

The genericity of obvious manipulations under the Egalitarian Walrasian rule when preferences are linear*

Pierre Bardier[†] Bach Dong-Xuan[‡] Van-Quy Nguyen[§]

May 2026

Abstract

In this note, we study the incentives properties of one of the central solutions to the problem of allocating divisible commodities among equally entitled agents, namely the Egalitarian Walrasian rule (EW), when agents have linear preferences. The EW rule being manipulable on this domain, we ask whether it is, at least, *immune to obvious manipulations*, in the sense of [Trojan and Morrill \(2020\)](#). Unfortunately, the answer is negative. This manipulability result is *generic* since we show that any agent for whom all the commodities are desirable has an obvious manipulation.¹

Keywords: Fair allocation, obvious manipulations, competitive equilibrium, Egalitarian Walrasian rule, linear preferences.

JEL classification: C72; D63; D71; D82.

1 Introduction

In this note, we study the incentives properties of the Egalitarian Walrasian (EW) allocation rule, when the divisible goods to be allocated among equally entitled agents are perfect substitutes —preferences are linear—, a situation in which the EW rule enjoys numerous

*We thank Susumu Cato, Marc Fleurbaey, Antonin Macé, and William Thomson for their valuable comments.

[†]Paris School of Economics, École Normale Supérieure de Paris (Paris Sciences et Lettres). E-mail: pierre.bardier@psemail.eu

[‡]Center for Intelligence Economic Science, Southwestern University of Finance and Economics. E-mail: bachdongx@gmail.com

[§]Faculty of Mathematical Economics, National Economics University, Vietnam. E-mail: quynv@neu.edu.vn

¹Rather than finding some specific preference relation for which an agent has an obvious manipulation.

desirable properties, but falls short of being *strategy-proof*. The EW rule being manipulable on this domain, we ask whether it is, at least, *immune to obvious manipulations*, in the sense of [Trojan and Morrill \(2020\)](#).

Linear preferences are plausible when the goods under consideration are “unrelated”: when an agent has a linear preference relation, her relative affinity for a good does not vary as the consumption of the other goods vary —the marginal rates of substitution are constant. As a matter of fact, online platforms dedicated to the implementation of fair division solutions work under this assumption, since they ask users to allocate a quota of points among goods ([Bogomolnaia et al. \(2017\)](#)).²

The incentives properties of the EW rule on this domain are all the more important that the rule is *efficient*, decisive in the sense that it is *single-valued in welfare terms*, and satisfies remarkable properties of punctual fairness and solidarity. Indeed, the EW rule meets the two central requirements of punctual fairness, namely *no-envy* and, thus, the *equal-division lower bound property*. It also meets two key solidarity principles, *resource monotonicity* and *population monotonicity*, each of which is, on more general domains,³ incompatible with *efficiency* and the *equal-division lower bound property* (see [Moulin and Thomson \(1988\)](#), [Kim \(2004\)](#), respectively).

It is known that, when preferences are linear, no *efficient* and “reasonably fair” allocation rule is *strategy-proof*. For two agents, the conjunction of *efficiency* and *strategy-proofness* implies *dictatorship* ([Schummer \(1996\)](#)). With at least three agents, no rule satisfies *efficiency*, *strategy-proofness* and *equal treatment of equals in welfare terms* ([Cho and Thomson \(2023\)](#)). Given this impossibility, any sense in which the manipulability of the EW rule is limited would unquestionably be good news.⁴ In this perspective, we rely on the concept of *obvious manipulation*, introduced in [Trojan and Morrill \(2020\)](#).

This concept builds on the idea of an agent who is not able to assess how any of her strategies precisely interplays with the strategies of others in a given mechanism, but still knows the set of possible outcomes this strategy can yield. In a direct revelation mechanism, a (profitable) manipulation is *obvious* if either the best outcome to which it may lead is strictly better than the best outcome to which truthful reporting may lead, *or* the worst outcome to which it may lead is strictly better than the worst outcome to which truthful reporting may lead. Inquiring into rules that are *immune to obvious manipulations* has proved a fruitful direction in many mechanism-design settings for which *strategy-proofness*

²See SPLIDDIT: www.spliddit.org/ and ADJUSTED WINNER: www.nyu.edu/projects/adjustedwinner/.

³All domains containing all monotonic, convex, and homothetic preferences.

⁴Along these lines, we note that [Bogomolnaia et al. \(2017\)](#) obtain a positive result, since they characterize the EW rule, on this domain of preferences, on the basis of efficiency, the *equal-division lower bound property* and a weakening of *Maskin invariance*, which they call *independence of lost bids*.

precludes meeting basic requirements pertaining to non-strategic dimensions such as fairness or efficiency.⁵

Unfortunately, as we show here, when preferences are linear, the EW rule is *subject to obvious manipulations*. Furthermore, it is strikingly so: we show that any agent for whom all the commodities are desirable has an obvious manipulation, rather than finding some specific preference relation for which an agent has an obvious manipulation.⁶

Section 2 describes the model and Section 3 presents the result.

2 Model

The set of real numbers, the set of non-negative real numbers, and the set of positive real numbers are $\mathbb{R}$, $\mathbb{R}_+$ and $\mathbb{R}_{++}$, respectively. A vector of equal coordinates may be denoted by a boldface character, *e.g.* $\mathbf{0} = (0, \dots, 0)$. The symbols $\geq, \geq, >$ denote the usual vector inequalities.⁷ The cardinality of a set A is denoted by $|A|$. The set A^B is the set of mappings from B to A . The set of subsets of a set A is denoted by 2^A . The Cartesian product of A and B is denoted by $A \times B$.

The population of agents involved in the allocation problem is a finite set N ($|N| \geq 2$). There is a finite set of L goods ($|L| \geq 2$). The social endowment is assumed —without loss of generality— to be $\mathbf{1}$.

Each agent i in N has a **preference relation** R_i , which is a complete reflexive and transitive binary relation on her consumption set $X = [0, 1]^L$. For two bundles $x_i, y_i \in X$, we write $x_i R_i y_i$ to denote that agent i is at least as well off at x_i as at y_i .⁸ We restrict attention to preferences R_i for which there exists a_i in the simplex of dimension $(|L|-1)$, denoted by Δ^L , such that, for all $x_i, y_i \in X$,

$$x_i R_i y_i \text{ if and only if } a_i \cdot x_i \geq a_i \cdot y_i,$$

⁵Listing all the papers following this direction is beyond the scope of this short paper.

⁶For indivisible goods, a related result is established by [Psomas and Verma \(2022\)](#), who find a preference relation for which an agent has an obvious manipulation in a mechanism which selects the allocations maximizing the “Nash welfare function”.

⁷For two vectors of same dimension $x = (x_1, \dots, x_K)$ and $y = (y_1, \dots, y_K)$, $x \geq y$ if $x_k \geq y_k$ for all $1 \leq k \leq K$; $x \geq y$ if $x \geq y$ and $x \neq y$; $x > y$ if $x_k > y_k$ for all $1 \leq k \leq K$.

⁸The corresponding strict preference and indifference relations, *i.e.* the asymmetric part of R_i and the symmetric part of R_i are denoted by P_i and I_i , respectively.

where $\cdot$ denotes the standard inner product operator.⁹ We refer to a_i as the **list of affinity parameters** for agent i .

An **allocation** is a list of bundles $x = (x_i)_{i \in N} \in X^N$, such that

$$\mathbf{1} \geq \sum_{i \in N} x_i.$$

The set of allocations is denoted by F .

For this domain of preferences, an **allocation rule** can simply be defined as a mapping

$$f : [\Delta^L]^N \rightarrow 2^F \setminus \emptyset,$$

which assigns a non-empty set of allocations to any profile of lists of affinity parameters $(a_i)_{i \in N}$.

3 Result

Definition 1. Let us denote the **Egalitarian Walrasian rule** by $EW : [\Delta^L]^N \rightarrow 2^F \setminus \emptyset$. For all $a \in [\Delta^L]^N$, for all $x \in F$,

$$x \in EW(a)$$

if and only if

$$\mathbf{1} = \sum_{i \in N} x_i, \text{ and there is } p \in \mathbb{R}_+^L \text{ such that, for all } i \in N, x_i \in \arg \max_{y_i \in \mathbb{R}_+^L} \{a_i \cdot y_i \mid p \cdot y_i \leq p \cdot \frac{\mathbf{1}}{N}\}.$$

It is a consequence of the well-known result from [Eisenberg and Gale \(1959\)](#), that, on this domain, EW is *single-valued in welfare terms*, as it selects all allocations maximizing the “Nash welfare function”, and only those allocations.

Lemma 1 ([Eisenberg and Gale \(1959\)](#)). For all $a \in [\Delta^L]^N$,

$$EW(a) = \arg \max_{x \in F} \prod_N (a_i \cdot x_i)$$

For all $x^*, y^* \in EW(a)$, for all $i \in N$, $a_i \cdot x_i^* = a_i \cdot y_i^*$.

⁹That is, we restrict attention to preferences that are continuous (*i.e.*, for all $x_i \in X$, the sets $U(x_i, R_i) = \{y_i \in X \mid y_i R_i x_i\}$ and $L(x_i, R_i) = \{y_i \in X \mid x_i R_i y_i\}$ are closed), monotonic (*i.e.*, for two bundles $x_i, y_i \in X$, if $x_i \geq y_i$, then $x_i R_i y_i$ and if $x_i > y_i$, then $x_i P_i y_i$), and linear (*i.e.*, for all three bundles $x_i, y_i, z_i \in X$, and all $\lambda \in [0, 1]$, $x_i R_i y_i$ if and only if $[\lambda x_i + (1 - \lambda) z_i] R_i [\lambda y_i + (1 - \lambda) z_i]$).

Because the rule we consider below, as most standard rules, is invariant to rescaling, it is without loss of generality, in the definition above, to assume that a_i lies in the simplex.

A **selection** (from EW) is a function

$$\begin{aligned}\varphi : [\Delta^L]^N &\rightarrow F \\ a &\mapsto (\varphi_i(a))_{i \in N} \in EW(a),\end{aligned}$$

which assigns an egalitarian Walrasian allocation to any profile of lists of affinity parameters $(a_i)_{i \in N}$. Because of Lemma 1, all the selections from EW are equivalent in welfare terms.

Any selection φ induces a **direct mechanism**, with abuse denoted by φ as well, according to which:

- 1) each $i \in N$ submits a vector $a_i \in \Delta^L$. Then,
- 2) each $i \in N$ is allocated $\varphi_i(a)$.

As is standard, in what follows, (a_i, a_{-i}) denotes the profile of lists of affinity parameters when agent i submits a_i and the other agents submit a_{-i} , and $\varphi_i(a_i, a_{-i})$ denotes the bundle allocated to i then. The notation $\varphi_i(a_i, a_j, a_{-ij})$ is analogously defined.

We study the manipulability of direct mechanisms induced by selections of EW , using the notion of *obvious manipulation* (Troyan and Morrill (2020)). The fact that the EW rule is single-valued in welfare terms on this domain greatly facilitates the interpretation and use of this notion.¹⁰

Definition 2. The direct mechanism φ is **subject to obvious manipulations** if there exist $i \in N$, $a_i \in \Delta^L$ and $a'_i \in \Delta^L$, such that:

$$\begin{aligned}(i) \quad & \min_{a_{-i} \in [\Delta^L]^{N \setminus \{i\}}} a_i \cdot \varphi_i(a'_i, a_{-i}) > \min_{a_{-i} \in [\Delta^L]^{N \setminus \{i\}}} a_i \cdot \varphi_i(a_i, a_{-i}), \\ \text{or,} \\ (ii) \quad & \max_{a_{-i} \in [\Delta^L]^{N \setminus \{i\}}} a_i \cdot \varphi_i(a'_i, a_{-i}) > \max_{a_{-i} \in [\Delta^L]^{N \setminus \{i\}}} a_i \cdot \varphi_i(a_i, a_{-i}).\end{aligned}$$

Then, following the interpretation in Troyan and Morrill (2020), Theorem 1 below suggests that a direct mechanism associated with any selection of EW is not only manipulable, but, for almost any linear preference relation, there exist, for an agent with this preference, manipulations that are easy to identify and enact:

Theorem 1. *Let φ be the direct mechanism associated with any selection of the egalitarian Walrasian rule EW .*

¹⁰Without this property, we would need to adapt the definition of obvious manipulations by choosing a particular way of comparing sets of outcomes that do not yield the same utility level.

- Inequality (i) in Definition 2 never holds.
- Let $i \in N$, and $a_i \in \Delta^L$ such that, for all $l \in L$, $a_i^l > 0$. Then, there is $a'_i \in \Delta^L$ such that inequality (ii) holds.

In particular, φ is subject to obvious manipulations.

The set of lists of parameters $a_i \in \Delta^L$ such that, for all $l \in L$, $a_i^l > 0$, is dense in Δ^L . That is why we say that φ is *generically* subject to obvious manipulations.

Before proving the result, let us provide some intuition in an example involving 2 agents and 2 commodities.

Example. Agent 1 and 2 have preferences over bundles $(x, y) \in \mathbb{R}_+^2$ represented by the mappings $u_1(x, y) = px + (1-p)y$ and $u_2(x, y) = qx + (1-q)y$, respectively, where $p, q \in [0, 1]$.

Assume that $p = 0.45$. In the Edgeworth box below, we represent the bundle of agent 1 in all the allocations selected by *EW* as q varies in $[0, 1]$, when agent 1 reports truthfully. In addition, we represent the bundle of agent 1 in all the allocations selected by *EW*, when agent 1 untruthfully reports $p' = 0.9$. Computations are omitted.

Truthful report from agent 1: The blue segment relating the points in which agent 1 receives $(0, \frac{10}{11})$ and $(1, \frac{1}{11})$ represents the selected allocations when $q = p = 0.45$. Its right endpoint, in which agent 1 gets $(1, \frac{1}{11})$, corresponds to the unique allocation selected by *EW* for $q < 0.45$. The green segment relating the points in which agent 1 receives $(0, \frac{10}{11})$ and $(0, 1)$ corresponds to the selected allocations for $0.45 < q \leq \frac{1}{2}$. Its top endpoint, in which agent 1 gets $(0, 1)$, corresponds to the unique allocation selected by *EW* for $q > \frac{1}{2}$. **Untruthful report from agent 1:** the range of attainable bundles in the mechanism when agent 1 reports $p' = 0.9$, as q varies, is represented by the orange and purple segments.

Then, the minimum and maximum attainable utility levels under truth-telling for agent 1 are 0.5, for example at $(0, \frac{10}{11})$, and 0.55, at $(0, 1)$, respectively. If individual 1 misreports $p' = 0.9$, she may get a utility level of 0.75 at $(\frac{4}{9}, 1)$, if agent 2 reports $q > 0.9$, and she may get a utility level of 0.45 at $(1, 0)$, if agent 2 reports $q \leq \frac{1}{2}$. Thus, compared to truthful reporting, her worst-case utility level has decreased, and her best-case utility level has increased.

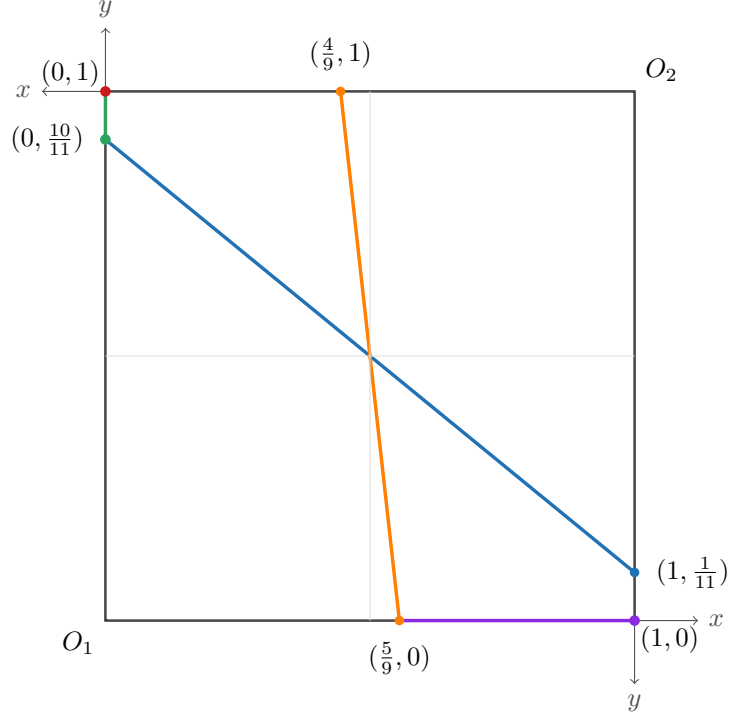


Figure 1. Allocations selected by EW

As her true relative affinity for good 1 is equal to 0.45, agent 1 slightly prefers good 2 to good 1. If she reports such “balanced” tastes, regardless of what agent 2 reports, *EW* will never select an allocation in which agent 1 gets the whole unit of good 2 and some positive quantity of good 1. On the other hand, when she untruthfully reports that she loves good 1 ($p' = 0.9$), if agent 2 reports that she loves good 1 even more than agent 1 does, *efficiency* and *no-envy* require *EW* to select an allocation in which agent 1 gets the whole unit of good 2 and some positive quantity of good 1. Agent 1 is thus better off in this case than in the best-case under truthful reporting. In the following figure, we illustrate the implications of *efficiency* and *no-envy* with such “unbalanced reports”.

The green line is the indifference curve of agent 1 passing through equal-division, when her affinity parameter is $p = 0.9$, the blue one is that of agent 2, when her affinity parameter is $q = 0.95$. The light gray area is the set of allocations without envy. For such an allocation, consider the allocation which is symmetric to it with respect to equal-division. For each agent, her bundle at the allocation lies on an indifference curve which is above the indifference curve passing through the bundle she gets under the symmetric allocation. The red segments represent the set of efficient allocations. The segment relating the two red points is the set of allocations that have both properties. In all these allocations agent 1 gets some positive quantity of good 1.

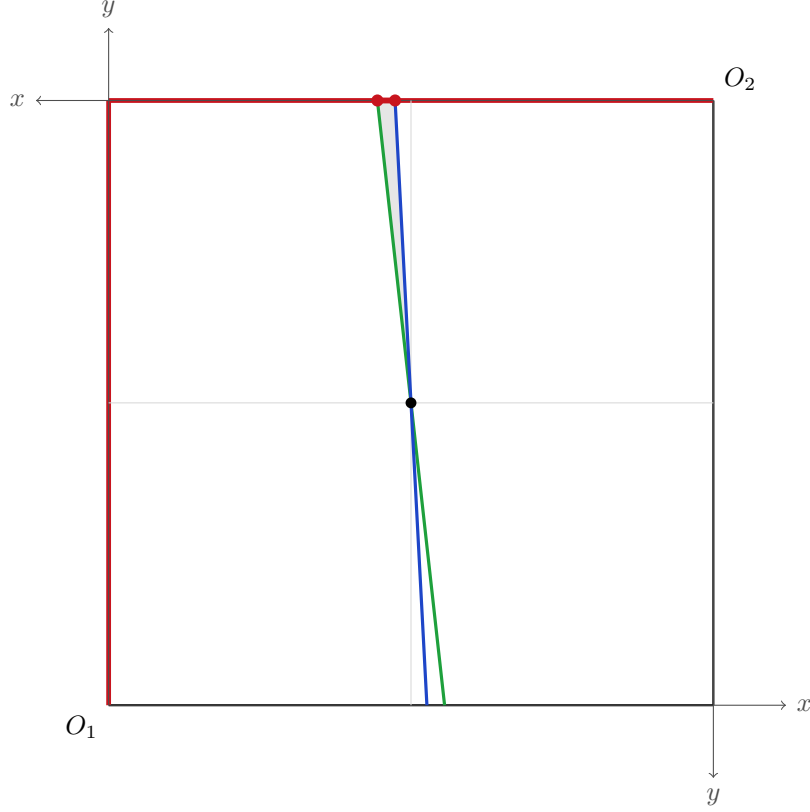


Figure 2. Efficient allocations without envy

The argument in the proof that inequality (ii) in Definition 2 holds is along these lines: for an agent for whom all the commodities are desirable, we prove the existence of an obvious manipulation which consists in inflating her affinity for the good she likes the least.

The proof uses a consequence of Lemma 1 in [Bogomolnaia et al. \(2017\)](#). Following their matrix formulation, for all $x \in F$, we let $[x_i^\ell]$ denote the $|N| \times |L|$ matrix to which x is identified.

Lemma 2 ([Bogomolnaia et al. \(2017\)](#)). *Let $a \in [\Delta^L]^N$ and $x \in EW(a)$. Let $U = (U_i)_{i \in N} = (a_i \cdot x_i)_{i \in N}$. Then, $U > \mathbf{0}$ and there is $z \in EW(a)$, with $(a_i \cdot z_i)_{i \in N} = U$, and*

- (a) *The matrix $[z_i^\ell]$ has at least $(|N|-1)(|L|-1)$ zeros,*
- (b) *for all $i \in N$, $\ell \in L$,*

$$z_i^\ell > 0 \text{ only if, for all } j \in N, \frac{a_i^\ell}{U_i} \geq \frac{a_j^\ell}{U_j}.$$

We decompose the proof of Theorem 1 into two lemmas, corresponding to inequality (i) and inequality (ii) in Definition 2, respectively.

Lemma 3 (Worst-case). *Let φ be the direct mechanism associated with any selection of the egalitarian Walrasian rule EW . Let $i \in N$. Then, for all $a_i \in \Delta^L$, and all $a'_i \in \Delta^L$,*

$$\min_{a_{-i} \in [\Delta^L]^{N \setminus \{i\}}} a_i \cdot \varphi_i(a'_i, a_{-i}) \leq \min_{a_{-i} \in [\Delta^L]^{N \setminus \{i\}}} a_i \cdot \varphi_i(a_i, a_{-i}).$$

Proof. Let $i \in N$. It follows from the definition of EW that, for all $a_i \in \Delta^L$, $a_{-i} \in [\Delta^L]^{N \setminus \{i\}}$, $a_i \cdot \varphi_i(a_i, a_{-i}) \geq a_i \cdot \frac{1}{|N|}$. Thus,

$$\min_{a_{-i} \in [\Delta^L]^{N \setminus \{i\}}} a_i \cdot \varphi_i(a_i, a_{-i}) \geq a_i \cdot \frac{1}{|N|}. \quad (1)$$

Let agent i submit $a'_i \in \Delta^L$, while her actual list of affinity parameters is $a_i = (a_i^\ell)_{\ell \in L}$, and all $j \neq i$ submit $a_j = a_i$. Then, for all $j \neq i$,

$$a_j \cdot \varphi_i(a_j, a'_i, a_{-ij}) \geq a_j \cdot \frac{1}{|N|}.$$

Yet,

$$\begin{aligned} a_i \cdot \varphi_i(a'_i, a_{-i}) &= a_i \cdot \left(1 - \sum_{j \neq i} \varphi_j(a_j, a'_i, a_{-ij}) \right) \\ &= a_i \cdot 1 - \sum_{j \neq i} a_j \cdot \varphi_j(a_j, a'_i, a_{-ij}) \\ &\leq a_i \cdot 1 - (|N| - 1)(a_i \cdot \frac{1}{|N|}) \\ &= a_i \cdot \frac{1}{|N|}. \end{aligned}$$

This implies that

$$\min_{a_{-i} \in [\Delta^L]^{N \setminus \{i\}}} a_i \cdot \varphi_i(a'_i, a_{-i}) \leq a_i \cdot \frac{1}{|N|}, \quad (2)$$

and combining (1) and (2) ends the proof. $\square$

Lemma 4 (Best-case). *Let φ be the direct mechanism associated with any selection of the egalitarian Walrasian rule EW . Let $i \in N$, and $a_i \in \Delta^L$ such that, for all $l \in L$, $a_i^l > 0$. Then, there is $a'_i \in \Delta^L$ such that*

$$\max_{a_{-i} \in [\Delta^L]^{N \setminus \{i\}}} a_i \cdot \varphi_i(a'_i, a_{-i}) > \max_{a_{-i} \in [\Delta^L]^{N \setminus \{i\}}} a_i \cdot \varphi_i(a_i, a_{-i}).$$

Proof. Let $i \in N$, and assume, without loss of generality, that her actual list of affinity

parameters $a_i = (a_i^l)_{l \in L}$ is such that $a_i^1 \geq a_i^2 \geq \dots \geq a_i^L$.¹¹

We first prove that

$$\max_{a_{-i} \in [\Delta^L]^{N \setminus \{i\}}} a_i \cdot \varphi_i(a_i, a_{-i}) = \sum_{l \neq L} a_i^l. \quad (3)$$

Let $\bar{U}_i = \sum_{l \neq L} a_i^l$. Assume that all $j \neq i$ submit $\bar{a}_j = (0, 0, \dots, 0, 1)$. By Lemma 1, $\varphi(a_i, (\bar{a}_j)_{j \neq i})$ is a solution of

$$\max_{x \in F} (a_i \cdot x_i) \times \prod_{j \neq i} (\bar{a}_j \cdot x_j),$$

which is equivalent to

$$\max_{x \in \Lambda} (a_i^L x_i^L + a_i^L x_i^L) \times \prod_{j \neq i} x_j^L,$$

where $\Lambda = \{x \in F \mid x_i^\ell = 1 \text{ for all } \ell \neq L\}$.

Since $\bar{U}_i > a_i^L$ (because $a_i^1 \geq a_i^2 \geq \dots \geq a_i^L$), the optimum x^* of this last program is such that $x_i^{*L} = 0$ and $x_j^{*L} = \frac{1}{|N|-1}$, for all $j \neq i$.

Therefore,

$$\max_{a_{-i} \in [\Delta^L]^{N \setminus \{i\}}} a_i \cdot \varphi_i(a_i, a_{-i}) \geq \bar{U}_i.$$

Now, consider any possible submitted list of affinity parameters for other agents than agent i , denoted by a_{-i} . Let $a = (a_i, a_{-i})$ and $x = \varphi(a)$. Let $U = (U_k)_{k \in N} = (a_k \cdot x_k)_{k \in N}$. By Lemma 2, there is $z \in EW(a)$, with $(a_k \cdot z_k)_{k \in N} = U$ and

- a) the $|N| \times |L|$ matrix $[z_k^\ell]$ has at least $(|N|-1)(|L|-1)$ zeros, and
- b) for all $k \in N$ and for all $\ell \in L$,

$$z_k^\ell > 0 \text{ only if, for all } j \in N, \frac{a_k^\ell}{U_k} \geq \frac{a_j^\ell}{U_j}.$$

There are two possible cases.

- Case 1: $z_i^\ell = 0$ for some $\ell \in L$. Then, clearly,

$$U_i = a_i \cdot z_i \leq \bar{U}_i.$$

- Case 2: $z_i^\ell > 0$ for all $\ell \in L$. For all $j \in N$, let b_j be the number of zeros in $[z_j^\ell]$. Then,

¹¹We do not assume that all the affinity parameters are positive for the first step below.

$b_i = 0$, $b_j \leq |L|-1$ for all $j \neq i$, and

$$\sum_N b_k \geq (|N|-1)(|L|-1).$$

Therefore, $b_j = |L|-1$ for all $j \neq i$. Consider $k \neq i$ and $\ell \in L$ such that $z_k^\ell > 0$. Then, $U_k = a_k^\ell z_k^\ell$. From point b) above, $\frac{a_i^\ell}{U_i} = \frac{a_k^\ell}{U_k} = \frac{1}{z_k^\ell}$, which implies that $U_i = a_i^\ell z_k^\ell \leq \bar{U}_i$.

We have thus proved (3).

Now, we assume that $a_i^1 \geq a_i^2 \geq \dots \geq a_i^L > 0$. Her best-case utility when submitting a_i , obtained from (3), is $\bar{U}_i = \sum_{\ell \neq L} a_i^\ell$.

Assume that she submits $a'_i = (\varepsilon, \dots, \varepsilon, 1 - (|N|-1)\varepsilon)$ for some $0 < \varepsilon < 1$. In addition, assume that all $j \neq i$ submit $a'_j = (0, \dots, 0, 1)$. Let $a' = (a'_i, a'_{-i})$. Once again, by Lemma 1, $\varphi(a')$ solves

$$\max_{x \in F} (a'_i \cdot x_i) \times \prod_{j \neq i} x_j^1.$$

Clearly, at the optimal solution x' , $x_i'^\ell = 1$, $x_j'^\ell = 0$ for all $\ell \neq L$, $j \neq i$, and the problem becomes

$$\max_{x^L \in \Delta^N} [(1 - (|N|-1)\varepsilon)x_i^L + (|N|-1)\varepsilon] \times \prod_{j \neq i} x_j^L.$$

By the *arithmetic mean–geometric mean inequality*, hereafter denoted AM-GM inequality, (see, e.g. [Steele \(2004\)](#)):

$$\prod_{j \neq i} x_j^L \leq \left(\frac{\sum_{j \neq i} x_j^L}{|N|-1} \right)^{|N|-1} = \left(\frac{1 - x_i^L}{|N|-1} \right)^{|N|-1}.$$

Applying once again the AM-GM inequality,

$$\left(\frac{1 - x_i^L}{|N|-1} \right)^{|N|-1} \left(x_i^L + \frac{\varepsilon(|N|-1)}{1 - \varepsilon(|N|-1)} \right) \leq \left(\frac{1}{|N|(1 - \varepsilon(|N|-1))} \right)^{|N|},$$

with equality if

$$\frac{1 - x_i^L}{|N|-1} = x_i^L + \frac{\varepsilon(|N|-1)}{1 - \varepsilon(|N|-1)}.$$

Then, with $\varepsilon > 0$ small enough, the optimal solution $(\tilde{x}_k)_{k \in N}$ is

$$\tilde{x}_i^L = \frac{1}{|N|} - \frac{|N|-1}{|N|} \frac{\varepsilon(|N|-1)}{1 - \varepsilon(|N|-1)} \text{ and, for all } j \neq i, \tilde{x}_j^L = \frac{1 - \tilde{x}_i^L}{|N|-1}.$$

Then, $a_i \cdot \varphi_i(a'_i, a'_{-i}) = \bar{U}_i + a_i^L(\frac{1}{|N|} - \frac{|N|-1}{|N|} \frac{\varepsilon(|N|-1)}{1-\varepsilon(|N|-1)}) > \bar{U}_i$ if ε is small enough, which implies that

$$\max_{a_{-i} \in [\Delta^L]^{N \setminus \{i\}}} a_i \cdot \varphi_i(a'_i, a_{-i}) > \max_{a_{-i} \in [\Delta^L]^{N \setminus \{i\}}} a_i \cdot \varphi_i(a_i, a_{-i}).$$

□

Conclusion. Given the important properties of efficiency, punctual fairness and solidarity that the Egalitarian Walrasian rule satisfies on the domain of linear preferences, and given that it violates *strategy-proofness*, it is important to understand “how manipulable” the EW rule is on this domain. In light of the interpretation of obvious manipulations in [Troyan and Morrill \(2020\)](#), we showed that, for almost any linear preference relation, there exist, for an agent with this preference, manipulations that are easy to identify and enact.

References

- Bogomolnaia, A., Moulin, H., Sandomirskiy, F., Yanovskaya, E., 2017. Dividing goods or bads under additive utilities. ArXiv .
- Cho, W.J., Thomson, W., 2023. Strategy-proofness in private good economies with linear preferences: an impossibility result. *Games and Economic Behavior* 142, 1012–1017.
- Eisenberg, E., Gale, D., 1959. Consensus of subjective probabilities: The pari-mutuel method. *The Annals of Mathematical Statistics* 30, 165–168.
- Kim, H., 2004. Population monotonic rules for fair allocation problems. *Social Choice and Welfare* 23, 59–70.
- Moulin, H., Thomson, W., 1988. Can everyone benefit from growth?: Two difficulties. *Journal of Mathematical Economics* 17, 339–345.
- Psomas, A., Verma, P., 2022. Fair and efficient allocations without obvious manipulations, in: *Proceedings of the 36th International Conference on Neural Information Processing Systems*, Curran Associates Inc., Red Hook, NY, USA.
- Schummer, J., 1996. Strategy-proofness versus efficiency on restricted domains of exchange economies. *Social Choice and Welfare* 14, 47–56.
- Steele, J.M., 2004. *The Cauchy-Schwarz Master Class: An Introduction to the Art of Mathematical Inequalities*. Cambridge University Press.
- Troyan, P., Morrill, T., 2020. Obvious manipulations. *Journal of Economic Theory* 185, 104970.