

Approximately Strategy-Proof Voting

Eleanor Birrell*
Cornell University
eleanor@cs.cornell.edu

Rafael Pass†
Cornell University
rafael@cs.cornell.edu

April 20, 2011

Abstract

The classic Gibbard-Satterthwaite Theorem establishes that only dictatorial voting rules are strategy-proof; under any other voting rule, players have an incentive to lie about their true preferences. We consider a new approach for circumventing this result: we consider randomized voting rules that only *approximate* a deterministic voting rule and only are *approximately* strategy-proof. We show that *any* deterministic voting rule can be approximated by an approximately strategy-proof randomized voting rule, and we provide asymptotically tight lower bounds on the parameters required by such voting rules.

1 Introduction

The classic Gibbard-Satterthwaite Theorem [7, 15] considers the question of when voters will honestly report their preferences. It shows that if the voting rule has at least three outcomes, then only *dictatorial* functions (i.e., one player determines the output) can be honestly computed by rational agents.

The earliest approach to circumventing this limitation, first suggested by Gibbard [8] and later advocated by Conitzer and Sandholm [3], consists of using randomized approximations as a means to bypass the limitations of deterministic rules. Unfortunately, the potential of this approach is limited by two negative results. Gibbard showed that the only strategy-proof randomized voting rules are *trivial* in that they consist of simple probability distributions over rules that depend on only one voter (*unilateral rules*) and rules that have at most two possible outputs (*dupe rules*). In recent work, Procaccia [14] quantified the quality of approximation that can be achieved by such trivial functions (for certain types of voting rules): in particular, he constructed a simple approximation of PLURALITY (the voting rule that returns the outcome that receives the most first-choice votes). However, his

*This material is based upon work supported under a National Science Foundation Graduate Research Fellowship

†Supported in part by a Alfred P. Sloan Fellowship, Microsoft New Faculty Fellowship, NSF CAREER Award CCF-0746990, AFOSR Award FA9550-08-1-0197, BSF Grant 2006317.

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE 20 APR 2011		2. REPORT TYPE		3. DATES COVERED 00-00-2011 to 00-00-2011	
4. TITLE AND SUBTITLE Approximately Strategy-Proof Voting				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Cornell University, Department of Computer Science, Ithaca, NY, 14853				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT The classic Gibbard-Satterthwaite Theorem establishes that only dictatorial voting rules are strategy-proof; under any other voting rule, players have an incentive to lie about their true preferences. We consider a new approach for circumventing this result we consider randomized voting rules that only approximate a deterministic voting rule and only are approximately strategy-proof. We show that any deterministic voting rule can be approximated by an approximately strategy-proof randomized voting rule and we provide asymptotically tight lower bounds on the parameters required by such voting rules.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 20	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

approximation only guarantees that the expected number of votes received by the returned output is $2/3$ the number received by the true winner, and he proves that his mechanism is asymptotically optimal.

An alternative approach that consists of restricting the class of preference functions. Notably, Moulin [13] considers voting rules defined over an output set with a natural total ordering with single-peaked preferences, that is utilities that have an optimal outcome (the *peak*) and decrease monotonically with distance from the peak. He constructs rules that are strategy-proof with respect to this (very restricted) class of utility functions.

A final intriguing approach to circumventing this limitation, first suggested by Bartholdi et al. [2], is to construct voting rules that are computationally difficult to manipulate. However, there are strong impossibility results that limit the potential of this approach. In their original work, Bartholdi et al. showed that many standard voting rules can be efficiently manipulated. Moreover, recent work demonstrates that for any voting rule, “bad inputs” (preference profiles that admit a successful non-honest strategy) are not rare [3, 5, 10]; these papers develop lower bounds for the number of preference profiles which admit some kind of manipulation and show that when the number of outputs is small it is easy to find successful manipulations.

In this work, we consider a new approach to circumventing these previous negative results: *approximately* strategy-proof voting rules. This approach is motivated by the observation that previous work assumes that people will deviate from the honest strategy even if the improvement in their utility is extremely small. In practice, people often don’t deviate if the gain is very small; a small gain in utility may be offset by a psychological cost associated with lying or by the computational cost of computing an effective deviation [9]. This phenomenon is compounded by the fact that when risk-averse individual voters are uncertain about the preferences of the other voters, they may not pursue a deviation that is only expected to yield a small benefit. In other contexts, these observations have led to the introduction of relaxed solution concepts (e.g., ε -Nash equilibria and ε -dominant strategies). In the context of voting, we thus consider ε -strategy-proof voting rules; that is voting rules for which no deviating strategy can improve a player’s expected utility by more than ε .

Unfortunately, a corollary of the Gibbard-Satterthwaite theorem shows that the only deterministic voting rules that are ε -strategy-proof are dictatorial rules. We thus consider approximate strategy-proofness in the context of randomized voting rules. As it turns out, the parameter ε quite sharply determines whether or not there exist non-trivial ε -strategy-proof voting rules.

For the remainder of this section, let n be the number of voters and let the number of outcomes be a fixed constant k .

$\varepsilon = \omega(1/n)$: In this regime, we show that there exist natural, non-trivial, ε -strategy-proof voting rules. Moreover, we show that *every* deterministic voting rule f can be approximated by a non-trivial, ε -strategy-proof voting rule g . Towards formalizing this, we require a notion of what it means for a randomized mechanism to approximate a deterministic voting rule. Intuitively, we want to capture the idea that with high probability, the output of g is “close” to that of f . To define closeness, we consider a specialized distance metric inspired by the idea of vote corruption, that is the observation that in real-life elections, votes get miscounted, recounted, lost, etc. We say that the output $y \in g(\vec{x})$ is δ -close to the right answer if there

exists some vector $\vec{x}'$ that differs from $\vec{x}$ in only δ positions, such that $f(\vec{x}') = y$; in other words, when the output is δ -correct, it means that we could have reached this output by flipping only δ votes. Our main theorem can now be informally stated as follows:

Theorem 1.1 (Upper Bound – Informal Statement). *Let $\varepsilon = \omega(1/n)$, $\beta > 0$, and let f be a deterministic voting rule over n players. For sufficiently large n , there exists an ε -strategy-proof randomized voting rule g that is a βn -approximation of f .*

Intuitively speaking, our mechanism guarantees that the outcome returned by g will be the correct one modulo a change in a few votes. When dealing with large election (e.g., national elections) this seems like a reasonable guarantee.

One may ask whether an alternative notion of approximation can be achieved: For instance, could we hope to get a mechanism that yields the *correct* output with high probability? In Section 6 we consider this alternative definition, and we show that this is impossible unless the underlying voting rule f is dictatorial. We believe this negative result highlights why our definition of approximation is both reasonable and minimal for circumventing the negative results of Gibbard and Satterthwaite.

$\varepsilon = o(1/n)$: In this setting, we show that the Gibbard’s result characterizing strategy-proof randomized voting rules [8] still applies:

Theorem 1.2 (Lower Bound – Informal Statement). *If g is a $o(1/n)$ -strategy-proof randomized voting rule, then g is trivial (i.e., a distribution over unilateral and dupe rules).*

We additionally show that natural voting rules (e.g., PLURALITY) cannot be well approximated by such mechanisms. This result can be informally stated as follows:

Theorem 1.3 (PLURALITY – Informal Statement). *Let g be a trivial voting rule over n players. There exists β such that for sufficiently large n , g cannot be a βn -approximation of PLURALITY.*

In fact, in Section 6 we extend this result to show that there is no trivial voting rule that approximates PLURALITY even with high probability.

The two theorems combined bound the approximation parameters that can be achieved for natural voting rules like PLURALITY. Note that Theorem 1.3 also implies that the approximations constructed in Theorem 1.1 are non-trivial.

Additional Properties Finally, we observe that there are many properties (other than strategy-proofness) that are desirable in a voting rule. We show that in addition to being ε -strategy-proof, the mechanisms we construct are *collusion-resistant*—a group of t players cannot increase their collective utility by more than a small amount. Moreover, when we consider a *neutral* voting rule—one whose outcome is independent of voter identities—and a constant number of outcomes, our mechanism is computationally efficient.

2 Definitions and Preliminaries

A voting rule f is a mapping from player “votes” to an outcome in the set $[k] = \{1, \dots, k\}$. Player votes are total preference orderings over the set of outcomes $[k]$; these preference orderings are represented as a permutation $\sigma_i \in \Sigma_k$ (here Σ_k denotes the set of permutations over $[k]$). We use the notation $\sigma_i(j) > \sigma_i(j')$ when the preference type $\sigma_i \in \Sigma_k$ ranks outcome j higher than outcome j' . For convenience, a preference profile $\vec{\sigma} = (\sigma_1, \dots, \sigma_n)$ may also be denoted by $(\sigma_i, \vec{\sigma}_{-i})$ for any $i \in [n]$.

Players are motivated by determinism—that is, they prefer that outcomes ranked higher in their preference ordering σ_i are chosen. More formally, each player i has a preference ordering σ_i and a utility function $u_i \in \mathcal{U}_{\sigma_i}$, the class of utility functions u_i that satisfy the following two properties: First, for all $\vec{\sigma}_{-i}$, $u_i(\vec{\sigma}, j) \geq u_i(\vec{\sigma}, j')$ if $\sigma_i(j) > \sigma_i(j')$. Second, for all input-output pairs $(\vec{\sigma}, j)$, $u_i(\vec{\sigma}, j) \in [0, 1]$.

Observe that although it is traditional to talk exclusively about the preference relations σ_i , the underlying utility function is necessary in order to quantify *approximate* strategy-proofness. The restriction to utilities in the range $[0, 1]$ is not strictly necessary (our techniques only rely on the fact that the utilities are bounded), however we believe that utilities in this range lend themselves to a more intuitive interpretation.

2.1 Approximately Strategy-proof Voting

We say that a voting rule is (approximately) strategy-proof if for all players, honestly reporting their preferences (approximately) dominates any other reporting strategy.

Definition 2.1 (ε -strategy-proof). A voting rule g is ε -strategy-proof if for all players i , all preference profiles $\vec{\sigma}$, all alternative preferences σ'_i , and all utility functions $u_i \in \mathcal{U}_{\sigma_i}$,

$$\mathbb{E}[u_i(\vec{\sigma}, g(\vec{\sigma}))] + \varepsilon \geq \mathbb{E}[u_i(\vec{\sigma}, g(\sigma'_i, \vec{\sigma}_{-i}))]$$

The notion of strategy-proof voting rules can also be extended to handle collusions between groups of t players.

Definition 2.2. A voting rule $g : (\Sigma_k)^n \rightarrow [k]$ is (t, ε) -strategy-proof if for all subsets $S \subseteq [n]$ such that $|S| \leq t$, all preference profiles $\vec{\sigma}$ and $\vec{\sigma}'$ such that $\sigma_i = \sigma'_i$ for all $i \notin S$, and all utility profiles $\vec{u} \in \mathcal{U}_{\vec{\sigma}}$,

$$\sum_{i \in S} \mathbb{E}[u_i(\vec{\sigma}, g(\vec{\sigma}))] + \varepsilon \geq \sum_{i \in S} \mathbb{E}[u_i(\vec{\sigma}, g(\vec{\sigma}'))].$$

Intuitively, we interpret ε -strategy-proofness to mean that if there is a small cost associated with deviating (e.g., a psychological cost to lying or a computational cost to finding a successful deviation) then the voters will honestly report their preferences.

2.2 Voter Max-Influence

We introduce a new concept inspired by the notion of variable influence [11] which we call *max-influence*; max-influence describes the maximum amount by which a single player can impact the probability that a particular output is returned by a voting rule.

Definition 2.3 (ξ -max-influence). A player $i \in [n]$ has ξ -max-influence over a voting rule $g : (\Sigma_k)^n \rightarrow [k]$ if there exists a preference profile $\vec{\sigma}$ and an alternative preference σ'_i such that the statistical distance between $g(\vec{\sigma})$ and $g(\sigma'_i, \vec{\sigma}_{-i})$ is at least ξ , that is if

$$\frac{1}{2} \sum_{j \in [k]} |\Pr[g(\vec{\sigma}) = j] - \Pr[g(\sigma'_i, \vec{\sigma}_{-i}) = j]| > \xi$$

We observe that if no player i has ξ -max-influence over a voting rule g then the rule is approximately strategy-proof.

Lemma 2.4. *Let $g : (\Sigma_k)^n \rightarrow [k]$ be a voting rule such that no player $i \in [n]$ has ξ -max-influence over g . Then g is 2ξ -strategy-proof.*

Proof. Let $\vec{\sigma}, \vec{\sigma}'$ be preference profiles that differ only on input i (that is $\sigma_j = \sigma'_j$ for all $j \neq i$), and let u_i be voter i 's utility function. Let $S_{\vec{\sigma}} = \{j : \Pr[g(\vec{\sigma}) = j] \geq \Pr[g(\vec{\sigma}') = j]\}$ and let $S_{\vec{\sigma}'}$ be its complement (the set of outputs that are more likely to be generated by $\vec{\sigma}'$).

$$\begin{aligned} E[u_i(\vec{\sigma}, g(\vec{\sigma}'))] - E[u_i(\vec{\sigma}, g(\vec{\sigma}))] &= \sum_j \Pr[g(\vec{\sigma}') = j] u_i(\vec{\sigma}, j) - \sum_j \Pr[g(\vec{\sigma}) = j] u_i(\vec{\sigma}, j) \\ &= \sum_{j \in S_{\vec{\sigma}'}} (\Pr[g(\vec{\sigma}') = j] - \Pr[g(\vec{\sigma}) = j]) u_i(\vec{\sigma}, j) \\ &\quad + \sum_{j \in S_{\vec{\sigma}}} (\Pr[g(\vec{\sigma}') = j] - \Pr[g(\vec{\sigma}) = j]) u_i(\vec{\sigma}, j) \\ &\leq \sum_{j \in S_{\vec{\sigma}'}} (\Pr[g(\vec{\sigma}') = j] - \Pr[g(\vec{\sigma}) = j]) \cdot 1 + 0 \\ &\leq 2\xi \end{aligned}$$

Therefore the constructed approximation g is 2ξ -strategy-proof. $\square$

We also present an analogous result for coalitions.

Lemma 2.5. *Let $g : (\Sigma_k)^n \rightarrow [k]$ be a voting rule such that no player $i \in [n]$ has ξ -max-influence over g . Then g is $(t, 2t^2\xi)$ -strategy-proof.*

Proof. For any subset of t players $S = \{s(1), \dots, s(t)\} \subseteq [n]$, let $\vec{\sigma}, \vec{\sigma}'$ be two preference profiles that differ only on components $i \in S$, that is for all $i \notin S$, $\sigma_i = \sigma'_i$. Define a sequence of hybrid profiles $\vec{\sigma}^0, \dots, \vec{\sigma}^t$ such that in profile $\vec{\sigma}^\ell$, players $j \in \{s(1), \dots, s(\ell)\}$ have preference $\sigma_j^\ell = \sigma'_j$, players $j \in \{s(\ell+1), \dots, s(t)\}$ have preference $\sigma_j^\ell = \sigma_j$ and all other players $j \notin S$ have preference $\sigma_j^\ell = \sigma_j = \sigma'_j$. Observe that $\vec{\sigma} = \vec{\sigma}^0$ and $\vec{\sigma}' = \vec{\sigma}^t$.

Observe that for all $i \in S$ and all $\ell \in [t]$, $\mathbb{E}[u_i(\vec{\sigma}, g(\vec{\sigma}^\ell))] - \mathbb{E}[u_i(\vec{\sigma}, g(\vec{\sigma}^{\ell-1}))] \leq 2\xi$ (the proof is identical to the proof of 2ξ -strategy-proofness in Lemma 2.4). It follows the difference in the group's collective utility between any two consecutive hybrid profiles is bounded by $\sum_{i \in S} \mathbb{E}[u_i(\vec{\sigma}, g(\vec{\sigma}^\ell))] - \sum_{i \in S} \mathbb{E}[u_i(\vec{\sigma}, g(\vec{\sigma}^{\ell-1}))] \leq 2t\xi$ and therefore that the total difference in collective utility between the honest strategy and the collective deviating strategy is at most

$$\sum_{i \in S} \mathbb{E}[u_i(\vec{\sigma}, g(\vec{\sigma}))] - \sum_{i \in S} \mathbb{E}[u_i(\vec{\sigma}, g(\vec{\sigma}'))] \leq 2t^2\xi \quad \square$$

Observe that Gibbard and Satterthwaite’s classic result shows that every non-dictatorial voting rule with at least three outputs has a player with 1-max-influence. By contrast, we will show that every voting rule can be approximated by a randomized voting rule over which no player has much max-influence.

2.3 Approximations

In order to formally define the notion of a randomized approximation, it is necessary to define a closeness metric d for voting rules. The appropriate choice of metric in such a context is not immediately clear. Since outcomes are assigned an arbitrary number $j \in [k]$ and votes are permutations σ_i over the outcomes (interpreted as a total preference ordering), standard notions of “closeness” are meaningless. While some particular voting rules have natural quality scores that can be used to define an approximation [14], such techniques are not fully generalizable.

We instead introduce a pseudometric d_v inspired by vote corruption, that is the observation that in real-life elections, votes get miscounted, recounted, lost, etc. This pseudometric d_v associates a number ℓ with each pair $(\vec{\sigma}, j) \in (\Sigma_k)^n \times [k]$ defined by

$$\ell(\vec{\sigma}, j) = \min_{\vec{\sigma}' \text{ s.t. } f(\vec{\sigma}') = j} \Delta(\vec{\sigma}, \vec{\sigma}')$$

where $\Delta(\vec{\sigma}, \vec{\sigma}')$ is the number of components which differ between $\vec{\sigma}$ and $\vec{\sigma}'$. We define an induced pseudometric $d_v((\vec{\sigma}, j), (\vec{\sigma}', j')) = |\ell(\vec{\sigma}, j) - \ell(\vec{\sigma}', j')|$. This says that an output is close to the correct answer if there exists an input close to the true input which generates that output.

Using this metric, we consider an approximation g to be a function whose value is *close* to that of f .

Definition 2.6 (δ -approximation). A (randomized) voting rule g is a δ -approximation of a voting rule f if for all inputs $\vec{\sigma}$ and all possible random coins,

$$d_v((\vec{\sigma}, g(\vec{\sigma})), (\vec{\sigma}, f(\vec{\sigma}))) \leq \delta.$$

Remark 2.7. In Section 6, we relax our definition of an approximation to consider functions that are close to the correct outcome with high probability and show that our lower bounds extend to the relaxed definition.

2.4 Trivial Voting Rules

There are two classes of simple voting rules, collectively referred to as trivial, that will be used to characterize the set of strategy-proof voting rules under certain conditions. The first is the class of rules that depend on only one player’s inputs, and the second is the class that returns at most two outputs.

Definition 2.8 (Gibbard [8]). A deterministic voting rule $f : (\Sigma_k)^n \rightarrow [k]$ is *unilateral* if there exists a player i such that for all preferences profiles $\vec{\sigma}, \vec{\sigma}' \in (\Sigma_k)^n$ satisfying $\sigma_i = \sigma'_i$, $f(\vec{\sigma}) = f(\vec{\sigma}')$. A voting rule that is unilateral and onto is also called *dictatorial*.

Definition 2.9 (Gibbard [8]). A deterministic voting rule $f : (\Sigma_k)^n \rightarrow [k]$ is a *duple* if the range is at most 2, that is if $|\{j \in [k] : \exists \vec{\sigma} \in (\Sigma_k)^n \text{ such that } f(\vec{\sigma}) = j\}| \leq 2$.

A voting rule is called *trivial* if it is a probability distribution over unilateral and duple rules, otherwise it is called *non-trivial*.

3 Previous Work

Randomized voting has been recently explored by Procaccia [14] who shows how to define 0-strategy-proof approximations of certain voting rules that are derived from natural quality scores. Our work differs from his approach both in our use of ε -strategy-proof rules and in the way we define and construct approximations. In particular, Procaccia’s work relies on a natural quality score to define an approximation, a technique that prevents discussion of certain types of voting rules including multi-layered rules like run-off elections or the electoral college-based system employed in U.S. presidential elections. Furthermore, as Procaccia demonstrates, the quality-score based approach is inherently limited in the quality of approximations that can be achieved; for example, it is impossible to construct a strategy-proof approximation of PLURALITY that will (in expectation) return an outcome with quality score greater $\Omega(1/\sqrt{k})$ times the optimal. Although he does construct an approximation for PLURALITY, if his approximation were employed during an election between four candidates, the expected number of votes received by the candidate it returns would be $2/3$ the number of votes received by the true winner.

There are two important negative results in the context of voting. The first, proven independently by Gibbard [7] and Satterthwaite [15], demonstrates that only a trivial collection of voting rules are strategy-proof. Restated with our definitions, they show that only *dictatorial* functions are 0-strategy-proof.

Theorem 3.1 (Gibbard-Satterthwaite). *Let $f : (\Sigma_k)^n \rightarrow [k]$ be a deterministic, onto voting rule with $k \geq 3$. Then f is 0-strategy-proof if and only if it is dictatorial.*

This result has been quantitatively extended to give lower bounds on the number of input profiles which admit manipulations [5, 10]. Their work shows that not only are manipulations relatively common, but they can also be found efficiently.

The Gibbard-Satterthwaite theorem has also been extended to characterize the class of strategy-proof randomized voting rules [8]. In terms of our definitions, this extension shows that only trivial voting rules are 0-strategy-proof.

Theorem 3.2 (Gibbard). *Let $g : (\Sigma_k)^n \rightarrow [k]$ be a randomized voting rule. Then g is 0-strategy-proof if and only if it is trivial.*

4 Approximate Voting

Leveraging our new notion of approximate strategy-proofness, we now show that *every* voting rule can be approximated by a randomized voting rule. Intuitively, we construct an approximation by adding noise to the original deterministic voting rule. The probability

that these approximations return a particular outcome decreases linearly with the distance between the outcome under consideration and the correct outcome; the slope of this linear function bounds the influence that a voter can have on the resulting approximation. If the slope is sufficiently steep, then we can guarantee that all “bad” outputs (ones that are δ -far from the correct output) are chosen with probability 0.

Our techniques can be seen as a linear analog of the exponential mechanism [12] that has been used to establish differential privacy [4].

Theorem 4.1. *For any deterministic voting rule $f : (\Sigma_k)^n \rightarrow [k]$, any $\varepsilon > 0$ and any $\delta \geq k(k+1+\varepsilon)/\varepsilon - 1$, f has a ε -strategy-proof δ -approximation g .*

Proof. We construct an approximation g of the voting rule f as follows: we assign each input-output pair $(\vec{\sigma}, j)$ a quality score $q(\vec{\sigma}, j) = -d_v((\vec{\sigma}, f(\vec{\sigma})), (\vec{\sigma}, j))$. Observe that $q(\vec{\sigma}, j)$ decreases linearly with the (minimal) number of votes that must be corrupted before f returns j instead of $f(\vec{\sigma})$. Let $\xi = \varepsilon/k(k+1+\varepsilon)$ —this value is chosen to guarantee ε -strategy-proofness. The mechanism g returns the value j with probability proportional to $\max\{1 + \xi q(\vec{\sigma}, j), 0\}$. Note that g never returns an outcome more than $1/\xi$ -far from $f(\vec{\sigma})$.

First, we bound the max-influence a voter can have over g . For any $\vec{\sigma}'$ that differs from $\vec{\sigma}$ in only one position and any outcome $j \in [k]$, the difference $|\Pr[g(\vec{\sigma}') = j] - \Pr[g(\vec{\sigma}) = j]|$ is equal to

$$\left| \frac{\max\{1 + \xi q(\vec{\sigma}', j), 0\}}{\sum_{\iota \in [k]} \max\{1 + \xi q(\vec{\sigma}', \iota), 0\}} - \frac{\max\{1 + \xi q(\vec{\sigma}, j), 0\}}{\sum_{\iota \in [k]} \max\{1 + \xi q(\vec{\sigma}, \iota), 0\}} \right|.$$

For the purpose of clarity, let $A = \max\{1 + \xi q(\vec{\sigma}, j), 0\}$ and let $B = \sum_{\iota \in [k]} \max\{1 + \xi q(\vec{\sigma}, \iota), 0\}$. Using this notation we observe that the difference can be bounded above by

$$\left| \frac{A + \xi}{B - k\xi} - \frac{A}{B} \right| = \left| \frac{AB + \xi B - AB + k\xi A}{B^2 - k\xi B} \right| = \left| \frac{\xi - k\xi A/B}{B - k\xi} \right|$$

Where the first expression follows from the definition of g , the second from cross multiplication, and the third from canceling terms.

Since $A = \max\{1 + \xi q(\vec{\sigma}, j), 0\} \leq 1$ (recall that quality scores are negative) and $B = \sum_{\iota \in [k]} \max\{1 + \xi q(\vec{\sigma}, \iota), 0\} \geq 1$ (since the correct output, which is included in the sum, contributes 1 and there are no negative components), we can maximize this bound by setting $A = B = 1$ giving

$$|\Pr[g(\vec{\sigma}') = j] - \Pr[g(\vec{\sigma}) = j]| \leq \frac{\xi + k\xi}{1 - k\xi} = \varepsilon/k$$

which implies that no player has max-influence greater than $\varepsilon/2$. Therefore by Lemma 2.4 the constructed approximation g is ε -strategy-proof.

Second, we claim that g is a good approximation for f according to the distance metric d_v . Observe that the approximation g returns a outcome with distance greater than $1/\xi$ from the correct answer with probability 0. Since we fixed $\delta \geq 1/\xi$, all “bad” outputs are sufficiently far from the correct answer that they are returned with probability zero, therefore g always returns an answer that is δ -close to the correct outcome. $\square$

Corollary 4.2. *Let $\varepsilon = \omega(1/n)$, $\beta > 0$, and let $f : (\Sigma_k)^n \rightarrow [k]$ be a deterministic voting rule. For sufficiently large n , there exists an ε -strategy-proof randomized voting rule g that is a βn -approximation of f .*

Proof. Fix $\beta > 0$ and let $\delta = \beta n$. Let g be defined as in the proof of Theorem 4.1 and recall that (as shown in the previous proof) g is an ε -strategy-proof βn -approximation of f if $\beta n \geq k(k+1+\varepsilon)/\varepsilon - 1$. Since $\varepsilon = \omega(1/n)$, the result follows immediately. $\square$

Example 4.3. For concreteness, consider an election with 100 million voters and three outputs (about the scale of a United States presidential election). Fixing $\varepsilon = .001$, we can construct an approximation g that is guaranteed to *always* return an answer within 12,500 votes of the correct answer, in practice, well within the vote corruption in such an election. Looked at in another way, this says that in any such election in which one outcome wins by at least 12,500 votes, this mechanism will always return the correct answer.

Observe that if $\varepsilon < 1/n$ then all outputs are chosen with positive probability. For such small values of ε , the approximation g is well-defined, but it is a trivial n -approximation of f . In Section 5 we show that this $\omega(1/n)$ restriction is an inherent bound on the achievable approximation parameters and not an artifact of the construction employed in Theorem 4.1. In contrast, when $\varepsilon = \omega(1/n)$, the approximations both offer good guarantees and are non-trivial.

Corollary 4.4. *Let $\varepsilon = \omega(1/n)$ and $k = O(1)$. For sufficiently large numbers of voters n , there exist non-trivial ε -strategy-proof randomized voting rules.*

Proof. Let f be PLURALITY, the voting rule that returns the outcome j that receives the most first-choice votes (using some deterministic tie-breaking rule). By Corollary 6.4 we know that there exist arbitrarily good approximations of f for $\varepsilon = \omega(1/n)$. By Theorem 5.5 this implies that for such ε , the mechanism from Theorem 4.1 is non-trivial. $\square$

In addition to being non-trivial and approximately strategy-proof, the randomized voting rule g constructed in the proof of Theorem 4.1 has several other nice properties. First, the voting rule g is collusion-resistant: its guarantees degrade gracefully if the mechanism is extended to protect against collusion by t players.

Corollary 4.5. *For any voting rule $f : (\Sigma_k)^n \rightarrow [k]$, any $t < n$, any $\varepsilon > 0$, and any $\delta \geq (k(k_1)t^2 + k\varepsilon)/\varepsilon - 1$, f has a (t, ε) -strategy-proof δ -approximation.*

Proof. The proof is equivalent to that of Theorem 4.1 except that we use a different parameter $\xi = \varepsilon/k(k+1)t^2 + k\varepsilon$. The proof that the randomized voting rule g —constructed as in the proof of Theorem 4.1 (except with the new value of ξ)—is (t, ε) -strategy-proof follows immediately from Lemma 2.5. The proof that g is a δ -approximation of f is identical to the proof in Theorem 4.1. $\square$

Second, we observe that for *neutral* voting rules—those whose outcome is independent of voter identities—our approximations are computationally efficient.

Corollary 4.6. *If $f : (\Sigma_k)^n \rightarrow [k]$ is a neutral, efficiently computable voting rule with a constant number of outputs k , then the approximation g defined in the proof of Theorem 4.1 can be computed in polynomial time.*

Proof. To compute the output of the approximation $g(\vec{\sigma})$, where g is defined as in the proof of Theorem 4.1, we begin by computing the quality score $q(\vec{\sigma}, j)$ for each possible outcome $j \in [k]$. Since f is neutral, the correct outcome $f(\vec{\sigma})$ depends only on the vote configuration—that is the unordered set of votes that were cast—and not on the full preference profile $\vec{\sigma}$ per se. We can therefore compute the quality scores of all outputs $j \in [k]$ simply by evaluating the original rule f on each of the possible vote configurations and setting the quality score of an output to be equal to the (negative) minimum distance between the configuration associated with $\vec{\sigma}$ and a configuration that returns the outcome j . There are $k!$ different ways a player could vote and each voting preference is cast by at most n voters, therefore there are $O(n^{k!})$ vote configurations that need to be considered.

Having computed all of the quality scores $q(\vec{\sigma}, j)$, we can define a distribution $D_{\vec{\sigma}}$ such that the probability that $D_{\vec{\sigma}}$ returns an outcome $j \in [k]$ is given by $D_{\vec{\sigma}}(j) = \max\{1 + \xi q(\vec{\sigma}, j), 0\} / \sum_{\iota \in [k]} \max\{1 + \xi q(\vec{\sigma}, \iota), 0\}$. Observe that this distribution is efficiently samplable, therefore the approximation g can be computed efficiently. $\square$

5 Optimality of our Approximations

In this section, we develop lower bounds on the approximation parameters that can be achieved for approximately strategy-proof voting rules. In particular, we demonstrate that only trivial voting rules can be ε -strategy-proof for small values of ε , and we show that such voting rules cannot provide good approximations of natural voting rules like PLURALITY.

We begin by showing that for small values of ε , only trivial voting rules are ε -strategy-proof. The general outline of the proof is as follows: we define a reduction between ε -strategy-proof voting rules and 0-strategy-proof voting rules by including punishments for players who misreport their preferences. Consider the following modified mechanism: with probability $1 - p$ we run the original mechanism and with probability p we choose a single player i and output his first choice (according to his *reported* preferences). This has the effect that for appropriate values of ε, p , no one wants to misreport their first choice. However, they can still successfully deviate by making more subtle changes further down their preference profile. We overcome this by modifying our behavior slightly: when we choose a player i , we return each of his choices with probability inversely proportional to their ranking in his reported preferences. For small values of ε , this style of punishment is sufficient to offset any benefits yielded by a deviating strategy under the original mechanism. The result then follows from Theorem 3.2. More formally:

Theorem 5.1. *Let $\varepsilon < 1/nk^3$. A randomized voting rule $g : (\Sigma_k)^n \rightarrow [k]$ is ε -strategy-proof if and only if it is trivial.*

Proof. Let u_i be uniformly distributed for all $i \in [n]$ — that is the utility assigned to the outcome ranked j th is $(k - j)/(k - 1)$. Define a new randomized voting rule g' that for each $i \in [n]$, returns i 's j th choice (according to i 's reported preference) with probability $k\varepsilon(k - j)$ and otherwise follows the original mechanism g . That is, we use one of the unilateral “punishment” mechanisms with probability $p = n \sum_{j \in [k]} k\varepsilon(k - j) < nk^3\varepsilon$ and we use the original mechanism with probability $1 - p$. Observe that both of these probabilities are in the range $[0, 1]$ because $\varepsilon < 1/nk^3$.

Let $\vec{\sigma}, \vec{\sigma}'$ be preference profiles that differ only in the i th component and let $\sigma_{i,j}$ denote the outcome ranked j th by σ_i . We begin by claiming that the expected loss of utility from deviating (due to the unilateral punishment mechanisms) is at least ε . This will offset any utility gained under the original mechanism g .

Claim 5.2. $\sum_j k\varepsilon(k-j)(u_i(\vec{\sigma}, \sigma_{i,j}) - u_i(\vec{\sigma}, \sigma'_{i,j})) \geq \varepsilon$.

We first show how to prove Theorem 5.1 assuming this claim, and then we prove the claim itself.

$$\begin{aligned}
E[u_i(\vec{\sigma}, g'(\vec{\sigma}))] - E[u_i(\vec{\sigma}, g'(\vec{\sigma}'))] &= \sum_j k\varepsilon(k-j)(u_i(\vec{\sigma}, \sigma_{i,j}) - u_i(\vec{\sigma}, \sigma'_{i,j})) \\
&\quad + (1-p)(E[u_i(\vec{\sigma}, g(\vec{\sigma}))] - E[u_i(\vec{\sigma}, g(\vec{\sigma}'))]) \quad (1) \\
&\geq \varepsilon + (1-p)(E[u_i(\vec{\sigma}, g(\vec{\sigma}))] - E[u_i(\vec{\sigma}, g(\vec{\sigma}'))]) \quad (2) \\
&\geq \varepsilon + (1-p)(-\varepsilon) \quad (3) \\
&\geq 0 \quad (4)
\end{aligned}$$

(1) follows by the definition of expectation, (2) follows immediately from Claim 5.2, (3) follows from the fact that the original mechanism g is ε -strategy proof, and (4) follows from the fact that $p \in [0, 1]$. Since g' is 0-strategy-proof, by Theorem 3.2 it must be trivial. Since g' follows mechanism g with probability $1-p > 0$, it follows that g is also trivial.

Proof of Claim 5.2: We begin by considering the simplified case where σ_i and σ'_i are identical except that the positions of two outcomes j_1, j_2 are swapped. Assume without loss of generality that $\sigma_i(j_1) > \sigma_i(j_2)$. In this case, most of the k unilateral punishment rules cancel (because the preferences only differ in two positions). Let $r_\sigma(j)$ indicate the rank of j according to preference σ . The expected difference in utility contributed by the punishment mechanism is

$$(k\varepsilon(k - r_{\sigma_i}(j_1))u_i(j_1) + k\varepsilon(k - r_{\sigma_i}(j_2))u_i(j_2)) - (k\varepsilon(k - r_{\sigma'_i}(j_1))u_i(j_1) + k\varepsilon(k - r_{\sigma'_i}(j_2))u_i(j_2))$$

Since the difference between σ_i, σ'_i is that the two ranks are switched— $r_{\sigma'_i}(j_1) = r_{\sigma_i}(j_2)$ and $r_{\sigma'_i}(j_2) = r_{\sigma_i}(j_1)$ —the expected difference in utility can be written as

$$k\varepsilon((k - r_{\sigma_i}(j_1)) - (k - r_{\sigma_i}(j_2)))(u_i(j_1) - u_i(j_2))$$

The difference in rank between j_1 and j_2 under preference σ_i is at least 1. Since u_i is uniformly distributed, the difference in utility u_i between j_1, j_2 is at least $1/k$, therefore $\sum_j k\varepsilon(k-j)(u_i(\vec{\sigma}, \sigma_{i,j}) - u_i(\vec{\sigma}, \sigma'_{i,j})) \geq k\varepsilon \cdot 1 \cdot (1/k) = \varepsilon$.

We now consider the general case where σ_i and σ'_i differ arbitrarily. We define a sequence $\vec{\sigma}' = \vec{\sigma}^{(1)}, \vec{\sigma}^{(2)}, \dots, \vec{\sigma}^{(k^2)} = \vec{\sigma}$ iteratively as follows: given $\vec{\sigma}^{(\iota)}$, let $\sigma_\ell^{(\iota+1)} = \sigma_\ell^{(\iota)}$ for all $\ell \neq i$. Define the preferences $\sigma_i^{(\iota)}$ iteratively using bubble sort [6]: more specifically, let $\tau(\iota) = \iota \bmod k$. If $\tau(\iota) \neq 1$ and $\sigma_i(\sigma_{i,\tau(\iota)-1}^{(\iota-1)} < \sigma_i(\sigma_{i,\tau(\iota)}^{(\iota-1)})$ then define $\sigma_i^{(\iota)}$ to be the same as $\sigma_i^{(\iota-1)}$ except with the outcomes ranked $\tau-1$ and τ flipped. Else let $\sigma_i^{(\iota)} = \sigma_i^{(\iota-1)}$. The fact that this sequence is well defined (and terminates with $\vec{\sigma}^{(m)} = \vec{\sigma}$) follows immediately from the correctness of bubble sort. If two adjacent hybrid preferences are identical, the voter's

expected utility is clearly identical. If two adjacent hybrid preferences are not identical, the difference in voter i 's expected utility between two adjacent hybrid preferences is at least ε : this follows by the above argument (for the simple case) since two outcomes are flipped if and only if they start in the wrong relative order. Claim 5.2 then follows by the triangle inequality. $\square$

In spite of this limitation, some trivial voting rules still have good approximation for small values of ε .

Example 5.3. Define a class of voting rules $\text{SUBSET-}\delta : (\Sigma_k)^n \rightarrow [k]$ by $\text{SUBSET-}\delta(\vec{\sigma}) = \text{PLURALITY}(\sigma_1, \dots, \sigma_\delta)$. The rule g that returns an output $j \in [k]$ uniformly at random is a 0-strategy-proof δ -approximation of $\text{SUBSET-}\delta$.

Example 5.4. Define a class of voting rules $\text{MODULO-}k : (\Sigma_k)^n \rightarrow [k]$ that returns the *number* of first-choice votes given to the winner (under the voting rule PLURALITY) modulo k . Again the rule g that returns an output k uniformly at random is a 0-strategy-proof $\lfloor k/2 \rfloor$ -approximation.

However, neither of these voting rules satisfy our intuition concerning what constitutes a “good” voting rule. $\text{SUBSET-}\delta$ is not *neutral*—that is the outcome depends on the order of the players—and $\text{MODULO-}k$ is not monotonic.

We focus instead on a common, natural voting rule: PLURALITY . We show that under our definition of an approximation, no trivial voting rule is a good approximation of PLURALITY . The proof closely follows that of Procaccia [14]. Observe that this result implies that for the voting rule PLURALITY , the approximation constructed in Theorem 4.1 is asymptotically optimal.

Theorem 5.5. *There exists $\beta > 0$ such that for all n , PLURALITY does not have a trivial βn -approximation.*

Proof. Fix a constant value $c > 0$ (e.g., $c = \sqrt{k}$). Recall that by definition, a trivial function g is a probability distribution over unilateral and duple deterministic voting rules.

First consider the case where g selects a duple. Define $q_j = \Pr[j \in \text{Range}(g) | g \text{ selects a duple}]$ and observe that $\sum_{j \in [k]} q_j = 2$ which implies that there exists a set $A' \subseteq [k]$ of k/c alternatives such that $\sum_{j \in A'} q_j \leq 2/c \leq 2/c$.

Now consider the case where g selects a unilateral rule. Observe that there exists a set $N' \subseteq [n]$ of n/c players such that a unilateral rule that considers only an agent $i \in N'$ is selected with probability at most $1/c$ (conditioned on g selecting a unilateral rule). Construct a partial preference profile defined for players $i \in [n] \setminus N'$ that satisfies the property that each outcome $j \in [k] \setminus A'$ is ranked first by equally many agents $i \in [n] \setminus N'$ (and no such agent ranks any other outcome first). Since $|A'| \geq k/c$, there exist an output $x^* \in A'$ such that for all preference profiles $\vec{\sigma}$ consistent with the defined preferences σ_i , the probability that $g(\vec{\sigma}) = x^*$ (conditioned on g selecting a unilateral rule that considers only an agent $i \in [n] \setminus N'$) is at most c/k . Complete the preference profile $\vec{\sigma}$ so that it satisfies the property

that for all $i \in N'$, x^* is ranked first. Observe that

$$\begin{aligned}
\Pr[g(\vec{\sigma}) = x^* | g \text{ selects unilateral rule}] &= \Pr[g(\vec{\sigma}) = x^* | g \text{ selects unilateral rule in } N'] \\
&\quad \cdot \Pr[g \text{ selects unilat. rule in } N' | g \text{ selects unilat. rule}] \\
&\quad + \Pr[g(\vec{\sigma}) = x^* | g \text{ selects unilateral rule in } [n] \setminus N'] \\
&\quad \cdot \Pr[g \text{ selects unilat. rule in } [n] \setminus N' | g \text{ selects unilat.}] \\
&\leq 1 \cdot 1/c + c/k \cdot 1 \\
&\leq 1/c + c/k
\end{aligned}$$

Since $x^* \in A'$, by combining the two cases, we get that $\Pr[g(\vec{\sigma}) = x^*] \leq \max\{2/c, 1/c + c/k\}$.

Finally, we observe that the number of first-choice votes given to x^* is at least n/c . and the number of first-choice votes given to any output $j \neq x^*$ is at most $(n - n/c)/(k - k/c) = n/k$.

We interpret this as saying that simultaneously, all outcomes $j \neq x^*$ are bad (at least $(n/c - n/k)/2$ corrupted votes are required to effect the output of f) and the probability that g returns the one good output x^* is low. This means that f does not have a trivial $(n/c - n/k)/2 - 1$ -approximation. $\square$

6 Relaxed Approximations

In this section, we consider a relaxed notion of approximation, called a (δ, μ) -approximation, under which we require only that the approximation g return an outcome that is close to the correct outcome *with high probability*. We explore how this relaxed definition can be leveraged to establish a trade-off between the proximity of “good answers” to the true outcome and the probability of returning a good answer. In particular, we demonstrate two distinct constructions of such relaxed approximations: the first construction extends the linear approach employed in Theorem 4.1, the second construction uses an exponential approach inspired by differential privacy and provides improved guarantees for certain values of δ . We also develop lower bounds that shown the asymptotic optimality of our constructions.

6.1 Defining (δ, μ) -Approximations

We begin by formalizing this weaker notion of approximation.

Definition 6.1 ((δ, μ) -approximation). A (randomized) voting rule g is a (δ, μ) -approximation of a voting rule f if for all inputs $\vec{\sigma}$,

$$\Pr[d_v((\vec{\sigma}, g(\vec{\sigma})), (\vec{\sigma}, f(\vec{\sigma}))) > \delta] \leq \mu.$$

Observe that a $(\delta, 0)$ -approximation is equivalent to a δ -approximation as defined in Section 2.

6.2 Constructing Relaxed Approximations

By relaxing our definition, we inherently facilitate the construction of randomized ε -strategy-proof approximations for any deterministic function. In this section, we show two such constructions. In Theorem 6.2 we extend the linear construction from Theorem 4.1 to quantify the μ -values received for smaller values of δ . In Theorem 6.3 we provide an alternative exponential construction that provides improved guarantees for small values of δ .

Theorem 6.2. *For any deterministic voting rule $f : (\Sigma_k)^n \rightarrow [k]$ and any $\varepsilon > 0, \delta \geq 0$, f has a ε -strategy-proof (δ, μ) -approximation g , where $\mu = 0$ if $\varepsilon(\delta + 1)/k(k + 1 + \varepsilon) \geq 1$ and $\mu = (k - 1)(1 - (\delta + 1)\varepsilon/k(k + 1 + \varepsilon))/(1 + (k - 1)(1 - (\delta + 1)\varepsilon/k(k + 1 + \varepsilon)))$ else.*

Proof. The construction is identical to that employed in the proof of Theorem 4.1: we assign each input-output pair $(\vec{\sigma}, j)$ a quality score $q(\vec{\sigma}, j) = -d_v((\vec{\sigma}, f(\vec{\sigma})), (\vec{\sigma}, j))$. As before, let $\xi = \varepsilon/k(k + 1 + \varepsilon)$ —this value is chosen to guarantee ε -strategy-proofness. The mechanism g returns the value j with probability proportional to $\max\{1 + \xi q(\vec{\sigma}, j), 0\}$.

The proof that the resulting mechanism is ε -strategy-proof is identical to that employed in the proof of Theorem 4.1. It is only necessary to demonstrate the quality of approximation achieved, using the relaxed notion of (δ, μ) -approximations.

As before, let $M(\vec{\sigma}, \iota) = |\{j \in [k] : q(\vec{\sigma}, j) = \iota\}|$, that is $M(\vec{\sigma}, \iota)$ is the number of outputs with quality score $q(\vec{\sigma}, j) = \iota$. We again observe that for all profiles $\vec{\sigma}$,

$$\Pr[d_v((\vec{\sigma}, f(\vec{\sigma})), (\vec{\sigma}, g(\vec{\sigma}))) > \delta] = \frac{\sum_{\iota=-\delta-1}^{-n} \max\{1 + \xi \iota, 0\} M(\vec{\sigma}, \iota)}{\sum_{\iota=0}^{-n} \max\{1 + \xi \iota, 0\} M(\vec{\sigma}, \iota)}$$

We now wish to bound this probability for smaller values of δ , so we express the denominator as the sum of two terms $\sum_{\iota=-\delta-1}^{-n} \max\{1 + \xi q(\vec{\sigma}, \iota), 0\} M(\vec{\sigma}, \iota) + \sum_{\iota=0}^{-\delta-1} \max\{1 + \xi q(\vec{\sigma}, \iota), 0\} M(\vec{\sigma}, \iota)$ which is minimized (thereby maximizing the total fraction) when the second term (the “good” outputs) consists only of the one correct answer.

$$\Pr[d_v((\vec{\sigma}, f(\vec{\sigma})), (\vec{\sigma}, g(\vec{\sigma}))) > \delta] \leq \frac{\sum_{\iota=-\delta-1}^{-n} \max\{1 + \xi \iota, 0\} M(\vec{\sigma}, \iota)}{\sum_{\iota=-\delta-1}^n \max\{1 + \xi \iota, 0\} M(\vec{\sigma}, \iota) + 1}$$

Finally, it is clear that the right hand side of the inequality is maximized when the numerator is maximized, and this occurs when there are $k - 1$ outputs with quality $q(\vec{\sigma}, j) = -\delta - 1$. Therefore

$$\forall \vec{\sigma}, \Pr[d_v((\vec{\sigma}, f(\vec{\sigma})), (\vec{\sigma}, g(\vec{\sigma}))) > \delta] \leq \frac{(k - 1) \max\{1 + \xi(-\delta - 1), 0\}}{1 + (k - 1) \max\{1 + \xi(-\delta - 1), 0\}}$$

The result follows. □

Observe that Theorem 6.2 is a simple extension of Theorem 4.1 that quantifies the guarantees achieved for smaller values of δ , that is for $\delta < k(k + 1 + \varepsilon)/\varepsilon - 1$. However, for such small values, this linear construction is not necessarily optimal. We instead present an alternative method for constructing (δ, μ) -approximations in which $\mu > 0$ but which offers improved guarantees for small values of δ .

Theorem 6.3. *For any deterministic voting rule $f : (\Sigma_k)^n \rightarrow [k]$ and any $\varepsilon > 0, \delta \geq 0$, f has a ε -strategy-proof (δ, μ) -approximation g , where $\mu = (k - 1)/((\varepsilon + 1)^{(\delta+1)/2} + k - 1)$.*

Proof. We employ the exponential mechanism of Talwar and McSherry [12] to construct an approximation g of the voting rule f as follows: we assign each input-output pair $(\vec{\sigma}, j)$ a quality score $q(\vec{\sigma}, j) = n - d_v((\vec{\sigma}, f(\vec{\sigma})), (\vec{\sigma}, j))$. Observe that $q(\vec{\sigma}, j)$ decreases linearly with the (minimal) number of votes that must be corrupted before f returns j instead of $f(\vec{\sigma})$. Let $\xi = (1/2) \ln(\varepsilon + 1)$ —this value is chosen to guarantee ε -strategy-proofness as discussed in Lemma ???. The mechanism g returns the value j with probability proportional to $\exp(\xi q(\vec{\sigma}, j))$.

First, we claim that the resulting mechanism is ε -strategy-proof. From [12] we know that this mechanism gives 2ξ -differential privacy. We reproduce the proof for the sake of completeness: The probability that an output j is chosen is given by $(\exp(\xi q(\vec{\sigma}, j)) / (\sum_{\iota \in [k]} \exp(\xi q(\vec{\sigma}, \iota))))$. A single change in the input profile $\vec{\sigma}$ can, by definition, change the quality score by at most 1, giving a factor of at most $\exp(\xi)$ in the numerator and at least $\exp(-\xi)$ in the denominator, yielding a total change in probability of at most $\exp(2\xi)$. It follows by Lemma ??? that the approximation g is ε -strategy-proof.

Second, we claim that g is a good approximation for f according to the distance metric d_v . We note that Talwar and McSherry present a general accuracy bound for the exponential mechanism [12], however, that bound is too loose for our purposes.¹ Instead, we take advantage of the particular distance pseudometric d_v to develop a tighter bound. Let $M(\vec{\sigma}, \iota) = |\{j \in [k] : q(\vec{\sigma}, j) = \iota\}|$, that is $M(\vec{\sigma}, \iota)$ is the number of outputs with quality score $q(\vec{\sigma}, j) = \iota$. We observe that for all profiles $\vec{\sigma}$,

$$\begin{aligned} \Pr[d_v((\vec{\sigma}, f(\vec{\sigma})), (\vec{\sigma}, g(\vec{\sigma}))) > \delta] &= \frac{\sum_{j: d_v((\vec{\sigma}, f(\vec{\sigma})), (\vec{\sigma}, g(\vec{\sigma}))) > \delta} \exp(\xi q(\vec{\sigma}, g(\vec{\sigma})))}{\sum_{j \in [k]} \exp(\xi q(\vec{\sigma}, g(\vec{\sigma})))} \\ &= \frac{\sum_{\iota=0}^{n-\delta-1} \exp(\xi \iota) M(\vec{\sigma}, \iota)}{\sum_{\iota=0}^n \exp(\xi \iota) M(\vec{\sigma}, \iota)} \end{aligned}$$

That is, the probability that g returns an output greater than δ from the true output is equal to sum over such outputs j of the probability that j is chosen divided by the equivalent sum over all outputs. Exploiting our metric d_v , this is then re-indexed over the set of possible quality scores.

We express the denominator as the sum of two terms $\sum_{\iota=0}^{n-\delta-1} \exp(\xi \iota) M(\vec{\sigma}, \iota) + \sum_{\iota=n-\delta}^n \exp(\xi \iota) M(\vec{\sigma}, \iota)$ which is minimized (thereby maximizing the total fraction) when the second term (the “good” outputs) consists only of the one correct answer.

$$\Pr[d_v((\vec{\sigma}, f(\vec{\sigma})), (\vec{\sigma}, g(\vec{\sigma}))) > \delta] \leq \frac{\sum_{\iota=0}^{n-\delta-1} \exp(\xi \iota) M(\vec{\sigma}, \iota)}{\sum_{\iota=0}^{n-\delta-1} \exp(\xi \iota) M(\vec{\sigma}, \iota) + \exp(\xi n)}$$

Finally, it is clear that the right hand side of the inequality is maximized when the numerator is maximized, and this occurs when there are $k - 1$ outputs with quality $q(\vec{\sigma}, j) = n - \delta - 1$.

¹In particular, for small values of δ (e.g., $\delta < 178$ when $\varepsilon = .05$ and $k = 3$) the original bound makes the trivial statement that $\Pr[d_v((\vec{\sigma}, f(\vec{\sigma})), (\vec{\sigma}, g(\vec{\sigma}))) > \delta] \leq 1$.

Therefore

$$\begin{aligned} \forall \vec{\sigma}, \Pr[d_v((\vec{\sigma}, f(\vec{\sigma})), (\vec{\sigma}, g(\vec{\sigma}))) > \delta] &\leq \frac{(k-1) \exp(\xi(n - \delta - 1))}{\exp(\xi n) + (k-1) \exp(\xi(n - \delta - 1))} \\ &= \frac{(k-1)}{(\varepsilon + 1)^{(\delta+1)/2} + (k-1)} \quad \square \end{aligned}$$

As before, we can express the quality of the exponential-approximations in terms of their asymptotic behavior.

Corollary 6.4. *Let $\varepsilon = \omega(1/n)$, $\beta, \mu > 0$, and let $f : (\Sigma_k)^n \rightarrow [k]$ be a deterministic voting rule. For sufficiently large n , there exists an ε -strategy-proof randomized voting rule g that is a $(\beta n, \mu)$ -approximation of f .*

Proof. Fix $\beta > 0$ and let $\delta = \beta n$. Let g be defined as in the proof of Theorem 4.1 and recall that (as shown in the previous proof) g is an ε -strategy-proof $(\beta n, \frac{(k-1)}{(\varepsilon+1)^{(\beta n+1)/2} + (k-1)})$ -approximation of f . Observe that

$$\begin{aligned} \lim_{n \rightarrow \infty} \frac{(k-1)}{(\varepsilon + 1)^{(\beta n+1)/2} + (k-1)} &= \lim_{n \rightarrow \infty} \frac{(k-1)}{((\varepsilon + 1)^{1/\varepsilon})^{\varepsilon(\beta n+1)/2} + (k-1)} \quad (1) \\ &= 0 \quad (2) \end{aligned}$$

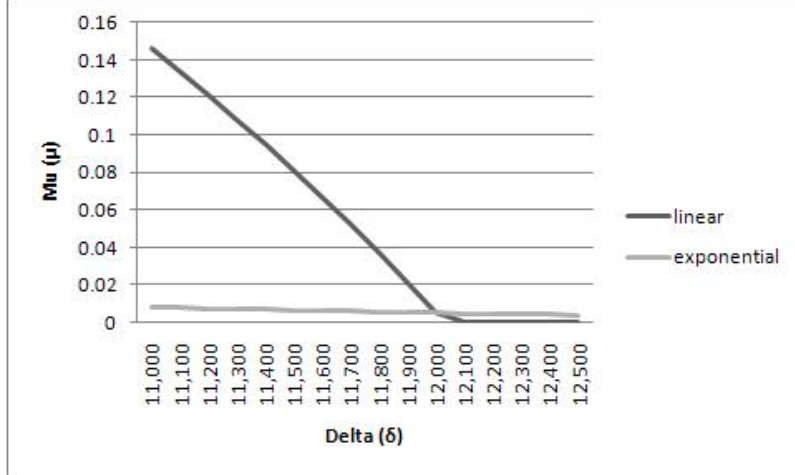
The first step is simple arithmetic. (2) follows by the following argument, which demonstrates that the denominator goes to infinity as n grows: If $\varepsilon = \Omega(1)$ then by definition $\varepsilon + 1$ is bounded below by some constant, so it suffices that $(\beta n + 1)/2$ goes to infinity, which is obvious. In the case where $\varepsilon = o(1)$, observe that since $\varepsilon = \omega(1/n)$, $\lim_{n \rightarrow \infty} \varepsilon(\beta n + 1)/2 = \infty$. It therefore suffices to show that $(\varepsilon + 1)^{1/\varepsilon}$ is bounded below by some constant. Since $\varepsilon = o(1)$, $\lim_{n \rightarrow \infty} (\varepsilon + 1)^{1/\varepsilon} = \lim_{\varepsilon \rightarrow 0} (\varepsilon + 1)^{1/\varepsilon} = e$.

It follows that for any β, μ , for sufficiently large n there exists an ε -strategy-proof $(\beta n, \mu)$ -approximation. $\square$

Example 6.5. For concreteness, consider the case where you have one hundred players and three outputs. Set $\varepsilon = .005$. The probability that g returns an incorrect answer is at most .67 (that is, any voting scheme $f : (\Sigma_3)^{100} \rightarrow [3]$ has a .05-strategy-proof $(0, .67)$ -approximation).

Alternatively, consider an election with 100 million voters and three outputs (about the scale of a United States presidential election). Again fixing $\varepsilon = .005$, the probability that g returns an answer further than 5000 votes from the correct answer (in practice, well within the vote corruption in such an election) is at most .00001 (i.e., any voting rule $f : (\Sigma_3)^{100,000,000} \rightarrow [3]$ has a .005-strategy-proof $(5000, .00001)$ -approximation). Looked at in another way, this says that in any election with 100 million voters and three outputs in which one outcome wins by at least 5000 votes (e.g., a typical national election), this mechanism will return the correct answer with probability at least .99999.

The relative guarantees achieved by these two constructions are shown in the following table; numbers are calculated with $\varepsilon = .001$, $n = 100,000,000$, and $k = 3$.



6.3 Extended Lower Bounds

In this section, we justify our original definition of approximation by showing two lower bounds that strictly limit the utility of alternative definitions; First, we extend the lower bound from Theorem 5.5 to show that our relaxed definition does not admit asymptotically better approximations. In particular, when $\varepsilon = o(1/n)$, we show that PLURALITY cannot be well-approximated even for $\mu > 0$. We then consider an alternative definition that has been considered previously in the context of 0-strategy-proofness: the requirement that the approximation g be equal to the original function with high probability (in our notation, approximations with $\delta = 0$). We extend the work of Gibbard and Satterthwaite to show that only dictatorial voting rules have ε -strategy-proof approximations with $\delta = 0$.

We first bound the asymptotic behavior of the relaxed approximations introduced in Appendix 6.

Theorem 6.6. *Let $n \geq k$, let $\beta > 0$. For sufficiently large n , PLURALITY does not have a trivial $(\beta n, 1 - 2/\sqrt{k})$ -approximation.*

Proof. This extended result follows immediately from the proof of Theorem 5.5 by setting $c = \sqrt{k}$. $\square$

We now consider the case where $\delta = 0$ and show that only dictatorial functions can be approximated in this manner. The key insight is to develop quantitative analogs of monotonicity and Pareto optimality and to show that any ε -strategy-proof voting rule g satisfies their guarantees. First employed introduced by Barberà and Peleg [1], these properties can be summarized as follows: *monotonicity* is the property that if a given preference profile $\vec{\sigma}$ returns an outcome j then any preference profile $\vec{\sigma}'$ with the property that all players i rank j at least as highly under $\vec{\sigma}'$ as under $\vec{\sigma}$ will also return the outcome j . *Pareto optimality* is the property that if all players rank outcome j above outcome j' then the voting rule will not return outcome j' . In their quantitative generalizations, the strict guarantees are relaxed to upper and lower bounds on the probability that a rule g returns the outcome j .

Lemma 6.7 (Monotonicity). *Let g be a ε -strategy-proof $(0, 1 - p)$ -approximation of the deterministic voting rule f . If $g(u) = a$ with probability at least p on some strategy profile*

$\vec{\sigma} \in (\Sigma_k)^n$ then for all strategy profiles $\vec{\sigma}' \in (\Sigma_k)^n$ such that $\forall x \in [k], i \in [n], \sigma'_i(a) \geq \sigma'_i(x)$ if $\sigma_i(a) \geq \sigma_i(x)$, $g(\vec{\sigma}') = a$ with probability at least p .

Proof. For all players i fix the u_i to be uniformly distributed between 0 and 1. That is the utility a player i receives if the outcome they ranked j th wins is $(k-j)/(k-1)$ (recall we can do this since the definition of strategy-proofness quantifies over all possible utility functions). Set $p = (1 + \varepsilon)(k-1)/k$ and observe that the following two properties hold for all $k \geq 3$:

(P1) If $\sigma_i(a) > \sigma_i(b)$ then $p \cdot u_i(a) > p \cdot u_i(b) + (1-p)u_i(x) + \varepsilon$ for all outputs $x \in [k]$

(P2) $p > 1/2$

Consider a sequence of preference profiles $\vec{v}_1, \dots, \vec{v}_\ell \in (\Sigma_k)^n$ ($\ell \leq n$) with the property that $\vec{v}_1 = \vec{\sigma}$, $\vec{v}_\ell = \vec{\sigma}'$, for all $j \in \{1, \dots, \ell-1\}$ $\vec{v}_j$ differs from $\vec{v}_{j+1}$ in exactly one place, and for all j , $v_{j+1,i}(a) \geq v_{j+1,i}(x)$ if $v_{j,i}(a) \geq v_{j,i}(x)$. Assume for contradiction that $\Pr[g(\vec{v}_\ell) = a] < p$. This implies that there exist $\vec{v}_j, \vec{v}_{j+1}$ such that $\Pr[g(\vec{v}_j) = a] \geq p$ and $\Pr[g(\vec{v}_{j+1}) = a] < p$. Since g is equal to f with probability $p > 1/2$ on all inputs, it follows that $\Pr[g(\vec{v}_{j+1}) = b] \geq p$ for some outcome $b \neq a$. Let i be the voter whose preferences change between $\vec{v}_j, \vec{v}_{j+1}$. Since g is ε -strategy proof and the parameters satisfy (P1), $v_{j,i}(b) \leq v_{j,i}(a)$ and hence by assumption $v_{j+1,i}(b) \leq v_{j+1,i}(a)$. On the other hand, ε -strategy-proofness combined with (P1) also implies that $v_{j+1,i}(b) \geq v_{j+1,i}(a)$, yielding a contradiction. $\square$

Lemma 6.8 (Pareto Optimality). *Let g be a ε -strategy-proof $(0, 1-p)$ -approximation of the deterministic voting rule f . If $\sigma_i(a) > \sigma_i(b)$ for all $i \in [n]$ then $\Pr[g(\vec{\sigma}) = b] < 1-p$.*

Proof. As in the proof of Lemma 6.7, fix the utility functions u_i to be uniformly distributed between 0 and 1 and observe that they satisfy properties (P1) and (P2) for $p = (1 + \varepsilon)(k-1)/k$. Assume for contradiction that $\Pr[g(\sigma) = b] \geq 1-p$. Since g is a $(0, 1-p)$ -approximation of f (that is $\Pr[g(\sigma) = f(\sigma)] \geq p$), this implies that $\Pr[g(\sigma) = b] \geq p$. Let $\sigma' \in (\Sigma_k)^n$ be a preference profile such that $\Pr[g(\sigma') = a] \geq p$ (such exists since f is onto). Define a third preference profile σ'' satisfying

$$\sigma''_i(a) > \sigma''_i(b) > \sigma''_i(x) \forall x \neq a, b, \forall i \in N$$

By monotonicity, $\Pr[g(\sigma'') = a] \geq \Pr[g(\sigma') = a] \geq p$. On the other hand, since $\sigma'_i(a) > \sigma'_i(b) \forall i \in N$, monotonicity also implies that $\Pr[g(\sigma'') = b] \geq \Pr[g(\sigma) = b] \geq p$. Since $p > 1/2$ (P2), this yields a contradiction. $\square$

Using these two lemmas, the lower bound follows by a proof analogous to that introduced by Svensson [16] for the exact case.

Lemma 6.9 ($n=2$). *Let $f : (\Sigma_k)^2 \rightarrow [k]$ be a deterministic, onto voting rule with $k \geq 3$ and let g be a $(0, 1-p)$ -approximation of f . If g is ε -strategy-proof then f has a dictator.*

Proof. Initialize $A = [k]$, $A_1 = A_2 = \emptyset$. While $|A| > 1$, let $a_0, a_1 \in A$ ($a_0 \neq a_1$), and define a preference profile $\vec{u} = (u_0, u_1)$ such that $u_i(a_i) > u_i(a_{1-i}) > u_i(x)$. By Pareto optimality, $\Pr[g(\vec{u}) = x] < 1-p$ for all $x \neq a_0, a_1$. Since g approximates f , it follows that $\exists i \in \{0, 1\}$ such that $\Pr[g(u) = a_i] \geq p$. Now define an alternative preference for the other player v_{1-i}

such that $v_{1-i}(a_{1-i}) > v_{1-i}(x) > v_{1-i}(a_i)$. By Pareto optimality, there still exists $j \in \{0, 1\}$ such that $\Pr[g(u_i, v_{1-i}) = a_j] \geq p$. By strategy-proofness, $j = i$. Monotonicity therefore implies that for all w such that $w_i(a_i) > w_i(x)$, $\Pr[g(w) = a_i] \geq p$. Add a_i to A_i and remove it from A .

After this process is complete, each set A_i contains outcomes for which player i is a p -dictator, this outcomes that satisfy the property that if player i ranks that outcome first, then g returns that outcome with probability at least p . Since $p > 1/2$, it is clear that one of the sets A_0, A_1 must be empty. Let A_i be the nonempty set.

We now claim that player i is a dictator for f . In order to show this, it is sufficient that for all $c \in A' \setminus A_i$, i is a p -dictator for c . Repeat the above argument setting $u_i(c) > u_i(a_i) > u_i(x)$ and $u_{1-i}(a_i) > u_{1-i}(c) > u_{1-i}(x)$. As before, either i is a p -dictator for c or $1 - i$ is a p -dictator for a . The later is impossible since we already showed that a_i is p -dominated by i . Since g is equal to f with probability at least p , the lemma follows. $\square$

Theorem 6.10. *If a deterministic, onto voting rule f with $k \geq 3$ outputs has a ε -strategy-proof $(0, 1 - (1 + \varepsilon)(k - 1)/k)$ -approximation, then it has a dictator.*

Proof. Let $p = (1 + \varepsilon)(k - 1)/k$ and let g be a ε -strategy-proof $(0, 1 - p)$ -approximation of the function f (the existence of such a function is guaranteed by the definition of approximate strategy-proofness). We proceed by induction, using the Lemma 6.9 as a base case. Define a function $h_1(u_1, v) = g(u_1, v, v, \dots, v)$. Since f is onto, for all outputs x there exists an input profile u such that $\Pr[g(u) = x] \geq p$, so it follows that there exists a preference ordering v that ranks x first such that $\Pr[h_1(u_1, v) = x] \geq p$ (by monotonicity). Moreover, since g is ε -strategy-proof and a p approximation of a deterministic function h_1 is also ε -strategy-proof and a p -approximation of a deterministic function. It follows by the previous lemma that h_1 has a p -dictator.

If the p -dictator is player 1: By monotonicity, 1 is also a p -dictator for g ; since g is a $(0, 1 - p)$ -approximation of f it follows that p is also a dictator for f .

If the p -dictator is player 2: Let u_1^* be an arbitrary fixed preference ordering for player 1 and define a probabilistic voting rule $h_2(u_2, \dots, u_n) = g(u_1^*, u_2, \dots, u_n)$. h_2 is also p -onto and ε -strategy-proof (this follows by the same argument as made for h_1). It follows by our induction hypothesis that h_2 has a p -dictator. Assume without loss of generality that 2 is a p -dictator for h . Define a final voting rule $q(u_1, u_2) = g(u_1, u_2, u_3^*, \dots, u_n^*)$, where $u_3^*, \dots, u_n^*$ are again fixed preference profiles. As before, onto and strategy-proofness follow from g . Player 1 cannot be a dictator for q since player 2 dictates when $u_1 = u_1^*$, so it follows that player 2 is a p -dictator for q . Since the fixed preferences were arbitrary, by monotonicity, 2 is a p -dictator for g and hence a dictator for f . $\square$

Example 6.11. Numerically, any voting rule with three outputs that has a .05-strategy-proof $(0, .3)$ -approximation has a dictator. By comparison, recall that Corollary 4.1 showed that every voting rule with three outputs has a .05-strategy-proof $(0, .67)$ -approximation.

References

- [1] S. Barberà and B. Peleg. Strategy-proof voting schemes with continuous preferences. *Social Choice and Welfare*, 7(1):31–38, 1990.

- [2] J. Bartholdi, C. Tovey, and M. Trick. The computational difficulty of manipulating an election. *Social Choice and Welfare*, 6:227–241, 1989.
- [3] V. Conitzer and T. Sandholm. Nonexistence of voting rules that are usually hard to manipulate. In *Proc. 21st AAAI Conference*, pages 627–634, 2006.
- [4] C. Dwork, F. McSherry, K. Nissim, and A. Smith. Calibrating noise to sensitivity in private data analysis. In *3rd TCC*, pages 265–284, 2006.
- [5] E. Friedgut, G. Kalai, and N. Nisan. Elections can be manipulated often. In *Proc. 49th FOCS*, pages 243–249, 2009.
- [6] Edward H. Friend. Sorting on electronic computer systems. *J. ACM*, 3:134–168, July 1956.
- [7] A. Gibbard. Manipulation of voting schemes: a general result. *Econometrica*, 41(4):587–601, 1973.
- [8] A. Gibbard. Manipulation of schemes that mix voting with chance. *Econometrica*, 45:665–681, 1977.
- [9] J. Halpern and R. Pass. Game theory with costly computation: Formulation and application to protocol security. In *ICS*, pages 120–142, 2010.
- [10] M. Isaksson, G. Kindler, and E. Mossel. The geometry of manipulation - a quantitative proof of the Gibbard Satterthwaite theorem. In *Proc. 50th FOCS*, 2010.
- [11] Jeff Kahn, Gil Kalai, and Nathan Linial. The influence of variables on boolean functions (extended abstract). In *FOCS*, pages 68–80, 1988.
- [12] F. McSherry and K. Talwar. Mechanism design via differential privacy. In *Proc. 48th FOCS*, 2007.
- [13] H. Moulin. On strategy-proofness and single peakedness. *Public Choice*, 35(4):437–455, 1980.
- [14] A. Procaccia. Can approximation circumvent Gibbard-Satterthwaite? In *Proc. 24th AAAI*, pages 836–841, 2010.
- [15] M. Satterthwaite. Strategy-proofness and arrow’s conditions: Existence and correspondence theorems for voting procedures and social welfare functions. *Journal of Economic Theory*, 10:187–217, 1975.
- [16] L.-G. Svensson. The proof of the Gibbard-Satterthwaite theorem revisited. *Working Paper No. 1999:1, Department of Economics, Lund University*, 1999.