

Random Dictators with a Random Referee: Constant Sample Complexity Mechanisms for Social Choice

Brandon Fain
Duke University

Ashish Goel
Stanford University

Kamesh Munagala
Duke University

Nina Prabhu
Stanford University

Abstract

We study social choice mechanisms in an implicit utilitarian framework with a metric constraint, where the goal is to minimize *Distortion*, the worst case social cost of an ordinal mechanism relative to underlying cardinal utilities. We consider two additional desiderata: *Constant sample complexity* and *Squared Distortion*. Constant sample complexity means that the mechanism (potentially randomized) only uses a constant number of ordinal queries regardless of the number of voters and alternatives. Squared Distortion is a measure of variance of the Distortion of a randomized mechanism.

Our primary contribution is the first social choice mechanism with constant sample complexity *and* constant Squared Distortion (which also implies constant Distortion). We call the mechanism *Random Referee*, because it uses a random agent to compare two alternatives that are the favorites of two other random agents. We prove that the use of a comparison query is necessary: no mechanism that only elicits the top- k preferred alternatives of voters (for constant k) can have Squared Distortion that is sublinear in the number of alternatives. We also prove that unlike any top- k only mechanism, the Distortion of Random Referee meaningfully improves on benign metric spaces, using the Euclidean plane as a canonical example. Finally, among top-1 only mechanisms, we introduce *Random Oligarchy*. The mechanism asks just 3 queries and is essentially optimal among the class of such mechanisms with respect to Distortion.

In summary, we demonstrate the surprising power of constant sample complexity mechanisms generally, and just three random voters in particular, to provide some of the best known results in the implicit utilitarian framework.

1 Introduction

Consider the social choice problem of deciding on an allocation of public tax dollars to public projects. This is a voting problem over budgets. Clearly, the number of voters in such situations can be large. More interestingly, unlike in traditional social choice theory, there is no reason to believe that the number of alternatives (budgets) is small. It is therefore unreasonable to assume that we can elicit full ordinal preferences over alternatives from every agent. For a voting mechanism to be practical in such a setting, one would ideally like it to require only an absolute constant number of

simple queries, regardless of the number of voters and alternatives. We call this property *constant sample complexity*, and we explore mechanisms of this sort in this paper.

We define our model more formally in Section 4, but at a high level, we have a set N of agents (or voters) and a set of alternatives S , from which we must choose a single outcome. We assume that N and S are both large, and that eliciting the full ordinal rankings may be prohibitively difficult. Instead, we work with an ordinal query model, and a constant sample complexity mechanism uses only a constant number of these queries.

Top- k Query. “What are your k favorite alternatives, in order?” (We call a top-1 query a **favorite** query); and

Comparison Query. “Which of two given alternatives do you prefer?”

Query models are not just of theoretical interest. They can be used to reduce cognitive overload in voting. For example, in the context of Participatory Budgeting (Goel et al. 2015), the space of possible budget allocations is large, and one mechanism is to ask voters to compare two proposed budgets. Similarly, in a context like transportation policy for a city, a single alternative can be an entire transportation plan. In such examples, not only are there many alternatives, but it may be infeasible to expect voters to compare more than two alternatives at the same time. We stress that constant sample complexity is particularly important in settings where there may be a large number of possibly complex alternatives.

To evaluate the quality of our mechanisms, we adopt the implicit utilitarian perspective with metric constraints (Boutilier et al. 2015; Cheng, Dughmi, and Kempe 2017; Anshelevich and Postl 2017; Goel, Krishnaswamy, and Munagala 2017; Anshelevich et al. 2018; Feldman, Fiat, and Golomb 2016). That is, we assume that agents have cardinal costs over alternatives, and these costs are constrained to be metric, but asking agents to work with or report cardinal costs is impractical or impossible. We want to design social choice mechanisms to minimize the total social cost, but our mechanism is constrained only to use ordinal queries, that is, those that can be answered given a total order over alternatives. We therefore measure the efficiency of a mechanism as its *Distortion* (see Section 4), the worst case approximation to the total social cost.

2 Results

The starting point for our inquiry is the constant sample complexity *Random Dictatorship* mechanism. The algorithm asks a single favorite query from an agent chosen uniformly at random, and has a tight Distortion bound of 3 (Anshelevich and Postl 2017). In this paper, we provide two new mechanisms (Random Referee and Random Oligarchy) that improve on this simple baseline in three different ways, outlined in each of our three technical sections. We hope that our work inspires future research on similarly lightweight mechanisms for social choice in large decision spaces. Any omitted proofs for our results can be found in the full version of this paper (Fain et al. 2018).

Random Referee: Comparison Queries and Squared Distortion. In Section 5, we show that one disadvantage of Random Dictatorship lies not in its Distortion, but in its variance. For randomized mechanisms, Distortion is measured as the expected approximation to the first moment of social cost. However, in many social choice problems, we might want a bound on the risk associated with a given mechanism. We capture this via *Squared Distortion*, as suggested in (Fain et al. 2017). The Squared Distortion (see Definition 2 in Section 4) is the expected approximation to the second moment of social cost. A mechanism with constant Squared Distortion has both constant Distortion and constant coefficient of variation of the Distortion.

We show that mechanisms using only top- k queries (including Random Dictatorship) have Squared Distortion $\Omega(|\mathcal{S}|)$. This motivates us to expand our query model to incorporate information about the relative preferences of agents between alternatives, *i.e.* use comparison queries. We define a novel mechanism called *Random Referee* (RR) that uses a random voter as a referee to compare the favorite alternatives of two other random voters (see Definition 3 in Section 5). Our main result in Section 5 is Theorem 2: The Squared Distortion of RR is at most 21. This also immediately implies that the Distortion of RR is at most 4.583.

Random Referee: Euclidean Plane. In Section 6, we show that top- k only mechanisms (including Random Dictatorship) achieve their worst case Distortion even on benign metrics like low dimensional Euclidean spaces. We analyze a special case on the Euclidean plane and prove that the Distortion of Random Referee beats that of any top- k only mechanism. While the improvement is quantitatively small, it is qualitatively interesting: we demonstrate that by using a *single* comparison query, Random Referee can exploit the structure of the metric space to improve Distortion, whereas Random Dictatorship or any other top- k only mechanism cannot. We conjecture this result extends to Euclidean spaces in any dimension, and present some evidence to support this conjecture in Section 8.

Random Oligarchy: Favorite Only Mechanisms. In Section 7, we consider mechanisms that are restricted to favorite queries and show that constant complexity mechanisms are nearly optimal. We present a mechanism that

uses only three favorite queries that has Distortion at most 3 for arbitrary $|\mathcal{S}|$; however, it also has Distortion for small $|\mathcal{S}|$ that improves upon the best known favorite only mechanism from (Gross, Anshelevich, and Xia 2017) that uses at most $\min(|N| + 1, |\mathcal{S}| + 1)$ favorite queries. Comparing with a lower bound for favorite only mechanisms, Random Oligarchy has nearly optimal distortion and constant sample complexity. Though this mechanism does not have constant Squared Distortion like Random Referee, we present it to demonstrate again the surprising power of constant sample complexity randomized social choice mechanisms in general, and of queries to just three voters in particular.

Techniques. We use different techniques to prove our different positive results. The proof of Squared Distortion (Theorem 2 in Section 5) relies heavily on Lemma 1, in which we prove (essentially) that Random Referee chooses a low social cost alternative as long as at least two of the three agents chosen at random are near the social optimum.

The proof of Distortion for Euclidean spaces (Theorem 5 in Section 6) is the most technical result. We upper bound the Distortion of a mechanism by the worst case “pessimistic distortion,” of just a constant size tuple of points, where the “pessimistic distortion” considers all permutations of the points as participating in Random Referee and allows OPT to choose the optimal point on just this tuple. This allows us to employ a computer assisted analysis by arguing that if a high Distortion instance exists, we can detect it as an instance with a small constant number of points on a sufficiently fine (but finite) grid. This approach may be of independent interest for providing tighter Distortion bounds for mechanisms in specific structured metric spaces.

3 Related Work

Distortion of Randomized Social Choice Mechanisms in Metrics. The Distortion of randomized social choice mechanisms in metrics has been studied in (Boutilier et al. 2015; Anshelevich and Postl 2017; Goel, Krishnaswamy, and Munagala 2017; Gross, Anshelevich, and Xia 2017). Of particular interest to us are the Random Dictatorship mechanism that uses a single favorite query and the 2-Agree mechanism (Gross, Anshelevich, and Xia 2017) that uses at most $\min(|N| + 1, |\mathcal{S}| + 1)$ favorite queries. Random Dictatorship has an upper bound on Distortion of 3 (Anshelevich and Postl 2017), and 2-Agree provides a strong guarantee on Distortion when $|\mathcal{S}|$ is small (better than Random Dictatorship for $|\mathcal{S}| \leq 6$). There is ongoing work on analyzing the Distortion of randomized ordinal mechanisms for other classic optimization problems like graph optimization (Abramowitz and Anshelevich 2018) and facility location (Anshelevich and Zhu 2018).

Squared Distortion and Variance. We are aware of two papers in mechanism design that consider the variance of mechanisms for facility location on the real line (Procaccia, Wajc, and Zhang 2018) and kidney exchange (Esfandiari and Kortsarz 2016). Our work is more related to the former, but is not restricted to the real line, and does not

focus on characterizing the tradeoff between welfare and variance. Using Squared Distortion as a proxy for risk was introduced in (Fain et al. 2017) along with the sequential deliberation protocol. Unlike sequential deliberation, Random Referee makes a constant number of ordinal queries. The most important baseline for Squared Distortion is the deterministic Copeland rule, which has Distortion 5 (Anshelevich et al. 2018) and therefore Squared Distortion 25. However, Copeland requires the communication of $\Omega(|N||\mathcal{S}|\log(|\mathcal{S}|))$ bits (Conitzer and Sandholm 2005), essentially the entire preference profile. Our Random Referee mechanism has constant sample complexity, and has better bounds on Squared Distortion (21) and Distortion (4.583).

Communication Complexity. For a survey on the complexity of eliciting ordinal preferences to implement social choice rules, we refer the interested reader to (Brandt et al. 2016). Of particular interest to us is (Conitzer and Sandholm 2005), in which the authors comprehensively characterize the *communication complexity* (in terms of the number of bits communicated) of common deterministic voting rules. A favorite query requires $O(\log(|\mathcal{S}|))$ bits of communication, so our mechanisms have constant sample complexity, but logarithmic communication complexity. (Bouveret et al. 2017) and (Caragiannis and Procaccia 2010) design social choice mechanisms with low communication complexity when there are a small number of voters, but potentially a large number of alternatives. All of our mechanisms have guarantees that are independent of the number of voters.

Strategic Incentives. We do not consider truthfulness in this paper, and we do not use the term mechanism to imply any such property. While strategic incentives are not the focus of this work, we note that any truthful mechanism must have Distortion at least 3 (Feldman, Fiat, and Golomb 2016). Random Dictatorship has a Distortion of 3, and is therefore in some sense optimal among exactly truthful mechanisms. Other works suggest that truthfulness is also incompatible with the weaker notion of Pareto efficiency in randomized social choice (Brandl et al. 2018; Aziz, Brandl, and Brandt 2014). Still other authors have considered the problem of truthful welfare maximization under range voting (Filos-Ratsikas and Miltersen 2014) and threshold voting (Bhaskar, Dani, and Ghosh).

4 Preliminaries

We have a set N of agents (or voters) and a set $\mathcal{S}$ of alternatives, from which we must choose a single outcome. For each agent $u \in N$ and alternative $a \in \mathcal{S}$, there is some dis-utility $d_u(a)$. Let $p_u = \operatorname{argmin}_{a \in \mathcal{S}} d_u(a)$, i.e., p_u is the most preferred alternative for agent u . Ordinal preferences are specified by a total order σ_u consistent with these disutilities (i.e., an alternative is ranked above another only if it has lower dis-utility). A preference profile $\sigma^{(N)}$ specifies the ordinal preferences of all agents. A deterministic social choice rule is a function f that maps a preference profile $\sigma^{(N)}$ to an alternative $a \in \mathcal{S}$. A randomized social choice rule maps a preference profile $\sigma^{(N)}$ to a distribution over $\mathcal{S}$.

We consider mechanisms that implement a randomized social choice rule using a constant number of queries of two types. A *top- k query* asks an agent $u \in N$ for the first k preferred alternatives according to the order σ_u (ties can be broken arbitrarily). We refer to a top-1 query as a *favorite query*, that asks an agent $u \in N$ for her most preferred alternative p_u . A *top- k only mechanism* f uses only top- k queries, for some constant k (constant with respect to $|N|$ and $|\mathcal{S}|$). Most of our lower bounds or impossibilities will be for any top- k only mechanism (for constant k), whereas our positive results will only need favorite and comparison queries. A *comparison query* with alternatives $a \in \mathcal{S}$ and $b \in \mathcal{S}$ asks an agent $u \in N$ for $\operatorname{argmin}_{x \in \{a,b\}} d_u(x)$.

We use the term mechanism to clarify that our algorithms are in a query model. However, it is important to note that mechanisms so defined are still randomized social choice rules in the formal sense as long as they do not make queries based on exogenous information (e.g., names of participants). Our mechanisms will in fact be randomized social choice rules, and thus can be appropriately compared to other such rules in the literature that do not explicitly use a query model. By using the term mechanism, we do not mean to imply any strategic properties.

Distortion and Sample Complexity. We measure the quality of an alternative $a \in \mathcal{S}$ by its *social cost*, given by $SC(a) = \sum_{u \in N} d_u(a)$. Let $a^* \in \mathcal{S}$ be the minimizer of social cost. We define the commonly studied approximation factor called Distortion (Procaccia and Rosenschein 2006), which measures the worst case approximation to the optimal social cost of a given mechanism. We use the expected social cost if a is the outcome of a randomized mechanism, and we seek to minimize Distortion.

Definition 1. The *Distortion* of an alternative a is $\text{Distortion}(a) = \frac{SC(a)}{SC(a^*)}$. The *Distortion* of a social choice mechanism f is

$$\text{Distortion}(f) = \sup_{\{d_u(a)\}} \mathbb{E}_{f(\sigma^{(N)})} \text{Distortion}(a)$$

where $\sigma^{(N)}$ is a preference profile consistent with $\{d_u(a)\}$.

We assume that $\mathcal{S}$ is a set of points in a metric space such that dis-utility can be measured by the distance from an agent. Specifically, we assume there is a distance function $d : (N \cup \mathcal{S}) \times (N \cup \mathcal{S}) \rightarrow \mathbb{R}_{\geq 0}$ satisfying the triangle inequality such that $d_u(a) = d(u, a)$. The metric assumption is common in the implicit utilitarian literature (Anshelevich and Postl 2017; Goel, Krishnaswamy, and Munagala 2017; Fain et al. 2017; Gross, Anshelevich, and Xia 2017; Cheng, Dughmi, and Kempe 2017; Anshelevich et al. 2018; Feldman, Fiat, and Golomb 2016). It is also a natural assumption for capturing social choice problems for which there is a natural notion of distance between alternatives. For example, in our original motivating example of public budgets, there are natural notions of distance between alternatives in terms of dollars.

We do not assume access to $\sigma^{(N)}$ directly, which may be prohibitively difficult to elicit when there are many alternatives. Instead, we work with a query model. The queries are

ordinal in the sense that they can be answered given only the information in $\sigma^{(N)}$. A mechanism f has *constant sample complexity* if there is an absolute constant c such that for all S , N , and $\sigma^{(N)}$, f can be implemented using at most c queries. In this paper, we consider top- k (and the special case of favorite) and comparison queries, and explore mechanisms with constant sample complexity.

Squared Distortion. It is easy to see that randomization is necessary for constant sample complexity mechanisms to achieve constant Distortion. As N grows large, any deterministic mechanism with constant sample complexity deterministically ignores (asks no queries of and receives no information from) an arbitrarily large fraction of N . An adversary can therefore place an alternative with 0 dis-utility for arbitrarily many agents; this gives a lower bound for Distortion approaching N as N becomes large.

This naturally leads us to ask: If we look at the distribution of outcomes produced by the mechanism, is this distribution well behaved? Following (Fain et al. 2017), we capture this notion via *Squared Distortion*: essentially the approximation to the optimal second moment of social cost.

Definition 2. The *Squared Distortion* of an alternative $a \in S$ is $\text{Distortion}^2(a) = \left(\frac{SC(a)}{SC(a^*)} \right)^2$. The *Distortion* of a social choice mechanism f is

$$\text{Distortion}^2(f) = \sup_{\{d_u(a)\}} \mathbb{E}_{f(\sigma^{(N)})} \text{Distortion}^2(a)$$

where $\sigma^{(N)}$ is a preference profile consistent with $\{d_u(a)\}$.

Note that a mechanism with constant Squared Distortion has both constant Distortion (by Jensen's inequality), and constant coefficient of variation. One way to interpret having constant Squared Distortion is that the deviation of the social cost around its mean falls off quadratically instead of linearly, which means that the social cost of such a mechanism is well concentrated around its mean value. We note that this interpretation gives randomized social choice mechanisms with constant Squared Distortion an interesting application in candidate selection. In particular, one can imagine running such a mechanism (like our Random Referee) to generate a candidate list on which one can use a deterministic (but potentially complex) voting mechanism like Copeland.

A related approach to understanding the distributional properties of a randomized mechanism is to characterize the tradeoff between Distortion (approximation to the first moment) and the variance of randomized mechanisms (Procaccia, Wajc, and Zhang 2018). Our specific goal in this paper is to develop a mechanism that achieves constant Distortion and constant variance, and this combination is captured by having constant Squared Distortion. We leave characterizing the exact tradeoff between the quantities as an interesting open direction.

5 Random Referee and Squared Distortion

Our first result is Theorem 1: Mechanisms that only elicit top- k preferences, for constant k , must necessarily have

Squared Distortion that grows linearly in the size of the instance. This holds even for mechanisms that elicit the top- k preferences of all of the voters, mechanisms which would not have constant sample complexity.

Theorem 1. Any top- k only social choice mechanism has Squared Distortion $\Omega(|S|)$.

The problem with top- k only mechanisms, exploited in the proof of Theorem 1, is that they treat agents as indifferent between their $(k+1)$ st favorite alternative and their least favorite alternative. This motivates the expansion of our query model to include *comparison queries*. Recall that a comparison query with alternatives $a \in S$ and $b \in S$ asks an agent $u \in N$ for $\arg\min_{x \in \{a,b\}} d_u(x)$. We use a single comparison query in our Random Referee mechanism.

Definition 3. The *Random Referee (RR)* mechanism samples three agents $u, v, w \in N$ independently and uniformly at random with replacement. u and v are asked for their favorite feasible alternatives p_u and p_v in S , and then w is asked to compare p_u and p_v . Output whichever of the two alternatives w prefers.

Upper Bound for Random Referee. Our main result in this section is Theorem 2: The Squared Distortion of RR is at most 21. This also implies that the Distortion of RR is at most 4.583. As far as we are aware, no randomized mechanism has lower Squared Distortion in general.

We will need the following technical lemma. Let $a^* \in S$ be the social cost minimizer. The basic intuition behind Random Referee is that it will choose a low social cost alternative as long as any two out of the three agents selected are near the optimal alternative a^* . Lemma 1 makes this intuition formal. For convenience, let $Z_u = d(u, a^*)$ for $u \in N$, i.e., the dis-utility of a^* for agent u . Let $C(\{u, v\}, w) = \arg\min_{y \in \{p_u, p_v\}} d(w, y)$, that is, the alternative Random Referee outputs when u and v are selected to propose alternatives and w chooses between them.

Lemma 1. For all $u, v, w, x \in N$,

$$d(C(\{u, v\}, w), x) \leq Z_x + 2 \min \left\{ \max(Z_u, Z_v), Z_w + \min(Z_u, Z_v) \right\}.$$

Proof. One should think of $u, v \in N$ as the agents drawn to present their favorite alternative, w as the referee, and x as the agent from whom we measure the dis-utility of the resulting outcome. The triangle inequality implies that

$$d(C(\{u, v\}, w), x) \leq Z_x + d(a^*, C(\{u, v\}, w)).$$

We give two separate upper bounds on $d(a^*, C(\{u, v\}, w))$, thus yielding the min. We will frequently use the fact that for all $u \in N$, $d(u, a^*) \leq d(p_u, u) + d(u, a^*) \leq 2Z_u$. This follows from the definition of p_u : the favorite alternative of u , and thus no greater in distance from u than a^* . The first bound is straightforward: $C(\{u, v\}, w) \in \{p_u, p_v\}$ by definition of Random Referee, so

$$\begin{aligned} d(a^*, C(\{u, v\}, w)) &\leq \max(d(a^*, p_u), d(a^*, p_v)) \\ &\leq 2 \max(Z_u, Z_v) \end{aligned}$$

In a sense, this bound concerns the situation where both u and v are near a^* , but w is far away from a^* . Now we argue for the second bound. Suppose without loss of generality that $Z_u \leq Z_v$. w chooses either p_u or p_v . If w chooses p_u , then $d(a^*, C(\{u, v\}, w)) = d(a^*, p_u) \leq 2d(a^*, u) = 2\min(Z_u, Z_v)$, and so the bound holds. If w chooses p_v , then $d(w, p_v) \leq d(w, p_u)$, which implies

$$\begin{aligned} d(a^*, C(\{u, v\}, w)) &= d(a^*, p_v) \\ &\leq d(a^*, w) + d(w, p_u) \\ &\leq Z_w + d(w, a^*) + d(a^*, p_u) \\ &\leq 2Z_w + 2Z_u \\ &= 2(Z_w + \min(Z_u, Z_v)). \end{aligned}$$

In either case, $d(a^*, C(\{u, v\}, w))$ is at most $2(Z_w + \min(Z_u, Z_v))$. Intuitively, this bound concerns the situation where w and u are close to a^* , but v is far away from a^* . Taking the better of this bound with the $d(a^*, C(\{u, v\}, w)) \leq 2\max(Z_u, Z_v)$ bound through the min and factoring out the 2 yields the lemma. $\square$

Using Lemma 1, we can upper bound the Squared Distortion of Random Referee by 21. This is in contrast to the $\Omega(|\mathcal{S}|)$ Squared Distortion of any favorite only mechanism (see Theorem 1).

Theorem 2. *The Squared Distortion of Random Referee is at most 21.*

Proof. Let OPT be the optimal squared social cost, that is, the squared social cost of a^* , where $a^* \in \mathcal{S}$ is the social cost minimizer. Recall that $Z_u = d(u, a^*)$. Then

$$OPT = \left(\sum_{u \in N} Z_u \right)^2 = \sum_{u \in N} \sum_{v \in N} Z_u Z_v.$$

Let ALG be the expected squared social cost of Random Referee. The expectation can be written out as

$$ALG = \frac{1}{|N|^3} \sum_{u, v, w \in N} \left(\sum_{x \in N} d(C(\{u, v\}, w), x) \right)^2.$$

Let $\alpha = \min(\max(Z_u, Z_v), Z_w + \min(Z_u, Z_v))$. We apply Lemma 1 and simplify.

$$ALG \leq \frac{1}{|N|^3} \sum_{u, v, w \in N} \left(\sum_{x \in N} (Z_x + 2\alpha) \right)^2$$

Noting that α does not depend on x , we can expand the square and simplify to find

$$ALG \leq OPT + \frac{4}{|N|^3} \sum_{u, v, w \in N} \left(\alpha^2 |N|^2 + \alpha |N| \sum_{x \in N} Z_x \right).$$

Now we sum each term separately. Let

$$\begin{aligned} T_1 &= \frac{4}{|N|^3} \sum_{u, v, w \in N} \alpha^2 |N|^2 \\ T_2 &= \frac{4}{|N|^3} \sum_{u, v, w \in N} \alpha |N| \sum_{x \in N} Z_x \end{aligned}$$

We will use the following basic facts: for any real numbers $a, b \geq 0$, $(\min(a, b))^2 \leq a \cdot b$, $\max(a, b) \leq a + b$, and $\min(a, b) \cdot \max(a, b) = a \cdot b$.

$$\begin{aligned} T_1 &\leq \frac{4}{|N|} \sum_{u, v, w \in N} (\max(Z_u, Z_v) (Z_w + \min(Z_u, Z_v))) \\ &\leq \frac{4}{|N|} \sum_{u, v, w \in N} ((Z_u + Z_v) Z_w + Z_u Z_v) \\ &= 12 OPT \end{aligned}$$

Similarly, we analyze the second term using the fact that $\alpha \leq Z_u + Z_v$.

$$\begin{aligned} T_2 &= \frac{4}{|N|^3} \sum_{u, v, w \in N} \alpha |N| \sum_{x \in N} Z_x \\ &\leq \frac{4}{|N|^2} \sum_{u, v, w \in N} (Z_u + Z_v) \sum_{x \in N} Z_x \\ &= 8 OPT \end{aligned}$$

Adding together all of the terms, $ALG \leq 21 OPT$. $\square$

This immediately yields the root mean square Distortion bound via Jensen's inequality.

Corollary 1. *The Distortion of Random Referee is at most $\sqrt{21} \approx 4.583$.*

6 Distortion of Random Referee on the Euclidean Plane

Though the upper bound on the Distortion of Random Referee is slightly worse than that of Random Dictatorship, we now show another advantage of using a comparison query: Such mechanisms can exploit structure in specific metric spaces that favorite-only mechanisms cannot. In other words, we show that the Distortion of Random Referee improves significantly for more structured metric spaces, while top- k only mechanisms do not share this property.

We examine the Distortion of Random Referee on a specific canonical metric of interest: the Euclidean plane when $d(u, p_u) = 0$ for every $u \in N$. The second assumption, functionally equivalent to assuming $N \subseteq \mathcal{S}$, simplifies our analysis considerably, and corresponds to the 0-decisive case in (Anshelevich and Postl 2017; Gross, Anshelevich, and Xia 2017). We note that in examples where $|\mathcal{S}|$ is very large, the assumption becomes more innocuous. If we consider our opening example of public budgets, the assumption is something like this: every agent is allowed to propose their absolute favorite over all budgets, and we assume that this budget p_u has dis-utility of 0. We consider the Euclidean plane because the problem of minimizing distortion on the real line can be solved exactly (Anshelevich and Postl 2017), whereas we are unaware of any results for the Euclidean plane that are stronger than those for general metric spaces.

Lower Bounds for Distortion in the Restricted Model.

We begin by giving lower bounds to demonstrate that our simplifying assumptions still result in a nontrivial problem

in two senses: (1) the Distortion of randomized social choice mechanisms are still bounded away from 1 and (2) any top- k only mechanism (one that elicits the top k preferred alternatives of agents) for constant k has Distortion at least 2 as $|N|$ and $|S|$ become large.

Theorem 3. *The Distortion of a randomized social choice mechanism is at least 1.2 generally, and at least 1.118 for the Euclidean plane, even when $d(u, p_u) = 0$ for every $u \in N$.*

Theorem 4. *The Distortion of any top- k only mechanism goes to 2 as $|S|$ and $|N|$ become large, even on the Euclidean plane when $d(u, p_u) = 0$ for every $u \in N$.*

Upper Bound for Distortion of Random Referee in the Restricted Model. Our positive result in this section demonstrates that a single comparison query is sufficient to construct a mechanism (Random Referee) with Distortion bounded below 1.97 for arbitrary $|S|$ and $|N|$. We note that the bound in the theorem seems very slack; our goal is simply to show that using a comparison query provably decreases Distortion. We conjecture the actual bound is below 1.75 based on computer assisted search, but leave proving this stronger bound as an interesting open question.

Theorem 5. *The worst case distortion of Random Referee is less than 1.97 when $d(u, p_u) = 0$ for every $u \in N$, $S \subseteq \mathbb{R}^2$, and for $x, y \in \mathbb{R}^2$ $d(x, y) = \|x - y\|_2$.*

In the remainder of this section, we sketch the proof of Theorem 5. Suppose for a contradiction that there is a set N of agents in $\mathbb{R}^2$ such that the Distortion of Random Referee is at least 1.97 under the Euclidean metric. We will successively refine this hypothesis for a contradiction, finally arguing that it implies that some “bad” instance would appear on an exhaustive computer assisted search over a finite grid.

The crucial lemmas bound the *pessimistic distortion* of any set of five points in $\mathbb{R}^2$ and relate this to the actual distortion. The quantity is pessimistic because it allows “OPT” to choose a separate point for every 5-tuple; the numerator of each 5-tuple is just a rewriting of the expected social cost of Random Referee. In this section, since we assume $d(u, p_u) = 0$, it will not be necessary to refer to u and p_u separately, and it will be convenient to let $C(\{p_u, p_v\}, p_w) = \argmin_{a \in \{p_u, p_v\}} d(p_w, a)$.

Definition 4. *The **pessimistic distortion**¹ of $\mathcal{P} = \{x_1, x_2, x_3, x_4, x_5\} \subset \mathbb{R}^2$ is defined as*

$$PD(\mathcal{P}) := \frac{SCRR_{avg}(\mathcal{P})}{OPT_{avg}(\mathcal{P})}$$

where $SCRR_{avg}(\mathcal{P})$ is the average social cost of Random Referee, which we can write as

$$\frac{1}{30} \sum_{i=1}^5 \sum_{j>i}^5 \sum_{k \neq i,j}^5 \sum_{l \neq i,j,k}^5 d(x_l, C(\{x_i, x_j\}, x_k)),$$

¹Note that pessimistic distortion is technically a function of the mechanism used; we will use it exclusively in reference to Random Referee.

and the average cost of the optimal solution is

$$OPT_{avg}(\mathcal{P}) := \min_{y \in \mathbb{R}^2} \frac{1}{5} \sum_{r=1}^5 d(x_r, y).$$

Finally, $PD(\mathcal{P}) = 1$ if $x_1 = x_2 = x_3 = x_4 = x_5$.

The first lemma relates this pessimistic distortion to the actual Distortion of Random Referee. The worst case (over 5-tuples) pessimistic distortion upper bounds the Distortion of Random Referee. Interestingly, this statement is not specific to the Euclidean plane, suggesting that our approach may be broadly applicable for proving stronger Distortion bounds on other specific metrics of interest.

Lemma 2. *If $PD(\mathcal{P}) \leq \beta$ for all $\mathcal{P} = \{x_1, \dots, x_5\} \subset \mathbb{R}^2$ then the Distortion of Random Referee is at most β on $\mathbb{R}^2$.*

Proof. Observe that we can rewrite the Distortion of Random Referee as a summation over all possible 5-tuples of points. Let $\mathcal{P}$ be a multiset of points (i.e., possibly with repeats). Let $\rho(\mathcal{P})$ be an ordering of $\mathcal{P}$. Let

$$C(\rho(\mathcal{P})) = C(\{\rho_1(\mathcal{P}), \rho_2(\mathcal{P})\}, \rho_3(\mathcal{P})).$$

We can rewrite the Distortion of Random Referee over permutations of 5-tuples as

$$\begin{aligned} & \frac{\frac{1}{|N|^4} \sum_{i,j,k,l \in N} d(p_l, C(\{p_i, p_j\}, p_k))}{\frac{1}{|N|} \sum_{i \in N} d(p_i, a^*)} \\ &= \frac{\sum_{\mathcal{P} \subset \mathcal{P}_N} \Pr(\mathcal{P}) \frac{1}{5!} \sum_{\rho(\mathcal{P})} \sum_{l \in \{4,5\}} d(\rho_l(\mathcal{P}), C(\rho(\mathcal{P})))}{\frac{1}{|N|} \sum_{i \in N} d(p_i, a^*)}. \end{aligned}$$

In words, we are considering all 5-tuples (with replacement) of agent points, and for each we consider the average distance over all orderings of the five agent points of the distance between the last two points and the outcome of Random Referee when the first three agents participate. This is in turn upper bounded by allowing OPT to choose a different a^* for every 5-tuple, so that the Distortion is at most

$$\begin{aligned} & \sum_{\mathcal{P} \subset \mathcal{P}_N} \Pr(\mathcal{P}) \frac{\frac{1}{5!} \sum_{\rho(\mathcal{P})} \sum_{l \in \{4,5\}} d(\rho_l(\mathcal{P}), C(\rho(\mathcal{P})))}{OPT_{avg}(\mathcal{P})} \\ & \leq \max_{\mathcal{P} \subset \mathcal{P}_N} \Pr(\mathcal{P}) \frac{\frac{1}{5!} \sum_{\rho(\mathcal{P})} \sum_{l \in \{4,5\}} d(\rho_l(\mathcal{P}), C(\rho(\mathcal{P})))}{OPT_{avg}(\mathcal{P})}. \end{aligned}$$

To complete the proof, note that $SCRR_{avg}(\mathcal{P})$ is equal to the numerator. We start with all 120 orderings of the 5 points and avoid double counting the symmetric cases that arise from swapping the two points from which we take the argmin and swapping the two points from which we measure distance. $\square$

Now we can refine our original hypothesis: without loss of generality, assume for a contradiction that there is a multiset $\mathcal{P} = \{x_1, \dots, x_5\} \subseteq \mathbb{R}^2$ with $PD(\mathcal{P}) \geq 1.97$.

Lemma 3. $PD(\mathcal{P}) < 1.97$.

Lemma 3 provides the contradiction to our hypothesis and establishes Theorem 5. We outline the main ideas of the proof. First, we consider a grid on the Euclidean plane, and argue that we can assume a certain canonical structure of $\mathcal{P}$ in relation to that grid by scaling, translating, and rotating. Next, we carefully argue for how much the pessimistic distortion would change if every point in $\mathcal{P}$ were snapped to a grid point. We show that if there is some $\mathcal{P}$ with pessimistic distortion at least 1.97, there must be a 5-tuple of points on a sufficiently fine (but finite) grid with a sufficiently large pessimistic distortion. However, we employ computer assisted analysis to brute force search over such a grid, and find no such bad example.

Discussion. Our analysis is slack in two ways: (1) we have to interpolate between grid points, and (2) we consider pessimistic distortion over 5-tuples. Both are computational constraints: (1) because we cannot simulate an arbitrarily fine grid, and (2) because there is a combinatorial blow up in the search space when considering larger tuples (note that by considering 5-tuples, we implicitly allow OPT to choose a separate optimal solution for each 5-tuple). For Random Referee, the worst case example found by computer simulations for grid points is fairly simple: The 5 points lie on a straight line with pessimistic distortion 1.75. We conjecture the same example is the worst case even in the continuous plane, which suggests a distortion bound of at most 1.75. We leave finding the exact bound as an open question, as our result is sufficient to demonstrate that comparison queries can take advantage of structure that top- k queries cannot.

7 Favorite Only Mechanisms: Random Oligarchy

Recall that a favorite query asks an agent $u \in N$ for her most preferred feasible point p_u . In this section, we return to the general model (arbitrary metrics and not assuming that $d(u, p_u) = 0$) and study mechanisms that are restricted to only use favorite queries. We show that essentially optimal Distortion as a function of $|\mathcal{S}|$ is achieved by a simple mechanism that uses just 3 queries. We call this mechanism Random Oligarchy.

Definition 5. *The **Random Oligarchy (RO)** mechanism samples three agents $u, v, w \in N$ independently and uniformly at random with replacement. All three are asked for their favorite alternatives p_u, p_v , and p_w in $\mathcal{S}$. If the same alternative is reported at least twice, output that, else output one of the three alternatives uniformly at random.*

We prove that Random Oligarchy has the best of both worlds with respect to the other favorite only mechanisms of Random Dictatorship and 2-Agree. Unlike 2-Agree, Random Oligarchy has constant sample complexity and the same Distortion bound of 3 for large $|\mathcal{S}|$ as Random Dictatorship. However, like 2-Agree, it outperforms Random Dictatorship for small $|\mathcal{S}|$.

Theorem 6. *The Distortion of Random Oligarchy is upper bounded by 3 for arbitrary $|\mathcal{S}|$, and by the following expres-*

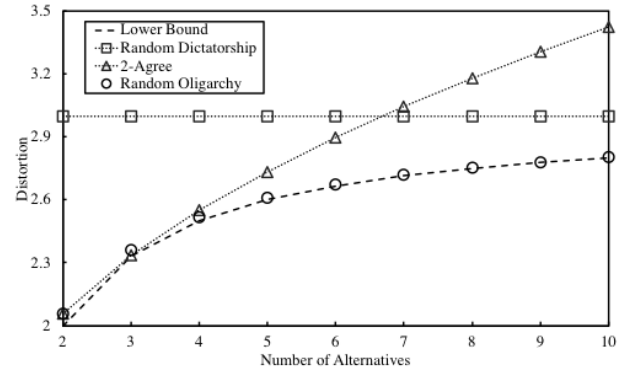


Figure 1: Distortion of Favorite Only Mechanisms. Upper bounds for Random Oligarchy are from Theorem 6. Lower bounds and 2-Agree upper bounds are from (Gross, Anshelevich, and Xia 2017). Random Dictatorship is analyzed in (Anshelevich and Postl 2017).

sion for particular $|\mathcal{S}|$.

$$1 + 2 \max_{p \in [0,1]} \left(1 + p^2(p-2) + \frac{(p-1)^3}{|\mathcal{S}| - 1} \right)$$

Figure 1 shows the Distortion bounds of favorite only mechanisms. Comparing against the lower bound for any favorite only mechanism from (Gross, Anshelevich, and Xia 2017) allows us to see that Random Oligarchy is essentially optimal among all favorite only mechanisms. In comparison to existing mechanisms, Random Oligarchy outperforms Random Dictatorship for small $|\mathcal{S}|$, and outperforms 2-Agree for large $|\mathcal{S}|$, while only using three favorite queries.

8 Open Directions

In this paper, we have considered constant sample complexity mechanisms. At a high level, we hope that our work inspires future research on lightweight mechanisms for social choice in large decision spaces. We also mention some natural technical questions raised by our work.

Compared to Distortion, there is much less understanding of the Squared Distortion of mechanisms. The only universal lower bound for Squared Distortion we are aware of is 4, a consequence of the lower bound of 2 for Distortion (Anshelevich and Postl 2017). Our Random Referee mechanism achieves Squared Distortion of at most 21 using just three ordinal queries. Closing this gap remains an interesting question even for mechanisms that elicit full ordinal information.

The analysis in Section 6 of Random Referee generalizes to higher dimensional Euclidean space, still only reasoning about 5-tuples of points. Computer search over a coarse grid in four dimensions again shows that the pessimistic distortion bound is better than 2. However, we have not been able to run the search on a fine enough grid to prove a Distortion bound formally. We leave proving this for higher dimensional Euclidean spaces as an interesting open question. Additionally, we hope that related methods may be of gen-

eral interest for proving tighter Distortion bounds of mechanisms on restricted metric spaces.

It remains unclear whether using a constant number of queries greater than 3 would meaningfully improve our results. E.g., consider the natural extension of Random Referee: sample the favorite points of k agents and ask a random referee to choose their favorite from among these. Using $k > 2$ does not straightforwardly decrease the Squared Distortion bound of 21 for $k = 2$. Also, as k becomes large, this mechanism devolves to Random Dictatorship. Another natural question is whether $O(k)$ comparison queries are necessary/sufficient to bound the k 'th moment of Distortion. We leave these as additional open questions.

Acknowledgments

Brandon Fain is supported by NSF grants CCF-1637397 and IIS-1447554. Ashish Goel is supported by NSF grant CCF-1637418 and ONR grant N00014-15-1-2786. Kamesh Munagala is supported by NSF grants CCF-1408784, CCF-1637397, and IIS-1447554. Much of this work was done while Nina Prabhu was a student at The North Carolina School of Science and Mathematics, in cooperation with Duke University.

References

- Abramowitz, B., and Anshelevich, E. 2018. Utilitarians without utilities: Maximizing social welfare for graph problems using only ordinal preferences. In *Proceedings of the Thirty-Second AAAI Conference on Artificial Intelligence*, 894–901. AAAI Press.
- Anshelevich, E., and Postl, J. 2017. Randomized social choice functions under metric preferences. *Journal of Artificial Intelligence Research* 58(1):797–827.
- Anshelevich, E., and Zhu, W. 2018. Ordinal approximation for social choice, matching, and facility location problems given candidate positions. In *The 14th Conference on Web and Internet Economics (WINE)*, 3–20. Springer.
- Anshelevich, E.; Bhardwaj, O.; Elkind, E.; Postl, J.; and Skowron, P. 2018. Approximating optimal social choice under metric preferences. *Artificial Intelligence* 264:27 – 51.
- Aziz, H.; Brandl, F.; and Brandt, F. 2014. On the incompatibility of efficiency and strategyproofness in randomized social choice. In *Proceedings of the Twenty-Eighth AAAI Conference on Artificial Intelligence*, 545–551. AAAI Press.
- Bhaskar, U.; Dani, V.; and Ghosh, A. Truthful and near-optimal mechanisms for welfare maximization in multi-winner elections. In *Proceedings of the Thirty-Second AAAI Conference on Artificial Intelligence*, 925–932.
- Boutilier, C.; Caragiannis, I.; Haber, S.; Lu, T.; Procaccia, A. D.; and Sheffet, O. 2015. Optimal social choice functions: A utilitarian view. *Artificial Intelligence* 227:190–213.
- Bouveret, S.; Chevaleyre, Y.; Durand, F.; and Lang, J. 2017. Voting by sequential elimination with few voters. In *Proceedings of the Twenty-Sixth International Joint Conference on Artificial Intelligence*, 128–134.
- Brandl, F.; Brandt, F.; Eberl, M.; and Geist, C. 2018. Proving the incompatibility of efficiency and strategyproofness via smt solving. *Journal of the ACM* 65(2):6:1–6:28.
- Brandt, F.; Conitzer, V.; Endriss, U.; Lang, J.; and Procaccia, A. D. 2016. *Handbook of Computational Social Choice*. Cambridge University Press, 1st edition.
- Caragiannis, I., and Procaccia, A. D. 2010. Voting almost maximizes social welfare despite limited communication. In *Proceedings of the Twenty-Fourth AAAI Conference on Artificial Intelligence*, 743–748. AAAI Press.
- Cheng, Y.; Dughmi, S.; and Kempe, D. 2017. Of the people: Voting is more effective with representative candidates. In *Proceedings of the 2017 ACM Conference on Economics and Computation (EC)*, 305–322. ACM.
- Conitzer, V., and Sandholm, T. 2005. Communication complexity of common voting rules. In *Proceedings of the 6th ACM Conference on Electronic Commerce*, 78–87. ACM.
- Esfandiari, H., and Kortsarz, G. 2016. A bounded-risk mechanism for the kidney exchange game. In Kranakis, E.; Navarro, G.; and Chávez, E., eds., *LATIN 2016: Theoretical Informatics*, 416–428. Springer.
- Fain, B.; Goel, A.; Munagala, K.; and Sakshuwong, S. 2017. Sequential deliberation for social choice. In *The 13th Conference on Web and Internet Economics (WINE)*, 177–190. Springer.
- Fain, B.; Goel, A.; Munagala, K.; and Prabhu, N. 2018. Random Dictators with a Random Referee: Constant Sample Complexity Mechanisms for Social Choice - Full Version. *ArXiv e-prints*.
- Feldman, M.; Fiat, A.; and Golomb, I. 2016. On voting and facility location. In *Proceedings of the 2016 ACM Conference on Economics and Computation*, 269–286. ACM.
- Filos-Ratsikas, A., and Miltersen, P. B. 2014. Truthful approximations to range voting. In *The 10th Conference on Web and Internet Economics (WINE)*, 175–188. Springer.
- Goel, A.; Krishnaswamy, A. K.; Sakshuwong, S.; and Aitamurto, T. 2015. Knapsack voting. *Collective Intelligence*.
- Goel, A.; Krishnaswamy, A. K.; and Munagala, K. 2017. Metric distortion of social choice rules: Lower bounds and fairness properties. In *Proceedings of the 2017 ACM Conference on Economics and Computation*, 287–304. ACM.
- Gross, S.; Anshelevich, E.; and Xia, L. 2017. Vote until two of you agree: Mechanisms with small distortion and sample complexity. In *Proceedings of the Thirty-First AAAI Conference on Artificial Intelligence*, 544–550. AAAI Press.
- Procaccia, A. D., and Rosenschein, J. S. 2006. The distortion of cardinal preferences in voting. In Klusch, M.; Rovatsos, M.; and Payne, T. R., eds., *Cooperative Information Agents X*, 317–331. Springer.
- Procaccia, A.; Wajc, D.; and Zhang, H. 2018. Approximation-variance tradeoffs in facility location games. In *Proceedings of the Thirty-Second AAAI Conference on Artificial Intelligence*, 1185–1192. AAAI Press.