

Strategyproof Randomized Social Choice for Restricted Sets of Utility Functions

Patrick Lederer

Technische Universität München

ledererp@in.tum.de

Abstract

When aggregating preferences of multiple agents, strategyproofness is a fundamental requirement. For randomized voting rules, so-called social decision schemes (SDSs), strategyproofness is usually formalized with the help of utility functions. In a central result, Gibbard [1977] characterizes the set of SDSs that are strategyproof with respect to all utility functions and shows that these SDSs are either indecisive or unfair. For finding more insights into the trade-off between strategyproofness and decisiveness, we propose the notion of U -strategyproofness which requires that only voters with a utility function in the set U cannot manipulate. In particular, we show that if the utility functions in U value the best alternative much more than other alternatives, there are U -strategyproof SDSs that choose an alternative with probability 1 whenever all but k voters rank it first. We also prove for rank-based SDSs that this large gap in the utilities is required to be strategyproof and that the gap must increase in k . On the negative side, we show that U -strategyproofness is incompatible with Condorcet-consistency if U satisfies minimal symmetry conditions and there are at least four alternatives. For three alternatives, the Condorcet rule can be characterized based on U -strategyproofness for the set U containing all equi-distant utility functions.

1 Introduction

When a group of agents wants to find a joint decision in a structured way, they can choose from a multitude of different voting rules. However, it is not clear which rule is the best one as each one has its benefits. This problem lies at the core of social choice theory which draws increased attention by computer scientists because it can be used to reason about computational multi-agent systems (see, e.g., [Chevalleyre *et al.*, 2007; Brandt *et al.*, 2013; Brandt *et al.*, 2016b; Endriss, 2017]). A fundamental requirement for voting rules is strategyproofness, i.e., agents should not be able to benefit by lying about their preferences. In a seminal result, Gibbard [1973] and Satterthwaite [1975] have shown that every

deterministic strategyproof voting rule is dictatorial if there are at least three different outcomes possible.

Randomization allows to escape this impossibility theorem, and we analyze therefore social decision schemes (SDSs). These functions aggregate the preferences of agents to lotteries over alternatives which determine for every alternative its winning chances. The final winner is then decided by chance according to these probabilities. While this model allows to circumvent many impossibilities, it is not straightforward how to define strategyproofness because the voters' preferences over lotteries are unclear. Maybe the most prominent approach is to assume that voters use cardinal utility functions on the alternatives to compare lotteries with respect to their expected utilities. However, voters still report ordinal preference relations to the SDS and hence, strategyproofness is defined by quantifying over utility functions: an SDS is strategyproof if voting honestly maximizes the expected utility for every voter and every utility function that is consistent with his true preferences. This strategyproofness notion, often called SD -strategyproofness, has been analyzed by Gibbard [1977] and Barberà [1979] who prove that all SD -strategyproof SDSs are indecisive because they almost always randomize over multiple alternatives. Even more, Benoît [2002] has shown that SD -strategyproofness is incompatible with the basic democratic idea that an alternative should be the winner of an election if an absolute majority of the voters report it as their best alternative.

While it is unfortunate that SD -strategyproofness does not allow for decisive SDSs, this strategyproofness notion seems also too demanding for because in many applications not all utility functions are plausible. For instance, when a representative body votes about budget proposals, it seems reasonable that similar proposals have similar utilities. Thus, we might neglect utility functions with a large gap between such options when discussing strategyproofness. This observation leads to the new notion of U -strategyproofness which requires that truth telling only maximizes the expected utility of a voter if his utility function is in the set U . Note that U -strategyproofness does not forbid utility functions $u \notin U$, but voters with such utility functions might be able to manipulate.

U -strategyproofness allows for a more detailed analysis than SD -strategyproofness because we can analyze the exact set of utility functions U for which an SDS is U -strategyproof. Conversely, we can also formulate strong im-

possibility results based on U -strategyproofness for severely restricted sets U and thus, we can pinpoint the source of manipulability far more detailed than with other strategyproofness notions. Hence, U -strategyproofness offers both the possibility of positive results by finding U -strategyproof SDSs for large sets U , and of strong impossibility results by using only a small number of utility functions. Furthermore, information about U -strategyproofness can also be valuable in practice: if the social planner can roughly guess the utility functions of the voters, he might be able to choose an SDS preventing manipulations. Even if the social planner does not have such insights, he might opt for an SDS that is U -strategyproof for a large set U as such an SDS is immune to manipulations from most voters.

Other than introducing U -strategyproofness, we use this new notion to investigate the trade-off between strategyproofness and decisiveness. On the positive side, we show that there are U -strategyproof SDSs that assign an alternative probability 1 whenever all but $k > 0$ voters agree that it is the best option if the utility functions in U value the best alternative much more than the other alternatives. Moreover, we prove for rank-based SDSs that this gap in the utility functions is required to be strategyproof and that it must increase in k . On the other hand, we show that Condorcet-consistency is incompatible with U -strategyproofness if the set U satisfies minimal symmetry conditions between preference relations and there are $m \geq 4$ alternatives. If there are only three alternatives and an odd number of voters, the Condorcet rule is characterized by U -strategyproofness for the set U of all equi-distant utility functions and Condorcet-consistency. The proofs of these theorems and of all propositions are omitted because of space limitations.

2 Related Work

To our knowledge, we are the first authors who explicitly investigate U -strategyproofness. Nevertheless, ideas similar to U -strategyproofness have been used before. For instance, Sen [2011] and Mennle and Seuken [2021] define strategyproofness by considering restricted sets of utility functions and thus, their works can be interpreted as first results on U -strategyproofness. Moreover, in set-valued social choice (where the outcome of an election is a non-empty set of alternatives instead of a lottery) preferences over sets of alternatives are often derived from utility functions. For instance, Duggan and Schwartz [2000] and Benoît [2002] employ this approach to motivate their strategyproofness notions. The relationship between these results and U -strategyproofness is discussed in more detail in Section 4.

There are also various results on other strategyproofness notions in randomized social choice (see, e.g., [Gibbard, 1977; Hoang, 2017; Aziz *et al.*, 2018; Brandl *et al.*, 2018]), many of which are surveyed by Brandt [2017]. These results either prove the incompatibility of strategyproofness with other axioms or characterize specific SDSs. Our results differ from previous ones as we investigate a different question: instead of asking whether an SDS is strategyproof according to some definition, we ask for which utility functions it is strategyproof.

Moreover, strategyproofness is often considered for restricted domains of preference profiles (see, e.g., [Ehlers *et al.*, 2002; Bogomolnaia *et al.*, 2005; Elkind *et al.*, 2017; Chatterji and Zeng, 2018]). For instance, Bogomolnaia *et al.* [2005] discuss an attractive SD -strategyproof SDS for dichotomous preferences. U -strategyproofness can be interpreted similarly, but we focus on utility functions instead of preference profiles: U -strategyproof SDSs are immune to manipulations if we only allow utility functions in U .

Another field related to U -strategyproofness is cardinal social choice, where the input of social decision schemes consists of the utility functions of the voters. If we allow all utility functions as input, every strategyproof cardinal SDS is, under mild additional assumptions, a variant of a random dictatorship (see, e.g., [Hylland, 1980; Dutta *et al.*, 2007; Nandeibam, 2013]). As noted by Dutta *et al.* [2007], these negative results break down if the domain of cardinal SDSs is restricted, but this setting is not well understood. Our results provide insights in this problem because every U -strategyproof SDS can be interpreted as a cardinal SDS that is strategyproof on the domain U .

Finally, note that our model assumptions are quite similar to those used in the analysis of the distortion of SDSs (see, e.g., [Procaccia and Rosenschein, 2006; Gross *et al.*, 2017; Abramowitz *et al.*, 2019]). Just as these authors, we assume that voters only report ordinal preferences but use utility functions to evaluate the quality of a lottery. Whereas distortion focuses on the welfare of SDSs, we investigate their resistance to strategic behavior of voters.

3 Preliminaries

Let $N = \{1, \dots, n\}$ be a finite set of voters and let A be a set containing m alternatives. A *preference relation* is an anti-symmetric, transitive, complete, and reflexive binary relation on A and R_i denotes the preference relation of voter i . We compactly represent preference relations as comma-separated lists. Let $\mathcal{R}$ denote the set of all preference relations on A . A *preference profile* R is an n -tuple containing the preference of every voter $i \in N$, i.e., $R \in \mathcal{R}^n$. When writing preference profiles, we indicate the corresponding voter directly before the preference relation to clarify which voter submits which preference relation. For example, $1 : a, b, c$ indicates that voter 1 reports that he prefers a to b to c .

In this paper, we discuss *social decision schemes* (SDSs), which are functions that map preference profiles to lotteries on A . A *lottery* p is a function from the set of alternatives A to the interval $[0, 1]$ such that $\sum_{x \in A} p(x) = 1$. Let $\Delta(A)$ denote the set of all lotteries on A . Formally, a social decision scheme is a function $f : \mathcal{R}^n \rightarrow \Delta(A)$ and we denote with $f(R, x)$ the probability assigned to x by the lottery $f(R)$.

The definition of SDSs allows for a huge variety of functions, some of which seem not desirable. Therefore, we introduce axioms to narrow down the set of SDSs. Two basic fairness axioms are anonymity and neutrality, which require that voter and alternatives, respectively, are treated equally. More formally, an SDS f is *anonymous* if $f(R) = f(\pi(R))$ for all profiles R and permutations $\pi : N \rightarrow N$, and *neutral* if $f(R, x) = f(\tau(R), \tau(x))$ for all alternatives $x \in A$,

profiles R , and permutations $\tau : A \rightarrow A$. Another natural axiom is *unanimity*, which requires of an SDS f that $f(R, x) = 1$ for all preference profiles R in which all voters agree that x is the best choice. While this axiom is so weak that is often considered indisputable, it is also irrelevant in practice as ballots are usually not unanimous. Therefore, we introduce the stronger notion of *k-unanimity*: an SDS f is *k-unanimous* if $f(R, x) = 1$ whenever $n - k$ or more voters report x as the best alternative. By definition, unanimity is equal to 0-unanimity and note that *k-unanimity* is only well-defined if $k < \frac{n}{2}$. A well-known strengthening of *k-unanimity* is Condorcet-consistency. For defining this axiom, let $n_{xy}(R) = |\{i \in N : xR_i y\}| - |\{i \in N : yR_i x\}|$ denote the *majority margin* between two alternatives $x, y \in A$ in the preference profile R . An alternative x is the *Condorcet winner* in a preference profile R if $n_{xy}(R) > 0$ for all other alternatives $y \in A \setminus \{x\}$. Less formally, an alternative x is the Condorcet winner if it is preferred to every other alternative by a majority of the voters. Finally, an SDS f is *Condorcet-consistent* if $f(R, x) = 1$ for all profiles R and alternatives $x \in A$ such that x is the Condorcet winner in R .

An important class of SDSs are rank-based SDSs. The basic idea of these schemes is that voters assign ranks to the alternatives and that an SDS should only rely on these ranks, but not on which voter assigns which rank to an alternative. For formalizing this concept, we denote with $r(R_i, x) = |\{y \in A : yR_i x\}|$ the *rank* of alternative x in voter i 's preference relation. Moreover, we define the *rank vector* $r^*(R, x)$ as the vector that contains the rank of x with respect to every voter in increasing order, i.e., $r^*(R, x)_i \leq r^*(R, x)_{i+1}$ for all $i \in \{1, \dots, n-1\}$, and the *rank matrix* $r^*(R)$ as the matrix that contains the rank vectors of all alternative as rows. Finally, we call an SDS f *rank-based* if it only depends on the rank matrix, i.e., $f(R) = f(R')$ for all preference profiles R, R' with $r^*(R) = r^*(R')$. The set of rank-based SDSs contains many prominent functions such as point scoring rules and anonymous SDSs that only depend on the first-ranked alternatives of the voters.

4 U-Strategyproofness

A central problem in social choice is that of manipulability: voters may lie about their preferences to achieve a better outcome. While the definition of a manipulation is easy if an SDS never randomizes between multiple alternatives, it is not clear how to compare non-degenerate lotteries. A classical approach for this problem is to assume that voters are endowed with *utility functions* $u_i : A \rightarrow \mathbb{R}$. We impose the constraint that no voter assigns the same utility to two alternatives, i.e., $u_i(x) \neq u_i(y)$ for all voters $i \in N$ and alternatives $x, y \in A$, to ensure that the ordinal preference relation induced by a utility function is anti-symmetric. We denote with $\mathcal{U}$ the set of all such utility functions and say that a utility function $u \in \mathcal{U}$ is *consistent* with a preference relation R if $u(x) \geq u(y)$ iff xRy for all alternatives $x, y \in A$. Finally, each voter i uses his utility function u_i to compare lotteries by their *expected utilities* $\mathbb{E}[p]_{u_i} = \sum_{x \in A} p(x)u_i(x)$, i.e., voter i prefers lottery p weakly to lottery q if $\mathbb{E}[p]_{u_i} \geq \mathbb{E}[q]_{u_i}$.

Even though we assume the existence of utility functions,

voters only report ordinal preferences. Consequently, strategyproofness is often defined by quantifying over utility functions. In particular, Gibbard [1977] employs this approach to define *SD-strategyproofness*: an SDS f is *SD-strategyproof* if $\mathbb{E}[f(R)]_{u_i} \geq \mathbb{E}[f(R')]_{u_i}$ for all voters $i \in N$, preference profiles R, R' , and *utility functions* $u_i \in \mathcal{U}$ such that u_i is consistent with R_i and $R_j = R'_j$ for all $j \in N \setminus \{i\}$. While *SD-strategyproofness* allows for strong negative results (see, e.g. [Gibbard, 1977; Barberà, 1979]), it lacks relevance for many practical applications as not all utility functions are plausible. Also, *SD-strategyproofness* provides often only shallow theoretical insights as it is not possible to pinpoint the source of manipulability.

In order to address these problems, we introduce a new strategyproofness notion by restricting the set of feasible utility functions U beforehand: an SDS f is *U-strategyproof* if $\mathbb{E}[f(R)]_{u_i} \geq \mathbb{E}[f(R')]_{u_i}$ for all voters $i \in N$, preference profiles R, R' , and *utility functions* $u_i \in U$ such that u_i is consistent with R_i and $R_j = R'_j$ for all $j \in N \setminus \{i\}$. Less formally, *U-strategyproofness* only requires that voters with a utility function in U cannot increase their expected utility by misrepresenting their preferences. Hence, *U-strategyproofness* is equal to *SD-strategyproofness* and smaller sets of utility functions result in less demanding strategyproofness notions. Note that *U-strategyproofness* solves both problems of *SD-strategyproofness*: we can investigate whether an SDS is manipulable in practice by dismissing implausible utility functions, and we can find the core of impossibility results by determining the minimally required set of utility functions. Next, we discuss an example to illustrate the difference between *U-strategyproofness* and *SD-strategyproofness*.

Example 1. Consider the profiles R^1 and R^2 shown below and let f denote an SDS such that $f(R^1, x) = \frac{1}{3}$ for $x \in \{a, b, c\}$ and $f(R^2, b) = 1$. Moreover, consider the utility functions u_1, u_2 , and u_3 with $u_1(a) = 2, u_1(b) = 1, u_1(c) = 0, u_2(a) = 3, u_2(b) = 1, u_2(c) = 0, u_3(a) = 3, u_3(b) = 2$, and $u_3(c) = 0$. These utility functions are only consistent with voter 1's preference relation in R^1 , and thus, we can check whether this voter can benefit by deviating to R^2 . A quick calculation shows that $\mathbb{E}[f(R^1)]_{u_1} = 1 = \mathbb{E}[f(R^2)]_{u_1}$, $\mathbb{E}[f(R^1)]_{u_2} = \frac{4}{3} > 1 = \mathbb{E}[f(R^2)]_{u_2}$, and $\mathbb{E}[f(R^1)]_{u_3} = \frac{5}{3} < 2 = \mathbb{E}[f(R^2)]_{u_3}$. Hence, voter 1 can increase his expected utility if his utility function is u_3 and thus, f is *SD-manipulable*. In contrast, voter 1 does not benefit from deviating to R^2 if his utility function is u_1 or u_2 . Since the preferences of the other voters are not consistent with u_1, u_2 , and u_3 , it follows that f is $\{u_1, u_2\}$ -strategyproof on these two profiles.

R^1 :	1: a, b, c	2: b, c, a	3: c, a, b
R^2 :	1: b, a, c	2: b, c, a	3: c, a, b

In our results, we always consider *U-strategyproofness* for *symmetric* sets U , i.e., we assume that $u \in U$ implies that $u^\pi = u \circ \pi \in U$ for every permutation π on A . This formalizes the natural condition that all preference relations should be treated equally. Moreover, the symmetry condition is rather weak since every neutral SDS is *U'-strategyproof* for a symmetric set U' if it is *U-strategyproof* for a set $U \neq \emptyset$.

Proposition 1. *If a neutral SDS is U -strategyproof for a set $U \neq \emptyset$, it is U' -strategyproof for a symmetric set U' with $U \subseteq U'$.*

A special case of our symmetry assumption is that U consists of a single utility function u and its renamings, i.e., that $U = \{u \circ \pi : \pi \in \Pi\}$, where Π denotes the set of all permutations on A . In this case, we write u^Π -strategyproofness instead of U -strategyproofness. Note that u^Π -strategyproofness associates every preference relation with exactly one utility function, whereas $\{u\}$ -strategyproofness, i.e., strategyproofness for a single utility function u , only affects a single preference relation. Since the utility of an alternative only depends on its rank for u^Π -strategyproofness, we often write $u(k)$ to denote the utility of the k -th best alternative of a voter. As the next proposition shows, it suffices to consider u^Π -strategyproofness or even $\{u\}$ -strategyproofness because for every SDS f and every preference relation R_i , the set of utility functions u that are consistent with R_i and for which f is strategyproof is convex.

Proposition 2. *For every SDS f and preference relation R_i , the set $U_{R_i} = \{u \in \mathcal{U} : u \text{ is consistent with } R_i \text{ and } f \text{ is } \{u\}\text{-strategyproof}\}$ is convex.*

We can use this proposition to show that an SDS is U -strategyproof for a large set U by proving that it is u_i^Π -strategyproof for a few utility functions $u_i \in \{u_1, \dots, u_l\}$. Assuming that $u_1, \dots, u_l$ are all consistent with a preference relation R_i , it follows then from Proposition 2 that the SDS is $\hat{u}^\Pi$ -strategyproof for every utility function $\hat{u}$ that can be represented as a convex mixture of $u_1, \dots, u_l$, which means that it is U -strategyproof for a large set U .

Next, note that U -strategyproofness inherits many attractive properties from SD -strategyproofness: for instance, the convex combination of U -strategyproof SDSs is itself U -strategyproof, i.e., the set of U -strategyproof SDSs is convex for every set U . As a consequence of this observation, it is often possible to construct an anonymous U -strategyproof SDS based on a non-anonymous U -strategyproof SDS. Another similarity between U -strategyproofness and SD -strategyproofness is that both axioms disincentivize even manipulations from groups of voters with the same preferences.

Finally, observe that U -strategyproofness can be used to transfer results from set-valued social choice to the probabilistic setting. We explain this relation using the impossibility result of Benoît [2002] as example. This theorem states that strategyproofness is incompatible with 1-unanimity for set-valued social choice functions if voters prefer every subset of their best two alternatives to every other set and other in our model negligible conditions are satisfied. For formulating this result for SDSs, we have to compare lotteries only based on their support $\text{supp}(p) = \{x \in A : p(x) > 0\}$. Hence, let $\epsilon_f = \min_{x \in A, R \in \mathcal{R}^n : f(R, x) > 0} f(R, x)$ denote the smallest non-zero probability assigned to an alternative by the SDS f and note that ϵ_f is well-defined since SDSs are defined for a fixed set of alternatives and voters. Given this probability, we derive that every voter whose utility function u satisfies $u(2) > (1 - \epsilon_f)u(1) + \epsilon_f u(3)$ prefers every lottery that randomizes only over his best two alternatives to every other lottery. After rearranging this equation, we can formu-

late Benoît's impossibility as follows.

Proposition 3. *No SDS f satisfies both u^Π -strategyproofness and 1-unanimity if $u(1) - u(2) < \frac{\epsilon_f}{1 - \epsilon_f}(u(2) - u(3))$, $m \geq 3$, and $n \geq 3$.*

Note that Proposition 3 highlights the central requirement of Benoît's impossibility theorem: voters must be close to indifferent between their best two alternatives. This refines Benoît's reasoning who justifies his strategyproofness notion with voters who "like his or her two favorite alternatives "much more" than the rest of the alternatives".¹ Based on this approach, we can also formalize other impossibility results from set-valued social choice with U -strategyproofness.

5 Results

In the sequel, we employ U -strategyproofness to analyze the trade-off between strategyproofness and decisiveness. In particular, we investigate two decisiveness axioms: k -unanimity and Condorcet-consistency. The first axiom allows for positive results if suitable utility functions are considered, whereas Condorcet-consistency is incompatible with u^Π -strategyproofness for every utility function $u \in \mathcal{U}$.

5.1 k -unanimity

A central result of Gibbard [1977], who attributes it to Hugo Sonnenschein, is that the SDS called random dictatorship (henceforth RD) is the only SD -strategyproof SDS that satisfies unanimity and anonymity. This SDS assigns an alternative x in a profile R the probability $\frac{PL(R, x)}{n}$, where $PL(R, x) = |\{i \in N : \forall y \in A : x R_i y\}|$ denotes the plurality score of alternative x . A common method for executing RD is to choose a voter uniformly at random and to return his most preferred alternative as winner. While RD is one of the most attractive SD -strategyproof SDSs, it violates k -unanimity for $k > 0$. Even more, Benoît [2002] has shown that every SD -strategyproof SDS fails k -unanimity for $k > 0$.

However, we can define a variant of RD that satisfies both k -unanimity for an arbitrary $k \in \{0, \dots, \lfloor \frac{n-1}{2} \rfloor\}$ and U -strategyproofness for a large set of utility functions U . Hence, consider the following SDS, which we call k -random dictatorship (abbreviated by RD^k): if at least $n - k$ voters agree that alternative x is the best choice, assign alternative x a probability of 1; otherwise, return the outcome of RD . As we show in Theorem 1, RD^k satisfies U -strategyproofness for $U = \{u \in \mathcal{U} : u(1) - u(2) \geq k(u(2) - u(m))\}$, i.e., if voters have a strong preference for the first alternative, RD^k is strategyproof. Unfortunately, the definition of U depends on k , i.e., for large values of k , there must be an extremely large gap between $u(1)$ and $u(2)$. Another variant of RD , which we refer to as $OMNI^*$, solves this problem. This SDS assigns probability 1 to an alternative x if more than half of the voters report x as their best alternative, and otherwise randomizes uniformly among all alternatives that are

¹Benoît [2002] also discusses a variant for SDSs in which he uses the minimal non-zero probability assigned to an alternative. However, Benoît only gives an example showing that there is a suitable utility function such that the required preferences over sets extend to preferences over lotteries.

at least once top-ranked. This SDS is U -strategyproof for $U = \{u \in \mathcal{U}: u(1) - u(2) \geq \sum_{i=3}^m u(2) - u(i)\}$. While $OMNI^*$ satisfies $\lfloor \frac{n-1}{2} \rfloor$ -unanimity for all numbers of voters and alternatives, the condition on U seems only realistic if there are few alternatives.

Theorem 1. For every $k \in \{1, \dots, \lfloor \frac{n-1}{2} \rfloor\}$, RD^k satisfies U -strategyproofness for $U = \{u \in \mathcal{U}: u(1) - u(2) \geq k(u(2) - u(m))\}$ and violates $\{u\}$ -strategyproofness for every utility function $u \notin U$. Moreover, $OMNI^*$ satisfies U -strategyproofness for $U = \{u \in \mathcal{U}: u(1) - u(2) \geq \sum_{i=3}^m u(2) - u(i)\}$ and violates $\{u\}$ -strategyproofness for every utility function $u \notin U$.

The constraint on the set U for RD^k arises naturally by considering the preference profile in which $n - k - 1$ voters top-rank the second best alternative of voter i and the remaining k voters top-rank voter i 's least preferred alternative. In this situation, voter i can ensure that his second best alternative is chosen with probability 1 by reporting it as his best one. Solving the corresponding inequality required by U -strategyproofness leads to the bound on U . A similar worst-case analysis can be applied for $OMNI^*$.

While it is positive that k -unanimity and U -strategyproofness can be simultaneously satisfied at all, the bounds on the sets U in Theorem 1 become increasingly worse with large k and m . This raises the question for less demanding bounds on the utility functions. As our next theorem shows, the approach used for defining RD^k and $OMNI^*$ has not much space for improvement as both SDSs are rank-based.

Theorem 2. There is no rank-based SDS that satisfies u^Π -strategyproofness and k -unanimity for $0 < k < \frac{n}{2}$ if $m \geq 3$, $n \geq 3$, and $u(1) - u(2) < \sum_{i=\max(3, m-k+1)}^m u(2) - u(i)$.

The proof of Theorem 2 works by contradiction: we assume that there is a k -unanimous rank-based SDS f that satisfies u^Π -strategyproofness for a utility function u with $u(1) - u(2) < \sum_{i=\max(3, m-k+1)}^m u(2) - u(i)$. Moreover, let $k^* = \min(k, m-2)$. Our analysis then starts at a profile R where $n - k^*$ voters favor a the most, which implies that $f(R, a) = 1$ due to k -unanimity. The central argument is a rather involved construction that shows that a voter can weaken alternative a from the first rank to the second one without affecting the outcome. By repeatedly applying this construction, we eventually arrive at a profile R' where only k^* voters top-rank a and the remaining voters top-rank b , but $f(R', a) = 1$. This is in conflict with k -unanimity as $n - k^* \geq n - k$ voters report b as best choice but $f(R', b) \neq 1$.

Remark 1. A computer-aided approach has shown that there are rather technical SDSs that satisfy k -unanimity and u^Π -strategyproofness for utility functions u with $u(1) - u(2) < \sum_{i=\max(3, m-k+1)}^m u(2) - u(i)$ if we dismiss rank-basedness and $m \leq 4$. Hence, rank-basedness is required for Theorem 2. Moreover, most bounds of the theorem are tight: if $m = 2$, $OMNI^*$ and RD^k are even SD -strategyproof, and if $n = 2$, k -unanimity is not well-defined for $k > 0$. Furthermore, the condition on the utility functions is almost tight: RD^1 shows that the bound is tight for 1-unanimity, and $OMNI^*$ shows that the bound is tight if $k \geq m - 2$.

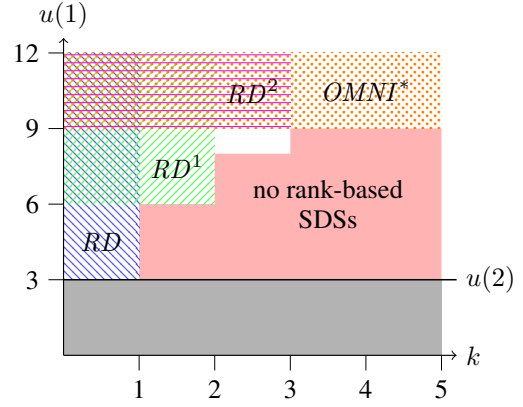


Figure 1: Illustration of Theorem 1 and Theorem 2. We assume that there are 5 alternatives and consider a utility function u with $u(2) = 3$, $u(3) = 2$, $u(4) = 1$, and $u(5) = 0$. The figure shows for which values of $u(1)$ the SDSs RD (blue area), RD^1 (green area), RD^2 (magenta area), and $OMNI^*$ (orange area) are u^Π -strategyproof on the vertical axis. The horizontal axis illustrates the values of k for which these SDSs are k -unanimous. The red area displays the impossibility of Theorem 2 and the gray area marks the values of $u(1)$ with $u(1) < u(2)$.

Finally, RD^k shows that no constraint of the type $u(1) - u(2) \leq \sum_{i=m-k+1}^m u(2) - u(i) + \epsilon$ with $\epsilon > 0$ can result in an impossibility because we can always find a utility function u such that $\sum_{i=m-k+1}^m u(2) - u(i) + \epsilon \geq u(1) - u(2) \geq k(u(2) - u(m))$ by making the difference between $u(i)$ and $u(m)$ for $i \geq 3$ sufficiently small. Nevertheless, it remains open to find rank-based SDSs that satisfy U -strategyproofness and k -unanimity for $U = \{u \in \mathcal{U}: u(1) - u(2) = \sum_{i=m-k+1}^m u(2) - u(i)\}$ and $2 \leq k \leq m - 3$.

Remark 2. Theorem 1 and Theorem 2 have an intuitive interpretation: if voters strongly prefer their best alternative, it becomes possible to achieve strategyproofness and decisiveness. This follows as strategyproofness is compatible with k -unanimity if there is a sufficiently large gap between $u(1)$ and $u(2)$. In contrast, it is impossible that an SDS satisfies both axioms if voters are close to indifferent between their best two alternatives. For the class of general SDSs, this is shown by Benoît [2002], and for the class of rank-based SDSs, Theorem 2 significantly weakens the requirements on the utility functions.

Remark 3. Figure 1 illustrates the results of this section. For this figure, we assume that there are 5 alternatives and a large number of voters $n \geq 11$, and we fix all utilities but $u(1)$. Hence, we can compute the values of $u(1)$ for all SDSs of Theorem 1 such that the considered SDS is u^Π -strategyproof. The figure shows that for RD^k , the required value of $u(1)$ increases in k and the bound of $OMNI^*$ is independent of k . Moreover, the required values of $u(1)$ are quite large compared to $u(2)$ for all SDSs but RD . However, the red area shows the values of $u(1)$ for which Theorem 2 applies and hence, these large values are indeed required. The white area shows that there is a small gap between the positive results in Theorem 1 and the impossibility in Theorem 2.

5.2 Condorcet-consistency

As there are even rank-based SDS that are k -unanimous and U -strategyproof for large sets U , the question arises whether stronger decisiveness notions can be achieved by dismissing rank-basedness. Unfortunately, we find a negative answer to this question by considering Condorcet-consistency.

Theorem 3. *There is no Condorcet-consistent SDS that satisfies u^Π -strategyproofness regardless of the utility function u if $m \geq 4$, $n \geq 5$ and $n \neq 6$, $n \neq 8$.*

The proof of this result works by contradiction and relies on a case distinction on the utility function u . If $u(1) - u(2) < u(2) - u(m)$, the utility of the second best alternative is larger than the average utility, which means that a voter can manipulate by making his second best alternative into the Condorcet winner. If $u(1) - u(m-1) > u(m-1) - u(m)$, voters value their second worst alternative less than the uniform lottery. As a consequence, there is a voter who can manipulate by weakening his second worst alternative such that it is no longer the Condorcet winner. Finally, note that these two cases are exhaustive: the strictness of the utility function u entails that $u(m-1) - u(m) < u(1) - u(m-1)$ if $u(1) - u(2) \geq u(2) - u(m)$ and $m \geq 4$.

A close inspection of the proof shows that the impossibility also holds if $m = 3$ unless U only contains equi-distant utility functions, i.e., utility functions with $u(1) - u(2) = u(2) - u(3)$. This raises the question whether there is a U -strategyproof SDS that satisfies Condorcet-consistency in this special case. Indeed, the Condorcet rule (abbreviated by *COND*), which assigns probability 1 to the Condorcet winner whenever it exists and returns the uniform lottery over all alternatives otherwise, satisfies U -strategyproofness for this set. Even more, the Condorcet rule is uniquely characterized by these axioms if n is odd.

Theorem 4. *COND is the only Condorcet-consistent SDS that satisfies U -strategyproofness for $U = \{u \in \mathcal{U} : u(1) - u(2) = u(2) - u(3)\}$ if $m = 3$ and n is odd.*

It is easy to show that the Condorcet-rule is U -strategyproof for $U = \{u \in \mathcal{U} : u(1) - u(2) = u(2) - u(3)\}$ if $m = 3$ because the uniform lottery on all three alternatives has for every voter the expected utility of $u(2)$. Hence, the proof mainly focuses on why no other Condorcet-consistent SDS f satisfies U -strategyproofness for this set U . For this, we show that there is a profile R and a voter i such that voter i 's expected utility $\mathbb{E}[f(R)]_u$ is less than $u(2)$. Moreover, this voter can either make his second best alternative into the Condorcet winner or revert to a preference profile in which each alternative is chosen with a probability of $\frac{1}{3}$. As both cases yield an expected utility of $u(2)$ for voter i , we have found a contradiction to U -strategyproofness.

Remark 4. The Condorcet rule is also U -strategyproof for the set of equi-distant utility functions if $m = 3$ and n is even. However, other SDSs satisfy Condorcet-consistency and U -strategyproofness for even n , too. For instance, the SDS that assigns the Condorcet winner probability 1 whenever it exists and uniformly randomizes among the top-ranked alternatives otherwise satisfies also all required axioms. The proof for this claim relies on the insight that every voter has a utility of at least $u(2)$ in the absence of a Condorcet winner.

Remark 5. A well-known class of SDSs are tournament solutions which only depend on the majority relation $R_M = \{(x, y) \in A^2 : n_{xy}(R) \geq n_{yx}(R)\}$ of the input profile R to compute the outcome. For these SDSs, unanimity and u^Π -strategyproofness entail Condorcet-consistency. Thus, there are no unanimous and u^Π -strategyproof tournament solutions, regardless of the utility function u , if $m \geq 4$. This is in harsh contrast to results for set-valued social choice, where attractive tournament solutions satisfy various strategyproofness notions (see, e.g., [Brandt et al., 2016a]).

Remark 6. The proof of Theorem 3 also reveals more insights about the compatibility of k -unanimity and u^Π -strategyproofness for general SDSs. In particular, the first case shows that no $\lceil \frac{n}{3} \rceil$ -unanimous SDS can be u^Π -strategyproof for a utility function u with $u(1) - u(2) < u(2) - u(m)$ if $m \geq 4$ and $n \geq 3$.

6 Conclusion and Discussion

We study a new strategyproofness notion called U -strategyproofness. Whereas the common notion of SD -strategyproofness is derived by quantifying over all utility functions, U -strategyproofness is derived by quantifying only over the utility functions in a specified set U . This new strategyproofness notion arises from practical observations as often not all utility functions are plausible, and also has theoretical advantages because it allows for a much finer analysis than SD -strategyproofness. Furthermore, we analyze the compatibility of U -strategyproofness and decisiveness axioms such as k -unanimity and Condorcet-consistency. In particular, we discuss SDSs that satisfy k -unanimity for any k with $0 < k < n/2$ and U -strategyproofness if the set U only contains utility functions u for which $u(1) - u(2)$ is sufficiently large. Moreover, we show for rank-based SDSs that the large gap between $u(1)$ and $u(2)$ is required to be strategyproof and has to increase in k . We also prove that U -strategyproofness is incompatible with Condorcet-consistency if the set U is symmetric and $m \geq 4$. This impossibility also holds if $m = 3$ unless the utility functions in U are equi-distant. In this special case and if n is odd, the Condorcet rule can be characterized by U -strategyproofness and Condorcet-consistency.

Our results have a very intuitive interpretation: strategyproofness is only compatible with decisiveness if each voter has a clear best alternative. Even more, the more decisiveness is required, the stronger voters have to favor their best alternative. This conclusion is highlighted by Theorems 1 and 2 as well as the impossibility of Benoît [2002]. Moreover, it coincides with the informal argument that it is easier to manipulate for a voter who deems many alternatives acceptable as he can just report another acceptable alternative as his best one. Hence, our results show that the main source of manipulability are voters who are close to indifferent between some alternatives.

Acknowledgments

This work was supported by the Deutsche Forschungsgemeinschaft under grant BR 2312/12-1. I thank the anonymous reviewers and Felix Brandt for helpful comments.

References

- [Abramowitz *et al.*, 2019] B. Abramowitz, E. Anshelevich, and W. Zhu. Awareness of voter passion greatly improves the distortion of metric social choice. In *Proceedings of the 15th International Conference on Web and Internet Economics*, pages 3–16. Springer, 2019.
- [Aziz *et al.*, 2018] H. Aziz, F. Brandl, F. Brandt, and M. Brill. On the tradeoff between efficiency and strategyproofness. *Games and Economic Behavior*, 110:1–18, 2018. Preliminary results appeared in the Proceedings of AAAI-2013 and AAMAS-2014.
- [Barberà, 1979] S. Barberà. Majority and positional voting in a probabilistic framework. *Review of Economic Studies*, 46(2):379–389, 1979.
- [Benoît, 2002] J.-P. Benoît. Strategic manipulation in voting games when lotteries and ties are permitted. *Journal of Economic Theory*, 102(2):421–436, 2002.
- [Bogomolnaia *et al.*, 2005] A. Bogomolnaia, H. Moulin, and R. Stong. Collective choice under dichotomous preferences. *Journal of Economic Theory*, 122(2):165–184, 2005.
- [Brandl *et al.*, 2018] F. Brandl, F. Brandt, M. Eberl, and C. Geist. Proving the incompatibility of efficiency and strategyproofness via SMT solving. *Journal of the ACM*, 65(2):1–28, 2018. Preliminary results appeared in the Proceedings of IJCAI-2016.
- [Brandt *et al.*, 2013] F. Brandt, V. Conitzer, and U. Endriss. Computational social choice. In G. Weiß, editor, *Multiagent Systems*, chapter 6, pages 213–283. MIT Press, 2nd edition, 2013.
- [Brandt *et al.*, 2016a] F. Brandt, M. Brill, and P. Harrenstein. Tournament solutions. In F. Brandt, V. Conitzer, U. Endriss, J. Lang, and A. D. Procaccia, editors, *Handbook of Computational Social Choice*, chapter 3. Cambridge University Press, 2016.
- [Brandt *et al.*, 2016b] F. Brandt, V. Conitzer, U. Endriss, J. Lang, and A. Procaccia, editors. *Handbook of Computational Social Choice*. Cambridge University Press, 2016.
- [Brandt, 2017] F. Brandt. Rolling the dice: Recent results in probabilistic social choice. In U. Endriss, editor, *Trends in Computational Social Choice*, chapter 1, pages 3–26. AI Access, 2017.
- [Chatterji and Zeng, 2018] S. Chatterji and H. Zeng. On random social choice functions with the tops-only property. *Games and Economic Behavior*, 109:413–435, 2018.
- [Chevaileyre *et al.*, 2007] Y. Chevaileyre, U. Endriss, J. Lang, and N. Maudet. A short introduction to computational social choice. In *Proceedings of the 33rd Conference on Current Trends in Theory and Practice of Computer Science (SOFSEM)*, volume 4362 of *Lecture Notes in Computer Science (LNCS)*, pages 51–69. Springer-Verlag, 2007.
- [Duggan and Schwartz, 2000] J. Duggan and T. Schwartz. Strategic manipulability without resoluteness or shared beliefs: Gibbard-Satterthwaite generalized. *Social Choice and Welfare*, 17(1):85–93, 2000.
- [Dutta *et al.*, 2007] B. Dutta, H. Peters, and A. Sen. Strategy-proof cardinal decision schemes. *Social Choice and Welfare*, 28(1):163–179, 2007.
- [Ehlers *et al.*, 2002] L. Ehlers, H. Peters, and T. Storcken. Strategy-proof probabilistic decision schemes for one-dimensional single-peaked preferences. *Journal of Economic Theory*, 105(2):408–434, 2002.
- [Elkind *et al.*, 2017] E. Elkind, M. Lackner, and D. Peters. Structured preferences. In U. Endriss, editor, *Trends in Computational Social Choice*, chapter 10. 2017.
- [Endriss, 2017] U. Endriss, editor. *Trends in Computational Social Choice*. AI Access, 2017.
- [Gibbard, 1973] A. Gibbard. Manipulation of voting schemes: A general result. *Econometrica*, 41(4):587–601, 1973.
- [Gibbard, 1977] A. Gibbard. Manipulation of schemes that mix voting with chance. *Econometrica*, 45(3):665–681, 1977.
- [Gross *et al.*, 2017] S. Gross, E. Anshelevich, and L. Xia. Vote until two of you agree: Mechanisms with small distortion and sample complexity. In *Proceedings of the 31st AAAI Conference on Artificial Intelligence (AAAI)*, pages 544–550, 2017.
- [Hoang, 2017] L. N. Hoang. Strategy-proofness of the randomized Condorcet voting system. *Social Choice and Welfare*, 48(3):679–701, 2017.
- [Hylland, 1980] A. Hylland. Strategyproofness of voting procedures with lotteries as outcomes and infinite sets of strategies. Mimeo, 1980.
- [Mennle and Seuken, 2021] T. Mennle and S. Seuken. Partial strategyproofness: Relaxing strategyproofness for the random assignment problem. *Journal of Economic Theory*, 191:105–144, 2021.
- [Nandeibam, 2013] S. Nandeibam. The structure of decision schemes with cardinal preferences. *Review of Economic Design*, 17(3):205–238, 2013.
- [Procaccia and Rosenschein, 2006] A. D. Procaccia and J. S. Rosenschein. The distortion of cardinal preferences in voting. In *Proceedings of 10th International Workshop on Cooperative Information Agents*, pages 317–331. Springer, 2006.
- [Satterthwaite, 1975] M. A. Satterthwaite. Strategy-proofness and Arrow’s conditions: Existence and correspondence theorems for voting procedures and social welfare functions. *Journal of Economic Theory*, 10(2):187–217, 1975.
- [Sen, 2011] A. Sen. The Gibbard random dictatorship theorem: a generalization and a new proof. *SERIEs*, 2(4):515–527, 2011.

Appendix: Omitted Proofs

In this appendix, we provide the proofs omitted in the main body. Note that we use additional notation here for presenting preference profiles. In particular, we use the $*$ -symbol to represent all missing alternatives. For instance, the preference relation $a, *, b$ means that a is preferred to every other alternative, b is the least preferred alternative, and the remaining alternatives can be ordered arbitrarily.

We start by proving the propositions in Section 4.

Proposition 1. *If a neutral SDS is U -strategyproof for a set $U \neq \emptyset$, it is U' -strategyproof for a symmetric set U' with $U \subseteq U'$.*

Proof. Let f denote a neutral SDS that is U -strategyproof for a non-empty set U and let $U' = \{u \circ \pi : u \in U, \pi \in \Pi\}$ denote the smallest symmetric set that contains U . We suppose in the sequel that $U \neq U'$ as otherwise, there is nothing to show. Moreover, assume for contradiction that f is not U' -strategyproof. This means that there are two preference profiles R and R' , a voter i , a utility function $u \in U$, and a permutation $\pi : A \rightarrow A$ such that $R_j = R'_j$ for all $j \in N \setminus \{i\}$, $u^\pi = u \circ \pi$ is consistent with R_i , and $\mathbb{E}[f(R')]_{u^\pi} > \mathbb{E}[f(R)]_{u^\pi}$. Note that $u^\pi \notin U$ as otherwise, this assumption is in direct conflict with the U -strategyproofness of f .

Next, let $\bar{R} = \pi(R)$, i.e., $x\bar{R}_jy$ if and only if $\pi(x)R_j\pi(y)$ for all $x, y \in A$ and $j \in N$, and $\bar{R}' = \pi(R')$, i.e., $x\bar{R}'_jy$ if and only if $\pi(x)R'_j\pi(y)$ for all $x, y \in A$ and $j \in N$, denote the profiles derived by permuting R and R' with π . Moreover, let π^{-1} denote the inverse permutation of π , i.e., $\pi^{-1}(\pi(x)) = x$ for all $x \in A$. Note that u is consistent with $\bar{R}_i$ because $x\bar{R}_iy \iff \pi^{-1}(x)R_i\pi^{-1}(y) \iff u^\pi(\pi^{-1}(x)) \geq u^\pi(\pi^{-1}(y)) \iff u(x) \geq u(y)$ for all $x, y \in A$. Furthermore, it follows from neutrality that $f(\bar{R}, \pi(x)) = f(R, x)$ and $f(\bar{R}', \pi(x)) = f(R', x)$ for all $x \in A$. Hence, we can calculate that

$$\begin{aligned} \mathbb{E}[f(\bar{R})]_u &= \sum_{x \in A} f(\bar{R}, x)u(x) \\ &= \sum_{x \in A} f(\bar{R}, \pi(x))u(\pi(x)) = \sum_{x \in A} f(R, x)u^\pi(x) \\ &< \sum_{x \in A} f(R', x)u^\pi(x) = \sum_{x \in A} f(\bar{R}', \pi(x))u(\pi(x)) \\ &= \sum_{x \in A} f(\bar{R}', x)u(x) = \mathbb{E}[f(\bar{R}')]_{u^\pi}. \end{aligned}$$

However, this contradicts that f is U -strategyproof as there is a utility function $u \in U$ with which a voter can manipulate. Hence, the assumption that f violates U' -strategyproofness is wrong. $\square$

Proposition 2. *For every SDS f and preference relation R_i , the set $U_{R_i} = \{u \in U : u \text{ is consistent with } R_i \text{ and } f \text{ is } \{u\}\text{-strategyproof}\}$ is convex.*

Proof. Let f denote an SDS and consider an arbitrary preference relation R_i . We need to show that the set $U_{R_i} = \{u \in U : u \text{ is consistent with } R_i \text{ and } f \text{ is } \{u\}\text{-strategyproof}\}$

is convex. First, note that if $|U_{R_i}| \leq 1$, the set is trivially convex. Hence, assume that $|U_{R_i}| \geq 2$ and consider two arbitrary utility functions $u, u' \in U_{R_i}$ with $u \neq u'$. We need to show that $u'' = \lambda u + (1 - \lambda)u' \in U_{R_i}$ for every $\lambda \in (0, 1)$. First note that u'' is consistent with R_i as xR_iy entails for all alternatives $x, y \in A$ that $u(x) \geq u(y)$ and $u'(x) \geq u'(y)$. As a consequence, $u''(x) = \lambda u(x) + (1 - \lambda)u'(x) \geq \lambda u(y) + (1 - \lambda)u'(y) = u''(y)$.

Next, we need to show that f is also $\{u''\}$ -strategyproof. Assume for contradiction that this is not the case. Then, there are two preference profiles R and R' and a voter $i \in N$ such that $R_j = R'_j$ for all $j \in N \setminus \{i\}$, R_i is consistent with u'' , and $\mathbb{E}[f(R')]_{u''} > \mathbb{E}[f(R)]_{u''}$. By the definition of u'' , this means that $\lambda \mathbb{E}[f(R')]_u + (1 - \lambda) \mathbb{E}[f(R')]_{u'} > \lambda \mathbb{E}[f(R)]_u + (1 - \lambda) \mathbb{E}[f(R)]_{u'}$. This inequality is only true if $\mathbb{E}[f(R')]_u > \mathbb{E}[f(R)]_u$ or $\mathbb{E}[f(R')]_{u'} > \mathbb{E}[f(R)]_{u'}$. However, as u and u' are also both consistent with R_i , this implies that f violates either $\{u\}$ -strategyproofness or $\{u'\}$ -strategyproofness. Hence, $\{u, u'\} \not\subseteq U_{R_i}$ which contradicts our initial assumption, and thus, f must be $\{u''\}$ -strategyproof. This means that $u'' \in U_{R_i}$ and that U_{R_i} is indeed convex. $\square$

Proposition 3. *No SDS f satisfies both u^Π -strategyproofness and 1-unanimity if $u(1) - u(2) < \frac{\epsilon_f}{1 - \epsilon_f}(u(2) - u(3))$, $m \geq 3$, and $n \geq 3$.*

Proof. As explained in the main body, this proposition is a variant of Benoît's impossibility [Benoît, 2002]. Formally, this impossibility theorem states that 1-unanimity is incompatible with strategyproofness for set-valued voting rules if voters have the following preferences over sets of alternatives (a_i denotes the i -th best alternative of a voter):

1. The singleton set $\{a_1\}$ is preferred to every other outcome.
2. The set $\{a_1, a_2\}$ is preferred to every other outcome but $\{a_1\}$.
3. The singleton set $\{a_2\}$ is preferred to every other outcome but $\{a_1\}$ and $\{a_1, a_2\}$.
4. Every other set is preferred to the singleton set $\{a_m\}$.

The central idea for proving our proposition is to compare lotteries only with respect to their support $\text{supp}(p) = \{x \in A : p(x) > 0\}$. In particular, we want to find utility functions such that the voters' preferences on lotteries agree with Benoît's preferences over sets of alternatives if we only consider the support of the lotteries. Then, the impossibility of Benoît also applies in our randomized setting. Note therefore that the first and fourth condition on the preferences over sets are trivially true in our model as we assume that $u(1) > \dots > u(m)$. Moreover, the third condition entails the second one because $u(2) < \lambda u(1) + (1 - \lambda)u(2) < u(1)$ for every $\lambda \in (0, 1)$. Hence, we only need to ensure that voters strictly prefer the lottery that chooses their second best alternative with probability 1 to every lottery that assigns a positive probability to a worse alternative.

For this, define $\epsilon_f = \min_{x \in A, R \in \mathcal{R}^n : f(R, x) > 0} f(R, x)$ as the minimum non-zero probability that an SDS f assigns to

an alternative. Note that ϵ_f is well-defined since SDSs are defined for fixed numbers of alternatives and voters. It follows from this definition that $\mathbb{E}[f(R)]_u \leq (1 - \epsilon_f)u(1) + \epsilon_f u(3)$ for all SDSs f , voters i with utility function u , and all preference profiles R such that $\text{supp}(f(R))$ is not a subset of voter i 's best two alternatives. This means that voter i prefers the lottery that assigns probability 1 to his second best alternative to every outcome of f that assigns probability to a worse alternative if $u(2) > (1 - \epsilon_f)u(1) + \epsilon_f u(3)$. This inequality is equivalent to $u(1) - u(2) < \frac{\epsilon_f}{1 - \epsilon_f}(u(2) - u(3))$.

If the utility function of every voter satisfies this condition, the voters' preferences over lotteries are consistent with the preferences over sets of alternatives required by Benoit's impossibility if we only consider the supports of the lotteries. Therefore, Benoit's impossibility applies also in the randomized setting and shows that 1-unanimity is incompatible with u^Π -strategyproofness if $u(1) - u(2) < \frac{\epsilon_f}{1 - \epsilon_f}(u(2) - u(3))$. $\square$

Finally, we prove the claim of the main body that U -strategyproofness also disincentivizes manipulations of groups of voters with the same preferences. For formalizing this observation, we introduce U -group-manipulability and U -group-strategyproofness. We say that an SDS f is U -group-manipulable if there is a subset of the voters $I \subseteq N$, two preference profiles R, R' , and a utility function $u_i \in U$ for every voter $i \in I$ such that $R_j = R'_j$ for all $j \in N \setminus I$, $R_i = R_j$ for all $i, j \in I$, u_i is consistent with R_i for every voter $i \in I$, $\mathbb{E}[f(R')]_{u_i} \geq \mathbb{E}[f(R)]_{u_i}$ for all $i \in I$, and the last inequality is strict for at least one voter $i^* \in I$. Less formally, this means that a group of voters with the same preferences can deviate such that each voter is weakly better off and at least one voter strictly increases his expected utility. Inversely, we call an SDS U -group-strategyproof if it cannot be U -group-manipulated by any subset of the voters $I \subseteq N$. Using this terminology, our observation states that U -strategyproofness and U -group-strategyproofness are equivalent.

Proposition 4. *An SDS is U -strategyproof if and only if it is U -group-strategyproof.*

Proof. Let f denote an arbitrary SDS. It follows immediately that f is U -strategyproof if it is U -group-strategyproof because U -group-strategyproofness is also defined for singleton sets of voters. Hence, we focus on the direction from left to right and assume that f is U -strategyproof for some set U . Moreover, assume for contradiction that f is U -group-manipulable, i.e., that there are two preference profiles R, R' , a set of voters $I \subseteq N$, and a utility function u_i for every voter $i \in I$ such that $R_j = R'_j$ for all $j \in N \setminus I$, $R_i = R_j$ for all $i, j \in I$, u_i is consistent with R_i for all $i \in I$, $\mathbb{E}[f(R')]_{u_i} \geq \mathbb{E}[f(R)]_{u_i}$ for all $i \in I$, and the last inequality is strict for some voter $i^* \in I$. The last assumption means that there is a utility function $u^* \in U$ such that u^* is consistent with R_i and $\mathbb{E}[f(R')]_{u^*} > \mathbb{E}[f(R)]_{u^*}$.

Next, consider the preference profiles $R^0, \dots, R^{|I|}$ such that $R^0 = R$, $R^{|I|} = R'$, and R^{k+1} differs from R^k for all $k \in \{0, \dots, |I| - 1\}$ by replacing the preference relation R_i of a voter in I with his preference relation in R' .

U -strategyproofness entails for each k that $\mathbb{E}[f(R^k)]_{u^*} \geq \mathbb{E}[f(R^{k+1})]_{u^*}$ as these profiles only differ in the preference of a single voter and $u^* \in U$ is consistent with R_i for all $i \in I$. It follows from this observation that $\mathbb{E}[f(R)]_{u^*} \geq \mathbb{E}[f(R')]_{u^*}$ contradicting our assumption that $\mathbb{E}[f(R)]_{u^*} < \mathbb{E}[f(R')]_{u^*}$. This means that the initial assumption is wrong and f is U -group-strategyproof. $\square$

Next, we prove the theorems in Section 5.

Theorem 1. *For every $k \in \{1, \dots, \lfloor \frac{n-1}{2} \rfloor\}$, RD^k satisfies U -strategyproofness for $U = \{u \in \mathcal{U}: u(1) - u(2) \geq k(u(2) - u(m))\}$ and violates $\{u\}$ -strategyproofness for every utility function $u \notin U$. Moreover, $OMNI^*$ satisfies U -strategyproofness for $U = \{u \in \mathcal{U}: u(1) - u(2) \geq \sum_{i=3}^m u(2) - u(i)\}$ and violates $\{u\}$ -strategyproofness for every utility function $u \notin U$.*

Proof. First, we show that RD^k , $k \in \{1, \dots, \frac{n-1}{2}\}$, is U -strategyproof for $U = \{u \in \mathcal{U}: u(1) - u(2) \geq k(u(2) - u(m))\}$. Assume for contradiction that it is not the case, i.e., that there are a utility function u with $u(1) - u(2) \geq k(u(2) - u(m))$, profiles R and R' , and a voter $i \in N$ such that $R_j = R'_j$ for all $j \in N \setminus \{i\}$, u is consistent with R_i , and $\mathbb{E}[RD^k(R')]_u = \sum_{x \in A} u(x)RD^k(R', x) > \sum_{x \in A} u(x)RD^k(R, x) = \mathbb{E}[RD^k(R)]_u$. If neither R nor R' contain $n - k$ voters who agree on a most preferred alternative, RD^k is equal to RD for both profiles. As RD is even SD -strategyproof, it follows that it is also $\{u\}$ -strategyproof and hence, voter i cannot manipulate RD^k in this case. Moreover, RD^k can also not be manipulated if $n - k$ voters agree on a most preferred alternative in R : if voter i is one of those voters he obtains already his maximal utility and if voter i prefers another alternative the most, he cannot affect the outcome.

The only remaining case is that $n - k - 1$ voters agree that an alternative a is the best choice in R , voter i prefers another alternative b the most, and the remaining k voters prefer some other alternatives the most. Then, voter i might try to manipulate by submitting a as his best alternative in R' . We assume in the sequel that the last k voters top-rank voter i 's worst alternative c in R as this minimizes voter i 's expected utility. Based on these insights, we derive the following inequality for voter i 's expected utility in R .

$$\begin{aligned} \mathbb{E}[RD^k(R)]_u &\geq \frac{n-k-1}{n}u(a) + \frac{1}{n}u(b) + \frac{k}{n}u(c) \\ &= \frac{n-k-1}{n}u(a) + \frac{1}{n}u(1) + \frac{k}{n}u(m) \end{aligned}$$

Moreover, RD^k assigns a probability of 1 to a in R' because $n - k$ voters report a as their best alternative. Hence, it follows directly that voter i 's expected utility is $\mathbb{E}[RD^k(R')]_u = u(a)$. Finally, we compare the expected utilities of voter i in R and R' .

$$\begin{aligned} \frac{n-k-1}{n}u(a) + \frac{1}{n}u(1) + \frac{k}{n}u(m) &\geq u(a) \\ \iff \frac{1}{n}(u(1) - u(a)) &\geq \frac{k}{n}(u(a) - u(m)) \end{aligned}$$

The second line is derived by reformulating the equation in the first line. Note that the left side of the simplified inequality is minimized and the right side is maximized if $u(a) = u(2)$. Hence, the assumption that $u(1) - u(2) \geq k(u(2) - u(m))$ entails that $\mathbb{E}[RD^k(R)]_u \geq \mathbb{E}[RD^k(R')]_u$ and no manipulation is possible. Consequently, RD^k is $\{u\}$ -strategyproof for every $u \in U = \{u \in \mathcal{U}: u(1) - u(2) \geq k(u(2) - u(m))\}$ and therefore also U -strategyproof.

Also note that the last inequality as well as and Proposition 1 immediately entail that RD^k violates $\{u\}$ -strategyproofness for every utility function $u \notin U$ as these utility functions satisfy $u(1) - u(2) < k(u(2) - u(m))$. The reason for this is that $\frac{n-k-1}{n}u(2) + \frac{1}{n}u(1) + \frac{k}{n}u(m) < u(2)$ is then true and voter i can manipulate if $n - k - 1$ voters top-rank his second best alternative and the remaining k voters top-rank his worst alternative.

Next, we show that $OMNI^*$ is U -strategyproof for $U = \{u \in \mathcal{U}: u(1) - u(2) \geq \sum_{i=3}^m (u(2) - u(i))\}$. Assume again for contradiction that this is not true, i.e., that there are a utility function u with $u(1) - u(2) \geq \sum_{i=3}^m (u(2) - u(i))$, preference profiles R and R' , and a voter i such that $R_j = R'_j$ for all $j \in N \setminus \{i\}$, u is consistent with R_i , and $\mathbb{E}[OMNI^*(R')]_u > \mathbb{E}[OMNI^*(R)]_u$. We proceed with a case distinction on the outcomes of R and R' . First, assume that $OMNI^*(R, a) = 1$ for some alternative $a \in A$. This means that a majority of the voters reports a as their best alternative and consequently, these voters receive the best possible outcome. Moreover, the remaining voters cannot influence the outcome and hence, $OMNI^*$ is U -strategyproof in this case.

Next, consider the case that the supports of both $OMNI^*(R)$ and $OMNI^*(R')$ consist of at least two alternatives, i.e., $OMNI^*$ returns for both R and R' the uniform lottery over the top-ranked alternatives of the respective profiles. Let $S = \{x \in A: OMNI^*(R, x) > 0\}$ denote the set of alternatives with positive winning chance in R , let a denote the most preferred alternative of voter i in R , and let b denote his most preferred alternative in R' . If voter i is the only voter who top-ranks alternative a in R , he cannot manipulate because alternative a receives no probability anymore if he misreports another alternative as his top choice. As a consequence, either $OMNI^*(R, a) = OMNI^*(R', b)$ and $OMNI^*(R, x) = OMNI^*(R', x)$ for all $x \in A \setminus \{a, b\}$ if b has not been top-ranked in R , or $OMNI^*(R', x) = \frac{|S|}{|S|-1} OMNI^*(R, x)$ for all $x \in A \setminus \{a\}$ otherwise. Both cases are no manipulation as only the probability assigned to a has been redistributed, but voter i assigns the most utility to alternative a .

If another voter top-ranks voter i 's best alternative a , voter i can only change the outcome by top-ranking an alternative that no voter reports as his best one. Hence, the difference between $OMNI^*(R)$ and $OMNI^*(R')$ is that $OMNI^*(R', b) = \frac{1}{|S|+1}$ instead of $OMNI^*(R, b) = 0$ and $OMNI^*(R, x) = \frac{1}{|S|+1}$ instead of $OMNI^*(R, x) = \frac{1}{|S|}$ for all $x \in S$. As $OMNI^*$ returns the uniform lottery on the top-ranked alternatives, a voter's expected utility is his average utility of the top-ranked alternatives. Hence, reporting b as best alternative is only a $\{u\}$ -manipulation for voter i if

$u(b) > \frac{\sum_{x \in S} u(x)}{|S|}$; otherwise, the average utility does not increase. However, this is not possible due to our condition on u , which is equivalent to $u(2) \leq \frac{1}{m} \sum_{k=1}^m u(k)$. This means that $u(2)$ has at most as much utility as the uniform lottery over all alternatives. As a consequence, the average utility of a set X that contains voter i 's best alternative is at least $\frac{1}{m} \sum_{k=1}^m u(k)$ because

$$\frac{1}{m} \sum_{k=1}^m u(k) = \frac{|X|}{m} \sum_{x \in X} \frac{u(x)}{|X|} + \frac{|A \setminus X|}{m} \sum_{x \in A \setminus X} \frac{u(x)}{|A \setminus X|}.$$

Since the last sum only contains alternatives with a utility of at most $u(2)$, it follows that $\sum_{x \in A \setminus X} \frac{u(x)}{|A \setminus X|} \leq u(2) \leq \frac{1}{m} \sum_{i=1}^m u(i)$. This entails that $\sum_{x \in X} \frac{u(x)}{|X|} \geq \frac{1}{m} \sum_{i=1}^m u(i)$ as the above equation cannot be true otherwise. Since $OMNI^*(R)$ puts positive probability on voter i 's best alternative, we derive that $\mathbb{E}[OMNI^*(R)]_u \geq \frac{1}{m} \sum_{k=1}^m u(k) \geq u(2)$. Finally, as $u(2) \geq u(b)$, it follows that voter i cannot $\{u\}$ -manipulate in this case.

As last case, assume that $OMNI^*$ randomizes for R over at least two alternatives and for R' only over a single alternative. This is only possible if voter i misreports an alternative b as his best choice in R' and $OMNI^*(R', b) = 1$. Hence, the expected utility of voter i for R' is at most $u(2)$. However, by the same line of argumentation as in the previous paragraph, we derive that $\mathbb{E}[OMNI^*(R)]_u \geq \frac{1}{m} \sum_{i=1}^m u(i) \geq u(2)$. Consequently, voter i cannot manipulate in this case either, which means that $OMNI^*$ is U -strategyproof for $U = \{u \in \mathcal{U}: u(1) - u(2) \geq \sum_{i=3}^m (u(2) - u(i))\}$.

Finally, note that $OMNI^*$ violates $\{u\}$ -strategyproofness for every utility function $u \notin U$, i.e., for every utility function u with $u(1) - u(2) < \sum_{i=3}^m (u(2) - u(i))$. Note for this that the condition on u is equivalent to $u(2) > \frac{\sum_{k=1}^m u(k)}{m}$, and consider a preference profile in which every alternative is top-ranked and voter i 's second ranked alternative is top-ranked by $\lfloor \frac{n}{2} \rfloor$ voters. Note that we can assume without loss of generality that u is consistent with voter i 's preference as $OMNI^*$ is neutral. In this situation, each alternative has a winning chance of $\frac{1}{m}$ and thus, voter i 's expected utility is $\frac{\sum_{k=1}^m u(k)}{m}$. On the other side, voter i can report his second best alternative as best one, which results in the fact that it is chosen with probability 1 as it is now top-ranked by an absolute majority of the voters. As $u(2) > \frac{\sum_{k=1}^m u(k)}{m}$, this is a successful $\{u\}$ -manipulation for voter i . $\square$

Theorem 2. *There is no rank-based SDS that satisfies u^Π -strategyproofness and k -unanimity for $0 < k < \frac{n}{2}$ if $m \geq 3$, $n \geq 3$, and $u(1) - u(2) < \sum_{i=\max(3, m-k+1)}^m (u(2) - u(i))$.*

Proof. Consider fixed values of $n \geq 3$, $m \geq 3$, and $0 < k < \frac{n}{2}$, and let $k^* = \min(k, m-2)$. Observe that the last definition entails that $\sum_{i=m-k^*+1}^m u(2) - u(i) = \sum_{i=\max(3, m-k+1)}^m (u(2) - u(i))$ and we use from now on the left hand side of the equation to avoid the maximum. We assume for contradiction that there is a rank-based SDS f for m alternatives and n voters that satisfies k -unanimity and u^Π -strategyproofness for some utility function u with

$u(1) - u(2) < \sum_{i=m-k^*+1}^m u(2) - u(i)$. For deriving a conflict, we proceed in two steps: first, we discuss a general construction that allows to weaken an alternative a that is currently assigned probability 1 from first place to second place without affecting the outcome if sufficiently many voters top-rank a . Secondly, we use this construction repeatedly to derive a profile R^* in which a gets probability 1 even though only k^* voters report it as their best choice. Moreover, we can ensure that the remaining $n - k^* \geq n - k$ voters agree that another alternative b is the best outcome, and thus, $f(R^*, a) = 1$ contradicts k -unanimity because this axiom requires that $f(R^*, b) = 1$.

Step 1: Let $\{x_0, \dots, x_{k^*}\}$ denote a set of $k^* + 1$ alternatives and let $\hat{x}_i = x_{i \bmod k^* + 1}$ to simplify notation. In this step, our goal is to find profiles $R^0, \dots, R^{k^*}$ such that (i) $r^*(R^i) = r^*(R^j)$ for all $i, j \in \{0, \dots, k^*\}$, and (ii) in every profile R^i , there is a voter j^* with preference $\hat{x}_i, a, *, \hat{x}_{i+1}, \hat{x}_{i+2}, \dots, \hat{x}_{i+k^*}$. Given these profiles, we show that $f(R^i, a) = 1$ if $f(\hat{R}^i, a) = 1$ for all $i \in \{0, \dots, k^*\}$, where $\hat{R}^i$ denotes the profile derived from R^i by letting voter j^* swap his best alternative $\hat{x}_i$ with a . For the sake of simplicity, we focus in this step on the case that there are $n = 2k^* + 1$ voters. If there are more voters, we can just pick a suitable subset of $2k^* + 1$ voters and apply our construction to these voters while keeping the preferences of the other voters constant.

Next, we explain how to construct the profiles $R^0, \dots, R^{k^*}$. In the profile R_i , the voters j with $1 \leq j \leq k^* + 1$ and $j \bmod k^* + 1 \neq i$ have the preference $a, \hat{x}_j, *, \hat{x}_{j+1}, \dots, \hat{x}_{j+k^*-1}$. Moreover, voter j^* with $j^* \leq k^* + 1$ and $i = j^* \bmod k^* + 1$ has the preference $\hat{x}_{j^*}, a, *, \hat{x}_{j^*+1}, \dots, \hat{x}_{j^*+k^*}$. Note that the construction of the preference of voter j^* differs from the previous preferences only in the swap of his best two alternatives. Moreover, if we restrict the preferences of these voters to the alternatives in $\{x_0, \dots, x_{k^*}\}$, these voters submit a cyclone. Next, the voters j with $k^* + 2 \leq j \leq 2k^* + 1$ and $j \neq k^* + 1 + i$ have the preference $\hat{x}_j, \hat{x}_0, *, \hat{x}_1, \dots, \hat{x}_{j-1}, \hat{x}_{j+1}, \dots, \hat{x}_{j+k^*}, a$ and voter $j = k^* + 1 + i$ only swaps $\hat{x}_j$ with x_0 , i.e., his preference relation is $\hat{x}_0, \hat{x}_j, *, \hat{x}_1, \dots, \hat{x}_{j-1}, \hat{x}_{j+1}, \dots, \hat{x}_{j+k^*}, a$. It should be mentioned that in R_0 , alternative $x_0 = \hat{x}_0$ is the second ranked by all voters j with $j \geq k^* + 2$ because $j = k^* + 1$ is always wrong in this case. Finally, we assume for simplicity that all voters have the same preferences on the alternatives in $Y = A \setminus \{a, x_0, \dots, x_{k^*}\}$ (these alternatives were abbreviated by the $*$ -symbol in all previous preferences).

Note that all profiles R^i have the same rank matrix because $r^*(R^i, x) = r^*(R^j, x)$ for all $x \in A$ and $i, j \in \{0, \dots, k^*\}$. For the alternatives in Y , this claim holds since the preferences involving these alternatives are always the same during the construction. For the alternatives in $A \setminus Y$, this follows because because the profile R^0 differs from every other profile R^i only in the preferences of voters $i, k^* + 1$, and $k^* + i + 1$. Moreover, the preferences of these voters also only differ in swaps between a, x_0 , and x_i . In more detail, R_i^i is derived from R_i^0 by reinforcing x_i against a , $R_{k^*+1}^i$ is derived from

$R_{k^*+1}^0$ by reinforcing a against x_0 , and $R_{k^*+i+1}^i$ is derived from $R_{k^*+i+1}^0$ by reinforcing x_0 against x_i . As all these swaps happen between first and second ranked alternatives, the rank vectors of the alternatives in $A \setminus Y$ are equal in the profiles R^0 and R^i . Thus, it holds that $r^*(R^0) = r^*(R^i)$ for all profiles R^i , and therefore also that $r^*(R^i) = r^*(R^j)$ for all $i, j \in \{0, \dots, k^*\}$. Consequently, rank-basedness implies that $f(R^i) = f(R^j)$ for all $i, j \in \{0, \dots, k^*\}$.

Finally, it remains to show that $f(R^i, a) = 1$ for all $i \in \{0, \dots, k^*\}$. We suppose therefore that $f(\hat{R}^i, a) = 1$ for all $i \in \{0, \dots, k^*\}$, where $\hat{R}^i$ denotes the profiles derived from R^i by reinforcing a against $\hat{x}_i$ in the preference of voter j^* ($j^* = i$ if $i > 0$ or $j^* = k^* + 1$ else). As a consequence, voter j^* can ensure that his expected utility is $u(2)$ if he deviates from R^i to $\hat{R}^i$. Hence, u^Π -strategyproofness entails that the expected utility of voter j^* in R^i must be at least $u(2)$, which means that the following inequality must be true.

$$u(2) \leq f(R^i, \hat{x}_i)u(1) + f(R^i, a)u(2) + \sum_{y \in Y} f(R^i, y)u(y) + \sum_{j=1}^{k^*} f(R^i, \hat{x}_{i+j})u(\hat{x}_{i+j})$$

We reformulate this inequality to highlight the similarity to our condition on the utility function u .

$$f(R^i, \hat{x}_i)(u(1) - u(2)) \geq \sum_{y \in Y} f(R^i, y)(u(2) - u(y)) + \sum_{j=1}^{k^*} f(R^i, \hat{x}_{i+j})(u(2) - u(\hat{x}_{i+j}))$$

Furthermore, note that we get for every profile R^j a symmetric inequality to the one shown above. Also recall that $f(R^i) = f(R^j)$ for all $i, j \in \{0, \dots, k^*\}$ because of rank-basedness and thus, we can substitute $f(R^j, x)$ in all inequalities and for all alternatives $x \in A$ with $f(R^i, x)$. Moreover, for every $i \in \{1, \dots, k^* + 1\}$, alternative $\hat{x}_i$ is top-ranked by the manipulator in R^i , and for every $r \in \{m - k^* + 1, \dots, m\}$, there is a single profile R^j such that the manipulator ranks $\hat{x}_i$ at position r because of the symmetry of the considered profiles. Hence, we derive the following inequality by summing up over all constraints on $f(R^i)$.

$$\sum_{j=0}^{k^*} f(R^i, x_j)(u(1) - u(2)) \geq (k^* + 1) \sum_{y \in Y} f(R^i, y)(u(2) - u(y)) + \sum_{j=0}^{k^*} f(R^i, x_j) \sum_{l=m-k^*+1}^m (u(2) - u(l))$$

However, note that this inequality can only be true if $f(R^i, a) = 1$ because $u(1) - u(2) < \sum_{j=m-k^*+1}^m (u(2) - u(j))$ by assumption and $u(2) \geq u(y)$ for all $y \in Y$ implies

that $\sum_{y \in Y} f(R^i, y)(u(2) - u(y)) \geq 0$. Hence, voter j^* can swap his best and second best alternatives in $\hat{R}^i$ without affecting the outcome.

Step 2: Our next goal is to use the construction in the last step to derive a profile R in which a is top-ranked by only k^* voters but assigned probability 1. We therefore start at the profile $\bar{R}^0$ in which the first $n - k^*$ voters prefer alternative a the most and the remaining voters prefer a uniquely the least. It follows from k -unanimity that $f(\bar{R}^0, a) = 1$ as $k^* \leq k$. Moreover, u^Π -strategyproofness entails that all voters can reorder the alternatives in $A \setminus \{a\}$ arbitrarily without affecting the outcome. If a voter who prefers a the most reorders the remaining alternatives and a does not obtain probability 1 anymore, he can u^Π -manipulate by reverting back, and if a voter who prefers a the least reorders his alternatives and a does not obtain probability 1 anymore, he u^Π -manipulates by applying this modification. In particular, we can pick a subset I of the voters who prefer a the most with $|I| = k^* + 1$ and the k^* voters who prefer a the least and assign them the preferences $\hat{R}^i$ for every $i \in \{0, \dots, k^*\}$ without affecting the outcome. Consequently, we can use the results of the last step and derive a profile $\bar{R}^1$ such that $f(\bar{R}^1, a) = 1$, the first $n - k^* - 1$ voters prefer a the most, voter $n - k^*$'s preference relation is $x_0, a, *, x_1, \dots, x_{k^*}$, and the remaining voters prefer a the least. Moreover, it is easy to see that we can repeat this step as long as at least $k^* + 1$ voters top-rank a as the construction in step 1 is independent of the voters that are not used. Hence, by repeatedly applying this construction, we derive a profile $\bar{R}$ such that $f(\bar{R}, a) = 1$, k^* voters prefer a the most, $n - 2k^*$ voters report $x_0, a, *, x_1, \dots, x_{k^*}$, and k^* voters report a as their least preferred outcome.

Finally, recall that the voters who prefer a the least can reorder the alternatives in $A \setminus \{a\}$ arbitrarily without affecting the outcome. Thus, these voters can also ensure that x_0 is their best alternative without changing the resulting lottery. However, this leads to a profile R^* in which $n - k^*$ voters report x_0 as their best alternative and therefore k -unanimity requires that $f(R^*, x_0) = 1$. This is in conflict with the observation that $f(R^*, a) = f(\bar{R}, a) = 1$ and therefore, we have derived a contradiction. $\square$

Next, we present an example for the constructions in the proof of Theorem 2. Therefore, assume that f is a rank-based SDS for $m = 4$ alternatives and $n = 5$ voters that satisfies 2-unanimity and u^Π -strategyproofness for a utility function u with $u(1) - u(2) < u(2) - u(3) + u(2) - u(4)$. By 2-unanimity, we know that $f(R^1, a) = f(R^2, a) = f(R^3, a) = 1$ for the profiles shown in the sequel.

R^1 :	1: a, x_1, x_2, x_0 4: x_0, x_1, x_2, a	2: a, x_2, x_0, x_1 5: x_2, x_0, x_1, a	3: a, x_0, x_1, x_2
R^2 :	1: a, x_1, x_2, x_0 4: x_1, x_0, x_2, a	2: a, x_2, x_0, x_1 5: x_0, x_2, x_1, a	3: a, x_0, x_1, x_2
R^3 :	1: a, x_1, x_2, x_0 4: x_1, x_0, x_2, a	2: a, x_2, x_0, x_1 5: x_2, x_0, x_1, a	3: a, x_0, x_1, x_2

Next, consider the profiles R^4 , R^5 , and R^6 , which correspond to the profiles R^1 , R^2 , and R^0 discussed in step 1 of the proof, respectively.

R^4 :	1: x_1, a, x_2, x_0 4: x_0, x_1, x_2, a	2: a, x_2, x_0, x_1 5: x_2, x_0, x_1, a	3: a, x_0, x_1, x_2
R^5 :	1: a, x_1, x_2, x_0 4: x_1, x_0, x_2, a	2: x_2, a, x_0, x_1 5: x_0, x_2, x_1, a	3: a, x_0, x_1, x_2
R^6 :	1: a, x_1, x_2, x_0 4: x_1, x_0, x_2, a	2: a, x_2, x_0, x_1 5: x_2, x_0, x_1, a	3: x_0, a, x_1, x_2

Note that R^1 differs only in the preference of voter 1 from R^4 , R^2 differs only in the preference of voter 2 from R^5 , and R^3 only differs in the preference of voter 3 from R^6 . Hence, we derive the following inequalities from u^Π -strategyproofness.

$$u(2) \leq f(R^4, x_1)u(1) + f(R^4, a)u(2) + f(R^4, x_2)u(3) + f(R^4, x_0)u(4)$$

$$u(2) \leq f(R^5, x_2)u(1) + f(R^5, a)u(2) + f(R^5, x_0)u(3) + f(R^5, x_1)u(4)$$

$$u(2) \leq f(R^6, x_0)u(1) + f(R^6, a)u(2) + f(R^6, x_1)u(3) + f(R^6, x_2)u(4)$$

Moreover, all of these profiles have the same rank matrix and thus, $f(R^4) = f(R^5) = f(R^6)$. This means that we can substitute $f(R^5, x)$ and $f(R^6, x)$ with $f(R^4, x)$ for all $x \in A$ in the second and third inequality. We derive the following equation by using this observation and summing up all three inequalities.

$$3u(2) \leq \sum_{x \in \{x_0, x_1, x_2\}} f(R^4, x)(u(1) + u(3) + u(4)) + 3f(R^4, a)u(2)$$

Finally, we can reformulate this expression as shown in the proof of Theorem 2 to derive from our assumption on u and rank-basedness that $f(R^4, a) = f(R^5, a) = f(R^6, a) = 1$.

As last step, observe that voters 4 and 5 can change their preferences without affecting the outcome in R^6 as any other lottery is a manipulation for them. Hence, it holds for the profile $\bar{R}^7$ that $f(\bar{R}^7, a) = 1$ because of u^Π -strategyproofness. However, this is in conflict with 2-unanimity because 3 voters report x_0 as their best alternative.

$\bar{R}^7$:	1: a, x_1, x_2, x_0 4: x_0, x_1, x_2, a	2: a, x_2, x_0, x_1 5: x_0, x_1, x_2, a	3: x_0, a, x_1, x_2
---------------	--	--	-----------------------

Theorem 3. *There is no Condorcet-consistent SDS that satisfies u^Π -strategyproofness regardless of the utility function u if $m \geq 4$, $n \geq 5$ and $n \neq 6, n \neq 8$.*

Proof. Assume for contradiction that there is a Condorcet-consistent SDS f for $m \geq 4$ alternatives and $n \geq 5$ voters (and $n \neq 6$, $n \neq 8$) that satisfies u^Π -strategyproofness for some utility function u . The proof works by a case distinction: first we show that there is no u^Π -strategyproof SDS that is Condorcet-consistent if $m \geq 4$, $n = 3$, and $u(1) - u(2) < u(2) - u(m)$. Next, we show that there is no u^Π -strategyproof SDS that satisfies Condorcet-consistency if $m \geq 4$, $n = 5$, and $u(1) - u(m-1) > u(m-1) - u(m)$. These two cases are exhaustive with respect to the utility functions, i.e., every utility function on at least 4 alternatives either satisfies $u(1) - u(2) < u(2) - u(m)$ or $u(1) - u(m-1) > u(m-1) - u(m)$: the strictness of the utility function entails that $u(1) - u(m-1) > u(m-1) - u(m)$ if $u(1) - u(2) \geq u(2) - u(m)$. Note that both cases are proven for a fixed value of n . Hence, we provide as last step arguments for generalizing the impossibility from a fixed number of voters to larger numbers of voters. Just as in the proof of Theorem 2, we assume in the sequel that all voters have the same preferences on the alternatives that are abbreviated by the $*$ -symbol.

Case 1: $u(1) - u(2) < u(2) - u(m)$

As first case, we assume that f is defined for $n = 3$ voters and satisfies u^Π -strategyproofness for a utility function u with $u(1) - u(2) < u(2) - u(m)$. Consider in this case the following preference profiles and note that b is the Condorcet winner in R^2 , a in R^3 , and c in R^4 . Consequently, Condorcet-consistency entails that $f(R^2, b) = f(R^3, a) = f(R^4, c) = 1$.

R^1 :	1: $a, b, *, c$	2: $c, a, *, b$	3: $b, c, *, a$
R^2 :	1: $b, a, *, c$	2: $c, a, *, b$	3: $b, c, *, a$
R^3 :	1: $a, b, *, c$	2: $a, c, *, b$	3: $b, c, *, a$
R^4 :	1: $a, b, *, c$	2: $c, a, *, b$	3: $c, b, *, a$

Moreover, R^1 differs from R^2 only in the preference of the first voter, from R^3 in the preference of the second voter, and from R^4 in the preference of the third voter. Hence, we can use u^Π -strategyproofness to derive constraints on $f(R^1)$. In particular, we derive the following inequality from u^Π -strategyproofness between R^1 and R^2 .

$$u(2) \leq f(R^1, a)u(1) + f(R^1, b)u(2) + f(R^1, c)u(m) + \sum_{x \in A \setminus \{a, b, c\}} f(R^1, x)u(x)$$

We reformulate this inequality such that it becomes more similar to our assumption on u . Moreover, we derive symmetric conditions from u^Π -strategyproofness between R^1 and R^3 , and between R^1 and R^4 . Hence, we deduce the following three inequalities.

$$\begin{aligned} f(R^1, a)(u(1) - u(2)) &\geq f(R^1, c)(u(2) - u(m)) \\ &\quad + \sum_{x \in A \setminus \{a, b, c\}} f(R^1, x)(u(2) - u(x)) \\ f(R^1, c)(u(1) - u(2)) &\geq f(R^1, b)(u(2) - u(m)) \\ &\quad + \sum_{x \in A \setminus \{a, b, c\}} f(R^1, x)(u(2) - u(x)) \end{aligned}$$

$$\begin{aligned} f(R^1, b)(u(1) - u(2)) &\geq f(R^1, a)(u(2) - u(m)) \\ &\quad + \sum_{x \in A \setminus \{a, b, c\}} f(R^1, x)(u(2) - u(x)) \end{aligned}$$

By summing up these inequalities, we derive the following equation.

$$\begin{aligned} \sum_{x \in \{a, b, c\}} f(R^1, x)(u(1) - u(2)) &\geq \sum_{x \in \{a, b, c\}} f(R^1, x)(u(2) - u(m)) \\ &\quad + 3 \sum_{x \in A \setminus \{a, b, c\}} f(R^1, x)(u(2) - u(x)) \end{aligned}$$

Recall that we assume that $u(1) - u(2) < u(2) - u(m)$, and note that every alternative $x \in A \setminus \{a, b, c\}$ is at most the third best alternative of a voter. Hence, our assumption on u and the above inequality are in conflict. Therefore, no u^Π -strategyproof SDS can satisfy Condorcet-consistency if $n = 3$, $m \geq 4$, and $u(1) - u(2) < u(2) - u(m)$.

Case 2: $u(1) - u(m-1) > u(m-1) - u(m)$

Next, assume that f denotes a Condorcet-consistent and u^Π -strategyproof SDS for $n = 5$ voters and a utility function u with $u(1) - u(m-1) > u(m-1) - u(m)$. Consider for this case the following profiles and note that c is the Condorcet winner in R^2 , b in R^3 , and a in R^4 . Hence, Condorcet-consistency entails that $f(R^2, c) = f(R^3, b) = f(R^4, a) = 1$.

R^1 :	1: $a, *, b, c$	2: $c, *, a, b$	3: $b, *, c, a$
	4: $a, b, c, *$	5: $c, b, a, *$	
R^2 :	1: $a, *, c, b$	2: $c, *, a, b$	3: $b, *, c, a$
	4: $a, b, c, *$	5: $c, b, a, *$	
R^3 :	1: $a, *, b, c$	2: $c, *, b, a$	3: $b, *, c, a$
	4: $a, b, c, *$	5: $c, b, a, *$	
R^4 :	1: $a, *, b, c$	2: $c, *, a, b$	3: $b, *, a, c$
	4: $a, b, c, *$	5: $c, b, a, *$	

Just as in the last case, the profile R^1 differs from the profile R^2 only in the preference of the first voter, from the profile R^3 only in the preference of the second voter, and from the profile R^4 only in the preference of the third voter. Hence, we can use u^Π -strategyproofness to derive conditions on $f(R^1)$. In particular, u^Π -strategyproofness between R^1 and R^2 entails the following inequality. The left hand side of this inequality is voter 1's expected utility in R^2 and the right hand side is his expected utility if he reports his preference dishonestly as R^1 .

$$\begin{aligned} u(m-1) &\geq f(R^1, a)u(1) + f(R^1, c)u(m-1) \\ &\quad + f(R^1, b)u(m) + \sum_{x \in A \setminus \{a, b, c\}} f(R^1, x)u(x) \end{aligned}$$

Next, we reformulate again the inequality such that our assumption on u can be used in the end. Moreover, we derive

symmetric conditions from R^3 and R^4 resulting in the following inequalities.

$$f(R^1, b)(u(m-1) - u(m)) \geq f(R^1, a)(u(1) - u(m-1)) + \sum_{x \in A \setminus \{a, b, c\}} f(R^1, x)(u(x) - u(m-1))$$

$$f(R^1, a)(u(m-1) - u(m)) \geq f(R^1, c)(u(1) - u(m-1)) + \sum_{x \in A \setminus \{a, b, c\}} f(R^1, x)(u(x) - u(m-1))$$

$$f(R^1, c)(u(m-1) - u(m)) \geq f(R^1, b)(u(1) - u(m-1)) + \sum_{x \in A \setminus \{a, b, c\}} f(R^1, x)(u(x) - u(m-1))$$

By summing up the last three inequalities, we derive the following equation.

$$\sum_{x \in \{a, b, c\}} f(R^1, x)(u(m-1) - u(m)) \geq \sum_{x \in \{a, b, c\}} f(R^1, x)(u(1) - u(m-1)) + 3 \sum_{x \in A \setminus \{a, b, c\}} f(R^1, x)(u(x) - u(m-1))$$

Every alternative $x \in A \setminus \{a, b, c\}$ is preferred to at least two other alternatives, and thus, $u(x) - u(m-1) > 0$. As a consequence, this inequality and our assumption that $u(1) - u(m-1) > u(m-1) - u(m)$ cannot be simultaneously true. Thus, no SDS satisfies both Condorcet-consistency and u^Π -strategyproofness if $n = 5$, $m \geq 4$, and $u(1) - u(m-1) > u(m-1) - u(m)$.

Case 3: Generalizing the impossibility

Finally, we explain why the impossibility also applies for larger values of n . For the case that n is odd, this is simple: we can just add pairs of voters with inverse preferences to the construction of the required case. These voters do not affect the Condorcet winner as they cancel each other out with respect to the majority margins and the remaining analysis only depends on u^Π -strategyproofness and therefore only on specific voters. Hence, no Condorcet-consistent SDS can satisfy u^Π -strategyproofness, regardless of the utility function u , if $m \geq 4$, $n \geq 5$, and n is odd.

For even n , we use Proposition 4: u^Π -strategyproofness entails u^Π -group-strategyproofness. This observation means that we can just duplicate each voter in the profiles used to reason about odd n , and the analysis stays intact. Hence, the impossibility also generalizes to even n once $n \geq 10$. $\square$

Theorem 4. *COND is the only Condorcet-consistent SDS that satisfies U-strategyproofness for $U = \{u \in \mathcal{U}: u(1) - u(2) = u(2) - u(3)\}$ if $m = 3$ and n is odd.*

Proof. First note that *COND* is by definition Condorcet-consistent, independently of the numbers of voters or alternatives. Next, we show that it also satisfies *U*-strategyproofness for $U = \{u \in \mathcal{U}: u(1) - u(2) = u(2) - u(3)\}$ if $m = 3$. Assume for contradiction that this is not true, i.e., that there are preference profiles R and R' , a voter i , and a utility function $u \in U$ such that u is consistent with R_i , $R_j = R'_j$ for all $j \in N \setminus \{i\}$, and $\mathbb{E}[\text{COND}(R')]_u > \mathbb{E}[\text{COND}(R)]_u$. We employ a case distinction with respect to the existence of a Condorcet winner. First, assume that there is a Condorcet winner a in R , i.e., $\text{COND}(R, a) = 1$. If another alternative b is the Condorcet winner in R' , voter i prefers a to b because he cannot make b into the Condorcet winner otherwise. As $\text{COND}(R', b) = 1$ in this case, this is no manipulation as $u(a) > u(b)$. Hence, assume that there is no Condorcet winner in R' . Then, we have that $\text{COND}(R', a) = \text{COND}(R', b) = \text{COND}(R', c) = \frac{1}{3}$, which means that voter i 's expected utility is $u(2)$ since $u(1) = 2u(2) - u(3)$. This is only a manipulation if the Condorcet winner a is voter i 's least preferred alternative, i.e., if $u(a) = u(3)$. However, then voter i cannot change that a is the Condorcet winner, and therefore no manipulation is possible in this case. Finally, assume that there is no Condorcet winner in R . Hence, voter i 's expected utility of $\text{COND}(R)$ is again $u(2)$, which means that he can only manipulate by making his best alternative into the Condorcet winner. This is again not possible and consequently, *COND* is *U*-strategyproof for $U = \{u \in \mathcal{U}: u(1) - u(2) = u(2) - u(3)\}$.

Next, we show that no other Condorcet-consistent SDS is *U*-strategyproof for $U = \{u \in \mathcal{U}: u(1) - u(2) = u(2) - u(3)\}$ if $m = 3$ and n is odd. Note that we assume in the sequel that $n \geq 3$ as the claim is trivial if $n = 1$ due to Condorcet-consistency. Assume for contradiction that there is another SDS f that satisfies these axioms and note that f coincides with *COND* on profiles with a Condorcet winner because of Condorcet-consistency. Since n is odd, f differs from *COND* in a profile R with a majority cycle, i.e., the alternatives in R can be relabeled such that $n_{xy}(R) > 0$, $n_{yz}(R) > 0$, and $n_{zx}(R) > 0$.

First, we show that there must be voters with specific preferences in R . In particular, for each of $R_1 = x, y, z$, $R_2 = y, z, x$, and $R_3 = z, x, y$, there is at least one voter who reports the preference relation. Assume for contradiction that this is not true, i.e., we have $n_{xy}(R) > 0$, $n_{yz}(R) > 0$, and $n_{zx}(R) > 0$ and one of the above preferences is not reported by any voter. Let n_1 , n_2 , and n_3 denote variables that correspond to the numbers of voters in R who submit preference R_1 , R_2 , and R_3 , respectively. Moreover, let n_4 , n_5 , and n_6 denote variables with the same meaning for the preferences $R_4 = x, z, y$, $R_5 = y, x, z$, and $R_6 = z, y, x$. Our contradiction assumption entails that one of n_1 , n_2 , or n_3 is zero, and due to symmetry, we suppose without loss of generality that $n_1 = 0$. Since we need the majority cycle in the preference of the voters, we derive the following inequalities.

$$\begin{aligned} n_{xy}(R) &= n_3 + n_4 - n_2 - n_5 - n_6 > 0 \\ n_{yz}(R) &= n_2 + n_5 - n_3 - n_4 - n_6 > 0 \\ n_{zx}(R) &= n_2 + n_3 + n_6 - n_4 - n_5 > 0 \end{aligned}$$

Summing up the first two inequalities results in $n_{xy}(R) +$

$n_{yz}(R) = -2n_6 > 0$, which cannot be true since $n_6 \geq 0$ by definition. Hence, the assumption that $n_1 = 0$ is wrong and by applying symmetric arguments, it follows that $n_1 > 0$, $n_2 > 0$, and $n_3 > 0$.

We use this observation to prove that $f(R, x) = f(R, y) = f(R, z) = \frac{1}{3}$ for all profiles R that induce a majority cycle, i.e., that have $n_{xy}(R) > 0$, $n_{yz}(R) > 0$, and $n_{zx}(R) > 0$. We therefore introduce the cycle weight $c(R) = n_{xy}(R) + n_{yz}(R) + n_{zx}(R) = 3 + 2k$ for some $k \geq 0$ and let R denote a profile with $f(R) \neq COND(R)$ that minimizes $c(R)$. As $f(R) \neq COND(R)$, one of the following inequalities is true: either $f(R, x) > f(R, y)$, $f(R, y) > f(R, z)$, or $f(R, z) > f(R, x)$; otherwise, we have that $f(R, x) \leq f(R, y) \leq f(R, z) \leq f(R, x)$, which implies that all alternatives receive probability $\frac{1}{3}$. We focus in the sequel on the case that $f(R, x) > f(R, y)$ as the remaining cases are symmetric. By our previous observation, there is a voter i with

preference y, z, x in R . Moreover, let $u \in U$ denote an arbitrary utility function that is consistent with R_i . Since u is equi-distant and $f(R, x) > f(R, y)$, it follows that the expected utility of this voter is less than $u(2)$. Next, consider the profile R' in which voter i reports his preference non-truthfully as z, y, x . This manipulation either results in the fact that z is the Condorcet winner, or it results in a profile R' with $c(R') = c(R) - 2$. In both cases, the expected utility of voter i is $u(2)$ because $f(R', z) = 1$ if z is the Condorcet winner and $f(R', x) = f(R', y) = f(R', z) = \frac{1}{3}$ otherwise. The latter observation is true as R minimizes the circle weights among all profiles in which f differs from $COND$ and $c(R') < c(R)$. Hence, voter i can $\{u\}$ -manipulate, contradicting the U -strategyproofness of f . This means that $COND$ is indeed the only Condorcet-consistent SDS that satisfies U -strategyproofness for $U = \{u \in \mathcal{U} : u(1) - u(2) = u(2) - u(3)\}$ if $m = 3$ and n is odd. $\square$