

Every Bit Helps: Achieving the Optimal Distortion with a Few Queries (Extended Abstract)*

Soroush Ebadian and Nisarg Shah

University of Toronto

{soroush, nisarg}@cs.toronto.edu

Abstract

A fundamental task in multi-agent systems is to match n agents to n alternatives (e.g., resources or tasks). This is often done by eliciting agents' ordinal rankings over the alternatives rather than their exact numerical utilities. While this simplifies elicitation, the incomplete information leads to inefficiency, captured by a worst-case measure called *distortion*. Recent work shows that making just a few cardinal utility queries per agent can significantly improve the distortion, and in particular that $O(\sqrt{n})$ distortion is achievable with two queries per agent. We generalize this result by achieving $O(n^{1/\lambda})$ distortion with λ queries per agent, for any constant λ , which is optimal up to a constant factor given a previously known lower bound. We extend this finding to the general social choice problem of selecting one of m alternatives based on n agents' preferences, achieving $O((\min\{n, m\})^{1/\lambda})$ distortion with λ queries per agent.

1 Introduction

Imagine you are tasked with allocating office spaces in a medical building to a group of doctors. Each doctor provides a ranked list of their preferred offices, and your goal as the manager is to find a matching that maximizes the overall satisfaction, or social welfare, of all doctors. While rankings give you a general sense of their preferences, they lack information about *intensity* of preferences. If you could ask a few targeted questions to obtain their exact numerical utilities, known as cardinal queries, you may be able to improve the allocation significantly. However, these queries are cognitively burdensome for the doctors to answer, so they must be designed carefully in order to maximize their benefit while minimizing burden on the doctors.

The concept of distortion, introduced by Procaccia and Rosenschein [2006], quantifies the loss in social welfare due to the lack of exact numerical preferences. Traditionally, distortion has been studied in settings where only ordinal

information is available (see the survey by Anshelevich *et al.* [2021]). For instance, in one-sided matching problems [Hylland and Zeckhauser, 1979], where n agents are matched to n items based on their preferences, the best achievable distortion is known to be $\Theta(\sqrt{n})$, though this requires randomization and normalized values [Filos-Ratsikas *et al.*, 2014].

Recent work has expanded this framework by considering the trade-offs between the amount of cardinal information gathered and the efficiency of the resulting allocation. Amanatidis *et al.* [2021; 2022] demonstrated that by allowing mechanisms to ask $\lambda \cdot \log n$ cardinal queries per agent, $O(n^{1/(\lambda+1)})$ distortion can be achieved, even without relying on randomization or normalization. Building on this, Amanatidis *et al.* [2024] showed that a distortion of $\Theta(\sqrt{n})$ can be achieved by asking only two queries per agent (instead of $\log n$), matching the performance of the best randomized mechanism that assumes normalized values.

Despite these advancements, our understanding of the trade-offs between using two cardinal queries per agent and a logarithmic number of queries remains limited. This raises an important question:

As we go from just two to three or more (but a constant number of) queries per agent, how well can we approximate the optimal social welfare of any one-sided matching?

Shifting from matching to voting, consider another important decision in the medical building: selecting the location for a new specialized clinic. The doctors must vote on several potential locations, each offering distinct advantages, such as proximity to the emergency room, patient accessibility, or the available space. Each doctor has their own underlying utilities for these locations, influenced by their specialty and the needs of their patients. While rankings provide a general understanding of their preferences, they once again fail to capture preference intensities. By eliciting a few cardinal queries, the decision-maker could potentially improve the selection process, aiming to maximize the collective satisfaction of the group, just as before. However, this voting problem is even less understood than the matching problem. When only two queries per agent are allowed, despite the attempt of Amanatidis *et al.* [2024], the problem of optimal achievable distortion remains unresolved. This again raises the same question posed earlier, but now for the setting more

*This extended abstract is based on the AAAI 2025 paper of the same title [Ebadian and Shah, 2025].

# Queries	Upper Bounds	Lower Bounds
0 (Ordinal)	-	Unbounded [†]
1	$O(n)$ [†]	$\Omega(n)$ [†]
2	$O(\sqrt{n})$ [§]	$\Omega(\sqrt{n})$ [†]
$\lambda = \Theta(1)$	$O(n^{1/\lambda})$	$\Omega(n^{1/\lambda})$ [†]
$\lambda \leq \log n$	$\lambda \cdot n^{1/\lambda}$	$\Omega(n^{1/\lambda}/\lambda)$ [†]
$\log n$	$O(\log n)$	$\Omega(1)$ [†]
$z \cdot \log n, z \geq 1$	$O(n^{1/(z+1)})$ [†]	
$O(\log^2 n)$	$O(1)$ [†]	

Table 1: Summary of our results and prior distortion bounds for one-sided matching with n agents and n alternatives. Gray cells highlight results from Theorem 5. [†][Amanatidis *et al.*, 2022], [§][Amanatidis *et al.*, 2024].

precisely called *single-winner elections*.

1.1 Our Contributions

We present a novel algorithm for both one-sided matching and single-winner elections, which leverages a limited number of λ cardinal queries per agent to achieve asymptotically optimal distortion bounds when λ is a constant. Table 1 provides a summary of our results alongside relevant prior work in one-sided matching. We impose no restrictions on the utilities, and all our algorithms are deterministic and run in poly-time. Our work builds on the ideas of Amanatidis *et al.* [2024] and introduces novel applications of the notion of *stable committees* in multi-winner elections [Aziz *et al.*, 2017; Cheng *et al.*, 2020].

One-sided matching. For the one-sided matching problem, we establish that the optimal distortion with a constant number λ of queries per agent is $O(n^{1/\lambda})$. This significantly generalizes the work of Amanatidis *et al.* [2024], who only considered two-query algorithms. Our algorithm achieves $O(n^{1/3})$ distortion with just three queries, a result that previously required $O(\log n)$ queries [Amanatidis *et al.*, 2022].

More generally, for any $1 \leq \lambda \leq \log n$, we achieve distortion of $\lambda \cdot n^{1/\lambda}$ with λ queries. This implies $O(\log n)$ distortion with $\log n$ queries, improving upon the previous best result that required $O(\log^2 n / \log \log n)$ queries [Amanatidis *et al.*, 2022]. The (asymptotic) optimality of our results for constant λ follows from the $\Omega(n^{1/\lambda})$ lower bound established by Amanatidis *et al.* [2022].

We also establish the existence of an *exactly* stable committee of matchings by modifying the serial dictatorship algorithm. This result may be of independent interest, as only approximately stable committees exist in committee selection with ranked preferences [Jiang *et al.*, 2020].

Single-winner elections. Applying similar techniques, we achieve comparable results for single-winner elections with a distortion bound of $O((\min\{n, m\})^{1/\lambda})$ for constant λ . While Amanatidis *et al.* [2024] showed $O(\sqrt{m})$ distortion for $m = \Omega(n)$ with two queries, the case of $m = o(n)$, with the number of agents far exceeding the number of candidates, remained open. Notably, our bounds depend on

# Queries	Upper Bounds	Lower Bounds
0 (Ordinal)	-	Unbounded [‡]
1	$O(\alpha_{n,m}), O(m)$ [‡]	$\Omega(m)$ [‡]
2	$O(\sqrt{\alpha_{n,m}})$	$\Omega(\sqrt{m})$ [§]
$\lambda = \Theta(1)$	$O(\alpha_{n,m}^{1/\lambda})$	$\Omega(m^{1/\lambda})$ [§]
$\lambda \leq \log(\alpha_{n,m})$	$O(\lambda \cdot \alpha_{n,m}^{1/\lambda})$	$\Omega(m^{1/(3\lambda)})$ [*]
$\log(\alpha_{n,m})$	$O(\log(\alpha_{n,m}))$	$\Omega(1)$
$z \cdot \log m, z \geq 1$	$O(m^{1/(z+1)})$ [‡]	
$O(\log^2 m)$	$O(1)$ [‡]	

Table 2: Summary of our results and prior distortion bounds for single-winner elections with n agents and m candidates. Gray cells (i.e., bounds with $\alpha_{n,m}$) indicate results from Theorem 6, where $\alpha_{n,m} = \min\{n, m\}$. All lower bounds assume $n \geq \Omega(m)$. [‡][Amanatidis *et al.*, 2021], [§][Amanatidis *et al.*, 2024], ^{*}[Caragiannis and Fehr, 2024].

$\min\{n, m\}$ rather than just m as in typical voting distortion results. For example, this allows modeling one-sided matching as a single-winner election with $n!$ possible matchings as alternatives, yet still achieve an appealing bound of $O(n^{1/\lambda})$ with only a constant factor increase.

1.2 Related Work

The study of querying beyond ordinal preferences builds on earlier work on distortion using only ordinal information, initiated by Procaccia and Rosenschein [2006] in the context of single-winner elections. Caragiannis and Procaccia [2011] and Caragiannis *et al.* [2017] demonstrate that the optimal distortion achievable by *deterministic* voting rules under normalized valuations is $\Theta(m^2)$. For *randomized* voting rules, Boutilier *et al.* [2015] and Ebadian *et al.* [2022] identify the best achievable distortion as $\Theta(\sqrt{m})$. Additionally, Caragiannis *et al.* [2017] explore distortion in multi-winner elections. Borodin *et al.* [2022] consider scenarios where even less information than ranked votes is used. Distortion is also examined in the *metric voting* framework [Anshelevich and Postl, 2017; Anshelevich *et al.*, 2018], with a comprehensive overview provided by Anshelevich *et al.* [2021].

Some studies relax query-type restrictions, allowing any query that elicits a fixed number of bits from voters, including ordinal elicitation, and focus on optimizing the query format itself [Mandal *et al.*, 2019; Mandal *et al.*, 2020] (and Kempe [2020] for metric voting). For single-winner elections, deterministic elicitation requires $\tilde{\Theta}(m/d)$ bits, while randomized elicitation requires $\tilde{\Theta}(m/d^3)$.

Latifian and Voudouris [2024] and Ma *et al.* [2021] investigate threshold approval queries for one-sided matching, while Ebadian *et al.* [2023] apply the same elicitation to voting. In the metric voting framework, Ebadian *et al.* [2024] propose nearly-optimal algorithms using limited pairwise comparisons per agent, with various constraints on query adaptivity, while Anagnostides *et al.* [2022] study adaptive algorithms that ask the same set of queries to all agents.

2 Model

Let $[t] := \{1, \dots, t\}$ for $t \in \mathbb{N}$.

One-sided matching. In one-sided matching, there is a set N of n agents and a set A of n alternatives. A matching M is a one-to-one mapping from N to A , pairing each agent with a unique alternative. Each agent $i \in N$ has a valuation function $u_i : A \rightarrow \mathbb{R}_{\geq 0}$, where $u_i(a)$ denotes the utility agent i receives when matched to alternative a . The utilitarian social welfare of a matching M is denoted by $sw(M) = \sum_{i \in N} u_i(M(i))$.

Implicitly utilitarian ordinal matching. For each agent $i \in N$, let $\sigma_i : [n] \rightarrow A$ be the preference ranking of agent i induced by the valuation function u_i (ties broken arbitrarily); that is, $u_i(\sigma_i(1)) \geq \dots \geq u_i(\sigma_i(n))$. The preference profile $\vec{\sigma} = \{\sigma_i\}_{i \in N}$ is the collection of all agents' preferences. We denote by $\vec{u} \triangleright \vec{\sigma}$ that u_i aligns with σ_i for each i . By $a \succ_i a'$ we mean i strictly prefers a to a' . An ordinal matching mechanism returns a matching based only on $\vec{\sigma}$.

Value queries. A λ -query ordinal matching mechanism, in addition to the preference profile $\vec{\sigma}$, is allowed to ask up to λ value queries per agent. Through a value query $Q(i, a)$, the algorithm learns $u_i(a)$.

Distortion. The distortion of a λ -query mechanism $\mathcal{M}$ w.r.t. a preference profile $\vec{\sigma}$ is its worst-case approximation ratio to the optimal welfare-maximizing matching $\text{OPT}(\vec{u})$ over all consistent utility profiles. That is,

$$\text{dist}(\mathcal{M}) = \max_{\vec{u} \triangleright \vec{\sigma}} \frac{sw(\text{OPT}(\vec{u}))}{sw(\mathcal{M}(\vec{\sigma}, \{Q(i, a_{i,j})\}))}.$$

3 Ordinal Matching with Value Queries

In this section, we present a λ -query mechanism with $\lambda \cdot n^{1/\lambda}$ distortion. We first describe a simple algorithm to find a stable “committee” of k matchings, which is a key subprocedure of the final matching algorithm.

3.1 Stable Matching Sets

Amanatidis *et al.* [2024] define a particular type of assignment called a “sufficiently representative set” that is specific to their two-query mechanism. The crux of this notion and its role in the algorithm, we believe, is closely related to finding a *stable committee* of $\sqrt{n}$ matchings. We extend their algorithm and show a weighted form of k -stable committee of matchings exists and can be computed by a simple algorithm for all $k \in [n]$.

Definition 1 (Stable set of matchings). *For an ordinal matching instance with agent weights w , a set of k matchings $X = \{M_1, \dots, M_k\}$ is stable if there does not exist another matching M' and agents $N' \subseteq N$ such that $w(N') \geq w(N)/k$ and $M'(i) \succ_i M_\ell(i)$ for all $i \in N'$ and $M_\ell \in X$.*

In contrast to committee selection from a set of candidates, where even 2-stable committees may fail to exist [Jiang *et al.*, 2020], we show a stable set of matchings always exist.

Theorem 2. *For an ordinal matching instance and agent weights w , there always exists a set of k matchings that is stable and can be computed in $\text{poly}(n)$ time.*

Algorithm 1 k -Capacity Serial Dictatorship

Input: Preference profile $\vec{\sigma}$, weights $\{w(i)\}_{i \in N}$, and k
Output: A k -capacity mapping

- 1: $B \leftarrow$ multiset of k copies of each alternative $a \in A$
- 2: **for** $i \in N$ in decreasing order of weights w_i **do**
- 3: Match i to their most preferred alternative $a_i \in B$
- 4: Remove one copy of a_i from B
- 5: **end for**
- 6: **return** $g = \{i \rightarrow a_i\}_{i \in N}$

Algorithm 2 λ -Query Matching Algorithm

Input: Preference profile $\vec{\sigma}$ and λ
Output: Matching $\tilde{M}$

- 1: Let $w_1(i) = 1$ for all $i \in N$
- 2: **for** $\ell \in \{1, \dots, \lambda\}$ **do**
- 3: Set $k_\ell \leftarrow n^{1-(\ell-1)/\lambda}$
- 4: $g_\ell \leftarrow k_\ell$ -Capacity Serial Dictatorship($\vec{\sigma}, \{w_\ell(i)\}_{i \in N}$)
- 5: Query each agent $i \in N$ of $g_\ell(i)$
- 6: $w_{\ell+1}(i) \leftarrow u_i(g_\ell(i))$ for all $i \in N$
- 7: **end for**
- 8: $\tilde{u}_i(a) \leftarrow \max\{u_i(g_\ell(i)) \mid a \succ_i g_\ell(i), \ell \in [\lambda]\}$, or 0 if no such queries exist
- 9: **return** social welfare maximizing matching $\tilde{M}$ based on $\{\tilde{u}_i\}_{i \in N}$.

We achieve Theorem 2 via Algorithm 1 that closely follows the $\sqrt{n}$ -Serial Dictatorship algorithm of Amanatidis *et al.* [2024]. The main difference is that we go over the agents in the order of the weights, which enables proving Theorem 2 for arbitrary agent weights and any $k \in [n]$.

k -Capacity serial dictatorship. Algorithm 1 is a serial dictatorship process that starts with a multiset containing k copies of each alternative. Agents then make their selections, by picking their favourite among the remaining alternatives, in an order determined by weights in non-increasing order. Since each alternative is mapped to at most k agents, we can decompose such a mapping g into a set of k matchings where each agent i is mapped to $g(i)$ in at least one of the k matchings.

Lemma 3. *Let $g : N \rightarrow A$ be a k -capacity mapping. There exists a set of k matchings $M_1, \dots, M_k$ such that for all $i \in N$, $g(i) = M_\ell(i)$ for some $\ell \in [k]$, and it can be computed in $\text{poly}(n)$ time.*

3.2 The λ -Query Algorithm

Next, we present our λ -query mechanism that achieves a distortion of $\lambda \cdot n^{1/\lambda}$, which uses Algorithm 1 and its guarantee (Theorem 2) to inform the queries.

The first round of queries. In the first round, since $k = n$, the capacity is always unexhausted, so the returned map satisfies $g_1(i) = \sigma_i(1)$, and the first query to each agent is about their favourite alternative.

Subsequent rounds. For round $\ell \geq 2$, the algorithm uses the answers to the previous queries as the weights, i.e.,

$w_\ell(i) = u_i(g_{\ell-1}(i))$. It computes g_ℓ via a $(n^{1-(\ell-1)/\lambda})$ -capacity serial dictatorship, and makes the ℓ -th set of queries based on g_ℓ by asking each agent i of their utility for $g_\ell(i)$.

The final matching. Finally, after making all the λ queries, the algorithm creates a proxy utility profile $\tilde{u}$ based on the queried utilities. The algorithm sets $\tilde{u}_i(a)$ to the maximum guaranteed utility learnt by either directly asking $u_i(a)$ or an alternative $a' \succ_i a$ ranked below a . The algorithm returns a matching with maximum social welfare w.r.t. $\tilde{u}$.

Minimum welfare guarantee. Let $\widetilde{sw}$ denote the social welfare with respect to the proxy utilities $\tilde{u}$. By construction, $\tilde{u}$ is an underestimation of u , so $sw(M) \geq \widetilde{sw}(M)$ for every matching M . The matching $\widetilde{M}$ returned by Algorithm 2 comes with the following lower bound.

Lemma 4 (Minimum Welfare Guarantee). *For an ordinal matching instance, the matching $\widetilde{M}$ returned by Algorithm 2 achieves*

$$\widetilde{sw}(\widetilde{M}) \geq \frac{1}{n^{1-(\ell-1)/\lambda}} \cdot \sum_{i \in N} u_i(g_\ell(i)), \quad \forall \ell \in [\lambda].$$

The lemma follows from Theorem 3: for each ℓ , the mapping g_ℓ decomposes into $\alpha_\ell = n^{1-(\ell-1)/\lambda}$ matchings whose proxy welfares sum to at least $\sum_{i \in N} u_i(g_\ell(i))$, and the best of them is dominated by the proxy-optimal matching $\widetilde{M}$.

Main result. We now have all the ingredients to prove the main distortion bound.

Theorem 5. *There is a λ -query ordinal matching mechanism that achieves a distortion of $\lambda \cdot n^{1/\lambda}$ and runs in $\text{poly}(n)$ time.*

Proof sketch. Let OPT be the welfare-maximizing matching. At a high level, we partition the agents into λ groups $N_1, \dots, N_\lambda$ and show that $sw(\text{OPT} \mid N_\ell) \leq n^{1/\lambda} \cdot sw(\widetilde{M})$ for each $\ell \in [\lambda]$. Summing over the λ groups yields $sw(\text{OPT}) \leq \lambda \cdot n^{1/\lambda} \cdot sw(\widetilde{M})$, which proves the sought distortion bound.

Partitioning by deviation. Recall that $g_1(i) = \sigma_i(1)$, so no agent strictly prefers OPT to g_1 . For each $\ell \in [\lambda - 1]$, let

$$N_\ell = \{i \in N \setminus (N_1 \cup \dots \cup N_{\ell-1}) \mid \text{OPT}(i) \succ_i g_{\ell+1}(i)\}$$

be the set of agents who “deviate” from $g_{\ell+1}$ to OPT but not from any of $g_1, \dots, g_\ell$. Finally, let $N_\lambda = N \setminus (\bigcup_{\ell \in [\lambda-1]} N_\ell)$ be the remaining agents.

Bounding the welfare of N_ℓ . Fix $\ell \in [\lambda - 1]$. For each $i \in N_\ell$, $g_\ell(i) \succ_i \text{OPT}(i) \succ_i g_{\ell+1}(i)$, so the agents in N_ℓ collectively prefer OPT to every matching in the stable set associated with $g_{\ell+1}$. By Theorem 2, this group has small $w_{\ell+1}$ -weight:

$$sw(\text{OPT} \mid N_\ell) \leq w_{\ell+1}(N_\ell) \leq \frac{w_{\ell+1}(N)}{n^{1-\ell/\lambda}} = \frac{\sum_{i \in N} u_i(g_\ell(i))}{n^{1-\ell/\lambda}},$$

where we used $w_{\ell+1}(i) = u_i(g_\ell(i))$ and $g_\ell(i) \succ_i \text{OPT}(i)$ for the first inequality. Combined with $sw(\widetilde{M}) \geq \widetilde{sw}(\widetilde{M})$ and Theorem 4, we obtain $sw(\text{OPT} \mid N_\ell) \leq n^{1/\lambda} \cdot sw(\widetilde{M})$. The same bound for N_λ follows directly: since $g_\lambda(i) \succ_i \text{OPT}(i)$ for all $i \in N_\lambda$, we have $sw(\text{OPT} \mid N_\lambda) \leq \sum_{i \in N} u_i(g_\lambda(i))$, and applying Theorem 4 for $\ell = \lambda$ gives the bound. $\square$

Implications. For a constant number of queries, i.e., $\lambda = O(1)$, given the $\Omega(n^{1/\lambda})$ lower bound of Amanatidis *et al.* [2021], Theorem 5 settles the optimal distortion using λ queries to be $\Theta(n^{1/\lambda})$. For $\lambda = \log n$, since $n^{1/\log n} = O(1)$, the bound translates to $O(\log n)$ distortion. It is notable that the distortion bound is minimized at $O(\log n)$. Previously, Amanatidis *et al.* [2021] required $O(\log^2 n / \log \log n)$ queries to achieve $O(\log n)$ distortion, and $O(\log^2 n)$ queries to achieve $O(1)$ distortion. Determining the optimal number of queries needed for constant distortion remains an exciting open problem.

4 Single-Winner Elections

We turn briefly to the general social choice setting where the goal is to select one out of m candidates. Existing results were weaker than their counterparts for one-sided matching: the optimal distortion bound was open even for two queries when $m = o(n)$. We settle the question for a constant number of queries.

Theorem 6. *There is a λ -query single-winner voting rule that achieves a distortion of $O(\lambda \cdot (\min\{n, m\})^{1/\lambda})$ and runs in $\text{poly}(n, m)$ time.*

The algorithm closely follows Algorithm 2. The main difference is that, instead of invoking Algorithm 1, the voting rule finds a β -approximately stable committee C_ℓ of size $(\min\{n, m\})^{1-(\ell-1)/\lambda}$ in round ℓ using Jiang *et al.* [2020] (with $\beta = 32 + \epsilon$), and queries each agent of their utility for their favourite candidate in C_ℓ . When λ is a constant, Amanatidis *et al.* [2024] show a matching $\Omega(m^{1/\lambda})$ lower bound, settling the question of optimal distortion in single-winner voting up to a constant factor.

The dependence on $\min\{n, m\}$ rather than m has further consequences. As an instructive example, modeling one-sided matching as a single-winner election with the $n!$ possible matchings as “candidates” yields the bound $O(\lambda \cdot n^{1/\lambda})$ via this generic black-box reduction (recovering Theorem 5 up to a constant). The same reduction also provides nontrivial distortion bounds for query-efficient resource allocation by viewing each allocation as a candidate, even with non-monotone utilities and externalities; see [Ebadian and Shah, 2025] for details.

5 Discussion

We settle the optimal distortion achievable in one-sided matching (and single-winner elections, up to constants) using a fixed number of cardinal value queries per agent. The approaches taken by Mandal *et al.* [2019; 2020] and Kempe [2020], which do not differentiate between ordinal and cardinal queries, and the research assuming that ordinal elicitation is free while cardinal elicitation is very costly, represent more extreme viewpoints. A more balanced approach would involve assigning a cost to ordinal elicitation, with a higher cost for cardinal elicitation, and then optimizing mechanisms within this framework subject to an elicitation budget. Another limitation of our algorithm is its strong dependence between rounds; future research could investigate the trade-offs associated with non-adaptive mechanisms.

Acknowledgements

This research was partially supported by an NSERC Discovery grant.

References

- [Amanatidis *et al.*, 2021] 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.
- [Amanatidis *et al.*, 2022] 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.
- [Amanatidis *et al.*, 2024] 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. *SIAM Journal on Discrete Mathematics*, 38(1):1007–1029, 2024.
- [Anagnostides *et al.*, 2022] Ioannis Anagnostides, Dimitris Fotakis, and Panagiotis Patsilinakos. Metric-distortion bounds under limited information. *Journal of Artificial Intelligence Research*, 74:1449–1483, 2022.
- [Anshelevich and Postl, 2017] Elliot Anshelevich and John Postl. Randomized social choice functions under metric preferences. *Journal of Artificial Intelligence Research*, 58:797–827, 2017.
- [Anshelevich *et al.*, 2018] Elliot Anshelevich, Onkar Bhardwaj, Edith Elkind, John Postl, and Piotr Skowron. Approximating optimal social choice under metric preferences. *Artificial Intelligence*, 264:27–51, 2018.
- [Anshelevich *et al.*, 2021] 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 13th International Joint Conference on Artificial Intelligence Survey Track*, pages 4294–4301, 2021.
- [Aziz *et al.*, 2017] Haris Aziz, Edith Elkind, Piotr Faliszewski, Martin Lackner, and Piotr Skowron. The Condorcet principle for multiwinner elections: from shortlisting to proportionality. In *Proceedings of the 26th International Joint Conference on Artificial Intelligence (IJCAI)*, pages 84–90, 2017.
- [Borodin *et al.*, 2022] Allan Borodin, Daniel Halpern, Mohammad Latifian, and Nisarg Shah. Distortion in voting with top-t preferences. In *Proceedings of the 31st International Joint Conference on Artificial Intelligence (IJCAI)*, page 4, 2022.
- [Boutilier *et al.*, 2015] Craig Boutilier, Ioannis Caragiannis, Simi Haber, Tyler Lu, Ariel D Procaccia, and Or Sheffet. Optimal social choice functions. *Artificial Intelligence*, 227(C):190–213, 2015.
- [Caragiannis and Fehrs, 2024] Ioannis Caragiannis and Karl Fehrs. Beyond the worst case: Distortion in impartial culture electorate. In *Proceedings of the 20th Conference on Web and Internet Economics (WINE)*, pages 420–437, 2024.
- [Caragiannis and Procaccia, 2011] Ioannis Caragiannis and Ariel D Procaccia. Voting almost maximizes social welfare despite limited communication. *Artificial Intelligence*, 175(9-10):1655–1671, 2011.
- [Caragiannis *et al.*, 2017] 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.
- [Cheng *et al.*, 2020] Yu Cheng, Zhihao Jiang, Kamesh Munagala, and Kangning Wang. Group fairness in committee selection. *ACM Transactions on Economics and Computation (TEAC)*, 8(4):1–18, 2020.
- [Ebadian and Shah, 2025] Soroush Ebadian and N. Shah. Every bit helps: Achieving the optimal distortion with a few queries. In *Proceedings of the 39th AAAI Conference on Artificial Intelligence (AAAI)*, pages 13788–13795, 2025.
- [Ebadian *et al.*, 2022] 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*, pages 563–600, 2022.
- [Ebadian *et al.*, 2023] Soroush Ebadian, Mohamad Latifian, and Nisarg Shah. The distortion of approval voting with runoff. In *Proceedings of the 2023 International Conference on Autonomous Agents and Multiagent Systems*, pages 1752–1760, 2023.
- [Ebadian *et al.*, 2024] Soroush Ebadian, Daniel Halpern, and Evi Micha. Metric distortion with elicited pairwise comparisons. In *Proceedings of the 33rd International Joint Conference on Artificial Intelligence (IJCAI)*, pages 2791–2798, 2024.
- [Filos-Ratsikas *et al.*, 2014] Aris Filos-Ratsikas, Søren Kristoffer Stiil Frederiksen, and Jie Zhang. Social welfare in one-sided matchings: Random priority and beyond. In *Proceedings of the 7th Symposium on Algorithmic Game Theory*, pages 1–12, 2014.
- [Hylland and Zeckhauser, 1979] Aanund Hylland and Richard Zeckhauser. The efficient allocation of individuals to positions. *Journal of Political economy*, 87(2):293–314, 1979.
- [Jiang *et al.*, 2020] Zhihao Jiang, Kamesh Munagala, and Kangning Wang. Approximately stable committee selection. In *Proceedings of the 52nd Annual ACM SIGACT Symposium on Theory of Computing*, pages 463–472, 2020.
- [Kempe, 2020] David Kempe. Communication, distortion, and randomness in metric voting. In *Proceedings of the 34th AAAI Conference on Artificial Intelligence (AAAI)*, pages 2087–2094, 2020.

- [Latifian and Voudouris, 2024] Mohamad Latifian and Alexandros A. Voudouris. The distortion of threshold approval matching. In *Proceedings of the Thirty-Third International Joint Conference on Artificial Intelligence, IJCAI-24*, pages 2851–2859, 8 2024. Main Track.
- [Ma *et al.*, 2021] 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-21*, pages 321–327, 8 2021. Main Track.
- [Mandal *et al.*, 2019] Debmalya Mandal, Ariel D Procaccia, Nisarg Shah, and David Woodruff. Efficient and thrifty voting by any means necessary. *Advances in Neural Information Processing Systems*, 32, 2019.
- [Mandal *et al.*, 2020] Debmalya Mandal, Nisarg Shah, and David P Woodruff. Optimal communication-distortion tradeoff in voting. In *Proceedings of the 21st ACM Conference on Economics and Computation*, pages 795–813, 2020.
- [Procaccia and Rosenschein, 2006] Ariel D Procaccia and Jeffrey S Rosenschein. The distortion of cardinal preferences in voting. In *Cooperative Information Agents X: 10th International Workshop, CIA 2006 Edinburgh, UK, September 11-13, 2006 Proceedings 10*, pages 317–331. Springer, 2006.