(C_i, a, P)$ denotes the *score increase* when including alternative a in the subcommittee C_i on a profile P . Thus, alternative s_{j^*} is included in the winning committee before the alternative c_{j^*} . $\square$

So far, we have shown that committee A is not winning on profile P under the sequential CC rule if there exists no assignment that satisfies ϕ . Now, assume that there is an assignment α to the variables $x_1, \dots, x_r$ that satisfies ϕ . We rename the literals according to their values under the assignment α such that

$$l_i^1 = \begin{cases} x_i & \text{if } \alpha(x_i) = 1 \\ \bar{x}_i & \text{otherwise,} \end{cases}$$

and

$$l_i^0 = \begin{cases} x_i & \text{if } \alpha(x_i) = 0 \\ \bar{x}_i & \text{otherwise.} \end{cases}$$

We conclude the proof by showing that there is an ordering of the alternatives in A such that the sequential CC rule returns A as the winning committee. The feasibility of the following sequence of inclusions can easily be verified.

First, the literal alternatives $l_1^0, \dots, l_r^0$ are added to the committee. Conversely, none of the alternatives $l_1^1, \dots, l_r^1$ have been included in the winning committee yet. Since α is a satisfying assignment, each clause alternative c_i for $i \in [t]$ is included in at least one vote $\{l, c_i\}$ where $l \in \{l_1^1, \dots, l_r^1\}$. Thus, every clause alternative c_i has a score increase of at least 3. This score increase is at least as high as the score increase

of the corresponding special alternative s_i and of any remaining literal alternative. This implies that there is an ordering of the clause alternatives such that each of these alternatives is included in the winning committee. At this point, all the remaining alternatives have a score increase of at most 1. All alternatives from A that have not yet been selected for the winning committee have a score increase of exactly 1. Hence, the alternatives $l_1^1, \dots, l_r^1$ can be included in the winning committee as well. Thereby, A is indeed winning under rule f on profile P . $\square$

Bibliography

- [1] Ben Abramowitz and Elliot Anshelevich. Utilitarians without utilities: Maximizing social welfare for graph problems using only ordinal preferences. In *Proceedings of the 32nd AAAI Conference on Artificial Intelligence (AAAI)*, pages 894–901, 2018. 82
- [2] Sara Ahmadian, Ashkan Norouzi-Fard, Ola Svensson, and Justin Ward. Better guarantees for k-means and euclidean k-median by primal-dual algorithms. *SIAM Journal on Computing*, 49(4), 2020. 82
- [3] Noga Alon. Tools from higher algebra. In *Handbook of Combinatorics (Vol. 2)*, pages 1749–1783. MIT Press, 1996. 34, 100, 101, 105
- [4] Georgios Amanatidis, Georgios Birmpas, Aris Filos-Ratsikas, and Alexandros A. Voudouris. Peeking behind the ordinal curtain: Improving distortion via cardinal queries. *Artificial Intelligence*, 296:103488, 2021. 5, 8, 9, 12, 42, 44, 62, 65, 69, 82
- [5] Georgios Amanatidis, Georgios Birmpas, Aris Filos-Ratsikas, and Alexandros A. Voudouris. Don’t roll the dice, ask twice: The two-query distortion of matching problems and beyond. In *Proceedings of the 36th Annual Conference on Neural Information Processing Systems (NeurIPS)*, 2022. 9, 45, 62, 81, 82
- [6] Georgios Amanatidis, Georgios Birmpas, Aris Filos-Ratsikas, and Alexandros A. Voudouris. A few queries go a long way: Information-distortion tradeoffs in matching. *Journal of Artificial Intelligence Research*, 74:227–261, 2022. 45, 81, 82
- [7] Haris Angelidakis, Konstantin Makarychev, and Yury Makarychev. Algorithms for stable and perturbation-resilient problems. In *Proceedings of the 49th Annual ACM Symposium on Theory of Computing (STOC)*, pages 438–451, 2017. 82
- [8] Elliot Anshelevich and Wennan Zhu. Tradeoffs between information and ordinal approximation for bipartite matching. In *Proceedings of the 10th International Symposium on Algorithmic Game Theory (SAGT)*, pages 267–279, 2017. 80, 81

- [9] Elliot Anshelevich and Wennan Zhu. Ordinal approximation for social choice, matching, and facility location problems given candidate positions. *ACM Transactions on Economics and Computation*, 9(2):art. 9, 2021. 95
- [10] Elliot Anshelevich, Onkar Bhardwaj, Edith Elkind, John Postl, and Piotr Skowron. Approximating optimal social choice under metric preferences. *Artificial Intelligence*, 264:27–51, 2018. 18, 27, 45, 80
- [11] Elliot Anshelevich, Aris Filos-Ratsikas, Nisarg Shah, and Alexandros A. Voudouris. Distortion in social choice problems: The first 15 years and beyond. In *Proceedings of the 30th International Joint Conference on Artificial Intelligence (IJCAI)*, pages 4294–4301, 2021. 8, 42, 45, 80
- [12] Elliot Anshelevich, Aris Filos-Ratsikas, and Alexandros A. Voudouris. The distortion of distributed metric social choice. *Artificial Intelligence*, 308: 103713, 2022. 80
- [13] Jose Apesteguia, Miguel A Ballester, and Rosa Ferrer. On the justice of decision rules. *The Review of Economic Studies*, 78(1):1–16, 2011. 44
- [14] Kenneth J. Arrow. A difficulty in the concept of social welfare. *Journal of Political Economy*, 58(4):328–346, 1950. 3
- [15] Kenneth J. Arrow. *Social Choice and Individual Values*. Wiley, 1951. 98
- [16] Vijay Arya, Naveen Garg, Rohit Khandekar, Adam Meyerson, Kamesh Mungala, and Vinayaka Pandit. Local search heuristics for k-median and facility location problems. *SIAM Journal on Computing*, 33(3):544–562, 2004. 82
- [17] Pranjal Awasthi, Avrim Blum, and Or Sheffet. Stability yields a PTAS for k-median and k-means clustering. In *Proceedings of the 51th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*, pages 309–318, 2010. 82
- [18] Haris Aziz, Serge Gaspers, Joachim Gudmundsson, Simon Mackenzie, Nicholas Mattei, and Toby Walsh. Computational aspects of multi-winner approval voting. In *Proceedings of the 14th International Conference on Autonomous Agents and Multiagent Systems*, page 107–115, 2015. 32
- [19] Dorothea Baumeister and Tobias Högrebe. How hard is the manipulative design of scoring systems? In *Proceedings of the 28th International Joint Conference on Artificial Intelligence (IJCAI)*, pages 74–80, 2019. 101
- [20] Gerdus Benadè, Swaprava Nath, Ariel D. Procaccia, and Nisarg Shah. Preference elicitation for participatory budgeting. *Management Science*, 67(5): 2813–2827, 2021. 45, 60

- [21] Jeremy Bentham. *An introduction to the principles of morals and legislation*. Clarendon Press, 1907. 3, 44
- [22] Nadja Betzler, Arkadii Slinko, and Johannes Uhlmann. On the computation of fully proportional representation. *Journal of Artificial Intelligence Research*, 47:475–519, 2013. 101
- [23] Niclas Boehmer, Markus Brill, Alfonso Cevallos, Jonas Gehrlein, Luis Sánchez-Fernández, and Ulrike Schmidt-Kraepelin. Approval-based committee voting in practice: A case study of (over-)representation in the polkadot blockchain. In *Proceedings of the 38th AAAI Conference on Artificial Intelligence (AAAI)*, 9, pages 9519–9527, 2024. 35
- [24] Allan Borodin, Omer Lev, Nisarg Shah, and Tyrone Strangway. Primarily about primaries. In *Proceedings of the 33rd AAAI Conference on Artificial Intelligence (AAAI)*, pages 1804–1811, 2019. 82
- [25] Craig Boutilier, Ioannis Caragiannis, Simi Haber, Tyler Lu, Ariel D. Procaccia, and Or Sheffet. Optimal social choice functions: A utilitarian view. *Artificial Intelligence*, 227:190–213, 2015. 9, 10, 11, 33, 34, 36, 37, 42, 44, 48, 101
- [26] Steven J. Brams and Peter C. Fishburn. Going from theory to practice: The mixed success of approval voting. *Social Choice and Welfare*, 25:457–474, 2005. 29
- [27] Felix Brandt, Vincent Conitzer, Ulle Endress, Jérôme Lang, and Ariel D. Procaccia, editors. *Handbook of Computational Social Choice*. Cambridge University Press, 2016. 3, 97
- [28] Jakob Burkhardt, Ioannis Caragiannis, Karl Fehrs, Matteo Russo, Chris Schwiegelshohn, and Sudarshan Shyam. Low-distortion clustering with ordinal and limited cardinal information. In *Proceedings of the 38th AAAI Conference on Artificial Intelligence (AAAI)*, 9, pages 9555–9563, 2024. 5, 45, 81, 96
- [29] Jaroslaw Byrka, Thomas W. Pensyl, Bartosz Rybicki, Aravind Srinivasan, and Khoa Trinh. An improved approximation for k -median and positive correlation in budgeted optimization. *ACM Trans. Algorithms*, 13(2):23:1–23:31, 2017. doi: 10.1145/2981561. URL <https://doi.org/10.1145/2981561>. 26
- [30] Jaroslaw Byrka, Krzysztof Sornat, and Joachim Spoerhase. Constant-factor approximation for ordered k -median. In *Proceedings of the 50th Annual ACM Symposium on Theory of Computing (STOC)*, pages 620–631, 2018. 96
- [31] Ioannis Caragiannis and Karl Fehrs. The complexity of learning approval-based multiwinner voting rules. In *Proceedings of the 36th AAAI Conference on Artificial Intelligence (AAAI)*, 5, pages 4925–4932, 2022. 5, 104

- [32] Ioannis Caragiannis and Karl Fehrs. Beyond the worst case: Distortion in impartial culture electorates. In *Proceedings of the 20th Conference on Web and Internet Economics (WINE)*, 2024. Forthcoming. 5
- [33] Ioannis Caragiannis and Ariel D. Procaccia. Voting almost maximizes social welfare despite limited communication. *Artificial Intelligence*, 175(9-10): 1655–1671, 2011. 4, 27, 42, 82
- [34] Ioannis Caragiannis, George A. Krimpas, and Alexandros A. Voudouris. Aggregating partial rankings with applications to peer grading in massive online open courses. In *Proceedings of the 14th International Joint Conference on Autonomous Agents and Multiagent Systems (AAMAS)*, pages 675–683, 2015. 101
- [35] Ioannis Caragiannis, Swaprava Nath, Ariel D. Procaccia, and Nisarg Shah. Subset selection via implicit utilitarian voting. *Journal of Artificial Intelligence Research*, 58:123–152, 2017. 45
- [36] Ioannis Caragiannis, Xenophon Chatzigeorgiou, George A. Krimpas, and Alexandros A. Voudouris. Optimizing positional scoring rules for rank aggregation. *Artificial Intelligence*, 267:58–77, 2019. 101
- [37] Ioannis Caragiannis, George A. Krimpas, and Alexandros A. Voudouris. How effective can simple ordinal peer grading be? *ACM Transactions on Economics and Computation*, 8(3), 2020. 101
- [38] Ioannis Caragiannis, Nisarg Shah, and Alexandros A. Voudouris. The metric distortion of multiwinner voting. *Artificial Intelligence*, 313:103802, 2022. 19, 20, 21, 45, 80
- [39] Ioannis Caragiannis, Evi Micha, and Jannik Peters. Can a few decide for many? The metric distortion of sortition. In *Proceedings of the 41st International Conference on Machine Learning (ICML)*, 2025. 22
- [40] Deeparnab Chakrabarty and Chaitanya Swamy. Interpolating between k-median and k-center: Approximation algorithms for ordered k-median. In *Proceedings of the 45th International Colloquium on Automata, Languages, and Programming (ICALP)*, pages 29:1–29:14, 2018. 26, 96
- [41] John R. Chamberlin and Paul N. Courant. Representative deliberations and representative decisions: Proportional representation and the Borda rule. *American Political Science Review*, 77(3):718–733, 1983. 100
- [42] Moses Charikar and Prasanna Ramakrishnan. Metric distortion bounds for randomized social choice. In *Proceedings of the 33rd ACM-SIAM Symposium on Discrete Algorithms (SODA)*, pages 2986–3004, 2022. 45, 80

- [43] Moses Charikar, Prasanna Ramakrishnan, Kangning Wang, and Hongxun Wu. Breaking the metric voting distortion barrier. *Journal of the ACM*, 71(6), 2024. 19, 45, 80
- [44] Ning Chen, Nicole Immorlica, Anna R. Karlin, Mohammad Mahdian, and Atri Rudra. Approximating matches made in heaven. In *Proceedings of the 36th International Colloquium on Automata, Languages and Programming (ICALP), Part I*, volume 5555 of *Lecture Notes in Computer Science*, pages 266–278. Springer, 2009. 45
- [45] Yu Cheng, Shaddin Dughmi, and David Kempe. Of the people: Voting is more effective with representative candidates. In *Proceedings of the 2017 ACM Conference on Economics and Computation (EC)*, pages 305–322, 2017. 45, 82
- [46] Yu Cheng, Shaddin Dughmi, and David Kempe. On the distortion of voting with multiple representative candidates. *Proceedings of the AAAI Conference on Artificial Intelligence (AAAI)*, 2018. 45, 82
- [47] Vincent Cohen-Addad and Chris Schwiegelshohn. On the local structure of stable clustering instances. In *Proceedings of the 58th IEEE Annual Symposium on Foundations of Computer Science (FOCS)*, pages 49–60, 2017. 82
- [48] Vincent Cohen-Addad, Philip N. Klein, and Claire Mathieu. Local search yields approximation schemes for k-means and k-median in euclidean and minor-free metrics. *SIAM Journal on Computing*, 48(2):644–667, 2019. 82
- [49] Vincent Cohen-Addad, Karthik C. S., and Euiwoong Lee. On approximability of clustering problems without candidate centers. In *Proceedings of the 32nd ACM-SIAM Symposium on Discrete Algorithms (SODA)*, pages 2635–2648, 2021. 82
- [50] Vincent Cohen-Addad, Andreas Emil Feldmann, and David Saulpic. Near-linear time approximation schemes for clustering in doubling metrics. *Journal of the ACM*, 68(6):44:1–44:34, 2021. 82
- [51] Vincent Cohen-Addad, David Saulpic, and Chris Schwiegelshohn. Improved coresets and sublinear algorithms for power means in euclidean spaces. In *Proceedings of the 34th Annual Conference on Neural Information Processing Systems (NeurIPS)*, pages 21085–21098, 2021. 82
- [52] Vincent Cohen-Addad, Anupam Gupta, Lunjia Hu, Hoon Oh, and David Saulpic. An improved local search algorithm for k-median. In *Proceedings of the 33rd ACM-SIAM Symposium on Discrete Algorithms (SODA)*, pages 1556–1612, 2022. 82

- [53] Vincent Cohen-Addad, Fabrizio Grandoni, Euiwoong Lee, and Chris Schwiegelshohn. Breaching the 2 LMP approximation barrier for facility location with applications to k -median. In *Proceedings of the 34th ACM-SIAM Symposium on Discrete Algorithms (SODA)*, pages 940–986, 2023. 82
- [54] Marek Cygan, Fedor V. Fomin, Lukasz Kowalik, Daniel Lokshtanov, Daniel Marx, Marcin Pilipczuk, Michał Pilipczuk, and Saket Saurabh. *Parameterized Algorithms*. Springer, 2015. 30, 100, 110
- [55] Amit Daniely and Shai Shalev-Shwartz. Optimal learners for multiclass problems. In *Proceedings of The 27th Conference on Learning Theory*, volume 35, pages 287–316, 2014. 35
- [56] Amit Daniely, Sivan Sabato, Shai Ben-David, and Shai Shalev-Shwartz. Multiclass learnability and the ERM principle. *Journal of Machine Learning Research*, 16:2377–2404, 2015. 34, 104
- [57] Soroush Ebadian and Evi Micha. Boosting sortition via proportional representation. In *Proceedings of 24th International Conference on Autonomous Agents and Multi-Agent Systems (AAMAS)*, 2025. Forthcoming. 22
- [58] Soroush Ebadian and Nisarg Shah. Every bit helps: Achieving the optimal distortion with a few queries. In *Proceedings of the 39th AAAI Conference on Artificial Intelligence (AAAI)*, 2025. Forthcoming. 14, 15
- [59] Soroush Ebadian, Anson Kahng, Dominik Peters, and Nisarg Shah. Optimized distortion and proportional fairness in voting. In *Proceedings of the 23rd ACM Conference on Economics and Computation (EC)*, pages 563–600, 2022. 42
- [60] Soroush Ebadian, Gregory Kehne, Evi Micha, Ariel D Procaccia, and Nisarg Shah. Is sortition both representative and fair? In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 35, pages 3431–3443, 2022. 22, 26
- [61] Roy Fairstein, Dan Vilenchik, and Kobi Gal. Learning aggregation rules in participatory budgeting: A data-driven approach, 2024. URL <https://arxiv.org/abs/2412.01864>. 35
- [62] Piotr Faliszewski and Nimrod Talmon. Between proportionality and diversity: Balancing district sizes under the Chamberlin-Courant rule. In *Proceedings of the 17th International Conference on Autonomous Agents and MultiAgent Systems (AAMAS)*, pages 14–22, 2018. 98
- [63] Piotr Faliszewski, Piotr Skowron, Arkadii Slinko, and Nimrod Talmon. Multi-winner rules on paths from k -Borda to Chamberlin-Courant. In *Proceedings of the 26th International Joint Conference on Artificial Intelligence (IJCAI)*, pages 192–198, 2017. 98

- [64] Piotr Faliszewski, Piotr Skowron, Arkadii M. Slinko, and Nimrod Talmon. Multiwinner voting: A new challenge for social choice theory. In Ulle Endriss, editor, *Trends in Computational Social Choice*, pages 27–47. AI Access, 2017. 98
- [65] Piotr Faliszewski, Piotr Skowron, Arkadii Slinko, and Nimrod Talmon. Committee scoring rules: Axiomatic characterization and hierarchy. *ACM Transactions on Economics and Computation*, 7(1):3:1–3:39, 2019. 30, 36, 122
- [66] Piotr Faliszewski, Stanislaw Szufa, and Nimrod Talmon. *Optimization-Based Voting Rule Design: The Closer to Utopia the Better*, pages 17–51. Springer International Publishing, 2022. 36, 101
- [67] Michal Feldman, Amos Fiat, and Iddan Golomb. On voting and facility location. In *Proceedings of the 17th ACM Conference on Economics and Computation (EC)*, pages 269–286, 2016. 80
- [68] Aris Filos-Ratsikas, Evi Micha, and Alexandros A. Voudouris. The distortion of distributed voting. *Artificial Intelligence*, 286:103343, 2020. 45, 82
- [69] Bailey Flanigan, Paul Gözl, Anupam Gupta, Brett Hennig, and Ariel Procaccia. Fair algorithms for selecting citizens’ assemblies. *Nature*, 596:548–552, 2021. 22
- [70] Dimitris Fotakis. On the competitive ratio for online facility location. *Algorithmica*, 50(1):1–57, 2008. 94
- [71] Zachary Friggstad, Mohsen Rezapour, and Mohammad R. Salavatipour. Local search yields a PTAS for k-means in doubling metrics. *SIAM Journal on Computing*, 48(2):452–480, 2019. 82
- [72] Mohammad Ghodsi, Mohamad Latifian, and Masoud Seddighin. On the distortion value of the elections with abstention. In *Proceedings of the 33rd AAAI Conference on Artificial Intelligence (AAAI)*, pages 1981–1988, 2019. 45
- [73] Allan Gibbard. Manipulation of voting schemes: A general result. *Econometrica*, 41(4):587–601, 1973. 3
- [74] Vasilis Gkatzelis, Daniel Halpern, and Nisarg Shah. Resolving the optimal metric distortion conjecture. In *Proceedings of the 61st IEEE Annual Symposium on Foundations of Computer Science (FOCS)*, pages 1427–1438, 2020. 18, 19, 45, 80
- [75] Ashish Goel, Anilesh Kollagunta Krishnaswamy, and Kamesh Munagala. Metric distortion of social choice rules: Lower bounds and fairness properties. In *Proceedings of the 18th ACM Conference on Economics and Computation (EC)*, pages 287–304, 2017. 80

- [76] Ashish Goel, Anilesh K. Krishnaswamy, Sukolsak Sakshuwong, and Tanja Aitamurto. Knapsack voting for participatory budgeting. *ACM Transactions on Economics and Computation (TEAC)*, 7(2):1–27, 2019. 29
- [77] Yannai A. Gonczarowski, Gregory Kehne, Ariel D. Procaccia, Ben Schiffer, and Shirley Zhang. The distortion of binomial voting defies expectation. In *Proceedings of the 37th International Conference on Neural Information Processing Systems (NeurIPS)*, 2024. 10, 44
- [78] Teofilo F. Gonzalez. Clustering to minimize the maximum intercluster distance. *Theoretical Computer Science*, 38:293–306, 1985. 21, 24, 82, 84
- [79] Kishen N. Gowda, Thomas W. Pensyl, Aravind Srinivasan, and Khoa Trinh. Improved bi-point rounding algorithms and a golden barrier for k -median. In *Proceedings of the 34th ACM-SIAM Symposium on Discrete Algorithms (SODA)*, pages 987–1011, 2023. 21, 82
- [80] Mohak Goyal and Sahasrajit Sarma Sarkar. Metric distortion under probabilistic voting, 2025. URL <https://arxiv.org/abs/2405.14223>. 26, 27
- [81] Sudipto Guha and Samir Khuller. Greedy strikes back: Improved facility location algorithms. *Journal of Algorithms*, 31(1):228–248, 1999. 20
- [82] Kamal Jain and Vijay V. Vazirani. Approximation algorithms for metric facility location and k -median problems using the primal-dual schema and Lagrangian relaxation. *Journal of the ACM*, 48(2):274–296, 2001. 82
- [83] Kamal Jain, Mohammad Mahdian, and Amin Saberi. A new greedy approach for facility location problems. In *Proceedings on 34th Annual ACM Symposium on Theory of Computing (STOC)*, pages 731–740, 2002. 82
- [84] Michał Jaworski and Piotr Skowron. Phragmén rules for degressive and regressive proportionality. In *Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence (IJCAI)*, pages 328–334, 2022. 98
- [85] Tushant Jha and Yair Zick. A learning framework for distribution-based game-theoretic solution concepts. In *Proceedings of the 21st ACM Conference on Economics and Computation (EC)*, pages 355–377, 2020. 101
- [86] Haim Kaplan, David Naori, and Danny Raz. Almost tight bounds for online facility location in the random-order model. In *Proceedings of the 34th ACM-SIAM Symposium on Discrete Algorithms (SODA)*, pages 1523–1544, 2023. 94
- [87] David Kempe. An analysis framework for metric voting based on LP duality. In *Proceedings of the 34th AAAI Conference on Artificial Intelligence (AAAI)*, pages 2079–2086, 2020. 80

- [88] Fatih Erdem Kizilkaya and David Kempe. Plurality veto: A simple voting rule achieving optimal metric distortion. In *Proceedings of the 31st International Joint Conference on Artificial Intelligence (IJCAI)*, pages 349–355, 2022. 19, 45
- [89] Nirman Kumar and Benjamin Raichel. Fault tolerant clustering revisited. In *Proceedings of the 25th Canadian Conference on Computational Geometry (CCCG)*, 2013. 21
- [90] Martin Lackner and Piotr Skowron. Approval-based multi-winner rules and strategic voting. In *Proceedings of the 27th International Joint Conference on Artificial Intelligence (IJCAI)*, pages 340–346, 2018. 30
- [91] Martin Lackner and Piotr Skowron. Utilitarian welfare and representation guarantees of approval-based multiwinner rules. *Artificial Intelligence*, 288: 103366, 2020. 98
- [92] Martin Lackner and Piotr Skowron. Consistent approval-based multi-winner rules. *Journal of Economic Theory*, 192:105173, 2021. 29, 98, 101, 102
- [93] Martin Lackner and Piotr Skowron. *Multi-Winner Voting with Approval Preferences*. SpringerBriefs in Intelligent Systems. Springer International Publishing, 2023. 30, 98, 101
- [94] Jean-Francois Laslier and M. Remzi Sanver, editors. *Handbook on Approval Voting*. Springer, 2010. 29, 32, 98
- [95] Mohamad Latifian and Alexandros A. Voudouris. The distortion of threshold approval matching, 2024. 45
- [96] Shi Li. A 1.488 approximation algorithm for the uncapacitated facility location problem. *Information and Computation*, 222:45–58, 2013. 82
- [97] Shi Li and Ola Svensson. Approximating k-median via pseudo-approximation. *SIAM Journal on Computing*, 45(2):530–547, 2016. 82
- [98] Thomas Ma, Vijay Menon, and Kate Larson. Improving welfare in one-sided matchings using simple threshold queries. In *Proceedings of the Thirtieth International Joint Conference on Artificial Intelligence (IJCAI)*, pages 321–327, 2021. 45
- [99] Pasin Manurangsi and Warut Suksompong. When do envy-free allocations exist? *SIAM Journal on Discrete Mathematics*, 34(3):1505–1521, 2020. 45
- [100] Samuel Merrill III. *A Unified Theory of Voting: Directional and Proximity Spatial Models*. Cambridge University Press, 1999. 17

- [101] Adam Meyerson. Online facility location. In *Proceedings of the 42nd IEEE Symposium on Foundations of Computer Science (FOCS)*, pages 426–431, 2001. 81, 94
- [102] Rajeev Motwani and Prabhakar Raghavan. *Randomized Algorithms*. Cambridge University Press, 1995. 49
- [103] Kamesh Munagala and Kangning Wang. Improved metric distortion for deterministic social choice rules. In *Proceedings of the 2019 ACM Conference on Economics and Computation (EC)*, pages 245–262, 2019. 18, 80
- [104] Roger B. Myerson. Optimal auction design. *Mathematics of Operations Research*, 6(1):58–73, 1981. 45
- [105] Balas K. Natarajan. On learning sets and functions. *Machine Language*, 4(1): 67–97, 1989. 32, 104
- [106] Balas K. Natarajan. *Machine Learning: A Theoretical Approach*. Morgan Kaufmann Publishers Inc., 1991. 33
- [107] Athanasios Papoulis and S. Unnikrishna Pillai. *Probability, Random Variables and Stochastic Processes*. McGraw-Hill Higher Education, 2002. 59
- [108] Grzegorz Pierczynski and Piotr Skowron. Approval-based elections and distortion of voting rules. In *Proceedings of the 28th International Joint Conference on Artificial Intelligence (IJCAI)*, pages 543–549, 2019. 36, 82
- [109] Marcus Pivato. Asymptotic utilitarianism in scoring rules. *Social Choice & Welfare*, 47(2):431–458, 2016. 44
- [110] Geoffrey Pritchard and Mark C. Wilson. Asymptotics of the minimum manipulating coalition size for positional voting rules under impartial culture behaviour. *Mathematical Social Sciences*, 58(1):35–57, 2009. 43
- [111] Ariel D. Procaccia and Jeffrey S. Rosenschein. The distortion of cardinal preferences in voting. In *Proceedings of the 10th International Workshop on Cooperative Information Agents (CIA)*, volume 4149 of *Lecture Notes in Computer Science*, pages 317–331. Springer, 2006. 4, 7, 42, 81, 82
- [112] Ariel D. Procaccia, Jeffrey S. Rosenschein, and Aviv Zohar. On the complexity of achieving proportional representation. *Social Choice and Welfare*, 30:353–362, April 2008. 32, 101
- [113] Ariel D. Procaccia, Aviv Zohar, Yoni Peleg, and Jeffrey S. Rosenschein. The learnability of voting rules. *Artificial Intelligence*, 173(12):1133–1149, 2009. 30, 33, 34, 99, 100
- [114] Haripriya Pulyassary. Algorithm design for ordinal settings. Master’s thesis, University of Waterloo, 2022. 81, 94

- [115] Tim Roughgarden, editor. *Beyond the Worst-Case Analysis of Algorithms*. Cambridge University Press, 2020. 43
- [116] Mark Allen Satterthwaite. Strategy-proofness and Arrow’s conditions: Existence and correspondence theorems for voting procedures and social welfare functions. *Journal of Economic Theory*, 10(2):187–217, 1975. 3
- [117] Shai Shalev-Shwartz and Shai Ben-David. *Understanding Machine Learning: From Theory to Algorithms*. Cambridge University Press, 2014. 104
- [118] Piotr Skowron and Piotr Faliszewski. Chamberlin-courant rule with approval ballots: Approximating the maxcover problem with bounded frequencies in FPT time. *Journal of Artificial Intelligence Research*, 60(1):687–716, 2017. 32
- [119] Chinmay Sonar, Palash Dey, and Neeldhara Misra. On the complexity of winner verification and candidate winner for multiwinner voting rules. In *Proceedings of the 29th International Joint Conference on Artificial Intelligence (IJCAI)*, pages 89–95, 2020. 32, 100, 101, 123
- [120] Thorvald N. Thiele. On multiple choice (in Danish). *Bulletin of the Royal Danish Academy of Sciences and Letters*, pp. 415–441, 1895. 30, 100
- [121] Panos Toulis and David C. Parkes. Design and analysis of multi-hospital kidney exchange mechanisms using random graphs. *Games & Economic Behaviour*, 91:360–382, 2015. 45
- [122] Ilia Tsetlin, Michel Regenwetter, and Bernard Grofman. The impartial culture maximizes the probability of majority cycles. *Social Choice & Welfare*, 21(3): 387–398, 2003. 43
- [123] V. Vapnik and A. Chervonenkis. On the uniform convergence of relative frequencies of events to their probabilities. *Theory of Probability and Its Applications*, 16(2):264–280, 1971. 32
- [124] Hugh E. Warren. Lower bounds for approximation by nonlinear manifolds. *Transactions of the American Mathematical Society*, 133(1):167–178, 1968. 34, 100, 105
- [125] Lirong Xia. Designing social choice mechanisms using machine learning. In *Proceedings of the 12th International Conference on Autonomous Agents and Multi-Agent Systems (AAMAS)*, pages 471–474, 2013. 101
- [126] Ryo Yoshinaka. Higher-order matching in the linear λ calculus in the absence of constants is np-complete. In *Proceedings of the 16th International Conference on Term Rewriting and Applications*, pages 235–249, 2005. 117
- [127] Peyton Young. Optimal voting rules. *The Journal of Economic Perspectives*, 9 (1):51–64, 1995. 3, 42

- [128] William S. Zwicker. Introduction to the theory of voting. In Felix Brandt, Vincent Conitzer, Ulle Endriss, Jérôme Lang, and Ariel D. Procaccia, editors, *Handbook of Computational Social Choice*, pages 23–56. Cambridge University Press, 2016. 41