

Best of Both Distortion Worlds

VASILIS GKATZELIS, Drexel University, USA

MOHAMAD LATIFIAN and NISARG SHAH, University of Toronto, Canada

We study the problem of designing voting rules that take as input the ordinal preferences of n agents over a set of m alternatives and output a single alternative, aiming to optimize the overall happiness of the agents. The input to the voting rule is each agent's ranking of the alternatives from most to least preferred, yet the agents have more refined (cardinal) preferences that capture the intensity with which they prefer one alternative over another. To quantify the extent to which voting rules can optimize over the cardinal preferences given access only to the ordinal ones, prior work has used the *distortion* measure, i.e., the worst-case approximation ratio between a voting rule's performance and the best performance achievable given the cardinal preferences.

The work on the distortion of voting rules has been largely divided into two "worlds": *utilitarian distortion* and *metric distortion*. In the former, the cardinal preferences of the agents correspond to general utilities and the goal is to maximize a normalized social welfare. In the latter, the agents' cardinal preferences correspond to costs given by distances in an underlying metric space and the goal is to minimize the (unnormalized) social cost. Several deterministic and randomized voting rules have been proposed and evaluated for each of these worlds separately, gradually improving the achievable distortion bounds, but none of the known voting rules perform well in both worlds simultaneously.

In this work, we prove that one can in fact achieve the "best of both worlds" by designing new voting rules, both deterministic and randomized, that simultaneously achieve near-optimal distortion guarantees in both distortion worlds. We also prove that this positive result does not generalize to the case where the voting rule is provided with the rankings of only the top- t alternatives of each agent, for $t < m$, and study the extent to which such best-of-both-worlds guarantees can be achieved.

CCS Concepts: • **Theory of computation** → *Approximation algorithms analysis*; • **Computing methodologies** → **Artificial intelligence**.

Additional Key Words and Phrases: voting, distortion, best of both worlds

ACM Reference Format:

Vasilis Gkatzelis, Mohamad Latifian, and Nisarg Shah. 2023. Best of Both Distortion Worlds. In *Proceedings of the 24th ACM Conference on Economics and Computation (EC '23)*, July 9–12, 2023, London, United Kingdom. ACM, New York, NY, USA, 21 pages. <https://doi.org/10.1145/3580507.3597739>

1 INTRODUCTION

A lot of recent work on computational social choice has focused on evaluating the *distortion* of voting rules. Informally, this captures the extent to which they can maximize the social welfare (or minimize the social cost) when making a decision using only limited information regarding the preferences of the agents. In the most fundamental setting, the voting rule is given a set $\mathcal{A}$ of m alternatives and a set $\mathcal{N}$ of n agents, each with their own preferences over the alternatives, and it needs to select an alternative. Even though the preferences of the agents can be complicated, the vast majority of the voting rules used in practice ask each agent to report just their *ranking* over the available alternatives (i.e., their *ordinal* preferences), or even just a prefix that ranks their

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

EC '23, July 9–12, 2023, London, United Kingdom

© 2023 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM ISBN 979-8-4007-0104-7/23/07...\$15.00

<https://doi.org/10.1145/3580507.3597739>

top- t alternatives for some $t < m$ (e.g., in many elections, the agents are asked to select their top choice, second choice, and third choice only). Although this provides very useful guidance for the voting rule, it does not capture the intensity with which an agent may prefer one alternative over another: two agents with the same ranking can be quite different in terms of how much they prefer their top choice over their second one. This limits the ability of the voting rule to optimize cardinal objectives such as the social welfare or the social cost. The distortion measures the ratio between the objective value achieved by the rule on an instance and the best possible objective value in that instance, in the worst case over all instances [Procaccia and Rosenschein, 2006]. Without appropriate modeling assumptions, it is impossible to achieve any bounded distortion using only the ordinal preferences. The distortion literature can be largely divided into two worlds that make different modeling assumptions: *utilitarian* and *metric*.

- The *utilitarian distortion* approach makes no assumptions regarding the relative intensity of the agents' preferences, captured by general von Neumann-Morgenstern utilities. The agents' happiness is measured by their expected utility, and the goal of the voting rule is to maximize the *normalized* social welfare (i.e., to maximize the sum of the agents' expected utilities after normalizing each agent's utility values over different alternatives to add up to 1). This normalization serves two primary purposes: i) it introduces a sense of fairness, by normalizing all the agents' utilities down to the same scale, while maintaining the relative intensity of their preferences, and ii) it enables the voting rule to achieve bounded distortion guarantees that would otherwise be impossible. We refer the reader to the work of Aziz [2020] for several additional justifications for this objective function.

- The *metric distortion* approach focuses on instances where the intensity of the agents' preferences can be captured by a distance function in a metric space, i.e., they obey the triangle inequality. A common motivating example is the classic facility location problem [Anshelevich and Zhu, 2021, Filos-Ratsikas and Voudouris, 2021]: the agents lie in some metric space and a facility needs to be opened at a location within that same metric space. Each agent prefers locations that are closer to them and they rank the potential locations accordingly. Another motivating application comes from ranking political candidates: each voter's preferences over different candidates are determined by their distance on different issues (which, e.g., can be captured as different dimensions). The metric space imposes enough structure on the underlying cardinal preferences of the agents to enable much stronger distortion guarantees whenever it holds, even without any normalization.

Recent work in the distortion literature has analyzed many classic voting rules and introduced several new ones to achieve (near-)optimal distortion bounds for both deterministic and randomized voting rules, focusing either on the utilitarian or on the metric setting. However, none of these voting rules perform well on both settings simultaneously. Therefore, if some designer chooses a voting rule optimized for metric distortion, they should be confident that the metric assumption holds in the setting at hand, otherwise they may get very high distortion. If, on the other hand, they choose a voting rule optimized for utilitarian distortion and the metric assumption happens to hold in their application domain, they may be getting optimal utilitarian distortion guarantees but missing out on much stronger metric distortion guarantees. In fact, it can be non-trivial to detect in advance whether the metric structure holds or not, even for the two main motivating applications discussed above. For example, in facility location the triangle inequality may not apply if the agents care about other features of the potential locations beyond the distance (e.g., proximity to some other facility that the agent often visits, or whether it is on their way to work), or if the agent's cost grows non-linearly in the distance to the location. Similarly, when voting over political candidates, analogous issues may arise if other factors can significantly affect the voters' preferences, e.g., who their parents or friends vote for, or whether they trust some candidate more than another.

In this paper, our goal is to investigate whether there exist (deterministic and randomized) voting rules that simultaneously achieve near-optimal distortion guarantees in both utilitarian and metric worlds, thus providing an obvious choice for someone looking to minimize distortion. Such rules would be perfect for deployment within an automated decision-making tool, providing uninitiated users with the opportunity to input their rankings over some alternatives and find a low-distortion alternative without having to seek expert input regarding which rule to use.

1.1 Our Results

We propose deterministic and randomized voting rules that take agents’ full rankings or top- t preferences as input, return a single alternative, and simultaneously provide near-optimal metric and utilitarian distortion guarantees. Our proposed rules add a safety net by ensuring that even when the metric assumption does not hold, approximately optimal welfare will be achieved with respect to all possible utility functions consistent with the voters’ ordinal preferences.

Full rankings. For deterministic rules, the optimal guarantees we seek are $\Theta(1)$ metric distortion [Gkatzelis et al., 2020] and $\Theta(m^2)$ utilitarian distortion [Caragiannis et al., 2017, Caragiannis and Procaccia, 2011]. For randomized rules, the optimal guarantees we seek are $\Theta(1)$ metric distortion [Anshelevich and Postl, 2017] and $\tilde{\Theta}(\sqrt{m})$ utilitarian distortion [Boutilier et al., 2015, Ebadian et al., 2022b].

We first observe that all the existing voting rules that achieve the desired guarantee for either the utilitarian or the metric world fail to achieve the desired guarantee for the other world.

Our main contribution is to deliver positive news: it is indeed possible to achieve the “best of both worlds” using both deterministic and randomized rules! Specifically, we design two novel voting rules: the deterministic *Pruned Plurality Veto* rule achieves the desired $\Theta(1)$ metric distortion and $\Theta(m^2)$ utilitarian distortion simultaneously, and the randomized *Truncated Harmonic* rule achieves $\Theta(1)$ metric distortion and $\tilde{\Theta}(\sqrt{m})$ utilitarian distortion simultaneously. Curiously, while Pruned Plurality Veto builds on Plurality Veto, which is optimal for metric distortion, Truncated

	Rules	Metric Distortion	Utilitarian Distortion
Deterministic	Plurality	$\Theta(m)$	$\Theta(m^2)$
	Copeland	$\Theta(1)$	Unbounded
	Plurality Matching, Plurality Veto	$\Theta(1)$	$\Omega(n/m)$
	Pruned Plurality Veto*	$\Theta(1)$	$\Theta(m^2)$
Randomized	Random Dictatorship	$\Theta(1)$	$\Theta(m\sqrt{m})$
	Harmonic Rule	Unbounded	$\tilde{\Theta}(\sqrt{m})$
	Stable Lottery Rule	Unbounded	$\Theta(\sqrt{m})$
	Truncated Harmonic Rule*	$\Theta(1)$	$\tilde{\Theta}(\sqrt{m})$

Table 1. Our novel voting rules, (deterministic) Pruned Plurality Veto and (randomized) Truncated Harmonic Rule, achieve asymptotically optimal metric and utilitarian distortions simultaneously (up to a log factor in one case). Existing rules that are asymptotically optimal in one of the two worlds are provided for comparison.

Harmonic builds on the Harmonic Rule, which is optimal for utilitarian distortion. While both rules employ intuitive modifications, the proofs of their distortion guarantees require several new technical lemmas, which may be of independent interest.

Table 1 shows the comparison between our voting rules that achieve asymptotically optimal distortion under both metric and utilitarian worlds simultaneously, and the known voting rules that achieve asymptotically optimal distortion under only one of the two worlds.

Top- t preferences. In stark contrast to the good news with full rankings, we deliver bad news when the input is top- t preferences with small t .

For deterministic rules, the best possible utilitarian distortion for any t is $\Theta(m^2)$ [Caragiannis et al., 2017, Caragiannis and Procaccia, 2011]. We prove that any deterministic rule with even bounded utilitarian distortion must incur $\Theta(m - t + 1)$ metric distortion, as opposed to the optimal $\Theta(m/t)$ metric distortion [Kempe, 2020]. This is sub-optimal whenever t is neither $\Theta(1)$ nor $m - \Theta(1)$.

For randomized rules, the optimal metric distortion is $\Theta(1)$ for any $t \geq 1$. We prove that $\Theta(1)$ metric distortion implies $\Omega(\max((m/t)^{1.5}, \sqrt{m}))$ utilitarian distortion, which is strictly worse than the optimal $\Theta(\max(m/t, \sqrt{m}))$ utilitarian distortion [Borodin et al., 2022] for all $t \in o(m^{2/3})$.

Our results chart a more general trade-off between metric and utilitarian distortions, and we complement our negative results with novel constructive upper bounds.

1.2 Related Work

There is a huge body of work on distortion, from both metric and utilitarian points of view; see the recent survey of Anshelevich et al. [2021] for a detailed summary. Procaccia and Rosenschein [2006] introduced the notion of distortion in the utilitarian setting. Caragiannis and Procaccia [2011] proved that the Plurality rule has $\Theta(m^2)$ utilitarian distortion, which later proved to be optimal over all deterministic rules [Caragiannis et al., 2017]. Boutilier et al. [2015] studied the utilitarian distortion of randomized voting rules and proposed the Harmonic Rule, which we use in our work. They proved that its distortion is $O(\sqrt{mH_m})$ and designed a more complex rule with distortion $O(\sqrt{m} \log^*(m))$, also proving a lower bound of $\Omega(\sqrt{m})$. This gap was closed by Ebadian et al. [2022b], who achieved the optimal $\Theta(\sqrt{m})$ distortion by introducing the Stable Lottery Rule.

On the other end, Anshelevich et al. [2018] initiated the study of distortion in the metric setting. They proved that the Copeland rule has a metric distortion of 5. They proved a lower bound of 3 on the metric distortion of any deterministic voting rule and conjectured that this lower bound is tight. Munagala and Wang [2019] improved the upper bound to 4.236, and Gkatzelis et al. [2020] finally resolved the conjecture by achieving a metric distortion of 3 by introducing the Plurality Matching rule. While the Plurality Matching rule is rather complicated, Kizilkaya and Kempe [2022] subsequently introduced a simpler refinement, Plurality Veto, which still achieves a metric distortion of 3. Very recent work has provided an even deeper understanding of the class of deterministic rules achieving a metric distortion of 3 [Kizilkaya and Kempe, 2023, Peters, 2023].

Even though a metric distortion better than 3 is not achievable by deterministic rules, Anshelevich and Postl [2017] proved that randomization can break this barrier. They proved that Random Dictatorship has metric distortion slightly better than 3 and provide a lower bound of 2 on the metric distortion of any randomized rule. Several works have since tried to close this gap, with Charikar and Ramakrishnan [2022] recently improving the lower bound to 2.0261. Nonetheless, the optimal metric distortion of randomized rules is still $\Theta(1)$.

While the aforementioned work considers full rankings as input, there is significant interest in analyzing the distortion achievable using other types of information. Amanatidis et al. [2021] considered eliciting restricted cardinal information regarding the agents' preferences in addition to

their ordinal preferences. In contrast, some work considers eliciting *less* information than ordinal preferences, focusing on top- t preferences. In the metric setting, Kempe [2020] proved a lower bound of $\frac{2m-t}{t}$ and an upper bound of $\frac{79m}{t}$ on the best metric distortion that any deterministic voting rule can achieve using top- t preferences; the upper bound was later improved to $6m/t + 1$ by Anagnostides et al. [2022]. While the exact bound is yet to be identified, this pins down the asymptotic bound as $\Theta(m/t)$. For randomized rules, the best bound is still $\Theta(1)$ for any t , achieved by Random Dictatorship using only top-1 preferences (i.e., plurality votes). In the utilitarian setting, Borodin et al. [2022] proved that the optimal utilitarian distortion achievable by randomized rules using top- t preferences is $\Theta(\max(m/t, \sqrt{m}))$. For deterministic rules, it is $\Theta(m^2)$ for any t because Plurality achieves it using top-1 preferences and is optimal even within rules using full rankings.

The distortion framework has also been extended to study many other problems such as the elicitation-distortion tradeoff [Kempe, 2020, Mandal et al., 2019, 2020], more complex elections [Benade et al., 2021, Borodin et al., 2019, Caragiannis et al., 2022, Ebadian et al., 2023b], and even problems beyond voting [Amanatidis et al., 2022, Anshelevich and Sekar, 2016, Ebadian et al., 2022a, Halpern and Shah, 2021].

2 PRELIMINARIES

Let $\mathcal{N}$ be a set of n agents and $\mathcal{A}$ be a set of m alternatives. Throughout the paper, agents are labeled i, j and alternatives are labeled X, Y, W . We use $\mathcal{L}(\mathcal{A})$ to denote the set of all rankings (linear orders) over the alternatives, and $\Delta(\mathcal{A})$ to denote the set of all distributions over $\mathcal{A}$.

Preference profile. Each agent $i \in \mathcal{N}$ has a preference ranking $\sigma_i \in \mathcal{L}(\mathcal{A})$ over the alternatives. Collectively, they form a preference profile $\vec{\sigma} = (\sigma_1, \sigma_2, \dots, \sigma_n)$. We use $r_i(X)$ to denote the rank of alternative X in agent i 's ranking, and say agent i prefers alternative X to alternative Y if $r_i(X) < r_i(Y)$, which we also denote as $X \succ_i Y$. The plurality score $\text{plu}(X, \vec{\sigma})$ of alternative X under $\vec{\sigma}$ is the number of agents who rank X first. We use $\mathcal{N}^X$ to denote the set of such agents.

Voting Rule. A (randomized) voting rule $f : \mathcal{L}(\mathcal{A})^n \rightarrow \Delta(\mathcal{A})$ is a function that takes as input a preference profile $\vec{\sigma}$ and outputs a probability distribution $f(\vec{\sigma})$ over alternatives. We call f deterministic if it always outputs a single alternative with probability 1; in this case, we slightly abuse the notation to let $f(\vec{\sigma})$ denote this alternative. The Plurality rule selects an alternative with the highest plurality score.

We assume that the preference profile stems from underlying cardinal preferences of the agents over the alternatives. We consider two frameworks for modeling these cardinal preferences. In this paper, our goal is to design voting rules that perform well under both frameworks simultaneously and prove corresponding lower bounds on such “best of both worlds” guarantees.

2.1 Metric Framework

In the metric framework, we assume that all the agents and alternatives are embedded in a (pseudo) metric space identified by a distance function $d : (\mathcal{N} \cup \mathcal{A})^2 \rightarrow \mathbb{R}_{\geq 0}$ satisfying the following.

- $d(a, a) = 0$ for all $a \in \mathcal{N} \cup \mathcal{A}$.
- Symmetry: $d(a, b) = d(b, a)$, for all $a, b \in \mathcal{N} \cup \mathcal{A}$.
- Triangle inequality: $d(a, b) \leq d(a, c) + d(c, b)$ for all $a, b, c \in \mathcal{N} \cup \mathcal{A}$.

We say that preference profile $\vec{\sigma}$ is consistent with metric space d , denoted as $d \triangleright_{\mathcal{M}} \vec{\sigma}$, if, for every $i \in \mathcal{N}$ and $X, Y \in \mathcal{A}$, we have $X \succ_i Y \Rightarrow d(i, X) \leq d(i, Y)$. In this framework, we capture the quality of each alternative by her total cost to the agents: define $\text{SC}(X, d) = \sum_{i \in \mathcal{N}} d(i, X)$ be the social cost of alternative X in metric space d .

Metric Distortion. In this setting, the rule that selects an alternative with a lower social cost is generally considered more efficient. Formally, we define the metric distortion of a distribution over alternatives $p \in \Delta(\mathcal{A})$ in metric space d as the ratio between the expected social cost of an alternative sampled from p to the minimum possible social cost, i.e.,

$$\text{dist}^M(p, d) = \frac{\mathbb{E}_{X \sim p}[\text{SC}(X, d)]}{\min_{X \in \mathcal{A}} \text{SC}(X, d)}.$$

The metric distortion of a voting rule f on preference profile $\vec{\sigma}$ is the worst-case metric distortion of $f(\vec{\sigma})$ on any metric $d \triangleright_M \vec{\sigma}$: $\text{dist}^M(f, \vec{\sigma}) = \sup_{d: d \triangleright_M \vec{\sigma}} \text{dist}^M(f(\vec{\sigma}), d)$. Finally, the metric distortion of f is obtained by taking the worst case over all preference profiles: $\text{dist}^M(f) = \max_{\vec{\sigma} \in \mathcal{L}(\mathcal{A})^n} \text{dist}^M(f, \vec{\sigma})$.

2.2 Utilitarian Framework

The utilitarian framework assumes that each agent $i \in \mathcal{N}$ has a utility function $u_i : \mathcal{A} \rightarrow \mathbb{R}_{\geq 0}$ over the alternatives. Collectively, they form the utility profile $\vec{u} = (u_1, u_2, \dots, u_n)$. Following the literature [Aziz, 2020], we assume that the utility functions are unit-sum, which means that for each agent i , we have $\sum_{X \in \mathcal{A}} u_i(X) = 1$. We say that preference profile $\vec{\sigma}$ is consistent with utility profile $\vec{u}$, denoted $\vec{u} \triangleright_U \vec{\sigma}$, if for each $i \in \mathcal{N}$ and $X, Y \in \mathcal{A}$, we have $X \succ_i Y \Rightarrow u_i(X) \geq u_i(Y)$.

Here, the quality of each alternative is determined based on the sum of the utilities of the agents for it: define the social welfare of alternative X in utility profile $\vec{u}$ as $\text{SW}(X, \vec{u}) = \sum_{i \in \mathcal{N}} u_i(X)$. We remark that studying the social welfare with unit-sum utilities is equivalent to studying the *normalized* social welfare with unrestricted utilities, which is the formulation described in the introduction. We use the former to be consistent with the literature. Similarly to the metric case, we can define distortion in the utilitarian framework.

Utilitarian Distortion. The utilitarian distortion of distribution over alternatives $p \in \Delta(\mathcal{A})$ in utility profile $\vec{u}$ is defined as:

$$\text{dist}^U(p, \vec{u}) = \frac{\max_{X \in \mathcal{A}} \text{SW}(X, \vec{u})}{\mathbb{E}_{X \sim p}[\text{SW}(X, \vec{u})]}.$$

The utilitarian distortion of voting rule f on preference profile $\vec{\sigma}$ is the worst-case distortion of $f(\vec{\sigma})$ on any utility profile $\vec{u} \triangleright_U \vec{\sigma}$: $\text{dist}^U(f, \vec{\sigma}) = \sup_{\vec{u}: \vec{u} \triangleright_U \vec{\sigma}} \text{dist}^U(f(\vec{\sigma}), \vec{u})$. Finally, the utilitarian distortion of f is obtained by taking the worst case over all preference profiles: $\text{dist}^U(f) = \max_{\vec{\sigma} \in \mathcal{L}(\mathcal{A})^n} \text{dist}^U(f, \vec{\sigma})$.

3 A DETERMINISTIC VOTING RULE

Our goal in this section is to check if there exists a deterministic voting rule that achieves near-optimal distortion guarantees in both metric and utilitarian frameworks simultaneously. Let us first review the state-of-the-art guarantees in each framework. For deterministic rules, it is known that the best possible metric distortion is 3 [Anshelevich et al., 2018, Gkatzelis et al., 2020], which can be achieved using Plurality Matching [Gkatzelis et al., 2020] or its refinement Plurality Veto [Kizilkaya and Kempe, 2022], and the best possible utilitarian distortion is $O(m^2)$, which can be achieved by Plurality [Caragiannis et al., 2017, Caragiannis and Procaccia, 2011].

While the aforementioned rules achieve the optimal distortion in one framework, we remark that they incur a terrible distortion in the other framework, thus failing to provide our desired best-of-both-worlds guarantee. Plurality Matching (and its refinement Plurality Veto) incurs a utilitarian distortion of $\Omega(n/m)$ (as opposed to the desired $O(m^2)$) and Plurality incurs a metric distortion of $\Omega(m)$ (as opposed to the desired $O(1)$).

Because Plurality Matching and Plurality Veto are not central to our work, we refer an interested reader to the works of Gkatzelis et al. [2020] and Kizilkaya and Kempe [2022] for their definitions.

Proposition 1. For $m \geq 3$, the utilitarian distortion of Plurality Matching (and consequently, its refinement Plurality Veto) is $\Omega(n/m)$.

PROOF. Let $\mathcal{A} = \{X^*, X_1, X_2, \dots, X_{m-1}\}$. Consider a preference profile $\vec{\sigma}$ in which, for $k \in [m-1]$, $n/(m-1) - 1$ agents rank X_k first, X^* second, and X_{k+1} last (with X_m redefined as X_1 for cyclicity) and one agent ranks X^* first and X_k last. The rest of the preference profile can be completed arbitrarily.

Gkatzelis et al. [2020] prove that an alternative can only be returned by Plurality Matching if it is ranked first at least as many times as it is ranked last. Note that every alternative except X^* appears last once more than it appears first. Hence, Plurality Matching and its refinement Plurality Veto must output X^* .

Consider a utility profile $\vec{u} \succ_{\cup} \vec{\sigma}$ in which every agent who ranks X^* first has utility $1/m$ for all alternatives, and every other agent has utility 1 for her top choice. Then, $\text{SW}(X^*, \vec{u}) = O(1)$ whereas $\text{SW}(X_1, \vec{u}) = \Omega(n/m)$, yielding the desired utilitarian distortion lower bound. $\square$

Proposition 2 ([Anshelevich et al., 2018]). The metric distortion of the Plurality rule is $2m - 1$.

This leaves open the question of whether there exists a deterministic rule that simultaneously achieves $O(1)$ metric distortion and $O(m^2)$ utilitarian distortion.

3.1 Pruned Plurality Veto

We settle this question positively by designing a new deterministic voting rule, *Pruned Plurality Veto*, and proving that it achieves the desired best-of-both-worlds guarantee. We remark that our rule does not critically rely on using Plurality Veto; any deterministic voting rule with $O(1)$ metric distortion can be used instead; the metric distortion of our rule will be a constant times greater than the metric distortion of this base rule. We use Plurality Veto because it is a simple voting rule that achieves the optimal metric distortion guarantee of 3.

Definition 1 (Pruned Plurality Veto (PPV)). Given any $\varepsilon > 0$ and a preference profile $\vec{\sigma}$, the (deterministic) Pruned Plurality Veto Rule f^{PPV} computes the set of alternatives $\mathcal{A}'$ with plurality score at least $\frac{\varepsilon n}{(6+\varepsilon)m}$ (note that $\mathcal{A}'$ must be non-empty by the Pigeonhole principle), restricts $\vec{\sigma}$ to the alternatives in $\mathcal{A}'$ to obtain preference profile $\vec{\sigma}'$, and runs Plurality Veto on $\vec{\sigma}'$.

The intuition behind the rule is simple. The argument used by Caragiannis et al. [2017] for establishing $O(m^2)$ utilitarian distortion of Plurality continues to hold when selecting any alternative with plurality score $\Omega(n/m)$. Hence, limiting our attention to such alternatives easily takes care of the utilitarian distortion guarantee. The trick is to prove that simply running a voting rule with constant metric distortion on the preference profile restricted to such alternatives still yields (slightly higher) constant metric distortion. For this, we need two lemmas.

The first one is due to Anshelevich et al. [2018]. We provide the short proof for completeness.

Lemma 1. For any metric space d , alternatives X and Y , and agent i who prefers Y to X (i.e., $Y \succ_i X$), we have $d(i, X) \geq d(Y, X)/2$.

PROOF. Note that

$$d(Y, X) \leq d(i, Y) + d(i, X) \leq 2d(i, X),$$

where the last inequality holds due to $Y \succ_i X$. $\square$

The next lemma, which is the crux of our proof, shows that ignoring alternatives with low plurality score and running a voting rule on the remaining alternatives can only increase its metric distortion by a small factor. This lemma may be of independent interest.

Lemma 2. Consider any preference profile $\vec{\sigma}$ and metric $d \triangleright_M \vec{\sigma}$. For some threshold $\tau < 1/m$, let $\mathcal{A}' \subseteq \mathcal{A}$ be the set of alternatives with plurality score at least τn . Then,

$$\frac{\min_{X \in \mathcal{A}'} \text{SC}(X, d)}{\min_{X \in \mathcal{A}} \text{SC}(X, d)} \leq 1 + \frac{2}{1 - \tau m}.$$

Consequently, for a constant d^M , applying a voting rule with a metric distortion of at most d^M on the preference profile obtained by restricting $\vec{\sigma}$ to $\mathcal{A}'$ incurs a metric distortion of at most $d^M \cdot (1 + 2/(1 - \tau m))$.

PROOF. Let $X^* \in \arg \min_{X \in \mathcal{A}} \text{SC}(X, d)$ be an optimal alternative and $X^+ \in \arg \min_{X \in \mathcal{A}'} d(X, X^*)$ be the closest alternative in $\mathcal{A}'$ to X^* . Recall that for an alternative Y , $\mathcal{A}^Y$ denotes the set of agents who rank Y first. For each alternative $Y \notin \mathcal{A}'$, $\text{plu}(Y, \vec{\sigma}) \leq \tau n$. Hence, $\sum_{Y \notin \mathcal{A}'} |\mathcal{A}^Y| \leq \tau mn$, implying that $\sum_{Y \in \mathcal{A}'} |\mathcal{A}^Y| \geq n \cdot (1 - \tau m)$.

We now prove a lower bound on the optimal social cost.

$$\begin{aligned} \text{SC}(X^*, d) &= \sum_{i \in N} d(i, X^*) \geq \sum_{Y \in \mathcal{A}'} \sum_{i \in \mathcal{A}^Y} d(i, X^*) \\ &\geq \sum_{Y \in \mathcal{A}'} \sum_{i \in \mathcal{A}^Y} d(X^*, Y)/2 && \text{(By Lemma 1)} \\ &\geq \sum_{Y \in \mathcal{A}'} |\mathcal{A}^Y| \cdot d(X^*, X^+)/2 && (\because X^+ \text{ is the closest member of } \mathcal{A}' \text{ to } X^*) \\ &\geq n \cdot (1 - \tau m) \cdot d(X^*, X^+)/2. \end{aligned}$$

On the other hand, using the triangle inequality for each agent and adding across agents, we get

$$\text{SC}(X^+, d) \leq \text{SC}(X^*, d) + n \cdot d(X^*, X^+).$$

Let $X' \in \arg \min_{X \in \mathcal{A}'} \text{SC}(X, d)$. Then,

$$\begin{aligned} \frac{\text{SC}(X', d)}{\text{SC}(X^*, d)} &\leq \frac{\text{SC}(X^+, d)}{\text{SC}(X^*, d)} && (\because X' \in \arg \min_{X \in \mathcal{A}'} \text{SC}(X, d)) \\ &\leq \frac{\text{SC}(X^*, d) + n \cdot d(X^*, X^+)}{\text{SC}(X^*, d)} \\ &\leq 1 + \frac{n \cdot d(X^*, X^+)}{n \cdot (1 - \tau m) \cdot d(X^*, X^+)/2} = 1 + \frac{2}{1 - \tau m}. \end{aligned}$$

Next, let d^M be a constant and f be a voting rule with a metric distortion of at most d^M . Let $\vec{\sigma}'$ be the preference profile restricted to alternatives in $\mathcal{A}'$ and $X = f(\vec{\sigma}')$. Then, we have

$$\text{SC}(X, d) \leq d^M \cdot \text{SC}(X', d) \leq d^M \cdot \left(1 + \frac{2}{1 - \tau m}\right) \cdot \text{SC}(X^*, d),$$

where the first transition is due to the metric distortion guarantee of f and the second transition is proven above. $\square$

We are now ready to prove the distortion guarantees of Pruned Plurality Veto.

Theorem 1. Given any $\varepsilon > 0$, Pruned Plurality Veto (PPV) with parameter ε has a metric distortion of $9 + \varepsilon$ and a utilitarian distortion of $O(m^2/\varepsilon)$.

PROOF. Let $\vec{\sigma}$ be any preference profile. Let $\mathcal{A}'$ be the set of alternatives with plurality score at least $\frac{\varepsilon n}{(6+\varepsilon)m}$, and $\bar{\mathcal{A}} = \mathcal{A} \setminus \mathcal{A}'$. Note that PPV always returns an alternative from $\mathcal{A}'$.

This is sufficient for the utilitarian distortion guarantee, since any alternative in $\mathcal{A}'$ is the top choice of at least $\frac{\varepsilon n}{(6+\varepsilon)m}$ agents, and each agent has utility at least $1/m$ for her top alternative due to the Pigeonhole principle. Hence, its social welfare is at least $\frac{\varepsilon n}{(6+\varepsilon)m^2}$. Due to unit-sum utilities, the maximum social welfare of any alternative is at most n . Hence, the utilitarian distortion of PPV is at most $\frac{(6+\varepsilon)m^2}{\varepsilon}$.

For the metric distortion, we simply need to apply Lemma 2 with $d^M = 3$ (the metric distortion of Plurality Veto) and $\tau = \frac{\varepsilon}{(6+\varepsilon)m}$. We get that the metric distortion of PPV is at most

$$d^M \cdot \left(1 + \frac{2}{1 - \tau m}\right) = 3 \cdot \left(1 + \frac{2}{1 - \frac{\varepsilon}{6+\varepsilon}}\right) = 3 \cdot \left(3 + \frac{\varepsilon}{3}\right) = 9 + \varepsilon. \quad \square$$

While Pruned Plurality Veto achieves asymptotically optimal distortion in both metric and utilitarian frameworks simultaneously, it does not achieve metric distortion better than 9, when the optimal metric distortion is 3. In the metric framework, since the optimal distortion is already a constant, it is common to care about what the exact constant is. Indeed, a metric distortion of 5 was established through Copeland's rule already in the seminal paper that introduced this framework [Anshelevich et al., 2018]. Reducing this to 3 was an open question that was settled recently [Gkatzelis et al., 2020].

This might lead one to wonder whether one can obtain a tighter best-of-both-worlds guarantee with a utilitarian distortion of $O(m^2)$ and a metric distortion of at most 3. We show that this is not possible and leave open the question of improving upon the metric distortion of $9 + \varepsilon$ in Theorem 1.

Theorem 2. Any deterministic rule with a utilitarian distortion of $o(m^2\sqrt{m})$ has a metric distortion strictly greater than 3.

PROOF. Let X be a specific alternative and partition the remaining alternatives into 3 sets, $\mathcal{A}_1$, $\mathcal{A}_2$, and $\mathcal{A}_3$, each with $(m-1)/3$ alternatives. Furthermore, partition the agents into 3 sets, $\mathcal{N}_1$ and $\mathcal{N}_2$ with $n(1-1/\sqrt{m})/2$ agents each, and $\mathcal{N}_3$ with $n/\sqrt{m}$ agents.

Consider the following preference profile $\vec{\sigma}$. For $k \in [3]$, each alternative in $\mathcal{A}_k$ appears as the top choice of $3/(m-1)$ fraction of the agents in $\mathcal{N}_i$, and X appears as the second choice of every agent. Each agent prefers every alternative in $\mathcal{A}_3$ to any alternative in $\mathcal{A}_1 \cup \mathcal{A}_2$ other than her top choice. Agents in $\mathcal{N}_1$ prefer every alternative in $\mathcal{A}_1$ to any alternative in $\mathcal{A}_2$, and agents in $\mathcal{N}_2$ prefer every alternative in $\mathcal{A}_2$ to any alternative in $\mathcal{A}_1$.

Suppose f is a deterministic rule with a utilitarian distortion of $o(m^2\sqrt{m})$. Let $f(\vec{\sigma}) = X^*$. First, we prove that $X^* \in \mathcal{A}_1 \cup \mathcal{A}_2$ on $\vec{\sigma}$.

If $X^* = X$, consider the consistent utility profile in which each agent has utility 1 for her top choice and zero for the rest. In this case, the utilitarian distortion of X is unbounded, which is a contradiction. If $X^* \in \mathcal{A}_3$, consider the consistent utility profile $\vec{u}$ in which agents who have X^* as their top choice have utility $1/m$ for every alternative and the other agents have utility $1/2$ for their top choice and $1/2$ for X . Then, $\text{SW}(X^*, \vec{u}) = O(n/(m^2\sqrt{m}))$ whereas $\text{SW}(X, \vec{u}) = \Omega(1)$, yielding a utilitarian distortion of $\Omega(m^2\sqrt{m})$, which is again a contradiction. Hence, we have $X^* \in \mathcal{A}_1 \cup \mathcal{A}_2$.

Without loss of generality, assume $X^* \in \mathcal{A}_1$. Consider a consistent one-dimensional distance metric d over $\mathbb{R}$ in which alternatives in $\mathcal{A}_1$ lie at 0, agents in $\mathcal{N}_1$ lie at 1, and the remaining agents

and alternatives lie at 2. Thus, the distances are as follows: for all agents i and alternatives Y ,

$$d(i, Y) = \begin{cases} 1 & i \in \mathcal{N}_1 \\ 2 & i \notin \mathcal{N}_1, Y \in \mathcal{A}_1 \\ 0 & i \notin \mathcal{N}_1, Y \notin \mathcal{A}_1. \end{cases}$$

Now, for $X^* \in \mathcal{A}_1$, we have:

$$SC(X^*, d) = \frac{n}{2} \left(1 - \frac{1}{\sqrt{m}}\right) + 2 \cdot \left(n - \frac{n}{2} \cdot \left(1 - \frac{1}{\sqrt{m}}\right)\right) = n \cdot \frac{3\sqrt{m} + 1}{2\sqrt{m}}.$$

On the other hand,

$$SC(X, d) = 1 \cdot \frac{n}{2} \cdot \left(1 - \frac{1}{\sqrt{m}}\right) = n \frac{\sqrt{m} - 1}{2\sqrt{m}}.$$

Thus, we have

$$\text{dist}^M(f, \vec{\sigma}) \geq \frac{SC(X^*, d)}{SC(X, d)} = \frac{3\sqrt{m} + 1}{\sqrt{m} - 1} > 3. \quad \square$$

4 A RANDOMIZED VOTING RULE

Next, we turn to randomized voting rules. In the metric world, the best distortion achievable by a randomized rule is known to be in $[2.0261, 3)$ [Anshelevich and Postl, 2017, Charikar and Ramakrishnan, 2022], and closing that gap remains a major open question. There exist randomized voting rules that guarantee a distortion slightly better than 3 (but their distortion converges to 3 as the number of agents or alternatives increases) [Anshelevich and Postl, 2017, Gkatzelis et al., 2020, Kempe, 2020], and notable among them is the Random Dictatorship, which returns the top alternative of an agent selected uniformly at random. In the utilitarian world, the best achievable distortion by a randomized rule was recently shown to be $\Theta(\sqrt{m})$ [Ebadian et al., 2022b] using the Stable Lottery Rule, improving upon a previous bound of $\Theta(\sqrt{m \log m})$ via the Harmonic Rule.

Like in the case of deterministic rules, one might wonder whether these rules, which achieve near-optimal distortion in one framework, also achieve near-optimal distortion in the other framework, thus achieving the best of both worlds. Unfortunately, this is yet again not true.

We omit defining the Stable Lottery Rule and the Harmonic Rule, but it suffices to know that they assign a positive probability to every alternative. In a preference profile where an alternative X is ranked last by every agent, it could have arbitrarily large social cost, and hence, assigning any positive probability to it yields unbounded metric distortion.

Proposition 3. The Stable Lottery Rule and the Harmonic Rule have unbounded metric distortion.

In concurrent work by a subset of the authors [Ebadian et al., 2023a], it has been shown that Random Dictatorship has $\Theta(m^{1.5})$ utilitarian distortion, which is significantly worse than the optimal bound of $\Theta(\sqrt{m})$.

Proposition 4 (Ebadian et al. 2023a). Random Dictatorship has $\Theta(m^{1.5})$ utilitarian distortion.

This leaves open the question of whether there exists a randomized rule that simultaneously achieves $\tilde{O}(\sqrt{m})$ utilitarian distortion and $O(1)$ metric distortion.

Our proposed randomized voting rule, Truncated Harmonic, selects an agent uniformly at random and then selects an alternative with probability that is inversely proportional to its rank within the selected agent's preferences (i.e., drops harmonically), but with one caveat: it assigns higher probability to the alternative $\hat{X}$ returned by the deterministic Plurality Veto voting rule and 0 probability to all alternatives that the agent ranks below $\hat{X}$.

Definition 2 (Truncated Harmonic voting rule). The Truncated Harmonic voting rule, $f_\varepsilon^{\text{TH}}$, is parameterized by a positive constant ε and uses the following steps:

- Let $\widehat{X}$ be the alternative that would be chosen by Plurality Veto
- Select an agent $i \in \mathcal{N}$ uniformly at random
- Return an alternative $Y \in \mathcal{A}$ with probability:

$$p(i, Y) = \begin{cases} \frac{\varepsilon}{6H_m r_i(Y)} & \text{if } Y \succ_i \widehat{X}, \\ 1 - \sum_{Y \succ_i \widehat{X}} p(i, Y) & \text{if } Y = \widehat{X}, \\ 0 & \text{otherwise.} \end{cases}$$

Our main result for this section show that the Truncated Harmonic mechanism can essentially achieve the best of both worlds:

Theorem 3. The Truncated Harmonic voting rule simultaneously guarantees a metric distortion of $3 + \varepsilon$ and a utilitarian distortion of $O(\sqrt{m}H_m/\varepsilon)$.

We provide the proof of this theorem in the following two subsections, first for the metric world and then for the utilitarian one.

4.1 Metric Distortion Guarantees

We start off by defining the class of X -truncated weight functions and proving a very useful lemma that extends Lemma 3 from [Anshelevich and Postl, 2017] to arbitrary X -truncated weight functions (the previous result applied only to weight functions that assigned all their weight to the top alternative of each agent).

Definition 3. Given an alternative $X \in \mathcal{A}$ and a preference profile $\vec{\sigma}$, we call a function $w : \mathcal{N} \times \mathcal{A} \rightarrow [0, 1]$ an X -truncated weight function (X -TWF) if for each agent $i \in \mathcal{N}$ it satisfies:

- For all $Y \in \mathcal{A}$ such that $X \succ_i Y$, the function assigns zero weight, i.e., $w(i, Y) = 0$,
- and for all other $Y \in \mathcal{A}$ the weight adds up to one: $\sum_{Y \succeq_i X} w(i, Y) = 1$.

For such a function we let $w^+(Y) = \sum_{i \in \mathcal{N}} w(i, Y)$ denote the total weight assigned to $Y \in \mathcal{A}$.

Lemma 3. Given an alternative $X \in \mathcal{A}$, a preference profile $\vec{\sigma}$, and any X -truncated weight function w_X , we have $\text{SC}(X, d) \geq \frac{1}{2} \sum_{Y \in \mathcal{A}} w_X^+(Y) d(X, Y)$ for any metric $d \triangleright_M \vec{\sigma}$.

PROOF.

$$\begin{aligned} \text{SC}(X, d) &= \sum_{i \in \mathcal{N}} d(i, X) = \sum_{i \in \mathcal{N}} \sum_{\substack{Y \in \mathcal{A}: \\ Y \succeq_i X}} w_X(i, Y) d(i, X) && (\because \forall i \in \mathcal{N}: \sum_{\substack{Y \in \mathcal{A}: \\ Y \succeq_i X}} w_X(i, Y) = 1) \\ &\geq \sum_{i \in \mathcal{N}} \sum_{\substack{Y \in \mathcal{A}: \\ Y \succeq_i X}} w_X(i, Y) d(X, Y) / 2 && (\text{by Lemma 1 and } Y \succeq_i X) \\ &= \sum_{i \in \mathcal{N}} \sum_{Y \in \mathcal{A}} w_X(i, Y) d(X, Y) / 2 && (\because w(i, Y) = 0 \text{ when } X \succ_i Y) \\ &= \sum_{Y \in \mathcal{A}} \sum_{i \in \mathcal{N}} w_X(i, Y) d(X, Y) / 2 \\ &= \sum_{Y \in \mathcal{A}} w_X^+(Y) d(X, Y) / 2. && \square \end{aligned}$$

The next lemma generalizes Lemma 4 from [Anshelevich and Postl, 2017] which is restricted to the special case where the weight function assigns all its weight to each agent's top alternative.

Lemma 4. Consider any preference profile $\vec{\sigma}$, any metric $d \triangleright_M \vec{\sigma}$, any alternative $X \in \mathcal{A}$, and any X -truncated weight function w_X . Then, if the metric distortion of X according to d is d^M , for any voting rule f that assigns probability $p(Y)$ to each alternative $Y \in \mathcal{A}$, we have

$$\text{dist}^M(f(\vec{\sigma}), d) \leq d^M \left(1 + \frac{2n \sum_{Y \in \mathcal{A}} p(Y) d(X, Y)}{\sum_{Y \in \mathcal{A}} w_X^+(Y) d(X, Y)} \right).$$

PROOF. If X^* is the optimal alternative in the underlying metric space, we have:

$$\begin{aligned} \text{dist}^M(f(\vec{\sigma}), d) &= \frac{\mathbb{E}_{Y \sim f(\vec{\sigma})} [\text{SC}(Y, d)]}{\text{SC}(X^*, d)} \\ &\leq \frac{d^M \mathbb{E}_{Y \sim f(\vec{\sigma})} [\text{SC}(Y, d)]}{\text{SC}(X, d)} && (\text{since } \text{SC}(X, d) \leq d^M \cdot \text{SC}(X^*, d)) \\ &= \frac{d^M \sum_{Y \in \mathcal{A}} p(Y) \text{SC}(Y, d)}{\text{SC}(X, d)} \\ &\leq \frac{d^M \sum_{Y \in \mathcal{A}} p(Y) (\text{SC}(X, d) + n \cdot d(X, Y))}{\text{SC}(X, d)} && (\text{by triangle inequality}) \\ &= d^M \left(1 + \frac{n \sum_{Y \in \mathcal{A}} p(Y) d(X, Y)}{\text{SC}(X, d)} \right) && (\text{since } \sum_{Y \in \mathcal{A}} p(Y) = 1) \\ &\leq d^M \left(1 + \frac{n \sum_{Y \in \mathcal{A}} p(Y) d(X, Y)}{\sum_{Y \in \mathcal{A}} \frac{1}{2} w_X^+(Y) d(X, Y)} \right) && (\text{by Lemma 3}) \\ &= d^M \left(1 + \frac{2n \sum_{Y \in \mathcal{A}} p(Y) d(X, Y)}{\sum_{Y \in \mathcal{A}} w_X^+(Y) d(X, Y)} \right). \quad \square \end{aligned}$$

Corollary 1. The metric distortion of the Truncated Harmonic voting rule is $3 + \varepsilon$.

PROOF. First, let us define $\widehat{X}$ -truncated weight function $\widehat{w}$ as follows:

$$\widehat{w}(i, Y) = \begin{cases} \frac{1}{H_m r_i(Y)} & Y \succ_i X, \\ 1 - \sum_{Y \succ_i X} \widehat{w}(i, Y) & Y = X, \\ 0 & o.w. \end{cases}$$

Note that by the definition of $f_\varepsilon^{\text{TH}}$, we have $p(i, Y) = \frac{\varepsilon}{6} \widehat{w}(i, Y)$ for $Y \neq \widehat{X}$. In addition, for any metric distance d the alternative $\widehat{X}$ used by the Truncated Harmonic voting rule to define its truncated weight function has metric distortion at most 3. Therefore, Lemma 4 implies that the distortion of the alternative chosen by this voting rule for any distance d , inducing preference profile $\vec{\sigma}$, and using $\widehat{w}$ as the truncated weight function, is:

$$\text{dist}^M(f_\varepsilon^{\text{TH}}(\vec{\sigma}), d) \leq 3 + \frac{6n \sum_{Y \in \mathcal{A}} p(Y) d(\widehat{X}, Y)}{\sum_{Y \in \mathcal{A}} \widehat{w}^+(Y) d(\widehat{X}, Y)},$$

where $p(Y)$ is the probability that the Truncated Harmonic rule returns alternative Y . Note that since we choose each agent i uniformly at random, this probability is $p(Y) = \frac{1}{n} \sum_i p(i, Y)$, so

$$\text{dist}^M(f_\varepsilon^{\text{TH}}(\vec{\sigma}), d) \leq 3 + \frac{6n \sum_{Y \in \mathcal{A}} \frac{d(\widehat{X}, Y) \sum_{i \in N} p(i, Y)}{n}}{\sum_{Y \in \mathcal{A}} \widehat{w}^+(Y) d(\widehat{X}, Y)}$$

$$\begin{aligned}
&= 3 + \frac{6 \sum_{Y \neq \widehat{X}} d(\widehat{X}, Y) \sum_{i \in \mathcal{N}} p(i, Y)}{\sum_{Y \in \mathcal{A}} \widehat{w}^+(Y) d(\widehat{X}, Y)} & (d(X, X) = 0) \\
&= 3 + \frac{6 \sum_{Y \neq \widehat{X}} \frac{\varepsilon}{6} \widehat{w}^+(Y) d(\widehat{X}, Y)}{\sum_{Y \in \mathcal{A}} \widehat{w}^+(Y) d(\widehat{X}, Y)} & (p(i, Y) = \frac{\varepsilon}{6} \widehat{w}(i, Y)) \\
&= 3 + \varepsilon. & \square
\end{aligned}$$

4.2 Utilitarian Distortion Guarantees

Theorem 4. The utilitarian distortion of the Truncated Harmonic voting rule is $O(\sqrt{m}H_m)$.

PROOF. Consider any preference profile $\vec{\sigma}$ and any utility profile $\vec{u} \triangleright_{\cup} \vec{\sigma}$. Let X^* be the optimal alternative according to $\vec{u}$. We partition the agents into two sets $\mathcal{N}_1 = \{i \in \mathcal{N} : \widehat{X} \succ_i X^*\}$, and $\mathcal{N}_2 = \mathcal{N} \setminus \mathcal{N}_1$. For each $k \in \{0, 1\}$, we use $\text{SW}_k(Y, \vec{u}) = \sum_{i \in \mathcal{N}_k} u_i(Y)$ to denote the total value of alternative Y for the agents in $\mathcal{N}_k$, and let $S = \sum_{i \in \mathcal{N}} \sum_{Y \succ_i \widehat{X}} u_i(Y)$ be the sum of the utilities of the agents for the alternatives they prefer to $\widehat{X}$. In addition for alternative Y let $\text{SW}_s(Y, \vec{u})$ be the utility that she gains when she appears above $\widehat{X}$, i.e., $\text{SW}_s(Y, \vec{u}) = \sum_{Y \succ_i \widehat{X}} u_i(Y)$. Note that

$p(Y) \geq \varepsilon \text{SW}_s(Y, \vec{u}) / (6nH_m)$, and the probability of $\widehat{X}$ being selected by this rule is at least $1 - \frac{\varepsilon}{6}$. We have:

$$\begin{aligned}
\mathbb{E}_{Y \sim f_{\varepsilon}^{\text{TH}}(\vec{\sigma})} [\text{SW}(Y, \vec{u})] &\geq \left(1 - \frac{\varepsilon}{6}\right) \text{SW}(\widehat{X}, \vec{u}) + \sum_{Y \neq \widehat{X}} p(Y) \text{SW}(Y, \vec{u}) \\
&\geq \left(1 - \frac{\varepsilon}{6}\right) \text{SW}(\widehat{X}, \vec{u}) + \sum_{Y \neq \widehat{X}} \frac{\varepsilon \text{SW}_s(Y, \vec{u})}{6nH_m} \text{SW}(Y, \vec{u}) \\
&\geq \left(1 - \frac{\varepsilon}{6}\right) \text{SW}(\widehat{X}, \vec{u}) + \sum_{Y \neq \widehat{X}} \frac{\varepsilon \text{SW}_s(Y, \vec{u})^2}{6nH_m} \\
&\geq \left(1 - \frac{\varepsilon}{6}\right) \text{SW}(\widehat{X}, \vec{u}) + \frac{\varepsilon}{6nmH_m} \left(\sum_{Y \neq \widehat{X}} \text{SW}_s(Y, \vec{u}) \right)^2 \\
&= \left(1 - \frac{\varepsilon}{6}\right) \frac{n - S}{m} + \frac{\varepsilon S^2}{6nmH_m} & (\because \text{Decreasing in terms of } S) \\
&\geq \frac{\varepsilon n}{6mH_m}.
\end{aligned}$$

That means

$$\text{dist}^U(f_{\varepsilon}^{\text{TH}}(\vec{\sigma}), \vec{u}) \leq \frac{\text{SW}(X^*, \vec{u})}{\varepsilon n / (6mH_m)} = \frac{6mH_m \text{SW}(X^*, \vec{u})}{\varepsilon n}. \quad (1)$$

On the other hand, we have

$$\begin{aligned}
\text{dist}^U(f_{\varepsilon}^{\text{TH}}(\vec{\sigma}), \vec{u}) &= \frac{\text{SW}(X^*, \vec{u})}{\mathbb{E}_{Y \sim f_{\varepsilon}^{\text{TH}}(\vec{\sigma})} [\text{SW}(Y, \vec{u})]} \\
&= \frac{\text{SW}_1(X^*, \vec{u})}{\mathbb{E}_{Y \sim f_{\varepsilon}^{\text{TH}}(\vec{\sigma})} [\text{SW}(Y, \vec{u})]} + \frac{\text{SW}_2(X^*, \vec{u})}{\mathbb{E}_{Y \sim f_{\varepsilon}^{\text{TH}}(\vec{\sigma})} [\text{SW}(Y, \vec{u})]} \\
&\leq \frac{\text{SW}_1(X, \vec{u})}{(1 - \frac{\varepsilon}{6}) \text{SW}(X, \vec{u})} + \frac{\text{SW}_2(X^*, \vec{u})}{\mathbb{E}_{Y \sim f_{\varepsilon}^{\text{TH}}(\vec{\sigma})} [\text{SW}(Y, \vec{u})]}
\end{aligned}$$

$$\begin{aligned}
&\leq \frac{6}{6-\varepsilon} + \frac{SW_2(X^*, \vec{u})}{\frac{\varepsilon SW_s(Y, \vec{u})}{6nH_m} SW(X^*, \vec{u})} \\
&= \frac{6}{6-\varepsilon} + \frac{6nH_m}{\varepsilon SW(X^*, \vec{u})} \quad (\because SW_2(X^*, \vec{u}) = SW_s(X^*, \vec{u})) \\
&\leq \frac{12nH_m}{\varepsilon SW(X^*, \vec{u})}.
\end{aligned}$$

Finally, if we put it together with Equation (1), and use the fact that $\min(a, b) \leq \sqrt{a \cdot b}$, we have

$$\text{dist}^U(f_\varepsilon^{\text{TH}}(\vec{\sigma}), \vec{u}) \leq \sqrt{\frac{12nH_m}{\varepsilon SW(X^*, \vec{u})} \times \frac{6mH_m SW(X^*, \vec{u})}{\varepsilon n}},$$

and since this holds for any $\vec{\sigma}$ and $\vec{u} \triangleright_U \vec{\sigma}$, and ε is a positive constant, we can conclude that $\text{dist}^U(f_\varepsilon^{\text{TH}}) \in O(\sqrt{mH_m})$. $\square$

5 PARTIAL ORDINAL PREFERENCES (TOP-t)

In this section, we consider the problem of designing a voting rule with a good distortion in both worlds while we only have each agent's ranking for their top- t preferences, where $t < m$.

Top- t distortion. Similar to the full ranking case, we can define the distortion of a voting rule f on top- t preference profile $\vec{\sigma}_t$ as the worst-case distortion of $f(\vec{\sigma}_t)$ on any metric (utility) profile that is consistent with $\vec{\sigma}_t$. Note that since $\vec{\sigma}_t$ contains less information than a full-ranking preference profile, a wider class of metrics (utility profiles) are considered in the definition of distortion.

5.1 Deterministic Rules with Top- t Preferences

When it comes to deterministic voting rules with top- t preferences, we can achieve $\Theta(m^2)$ utilitarian distortion by choosing the plurality winner, which is optimal. In the metric framework, the distortion of any deterministic voting rule is at least $2m/t$. [Kempe, 2020] and Anagnostides et al. [2022] propose a deterministic voting rule with $6m/t + 1$ metric distortion. We first show that achieving the best of both worlds in this case is infeasible.

Theorem 5. Any deterministic voting rule on top- t preferences with bounded utilitarian distortion has metric distortion of $\Omega(m - t + 1)$.

PROOF. Consider the preference profile $\vec{\sigma}$ where $m - t + 1$ of the alternatives each appear as the top choice of $n/(m-t+1)$ agents, and the remaining $t - 1$ alternatives appear with the same order in the second to the t^{th} position of all the agents. Let X be the selected candidates. If X does not appear as the top choice of any agent, then in the utility profile where each agent has utility 1 for his top choice and zero for the rest we have unbounded distortion. Now let us consider the case where X has plurality score greater than zero. In that case consider the 1 dimensional metric space d where X is located at point 1, all the agents who have her as their top choice are located at point 0.5, and all remaining agents and alternatives are located at point zero. You can see that X has distance of at least 0.5 to each agent and hence $SC(X, d) > n/2$ and for any $Y \neq X$, $SC(Y, d) = n/2(m-t+1)$ which means X has $\Omega(m - t + 1)$ metric distortion. $\square$

We remark the contrast between this metric distortion lower bound of $\Omega(m - t + 1)$ when bounded utilitarian distortion is imposed, and the optimal metric distortion of $\Theta(m/t)$ when we ignore utilitarian distortion [Kempe, 2020]. This shows that requiring utilitarian distortion to be bounded makes metric distortion strictly worse whenever t is neither $\Theta(1)$ nor $m - \Theta(1)$.

Next, we show that this lower bound is tight by proving a matching upper bound.

Theorem 6. A deterministic voting rule exists with metric distortion $O(m - t + 1)$ and utilitarian distortion $O(m^2)$.

PROOF. Consider any preference profile $\vec{\sigma}$, utility profile $\vec{u} \triangleright_{\mathcal{U}} \vec{\sigma}$, and metric $d \triangleright_{\mathcal{M}} \vec{\sigma}$. If $t \leq m/2$, we can use the Plurality rule which has $O(m^2)$ utilitarian and $O(m) = O(m - t + 1)$ metric distortion, giving us the desired bounds.

Suppose that $t > m/2$, and let $\mathcal{A}^+$ be the set of alternatives that appear as the top choice of at least $n/2m$ agents. If $|\mathcal{A}^+| < 2(m - t + 1)$, there exists an alternative X that appears as the top choice of at least $n/4(m - t + 1)$ agents. This alternative at least gets a $1/m$ utility for each time she appears as the top choice of an agent, so her social welfare would at least be $n/(4m(m - t + 1))$, and since the maximum social welfare of any alternative is n this gives us $O(m(m - t + 1))$ utilitarian distortion. In addition, we know that an alternative that appears as the top choice of n' agents at most has a metric distortion of $\frac{n - n'}{n}$ that means X at most has a metric distortion of $\frac{n - n/(4m(m - t + 1))}{n/(4m(m - t + 1))} \in O(m - t + 1)$.

In the other case, where $|\mathcal{A}^+| \geq 2(m - t + 1)$, $\mathcal{A}^+$ includes at least $|\mathcal{A}^+| - m + t$ of the top- t alternatives of each agent's list. That means we can run a top- $|\mathcal{A}^+| - m + t$ deterministic voting rule to select alternative X with $O(\frac{|\mathcal{A}^+|}{|\mathcal{A}^+| - m + t})$ metric distortion (by Theorem 4.5 of [Anagnostides et al., 2022]). Note that this distortion is only compared to the best alternative in $\mathcal{A}^+$ and not compared to the best alternative overall. Since each member of $\mathcal{A}^+$ has a plurality score of at least $n/2m$, we can apply Lemma 2 with $\tau = 1/2m$ and show that the $\text{dist}^{\mathcal{M}}(X, d) \in O(\frac{|\mathcal{A}^+|}{|\mathcal{A}^+| - m + t}) = O(1)$. Finally, since $X \in \mathcal{A}^+$, she has at least $n/2m$ top votes and at least gets $1/m$ utility for each of them. That means her social welfare is $O(n/m^2)$, and since the maximum possible social welfare is n , this yields $O(m^2)$ utilitarian distortion. $\square$

Remark 1. For each instance, if the mechanism described in the proof of Theorem 6 outputs an alternative with $d^{\mathcal{M}}$ metric and $d^{\mathcal{U}}$ utilitarian distortion, we have $d^{\mathcal{M}} \times d^{\mathcal{U}} \in O(\max(m(m - t + 1)^2, m^2))$.

5.2 Randomized Rules with Top- t Preferences

In the randomized case, we know that we can achieve $O(\max(\sqrt{m}, m/t))$ utilitarian and $3 - 2/n$ metric distortion. The following theorem shows that for $t < \sqrt{m}$, we cannot simultaneously achieve the (asymptotically) best possible distortion in both worlds.

Theorem 7. Any (randomized) voting rule with access only to the top- t prefix of each agent's ranking that achieves a metric distortion of $d^{\mathcal{M}}$ will have a utilitarian distortion of $\Omega(\frac{m\sqrt{m}}{d^{\mathcal{M}}t\sqrt{t}})$.

PROOF. We will construct a top- t preference profile $\vec{\sigma}$ with n agents and m alternatives. Assume that n is divisible by $\frac{d^{\mathcal{M}}m\sqrt{m}}{t\sqrt{t}}$ and m is divisible by $3t$. We partition the alternatives into $\frac{m}{3t} + 2$ sets: let $\mathcal{A}^+$ and $\mathcal{A}^-$ each include $m/3$ alternatives and, for $k \in [\frac{m}{3t}]$, let $\mathcal{A}_k$ include t alternatives. Choose any $X \in \mathcal{A}^-$.

In $\vec{\sigma}$, for $k \in [\frac{m}{3t}]$, let $\mathcal{N}_k$ be a set of $\frac{nt\sqrt{t}}{d^{\mathcal{M}}m\sqrt{m}}$ agents who rank alternatives in $\mathcal{A}_k$ as their top- t choices (in any order). The remaining agents rank alternatives in $\mathcal{A}^+$ as their top- t choices in such a way that each of $m/3$ alternatives in $\mathcal{A}^+$ appears in the top- t choices of a $t/(m/3) = 3t/m$ fraction of these agents. Consider any voting rule f with a metric distortion at most $d^{\mathcal{M}}$ on this preference profile, where $d^{\mathcal{M}}$ is finite. We want to prove that its utilitarian distortion must be $\Omega(\frac{m\sqrt{m}}{d^{\mathcal{M}}t\sqrt{t}})$.

First, it is easy to check that $f(\vec{\sigma})$ cannot assign any positive probability to any alternative in $\mathcal{A}^-$. This is because it is possible that all agents rank the alternatives in $\mathcal{A}^-$ at the bottom of their preference rankings, so $\vec{\sigma}$ is consistent with distance metrics in which every alternative in $\mathcal{A}^-$ has an arbitrarily large social cost. Thus, assigning any positive probability to any alternative in $\mathcal{A}^-$ would result in unbounded metric distortion.

Next, define p_k to be the sum of the probabilities assigned to the alternatives in $\mathcal{A}_k$ under $f(\vec{\sigma})$, and let $k^* \in \arg \max_{k \in [\frac{m}{3t}]} p_k$. Now, consider the simple one-dimensional metric d over $\mathbb{R}$ in which alternatives in $\mathcal{A}_{k^*}$ are at 0, agents in $\mathcal{N}_{k^*}$ are at 1, and all other agents and alternatives are at 2. Hence, the distances are as follows for all agents i and alternatives Y :

$$d(i, Y) = \begin{cases} 1 & i \in \mathcal{N}_{k^*} \\ 2 & i \notin \mathcal{N}_{k^*} \text{ and } Y \in \mathcal{A}_{k^*} \\ 0 & i \notin \mathcal{N}_{k^*} \text{ and } Y \notin \mathcal{A}_{k^*}. \end{cases}$$

For the metric distortion of f , we have:

$$\begin{aligned} \text{dist}^M(f(\vec{\sigma}), d) &= \frac{\mathbb{E}_{Y \sim f(\vec{\sigma})} [\text{SC}(Y, d)]}{\min_{Y \in \mathcal{A}} \text{SC}(Y, d)} \\ &\geq \frac{p_{k^*} (2n - \frac{nt\sqrt{t}}{d^M m \sqrt{m}}) + (1 - p_{k^*}) \min_{Y \in \mathcal{A}} \text{SC}(Y, d)}{\min_{Y \in \mathcal{A}} \text{SC}(Y, d)} \\ &= 1 - p_{k^*} + \frac{p_{k^*} (2n - \frac{nt\sqrt{t}}{d^M m \sqrt{m}})}{\frac{nt\sqrt{t}}{d^M m \sqrt{m}}} = 1 - p_{k^*} + p_{k^*} \cdot \left(\frac{2d^M m \sqrt{m}}{t\sqrt{t}} - 1 \right) \\ &= 1 + 2p_{k^*} \cdot \left(\frac{d^M m \sqrt{m}}{t\sqrt{t}} - 1 \right) \geq 1 + 2p_{k^*} \cdot \frac{(d^M - 1)m\sqrt{m}}{t\sqrt{t}}. \end{aligned}$$

Since f has metric distortion at most d^M , we know that $\text{dist}^M(f(\vec{\sigma}), d) \leq d^M$, which means

$$2p_{k^*} \cdot \frac{(d^M - 1)m\sqrt{m}}{t\sqrt{t}} \leq d^M - 1 \implies p_{k^*} \leq \frac{t\sqrt{t}}{2m\sqrt{m}}. \quad (2)$$

We use this to derive a lower bound on the utilitarian distortion of f . Suppose that in the underlying full preference profile, every agent has alternatives in $\mathcal{A}^-$ at ranks $t+1$ through $t+m/3$, with X appearing at rank $t+1$. Consider the consistent utility profile $\vec{u}$ in which agents who have an alternative in $\mathcal{A}^+$ as their top choice have a utility of $\frac{1}{t+m/3}$ for their top $t+m/3$ alternatives and the remaining agents have a utility of $\frac{1}{t+1}$ for their top $t+1$ alternatives.

Under this utility profile, we have the following:

$$\begin{aligned} \text{SW}(X, \vec{u}) &\geq \frac{1}{t+1} \cdot \frac{n\sqrt{t}}{3d^M \sqrt{m}} \geq \frac{n}{6d^M \sqrt{mt}}, \\ \text{SW}(Y, \vec{u}) &\leq \frac{n \cdot 3t/m}{t+m/3} \leq \frac{9nt}{m^2}, \quad \forall Y \in \mathcal{A}^+, \\ \text{SW}(Y, \vec{u}) &= \frac{1}{t+1} \cdot \frac{nt\sqrt{t}}{d^M m \sqrt{m}} \leq \frac{n\sqrt{t}}{d^M m \sqrt{m}}, \quad \forall Y \in \mathcal{A}_k, k \in [m/(3t)]. \end{aligned}$$

Additionally, since $p_k \leq p_{k^*}$ for $k \in [m/(3t)]$, we have

$$\text{dist}^M(f(\vec{\sigma}), \vec{u}) \geq \frac{\text{SW}(X, \vec{u})}{\mathbb{E}_{Y \sim f(\vec{\sigma})} [\text{SW}(Y, \vec{u})]} \geq \frac{\frac{n}{6d^M \sqrt{mt}}}{p_{k^*} \cdot \frac{m}{3t} \cdot \frac{n\sqrt{t}}{d^M m \sqrt{m}} + \frac{9nt}{m^2}} \geq \frac{\frac{n}{6d^M \sqrt{mt}}}{\frac{nt}{6d^M m^2} + \frac{9nt}{m^2}} \in \Omega \left(\frac{m\sqrt{m}}{d^M t\sqrt{t}} \right),$$

where the last inequality is implied by Equation (2). $\square$

Recall that Random Dictatorship (which only requires access to plurality votes) achieves metric distortion less than 3 [Anshelevich and Postl, 2017], so the optimal metric distortion of randomized rules given top- t preferences is $O(1)$. The following corollary shows that any rule achieving this

optimal metric distortion bound must have utilitarian distortion $\Omega(\max((m/t)^{1.5}, \sqrt{m}))$, which is strictly worse than the optimal bound of $O(\max(m/t, \sqrt{m}))$ [Borodin et al., 2022] when $t = o(m^{2/3})$.

Corollary 2. Any voting rule with constant metric distortion on top- t preference profiles has utilitarian distortion of $\Omega\left(\max\left(\frac{m\sqrt{m}}{t\sqrt{t}}, \sqrt{m}\right)\right)$.

Definition 4 (Top- t Truncated Harmonic voting rule). The Top- t Truncated Harmonic voting rule, f^{TTH} uses the following steps:

- Let $\widehat{X}$ be the alternative that would be chosen by the mechanism proposed by Anagnostides et al. [2022], which has $6m/t + 1$ metric distortion.
- Select an agent $i \in \mathcal{N}$ uniformly at random.
- Return an alternative $Y \in \mathcal{A}$ with probability:

$$p(i, Y) = \begin{cases} \frac{1}{2H_i r_i(Y)} & \text{if } Y \succ_i \widehat{X}, \\ 1 - \sum_{Y \succ_i \widehat{X}} p(i, Y) & \text{if } Y = \widehat{X}, \\ 0 & \text{otherwise.} \end{cases}$$

Theorem 8. The Top- t Truncated Harmonic rule has metric distortion of $O\left(\frac{m}{t}\right)$.

PROOF. First, let us define $\widehat{X}$ -truncated weight function $\widehat{w}$ as $\widehat{w}(i, Y) = p(i, Y)$.

Note that for *any* metric distance d the alternative $\widehat{X}$ used by the Top- t Truncated Harmonic voting rule to define its truncated weight function has metric distortion at most $7m/t$. Therefore, Lemma 4 implies that the distortion of the alternative chosen by this voting rule for any distance d , inducing preference profile $\vec{\sigma}$, and using $\widehat{w}$ as the truncated weight function, is:

$$\text{dist}^M(f_\varepsilon^{\text{TTH}}(\vec{\sigma}), d) \leq \frac{7m}{t} + \frac{14mn \sum_{Y \in \mathcal{A}} p(Y) d(\widehat{X}, Y)}{t \sum_{Y \in \mathcal{A}} \widehat{w}^+(Y) d(\widehat{X}, Y)},$$

where $p(Y)$ is the probability that the Top- t Truncated Harmonic rule returns alternative Y . Note that since we choose each agent i uniformly at random, this probability is $p(Y) = \frac{1}{n} \sum_i p(i, Y) = \frac{1}{n} \sum_i \widehat{w}(i, Y)$, which yields:

$$\text{dist}^M(f_\varepsilon^{\text{TTH}}(\vec{\sigma}), d) \leq \frac{7m}{t} + \frac{14mn \sum_{Y \in \mathcal{A}} \frac{\widehat{w}^+(Y)}{n} d(\widehat{X}, Y)}{t \sum_{Y \in \mathcal{A}} \widehat{w}^+(Y) d(\widehat{X}, Y)} \leq \frac{21m}{t} \in O\left(\frac{m}{t}\right). \quad \square$$

Theorem 9. The Top- t Truncated Harmonic rule has utilitarian distortion of $O\left(\frac{m\sqrt{m}H_t}{t}\right)$.

PROOF. Consider any preference profile $\vec{\sigma}$ and utility profile $\vec{u} \triangleright_{\mathcal{U}} \vec{\sigma}$. Let p be the distribution returned by the Top- t Truncated Harmonic rule on $\vec{\sigma}$. Let $X^* \in \max_{X \in \mathcal{A}} \text{SW}(X, \vec{u})$ be an optimal alternative under $\vec{u}$. Define T_i to be the set of top- t alternatives of agent i . Let $S_1 = \sum_{i \in \mathcal{N}} \sum_{Y \in T_i: Y \succ_i \widehat{X}} u_i(Y)$ be the sum of the utilities of the agents for the alternatives in their top- t choices that they prefer to $\widehat{X}$, and S_2 be the sum of the utilities of the agents for all of their top- t choices.

In addition, for alternative Y , let $\text{SW}_s(Y, \vec{u})$ be her total utility from the agents who rank her among their top- t choices and above $\widehat{X}$, i.e., $\text{SW}_s(Y, \vec{u}) = \sum_{i \in \mathcal{N}: Y \succ_i \widehat{X}} u_i(Y)$.

Note that $p(Y) \geq \frac{\text{SW}_s(Y, \vec{u})}{2nH_t}$ for each $Y \in \mathcal{A} \setminus \{\widehat{X}\}$ and $p(\widehat{X}) \geq 1/2$, so

$$\mathbb{E}_{Y \sim f^{\text{TTH}}(\vec{\sigma})} [\text{SW}(Y, \vec{u})] = \frac{1}{2} \text{SW}(\widehat{X}, \vec{u}) + \sum_{Y \neq \widehat{X}} p(Y) \text{SW}(Y, \vec{u})$$

$$\begin{aligned}
&\geq \frac{1}{2} \text{SW}(\widehat{X}, \vec{u}) + \sum_{Y \neq \widehat{X}} \frac{\text{SW}_s(Y, \vec{u})}{2nH_t} \cdot \text{SW}_s(Y, \vec{u}) \\
&\geq \frac{1}{2} \text{SW}(\widehat{X}, \vec{u}) + \frac{\left(\sum_{Y \neq \widehat{X}} \text{SW}_s(Y, \vec{u}) \right)^2}{2nmH_t} \quad (\because \text{AM-QM inequality}) \\
&\geq \frac{S_2 - S_1}{2m} + \frac{S_1^2}{2nmH_t} \geq \frac{S_2^2}{4nmH_t} \geq \frac{nt^2}{4m^3H_t}, \quad (\because S_2 \geq nt/m)
\end{aligned}$$

where the fifth transition holds because $\frac{S_2 - S_1}{2m} + \frac{S_1^2}{2nmH_t} \geq \frac{S_2 - S_1}{2m} + \frac{S_1^2}{4nmH_t} \geq \frac{S_2^2}{4nmH_t}$ (the last transition is equivalent to $S_1 + S_2 \leq 2nH_t$, which is true because $S_1 + S_2 \leq 2n$ trivially). Next, let $\text{plu}(Y)$ be the number of agents whose top choice is Y . We have

$$\begin{aligned}
\mathbb{E}_{Y \sim f^{\text{TTH}}(\vec{\sigma})} [\text{SW}(Y, \vec{u})] &\geq \sum_{Y \in \mathcal{A}} \frac{\text{plu}(Y)}{2nH_t} \cdot \text{SW}(Y, \vec{u}) \geq \frac{1}{2nH_t} \sum_{Y \in \mathcal{A}} \frac{\text{plu}(Y)^2}{m} \\
&\geq \frac{1}{2nH_t} \left(\sum_{Y \in \mathcal{A}} \frac{\text{plu}(Y)}{m} \right)^2 \geq \frac{n}{2m^2H_t}.
\end{aligned}$$

That means

$$\text{dist}^U(f^{\text{TTH}}(\vec{\sigma}), \vec{u}) \leq \frac{\text{SW}(X^*, \vec{u})}{\max\left(\frac{nt^2}{4m^3H_t}, \frac{n}{2m^2H_t}\right)} = \min\left(\frac{4m^3H_t \text{SW}(X^*, \vec{u})}{nt^2}, \frac{2m^2H_t \text{SW}(X^*, \vec{u})}{n}\right). \quad (3)$$

Now, let $\mathcal{N}_1$ be the set of agents that have X^* in their top- t choices and prefer X^* to $\widehat{X}$, $\mathcal{N}_2$ be the set of agents that have $\widehat{X}$ in their top- t choices and prefer $\widehat{X}$ to X^* , and $\mathcal{N}_3$ be the remaining agents. Also, for $k \in \{1, 2, 3\}$, let P_k be the sum of utilities of agents in $\mathcal{N}_k$ for X^* . Note that $\text{SW}(X^*, \vec{u}) = P_1 + P_2 + P_3$.

If $P_3 \geq P_1 + P_2$ we have $\text{SW}(X^*, \vec{u}) \leq 2P_3$. On the other hand,

$$\begin{aligned}
\mathbb{E}_{Y \sim f^{\text{TTH}}(\vec{\sigma})} [\text{SW}(Y, \vec{u})] &\geq \sum_{Y \in \mathcal{A}} p(Y) \cdot \sum_{i \in \mathcal{N}_3: Y \in T_i} u_i(Y) \geq \sum_{Y \in \mathcal{A}} \frac{(\sum_{i \in \mathcal{N}_3, Y \in T_i} u_i(Y))^2}{2nH_t} \\
&\geq \frac{(\sum_{Y \in \mathcal{A}} \sum_{i \in \mathcal{N}_3, Y \in T_i} u_i(Y))^2}{2nmH_t} \geq \frac{t^2 P_3^2}{2nmH_t} \geq \frac{t^2 \text{SW}(X^*, \vec{u})^2}{8nmH_t},
\end{aligned}$$

where the third transition is due to the AM-QM inequality. Thus, we get that $\text{dist}^U(f^{\text{TTH}}(\vec{\sigma}), \vec{u}) \leq \frac{8nmH_t}{t^2 \text{SW}(X^*, \vec{u})}$. On the other hand, if $P_3 \leq P_1 + P_2$, we have:

$$\begin{aligned}
\text{dist}^U(f^{\text{TTH}}(\vec{\sigma}), \vec{u}) &= \frac{\text{SW}(X^*, \vec{u})}{\mathbb{E}_{Y \sim f^{\text{TTH}}(\vec{\sigma})} [\text{SW}(Y, \vec{u})]} \leq \frac{2P_2}{\frac{1}{2} \text{SW}(\widehat{X}, \vec{u})} + \frac{2P_1}{p(X^*) \text{SW}(X^*, \vec{u})} \\
&\leq \frac{4P_2}{P_2} + \frac{2P_1}{\frac{P_1}{2nH_t} \text{SW}(X^*, \vec{u})} \leq 4 + \frac{4nH_t}{\text{SW}(X^*, \vec{u})} \leq \frac{8nH_t}{\text{SW}(X^*, \vec{u})}.
\end{aligned}$$

That means

$$\text{dist}^U(f^{\text{TTH}}(\vec{\sigma}), \vec{u}) \leq \max\left(\frac{8nmH_t}{t^2 \text{SW}(X^*, \vec{u})}, \frac{8nH_t}{\text{SW}(X^*, \vec{u})}\right).$$

Finally, combined with Equation (3), and using the fact that $\min(a, b) \leq \sqrt{a \cdot b}$, for $t \leq \sqrt{m}$ we have

$$\text{dist}^U(f^{\text{TTH}}(\vec{\sigma}), \vec{u}) \leq \sqrt{\frac{4nmH_t}{t^2 \text{SW}(X^*, \vec{u})} \times \frac{2m^2H_t \text{SW}(X^*, \vec{u})}{n}} \in O\left(\frac{m\sqrt{m}H_t}{t}\right),$$

and for $t \geq \sqrt{m}$ we have

$$\text{dist}^U(f^{\text{TTH}}(\vec{\sigma}), \vec{u}) \leq \sqrt{\frac{8nH_t}{\text{SW}(X^*, \vec{u})} \times \frac{4m^3H_t\text{SW}(X^*, \vec{u})}{nt^2}} \in O\left(\frac{m\sqrt{m}H_t}{t}\right). \quad \square$$

The Top- t Truncated Harmonic rule provides one possible tradeoff by achieving $O(m/t)$ metric distortion and $O(m\sqrt{m}H_t/t)$ utilitarian distortion. We already know that Random Dictatorship, which can be executed given only top-1 preferences (and thus, also given top- t preferences for any $t \geq 1$) provides a different tradeoff by achieving $O(1)$ metric distortion [Anshelevich and Postl, 2017] and $O(m^{1.5})$ utilitarian distortion (Proposition 4). Using the next result, one can combine these rules (or any two rules with different tradeoffs) to find a spectrum of possible tradeoffs, which may help in eventually charting out the Pareto frontier.

Theorem 10. Fix any $\beta \in [0, 1]$. If we have voting rule f^1 with metric distortion d_1^M and utilitarian distortion d_1^U , and voting rule f^2 with metric distortion d_2^M and utilitarian distortion d_2^U , we can combine these two rules to achieve $\beta d_1^M + (1 - \beta)d_2^M$ metric distortion and $\frac{d_1^U d_2^U}{\beta d_2^U + (1 - \beta)d_1^U} \leq \max\{d_1^U/\beta, d_2^U/(1 - \beta)\}$ utilitarian distortion.

6 DISCUSSION

In this work, we investigate the feasibility of achieving asymptotically (near-)optimal distortion under both metric and utilitarian worlds simultaneously. We prove that while this is possible given full rankings, this is not the case given partial rankings in the form of top- t preferences.

Our work leaves open a number of exciting open questions. For example, our deterministic rule, Pruned Plurality Veto, for achieving the best of both worlds given full rankings achieves a metric distortion of $9 + \epsilon$. Even though this is a constant, it is undesirably high, and improving it is an important direction. We prove that it cannot be improved all the way to 3, but our lower bound is only slightly greater than 3. Achieving tight distortion bounds for randomized rules with top- t preferences is another obvious open direction. For example, does there exist a randomized voting rule with top- t preferences achieving constant metric distortion and $O(\frac{m\sqrt{m}}{t\sqrt{m}})$ utilitarian distortion?

More broadly, it would be interesting to derive “best of both worlds” style results for other settings studied in the distortion literature, such as the utilitarian distortion with unit-range instead of unit-sum assumption (where the minimum and maximum utility of each agent is set to 0 and 1, respectively, instead of normalizing the sum to 1), alternative ballot designs such as approval ballots, and returning a committee or a ranking of alternatives instead of a single alternative.

Another compelling direction concerns the communication complexity of these rules. Our results indicate that it is not feasible to simultaneously achieve asymptotically optimal distortion in both worlds when given top- t preferences, for some values of t . However, in the utilitarian setting, it is known that one can achieve lower distortion while eliciting fewer bits of information than under top- t preferences by using more efficient elicitation methods [Mandal et al., 2019, 2020]. This line of work studies the communication-distortion tradeoff, i.e., the optimal distortion achievable while eliciting a given number of bits of information regarding agent preferences. Combining it with our best of both worlds guarantee adds another axis of metric versus utilitarian to that tradeoff.

ACKNOWLEDGMENTS

Gkatzelis was partially supported by NSF CAREER award CCF 2047907. Latifian and Shah were partially supported by an NSERC Discovery Grant.

REFERENCES

- Georgios Amanatidis, Georgios Birmpas, Aris Filos-Ratsikas, and Alexandros A Voudouris. 2021. Peeking behind the ordinal curtain: Improving distortion via cardinal queries. *Artificial Intelligence* 296 (2021), 103488.
- Georgios Amanatidis, Georgios Birmpas, Aris Filos-Ratsikas, and Alexandros A Voudouris. 2022. A few queries go a long way: Information-distortion tradeoffs in matching. *Journal of Artificial Intelligence Research* 74 (2022), 227–261.
- Ioannis Anagnostides, Dimitris Fotakis, and Panagiotis Patsilinas. 2022. Metric-distortion bounds under limited information. *Journal of Artificial Intelligence Research* 74 (2022), 1449–1483.
- Elliot Anshelevich, Onkar Bhardwaj, Edith Elkind, John Postl, and Piotr Skowron. 2018. Approximating optimal social choice under metric preferences. *Artificial Intelligence* 264 (2018), 27–51.
- Elliot Anshelevich, Aris Filos-Ratsikas, Nisarg Shah, and Alexandros A. Voudouris. 2021. Distortion in Social Choice Problems: The First 15 Years and Beyond. In *Proceedings of the 30th International Joint Conference on Artificial Intelligence (IJCAI)*. 4294–4301. Survey Track.
- Elliot Anshelevich and John Postl. 2017. Randomized social choice functions under metric preferences. *Journal of Artificial Intelligence Research* 58 (2017), 797–827.
- Elliot Anshelevich and Shreyas Sekar. 2016. Blind, greedy, and random: algorithms for matching and clustering using only ordinal information. In *Proceedings of the 30th AAAI Conference on Artificial Intelligence (AAAI)*. 383–389.
- Elliot Anshelevich and Wennan Zhu. 2021. Ordinal approximation for social choice, matching, and facility location problems given candidate positions. *ACM Transactions on Economics and Computation (TEAC)* 9, 2 (2021), 1–24.
- Haris Aziz. 2020. Justifications of welfare guarantees under normalized utilities. *ACM SIGecom Exchanges* 17, 2 (2020), 71–75.
- Gerdus Benade, Swaprava Nath, Ariel D Procaccia, and Nisarg Shah. 2021. Preference elicitation for participatory budgeting. *Management Science* 67, 5 (2021), 2813–2827.
- Allan Borodin, Daniel Halpern, Mohamad Latifian, and Nisarg Shah. 2022. Distortion in voting with top-t preferences. In *Proceedings of the 31st International Joint Conference on Artificial Intelligence (IJCAI)*. 116–122.
- Allan Borodin, Omer Lev, Nisarg Shah, and Tyrone Strangway. 2019. Primarily about primaries. In *Proceedings of the 33rd AAAI Conference on Artificial Intelligence (AAAI)*. 1804–1811.
- Craig Boutilier, Ioannis Caragiannis, Simi Haber, Tyler Lu, Ariel D Procaccia, and Or Sheffet. 2015. Optimal social choice functions: A utilitarian view. *Artificial Intelligence* 227 (2015), 190–213.
- Ioannis Caragiannis, Swaprava Nath, Ariel D Procaccia, and Nisarg Shah. 2017. Subset selection via implicit utilitarian voting. *Journal of Artificial Intelligence Research* 58 (2017), 123–152.
- Ioannis Caragiannis and Ariel D Procaccia. 2011. Voting almost maximizes social welfare despite limited communication. *Artificial Intelligence* 175, 9-10 (2011), 1655–1671.
- Ioannis Caragiannis, Nisarg Shah, and Alexandros A Voudouris. 2022. The metric distortion of multiwinner voting. *Artificial Intelligence* 313 (2022), 103802.
- Moses Charikar and Prasanna Ramakrishnan. 2022. Metric distortion bounds for randomized social choice. In *Proceedings of the 2022 Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*. SIAM, 2986–3004.
- Soroush Ebadian, Aris Filos-Ratsikas, Mohamad Latifian, and Nisarg Shah. 2023a. Explainable and Efficient Randomized Voting Rules. <https://www.cs.toronto.edu/~nisarg/papers/explainable-distortion.pdf>.
- Soroush Ebadian, Rupert Freeman, and Nisarg Shah. 2022a. Efficient Resource Allocation with Secretive Agents. In *Proceedings of the 31st International Joint Conference on Artificial Intelligence (IJCAI)*. 272–278.
- Soroush Ebadian, Anson Kahng, Dominik Peters, and Nisarg Shah. 2022b. Optimized Distortion and Proportional Fairness in Voting. In *Proceedings of the 23rd ACM Conference on Economics and Computation (EC)*. 563–600.
- Soroush Ebadian, Mohamad Latifian, and Nisarg Shah. 2023b. The Distortion of Approval Voting with Runoff. In *Proceedings of the 22nd International Conference on Autonomous Agents and Multi-Agent Systems (AAMAS)*. Forthcoming.
- Aris Filos-Ratsikas and Alexandros A Voudouris. 2021. Approximate mechanism design for distributed facility location. In *Proceedings of the 14th International Symposium on Algorithmic Game Theory (SAGT)*. 49–63.
- Vasilis Gkatzelis, Daniel Halpern, and Nisarg Shah. 2020. Resolving the optimal metric distortion conjecture. In *Proceedings of the 61st Symposium on Foundations of Computer Science (FOCS)*. 1427–1438.
- Daniel Halpern and Nisarg Shah. 2021. Fair and Efficient Resource Allocation with Partial Information. In *Proceedings of the 30th International Joint Conference on Artificial Intelligence (IJCAI)*. 224–230.
- David Kempe. 2020. Communication, distortion, and randomness in metric voting. In *Proceedings of the 34th AAAI Conference on Artificial Intelligence (AAAI)*. 2087–2094.
- Fatih Erdem Kizilkaya and David Kempe. 2022. Plurality Veto: A Simple Voting Rule Achieving Optimal Metric Distortion. In *Proceedings of the 31st International Joint Conference on Artificial Intelligence (IJCAI)*. 349–355.
- Fatih Erdem Kizilkaya and David Kempe. 2023. Generalized Veto Core and a Practical Voting Rule with Optimal Metric Distortion. In *Proceedings of the 24th ACM Conference on Economics and Computation (EC)*. Forthcoming.

- Debmalya Mandal, Ariel D. Procaccia, Nisarg Shah, and David P. Woodruff. 2019. Efficient and thrifty voting by any means necessary. In *Proceedings of the 32nd Annual Conference on Neural Information Processing Systems (NeurIPS)*. 7178–7189.
- Debmalya Mandal, Nisarg Shah, and David P. Woodruff. 2020. Optimal communication-distortion tradeoff in voting. In *Proceedings of the 21st ACM Conference on Economics and Computation (EC)*. 795–813.
- Kamesh Munagala and Kangning Wang. 2019. Improved metric distortion for deterministic social choice rules. In *Proceedings of the 20th ACM Conference on Economics and Computation (EC)*. 245–262.
- Jannik Peters. 2023. A Note on Rules Achieving Optimal Metric Distortion. arXiv:2305.08667.
- Ariel D Procaccia and Jeffrey S Rosenschein. 2006. The distortion of cardinal preferences in voting. In *Proceedings of the 10th International Workshop on Cooperative Information Agents (CIA)*. 317–331.

APPENDIX

A MISSING PROOF

PROOF OF THEOREM 10. Consider preference profile $\vec{\sigma}$, utility profile $\vec{u} \triangleright_U \vec{\sigma}$, and metric $d \triangleright_M \vec{\sigma}$. Define rule f to run rule f^1 with probability β and rule f^2 with probability $1 - \beta$. For the metric distortion we have:

$$\begin{aligned}
 \text{dist}^M(f(\vec{\sigma}), d) &= \frac{\mathbb{E}_{X \sim f(\vec{\sigma})} [\text{SC}(X, d)]}{\min_{X \in \mathcal{A}} \text{SC}(X, d)} \\
 &= \frac{\beta \mathbb{E}_{X \sim f^1(\vec{\sigma})} [\text{SC}(X, d)] + (1 - \beta) \mathbb{E}_{X \sim f^2(\vec{\sigma})} [\text{SC}(X, d)]}{\min_{X \in \mathcal{A}} \text{SC}(X, d)} \\
 &= \frac{\beta \mathbb{E}_{X \sim f^1(\vec{\sigma})} [\text{SC}(X, d)]}{\min_{X \in \mathcal{A}} \text{SC}(X, d)} + \frac{(1 - \beta) \mathbb{E}_{X \sim f^2(\vec{\sigma})} [\text{SC}(X, d)]}{\min_{X \in \mathcal{A}} \text{SC}(X, d)} \\
 &= \beta \text{dist}^M(f^1(\vec{\sigma}), d) + (1 - \beta) \text{dist}^M(f^2(\vec{\sigma}), d) \leq \beta d_1^M + (1 - \beta) d_2^M.
 \end{aligned}$$

On the other hand, for the utilitarian case we have:

$$\begin{aligned}
 \text{dist}^U(f(\vec{\sigma}), \vec{u}) &= \frac{\max_{X \in \mathcal{A}} \text{SW}(X, \vec{u})}{\mathbb{E}_{X \sim f(\vec{\sigma})} [\text{SW}(X, d)]} \\
 &= \frac{\max_{X \in \mathcal{A}} \text{SW}(X, \vec{u})}{\beta \mathbb{E}_{X \sim f^1(\vec{\sigma})} [\text{SW}(X, d)] + (1 - \beta) \mathbb{E}_{X \sim f^2(\vec{\sigma})} [\text{SW}(X, d)]} \\
 &= \frac{1}{\beta \frac{\mathbb{E}_{X \sim f^1(\vec{\sigma})} [\text{SW}(X, d)]}{\max_{X \in \mathcal{A}} \text{SW}(X, \vec{u})} + (1 - \beta) \frac{\mathbb{E}_{X \sim f^2(\vec{\sigma})} [\text{SW}(X, d)]}{\max_{X \in \mathcal{A}} \text{SW}(X, \vec{u})}} \\
 &\leq \frac{1}{\frac{\beta}{d_1^U} + \frac{1 - \beta}{d_2^U}} = \frac{d_1^U d_2^U}{\beta d_2^U + (1 - \beta) d_1^U}.
 \end{aligned}$$

□