

Approximately Optimal Mechanisms for Strategyproof Facility Location: Minimizing L_p Norm of Costs

Itai Feigenbaum ^{*} Jay Sethuraman [†] Chun Ye [‡]

Abstract

This paper is concerned with the problem of locating a facility on the line in the presence of strategic agents, also located on the line. Each agent incurs a cost equal to her distance to the facility whereas the planner wishes to minimize the L_p norm of the vector of agent costs. The location of each agent is only privately known, and the goal is to design a strategyproof mechanism that approximates the optimal cost well. It is shown that the median mechanism provides a $2^{1-\frac{1}{p}}$ approximation ratio, and that this is the optimal approximation ratio among all deterministic strategyproof mechanisms. For randomized mechanisms, two results are shown: First, for any integer $\infty > p > 2$, no mechanism—from a rather large class of randomized mechanisms—has an approximation ratio better than that of the median mechanism. This is in contrast to the case of $p = 2$ and $p = \infty$ where a randomized mechanism provably helps improve the worst case approximation ratio. Second, for the case of 2 agents, the LRM mechanism, first designed by Procaccia and Tennenholtz for the special case of L_∞ , provides the optimal approximation ratio among all randomized mechanisms.

1 Introduction

We consider the problem of locating a single facility on the real line. This facility serves a set of n agents, each of whom is located somewhere on the line as well. Each agent cares about his distance to the facility, and incurs a disutility (equivalently, cost) that is equal to his distance to access the facility. An agent’s location is assumed to be private information that is known only to him. Agents report their locations to a central planner who decides where to locate the facility based on the reports of the agents. The planner’s objective is to minimize a “social” cost function that depends on the vector of distances that the agents need to travel to access the

^{*}IEOR Department, Columbia University, New York, NY; itai@ieor.columbia.edu

[†]IEOR Department, Columbia University, New York, NY; jay@ieor.columbia.edu. Research supported by NSF grant CMMI-0916453 and CMMI-1201045.

[‡]IEOR Department, Columbia University, New York, NY; cy2214@columbia.edu

facility. It is natural for the planner to consider locating the facility at a point that minimizes her objective function, but in that case the agents may not have an incentive to report their locations truthfully. As an example, consider the case of 2 agents located at x_1 and x_2 respectively, and suppose the location that optimizes the planner’s objective is the mid-point $(x_1 + x_2)/2$. Then, assuming $x_1 < x_2$, agent 1 has an incentive to report a location $x'_1 < x_1$ so that the planner’s decision results in the facility being located closer to his true location. The planner can address this issue by restricting herself to a *strategyproof* mechanism: by this we mean that it should be a (weakly) dominant strategy for each agent to report his location truthfully to the central planner. This, of course, is an attractive property, but it comes at a cost: based on the earlier example, it is clear that the planner cannot hope to optimize her objective. One way to avoid this difficulty is to assume an environment in which agents (and the planner) can make or receive payments; in such a case, the planner selects the location of the facility, and also a payment scheme, which specifies the amount of money an agent pays (or receives) as a function of the reported locations of the agents as well as the location of the facility. This option gives the planner the ability to support the “optimal” solution as the outcome of a strategyproof mechanism by constructing a carefully designed payment scheme in which any potential benefit for a misreporting agent from a change in the location of the facility is offset by an increase in his payment.

There are many settings, however, in which such monetary compensations are either not possible or are undesirable. This motivated Procaccia and Tennenholtz [15] to formulate the notion of *Approximate Mechanism Design without Money*. In this model the planner restricts herself to strategyproof mechanisms, but is willing to settle for one that does not necessarily optimize her objective. Instead, the planner’s goal is to find a mechanism that effectively *approximates* her objective function. This is captured by the standard notion of approximation that is widely used in the CS literature: for a minimization problem, an algorithm is an α -approximation if the solution it finds is guaranteed to have cost at most α times that of the optimal cost ($\alpha \geq 1$).

Procaccia and Tennenholtz [15] apply the notion of approximate mechanism design without money to the facility location problem considered here for two different objectives: (i) *minisum*, where the goal is to minimize the sum of the costs of the agents; and (ii) *minimax*, where the goal is to minimize the maximum agent cost. They show that for the minimax objective choosing any k -th median—picking the k th largest reported location—is a strategyproof, 2-approximate mechanism. They design a randomized mechanism called LRM (Left-Right-Middle) and show that it is a strategyproof, $3/2$ -approximate mechanism; furthermore, they show that those mechanisms provide the optimal worst-case approximation ratio possible (among all deterministic and randomized strategyproof mechanisms, respectively). For the *minisum* objective, it is known that choosing the median reported location is optimal and strategyproof [14]. Feldman and Wilf [9] consider the same facility location problem on a line but with the social cost function being the

L_2 norm of the agents' costs (Feldman and Wilf actually used the sum of squares of the agents' costs, however most of their results can be easily converted to the L_2 norm. Of course, the approximation ratios they report need to be adjusted as well). They show that the median is a $\sqrt{2}$ -approximate strategyproof mechanism for this objective function, and provide a randomized $(1 + \sqrt{2})/2$ -approximate strategyproof mechanism. Feldman and Wilf also generalize the median mechanism to maintain strategyproofness and a $\sqrt{2}$ approximation ratio on trees; furthermore, they provide a family of randomized strategyproof mechanisms for trees, and in particular show that a member of this family reduces the approximation ratio to strictly below $\sqrt{2}$. A general survey of approximate mechanism design without money for facility location problems has been written by Cheng et al. [7].

Aside from the recent literature on approximate mechanism design, our work is loosely related to other different strands in the literature with a much longer history. First is the classical work on social choice, which deals with the aggregating the preferences of a set of voters over a set of alternatives [13]. The location problem we consider is a special case in which the alternatives are all possible points on the real line (the location of the facility), and agents have single-peaked preferences. An important difference, however, is the following: a typical social choice problem is to find an aggregation rule satisfying a desired set of properties, whereas in our case the planner wishes to optimize or approximate a given social objective function. Nevertheless, various techniques and results from this literature are useful in our setting as well. An important result along these lines is Moulin's characterization of strategyproof mechanisms on the line [14]. A parallel characterization result was developed by Schummer and Vohra [17] for general graphs. In both these papers, much like in our paper, generalized medians play an important role; also, despite not having a specific objective function, these characterizations assume less specific efficiency related properties, such as Pareto efficiency and onto range. Additional papers along these lines are [5, 8]. It is important to note that impossibility results abound in social choice models—our focus on the simple special case enables us to avoid impossibility results such as the Gibbard-Satterthwaite Theorem [10, 16], which implies the non-existence of a *reasonable* social choice function. Second is the classical work in operations research on graphical location problems that considers locating the facility at a Condorcet point [11, 12, 4, 3]. (A Condorcet point is one that is preferred by a majority of agents to any other location.) This literature seeks to establish bounds on the total cost to all the agents to access the facility divided by the minimum cost, with the understanding that smaller ratios are better. However, this literature does not model individual agent incentives, and moreover does not also explore other mechanisms. Finally, there is a rich literature on facility location problems and variations (such as the k -median and k -center problems) where agent incentives are not taken into account. In such problems, there is typically a single objective function (the planner's), and agent locations are known. In this literature, one re-

sorts to approximation algorithms for a different reason—often, these optimization problems turn out to be computationally intractable, and the focus is on developing computationally efficient heuristics for which a worst-case approximation guarantee can be proved (see [19], and chapters 25-26 of [18]). To our knowledge, most of the algorithms designed in this literature violate our (rather strong) strategyproofness requirements. In addition, some consideration has been given in literature to the circle topology, by Alon et al. [1, 2]. It is important to note that while the idea of using approximate mechanisms to induce strategyproofness was first proposed in 2009, the problem of finding strategyproof mechanisms has received attention beforehand. The papers that were written before 2009 allow much more generality in the preferences of the agents, but as a result do not have a specific objective function to optimize, and thus approximation is not of relevance there.

In our paper, we follow the suggestion of Feldman and Wilf [9] and study the problem of locating a single facility on a line, but with the objective function being the L_p norm of the vector of agent-costs (for general $p \geq 1$). In the context of real world facility location problems, where the agents must drive to and from the facility, the L_p norm can represent situations where travel time or other cost increases superlinearly with the distance (as suggested in [6]). For example, when driving over larger distances, there is an increased likelihood (depending on traffic) of the need to stop and refuel, or, in the case of electric cars, stop and recharge—which is even more costly since such recharging can be done at home, without wasting the driver’s time. As another example, certain hybrid cars increase their fuel consumption in longer drives— which is relevant if the cost represents fuel consumption rather than travel time. For such problems, our results provide strong lower bounds, robust to the topology of the road network (since they only require a line) and the value of p . We also hope that our results regarding the median will guide the construction of good mechanisms for more general topologies, similarly to the case of $p = 2$ in [9], where the optimality of the median on the line inspires the construction of a mechanism for tree networks using the appropriate adaptation of the median. Another use of the L_p norm is to strike a balance between efficiency and fairness. The cases of $p = 1$ and $p = \infty$, which were both studied in [15], can be viewed as representing the two extremes on the spectrum between maximizing efficiency (minimizing the total social cost) and maximizing fairness (minimizing the cost of the agent who is worst off). Thus, our definition of social cost allows for a controlled tradeoff between efficiency and fairness by varying the value of p . On the line, this interpretation of the L_p norm becomes particularly interesting in the context of voting. Public opinion on many issues is considered to be on a spectrum between political left and right, lending itself naturally to a one dimensional description. One of the common problems in democratic societies is to balance between majority rule and respecting minority rights; thus, the L_p measure allows for a quantitative exploration of this balance. Of course, this interpretation of the L_p norm can be

relevant to physical facility location problems as well.

We define the problem formally in section 2. In section 3, we show that the median mechanism (which is strategyproof) provides a $2^{1-\frac{1}{p}}$ approximation ratio, and that this is the optimal approximation ratio among all deterministic strategyproof mechanisms. We move onto randomized mechanisms in section 4. First, we present a negative result: we show that for integer $\infty > p > 2$, *no* mechanism—from a rather large class of randomized mechanisms—has an approximation ratio better than that of the median mechanism, as the number of agents goes to infinity. It is worth noting that all the mechanisms proposed in literature so far—for minimax, minisum, and the L_2 social cost functions—belong to this class of mechanisms. Next, we consider the case of 2 agents, and show that the LRM mechanism provides the optimal approximation ratio among all randomized strategyproof mechanisms (that satisfy certain mild assumptions) for this special case, for every $p \geq 1$. Our result for the special case of 2 agents also gives a lower bound on the approximation ratio for all randomized mechanisms. We briefly discuss some directions for further research in section 5. In the appendix we discuss some technical details omitted from the paper, as well as an additional negative result for an alternative definition of the agents' cost.

2 Model

Let $N = \{1, 2, \dots, n\}$, $n \geq 2$, be the set of agents. Each agent $i \in N$ reports a location $x_i \in \mathbb{R}$. A *deterministic* mechanism is a collection of functions $f = \{f_n \mid n \in \mathbb{N}, n \geq 2\}$ such that each $f_n : \mathbb{R}^n \rightarrow \mathbb{R}$ maps each location profile $\mathbf{x} = (x_1, x_2, \dots, x_n)$ to the location of a facility. We will abuse notation and let $f(\mathbf{x})$ denote $f_n(\mathbf{x})$. Under a similar notational abuse, a *randomized* mechanism is a collection of functions f that maps each location profile to a probability distribution over $\mathbb{R}$: if $f(x_1, x_2, \dots, x_n)$ is the distribution π , then the facility is located by drawing a single sample from π .

Our focus will be on deterministic and randomized mechanisms for the problem of locating a single facility when the location of any agent is *private* information to that agent and cannot be observed or otherwise verified. It is therefore critical that the mechanism be *strategyproof*—it should be optimal for each agent i to report his *true* location x_i rather than something else. To that end we assume that if the facility is located at y , an agent's disutility, equivalently cost, is simply his distance to y . Thus, an agent whose true location is x_i incurs a cost $C(x_i, y) = |x_i - y|$. If the location of the facility is random and according to a distribution π , then the cost of agent i is simply $C(x_i, \pi) = \mathbb{E}_{y \sim \pi} |x_i - y|$, where y is a random variable with distribution π . The formal definition of strategyproofness is now:¹

Definition 1. A mechanism f is strategyproof if for each $i \in N$, each $x_i, x'_i \in \mathbb{R}$, and for each

¹Note that for randomized mechanisms, we require strategyproofness in expectation, rather than ex-post.

$$\mathbf{x}_{-i} = (x_1, x_2, \dots, x_{i-1}, x_{i+1}, \dots, x_n) \in \mathbb{R}^{n-1},$$

$$C(x_i, f(x_i, \mathbf{x}_{-i})) \leq C(x_i, f(x'_i, \mathbf{x}_{-i})),$$

where $(\alpha, \mathbf{x}_{-i})$ denotes a vector with the i -th component being α and the j -th component being x_j for all $j \neq i$.

The class of strategyproof mechanisms is quite large: for example, locating the facility at agent 1's reported location is strategyproof, but is not particularly appealing because it fails almost every reasonable notion of fairness and could also be highly "inefficient". To address these issues, and to winnow down the class of acceptable mechanisms, we impose additional requirements that stem from efficiency or fairness considerations. In this paper we assume that locating a facility at y when the location profile is $\mathbf{x} = (x_1, x_2, \dots, x_n)$ incurs the *social cost*

$$sc(\mathbf{x}, y) = \left(\sum_{i \in N} |x_i - y|^p \right)^{1/p}, \quad p \geq 1.$$

For a randomized mechanism f that maps $\mathbf{x}$ to a distribution π , we define the social cost to be²

$$sc(\mathbf{x}, \pi) = \mathbb{E}_{y \sim \pi} \left[\left(\sum_{i \in N} |x_i - y|^p \right)^{1/p} \right].$$

For this definition of social cost, our goal now is to find a strategyproof mechanism that does well with respect to minimizing the social cost. A natural mechanism (and this is the approach taken in the classical literature on facility location) is the "optimal" mechanism: each location profile $\mathbf{x} = (x_1, x_2, \dots, x_n)$ is mapped to $OPT(\mathbf{x})$, defined as $OPT(\mathbf{x}) \in \arg \min_{y \in \mathbb{R}} sc(\mathbf{x}, y)$.³ This optimal mechanism is not strategyproof as shown in the following example.

Example. Suppose there are two agents located at the points 0 and 1 respectively on the real line. If they report their locations truthfully, the optimal mechanism will locate the facility at $y = 0.5$, for any $p > 1$. Assuming agent 2 reports $x_2 = 1$, if agent 1 reports $x'_1 = -1$ instead, the facility will be located at 0, which is best for agent 1.

Given that strategyproofness and optimality cannot be achieved simultaneously, it is necessary to find a tradeoff. In this paper we shall restrict ourselves to strategyproof mechanisms that

²For this definition of social cost, an alternative option is to let the agents' costs increase non-linearly with their distance from the facility, in particular $C(x_i, y) = |x_i - y|^p$. In Appendix A we provide an interesting result for this case.

³Strictly speaking, the mechanism is not well defined in cases where the social cost at $\mathbf{x}$ is minimized by multiple locations, but we could pick an exogenous tie-breaking rule to deal with such cases.

approximate the optimal social cost as best as possible. The notion of approximation that we use is standard in computer science: an α -approximation algorithm is one that is guaranteed to have cost no more than α times the optimal social cost. Formally, the approximation ratio of an algorithm A is $\sup_I \{A(I)/OPT(I)\}$, where the supremum is taken over all possible instances I of the problem, and $A(I)$ and $OPT(I)$ are, respectively, the costs incurred by algorithm A and the optimal algorithm on the instance I .⁴ Our goal then is to design strategyproof (deterministic or randomized) mechanisms whose approximation ratio is as close to 1 as possible.

3 The Median Mechanism

For the location profile $\mathbf{x} = (x_1, x_2, \dots, x_n)$, the median mechanism is a deterministic mechanism that locates the facility at the “median” of the reported locations. The median is unique if n is odd, but not when n is even, so we need to be more specific in describing the mechanism. For odd n , say $n = 2k - 1$ for some $k \geq 1$, the facility is located at $x_{[k]}$, where $x_{[k]}$ is the k th largest component of the location profile. For even n , say $n = 2k$, the “median” can be any point in the interval $[x_{[k]}, x_{[k+1]}]$; to ensure strategyproofness, we need to pick either $x_{[k]}$ or $x_{[k+1]}$, and as a matter of convention we take the median to be $x_{[k]}$. It is well known that the median mechanism is strategyproof.⁵ Furthermore, the median mechanism is *anonymous*.⁶ Thus we may assume, without loss of generality, that each agent reports her location truthfully.

Our main result in this section is that, for any $p \geq 1$, the median mechanism uniformly achieves the best possible approximation ratio among all deterministic strategyproof mechanisms. We start with two simple observations, which will be used in the proof of this main result.

Lemma 1. *For any real numbers a, b, c with $a \leq b \leq c$, and any $p \geq 1$,*

$$(c - a)^p \leq 2^{p-1}[(c - b)^p + (b - a)^p].$$

Proof. For any $p \geq 1$, $f(x) = x^p$ is a convex function on $[0, \infty)$, and so for any $\lambda \in [0, 1]$ and $x, y \geq 0$,

$$f(\lambda x + (1 - \lambda)y) \leq \lambda f(x) + (1 - \lambda)f(y). \quad (1)$$

Setting $\lambda = 1/2$, $x = c - b$, and $y = b - a$, we get:

$$\frac{1}{2^p}(c - a)^p \leq \frac{1}{2}[(c - b)^p + (b - a)^p]. \quad (2)$$

⁴For the case of randomized mechanisms, it should be noted that this is the approximation ratio in expectation rather than with high probability.

⁵A classical paper of Moulin [14] for a closely related model shows that all deterministic strategyproof mechanisms are essentially generalized median mechanisms.

⁶In an anonymous mechanism, the facility location is the same for two location profiles that are permutations of each other.

Multiplying both sides of the inequality by 2^p gives the result. $\square$

Lemma 2. *For any non-negative real numbers a and b , and any $p \geq 1$,*

$$(a + b)^p \geq a^p + b^p.$$

Proof. For integer p , the result is a direct consequence of the binomial theorem; the same argument covers the case of rational p as well. Continuity implies the result for all p . $\square$

Theorem 1. *Suppose there are n agents with the location profile $\mathbf{x} = (x_1, x_2, \dots, x_n)$. Define the social cost of locating a facility at y as $(\sum_{i=1}^n |y - x_i|^p)^{\frac{1}{p}}$ for $p \geq 1$. The social cost incurred by the median mechanism is at most $2^{1-\frac{1}{p}}$ times the optimal social cost.⁷*

Proof. We may assume that $x_1 \leq \dots \leq x_n$. Let OPT be a facility location that minimizes the social cost, and let m be the median. The inequality we need to prove is

$$\sum_{i=1}^n |m - x_i|^p \leq 2^{p-1} \sum_{i=1}^n |OPT - x_i|^p.$$

We do this by pairing each location x_i with its “symmetric” location x_{n+1-i} and arguing that the total cost of these two locations in the median mechanism is within the required bound of their total cost in an optimal solution. For even n , this completes the argument; for odd n the only location without such a pair is the median itself, which incurs *zero* cost in the median mechanism, and so the argument is complete. Formally, the result follows if we can show

$$|m - x_i|^p + |x_{n+1-i} - m|^p \leq 2^{p-1}(|OPT - x_i|^p + |OPT - x_{n+1-i}|^p), \quad \forall i \leq \lfloor n/2 \rfloor.$$

We consider two cases, depending on whether OPT is in the interval $[x_i, x_{n+1-i}]$ or not. In each of these cases, OPT may be above the median or below, but the proof remains identical in each subcase, so we give only one.

1. $x_i \leq m \leq OPT \leq x_{n+1-i}$ or $x_i \leq OPT \leq m \leq x_{n+1-i}$. We will prove the first of these subcases; the proof of the second is identical. Applying Lemma 1 by setting $a = m$, $b = OPT$, and $c = x_{n+1-i}$, we get

$$|x_{n+1-i} - m|^p \leq 2^{p-1}(|x_{n+1-i} - OPT|^p + |OPT - m|^p).$$

Thus,

$$\begin{aligned} |m - x_i|^p + |x_{n+1-i} - m|^p &\leq |m - x_i|^p + 2^{p-1}(|x_{n+1-i} - OPT|^p + |OPT - m|^p) \\ &\leq 2^{p-1}(|m - x_i|^p + |x_{n+1-i} - OPT|^p + |OPT - m|^p) \\ &\leq 2^{p-1}(|x_{n+1-i} - OPT|^p + |OPT - x_i|^p), \end{aligned}$$

⁷This is a generalization of the results for $p = 2$ [9], $p = 1$ and $p = \infty$ [15] (when $p = \infty$, the median mechanism provides a 2-approximation).

where the last inequality is obtained by applying Lemma 2 to the terms $|m - x_i|^p$ and $|OPT - m|^p$.

2. $OPT \leq x_i \leq m \leq x_{n+1-i}$ or $x_i \leq m \leq x_{n+1-i} \leq OPT$. Again, we prove only the first subcase. Note that

$$\begin{aligned} |x_{n+1-i} - m|^p + |m - x_i|^p &\leq |x_{n+1-i} - x_i|^p \\ &\leq |OPT - x_{n+1-i}|^p \\ &\leq 2^{p-1}(|OPT - x_i|^p + |OPT - x_{n+1-i}|^p) \end{aligned}$$

where the first inequality follows from Lemma 2. (Note that Lemma 1 is not used in the proof of this case.)

We end this section by showing that *no* deterministic and strategyproof mechanism can give a better approximation to the social cost.

Lemma 3. *Consider the case of two agents and suppose the location profile is (x_1, x_2) with $x_1 < x_2$. For $p \geq 1$, suppose the social cost of locating a facility at y is $(|x_1 - y|^p + |x_2 - y|^p)^{1/p}$. Any deterministic mechanism whose approximation ratio is better than $2^{1-\frac{1}{p}}$ for $p > 1$ must locate the facility at y for some $y \in (x_1, x_2)$.⁸*

Proof. The function $sc(\mathbf{x}, y)$ is strictly convex in y , and its unique minimizer is $y^* = (x_1 + x_2)/2$, with the corresponding value $sc(\mathbf{x}, y^*) = |x_2 - x_1|/2^{1-\frac{1}{p}}$. Moreover $sc(\mathbf{x}, x_1) = sc(\mathbf{x}, x_2) = |x_2 - x_1| = 2^{1-\frac{1}{p}}sc(\mathbf{x}, y^*)$. It follows that for the deterministic mechanism to do strictly better than the stated ratio, the facility cannot be located at the reported locations; locating the facility to the left of x_1 or to the right of x_2 only increases the cost of the mechanism, so the only option left for a mechanism to do better is to locate the facility in the interior, i.e., in (x_1, x_2) . $\square$

Theorem 2. *Any strategyproof deterministic mechanism has an approximation ratio of at least $2^{1-\frac{1}{p}}$ for the L_p social cost function for any $p \geq 1$.⁹*

Proof. Using Lemma 3, we can now argue similarly to the case of $p = \infty$ (theorem 3.2 in [15]).¹⁰ Suppose $p > 1$ (the bound holds trivially for $p = 1$), and suppose a deterministic strategyproof mechanism yields an approximation ratio strictly better than $2^{1-\frac{1}{p}}$ for the L_p social cost. For the two-agent location profile $x_1 = 0, x_2 = 1$, Lemma 3 implies the facility is located at some $y \in (0, 1)$. Now consider the location profile $x_1 = 0, x_2 = y$. Again, by Lemma 3, the mechanism

⁸Ex-post Pareto efficiency (as defined in section 4.2) requires the facility to be located in $[x_1, x_2]$; thus, this property is stronger.

⁹The lower bound of 2 on the approximation ratio holds when $p = \infty$, see Procaccia and Tennenholtz [15].

¹⁰Another argument along this line can be found in the proof of theorem 4.4 in [9].

must locate the facility at $y' \in (0, y)$ to guarantee the improved approximation. But if agent 2 is located at $y < 1$, he can misreport his location as 1, forcing the mechanism to locate the facility at y , his true location; this violates strategyproofness. $\square$

4 Randomized Mechanisms

Recall that when the social cost is measured by the L_2 norm or the L_∞ norm, randomization provably improves the approximation ratio. In the former case, Feldman and Wilf [9] describe an algorithm whose approximation ratio is $(\sqrt{2}+1)/2$; for the latter, Procaccia and Tennenholtz [15] design an algorithm with an approximation ratio of $3/2$. The mechanisms in both cases are simple and somewhat similar, placing non-negative probabilities *only* on the optimal location and generalized medians (defined shortly), where these probabilities are independent of the reported location profile. In this section we show that this is not enough in general; namely, randomizing over generalized medians and the optimal location does not improve the approximation ratio of the median mechanism for *any* integer $p \in (2, \infty)$. For the case of 2 agents we show that the best approximation ratio is given by the LRM mechanism among all strategyproof mechanisms. Extending this analysis even to the case of 3 agents appears to be non-trivial.

4.1 Mixing Dictatorships and Generalized Medians with the Optimal Location

We begin with a definition of generalized medians.

Definition 2. Let $\mathbf{x} \in \mathbb{R}^n$, $S \subseteq N$, and $m \in \{1, \dots, |S|\}$. Let $S = \{s_1, \dots, s_{|S|}\}$, where $x_{s_i} \leq x_{s_{i+1}}$. Then, the m th generalized median of subset S in location profile $\mathbf{x}$ is $x_{[m,S]} = x_{s_m}$.¹¹ If $S = N$, we allow for the shorthand $x_{[m]} = x_{[m,N]}$.

Next, we define the class of mechanisms currently used in literature:

Definition 3. Let f be a mechanism which satisfied the following. For every $n \in \mathbb{N}$, $S \subseteq N$, $m \in \{1, \dots, |S|\}$, there exist non-negative numbers $v_m^{S^n}$, and v_{OPT}^n with $v_{OPT}^n + \sum_{S \subseteq N, m \in \{1, \dots, |S|\}} v_m^{S^n} = 1$, such that for every profile $(x_1, x_2, \dots, x_n)$, f locates the facility at OPT with probability v_{OPT}^n and at $x_{[m,S]}$ with probability $v_m^{S^n}$ (where OPT is the optimal location for the profile $(x_1, x_2, \dots, x_n)$).¹² If f satisfies these properties, we say that f is a Mixed Generalized Medians Optimal (MGMO) mechanism.

We now show that for integer $p > 2$, MGMO mechanisms cannot beat the median.

¹¹That is, $x_{[m,S]}$ is the m th largest location among the locations of the agents in S , allowing for repetition.

¹²When a location appears more than once in OPT and $x_{[m,S]}$ for $S \subseteq N$ and $m \in \{1, \dots, |S|\}$, the probabilities add up.

Theorem 3. *Let f be a strategyproof MGMO mechanism. Then, for any finite integer $p > 2$, the approximation ratio of f is at least $2^{1-\frac{1}{p}}$.*

Proof. Fix $n = 2k$, with $k \in \mathbb{N}$. In all profiles in our proof, the relative order of agents locations remains the same: specifically, $i < j$ implies $x_i \leq x_j$ for all of our profiles $\mathbf{x}$. For every $S \subseteq N$, and every $j \in S$ let $S(j)$ be the number of agents with index weakly smaller than j in S (for example, if $S = \{2, 4, 9\}$, then $S(2) = 1$, $S(4) = 2$, and $S(9) = 3$). On our profiles, the probability that the location of agent $j \in N$ is chosen as a generalized median therefore is $v_j^n = \sum_{S \subseteq N: j \in S} v_{S(j)}^{S^n}$.

For $j = 1, \dots, k$, define the profile $\mathbf{x}^j$ as follows (where a_j is a parameter to be defined shortly): agents 1 through j are located at $-a_j$; agents $j+1$ through k are located at 0; agents $k+1$ through $2k-j+1$ are located at 1; and agents $2k-j+2$ through $2k$ are located at $1+a_j$ (note the *slight* asymmetry in the location of the agents: while k agents are at or below zero, and k agents are at or above 1, there is an additional agent at 1 compared to zero and so one less agent at $1+a_j$ compared to $-a_j$). Now, a_j is chosen to be the smallest positive root of the function $g_j(\alpha) = j\alpha^{p-1} - (k-j+1) - (j-1)(1+\alpha)^{p-1}$; such an a_j must exist by the intermediate value theorem, as $g_j(0) < 0$ and $g_j(\alpha)$ is a continuous function of α with $g_j(\alpha) \rightarrow \infty$ as $\alpha \rightarrow \infty$.

We show that the optimal mechanism locates the facility at zero for the profile $\mathbf{x}^j$, i.e., $OPT = 0$. Note that the social cost for this profile, when locating the facility at $z \in [0, 1]$, is $j(z+a_j)^p + (k-j)z^p + (k-j+1)(1-z)^p + (j-1)(1+a_j-z)^p$, and when $z \in (-a_j, 0)$ the social cost becomes $j(z+a_j)^p + (k-j)(-z)^p + (k-j+1)(1-z)^p + (j-1)(1+a_j-z)^p$. Note that the social cost function is differentiable for $z \in (0, 1)$ and for $z \in (-a_j, 0)$. The left and right derivatives at 0 are both $pja_j^{p-1} - p(k-j+1) - p(j-1)(1+a_j)^{p-1}$, and thus the social cost function is differentiable on $(-a_j, 1)$ with its derivative at $z = 0$ equal to zero (by our choice of a_j). The fact that this is a global minimum now follows from strict convexity of the social cost function $\|\mathbf{x}^j - z(1, \dots, 1)\|_p$ (for all $z \in \mathbb{R}$). Thus, indeed, $OPT = 0$.

We now attempt to bound v_{OPT}^n . For each profile $\mathbf{x}^j$, consider the profile $\mathbf{x}'^j$ that differs only in the location of agent j : namely, $x'^j_j = 0$ instead of $-a_j$. Note that on this profile, $OPT = 0.5$ by symmetry. Strategyproofness implies that a deviation from profile $\mathbf{x}'^j$ to profile $\mathbf{x}^j$ should not be beneficial for agent j , namely $a_j v_j^n - \frac{1}{2} v_{OPT}^n \geq 0$ (where a_j is the increase in agent j 's cost caused by that deviation when the facility is built in his reported location, and $\frac{1}{2}$ is the decrease in his cost caused by that deviation when the facility is located at OPT), which implies $v_j^n \geq \frac{v_{OPT}^n}{2a_j}$. Defining a_j for $j = k+1, \dots, 2k$ in a symmetric fashion, it follows that the same inequality holds for j in that range, and that $a_j = a_{2k-j+1}$. Summing those inequalities up, we get:

$$1 - v_{OPT}^n = \sum_{j=1}^{2k} v_j^n \geq \sum_{j=1}^{2k} \frac{v_{OPT}^n}{2a_j} = 2 \sum_{j=1}^k \frac{v_{OPT}^n}{2a_j} = \sum_{j=1}^k \frac{v_{OPT}^n}{a_j}$$

$$v_{OPT}^n \leq \frac{1}{1 + \sum_{j=1}^k \frac{1}{a_j}}$$

Now, we claim it is enough to show that as $n \rightarrow \infty$ (or equivalently, as $k \rightarrow \infty$), $\sum_{j=1}^k \frac{1}{a_j} \rightarrow \infty$. The inequality then implies that $v_{OPT}^n \rightarrow 0$. Consider the profile which locates k agents at 0 and k agents at 1. The social cost of locating the facility at OPT on this profile is $\sqrt[p]{n}/2$, while the social cost of locating the facility at an agent's location is $\sqrt[p]{n}2^{-\frac{1}{p}}$; thus, the approximation ratio of f on this profile is $\frac{v_{OPT}^n \sqrt[p]{n}/2 + (1 - v_{OPT}^n) \sqrt[p]{n}2^{-\frac{1}{p}}}{\sqrt[p]{n}/2} = 2^{1-\frac{1}{p}} - (2^{1-\frac{1}{p}} - 1)v_{OPT}^n$. Thus, as $n \rightarrow \infty$, the approximation ratio on these profiles approaches $2^{1-\frac{1}{p}}$, completing the proof.

We are left with the task of showing that $\lim_{k \rightarrow \infty} \sum_{j=1}^k \frac{1}{a_j} = \infty$. To do so, we first show that for $j \geq k^{\frac{1}{p-1}} + 1$, $2^{p-1}(j-1) > a_j$. Recall that a_j was defined as the smallest positive root of $g_j(\alpha)$, and that $g_j(0) < 0$. Thus, it is enough to show that for j in the appropriate range, $g_j(2^{p-1}(j-1)) > 0$. For notational convenience, we denote $Q = 2^{p-1}$.

$$\begin{aligned} g_j(Q(j-1)) &= jQ^{p-1}(j-1)^{p-1} - (k-j+1) - (j-1)(1+Q(j-1))^{p-1} \\ &= Q^{p-1}(j-1)^{p-1} - k - (j-1) \sum_{i=1}^{p-2} \binom{p-1}{i} (Q(j-1))^{p-1-i} \\ &\geq Q^{p-1}(j-1)^{p-1} - (j-1)^{p-1} - (j-1) \sum_{i=1}^{p-2} \binom{p-1}{i} (Q(j-1))^{p-1-i} \\ &\geq Q^{p-1}(j-1)^{p-1} - (j-1)^{p-1} - (j-1) \sum_{i=1}^{p-2} \binom{p-1}{i} (Q(j-1))^{p-2} \\ &> Q^{p-1}(j-1)^{p-1} - (j-1) \sum_{i=1}^{p-1} \binom{p-1}{i} (Q(j-1))^{p-2} \\ &= Q^{p-1}(j-1)^{p-1} - (j-1)(Q(j-1))^{p-2} \sum_{i=1}^{p-1} \binom{p-1}{i} \\ &> Q^{p-1}(j-1)^{p-1} - (j-1)(Q(j-1))^{p-2} 2^{p-1} = 0. \end{aligned}$$

Now,

$$\begin{aligned}
\lim_{k \rightarrow \infty} \sum_{j=1}^k \frac{1}{a_j} &> \lim_{k \rightarrow \infty} \sum_{j=\lceil k^{\frac{1}{p-1}} \rceil + 1}^k \frac{1}{2^{p-1}j} \\
&= \frac{1}{2^{p-1}} \lim_{k \rightarrow \infty} \sum_{j=\lceil k^{\frac{1}{p-1}} \rceil + 1}^k \frac{1}{j} \\
&\geq \frac{1}{2^{p-1}} \lim_{k \rightarrow \infty} \int_{\lceil k^{\frac{1}{p-1}} \rceil + 2}^k \frac{1}{t} dt \\
&= \frac{1}{2^{p-1}} \left(\lim_{k \rightarrow \infty} \int_{\lceil k^{\frac{1}{p-1}} \rceil}^k \frac{1}{t} dt - \lim_{k \rightarrow \infty} \int_{\lceil k^{\frac{1}{p-1}} \rceil}^{\lceil k^{\frac{1}{p-1}} \rceil + 2} \frac{1}{t} dt \right) \\
&= \frac{1}{2^{p-1}} \left(\left(\lim_{k \rightarrow \infty} \left(1 - \frac{1}{p-1} \right) \ln k \right) - 0 \right) = \infty
\end{aligned}$$

which completes our proof. $\square$

4.2 Optimality of the LRM Mechanism for 2 Agents

Procaccia and Tennenholtz [15] defined the mechanism Left-Right-Middle (LRM) as follows: place the facility with probability $\frac{1}{2}$ at OPT , and with probability $\frac{1}{4}$ at each of $x_{[1]}$ and $x_{[n]}$. They have shown that it is strategyproof, and that it provides a best-possible approximation ratio of $\frac{3}{2}$ when $p = \infty$. Our next result shows that the LRM mechanism provides the best possible approximation ratio among all shift and scale invariant (defined below) strategyproof mechanisms for the case of 2 agents for all L_p social cost functions for $p \geq 1$.

We begin with some definitions: we say that a mechanism f is *shift and scale invariant* if for every location profile $\mathbf{x} = (x_1, x_2)$ and every $c \in \mathbb{R}$, the following two properties are satisfied:¹³

1. Shift Invariance: the random variables $Y' \sim f(x_1 + c, x_2 + c)$ and $Y + c$ s.t. $Y \sim f(\mathbf{x})$ are equal in distribution.
2. Scale Invariance: the random variables $Y' \sim f(cx_1, cx_2)$ and cY s.t. $Y \sim f(\mathbf{x})$ are equal in distribution.¹⁴

¹³While these two properties are natural and reasonable to expect, it should be noted that they are not implied by strategyproofness- one example is the constant mechanism, which always locates the facility at the same point regardless of the reports. Requiring unanimity in addition to strategyproofness is also not sufficient to guarantee these properties; for example, the mechanism that runs LRM if $x_{[1]} = 0$, and otherwise locates the facility at $x_{[1]}$ and $x_{[2]}$ with probability $1/2$ each, is easily seen to be strategyproof and unanimous but neither shift nor scale invariant.

¹⁴It is possible to replace shift invariance with symmetry in our assumptions, and preserve our results; see appendix.

A convenient notation for a given location profile $\mathbf{x}$ is to denote its midpoint as $m_{\mathbf{x}} = \frac{x_1+x_2}{2}$. We say that a mechanism f is *symmetric* if for any location profile $\mathbf{x}$ and for any $y \in \mathbb{R}$, $\mathbb{P}(f(\mathbf{x}) \geq m_{\mathbf{x}} + y) = \mathbb{P}(f(\mathbf{x}) \leq m_{\mathbf{x}} - y)$.

The structure of the proof is as follows. Our goal is to show that within the class of strategyproof, shift invariant and scale invariant mechanisms, we can further limit ourselves to symmetric mechanism that locate the facility always at the agents' locations or the midpoint; within this further restricted class, it becomes easy to prove that LRM is optimal. We achieve this goal gradually. First we show that we may restrict ourselves to symmetric (and anonymous) mechanisms. We then provide a characterization of strategyproofness for such mechanisms, and use it to show that we can further restrict ourselves to mechanisms which, for each profile $\mathbf{x}$, do not locate the facility both at $(\min\{x_1, x_2\}, \max\{x_1, x_2\})$ and at $(-\infty, \min\{x_1, x_2\}) \cup (\max\{x_1, x_2\}, \infty)$ with positive probability. We then show that we can restrict ourselves to mechanisms that locate the facility always at the agents' locations or the midpoint.

The following lemma allows us to focus on symmetric mechanisms.

Lemma 4. *Given any strategyproof, shift and scale invariant mechanism, there exists a symmetric, strategyproof, shift and scale invariant mechanism with the same worst-case approximation ratio.*

Proof. Given a mechanism f , we define the *mirror mechanism* of f , f_{mirror} , to be such that for every profile $\mathbf{x}$, we have that $\mathbb{P}(f_{\text{mirror}}(\mathbf{x}) \geq m_{\mathbf{x}} + b) = \mathbb{P}(f(\mathbf{x}) \leq m_{\mathbf{x}} - b)$ for all $b \in \mathbb{R}$.¹⁵

We will need the following notation: For each profile $\mathbf{x} = (x_1, x_2)$, let $Y_{x_1, x_2} \sim f(\mathbf{x})$, and $Y'_{x_1, x_2} \sim f_{\text{mirror}}(\mathbf{x})$. We claim that f_{mirror} is shift invariant, scale invariant and strategyproof (all of the equalities below are in distribution):

1. Shift invariance: let $c \in \mathbb{R}$. Then $Y'_{x_1+c, x_2+c} = 2m_{x_1+c, x_2+c} - Y_{x_1+c, x_2+c} = 2m_{\mathbf{x}} + 2c - Y_{x_1, x_2} - c = Y'_{x_1, x_2} + c$.
2. Scale invariance: let $c \in \mathbb{R}$. Then $Y'_{cx_1, cx_2} = 2cm_{x_1, x_2} - Y_{cx_1, cx_2} = c(2m_{x_1, x_2} - Y_{x_1, x_2}) = cY'_{x_1, x_2}$.
3. Strategyproofness: assume f_{mirror} is not strategyproof, and assume without loss of generality that agent 2 has a profitable misreport: there exist profiles (w_1, w_2) and $(w_1, w_2 + \alpha)$ for some $\alpha \in \mathbb{R}$ such that $\mathbb{E}[|w_2 - Y'_{w_1, w_2}|] > \mathbb{E}[|w_2 - Y'_{w_1, w_2 + \alpha}|]$. However, note that $w_2 - Y'_{w_1, w_2 + \alpha} = -w_1 - \alpha + Y_{w_1, w_2 + \alpha} = Y_{w_1 - \alpha, w_2} - w_1$ (the second equality follows from shift

¹⁵Equivalently, the mirror mechanism can be thought of as follows: whenever f locates the facility at $y \in \mathbb{R}$ (that is, the single sampling of $f(\mathbf{x})$ yields y), f_{mirror} "mirrors" that location about $m_{\mathbf{x}}$, meaning it locates the facility at $2m_{\mathbf{x}} - y$.

invariance), and that $w_2 - Y'_{w_1, w_2} = Y_{w_1, w_2} - w_1$. Thus, it follows that $\mathbb{E}[|w_1 - Y_{w_1, w_2}|] > \mathbb{E}[|Y_{w_1 - \alpha, w_2} - w_1|]$, violating strategyproofness for f . Thus f_{mirror} must be strategyproof.

Therefore, the mechanism g that picks f with probability $1/2$ and f_{mirror} with probability $1/2$ is a strategyproof mechanism that is also symmetric; g trivially satisfies shift and scale invariance. Finally, note that g has the same approximation ratio as f for all location profiles, since f_{mirror} has the same approximation ratio as f . $\square$

Mechanisms which satisfy shift and scale invariance as well as symmetry also satisfy anonymity:

Lemma 5. *If a mechanism f is shift invariant, scale invariant and symmetric, it is also anonymous.*

Proof. Again, all equalities are in distribution. Let $\mathbf{x}$ be a location profile. We need to prove $Y_{x_1, x_2} = Y_{x_2, x_1}$. Shift and scale invariance gives $Y_{x_2, x_1} = -Y_{x_1, x_2} + x_1 + x_2$; thus, $\mathbb{P}(Y_{x_2, x_1} \leq b) = \mathbb{P}(x_1 + x_2 - b \leq Y_{x_1, x_2})$. But $\mathbb{P}(x_1 + x_2 - b \leq Y_{x_1, x_2}) = \mathbb{P}(Y_{x_1, x_2} \leq b)$ by symmetry about $m_{\mathbf{x}}$, thus $Y_{x_2, x_1} = Y_{x_1, x_2}$. $\square$

The next lemma deals with an equivalent condition for strategyproofness for symmetric, shift and scale invariant mechanisms.

Lemma 6. *A symmetric, shift and scale invariant mechanism f is strategyproof if and only if for any profile $\mathbf{x} \in \mathbb{R}^2$ with $x_1 = 0 < x_2$, the following conditions hold:*

1. $-\int_{(-\infty, x_2)} y dF(y) + \int_{(x_2, \infty)} y dF(y) + x_2 \mathbb{P}(Y = x_2) \geq 0$
2. $\int_{(-\infty, x_2)} y dF(y) - \int_{(x_2, \infty)} y dF(y) + x_2 \mathbb{P}(Y = x_2) \geq 0$

where $Y \sim f(\mathbf{x})$ with c.d.f. F .

Proof. The proof is in the appendix B. $\square$

Given a strategyproof, shift invariant, scale invariant and symmetric mechanism, the upcoming results demonstrate how to find another strategyproof, shift invariant, scale invariant and symmetric mechanism that restricts the probability assignment to x_1, x_2 , and $m_{\mathbf{x}}$ for every profile $\mathbf{x}$ and simultaneously gives a weakly better approximation than the original mechanism.

Lemma 7. *Let f be a strategyproof, shift invariant, scale invariant and symmetric mechanism. There exists another strategyproof, shift invariant, scale invariant and symmetric mechanism g with a weakly smaller expected social cost on every profile, such that at least one of the following two properties holds:*

(1) For every two-agent profile $\mathbf{x}$, $\mathbb{P}(g(\mathbf{x}) \in (x_1, x_2)) = 0$ for every two-agent profile $\mathbf{x}$. (Doesn't utilize interior)¹⁶

(2) For every two-agent profile $\mathbf{x}$, $\mathbb{P}(g(\mathbf{x}) \in (-\infty, x_1) \cup (x_2, \infty)) = 0$ for every two-agent profile $\mathbf{x}$. (Ex-post Pareto efficiency)

Proof. The proof is in the appendix B. □

Lemma 8. *Let f be a strategyproof, shift invariant, scale invariant, symmetric mechanism. Assume that f is either ex-post Pareto efficient or doesn't utilize interior. Then there exists another strategyproof mechanism g with a weakly smaller expected social cost on every profile, such that $\mathbb{P}(g(\mathbf{x}) \in \{x_1, x_2, m_{\mathbf{x}}\}) = 1$ for every location profile $\mathbf{x}$. Furthermore, g satisfies shift invariance, scale invariance and symmetry.*

Proof. We break the proof into two cases.

1. Assume f is ex-post Pareto efficient. Let g be the mechanism that satisfies $\mathbb{P}(g(\mathbf{x}) = x_1) = \mathbb{P}(f(\mathbf{x}) = x_1)$, $\mathbb{P}(g(\mathbf{x}) = x_2) = \mathbb{P}(f(\mathbf{x}) = x_2)$, $\mathbb{P}(g(\mathbf{x}) = m_{\mathbf{x}}) = 1 - \mathbb{P}(g(\mathbf{x}) = x_1) - \mathbb{P}(g(\mathbf{x}) = x_2)$. Note that since $m_{\mathbf{x}}$ minimizes the social cost function for the profile $\mathbf{x}$, g certainly provides a weakly better approximation ratio than f . Furthermore, symmetry, shift and scale invariance are preserved.

Let us prove that condition 1 in Lemma 6 holds for g ; the proof for condition 2 is similar. Since f is a strategyproof mechanism, the condition implies that for any profile $\mathbf{x} = (x_1, x_2)$ with $x_1 = 0 < x_2$,

$$\begin{aligned}
0 &\leq - \int_{[0, x_2)} y dF(y) + x_2 \mathbb{P}(f(\mathbf{x}) = x_2) \\
&= - \int_{(0, x_2)} y dF(y) + x_2 \mathbb{P}(f(\mathbf{x}) = x_2) \\
&= -\mathbb{E}[f(\mathbf{x}) \mathbb{1}(f(\mathbf{x}) \in (x_1, x_2))] + x_2 \mathbb{P}(f(\mathbf{x}) = x_2) \\
&= -m_{\mathbf{x}} \mathbb{P}(f(\mathbf{x}) \in (x_1, x_2)) + x_2 \mathbb{P}(f(\mathbf{x}) = x_2) \\
&= -m_{\mathbf{x}} \mathbb{P}(g(\mathbf{x}) = m_{\mathbf{x}}) + x_2 \mathbb{P}(g(\mathbf{x}) = x_2) \\
&= - \int_{(0, x_2)} y dG(y) + x_2 \mathbb{P}(g(\mathbf{x}) = x_2).
\end{aligned}$$

The third equality holds because the distribution is symmetric around $m_{\mathbf{x}}$. Hence, the condition is satisfied for the mechanism g .

¹⁶Note that it is possible for such a mechanism to still be ex-post Pareto efficient, if $\mathbb{P}(g(\mathbf{x}) \in \{x_1, x_2\})$.

2. Assume f doesn't utilize interior. Let g be the mechanism which, for every profile $\mathbf{x}$, locates $\mathbb{P}(g(\mathbf{x}) = x_1) = \mathbb{P}(g(\mathbf{x}) = x_2) = 0.5$, which is clearly strategyproof, shift invariant, scale invariant, and symmetric. $sc(\mathbf{x}, x_2)$ minimizes $sc(\mathbf{x}, y)$ among $y \geq x_2$ and $sc(\mathbf{x}, x_1)$ minimizes $sc(\mathbf{x}, y)$ among $y \leq x_1$. Hence, $\mathbb{E}[sc(\mathbf{x}, g(\mathbf{x}))] \leq \mathbb{E}[sc(\mathbf{x}, f(\mathbf{x}))]$.

□

Now we are ready to prove the main theorem.

Theorem 4. *The LRM mechanism gives the best approximation ratio among all strategyproof mechanisms that are shift invariant, scale invariant and ex-post Pareto efficient.*

Proof. By the previous lemma, it suffices to search among the class of strategyproof shift invariant, scale invariant and symmetric mechanisms where any element f of the class satisfies the property that $\mathbb{P}(f(\mathbf{x}) \in \{x_1, x_2, m_{\mathbf{x}}\}) = 1$ for every location profile $\mathbf{x}$. Clearly, for such mechanisms, the approximation ratio increases as $\mathbb{P}(f(\mathbf{x}) \in \{x_1, x_2\})$ increases. Assume $\mathbb{P}(f(\mathbf{x}) \in \{x_1, x_2\}) < 0.5$. Then $\mathbb{P}(f(\mathbf{x}) = m_{\mathbf{x}}) > 0.5$, and by symmetry, $\mathbb{P}(f(\mathbf{x}) = x_2) < 0.25$. But this gives, when $x_1 = 0$ and $x_2 > 0$, that $-m_{\mathbf{x}}\mathbb{P}(f(\mathbf{x}) = m_{\mathbf{x}}) + x_2\mathbb{P}(f(\mathbf{x}) = x_2) = -\frac{x_2}{2}\mathbb{P}(f(\mathbf{x}) = m_{\mathbf{x}}) + x_2\mathbb{P}(f(\mathbf{x}) = x_2) < 0$, violating strategyproofness by Lemma 6. Thus we must have that $\mathbb{P}(f(\mathbf{x}) \in \{x_1, x_2\}) \geq 0.5$, which implies that among all such mechanisms, LRM provides the best approximation ratio of $0.5(2^{1-\frac{1}{p}} + 1)$. □

An immediate consequence of Theorem 4 is the following corollary.

Corollary 1. *Any strategyproof shift and scale invariant mechanism has an approximation of at least $0.5(2^{1-\frac{1}{p}} + 1)$ in the worst case.*

5 Discussion

The most important open question in our view is whether or not randomization can help improve the worst-case approximation ratio for general L_p norm cost functions. The case of $p = 1$ is uninteresting because there is an optimal deterministic mechanism; for $p = 2$ and $p = \infty$ we already saw that randomization improves the worst-case approximation ratio, but we do not know if this is simply a happy coincidence, or if one can obtain similar results for all $p > 2$. Our negative result in Section 4 implies that any improvement by randomization would require a different approach than the existing mechanisms.

There are many other natural questions as well: for instance, what happens for more general topologies such as trees or cycles? Is it possible to characterize all randomized strategyproof mechanisms on specific topologies?

Finally, we believe it is of interest to consider more general cost functions for the individual agents. The properties established for LRM and many other randomized mechanisms depend on the assumption that agents incur costs that are *exactly equal* to the distance to access the facility. Clearly, this is a very restrictive assumption, and working with more general individual agent costs is an interesting direction to broaden the applicability of this class of models (see Appendix A for a result regarding this direction).¹⁷

6 Acknowledgements

The authors would like to thank the anonymous referees of this paper for their many valuable suggestions, as well as the idea behind appendix C.

References

- [1] Noga Alon, Michal Feldman, Ariel D. Procaccia, and Moshe Tennenholtz. Strategyproof approximation of the minimax on networks. *Math. Oper. Res.*, 35(3):513–526, 2010.
- [2] Noga Alon, Michal Feldman, Ariel D. Procaccia, and Moshe Tennenholtz. Walking in circles. *Discrete Mathematics*, 310(23):3432–3435, 2010.
- [3] H. J. Bandelt and M. Labbé. How bad can a voting location be. *Social Choice and Welfare*, 3(2):125–145.
- [4] Hans-Jrgen Bandelt. Networks with condorcet solutions. *European journal of operational research : EJOR*, 20(3):314–326, 1985.
- [5] Salvador Barberà, Jordi Massó, and Shigehiro Serizawa. Strategy-proof voting on compact ranges. *Games and Economic Behavior*, 25(2):272–291, 1998.
- [6] Margaret L Brandeau and Samuel S Chiu. Parametric facility location on a tree network with an l p-norm cost function. *Transportation Science*, 22(1):59–69, 1988.
- [7] Yukun Cheng and Sanming Zhou. A survey on approximation mechanism design without money for facility games. In *Advances in Global Optimization*, pages 117–128. Springer, 2015.
- [8] Vladimir I Danilov. The structure of non-manipulable social choice rules on a tree. *Mathematical Social Sciences*, 27(2):123–131, 1994.

¹⁷For deterministic mechanisms, our result continues to hold for arbitrary single peaked cost functions, as long as the social cost remains an L_p measure of the distances.

- [9] Michal Feldman and Yoav Wilf. Strategyproof facility location and the least squares objective. In *Proceedings of the Fourteenth ACM Conference on Electronic Commerce*, EC '13, pages 873–890, New York, NY, USA, 2013. ACM.
- [10] Allan Gibbard. Manipulation of voting schemes: A general result. *Econometrica*, 41(4):587–601, 1973.
- [11] Pierre Hansen and Jacques Thisse. Outcomes of voting and planning: Condorcet, weber and rawls locations. *Journal of Public Economics*, 16(1):1–15, 1981.
- [12] M. Labbé. Outcomes of voting and planning in single facility location problems. *European journal of operational research : EJOR*, 20(3):299–313, 1985.
- [13] Hervé Moulin. One dimensional mechanism design. *Working paper*.
- [14] Herve Moulin. On strategy-proofness and single-peakedness. *Public Choice*, 35:437–455, 1980.
- [15] Ariel D Procaccia and Moshe Tennenholtz. Approximate mechanism design without money. *ACM Transactions on Economics and Computation*, 1(4):18, 2013.
- [16] Mark Allen 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.
- [17] James Schummer and Rakesh V Vohra. Strategy-proof location on a network. *Journal of Economic Theory*, 104(2):405–428, 2002.
- [18] Vijay V. Vazirani. *Approximation Algorithms*. Springer-Verlag New York, Inc., New York, NY, USA, 2001.
- [19] Jens Vygen. Approximation algorithms for facility location problems. *Lecture Notes, University of Bonn*.

Appendix A An Alternative Definition of Individual Cost

Let g be a strictly increasing and convex $\mathcal{C}_1$ function on $[0, \infty)$ with $g(0) = g'(0) = 0$. Note that $g(x) = x^p$ satisfies this description for all $p > 1$. We consider a scenario where the cost of agent i is $C(x_i, y) = g(|x_i - y|)$ when the mechanism is deterministic and locates the facility at y . Similarly $C(x_i, \pi) = \mathbb{E}_{y \sim \pi}[g(|x_i - y|)]$ when the mechanism is randomized and locates the facility according to distribution π . The social cost function $h(|x_1 - y|, |x_2 - y|)$ is only assumed to be

(1) anonymous ($h(d, d') = h(d', d)$) and (2) satisfy that for all $a \in (\min\{x_1, x_2\}, \max\{x_1, x_2\})$ where $x_1 \neq x_2$, $h(|x_1 - a|) + h(|x_2 - a|) < h(|x_2 - x_1|)$. Note that for $p > 1$, the L_p norm of the distances and the L_p norm of the costs (for the general g above) both satisfy these conditions. We show that in this case, no randomized strategyproof mechanism satisfying shift invariance, scale invariance and ex-post Pareto efficiency for $n = 2$ can help us improve the approximation ratio relatively to the median mechanism.

Theorem 5. *Let f be a randomized mechanism satisfying shift invariance and scale invariance, and ex-post Pareto efficiency for $n = 2$. Assume f is strategyproof with respect to the individual cost function $C(x_i, y) = g(|x_i - y|)$, where g is a strictly increasing and convex $\mathcal{C}_1$ function on $[0, \infty)$ with $g(0) = g'(0) = 0$. If the social cost function satisfies (1) and (2), then the approximation ratio of f is at least as large as the median's.*

Proof. Using a proof similar to that of Lemma 4, we may assume without loss of generality that f is symmetric. Consider a profile where $n = 2$ and $x_1 = 0, x_2 = 1$. Let $Y = f(0, 1)$. We would like that $\mathbb{P}(Y \in (0, 1)) = 0$. Suppose for the sake of contradiction that there exists $x \in (0, \frac{1}{2})$ such that $\mathbb{P}(Y \in (x, 1 - x)) = q > 0$. Now suppose agent 2 now misreports his location to $1 + \epsilon$ for some small $\epsilon > 0$ such that $\frac{1}{1+\epsilon} > 1 - x$. By shift and scale invariance, $f(0, 1 + \epsilon) = (1 + \epsilon)Y$. Then the difference in cost for agent 2 between the two profile of reports is

$$\begin{aligned} \mathbb{E}[g(|1 - (1 + \epsilon)Y|)] - \mathbb{E}[g(|1 - Y|)] &= - \int_0^{\frac{1}{1+\epsilon}} (g(1 - y) - g(1 - (1 + \epsilon)y)) dF(y) \\ &\quad + \int_{\frac{1}{1+\epsilon}}^1 (g((1 + \epsilon)y - 1) - g(1 - y)) dF(y) \\ &\leq \mathbb{P}(Y \in [\frac{1}{1+\epsilon}, 1])g(\epsilon) - q(g(1 - x^*) - g(1 - (1 + \epsilon)x^*)) \end{aligned}$$

where $x^* \in \arg \min_{y \in [x, 1-x]} g(1 - y) - g(1 - (1 + \epsilon)y)$. The inequality follows from the fact that $g((1 + \epsilon)y - 1) - g(1 - y) \leq g(\epsilon)$ for all $y \in [\frac{1}{1+\epsilon}, 1]$ and that $g(1 - y) - g(1 - (1 + \epsilon)y) \geq g(1 - x^*) - g(1 - (1 + \epsilon)x^*)$ for all $y \in [x, 1 - x]$. Note that

$$\begin{aligned} \lim_{\epsilon \rightarrow 0^+} \frac{\mathbb{E}[g(|1 - (1 + \epsilon)Y|)] - \mathbb{E}[g(|1 - Y|)]}{\epsilon} &\leq \lim_{\epsilon \rightarrow 0^+} \mathbb{P}(Y \in [\frac{1}{1+\epsilon}, 1]) \frac{g(\epsilon)}{\epsilon} - q \frac{g(1 - x^*) - g(1 - (1 + \epsilon)x^*)}{\epsilon} \\ &\leq \mathbb{P}(Y = 1)g'(0) - qg'(1 - x^*)x^* < 0 \end{aligned}$$

The third inequality follows from $g'(0) = 0$ and $g'(1 - x^*) > 0$ (since g is strictly convex). This implies that $\mathbb{E}[g(|1 - (1 + \epsilon)Y|)] - \mathbb{E}[g(|1 - Y|)] < 0$ for ϵ sufficiently small, implying that there is a profitable deviation for agent 2. $\square$

Appendix B Omitted proofs from section 4.2

Proof of Lemma 6. First, let us prove that the two conditions imply strategyproofness. By shift invariance and anonymity, it suffices to check strategyproofness for profiles where $x_1 = 0$ and $x_2 \geq 0$. Moreover, any scale invariant mechanism is trivially strategyproof with respect to the profile $(0, 0)$ since scale invariance implies $f(0, 0) = 0$, which means that no agent has incentive to misreport his location.¹⁸ Thus, we can assume that $x_2 > 0$. It suffices to show that agent 2 cannot benefit by deviating from his true location if the two aforementioned conditions hold. Since $x_2 > 0$, we can denote agent 2's deviation x'_2 as cx_2 for some $c \in \mathbb{R}$. Moreover, we can assume that $c \geq 0$. This can be justified as follows. Assume $c < 0$. Note that by symmetry, in any fixed profile $\mathbf{z}$, the closer a point is to $m_{\mathbf{z}}$, the smaller the expected distance of the facility is from that point. In particular, this implies that $C(x_2, f(0, -cx_2)) \leq C(-x_2, f(0, -cx_2))$. But also note that by scale invariance, $C(-x_2, f(0, -cx_2)) = C(x_2, f(0, cx_2))$. Thus, $C(x_2, f(0, -cx_2)) \leq C(x_2, f(0, cx_2))$. Consequently, if reporting cx_2 is a profitable deviation for agent 2 for some $c < 0$, then reporting $-cx_2$ is also a profitable deviation for the agent.

When agent 2 reports his location to be cx_2 , where $c > 1$, the change in cost incurred by agent 2 is (where C_{orig} is the expected cost of agent 2 under truthful reporting and C_{dev} is the expected cost of agent 2 under misreporting):

$$\begin{aligned}
C_{dev} - C_{orig} &= -(c-1) \int_{(-\infty, \frac{x_2}{c})} y dF(y) + \int_{[\frac{x_2}{c}, x_2)} ((c+1)y - 2x_2) dF(y) + (c-1) \int_{(x_2, \infty)} y dF(y) + \\
&\quad + (c-1)x_2 \mathbb{P}(Y = x_2) \\
&= -(c-1) \int_{(-\infty, x_2)} y dF(y) + \int_{[\frac{x_2}{c}, x_2)} (2cy - 2x_2) dF(y) + (c-1) \int_{(x_2, \infty)} y dF(y) \\
&\quad + (c-1)x_2 \mathbb{P}(Y = x_2) \\
&\geq -(c-1) \int_{(-\infty, x_2)} y dF(y) + (c-1) \int_{(x_2, \infty)} y dF(y) + (c-1)x_2 \mathbb{P}(Y = x_2).
\end{aligned}$$

Hence, when condition 1 holds, we have that $-(c-1) \int_{(-\infty, x_2)} y dF(y) + (c-1) \int_{(x_2, \infty)} y dF(y) + (c-1)x_2 \mathbb{P}(Y = x_2) \geq 0$, which means that $C_{dev} - C_{orig} \geq 0$.

¹⁸ $f(0, 0) = 0$ follows from, say, $f(0, 0) = f(0 \cdot 1, 0 \cdot 1) = 0 \cdot f(1, 1) = 0$, where the second equality is by scale invariance.

Similarly, when $0 \leq c < 1$, the change in cost incurred by agent 2 is:

$$\begin{aligned}
C_{dev} - C_{orig} &= (1-c) \int_{(-\infty, x_2)} y dF(y) + \int_{(x_2, \frac{x_2}{c}]} (2x_2 - (c+1)y) dF(y) - (1-c) \int_{(\frac{x_2}{c}, \infty)} y dF(y) + \\
&\quad + (1-c)x_2 \mathbb{P}(Y = x_2) \\
&= (1-c) \int_{(-\infty, x_2)} y dF(y) + \int_{(x_2, \frac{x_2}{c}]} (2x_2 - 2cy) dF(y) - (1-c) \int_{(x_2, \infty)} y dF(y) \\
&\quad + (1-c)x_2 \mathbb{P}(Y = x_2) \\
&\geq (1-c) \int_{(-\infty, x_2)} y dF(y) - (1-c) \int_{(x_2, \infty)} y dF(y) + (1-c)x_2 \mathbb{P}(Y = x_2).
\end{aligned}$$

Hence, when condition 2 holds, we have that $(1-c) \int_{(-\infty, x_2)} y dF(y) - (1-c) \int_{(x_2, \infty)} y dF(y) + (1-c)x_2 \mathbb{P}(Y = x_2) \geq 0$, which means that $C_{dev} - C_{orig} \geq 0$. Hence, the mechanism is strategyproof for any profile $\mathbf{x}$ with $x_1 = 0 < x_2$.

To prove the other direction, suppose condition 1 does not hold for some profile $\mathbf{x}$ with $x_1 = 0 < x_2$. Then there exists $\epsilon > 0$ small enough such that $-\int_{(-\infty, x_2)} y dF(y) + \int_{(x_2, \infty)} y dF(y) + x_2 \mathbb{P}(Y = x_2) \leq -\epsilon$ for some $x_2 > 0$. We choose $c > 1$ s.t. $\mathbb{P}(Y \in [\frac{x_2}{c}, x_2)) < \frac{\epsilon}{4x_2}$, then we have that

$$\begin{aligned}
C_{dev} - C_{orig} &= -(c-1) \int_{(-\infty, x_2)} y dF(y) + \int_{[\frac{x_2}{c}, x_2)} (2cy - 2x_2) dF(y) + (c-1) \int_{(x_2, \infty)} y dF(y) \\
&\quad + (c-1)x_2 \mathbb{P}(Y = x_2) \\
&\leq (c-1) \left(- \int_{(-\infty, x_2)} y dF(y) + \int_{[\frac{x_2}{c}, x_2)} (2x_2) dF(y) + \int_{(x_2, \infty)} y dF(y) + x_2 \mathbb{P}(Y = x_2) \right) \\
&< -(c-1) \frac{\epsilon}{2} < 0,
\end{aligned}$$

which contradicts strategyproofness of the mechanism.

Similarly, suppose condition 2 does not hold for some profile $\mathbf{x}$ with $x_1 = 0 < x_2$. Then there exists $\epsilon > 0$ small enough such that $\int_{(-\infty, x_2)} y dF(y) - \int_{(x_2, \infty)} y dF(y) + x_2 \mathbb{P}(Y = x_2) \leq -\epsilon$ for some $x_2 > 0$. We choose $0 < c < 1$ s.t. $\mathbb{P}(Y \in (x_2, \frac{x_2}{c}]) < \frac{\epsilon}{4x_2}$, then we have that

$$\begin{aligned}
C_{dev} - C_{orig} &= (1-c) \int_{(-\infty, x_2)} y dF(y) + \int_{(x_2, \frac{x_2}{c}]} (2x_2 - 2cy) dF(y) - (1-c) \int_{(x_2, \infty)} y dF(y) \\
&\quad + (1-c)x_2 \mathbb{P}(Y = x_2) \\
&\leq (1-c) \left(\int_{(-\infty, x_2)} y dF(y) + \int_{[\frac{x_2}{c}, x_2)} (2x_2) dF(y) - \int_{(x_2, \infty)} y dF(y) + x_2 \mathbb{P}(Y = x_2) \right) \\
&< -(1-c) \frac{\epsilon}{2} < 0,
\end{aligned}$$

which contradicts strategyproofness of the mechanism. $\square$

Proof of Lemma 7. Let f be as given above. Assume f violates both (1) and (2) on some profile $\mathbf{x}$ (otherwise, there is nothing to prove: we can take $g = f$). By shift invariance we may assume without loss of generality that $x_1 = 0$. We may assume by anonymity and shift invariance that $x_1 = 0 < x_2$. Let $Y \sim f(\mathbf{x})$. Let $p_1 = \mathbb{P}(Y \in (m_{\mathbf{x}}, x_2)) + \frac{\mathbb{P}(Y=m_{\mathbf{x}})}{2} = \mathbb{P}(Y \in (x_1, m_{\mathbf{x}})) + \frac{\mathbb{P}(Y=m_{\mathbf{x}})}{2}$, $p_2 = \mathbb{P}(x_2, \infty) = \mathbb{P}(-\infty, x_1)$, $z_1 = \frac{\mathbb{E}[Y \mathbb{1}(Y \in (x_1, m_{\mathbf{x}}))] + m_{\mathbf{x}} \mathbb{P}(Y=m_{\mathbf{x}})/2}{p_1}$, and $z'_1 = \mathbb{E}[Y | Y \in (-\infty, x_1)]$, $z_2 = \frac{\mathbb{E}[Y \mathbb{1}(Y \in (m_{\mathbf{x}}, x_2))] + m_{\mathbf{x}} \mathbb{P}(Y=m_{\mathbf{x}})/2}{p_1}$, $z'_2 = \mathbb{E}[Y | Y \in (x_2, \infty)]$.¹⁹

Consider a random variable Y'' obtained from Y as follows: $\mathbb{P}(Y'' \in \{z'_1, x_1, z_1, z_2, x_2, z'_2\}) = 1$, $\mathbb{P}(Y'' = z'_1) = \mathbb{P}(Y'' = z'_2) = p_2$, $\mathbb{P}(Y'' = z_1) = \mathbb{P}(Y'' = z_2) = p_1$, and $\mathbb{P}(Y'' = x_1) = \mathbb{P}(Y'' = x_2) = \mathbb{P}(Y = x_1) = \mathbb{P}(Y = x_2)$. Clearly, Y'' is symmetric about the midpoint $m_{\mathbf{x}}$. Since the social cost function is convex, it follows that $\mathbb{E}[sc(\mathbf{x}, Y'')] \leq \mathbb{E}[sc(\mathbf{x}, Y)]$.

Now, consider a random variable Y' obtained from Y'' as follows. We construct Y' from Y'' by shifting parts of the probability mass at z_1 and z'_1 to x_1 as well as by shifting parts of the probability mass at z_2 and z'_2 to x_2 while ensuring that $\mathbb{E}[Y'] = \mathbb{E}[Y'']$. Specifically, since $z_1 < x_1 < z'_1$, we can write $x_1 = \lambda z_1 + (1 - \lambda)z'_1$ for some $0 < \lambda < 1$. One way to shift the probability mass is to subtract probability λp and $(1 - \lambda)p$ from z_1 and z'_1 respectively and add probability p to x_1 for p sufficiently small (do the same transformation for points z_2 , z'_2 , and x_2). This transformation ensures $\mathbb{E}[Y'] = \mathbb{E}[Y'']$ because

$$(p_1 - \lambda p)z_1 + (p_2 - (1 - \lambda)p)z_2 + (\mathbb{P}(Y'' = x_1) + p)x_1 = p_1 z_1 + p_2 z_2 + \mathbb{P}(Y'' = x_1)x_1.$$

In order to maximize the shift in probability mass, we choose the largest p possible or $p = \min(\frac{p_1}{\lambda}, \frac{p_2}{1-\lambda})$. If $p = \frac{p_1}{\lambda}$, then $\mathbb{P}(Y' \in \{z'_1, x_1, x_2, z'_2\}) = 1$, as $\mathbb{P}(Y' = z'_1) = \mathbb{P}(Y' = z'_2) = p_2 - (1 - \lambda)p$, and $\mathbb{P}(Y' = x_1) = \mathbb{P}(Y' = x_2) = \mathbb{P}(Y'' = x_1) + p$. Else if $p = \frac{p_2}{1-\lambda}$, then $\mathbb{P}(Y' \in \{x_1, z_1, z_2, x_2\}) = 1$, $\mathbb{P}(Y' = z_1) = \mathbb{P}(Y' = z_2) = p_1 - \lambda p$, and $\mathbb{P}(Y' = x_1) = \mathbb{P}(Y' = x_2) = \mathbb{P}(Y'' = x_1) + p$. It is clear from construction that Y' is symmetric about $m_{\mathbf{x}}$. Convexity implies $\mathbb{E}[sc(\mathbf{x}, Y')] \leq \mathbb{E}[sc(\mathbf{x}, Y'')]$, and so $\mathbb{E}[sc(\mathbf{x}, Y')] \leq \mathbb{E}[sc(\mathbf{x}, Y)]$.

Now, let g be a mechanism that locates the facility according to Y' given profile $\mathbf{x}$. Note that there is a unique way to extend the definition of g to all other two-agent profiles such that g is shift and scale invariant as well as symmetric; let us extend the definition of g that way. Furthermore, this extension is easily seen to imply the following:

¹⁹Note that if $\mathbb{P}(Y = m_{\mathbf{x}}) = 0$, then z_1 is the conditional expectation of Y given that $Y \in (x_1, m_{\mathbf{x}})$. When $\mathbb{P}(Y = m_{\mathbf{x}}) > 0$, imagine that whenever $Y = m_{\mathbf{x}}$, we flip a fair coin; then z_1 is the conditional expectation of Y given that $Y \in (x_1, m_{\mathbf{x}})$ or $Y = m_{\mathbf{x}}$ and the coin lands on heads. z_2 can be defined in a similar manner (replace $(x_1, m_{\mathbf{x}})$ with $(m_{\mathbf{x}}, x_2)$ and heads with tails). From this description it is clear that $z_1 \in (x_1, m_{\mathbf{x}}]$, $z_2 \in [m_{\mathbf{x}}, x_2)$, and that they are symmetric about $m_{\mathbf{x}}$.

1. Since $\mathbb{E}[sc(\mathbf{x}, g(\mathbf{x}))] \leq \mathbb{E}[sc(\mathbf{x}, f(\mathbf{x}))]$ for the profile $\mathbf{x}$, the social cost obtained by mechanism g via the extension is no more than the one obtained by mechanism f for all two-agent profiles.
2. If $\mathbb{P}(g(\mathbf{x}) \in (x_1, x_2)) = 0$, then $\mathbb{P}(g(\mathbf{q}) \in (q_1, q_2)) = 0$ for all two-agent profiles $\mathbf{q}$. Similarly, if $\mathbb{P}(g(\mathbf{x}) \in (-\infty, x_1) \cup (x_2, \infty)) = 0$, then $\mathbb{P}(g(\mathbf{q}) \in (-\infty, q_1) \cup (q_2, \infty)) = 0$ for all two-agent profiles $\mathbf{q}$.

Thus, all that is left for us to do is to show strategyproofness of g . We can do so by verifying the conditions in Lemma 6 (the fact that it holds for all the required profiles is then again immediate by shift and scale invariance). When $p = \frac{p_1}{\lambda}$, we claim that it suffices to show that:

$$-z'_1(p_2 - (1-\lambda)p) + z'_2(p_2 - (1-\lambda)p) + x_2(\mathbb{P}(Y' = x_2) + p) \geq -z'_1p_2 - z_1p_1 - z_2p_1 + z'_2p_2 + x_2\mathbb{P}(Y = x_2),$$

and that

$$z'_1(p_2 - (1-\lambda)p) - z'_2(p_2 - (1-\lambda)p) + x_2(\mathbb{P}(Y' = x_2) + p) \geq z'_1p_2 + z_1p_1 + z_2p_1 - z'_2p_2 + x_2\mathbb{P}(Y = x_2).$$

To justify this claim, we need to show that the right hand sides are always greater than or equal to 0. But note that $z_1, z_2, p_1, z'_1, z'_2$, and p_2 were defined so that the right hand sides amount exactly to the conditions of Lemma 6 for f on the profile $\mathbf{x}$, and thus must be greater than or equal to zero. After some algebra, the two inequalities above reduce to:

$$z'_1(1-\lambda)p - z'_2(1-\lambda)p + x_2p \geq -z_1p_1 - z_2p_1, \quad (3)$$

and

$$-z'_1(1-\lambda)p + z'_2(1-\lambda)p + x_2p \geq z_1p_1 + z_2p_1. \quad (4)$$

To show (3), we know that

$$z'_1(1-\lambda)p + z_1p_1 = (z'_1(1-\lambda) + z_1\lambda)p = x_1p = 0, \quad \text{that is} \quad z'_1(1-\lambda)p = -z_1p_1,$$

and that

$$x_2p = (z'_2(1-\lambda) + z_2\lambda)p \geq z'_2(1-\lambda)p - z_2p_1, \quad \text{that is} \quad x_2p - z'_2(1-\lambda)p \geq -z_2p_1.$$

Combining the two above expressions gives us the desired result. Similarly, (4) follows from the fact that $z'_1(1-\lambda)p + z_1p_1 = 0$ and that $x_2p = (z'_2(1-\lambda) + z_2\lambda)p \geq -z'_2(1-\lambda)p + z_2p_1$. The proof for the case where $p = \frac{p_2}{1-\lambda}$ is similar and so will be omitted. $\square$

Appendix C Alternative assumptions in section 4.2

Theorem 4 holds if we replace the assumption of shift invariance with symmetry. It is clear from the structure of the proof that it is enough to replace Lemma 4 with the following lemma:

Lemma 9. *Given any strategyproof, symmetric and scale invariant mechanism, there exists a strategyproof, symmetric, scale and shift invariant mechanism with a weakly smaller worst-case approximation ratio.*

Proof. Given a mechanism f , define $g(\mathbf{x}) = f(0, x_2 - x_1) + x_1$. Assume f is strategyproof, symmetric and scale invariant. We claim that g is strategyproof, symmetric, scale and shift invariant with a weakly smaller worst-case approximation ratio. The fact that g is shift invariant and has a weakly smaller worst-case approximation ratio than f is immediate. Let $Y_{x_1, x_2} \sim f(\mathbf{x})$ and $Y'_{x_1, x_2} \sim g(\mathbf{x})$; the relevant equalities below are in distribution.

1. g is symmetric: let $\mathbf{x} \in \mathbb{R}^2$, and let $b \in \mathbb{R}$. Then $\mathbb{P}(Y'_{x_1, x_2} \geq m_{\mathbf{x}} + b) = \mathbb{P}(Y_{0, x_2 - x_1} \geq m_{\mathbf{x}} + b - x_1) = \mathbb{P}(Y_{0, x_2 - x_1} \geq m_{(0, x_2 - x_1)} + b) = \mathbb{P}(Y_{0, x_2 - x_1} \leq m_{(0, x_2 - x_1)} - b) = \mathbb{P}(Y_{0, x_2 - x_1} \leq m_{\mathbf{x}} - b - x_1) = \mathbb{P}(Y_{0, x_2 - x_1} + x_1 \leq m_{\mathbf{x}} - b) = \mathbb{P}(Y'_{x_1, x_2} \leq m_{\mathbf{x}} - b)$.
2. g is scale invariant: let $\mathbf{x} \in \mathbb{R}^2$ and let $c \in \mathbb{R}$. Then $Y_{cx_1, cx_2} = Y_{0, c(x_2 - x_1)} + cx_1 = cY_{0, x_2 - x_1} + cx_1 = c(Y_{0, x_2 - x_1} + x_1) = cY'_{x_1, x_2}$. The second equality follows from scale invariance of f .
3. g is strategyproof: let $\mathbf{x} \in \mathbb{R}^2$, $b, x'_2 \in \mathbb{R}$. There are two cases:
 - (a) Assume $\mathbb{E}[|x_2 - Y'_{x_1, x_2}|] > \mathbb{E}[|x_2 - Y'_{x_1, x'_2}|]$. Note that $\mathbb{E}[|x_2 - Y'_{x_1, x_2}|] = \mathbb{E}[|(x_2 - x_1) - Y_{0, x_2 - x_1}|]$ and $\mathbb{E}[|x_2 - Y'_{x_1, x'_2}|] = \mathbb{E}[|(x_2 - x_1) - Y_{0, x'_2 - x_1}|]$. Thus, it follows that when agent 1's location is 0 and agent 2's location is $x_2 - x_1$, agent 2 can benefit under f when reporting $x'_2 - x_1$ instead, violating strategyproofness of f . Contradiction.
 - (b) Assume $\mathbb{E}[|x_1 - Y'_{x_1, x_2}|] > \mathbb{E}[|x_1 - Y'_{x_1 + b, x_2}|]$. Note that $\mathbb{E}[|x_1 - Y'_{x_1, x_2}|] = \mathbb{E}[|-Y_{0, x_2 - x_1}|] = \mathbb{E}[|(x_2 - x_1) - Y_{0, x_2 - x_1}|]$, where the last equality follows from symmetry of f . Also note that $\mathbb{E}[|x_1 - Y'_{x_1 + b, x_2}|] = \mathbb{E}[|-b - Y_{0, x_2 - x_1 - b}|] = \mathbb{E}[|(x_2 - x_1) - Y_{0, x_2 - x_1 - b}|]$, where again the last equality follows from symmetry of f . Thus, when agent 1's true location is 0 and agent 2's true location is $x_2 - x_1$, then agent 2 benefits under f by reporting $x_2 - x_1 - b$, violating strategyproofness of f . Contradiction.

□