

Constrained Serial Dictatorships can be Fair¹

Sylvain Bouveret, Hugo Gilbert, Jérôme Lang and Guillaume M  rou  

Abstract

When allocating indivisible items to agents, it is known that the only strategyproof mechanisms that satisfy a set of rather mild conditions are *constrained serial dictatorships*: given a fixed order over agents, at each step the designated agent chooses a given number of items (depending on her position in the sequence). Agents who come earlier in the sequence have a larger choice of items; however, this advantage can be compensated by a higher number of items received by those who come later. How to balance priority in the sequence and number of items received is a nontrivial question. We use a previous model, parameterized by a mapping from ranks to scores, a social welfare functional, and a distribution over preference profiles. For several meaningful choices of parameters, we show that the optimal sequence can be computed exactly in polynomial time or approximated using sampling. Our results hold for several probabilistic models on preference profiles, with an emphasis on the Plackett-Luce model. We conclude with experimental results showing how the optimal sequence is impacted by various parameters.

1 Introduction

In an ideal world, a mechanism for dividing a set of indivisible goods (or items, we use both terms interchangeably) should be at the same time efficient, fair, and insensitive to strategic behaviour. Now, strategyproofness is a very strong requirement that severely limits the choice of mechanisms. The question we address in this paper is, *how can we design strategyproof mechanisms while retaining an acceptable level of fairness and/or efficiency?*

It is known that *under mild conditions, the only strategyproof mechanisms are within the family of serial dictatorships* (although the landscape is less dramatic when there are only two agents, see our [related work](#) section). A standard serial dictatorship is defined by a permutation of the set of agents; at each step, the designated agent chooses all the items she likes from those that are still available. A *constrained serial dictatorship* (CSD), also called quota serial dictatorship, is similar except that at each step, the designated agent chooses a predefined number of items. Notice that constrained serial dictatorship is actually a particular case of the *picking sequence* protocol [e.g 11], restricted to *non-interleaving sequences* where every agent picks all her entitled items in a row.

(Constrained or unconstrained) serial dictatorships are strategyproof and elicitation-free: they do not require to know the agents preferences, which are only revealed through their picking choices. This is a major property, as in many contexts, it is not realistic to hope eliciting all the agents preferences, because it would be too big a communication burden, and also for privacy reasons. However, are they acceptable on efficiency and fairness grounds? Unconstrained serial dictatorships are clearly not: if the first agent likes all items then she will pick them all. Constrained serial dictatorships do better, at the price of the loss of Pareto-efficiency; but still, agents appearing early in the sequence have a much larger choice than those appearing late. This is patent in the case where there are as many items as agents, each agent being entitled to only one item, CSDs cannot do better than this: the first agent will get her preferred item, and the last agent will have no choice and might receive her least preferred item.

However, when there are more items than agents, and agents can receive several items, things become better, because the advantage towards agents who come early in the sequence can be compensated by a higher number of items received by those who come later. Suppose, as a simple example, that three

¹This paper has been accepted to IJCAI 2025

items have to be assigned to two agents, A(nn) and B(ob). Assuming that Ann picks first, there are three CSDs: (A:3,B:0) (Ann picks all items), (A:2,B:1) (Ann picks two, Bob one), and (A:1, B:2) (Ann picks one, Bob two). It is intuitively clear that (A:1,B:2) is optimal, but how can optimality be defined? With four items, things are less clear: (A:4,B:0) and (A:3,B:1) are clearly less desirable than (A:2,B:2) and (A:1,B:3), but which of these two should we choose? And what if we have five agents and seventeen items?

To sum up: strategyproofness leaves us almost no choice but (constrained) serial dictatorship; some are intuitively better than others. What remains to be done is to *define formal optimality criteria for choosing between CSDs*, and to *compute optimal ones*. Our paper addresses these questions.

A way of answering the first question has been suggested by Bouveret and Lang [11], and further studied by Kalinowski et al. [30] (in the more general context of *picking sequences*).² A standard way of estimating the efficiency and fairness of a CSD consists in evaluating the *expected social welfare*, according to some social welfare functional [22] – egalitarian, Nash, utilitarian³ – of the allocation resulting from the application of the serial dictatorship. Because the agents’ values for items are not known, Bouveret and Lang [11] estimate them from the ranks of items in an agent’s preference relation: for each agent i , the value of item ranked in position j is a fixed value s_j , independent from i . To estimate the expected social welfare, in addition to the non-increasing scoring vector $(s_1, \dots, s_m)$ one also needs to assume a probability distribution over ordinal preference profiles. These preferences can be drawn following different models as impartial culture, or more generally the Mallows [36], or Plackett-Luce models [35, 39].

These three components (scoring vector, probability over profiles, social welfare functional) allow to associate an expected social welfare with any CSD. We define optimal CSDs this way, for various scoring vectors, three social welfare functionals, and various probability distributions.

For egalitarian social welfare, we provide a simple algorithm which returns an optimal CSD given that one can compute the expected utility obtained by an agent when a CSD is used. This algorithm makes it possible to compute an optimal (respectively, close to optimal) CSD when this expected utility is polynomial-time computable (respectively, can be approximately evaluated, e.g., by sampling). We also provide a dynamic programming algorithm that computes an optimal CSD for utilitarian, Nash or egalitarian social welfare under a specific condition, which is met when preferences are fully correlated, or when they are fully independent and follow the impartial culture or more generally the Plackett-Luce model.

Sections 2 and 3 discuss related work and present our model. Section 4 presents our algorithms for computing an optimal CSD. These algorithms assume the existence of an oracle which can compute or estimate the expected utility of a picker given a CSD. Section 5 designs such oracles under various model assumptions. Section 6 gives results for small values of n , and depicts and comments on the evolution of the optimal sequences when all criteria except one are fixed.

2 Related Work

Strategyproof allocation of indivisible goods Various characterization theorems state that, under mild additional conditions, strategyproof allocation mechanisms all have a serial dictatorship flavour: with strict preferences over subsets, only serial dictatorships are strategyproof, neutral, and nonbossy [43], whereas only sequential dictatorships (a generalization of serial dictatorship where the identity of the agent picking in position k depends on the items assigned to the agents in positions 1 to $k - 1$) are strategyproof, Pareto-efficient, and nonbossy [41]. If preferences are quantity-monotonic (a bundle of

²Picking sequences are more general as agents don’t necessarily pick their items in a row; for instance, the sequence where Ann picks one item, Bob two, and Ann picks the last remaining item, is not a CSD. Round Robin (perfect alternation) is another example. CSDs coincide with *non-interleaving* picking sequences.

³If our main objective is *fairness*, utilitarian social welfare may not fit well. We will see further that it is the case indeed.

larger cardinality is always preferred to one of lower cardinality) then a mechanism is strategyproof, nonbossy, Pareto-efficient and neutral if and only if it is a CSD (also called a quota serial dictatorship) [38]. Similar characterizations hold replacing quantity-monotonic by lexicographic preferences [28, 29]. With standard monotonicity, only quasi-dictatorships remain, where only the first agent in the sequence is allowed to pick more than one item [38]. Variants of these characterizations have been established by Ehlers and Klaus [23], Bogomolnaia et al. [9] and Hatfield [26]. Ignoring Pareto-efficiency or neutrality opens the door to more complex strategyproof mechanisms; a full characterization in the two-agent case is given by Amanatidis et al. [2]. Amanatidis et al. [1] show that the CSD where all agents except the last one pick only one item is a $1/\lfloor \frac{n-m+2}{2} \rfloor$ -approximation to maxmin fair share. Weakening strategyproofness into non obvious manipulability opens the door for more possibilities [40].

Nguyen et al. [37] show that when agents have preferences over sets of items defined from preferences over single items by an extension principle, some scoring rules are strategyproof for some extension principles. Allowing randomized mechanisms offers more possibilities, but not much [16, 25, 28, 33].

CSDs are also considered in chore allocation [5].

Picking sequences Sequential allocation of indivisible goods, also known as picking sequences, originates from Kohler and Chandrasekaran [32], with a game-theoretic study of the alternating sequence for two agents. Still for two agents, Brams and Taylor [15] consider other particular sequences. Bouveret and Lang [11] define a more general class of sequences, for any number of agents, and argue that sequences can be compared with respect to their expected social welfare, using a scoring vector and a prior distribution over profiles. Kalinowski et al. [30] show that computing the expected utility of a sequence is polynomial under full independence, and that strict alternation is optimal for two agents, utilitarian social welfare and Borda scoring. The manipulation of picking sequences is studied by Bouveret and Lang [12], Tominaga et al. [44] and Aziz et al. [4]. Flammini and Gilbert [24] and Xiao and Ling [45] study the parameterized complexity of computing an optimal manipulation. Game-theoretic aspects of picking sequences are addressed by Kalinowski et al. [31]. Chakraborty et al. [19] study picking sequences for agents with different entitlements. While all these works are oblivious to agent identities, Caragiannis and Rathi [17] try to find an approximately optimum order of agents in a serial dictatorship with a limited number of queries.

Maximizing social welfare in allocation of indivisible goods A classic way of guaranteeing a level of fairness and/or efficiency consists in finding an allocation *maximizing social welfare*, under the assumption that the input contains, for each agent, her utility function over all bundles of goods (usually assumed additive). Egalitarian social welfare places fairness above all, utilitarian social welfare only cares about efficiency, and Nash social welfare is considered as a sweet spot in-between. See [3, 6, 13, 34] for surveys. These mechanisms are not strategyproof.

3 Preliminaries: The Model

Given $n \in \mathbb{N}^*$, we use $[n]$ to denote $\{1, \dots, n\}$ and $[n]_0$ to denote $\{0, 1, \dots, n\}$. Bold symbols represent vectors.

Let $\mathcal{A} = \{a_1, \dots, a_n\}$ be a set of n agents with a_i the i^{th} agent to intervene in the allocation process and $\mathcal{G} = \{g_1, \dots, g_m\}$ a set of m goods. A preference profile $\mathbf{P} = (\succ_{a_1}, \dots, \succ_{a_n})$ describes the preferences of the agents: $\succ_a$ is a ranking that specifies the preferences of agent a over the goods in $\mathcal{G}$. We denote by $\text{rk}_{\mathbf{P}}^a(g)$, the rank of item g in the ranking of a , given profile $\mathbf{P}$. *The preference profile is hidden, and therefore not part of the input*: we will assume that rankings are drawn independently according to some probabilistic model, that we denote by Ψ .

Two well-known probabilistic models are the *Mallows* and *Plackett-Luce* models [35, 36, 39]:

- The *Mallows model* is parameterized by a dispersion parameter $\phi \in [0, 1]$ and a ranking μ . We denote this model by $\text{M11}_{\mu, \phi}$. In this model, the probability of a ranking r is proportional to $\phi^{d_{\text{KT}}(r, \mu)}$, with $d_{\text{KT}}(r, \mu)$, the Kendall-Tau distance between rankings r and μ .
- The *Plackett-Luce (PL) model* is parameterized by a value vector $\nu = (\nu_1, \dots, \nu_m)$. Intuitively, $\nu_i > 0$ represents the social value of good g_i . In this model, which we denote by PL_{ν} , the probability of a ranking $r = g_{i_1} \succ g_{i_2} \succ \dots \succ g_{i_m}$ is:

$$\prod_{j=1}^m \frac{\nu_{i_j}}{\sum_{l=j}^m \nu_{i_l}}.$$

The Plackett-Luce model has proven particularly good for learning a preference relation over a set of items (a.k.a. label ranking) [20] so it fits particularly well here.

These models generalize the two following sub-cases:

- *Impartial Culture*, denoted by IC, in which each preference ranking is drawn u.a.r. from the set of all possible rankings. Impartial culture is obtained when $\phi = 1$ for the Mallows model and when all values in ν are equal for the Plackett-Luce model.
- The *Full Correlation* case, denoted by FC stipulates that all agents have exactly the same preference ranking. Full correlation is obtained when $\phi = 0$ for the Mallows model (and also as the limit of Plackett-Luce models $\nu^M = (M^{m-1}, \dots, M, 1)$ when $M \rightarrow \infty$).

In the sequel, we obtain different results for $\Psi \in \{\text{FC}, \text{IC}, \text{M11}_{\mu, \phi}, \text{PL}_{\nu}\}$.

The items are allocated to the different agents according to a CSD: given a vector $\mathbf{k} = (k_1, \dots, k_n)$ of n non-negative integers, agent a_1 will first pick k_1 goods, then a_2 will pick k_2 goods within the remaining ones, and so on until a_n picks k_n items. In most cases, we will consider complete CSDs, in the sense that $\sum_{i=1}^n k_i = m$. However, we may also consider incomplete CSDs such that $\sum_{i=1}^n k_i < m$. We assume that agents behave greedily by choosing their preferred goods within the remaining ones. This sequential process leads to an allocation that we denote by $\pi_{\mathbf{P}}^{\mathbf{k}}$. More formally, $\pi_{\mathbf{P}}^{\mathbf{k}}$ is a function such that $\pi_{\mathbf{P}}^{\mathbf{k}}(a)$ is the set of goods that agent a has obtained at the end of the sequential allocation process, given preference profile $\mathbf{P}$ and vector $\mathbf{k}$.

The utility of an agent for obtaining an item i will be derived using a scoring vector. Stated otherwise, there is a vector $\mathbf{s} = (s_1, \dots, s_m) \in \mathbb{Q}^{+m}$ such that $s_i \geq s_{i+1}$ for all $i \in [m-1]$. The value received by an agent for obtaining her j^{th} preferred item is s_j . Different scoring vectors can be considered. An important example is the *Borda* scoring vector, where $s_i = m - i + 1$. Using scores as a proxy for utilities is classic in social choice: this is exactly how positional scoring voting rules (e.g., the Borda rule) are defined, and they are also used in fair division settings [7, 14, 21].

We denote by $U_{\mathbf{P}}^{\mathbf{k}}(a) = \sum_{g \in \pi_{\mathbf{P}}^{\mathbf{k}}(a)} s_{\text{rk}_{\mathbf{P}}^a(g)}$ the utility obtained by a when receiving $\pi_{\mathbf{P}}^{\mathbf{k}}(a)$ and by $EU_{\Psi}^{\mathbf{k}}(a) = \mathbb{E}_{\mathbf{P} \sim \Psi}[U_{\mathbf{P}}^{\mathbf{k}}(a)]$ her expected utility given model Ψ . This assumes that agents have *additive* preferences, which is very common in fair division. The utilitarian social welfare (USW) $SW_{\Psi}^U(\mathbf{k})$, egalitarian social welfare (ESW) $SW_{\Psi}^E(\mathbf{k})$, and Nash social welfare (NSW) $SW_{\Psi}^N(\mathbf{k})$ are then defined by:

$$\begin{aligned} SW_{\Psi}^U(\mathbf{k}) &= \sum_{a \in \mathcal{A}} EU_{\Psi}^{\mathbf{k}}(a), & SW_{\Psi}^E(\mathbf{k}) &= \min_{a \in \mathcal{A}} EU_{\Psi}^{\mathbf{k}}(a), \\ SW_{\Psi}^N(\mathbf{k}) &= \prod_{a \in \mathcal{A}} EU_{\Psi}^{\mathbf{k}}(a). \end{aligned}$$

Note that our social welfare notions are meant *ex ante*, i.e., we define them on the expected utility values of the agents. This is different from the notion of *ex post* social welfare which considers the

utility of the agents once the profile $\mathbf{P}$ issued from Ψ is determined.

Our objective is to study the following class of optimization problems $\text{OptSD-}\Psi\text{-}x$ with $x \in \{U, E, N\}$.

$\text{OptSD-}\Psi\text{-}x$
Input: A number n of agents, a number m of goods, and a scoring vector $\mathbf{s}$. Find: A vector $\mathbf{k} = (k_1, \dots, k_n)$ of n non-negative integers with $\sum_{i=1}^n k_i = m$ maximizing $SW_{\Psi}^x(\mathbf{k})$.

The following easy observation will be useful:

Observation 1. For given n and m , the number of vectors $\mathbf{k} = (k_1, \dots, k_n)$ such that $\sum_{i=1}^n k_i = m$ equals $\binom{n+m-1}{n-1}$.

From this observation, we can deduce that the number of potential vectors is lower-bounded by $\frac{m^{n-1}}{(n-1)!}$. This number does not take into account a natural further assumption that the optimal sequence is *non-decreasing*, that is, that $k_1 \leq k_2 \leq \dots \leq k_n$. We will see further that this assumption holds for ESW (under a mild condition), *but not* for USW. When the assumption holds, we can restrict the search to non-decreasing vectors; their number is the number of integer partitions m into n numbers; it is still exponentially large, but no closed form expression is known.

4 Computing an Optimal CSD

We now investigate the problem $\text{OptSD-}\Psi\text{-}x$ with $x \in \{U, E, N\}$. All algorithms in this Section assume access to an oracle algorithm $\mathcal{T}_{\Psi}(\mathbf{k}, i)$ computing $EU_{\Psi}^{\mathbf{k}}(a_i)$ in time $K(n, m, \mathbf{s})$. The computation of expected utilities of agents for various models will be addressed in Section 5.

We start by a positive result for Egalitarian Social Welfare: the optimal CSD can be computed by the greedy-like Algorithm 1. $\text{Completion}(\mathbf{k})$ denotes, for any partial CSD $\mathbf{k}$, the complete CSD such that $\text{Completion}(\mathbf{k})_i = k_i$ for $i \in [n-1]$ and $\text{Completion}(\mathbf{k})_n = m - \sum_{i \in [n-1]} k_i$. In informal terms, $\mathbf{k}$ is completed by giving all remaining goods to the last agent.

Algorithm 1 GreedyESW

Require: the number of agents n , the number of goods m , the scoring vector $\mathbf{s}$, the oracle algorithm

```

 $\mathcal{T}_{\Psi}$ 
1:  $\mathbf{k} \leftarrow (0, \dots, 0)$  # empty CSD
2:  $\text{max\_k}, \text{max\_esw} \leftarrow \mathbf{k}, 0$ 
3: for  $t = 1$  to  $m$  do
4:    $i \leftarrow \text{any value in } \arg \min_{i \in [n]} EU_{\Psi}^{\mathbf{k}}(a_i)$ 
5:    $k_i \leftarrow k_i + 1$ 
6:   if  $SW_{\Psi}^E(\mathbf{k}) > \text{max\_esw}$  then
7:      $\text{max\_k}, \text{max\_esw} \leftarrow \mathbf{k}, SW_{\Psi}^E(\mathbf{k})$ 
8:   end if
9: end for
10: return  $\text{Completion}(\text{max\_k})$ 

```

At line 1, we start with an empty CSD, that we will modify in a greedy fashion. In the **for** loop (lines 3-9), we identify an agent with minimal expected utility (line 4) and increment the number of goods

that she gets (line 5). The CSD that is returned is not necessarily this CSD $\mathbf{k}$. During the algorithm, we keep in variables max_esw and max_k , the maximum ESW found so far and the corresponding (partial) CSD. The algorithm returns max_k completed by giving all remaining goods to the last agent (line 10). The completion step is not really necessary (the partial sequence obtained at line 9 already has maximum expected egalitarian social welfare); its role is to ensure that no good is left unallocated. The reason why one needs the test at line 6 is that letting the currently least happy agent pick one more good may decrease the ESW, as can be seen on the following example.

Example 1. Let $n = 2$, $m = 5$, $s = (50, 10, 4, 2, 1)$, and $\Psi = \text{IC}$. We show below the partial CSDs obtained in each iteration t together with the expected utilities of both agents (they can be computed easily, as we will see in Section 5) and the values of i and max_esw .

t	$\mathbf{k}$	max_k	$EU_{\Psi}^{\mathbf{k}}(a_1)$	$EU_{\Psi}^{\mathbf{k}}(a_2)$	i	max_esw
1	(0, 0)	(0, 0)	0	0	1	0
2	(1, 0)	(0, 0)	50	0	2	0
3	(1, 1)	(1, 1)	50	42	2	42
4	(1, 2)	(1, 2)	50	49.6	2	49.6
5	(1, 3)	(1, 3)	50	52.4	1	50
6	(2, 3)	(1, 3)	60	40.2	1	50

At iteration 5, the least happy agent is a_1 ; however, letting a_1 pick one more good, that is, $\mathbf{k} = (2, 3)$ gives $EU_{\Psi}^{\mathbf{k}}(a_1) = 60$ and $EU_{\Psi}^{\mathbf{k}}(a_2) = 40.2$ (iteration 6), decreasing the currently optimal expected ESW. Therefore, max_k is not replaced by $\mathbf{k} = (2, 3)$ at line 6 of the algorithm. The algorithm returns $\text{Completion}(\text{max_k}) = (1, 4)$ (with expected utilities 50 and 53.6) with the remaining good given to a_2 .

Proposition 1. Algorithm 1 returns a CSD $\mathbf{k}$ maximizing $SW_{\Psi}^E(\mathbf{k})$, solving problem $\text{OptSD-}\Psi\text{-E}$, in time $O(nmK(n, m, s))$.

The proof is based on the following lemma:

Lemma 1. Let $\hat{\mathbf{k}}$ be a CSD. Let max_esw^t , $\mathbf{k}^t$ and i^t denote max_esw , $\mathbf{k}$ and i after line 4 of iteration t of the **for** loop in Algorithm 1. For all t , a necessary condition for $SW_{\Psi}^E(\hat{\mathbf{k}}) > \text{max_esw}^t$ is that $\hat{k}_j \geq k_j^t$ for all $j \in [n]$, and $\hat{k}_{i^t} > k_{i^t}^t$.

Proof. By induction. At iteration 0, the claim is obvious. Assume that the claim holds for iteration t , and let $\hat{\mathbf{k}}$ be a CSD such that $SW_{\Psi}^E(\hat{\mathbf{k}}) > \text{max_esw}^{t+1}$. Then obviously $SW_{\Psi}^E(\hat{\mathbf{k}}) > \text{max_esw}^t$ as $\text{max_esw}^{t+1} \geq \text{max_esw}^t$. Because the condition holds for iteration t and by construction of $\mathbf{k}^{t+1}$ we have that $\hat{k}_j \geq k_j^{t+1}$ for all $j \in [n]$. Now suppose that $\hat{k}_{i^{t+1}} = k_{i^{t+1}}^{t+1}$. In that case, $SW_{\Psi}^E(\hat{\mathbf{k}}) \leq EU^{\mathbf{k}^{t+1}}(a_{i^{t+1}}) \leq \text{max_esw}^{t+1}$, a contradiction with the induction hypothesis. The first inequality is due to the fact that $a_{i^{t+1}}$ will get the same number of goods in $\hat{\mathbf{k}}$ and $\mathbf{k}^{t+1}$ while the agents picking before her will get at least as many goods in $\hat{\mathbf{k}}$ than in $\mathbf{k}^{t+1}$. The second inequality is due to the definition of i^{t+1} . $\square$

Proof of Proposition 1. Suppose that there exists a CSD $\hat{\mathbf{k}}$ such that $SW_{\Psi}^E(\hat{\mathbf{k}}) > \text{max_esw}$. Lemma 1 applied at iteration $t = m$ implies that each agent receives more objects with $\hat{\mathbf{k}}$ than with the greedily constructed complete CSD $\mathbf{k}$ obtained at the end of the **for** loop. As they both have m objects to allocate, they must be equal. This is a contradiction of the hypothesis as $\text{max_esw} \geq SW_{\Psi}^E(\mathbf{k})$. $\square$

We now go beyond egalitarian social welfare. For utilitarian and Nash social welfare, we do not know of an efficient algorithm which would work for any distribution. A general approach could be to sample a large but hopefully reasonable number of preference profiles from Ψ and find a CSD with maximal

social welfare considering the average utility of each agent. Yet, we prove in Appendix B that such an approach leads to an NP-hard problem for USW.

However, provided the distribution satisfies a natural condition, a CSD maximizing utilitarian and Nash social welfare can be computed by dynamic programming. This condition on Ψ states that $EU_{\Psi}^{\mathbf{k}}(a)$ only depends on the number of items picked by a , and the number of items that have been picked before a , but not on the number of agents who have picked before and how many items they have picked each.

Definition 1. A distribution Ψ satisfies **prefix independence** if for any sequence $\mathbf{k}$ and $i \in [n]$, if a is the i^{th} picker in $\mathbf{k}$, then $EU_{\Psi}^{\mathbf{k}}(a)$ only depends on (1) $\kappa = k_i$, the number of goods that she picks, and (2) $\tau = \sum_{j=1}^{i-1} k_j$, the number of goods that have been picked before she starts picking.

Under prefix independence, the utility that agent a gets when picking κ goods while τ have already been picked, $\text{eu}(\kappa, \tau)$, is well-defined, and is exactly equal to $EU_{\Psi}^{\mathbf{k}}(a)$ when a is the i^{th} picker $\kappa = k_i$ and $\tau = \sum_{j=1}^{i-1} k_j$.

For pedagogical purposes, let us first focus on maximising USW. When prefix independence is met, one can use the following dynamic programming equations:

$$\begin{aligned} F(i, \tau) &= \max_{\kappa \in [m-\tau]_0} (\text{eu}(\kappa, \tau) + F(i+1, \tau + \kappa)), \\ &\quad \forall i, \tau \in [n-1] \times [m]_0, \\ F(n, \tau) &= \text{eu}(m - \tau, \tau), \forall \tau \in [m]_0, \end{aligned} \tag{1}$$

where $F(i, \tau)$ corresponds to the maximum USW that can be obtained by agents $\{a_i, a_{i+1}, \dots, a_n\}$ in the situation in which τ goods have already been allocated and we allocate the $m - \tau$ remaining goods to them. Of course the optimal value is given by $F(1, 0)$.

The other problems can be solved similarly. For problem OptSD- Ψ - E (resp. OptSD- Ψ - N), one should adapt Equation 1 by replacing the sum operation between $\text{eu}(\kappa, \tau)$ and $F(i+1, \tau + \kappa)$ by a min (resp. multiplication) operation.

Proposition 2. If Ψ satisfies prefix independence, problems OptSD- Ψ - U , OptSD- Ψ - E and OptSD- Ψ - N can be solved in $O(nm^2K(n, m, \mathbf{s}))$ time.

We conclude by giving a structural property satisfied by an optimal CSD for ESW when prefix independence holds. We will see that such property does not necessarily hold for USW (see Appendix B and Section 6).

Proposition 3. Under prefix independence, there exists an optimal solution to OptSD- Ψ - E which is non-decreasing, i.e., in which the earlier an agent picks, the less goods she gets.

This property is not true with utilitarian social welfare. Take $n = 4$, $m = 10$, the IC model and the Borda scoring vector. Then $\mathbf{k} = (3, 3, 2, 2)$ is optimal for utilitarian social welfare, $\mathbf{k} = (2, 2, 3, 3)$ for Nash social welfare, and $\mathbf{k} = (2, 2, 2, 4)$ for egalitarian social welfare. Utilitarianism gives the first two agents more goods than the last two; the first agents have the cake and eat it, as they pick more goods and have more choice. The intuition is that late agents may end up with items of low utility; egalitarianism compensates by giving them more, while utilitarianism avoids this "waste" by favoring earlier agents, who are more likely to secure high-utility items.

5 Computing the Expected Utility of an Agent

In this section, we address the computation of $EU_{\Psi}^{\mathbf{k}}(a)$. Prefix independence again plays a crucial role: when it is satisfied, $EU_{\Psi}^{\mathbf{k}}(a)$ only depends on the number of items picked by a , and the number of items that have been picked before a , but not on the number of agents who have picked before and how many items they have picked each. We first investigate which of our different probabilistic models satisfy it.

Proposition 4. $\Psi \in \{\text{FC}, \text{IC}\}$ satisfy prefix independence.

Proof. Consider a situation where an agent starts picking while τ goods have previously been picked. When $\Psi = \text{FC}$ or $\Psi = \text{IC}$, the probability distribution on the set S of goods that have previously been picked only depends on τ : for $\Psi = \text{FC}$, this probability distribution assigns probability 1 to the set composed of the τ (unanimously) most preferred goods; for $\Psi = \text{IC}$, this probability distribution assigns equal probability to all sets of size τ and 0 to others. Note that, given the set S , the utility that the agents get is then determined by the number of goods she picks. $\square$

More interestingly, the PL_ν model, which generalizes FC and IC, also satisfies prefix independence.

Proposition 5. $\Psi = \text{PL}_\nu$ satisfies prefix independence.

To reason on the Plackett-Luce model, one can use the *vase model metaphor* [42] Consider a vase filled with m types of balls, the proportion of balls of type j being $f(j) = \frac{\nu_j}{\sum_{l=1}^m \nu_l}$. The ranking is then generated by the following sequential process. At each stage, a ball is taken from the vase such that a ball of type j is chosen with probability $f(j)$. If the ball is of a different type than the ones previously picked, it yields the next good in the ranking. In either case, the ball is put back in the vase and the process continues. Using this metaphor, one can prove the following lemma (whose formal proof is postponed to Appendix C).

Lemma 2. Let $I = (i_1, \dots, i_q)$ be a sequence of q different indices in $[m]$. Consider the following two cases:

- i) Agent a_1 picks q goods;
- ii) Agent a_1 picks q_1 goods and agent a_2 picks q_2 goods with $q_1 + q_2 = q$.

For the PL model, the probability that for all $t \in [q]$, g_{i_t} is picked at timestep t is the same in cases i and ii.

Proof of Proposition 5. We recall that the preference rankings of the agents are drawn independently from PL_ν . Using Lemma 2 and a simple induction argument, we get that the probability of a specific sequence of q consecutive picks is the same regardless of whether they were picked by one, two or more agents. This entails that the probability distribution on the set S of goods that have been picked after τ timesteps only depends on the value of τ . Hence, the expected utility that an agent gets when choosing κ goods once τ have been picked only depends on the values of κ and τ . $\square$

Unfortunately, things are different for the Mallows model:

Proposition 6. There exists $\phi \in (0, 1)$ and a ranking μ such that $\Psi = \text{Mll}_{\phi, \mu}$ does not satisfy prefix independence.

This holds even for 3 agents and 3 goods. See Appendix C for the proof.

Computation of $EU_\Psi^k(a)$ Under prefix independence, we show how to compute $\text{eu}(\kappa, \tau)$ efficiently, starting by FC.

Proposition 7. If $\Psi = \text{FC}$, $\text{eu}(\kappa, \tau) = \sum_{i=\tau+1}^{\tau+\kappa} s_i$. All values $\text{eu}(\kappa, \tau)$ can be computed in time $O(m^2)$ with the recursive formula $\text{eu}(\kappa, \tau) = \text{eu}(\kappa - 1, \tau) + s_{\kappa+\tau}$.

We then turn to $\Psi = \text{IC}$, and show that the values $\text{eu}(\kappa, \tau)$ can be computed using a recursive formula. Let $T(j, \kappa, \tau)$ denote the utility that an agent can get if she can pick κ goods within the ones of rank in $\{j, \dots, m\}$, given that τ of these goods have been picked by preceding agents. Then, it is clear that we have:

$$\text{eu}(\kappa, \tau) = T(1, \kappa, \tau), \forall \kappa, \tau \in [m]_0 \times [m - \kappa]_0$$

The key point is that there is a probability $1 - \frac{\tau}{m-j+1}$ that the good of rank j is free and in this case the agent will pick this good, and a probability of $\frac{\tau}{m-j+1}$ that this good is one of the τ goods that have previously been picked. In both cases, we move to goods of rank in $\{j+1, \dots, m\}$. In the first case, we decrease κ by one as the agent has picked a good. In the second case, we decrease τ by 1 as we have identified one of the goods picked within the ones of rank j to m . Hence, $\text{eu}(\kappa, \tau)$ can be computed by the following formula:

$$\begin{aligned} T(j, \kappa, \tau) &= \left(1 - \frac{\tau}{m-j+1}\right)(s_j + T(j+1, \kappa-1, \tau)) \\ &\quad + \frac{\tau}{m-j+1}T(j+1, \kappa, \tau-1), \\ \forall j, \kappa, \tau &\in [m-1] \times [m-j+1]_0 \times [m-j-\kappa+1], \end{aligned} \quad (2)$$

with the following base cases:

$$\begin{aligned} T(j, 0, \tau) &= 0, \forall j, \tau \in [m] \times [m-j+1]_0 \\ T(j, \kappa, 0) &= \sum_{j \leq i < j+\kappa} s_i, \forall j, \kappa \in [m] \times [m-j+1]_0. \end{aligned}$$

By computing all values $T(j, \kappa, \tau)$ in $O(m^3)$ operations, we obtain the following result.

Proposition 8. *If $\Psi = \text{IC}$, then all values $\text{eu}(\kappa, \tau)$ can be computed in time $O(m^3)$ by using Equation 2.*

Propositions 2, 7, and 8 imply that $\text{OptSD-}\Psi\text{-}x$ for $x \in \{U, E, N\}$ can be solved in polynomial time for $\Psi = \text{FC}$ and $\Psi = \text{IC}$, in $O(nm^2)$ for $\Psi = \text{FC}$ and $O(m^2 \max(n, m))$ for $\Psi = \text{IC}$, by precomputing all values $\text{eu}(\kappa, \tau)$ before running the dynamic programming algorithm.

For $\Psi \notin \{\text{IC}, \text{FC}\}$, one can still use GreedyESW and the dynamic programming algorithm with values $EU_{\Psi}^k(a)$ approximated by sampling, providing close-to optimal CSDs: the returned CSD is optimal with expected utility values replaced by their approximate values.⁴

For the general PL_{ν} model beyond FC and IC, we do not know whether values $\text{eu}(\kappa, \tau)$ can be computed exactly in polynomial time; however, they can be efficiently approximated by sampling preference profiles from Ψ and averaging the utility values obtained on the samples, with approximation guarantees from Hoeffding's (1963) inequality.

To present this guarantee, let $u_{\kappa, \tau}(\mathbf{P}, \mathbf{s})$ denote the utility value obtained by the second picker when she picks her κ preferred (available) goods, while the first picker has picked her τ preferred ones, given the preference profile $\mathbf{P}$.

Proposition 9. *Let $\epsilon > 0$ and $\delta \in (0, 1)$ two fixed values, and Υ an upper bound on values $\text{eu}(\kappa, \tau)$ (e.g., $\sum_{i=1}^m s_i$).*

Let $\widetilde{\text{eu}}_{\kappa, \tau}$ be the value computed by averaging the values $u_{\kappa, \tau}(\mathbf{P}_i, \mathbf{s})$ over N preference profiles $\mathbf{P}_i$ sampled independently from Ψ . If $N \geq (\Upsilon^2 \ln(2m^2/\delta))/2\epsilon^2$, then it holds with probability $1 - \delta$ that:

$$|\text{eu}(\kappa, \tau) - \widetilde{\text{eu}}_{\kappa, \tau}| \leq \epsilon, \forall \kappa, \tau \in [m] \times [m - \kappa].$$

⁴Some mild monotonicity conditions are required on the approximated $EU_{\Psi}^k(a)$ values for the validity of Algorithm 1.

Moreover, we show that these utility values can be computed exactly in time FPT (Fixed-Parameter Tractable) with respect to parameter m and XP (slice-wise polynomial) with respect to ρ , where ρ is the number of distinct values in ν . This seems particularly appealing as goods may often be partitioned in categories. When $\rho = 1$, all goods are in the same category and we obtain the IC model; when ρ equals 2 or 3 we obtain categories {high value, low value} or {high value, medium value, low value}.

Proposition 10. *If $\Psi = \text{PL}_\nu$, then all values $\text{eu}(\kappa, \tau)$ can be computed in time $O(4^m \text{Poly}(m))$.*

Proposition 11. *If $\Psi = \text{PL}_\nu$, then all values $\text{eu}(\kappa, \tau)$ can be computed in time $O(m^{2\rho} \text{Poly}(m))$.*

6 Numerical Tests

We performed several experiments to explore the properties of the CSDs obtained by maximizing either USW, NSW or ESW. More precisely, we explored the impact of increasing one of the parameters, all other parameters being fixed.

Impact of the number of goods Figure 1 displays the proportion of utility (left-hand side) and goods (right-hand side) obtained for $n = 5$ and increasing the number of goods m from 5 to 300 in steps of 5. To generate both figures, the IC model and the Borda scoring vector were used and we optimized either USW, ESW or NSW.

Several comments can be made. First, as expected, in the egalitarian case (middle of Figure 1), we observe that as m increases, the distribution of utility received by each agent converges towards equal share.⁵ In order to achieve this, the agents who arrive later in the sequence receive more items.

Second, with Borda and utilitarianism, the first agent in the sequence may pick more items than others (plots on top of Figure 1). More generally, on this plot, the utility of an agent seems to decrease with the position in the sequence.

Finally, for the Borda scoring vector, the Nash social welfare objective seems to yield somewhat intermediate results between the utilitarian and the egalitarian ones, but seems to be closer to the latter.

Impact of correlation We explore the impact of correlation, through the parameters ϕ and ν of models PL_ν and $\text{M11}_{\phi,\mu}$. We use the Borda scoring vector and maximize ESW. To run Algorithm 1, we approximate the expected utility values of the agents by sampling 10000 preference profiles from PL_ν and from $\text{M11}_{\phi,\mu}$ with the PrefSampling library [8]. Figure 2 displays the utility value (plots at the bottom) and the number of items (top) received by each of 5 agents with $m = 70$ goods, for models PL_ν (right) and $\text{M11}_{\phi,\mu}$ (left), for ESW and Borda scoring vector⁶. In the former model, we use $\nu^x = (x^m, x^{m-1}, \dots, x^1)$ and decrease x from 1.5 (which already yields very correlated preference profiles in similar to FC) to 1 (IC) in steps of 0.01. In the latter model, we increase ϕ from 0 (FC) to 1 (IC) in steps of 0.02.

Several comments are in order. First, as can be seen in Figure 2, the utility values of all agents (and hence their sum) increase when x decreases or ϕ increases. Indeed, as we come closer to IC, the preferences of the agents become more different, allowing some agents to receive some of their preferred items even if they pick late in the allocation process.

Second, the number of goods received by the first agents in the CSD increases while it decreases for the last ones. Indeed, as these latter agents can receive more preferred goods, the CSD needs less to compensate by giving them a high number of goods (recall that we optimize ESW).

⁵This observation is proven formally in Appendix D.

⁶This choice was motivated by the fact that Borda is the most standard scoring vector and the ESW naturally conveys fairness.

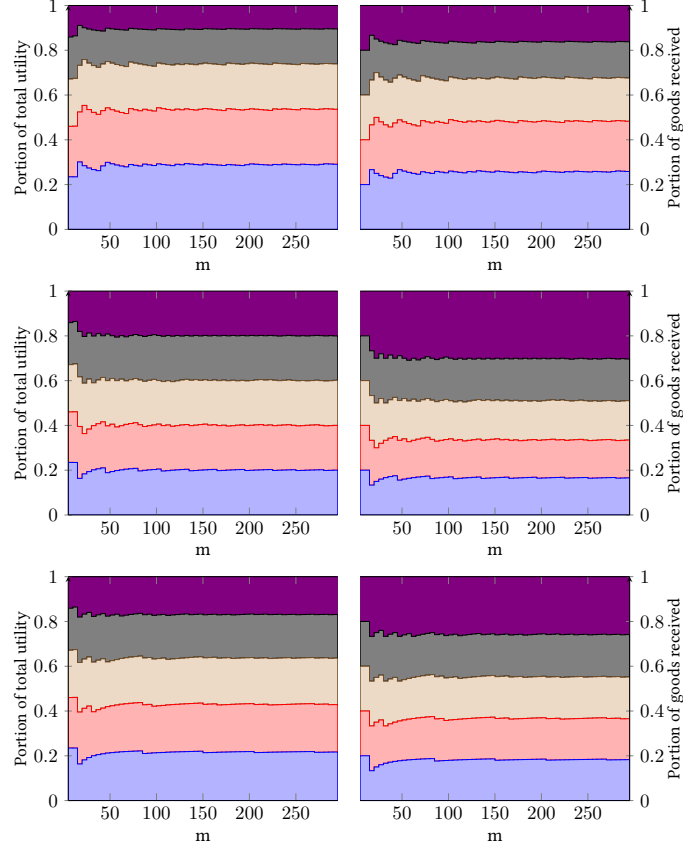


Figure 1: Portion of total utility (plots on the left) and of goods (right) received by each of 5 agents with m increasing from 5 to 300 in steps of 5. Maximizing USW (plots at the top), NSW (bottom), or ESW (middle), using Borda scoring vector and IC. The 1st picker corresponds to the color *blue* (at the bottom of each plot) while the 5th and last agent to pick corresponds to the color *purple* (at the top of each plot). Moreover, note that the values plotted are in fact cumulative values.

Third, we notice that both models PL_{ν} and $M11_{\phi,\mu}$ yield very similar plots as we decrease the level of correlation.

Code and an interactive demo are available at <https://github.com/GuillaumeMeroue/CSD-can-be-Fair> and <https://guillaumemeroue.github.io/IJCAI25>. This Web application makes it possible to explore the characteristics of optimal CSDs for various scoring vectors, probabilistic models, number of agents and goods.

7 Discussion

The practical use of our setting raises a few questions.

First, we need to choose a distribution. The choice has to be tailored to the domain at hand, and distributions can be learnt using some preference learning models and techniques. If computation time is an important issue then it is wise to learn a Plackett-Luce model [20].

Second, we need to choose a scoring vector as a proxy for agents' valuations over items. Again, this depends on the specific domain at hand. For each context, the scores can be estimated by an experiment where subjects are presented with a list of items to elicit their valuations; see Appendix E.

Third, we need to choose a social welfare functional. We have seen that, unsurprisingly, utilitarianism may lead to clearly unfair solutions and should be used only with care. As usual, egalitarianism may

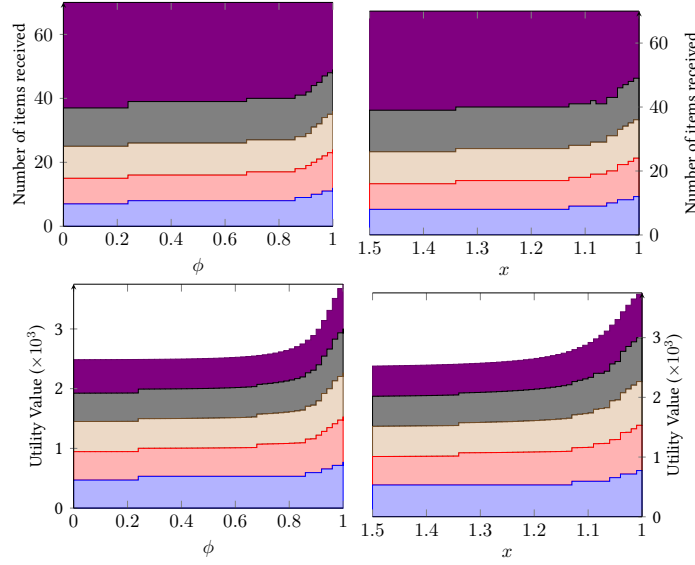


Figure 2: Number of goods received per agent (top); expected utility value per agent (bottom) as a function of ϕ for $M11_{\phi,\mu}$ and x for PL_{ν^x} . Maximizing ESW, Borda scoring vector, $n = 5$, $m = 70$.

lead to a loss of efficiency, but is easier to compute or approximate; Nash is a good trade-off (see [18] for a manifesto towards using Nash social welfare in fair division) but is hard to compute if the distribution does not satisfy prefix independence.

Four, once a CSD is found, it is anonymous: for instance, with two agents, if the output is $(1, 2)$, it does not say who should start picking. Assigning agents to positions in the sequence has no impact on *ex ante* social welfare, but it may have an impact on *ex post* social welfare (see Appendix F).

8 Conclusion

Our main messages are: (1) imposing strategyproofness does not leave much choice beyond constrained serial dictatorships; (2) some constrained serial dictatorships are fairer than others; (3) their efficiency and fairness can be measured by expected social welfare, defined by a scoring vector, a distribution over profiles, and a social welfare functional; (4) depending on the social welfare functional and the distribution, the optimal sequence can be polynomial-time computable, efficiently approximated by sampling, or hard to approximate by sampling. The following table summarizes the results obtained. PI means that prefix independence is satisfied, poly means “polynomial-time computable”, and approx means “efficiently approximable by sampling”.

Ψ	PI	$EU_{\Psi}^k(a_i)$	Egal	Nash	Uti
FC	yes	poly	poly	poly	poly
IC	yes	poly	poly	poly	poly
PL_{ν}	yes	approx	approx	approx	approx
$M11_{\phi,\mu}$	no	approx	approx	?	?

Even in the cases where we are able to compute optimal sequences in polynomial time, we do not know any closed-form formulas for these optimal sequences.

If items were bads (e.g., chores) instead of goods, a similar methodology would work, with values in the scoring vector representing costs. Of course, agents coming first in the sequence should now take *more* items than those coming later.

Acknowledgements

This work was supported in part by project ANR-22-CE26-0019 “Citizens” and project PR[AI]RIE-PSAI ANR-23-IACL-0008, both funded by the Agence Nationale de la Recherche.

References

- [1] Georgios Amanatidis, Georgios Birmpas, and Evangelos Markakis. On truthful mechanisms for maximin share allocations. In Subbarao Kambhampati, editor, *Proceedings of the 25th International Joint Conference on Artificial Intelligence, IJCAI 2016, New York, NY, USA, 9-15 July 2016*, pages 31–37. IJCAI/AAAI Press, 2016.
- [2] Georgios Amanatidis, Georgios Birmpas, George Christodoulou, and Evangelos Markakis. Truthful allocation mechanisms without payments: Characterization and implications on fairness. In *Proceedings of the 2017 ACM Conference on Economics and Computation, EC ’17, Cambridge, MA, USA, June 26-30, 2017*, pages 545–562. ACM, 2017.
- [3] Georgios Amanatidis, Haris Aziz, Georgios Birmpas, Aris Filos-Ratsikas, Bo Li, Hervé Moulin, Alexandros A. Voudouris, and Xiaowei Wu. Fair division of indivisible goods: Recent progress and open questions. *Artif. Intell.*, 322:103965, 2023.
- [4] Haris Aziz, Sylvain Bouveret, Jérôme Lang, and Simon Mackenzie. Complexity of manipulating sequential allocation. In *Proceedings of the 31st AAAI Conference on Artificial Intelligence, San Francisco, California, USA, February 4-9, 2017*, pages 328–334. AAAI Press, 2017.
- [5] Haris Aziz, Bo Li, and Xiaowei Wu. Strategyproof and approximately maximin fair share allocation of chores. In *Proceedings of the 28th International Joint Conference on Artificial Intelligence, IJCAI 2019, Macao, China, August 10-16, 2019*, pages 60–66. ijcai.org, 2019.
- [6] Haris Aziz, Bo Li, Hervé Moulin, and Xiaowei Wu. Algorithmic fair allocation of indivisible items: A survey and new questions. *CoRR*, abs/2202.08713, 2022.
- [7] Dorothea Baumeister, Sylvain Bouveret, Jérôme Lang, Nhan-Tam Nguyen, Trung Thanh Nguyen, Jörg Rothe, and Abdallah Saffidine. Positional scoring-based allocation of indivisible goods. *Auton. Agents Multi Agent Syst.*, 31(3):628–655, 2017.
- [8] Niclas Boehmer, Piotr Faliszewski, ukasz Janeczko, Andrzej Kaczmarczyk, Grzegorz Lisowski, Grzegorz Pierczynski, Simon Rey, Dariusz Stolicki, Stanisaw Szufa, and Tomasz Ws. Guide to numerical experiments on elections in computational social choice. In Kate Larson, editor, *Proceedings of the Thirty-Third International Joint Conference on Artificial Intelligence, IJCAI-24*, pages 7962–7970. International Joint Conferences on Artificial Intelligence Organization, 8 2024. doi: 10.24963/ijcai.2024/881. URL <https://doi.org/10.24963/ijcai.2024/881>. Survey Track.
- [9] Anna Bogomolnaia, Rajat Deb, and Lars Ehlers. Strategy-proof assignment on the full preference domain. *J. Econ. Theory*, 123(2):161–186, 2005.
- [10] Craig Boutilier, Ioannis Caragiannis, Simi Haber, Tyler Lu, Ariel D. Procaccia, and Or Sheffet. Optimal social choice functions: A utilitarian view. *Artif. Intell.*, 227:190–213, 2015.
- [11] Sylvain Bouveret and Jérôme Lang. A general elicitation-free protocol for allocating indivisible goods. In *Proceedings of the 22nd International Joint Conference on Artificial Intelligence, IJCAI 2011, Barcelona, Catalonia, Spain, July 16-22, 2011*, pages 73–78. IJCAI/AAAI Press, 2011.
- [12] Sylvain Bouveret and Jérôme Lang. Manipulating picking sequences. In *Proceedings of the 21st European Conference on Artificial Intelligence, ECAI 2014, Prague, Czech Republic, August 18-22, 2014*, volume 263 of *Frontiers in Artificial Intelligence and Applications*, pages 141–146. IOS Press, 2014.
- [13] Sylvain Bouveret, Yann Chevaleyre, and Nicolas Maudet. Fair allocation of indivisible goods. In *Handbook of Computational Social Choice*, pages 284–310. Cambridge University Press, 2016.
- [14] Steven Brams, Paul Edelman, and Peter Fishburn. Fair division of indivisible items. *Theory and Decision*, 55:147–180, 02 2003. doi: 10.1023/B:THEO.0000024421.85722.0a.

- [15] Steven J. Brams and Alan D. Taylor. *The Win-win Solution. Guaranteeing Fair Shares to Everybody*. W. W. Norton & Company, 2000.
- [16] Xiaolin Bu and Biaoshuai Tao. Truthful and almost envy-free mechanism of allocating indivisible goods: the power of randomness. *CoRR*, abs/2407.13634, 2024. doi: 10.48550/ARXIV.2407.13634. URL <https://doi.org/10.48550/arXiv.2407.13634>.
- [17] Ioannis Caragiannis and Nidhi Rathi. Optimizing over serial dictatorships. In *SAGT*, volume 14238 of *Lecture Notes in Computer Science*, pages 329–346. Springer, 2023.
- [18] Ioannis Caragiannis, David Kurokawa, Hervé Moulin, Ariel D. Procaccia, Nisarg Shah, and Junxing Wang. The unreasonable fairness of maximum nash welfare. *ACM Trans. Economics and Comput.*, 7(3):12:1–12:32, 2019.
- [19] Mithun Chakraborty, Ulrike Schmidt-Kraepelin, and Warut Suksompong. Picking sequences and monotonicity in weighted fair division. *Artif. Intell.*, 301:103578, 2021.
- [20] Weiwei Cheng, Eyke Hüllermeier, and Krzysztof J Dembczynski. Label ranking methods based on the plackett-luce model. In Johannes Fürnkranz and Thorsten Joachims, editors, *Proceedings of the 27th International Conference on Machine Learning (ICML-10), June 21-24, 2010, Haifa, Israel*, pages 215–222. Omnipress, 2010.
- [21] Andreas Darmann and Christian Klamler. Proportional borda allocations. *Soc. Choice Welf.*, 47(3): 543–558, 2016.
- [22] Claude d’Aspremont and Louis Gevers. Chapter 10 social welfare functionals and interpersonal comparability. In *Handbook of Social Choice and Welfare*, volume 1 of *Handbook of Social Choice and Welfare*, pages 459–541. Elsevier, 2002.
- [23] Lars Ehlers and Bettina Klaus. Coalitional strategy-proof and resource-monotonic solutions for multiple assignment problems. *Soc. Choice Welf.*, 21(2):265–280, 2003.
- [24] Michele Flammini and Hugo Gilbert. Parameterized complexity of manipulating sequential allocation. In *Proceedings of the 24th European Conference on Artificial Intelligence, ECAI 2020, Santiago de Compostela, Spain, August 29 - September 8, 2020*, volume 325 of *Frontiers in Artificial Intelligence and Applications*, pages 99–106. IOS Press, 2020.
- [25] Rohan Garg and Alexandros Psomas. Efficient mechanisms without money: Randomization won’t let you escape from dictatorships. Manuscript, 2022.
- [26] John William Hatfield. Strategy-proof, efficient, and nonbossy quota allocations. *Soc. Choice Welf.*, 33(3):505–515, 2009.
- [27] Wassily Hoeffding. Probability inequalities for sums of bounded random variables. *Journal of the American Statistical Association*, 1963.
- [28] Hadi Hosseini and Kate Larson. Multiple assignment problems under lexicographic preferences. In *Proceedings of the 18th International Conference on Autonomous Agents and MultiAgent Systems, AAMAS ’19, Montreal, QC, Canada, May 13-17, 2019*, pages 837–845. International Foundation for Autonomous Agents and Multiagent Systems, 2019.
- [29] Hadi Hosseini, Sujoy Sikdar, Rohit Vaish, and Lirong Xia. Fair and efficient allocations under lexicographic preferences. In *35th AAAI Conference on Artificial Intelligence, AAAI 2021, Virtual Event, February 2-9, 2021*, pages 5472–5480. AAAI Press, 2021.
- [30] Thomas Kalinowski, Nina Narodytska, and Toby Walsh. A social welfare optimal sequential allocation procedure. In *Proceedings of the 23rd International Joint Conference on Artificial Intelligence, IJCAI 2013, Beijing, China, August 3-9, 2013*, pages 227–233. IJCAI/AAAI Press, 2013.
- [31] Thomas Kalinowski, Nina Narodytska, Toby Walsh, and Lirong Xia. Strategic behavior when allocating indivisible goods sequentially. In *Proceedings of the 27th AAAI Conference on Artificial Intelligence, Bellevue, Washington, USA, July 14-18, 2013*. AAAI Press, 2013.
- [32] David A. Kohler and R. Chandrasekaran. A class of sequential games. *Operations Research*, 19(2): 270–277, 1971.
- [33] Fuhito Kojima. Random assignment of multiple indivisible objects. *Math. Soc. Sci.*, 57(1):134–142, 2009.
- [34] Jérôme Lang and Jörg Rothe. Fair division of indivisible goods. In *Economics and Computation*,

- Springer texts in business and economics, pages 493–550. Springer, 2nd edition, 2024.
- [35] R Duncan Luce. *Individual choice behavior*, volume 4. Wiley New York, 1959.
 - [36] Colin L Mallows. Non-null ranking models. i. *Biometrika*, 44(1/2):114–130, 1957.
 - [37] Nhan-Tam Nguyen, Dorothea Baumeister, and Jörg Rothe. Strategy-proofness of scoring allocation correspondences for indivisible goods. *Soc. Choice Welf.*, 50(1):101–122, 2018.
 - [38] Szilvia Pápai. original papers : Strategyproof multiple assignment using quotas. *Review of Economic Design*, 5(1):91–105, 2000.
 - [39] Robin L Plackett. The analysis of permutations. *Journal of the Royal Statistical Society Series C: Applied Statistics*, 24(2):193–202, 1975.
 - [40] Alexandros Psomas and Paritosh Verma. Fair and efficient allocations without obvious manipulations. In S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh, editors, *Advances in Neural Information Processing Systems*, volume 35, pages 13342–13354. Curran Associates, Inc., 2022. URL https://proceedings.neurips.cc/paper_files/paper/2022/file/57250222014c35949476f3f272c322d2-Paper-Conference.pdf.
 - [41] Szilvia Pápai. Strategyproof and Nonbossy Multiple Assignments. *Journal of Public Economic Theory*, 3(3):257–271, July 2001.
 - [42] Arthur Richard Silverberg. *Statistical models for q-permutations*. Princeton University, 1980.
 - [43] Lars-Gunnar Svensson. Strategy-proof allocation of indivisible goods. *Social Choice and Welfare*, 16(4):557–567, 1999.
 - [44] Yuto Tominaga, Taiki Todo, and Makoto Yokoo. Manipulations in two-agent sequential allocation with random sequences. In *Proceedings of the 2016 International Conference on Autonomous Agents & Multiagent Systems, Singapore, May 9-13, 2016*, pages 141–149. ACM, 2016.
 - [45] Mingyu Xiao and Jiaying Ling. Algorithms for manipulating sequential allocation. In *Proceedings of the 34th AAAI Conference on Artificial Intelligence, AAAI 2020, New York, NY, USA, February 7-12, 2020*, pages 2302–2309. AAAI Press, 2020.

Sylvain Bouveret
 Univ. Grenoble Alpes, CNRS, Grenoble INP, LIG,
 Grenoble, France
 Email: sylvain.bouveret@imag.fr

Hugo Gilbert
 Université Paris-Dauphine, Université PSL, CNRS, LAMSADE
 Paris, France
 Email: hugo.gilbert@lamsade.dauphine.fr

Jérôme Lang
 Université Paris-Dauphine, Université PSL, CNRS, LAMSADE
 Paris, France
 Email: lang@lamsade.dauphine.fr

Guillaume Méréroué
 Université Côte d’Azur, Inria, CNRS, I3S
 Sophia Antipolis, France
 Email: guillaume.meroué@inria.fr

Supplementary material to submission “Constrained Serial Dictatorships can be Fair”

A Omitted Proofs of Section 3

Observation 1. For given n and m , the number of vectors $\mathbf{k} = (k_1, \dots, k_n)$ such that $\sum_{i=1}^n k_i = m$ equals $\binom{n+m-1}{n-1}$.

Proof. Choosing n numbers $(k_1, \dots, k_n)$ matching the definition amounts to partition $[m]$ into n subintervals, which in turn amounts to choose $n - 1$ “separation bars”. Said otherwise, this comes down to choose $n - 1$ increasing numbers $l_1 \leq \dots \leq l_{n-1}$ among $m + 1$. Here, $k_i = l_i - l_{i-1}$, with the convention that $l_0 = 0$. This problem can be equivalently formulated as the one of drawing $n - 1$ different numbers among $n + m - 1$. For each such draw, we can obtain a set of increasing numbers $l'_1 \leq \dots \leq l'_{n-1}$ between 1 and $n + m - 1$, that can be cast to increasing numbers between 0 and m by choosing $l_i = l'_i - i$. Since there are $\binom{n+m-1}{n-1}$ subsets of $n - 1$ elements among $n + m - 1$, we obtain the result. $\square$

B Omitted Proofs of Section 4

We now show that prefix independence entails a specific property for the optimal solutions of OptSD- Ψ - E .

Proposition 3. Under prefix independence, there exists an optimal solution to OptSD- Ψ - E which is non-decreasing, i.e., in which the earlier an agent picks, the less goods she gets.

Proof. Let $\text{eu}(\kappa, \tau)$ denote the utility obtained by an agent if we allocate κ items to her knowing that τ items have already been allocated. As items have positive valuations and agent’s preference rankings are drawn independently from the same probabilistic model, it is easy to prove that $\text{eu}(\kappa, \tau)$ is non-decreasing in κ and non-increasing in τ . Let us consider a solution $\mathbf{k} = (k_1, \dots, k_n)$, which is not a non-decreasing vector. Then, there exists $i \in [n - 1]$ such that $k_i > k_{i+1}$. We set $\tau_i = \sum_{j=1}^{i-1} k_j$, $\tau_{i+1} = \tau_i + k_i$ and $\tau'_{i+1} = \tau_i + k_{i+1}$. From the properties of function eu , it is clear that $\min(\text{eu}(k_i, \tau_i), \text{eu}(k_{i+1}, \tau_{i+1})) = \text{eu}(k_{i+1}, \tau_{i+1}) \leq \min(\text{eu}(k_{i+1}, \tau_i), \text{eu}(k_i, \tau'_{i+1}))$. Hence, by swapping k_i and k_{i+1} in $\mathbf{k}$, we do not decrease the egalitarian score of $\mathbf{k}$ (note that this swap does not affect the utility values received by agents other than a_i and a_{i+1}). The repetition of this argument shows that there exists an optimal solution to OptSD- Ψ - E which is a non-decreasing vector. $\square$

We will see in the following example that this property fails for utilitarian social welfare. For Nash social welfare, we conjecture it holds, but so far we do not have a proof.

Example 2. Let $n = 3$, $m = 7$ and the Borda scoring vector. By dynamic programming we find the values $\text{eu}(\kappa, \tau)$ displayed on Table 1. For USW, we obtain an optimal vector $\mathbf{k} = (3, 2, 2)$, yielding expected social welfare 37.2. For maximizing ESW and NSW, we obtain $\mathbf{k} = (2, 2, 3)$, with expected social welfare 12 and 1872 respectively. Note that the optimal vectors for ESW and NSW may be different: with $n = 4$ and $m = 10$ and the Borda scoring vector, $\mathbf{k} = (2, 2, 2, 4)$ is optimal for ESW ; and $\mathbf{k} = (2, 2, 3, 3)$ for NSW.

We see on Example 2 that the optimal sequence for utilitarian social welfare, IC, and Borda scoring, is not non-decreasing, and thus clearly not fair: the first agent in the sequence not only has a larger choice of items but picks one more than the other two! It is not new that utilitarianism may clash fairness when looking for optimal CSDs: for instance, it is known that the optimal sequence for utilitarianism, Borda scoring, FI, $n = 2$ and m even is perfect alternation 1212...12, which is obviously not fair [30].

Table 1: Utilities $eu(\kappa, \tau)$ in Example 2; $m = 7$, IC model.

$\kappa \backslash \tau$	0	1	2	3	4	5	6	7
0	0	0	0	0	0	0	0	0
1	7	6.86	6.67	6.4	6	5.33	4	-
2	13	12.57	12	11.2	10	8	-	-
3	18	17.14	16	14.4	12	-	-	-
4	22	20.57	18.67	16	-	-	-	-
5	25	22.86	20	-	-	-	-	-
6	27	24	-	-	-	-	-	-
7	28	-	-	-	-	-	-	-

(Still, we continue to include utilitarianism in our study, first for the sake of comparison, and second because utilitarianism is relevant in some situations.)

The following result concerns utilitarian social welfare:

Proposition 12. *Given a scoring vector s , a finite set $\mathcal{P}$ of n -agent preference profiles over a set of goods and an integer K , the problem of determining whether there is a CSD $\mathbf{k}$ such that the average utilitarian social welfare of $\mathbf{k}$ over all profiles of $\mathcal{P}$ is greater than or equal to K is NP-complete.*

Proof. We will prove the proposition by reduction from Exact-Cover-By-3-Sets (X3C):

X3C
Input: A set $\mathcal{X} = \{x_1, \dots, x_n\}$ of n elements; a collection $\mathcal{S} = \{S_1, \dots, S_m\}$ of m subsets such that $\forall S \in \mathcal{S}$, S contains exactly three elements of $\mathcal{X}$. Question: Does there exist a subcollection $\mathcal{C} \subseteq \mathcal{S}$, such that $\bigcup_{S \in \mathcal{C}} S = \mathcal{X}$ and $S \cap S' = \emptyset, \forall S, S' \in \mathcal{C}$.

Let $(\mathcal{X}, \mathcal{S})$ be an X3C instance. From that instance, we create an instance of our problem with n goods and $3m$ agents such that the set of goods is exactly $\mathcal{X}$ (by notation abuse), and such that there are 3 agents a_S^1, a_S^2 , and a_S^3 for each set $S \in \mathcal{S}$.

For each set $S = \{a, b, c\} \in \mathcal{S}$, we create 3 rankings r_S^a, r_S^b, r_S^c such that r_S^a starts with a , r_S^b starts with b , and r_S^c starts with c (the rest of the ranking does not matter). Then for each pair of agents (a_S^i, a_T^j) with $S \neq T$, we create 9 profiles as follows:

- if $S = \{a, b, c\}$, agents a_S^1, a_S^2, a_S^3 have rankings from $\{(r_S^a, r_S^b, r_S^c), (r_S^b, r_S^c, r_S^a), (r_S^c, r_S^a, r_S^b)\}$;
- if $T = \{d, e, f\}$, agents a_T^1, a_T^2, a_T^3 have rankings from $\{(r_T^d, r_T^e, r_T^f), (r_T^e, r_T^f, r_T^d), (r_T^f, r_T^d, r_T^e)\}$;
- for each $X = \{x, y, z\}$ different from S and T , agents a_X^1, a_X^2, a_X^3 have rankings (r_X^x, r_X^y, r_X^z) .

This thus makes $\frac{81m(m-1)}{2}$ profiles in total. Now, the scoring vector is such that the top object has utility 1 while all the other items have utility 0. We will prove that there exists an exact cover iff there exists a CSD with average utility at least n .

($\Rightarrow$) If $\mathcal{C} \subseteq \mathcal{S}$ is an exact cover, let agents a_S^i for $S \in \mathcal{C}$ and $i \in \{1, 2, 3\}$ pick one item. By construction, for each profile, these agents will all have their top choices yielding a utility of n .

($\Leftarrow$) Conversely, let $\mathbf{k}$ be a CSD yielding utility n for all profiles. Necessarily $\mathbf{k}$ gives exactly one item to n agents. Let a_S^i and a_T^j be two such agents, with $S \neq T$. Suppose that $S \cap T \neq \emptyset$ and let $x \in S \cap T$. In the profile where a_S^i has ranking r_S^x and a_T^j has ranking r_T^x both agents have the same top object. Hence, the utility yielded by $\mathbf{k}$ is necessarily strictly lower than n for this profile, a contradiction with the hypothesis. This proves that all the agents a_S^i and a_T^j receiving an object in $\mathbf{k}$ are such that either $S = T$ or $S \cap T = \emptyset$. Hence, we necessarily obtain $n/3$ sets who are pairwise disjoint and hence provide an exact cover. $\square$

C Omitted Proofs of Section 5

To prove Lemma 2, we first need to prove the following Lemma.

Lemma 3. *Let $\tilde{r} : g_{i_1} \succ g_{i_2} \succ \dots \succ g_{i_q}$ be an incomplete ranking over $\mathcal{G}$ with $q \leq m$. Under the PL model, the probability to generate a ranking r which is a consistent extension of $\tilde{r}$ is equal to:*

$$\prod_{j=1}^q \frac{\nu_{i_j}}{\sum_{l=j}^q \nu_{i_l}}$$

Proof. Let $S \in \mathcal{G}$ be the set of goods on which $\tilde{r}$ express preferences. The lemma can easily be derived from the vase model metaphor. Consider the following slightly different sequential process. At each stage, a ball is taken from the vase such that a ball of type j is chosen with probability $f(j)$. If the ball is of a different type than the ones previously picked, *and is a ball of a type corresponding to an element of S* , then it yields the next item in the ranking. In either case, the ball is put back in the vase and the process continues. This process generates a ranking on the elements of S according to the original PL model. We get that the probability of $\tilde{r}$ is:

$$\prod_{j=1}^q \frac{\nu_{i_j}}{\sum_{l=j}^q \nu_{i_l}}.$$

$\square$

Lemma 2. *Let $I = (i_1, \dots, i_q)$ be a sequence of q different indices in $[m]$. Consider the following two cases:*

- i) *Agent a_1 picks q goods;*
- ii) *Agent a_1 picks q_1 goods and agent a_2 picks q_2 goods with $q_1 + q_2 = q$.*

For the PL model, the probability that for all $t \in [q]$, g_{i_t} is picked at timestep t is the same in cases i and ii.

Proof. We will show that in both cases, the probability that for all $t \in [q]$ g_{i_t} is picked at timestep t is:

$$p_I = \prod_{j=1}^q \frac{\nu_{i_j}}{\sum_{l=j}^q \nu_{i_l} + \sum_{p \in [m] \setminus I} \nu_p}.$$

In case i), g_{i_t} is picked at timestep t for all $t \in [q]$ if $\text{rk}_{\mathbf{P}}^{a_1}(g_{i_t}) = t$ for all $t \in [q]$. Under the PL model, this occurs with probability p_I .

In case ii), let I_1 (resp. I_2) be the subsequence composed of the q_1 first (resp. q_2 last) elements of I and $S_1 = \{g_{i_1}, \dots, g_{i_{q_1}}\}$ (resp. $S_2 = \{g_{i_{q_1+1}}, \dots, g_{i_q}\}$). In the PL model, the probability that a_1 picks g_{i_t} at timestep t for all $t \in [q_1]$ is:

$$p_I^1 = \prod_{j=1}^{q_1} \frac{\nu_{i_j}}{\sum_{l=j}^q \nu_{i_l} + \sum_{p \in [m] \setminus I} \nu_p}.$$

Then, the probability that a_2 picks g_{i_t} at timestep t for all $t \in [q] \setminus [q_1]$ corresponds to the probability that g_{i_t} is ranked at position $t - q_1$, when restricting ourselves to the goods in $\mathcal{G} \setminus S_1$. Put another way, goods in S_2 should be ranked in the top q_2 positions in the partial ranking which only ranks goods in $\mathcal{G} \setminus S_1$. Let $\tilde{r} : g_{i'_1} \succ g_{i'_2} \succ \dots \succ g_{i'_{m-q_1}}$ be one such ranking over $\mathcal{G} \setminus S_1$ with $g_{i'_l} = g_{i_{(q_1+l)}}$ for all $l \in [q_2]$. Resorting to Lemma 3, $\tilde{r}$ occurs with probability:

$$\prod_{j=1}^{q_2} \frac{\nu_{i'_j}}{\sum_{l=j}^{m-q_1} \nu_{i'_l}} \prod_{j=q_2+1}^{m-q_1} \frac{\nu_{i'_j}}{\sum_{l=j}^{m-q_1} \nu_{i'_l}}.$$

By marginalizing over all such rankings, the second product vanishes as we obtain the sum of probabilities over all rankings over $\mathcal{G} \setminus (S_1 \cup S_2)$ under PL_ν . Hence, we get probability:

$$p_I^2 = \prod_{j=1}^{q_2} \frac{\nu_{i'_j}}{\sum_{l=j}^{q_2} \nu_{i'_l} + \sum_{p \in [m] \setminus I} \nu_p}.$$

The product of p_I^1 and p_I^2 yields exactly p_I . □

Proposition 6. *There exists $\phi \in (0, 1)$ and a ranking μ such that $\Psi = \text{Mll}_{\phi, \mu}$ does not satisfy prefix independence.*

Proof. To see why, consider the case of $m = 3$ items and $n = 3$ agents, and a Mallows model with center $\mu = a \succ b \succ c$ and parameter ϕ . The probability of a ranking r to occur is $\phi^{d_{KT}(r, \mu)} / C$ where $C = 1 + 2\phi + 2\phi^2 + \phi^3$ is a normalization constant. The probabilities of the different rankings are described on the table below.

	Probability	Ranking
1	$1/C$	$a \succ b \succ c$
2	ϕ/C	$a \succ c \succ b$
3	ϕ/C	$b \succ a \succ c$
4	ϕ^2/C	$b \succ c \succ a$
5	ϕ^2/C	$c \succ a \succ b$
6	ϕ^3/C	$c \succ b \succ a$

If the expected utility that an agent gets in the allocation process only depend on the number of goods that she picks and that have been picked before she started picking, then it should be the same for a_3 if (1) a_1 and a_2 both picked one and if (2) a_1 picked two and a_2 picked zero items. We will see that it is not the case. Let us assume that $s = (1, 1, 0)$ so that a_3 only cares about not getting her least preferred item. The probability that a (resp. b , c) is her least preferred good is $\phi^2(1 + \phi)/C$ (resp. $\phi(1 + \phi)/C$, $(1 + \phi)/C$).

We can compute the probability that the first two goods picked (by agent 1 and 2) are a and b in the two cases. In the first case, either agent 1 picks a first (so it has ranking 1 or 2) and agent 2 picks b (so it has ranking 1, 3 or 4), or agent 1 picks b first (ranking 3 or 4) and agent 2 picks a (ranking 1 or 2 or 3). The probability of this is

$$\begin{aligned} & \frac{1+\phi}{C} \frac{1+\phi+\phi^2}{C} + \frac{\phi+\phi^2}{C} \frac{1+\phi+\phi}{C} \\ &= \frac{(1+\phi)(1+2\phi+3\phi^2)}{C^2} \end{aligned}$$

In the second case, the probability that agent 1 picks both a and b is the probability of rankings 1 and 3, which is:

$$\frac{1+\phi}{C} = \frac{(1+\phi)(1+2\phi+2\phi^2+\phi^3)}{C^2}.$$

Similarly:

- The probability that items a and c are picked by the first two agents is $\phi(1+\phi)/C$ in cases 1 and 2.
- The probability that items b and c are picked by the first two agents is $\phi^3(1+\phi)(3+2\phi+\phi^2)/C^2$ in case 1 and $\phi^2(1+\phi)/C$ in case 2.

As a result, the probability that a_3 gets her least preferred good is:

$$\begin{aligned} & \frac{(1+\phi)^2}{C^2} \frac{(1+2\phi+3\phi^2)}{C} + \phi^2 \frac{(1+\phi)^2}{C^2} \\ & \quad + \phi^4 \frac{(1+\phi)^2}{C^2} \frac{(3\phi+2\phi^2+\phi^3)}{C} \end{aligned}$$

in case 1 and

$$\frac{(1+\phi)^2}{C^2} + \phi^2 \frac{(1+\phi)^2}{C^2} + \phi^4 \frac{(1+\phi)^2}{C^2}$$

in case 2. These are two different values, e.g., for $\phi = 0.5$ we obtain 0.440 in case 1 and 0.429 in case 2. $\square$

Proposition 9. Let $\epsilon > 0$ and $\delta \in (0, 1)$ two fixed values, and Υ an upper bound on values $\text{eu}(\kappa, \tau)$ (e.g., $\sum_{i=1}^m s_i$).

Let $\widetilde{\text{eu}}_{\kappa, \tau}$ be the value computed by averaging the values $u_{\kappa, \tau}(\mathbf{P}_i, \mathbf{s})$ over N preference profiles $\mathbf{P}_i$ sampled independently from Ψ . If $N \geq (\Upsilon^2 \ln(2m^2/\delta))/2\epsilon^2$, then it holds with probability $1 - \delta$ that:

$$|\text{eu}(\kappa, \tau) - \widetilde{\text{eu}}_{\kappa, \tau}| \leq \epsilon, \forall \kappa, \tau \in [m] \times [m - \kappa].$$

Proof. Let $\text{eu}_{\kappa, \tau}^i$ represent the utility of an agent receiving t items after k items have already been taken, for the preference profile $\mathbf{P}_i$ such that $\widetilde{\text{eu}}_{\kappa, \tau} = \sum_{i=1}^n \text{eu}_{\kappa, \tau}^i$.

We aim to show that with probability at least $1 - \delta$, the sampled value of utility is close to the expected utility within ϵ . Formally, we want:

$$\Pr \left(\left| \frac{\widetilde{\text{eu}}_{\kappa, \tau}}{n} - \text{eu}_{\kappa, \tau} \right| \leq \epsilon \right) \geq 1 - \delta$$

However, this form is not directly suitable for applying Hoeffding's inequality, so we first perform some manipulations. We have:

$$\begin{aligned}\Pr\left(\left|\frac{\widetilde{\mathbf{e}\mathbf{u}}_{\kappa,\tau}}{n} - \mathbf{e}\mathbf{u}_{\kappa,\tau}\right| \leq \epsilon\right) &= 1 - \Pr\left(\left|\frac{\widetilde{\mathbf{e}\mathbf{u}}_{\kappa,\tau}}{n} - \mathbf{e}\mathbf{u}_{\kappa,\tau}\right| > \epsilon\right) \\ &\geq 1 - \Pr\left(\left|\frac{\widetilde{\mathbf{e}\mathbf{u}}_{\kappa,\tau}}{n} - \mathbf{e}\mathbf{u}_{\kappa,\tau}\right| \geq \epsilon\right)\end{aligned}$$

Hoeffding's inequality gives us:

$$\Pr\left(\left|\frac{\widetilde{\mathbf{e}\mathbf{u}}_{\kappa,\tau}}{n} - \mathbf{e}\mathbf{u}_{\kappa,\tau}\right| \geq \epsilon\right) \leq 2 \exp\left(-\frac{2n^2\epsilon^2}{\sum_{i=1}^n (b_i - a_i)^2}\right)$$

Given that $\mathbf{e}\mathbf{u}_{\kappa,\tau}^i$ are identically distributed, we can rewrite $\sum_{i=1}^n (b_i - a_i)^2$ as $n \times (b - a)^2$, where $a = \sum_{i=1}^t s(i)$ when the agent receives the items they like the least, and $b = \sum_{i=1}^t s(m - i)$ when the agent receives the items they prefer the most. For simplicity, we set $a = 0$ (no selected item) and $b = \Upsilon$ (all items are selected).

Substituting these expressions into the inequality, we get:

$$\Pr\left(\left|\frac{\widetilde{\mathbf{e}\mathbf{u}}_{\kappa,\tau}}{n} - \mathbf{e}\mathbf{u}_{\kappa,\tau}\right| \leq \epsilon\right) \geq 1 - 2 \exp\left(-\frac{2n\epsilon^2}{\Upsilon^2}\right)$$

We want this probability to be at least $1 - \delta$, which gives us:

$$1 - 2 \exp\left(-\frac{2n\epsilon^2}{\Upsilon^2}\right) \geq 1 - \delta$$

This rearranges to:

$$\exp\left(\frac{2n\epsilon^2}{\Upsilon^2}\right) \geq \frac{2}{\delta}$$

Thus, for the probability to be at least $1 - \delta$, the number of samples n must satisfy the inequality below.

$$n \geq \frac{\Upsilon^2 \ln\left(\frac{2}{\delta}\right)}{2\epsilon^2}$$

Then let $E_{\kappa,\tau}$ denote the event that $\widetilde{\mathbf{e}\mathbf{u}}_{\kappa,\tau}$ is an ϵ -additive approximation of $\mathbf{e}\mathbf{u}_{\kappa,\tau}$. This event occurs with probability $1 - \delta$.

We want to determine the probability that all events $E_{\kappa,\tau}$ hold simultaneously. This is equivalent to computing:

$$\Pr(\cap_{\kappa,\tau} E_{\kappa,\tau}) = 1 - \Pr(\cup_{\kappa,\tau} \overline{E_{\kappa,\tau}})$$

Applying the union bound, we get:

$$\Pr(\cap_{\kappa,\tau} E_{\kappa,\tau}) \geq 1 - \sum_{k,t \in [m]} \Pr(\overline{E_{\kappa,\tau}})$$

Given that $\Pr(\overline{E_{\kappa,\tau}}) \leq \delta$, we have:

$$\Pr(\cap_{\kappa,\tau} E_{\kappa,\tau}) \geq 1 - m^2\delta$$

For this probability to be at least $1 - \Delta$, we require:

$$1 - m^2\delta \geq 1 - \Delta$$

This rearranges to:

$$\delta \leq \frac{\Delta}{m^2}$$

Substituting this condition into the sample size inequality, we find that n must satisfy:

$$n \geq \frac{\Upsilon^2 \ln\left(\frac{2m^2}{\Delta}\right)}{2\epsilon^2}$$

□

Proposition 10. *If $\Psi = \text{PL}_\nu$, then all values $\text{eu}(\kappa, \tau)$ can be computed in time $O(4^m \text{Poly}(m))$.*

Proof. Let $S \subseteq G$ be a set of goods and $g \in S$, we denote by $\text{ft}(g, S) = \nu_g / (\sum_{g' \in S} \nu_{g'})$, the probability that g is ranked first among the elements of S according to PL_ν .

The proof relies on recursive equations. Let us consider the following setting. The picker under consideration should pick κ goods within a set S of goods occupying the $|S|$ last positions of her ranking. Goods in $G \setminus S$, occupy the top $m - |S|$ ranks in her ranking and have already been picked, either by her or by other agents. Moreover, the set $S' \subseteq S$ have already been picked by previous pickers. We are interested in computing the expected utility $U(\kappa, S, S')$ of the κ picks of the agent in such a situation. We argue that $U(\kappa, S, S')$ satisfies the following recursive equations:

$$\begin{aligned} U(\kappa, S, S) &= \sum_{g \in S \cap S} \text{ft}(g, S) U(\kappa, S \setminus \{g\}, S \setminus \{g\}) \\ &\quad + \sum_{g \in S \setminus S} \text{ft}(g, S) (s_{m-|S|+1} + U(\kappa - 1, S \setminus \{g\}, S')) \end{aligned} \quad (3)$$

$$U(0, S, S') = 0 \quad \forall S, S' \quad (4)$$

$$U(\kappa, S, \emptyset) = \sum_{i=1}^{\kappa} s_{m-|S|+i} \quad \forall \kappa, S \quad (5)$$

In Equation 3, we consider all possible goods which could be placed at rank $m - |S| + 1$. This occurs for good $g \in S$ with probability $\text{ft}(g, S)$. If $g \in S'$, this good has already been picked and the agent still has to pick κ goods within the goods in $S \setminus \{g\}$ which are ranked in last position, hence we consider $U(\kappa, S \setminus \{g\}, S \setminus \{g\})$. If $g \notin S'$, this good is picked by the agent leading to a utility $s_{m-|S|+1}$ and the agent still has to pick $\kappa - 1$ goods within the goods in $S \setminus \{g\}$ which are ranked in last position, hence we consider $U(\kappa - 1, S \setminus \{g\}, S)$. Equations 4 and 5 provide the base cases.

Next, we consider the probability $P(S, S')$ that goods in S are ranked in the top $|S|$ positions among a set S' of goods. $P(S, S')$ trivially satisfies the following recursive equation.

$$\begin{aligned} P(S, S') &= \sum_{g \in S} \text{ft}(g, S') P(S \setminus \{g\}, S' \setminus \{g\}) \\ P(\emptyset, S') &= 1 \end{aligned}$$

Once values $U(\kappa, S, S')$ and $P(S, S')$ have been computed, we use the fact that:

$$\text{eu}(\kappa, \tau) = \sum_{S' \subseteq G, |S'|=\tau} P(S', G) U(\kappa, G, S').$$

To do the computation, we use memoization to store the different values $U(\kappa, S, S')$ and $P(S, S')$ (which represents $O(m4^m)$ values), avoid redundant computation, and obtain the desired time complexity. □

Proposition 11. *If $\Psi = \text{PL}_\nu$, then all values $\text{eu}(\kappa, \tau)$ can be computed in time $O(m^{2\rho} \text{Poly}(m))$.*

Proof. Let $\bar{m} = (m_1, m_2, \dots, m_\rho)$ be a vector representing a set containing m_i goods of value ν_i . For $i \in [\rho]$, we denote by $\text{ft}(i, \bar{m}) = \nu_i m_i / (\sum_{j \in [\rho]} m_j \nu_j)$, the probability that a good with parameter ν_i is ranked first among the elements of the set represented by $\bar{m}$ according to PL_ν . We further define $\bar{m}[-i]$ as the vector defined as $\bar{m}[-i]_i = \bar{m}_i - 1$ and $\bar{m}[-i]_j = \bar{m}_j$ for $j \in [\rho] \setminus \{i\}$ and $\text{sum}(\bar{m}) = \sum_{i=1}^\rho m_i$.

The proof relies on recursive equations. Let us consider the following setting. The picker under consideration should pick κ goods within a set S of goods occupying the $|S|$ last positions of her ranking. This set is represented by a vector $\bar{m}_S$. Goods in $G \setminus S$, occupy the top $m - |S|$ ranks in her ranking and have already been picked, either by her or by other agents. Moreover, the set $S' \subseteq S$ have already been picked by previous pickers. The set S' is represented by a vector $\bar{m}_{S'}$ such that $\bar{m}_{S'} \leq \bar{m}_S$. We are interested in computing the expected utility $U(\kappa, \bar{m}_S, \bar{m}_{S'})$ of the κ picks of the agent in such a situation. We argue that $U(\kappa, \bar{m}_S, \bar{m}_{S'})$ satisfies the following recursive equations:

$$U(\kappa, \bar{m}_S, \bar{m}_{S'}) = \sum_{i \in [\rho]} \text{ft}(i, \bar{m}_S) \left(\frac{m'_i}{m_i} U(\kappa, \bar{m}_S[-i], \bar{m}_{S'}[-i]) \right. \\ \left. + \frac{m_i - m'_i}{m_i} (s_{m - \text{sum}(\bar{m}_S) + 1} + U(\kappa - 1, \bar{m}_S[-i], \bar{m}_{S'})) \right) \quad (6)$$

$$U(0, \bar{m}_S, \bar{m}_{S'}) = 0 \quad \forall S, S' \quad (7)$$

$$U(\kappa, \bar{m}_S, \bar{m}_\emptyset) = \sum_{i=1}^\kappa s_{m - \text{sum}(\bar{m}_S) + i} \quad \forall \kappa, S \quad (8)$$

In Equation 6, we consider all possible goods which could be placed at rank $m - |S| + 1 = m - \text{sum}(\bar{m}_S) + 1$, considering only their value in ν . This occurs for a good g with parameter ν_i with probability $\text{ft}(i, \bar{m}_S)$. Let us assume that this good has indeed value ν_i . This good is in (resp. out of) S' with probability m'_i/m_i (resp. $(m_i - m'_i)/m_i$). If $g \in S'$, this good has already been picked and the agent still has to pick κ goods within the goods in $S \setminus \{g\}$ which are ranked in last position, hence we consider $U(\kappa, \bar{m}_S[-i], \bar{m}_{S'}[-i])$. If $g \notin S'$, this good is picked by the agent leading to a utility $s_{m - |S| - 1}$ and the agent still has to pick $\kappa - 1$ goods within the goods in $S \setminus \{g\}$ which are ranked in last position, hence we consider $U(\kappa - 1, \bar{m}_S[-i], \bar{m}_{S'})$. Equations 7 and 8 provide the base cases.

Next, we consider the probability $P(\bar{m}, \bar{m}')$ that a set of good S with vector $\bar{m}_S = \bar{m}$ are ranked in the top $\text{sum}(\bar{m})$ positions among a set S' of goods with vector $\bar{m}_{S'} = \bar{m}'$. $P(\bar{m}, \bar{m}')$ trivially satisfies the following recursive equation.

$$P(\bar{m}, \bar{m}') = \sum_{i \in [\rho], \bar{m}_i \neq 0} \text{ft}(i, \bar{m}') P(\bar{m}[-i], \bar{m}'[-i])$$

$$P(\bar{m}_\emptyset, \bar{m}') = 1$$

Once values $U(\kappa, \bar{m}, \bar{m}')$ and $P(\bar{m}, \bar{m}')$ have been computed, we use the fact that:

$$\text{eu}(\kappa, \tau) = \sum_{\bar{m}' \leq \bar{m}_G, \text{sum}(\bar{m}') = \tau} P(\bar{m}', \bar{m}_G) U(\kappa, \bar{m}_G, \bar{m}').$$

To do the computation, we use memoization to store the different values $U(\kappa, \bar{m}_S, \bar{m}_{S'})$ and $P(\bar{m}, \bar{m}')$ (which represents $O(m \times m^{2\rho})$ values), avoid redundant computation, and obtain the desired time complexity. $\square$

D Omitted Proofs of Section 6

Figure 3 displays our results using the lexicographic scoring vector (where $s_i = 2^{m-i}$), showing the proportion of utility (left-hand side) and goods (right-hand side) obtained for $n = 5$ and increasing the

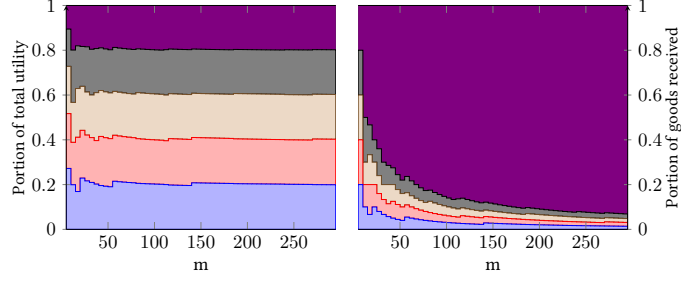


Figure 3: Portion of the total utility (plot on the left) and of goods (right) received by each of 5 agents with m increasing from 5 to 300 in steps of 5. Maximizing ESW and using the lexicographic scoring vector and FI.

number of goods m from 5 to 300 in steps of 5. The IC model was used and we optimized ESW. We see that because of the lexicographic scoring vector, almost all goods are given to the last picker to compensate for this disadvantageous position.

We now provide a property of optimal CSDs when maximizing ESW and for a large number of goods.

Proposition 13. *Assume there are n agents (n being fixed), and let $\mathcal{K}_m^*$ be the set of allocation vectors maximizing $SW_{IC}^E(\mathbf{k})$ when there are m items. Then, if one uses the lexicographic (resp. Borda) scoring vector, then for any value $\epsilon > 0$, there exists a value M (dependent on n) such that $m \geq M$ implies that $(\max_{a \in \mathcal{A}} EU_{IC}^{\mathbf{k}}(a) - \min_{a \in \mathcal{A}} EU_{IC}^{\mathbf{k}}(a)) / (\sum_{a \in \mathcal{A}} EU_{IC}^{\mathbf{k}}(a)) < \epsilon$ for any (resp. an element) $\mathbf{k} \in \mathcal{K}_m^*$.*

Proof. We treat the cases of the Borda and lexicographic scoring vectors using two different proofs.

The lexicographic case. Let us fix n the number of agents, a value $\epsilon > 0$, and l an integer such that $\epsilon/2 \geq 1/2^l$. As $m (> nl)$ increases, we can ensure with a probability tending towards one that each agent receives her l preferred items. Indeed, because of the independence and uniformity assumptions of the IC model, the probability that these sets of items do not intersect tends towards one. Hence, for $\epsilon > 0$, there exists a value M such that if $m > M$, this event (the sets being disjoint) occurs with probability $1 - \epsilon/2$. Using the lexicographic scoring vector, this event implies that each agent will receive a proportion greater than or equal to $(1 - 1/2^l)$ of the total utility she gives to items, i.e., $\sum_{j=1}^m s_m = 2^m - 1$. To sum up, if $m \geq M$, we can ensure that each agent receives an expected utility greater than:

$$\begin{aligned} (1 - \epsilon/2)(1 - 1/2^l)(2^m - 1) &\geq (1 - \epsilon/2)^2(2^m - 1) \\ &\geq (1 - \epsilon)(2^m - 1). \end{aligned}$$

Moreover, note that this expected utility is upper bounded by $\sum_{j=1}^m s_m = 2^m - 1$ and that $\sum_{a \in \mathcal{A}} EU_{IC}^{\mathbf{k}}(a)$ is lower bounded by $\sum_{j=1}^m s_m = 2^m - 1$. Hence, for any $\mathbf{k} \in \mathcal{K}_m^*$:

$$\begin{aligned} \frac{|EU_{IC}^{\mathbf{k}}(a_i) - EU_{IC}^{\mathbf{k}}(a_j)|}{\sum_{a \in \mathcal{A}} EU_{IC}^{\mathbf{k}}(a)} &\leq \frac{|EU_{IC}^{\mathbf{k}}(a_i) - EU_{IC}^{\mathbf{k}}(a_j)|}{2^m - 1} \\ &\leq \frac{(2^m - 1) - (1 - \epsilon)(2^m - 1)}{2^m - 1} \\ &\leq \epsilon. \end{aligned}$$

The Borda case. We wish to show that for any ϵ , there always exists an optimal egalitarian solution for which the difference in portions of total utility assigned to any two different agents is smaller than ϵ when m is high enough. However, instead of reasoning on the portion of total expected utility received by an agent, we will work on a proxy, denoted by $\tilde{P}^{\mathbf{k}}(a) = 2EU_{IC}^{\mathbf{k}}(a)/m(m+1)$. Note that, compared

to $EU_{\text{IC}}^{\mathbf{k}}(a)/(\sum_{a \in \mathcal{A}} EU_{\text{IC}}^{\mathbf{k}}(a))$, $\tilde{P}^{\mathbf{k}}(a)$ replaces the expected total utility received by the n agents by a lower bound on it given by $\sum_{j=1}^m s_j = m(m+1)/2$.

Let $\epsilon > 0$ be a positive value. We set $\epsilon' = \epsilon/n$, and $m = 2/\epsilon'$. We now show by induction on l the following: For any value $\tau \in \{i/m, i \in [m]_0\}$, there exists a solution $\mathbf{k}$ maximizing $\min_{a \in \{a_1, \dots, a_l\}} \tilde{P}^{\mathbf{k}}(a)$ under the constraint that a proportion τ of the picks are assigned to the l first pickers which ensures that $|\tilde{P}^{\mathbf{k}}(a_i) - \tilde{P}^{\mathbf{k}}(a_j)| \leq l\epsilon'$ for any $i, j \in [l]^2$. We denote by $\mathcal{P}_l^\tau$ the previous optimization problem and Γ_l^τ its optimal value.

The claim is trivially true for $l = 1$. Assume, it is true for $l \geq 1$. We seek a solution maximizing $\min_{a \in \{a_1, \dots, a_{l+1}\}} \tilde{P}^{\mathbf{k}}(a)$ given that they receive a proportion τ of the items. The l first agents will receive a proportion $\tau' \in [0, \tau]$ of the items, and we can assume wlog that the allocation to the l first agents is the one maximizing $\mathcal{P}_l^{\tau'}$ insuring our inductive property. Note that if τ' increases (resp. decrease) by $1/m$, this may only increase (resp. decrease) $\Gamma_l^{\tau'+1/m}$ by $2/m$, i.e., $\Gamma_l^{\tau'+1/m} \leq \Gamma_l^{\tau'} + 2/m$ and decrease (resp. increase) the $\tilde{P}^{\mathbf{k}}(a_{l+1})$ value of the $(l+1)^{\text{th}}$ picker by $2/m$ and that Γ_l^τ (resp. $\tilde{P}^{\mathbf{k}}(a_{l+1})$) is non-decreasing (resp. non-increasing) in τ . Hence, by adjusting the value of τ' , we can ensure that there exists a solution $\mathbf{k}$ optimal for $\mathcal{P}_{l+1}^\tau$ and such that $|\tilde{P}^{\mathbf{k}}(a_{l+1}) - \Gamma_l^{\tau'}| \leq \epsilon'$. Therefore, by the inductive property, $|\tilde{P}^{\mathbf{k}}(a_i) - \tilde{P}^{\mathbf{k}}(a_j)| \leq (l+1)\epsilon'$ for any $i, j \in [l+1]^2$. This proves the inductive property. Using $l = n$, and $\tau = 1$, we obtain the claimed result as

$$|\tilde{P}^{\mathbf{k}}(a_i) - \tilde{P}^{\mathbf{k}}(a_j)| \leq \frac{|EU_{\text{IC}}^{\mathbf{k}}(a_i) - EU_{\text{IC}}^{\mathbf{k}}(a_j)|}{\sum_{a \in \mathcal{A}} EU_{\text{IC}}^{\mathbf{k}}(a)}.$$

□

A note on computation times All tests were run with Python 3.10.12 on a personal computer with Ubuntu 22.04.4 LTS, 8 Intel(R) Core(TM) CPU i7-1185G7 3.00GHz cores and 32 GB RAM. With $n = 5$ and $m = 70$, and a sample size of 1000 profiles, the computation of the optimal allocation using Equation 1, given that prefix independence is satisfied, takes approximately 50 seconds. By contrast, the use of Algorithm 1 (GreedyESW) reduces the computation time to approximately 7 seconds. Furthermore, when applicable, the exact computation using Equation 2 is highly efficient, requiring only approximately 0.07 seconds.

E Finding a suitable scoring vector

A question that has been overlooked until now⁷ is, where does the vector of scores come from? In order to address it we suggest, and test, the following methodology. For the specific domain at hand, prepare a questionnaire where some users, considered representative of the population of users, are presented a set of items: for instance, if the problem is about allocating time slots for using a tennis court, users are presented several time slots. Once this scoring vector has been elicited, it is used to determine optimal CSDs, which can be applied many times, with different sets of users. A similar method has been used for voting by Boutilier et al. [10] (see Section 5.6).

We designed an online experiment. To each user taking part in it, we present 12 ice-cream flavours uniformly selected among 62 possible and elicit their utility on a scale [0,100].

We first ask the user to tell which is their preferred flavour (PF) among the 12, and we tell them that the value for PF is fixed to 100. Then, for each flavour F (including PF), we present the user a slider, with which they indicate the value of F between 0 and 100.

⁷Not only in this paper but also in previous papers on fair division who also use scoring vectors (e.g [7]).

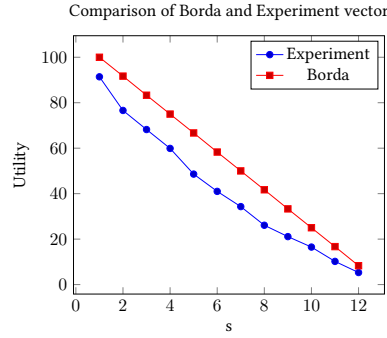


Figure 4: Comparison of the Borda scoring score (rescaled to $[0,100]$) with the scoring vector obtained through our experiment by taking the expectation of the participants' answers.

Once all vectors are collected, we rearrange them non-increasingly. Then all vectors are averaged among all users. We obtain a scoring vector $s = (s_1, \dots, s_{12})$: s_i is the average value, among all users, of their i th most preferred item.

We had 54 participants. Screenshots of the experiment, information on how consents and the data were collected, as well as the list of the 54 gathered vectors are included in the Appendix. Their average is

$$\begin{array}{llll} s_1 = 91.4 & s_2 = 76.6 & s_3 = 68.2 & s_4 = 56.9 \\ s_5 = 48.6 & s_6 = 41 & s_7 = 34.3 & s_8 = 26.1 \\ s_9 = 21.1 & s_{10} = 16.5 & s_{11} = 10.2 & s_{12} = 5.3 \end{array}$$

Figure 4 shows how this vector compares to the Borda vector (rescaled such that the score of one's preferred flavor is 100)⁸.

The ice-cream experiment

We provide additional information on the experiment used for estimating an adequate scoring vector for the application of allocating ice creams with different flavors.

The list of all ice-cream flavors used for the experiment is the following one :

[Kiwi, Litchi, Mango, Mandarin, Melon, Mirabelle, Blackberry, Blueberry, Orange, Blood orange, Apricot, Pineapple, Banana, Lemon, Lime, Cherry, Cassis, Raspberry, Coco, Fig, Strawberry, Passion fruit, Pear, Rhubarb, Grapefruit, Honey-Pine nuts, Tiramisu, Chocolate ginger, Tagada strawberry, Nougat, Speculoos, Coffee, Milk jam, Pistachio, Licorice, Lavender, Caramel, Dragibus, Avocado, Chewing gum, Olive, Chili chocolate, Tomato-Basil, Cinnamon, White chocolate, Chocolate, Almond, Poppy, Cookies, Gingerbread, Cactus, Beer, Oreo, Nutella, Vanilla, Candy floss, Rum-Raisin, Pumpkin, Chestnut, Wild pollen, Rice pudding, Salted butter caramel].

Each participant was presented a subset of 12 of these flavors, sampled uniformly at random with a seed set according to the `Math.Random()` Javascript method. Note that this method sets the seed for simulating randomness in a way that depends on the browser of the user.

We tested two different ways of explaining the experiment to participants. Indeed, we wanted to evaluate the impact of describing the experiment in one way or another.

⁸Note that by averaging we lost some interesting information about variance: some users have rapidly decreasing, and some others slowly decreasing valuations. Moreover, note that the highest score of the averaged vector is not 100 as several users did not respect this constraint.

- For half of the participants⁹, we presented the scores assigned to the ice-cream flavours as Von Neumann-Morgenstein utility values. We first ask the user to tell which is their preferred flavour (PF) among the 12, and we tell them that the value for PF is fixed to 100. Then, for each flavour F (including PF), we present the user a slider, with which they will indicate the value of F between 0 and 100. They are told that they can interpret the chosen value V as the exact point where they are indifferent between receiving PF with probability $\frac{V}{100}$ and nothing with probability $1 - \frac{V}{100}$, or receiving F for sure. A screenshot of this process is reported in Figure 6.
- For the other half of participants, we tested simpler directives. As done previously, we explain to users that the value for their preferred flavour should be fixed to 100. Then, for each flavour F , we present the user a slider, with which they should indicate the value of F between 0 and 100, without mentioning any probabilistic interpretation for these values. A screenshot of this process is reported in Figure 7.

The 27 scoring vectors of the participants who followed the first (resp. second) directives are displayed on Table 2 (resp. Table 3).

A comparison of the resulting averaged scoring vectors with the Borda scoring vector is provided in Figure 5. We observe that the two averaged scoring vectors obtained for each set of directives have a similar shape with one being slightly above the other one. Indeed, more participants following the simpler directives did not follow the constraint that their preferred flavour among the 12 should receive a value of 100, leading to smaller values in the averaged scoring vector. This is probably due to an ambiguity about the fact that the preferred flavour of the participant should be understood as the most preferred one among the ones which are presented. While this finding points out a possible improvement for our experiment, we insist on the fact that it should be understood here as a proof of concept illustrating the feasibility of such an approach to estimate a relevant scoring vector for the domain at hand. As the results were similar for the two ways of describing the experiment, we decided to merge the two list of vectors for plotting the Figure 4 presented in the main document of the submission.

Collect of Consents and Data The data which was collected was completely anonymous; indeed, we did not collect any piece of information about the participants besides the scores assigned to the ice-cream flavors. Moreover, we checked with the ethic committee of one of the authors' university that the experiment complies with the data protection regulation. Last, as illustrated in Figure 9, all the participants to the experiment had to check a box, confirming that they agreed that their answers would be stored and used for research purposes.

⁹In fact, each participant had a probability 0.5 of getting one set of directives or the other. Luckily this procedure split the set of participants in two sets of equal sizes.

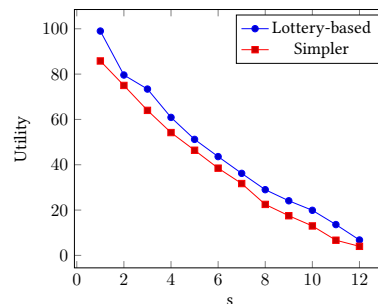


Figure 5: Scoring vectors obtained through our experiment. The scoring vector “Lottery-based” (resp. “Simpler”) was obtained by averaging the answers of the participants receiving the directives which mentioned (resp. did not mention) a probabilistic interpretation for values assigned to ice-cream flavours.

Which ice-cream flavor do you love ?

The aim of this survey is to observe the distribution of the intensity of your preferences in an interesting topic for many: ice-cream flavors. If you like neither ice-creams nor sorbets, this survey will probably bore you and we would suggest you not to do it.

First question : from the following twelve flavors:

Which one is your favorite?
Salted butter caramel

We will ask you to express your preferences among the other flavors, knowing the following constraint is imposed : Salted butter caramel worth 100 points. One way to evaluate the number of points to associate to one flavor is the following : x is the number such that you are indifferent between being sure of having this flavor, and having x% chance of getting your favorite flavor and (100-x)% chance of getting nothing at all.

For example :

- * You are indifferent between being sure of having Tiramisu, and having 50% chance of getting Salted butter caramel and 50% chance of getting nothing at all, then you can give 50 points to Tiramisu.
- * You are indifferent between being sure of having Pistachio, and having 30% chance of getting Salted butter caramel and 70% chance of getting nothing at all, then you can give 30 points to Pistachio.
- * You are indifferent between being sure of having Raspberry, and having 90% chance of getting Salted butter caramel and 10% chance of getting nothing at all, then you can give 90 points to Raspberry.

Figure 6: Screenshot of our questionnaire. Directives to the user mentioning a probabilistic interpretation for the scores assigned to ice-cream flavors.

Which ice-cream flavor do you love ?

The aim of this survey is to observe the distribution of the intensity of your preferences in an interesting topic for many: ice-cream flavors. If you like neither ice-creams nor sorbets, this survey will probably bore you and we would suggest you not to do it.

Give each flavour a score on a scale 0-100.

100 is the score that you give to your preferred flavour.

0 is the score you give to a flavour what you would definitely not eat.

The intermediate scores between 0 and 100 have to be interpreted as increasing taste for the flavour; for instance you may interpret 50 as "I like it half as much as my preferred flavour"

Figure 7: Screenshot of our questionnaire. Directives to the user not mentioning a probabilistic interpretation for the scores assigned to ice-cream flavors.

Licorice		5
Tiramisu		80
Pistachio		20
Raspberry		30
Gingerbread		10
Kiwi		15
Caramel		70
Rum - Raisin		0
Chocolate ginger		40
Salted butter caramel		100
Rhubarb		40
Dragibus		5

RESET

SEND

Figure 8: Screenshot of our questionnaire. The sliders make it possible for the participants to assign a value to each flavor.



By checking this box I agree that my preferences will be stored anonymously for two years and may be used for research purposes

SEND

Figure 9: Screenshot of our questionnaire. Participants had to check a box, agreeing that their answers could be used for research purposes.

F The price of the assignment of agents to positions

By abuse of notation, we may also use notation $SW_{\mathbf{P}}^x(\mathbf{k})$ (for $x \in \{U, E, N\}$) when Ψ is the degenerate probability distribution for which profile $\mathbf{P}$ occurs with probability 1.

So far, we have considered probability distributions over profiles that treat all agents in an interchangeable way. Hence, deciding who should be agent a_1 and pick first, who should be agent a_2 and pick second *et caetera*, has no impact on ex ante social welfare. What about its impact on ex post social welfare? We now study this impact, in terms of loss of social welfare between the best possible assignment and the worst possible assignments of agents to positions.

Let S_n denotes the set of permutations of $[n]$. Given $\pi \in S_n$ and $\mathbf{P}$ a preference profile, $\mathbf{P}_\pi$ denotes the preference profile obtained from $\mathbf{P}$ by permuting agents rankings according to π .

Definition 2. Given a CSD with vector $\mathbf{k}$ and a preference profile $\mathbf{P}$, the utilitarian, egalitarian and Nash price of assignment of agents to positions, denoted by $\mathcal{P}_{AtoP}^U$, $\mathcal{P}_{AtoP}^E$, and $\mathcal{P}_{AtoP}^N$ respectively, are defined by:

$$\begin{aligned}\mathcal{P}_{AtoP}^U &= \frac{\max_{\pi \in S_n} SW_{\mathbf{P}_\pi}^U(\mathbf{k})}{\min_{\pi \in S_n} SW_{\mathbf{P}_\pi}^U(\mathbf{k})}, \\ \mathcal{P}_{AtoP}^E &= \frac{\max_{\pi \in S_n} SW_{\mathbf{P}_\pi}^E(\mathbf{k})}{\min_{\pi \in S_n} SW_{\mathbf{P}_\pi}^E(\mathbf{k})}, \\ \mathcal{P}_{AtoP}^N &= \frac{\max_{\pi \in S_n} SW_{\mathbf{P}_\pi}^N(\mathbf{k})}{\min_{\pi \in S_n} SW_{\mathbf{P}_\pi}^N(\mathbf{k})}.\end{aligned}$$

We make the two following easy observations which hold whatever the notion of social welfare which is used:

1. The worst social welfare that can be obtained when allocating resources using a CSD is obtained when for all $j \in [m]$, the j^{th} good that is picked by an agent is her j^{th} preferred good. In particular, this results in a utilitarian social welfare of $\sum_{j=1}^m s_j$.
2. The best social welfare that can be obtained when allocating resources using a picking sequence (not necessarily non-interleaving) where a_i picks k_i goods for all $i \in [n]$ is obtained when each agent picks her k_i preferred goods. This results in a social welfare of $\star_{i=1}^n \sum_{j=1}^{k_i} s_j$, where $\star = \sum, \prod$ or $\min$ depending on the chosen notion of social welfare.

Hence, upper bounds on $\mathcal{P}_{AtoP}^U$, $\mathcal{P}_{AtoP}^E$, and $\mathcal{P}_{AtoP}^N$ are given by:

$$\frac{\sum_{i=1}^n \sum_{j=1}^{k_i} s_j}{\sum_{j=1}^m s_j}, \frac{\sum_{j=1}^{k_{\min}} s_j}{\min_{i \in [n]} \sum_{j=c_i+1}^{c_i+k_i} s_j}, \text{ and } \frac{\prod_{i=1}^n (\sum_{j=1}^{k_i} s_j)}{\prod_{i=1}^n (\sum_{j=c_i+1}^{c_i+k_i} s_j)}.$$

where $c_i = \sum_{l < i} k_l$ and $k_{\min} = \min\{k_i | i \in [n]\}$.

We show that there exists a preference profile and a CSD, such that this bound is closely matched.

Proposition 14. Assume $m = d \times n$ with $d \in \mathbb{N}^*$. There exists a preference profile and a CSD such that:

$$\begin{aligned}
\frac{(n-1) \sum_{j=1}^d s_j + \sum_{j=d+1}^{2d} s_j}{\sum_{j=1}^m s_j} &\leq \mathcal{P}_{AtoP}^U \leq \frac{n \sum_{j=1}^d s_j}{\sum_{j=1}^m s_j}, \\
\frac{\sum_{j=d+1}^{2d} s_j}{\sum_{j=(n-1)d+1}^m s_j} &\leq \mathcal{P}_{AtoP}^E \leq \frac{\sum_{j=1}^d s_j}{\sum_{j=(n-1)d+1}^m s_j}, \\
\frac{\left(\sum_{j=1}^d s_j \right)^{(n-1)} \times \sum_{j=d+1}^{2d} s_j}{\prod_{i=1}^n \sum_{j=(i-1)d+1}^{id} s_j} &\leq \mathcal{P}_{AtoP}^N \leq \frac{\left(\sum_{j=1}^d s_j \right)^n}{\prod_{i=1}^n \sum_{j=(i-1)d+1}^{id} s_j}.
\end{aligned}$$

Proof. Consider the CSD with vector $\mathbf{k}$ such that $k_1 = k_2 = \dots = k_n = d$. Given the previously defined vector $\mathbf{k}$, we build a preference profile $\mathbf{P}$ as follows. Let S_i be a set of d goods for $i \in \{2, \dots, n\}$ such that $S_i \cap S_j = \emptyset$ for all $i \neq j \in \{2, \dots, n\}$, $S_1 = S_2$, and $S_{n+1} = \mathcal{G} \setminus \bigcup_{i=2}^n S_i$. We let S_i be the preferred goods of agent a_i for $i \in [n]$. Additionally, each agent a_j with $j \in [n]$ prefers any good in S_s to any good in S_t if $s < t$ and $s, t \in [n] \setminus \{j\}$. Lastly, we assume each agent a_j with $j \in \{2, \dots, n\}$ ranks goods in S_{n+1} last and that a_1 ranks these goods just after the ones in S_1 .

One can easily check that for all $i \in [n]$, $U_{\mathbf{P}}^{\mathbf{k}}(a_i) = \sum_{j=(i-1)d+1}^{id} s_j$. If we otherwise consider the permutation $\pi = (2, 3, \dots, n, 1)$, we obtain that the first $n-1$ agents get utility $\sum_{j=1}^d s_j$ while the last agent of the sequence (agent a_1) gets utility $\sum_{j=d+1}^{2d} s_j$. The three lower bounds follow. $\square$

To give an example, let us take the Borda scoring vector. The price of assignment of agents to positions is reasonable for utilitarianism, as it tends to 2 when m grows; it is much larger for egalitarianism (it is in the order of m when m grows), with Nash being even worse (especially if n grows and d is kept constant, $\mathcal{P}_{AtoP}^N$ explodes). As a consequence, optimizing the CSD gives good *ex ante* fairness guarantees, but much less *ex post* fairness guarantees (it is consistent with the well-known general observation, in fair division, that *ex post* fairness guarantees are harder to obtain than *ex post* fairness guarantees).

G Examples

All examples were computed using a sample of $k = 1000$ preference profiles drawn from the distribution of the table. Note that with FC and utilitarianism, every allocation is optimal.

Table 4: Results for FC

(n,m)	SW	Best Policy	Best Utilities
(2,4)	ESW	(1, 3)	[4.0, 6.0]
	NSW	(1, 3)	[4.0, 6.0]
	USW	(4, 0)	[10.0, 0.0]
(2,7)	ESW	(2, 5)	[13.0, 15.0]
	NSW	(2, 5)	[13.0, 15.0]
	USW	(7, 0)	[28.0, 0.0]
(2,10)	ESW	(3, 7)	[27.0, 28.0]
	NSW	(3, 7)	[27.0, 28.0]
	USW	(10, 0)	[55.0, 0.0]
(3,4)	ESW	(1, 1, 2)	[4.0, 3.0, 3.0]
	NSW	(1, 1, 2)	[4.0, 3.0, 3.0]
	USW	(4, 0, 0)	[10.0, 0.0, 0.0]
(3,7)	ESW	(1, 2, 4)	[7.0, 11.0, 10.0]
	NSW	(1, 2, 4)	[7.0, 11.0, 10.0]
	USW	(7, 0, 0)	[28.0, 0.0, 0.0]
(3,10)	ESW	(2, 3, 5)	[19.0, 21.0, 15.0]
	NSW	(2, 3, 5)	[19.0, 21.0, 15.0]
	USW	(10, 0, 0)	[55.0, 0.0, 0.0]
(4,4)	ESW	(1, 1, 1, 1)	[4.0, 3.0, 2.0, 1.0]
	NSW	(1, 1, 1, 1)	[4.0, 3.0, 2.0, 1.0]
	USW	(4, 0, 0, 0)	[10.0, 0.0, 0.0, 0.0]
(4,7)	ESW	(1, 1, 2, 3)	[7.0, 6.0, 9.0, 6.0]
	NSW	(1, 1, 2, 3)	[7.0, 6.0, 9.0, 6.0]
	USW	(7, 0, 0, 0)	[28.0, 0.0, 0.0, 0.0]
(4,10)	ESW	(2, 2, 2, 4)	[19.0, 15.0, 11.0, 10.0]
	NSW	(1, 2, 2, 5)	[10.0, 17.0, 13.0, 15.0]
	USW	(10, 0, 0, 0)	[55.0, 0.0, 0.0, 0.0]

Table 5: Results for IC

(n,m)	SW	Best Policy	Best Utilities
(2,4)	ESW	(2, 2)	[7.0, 4.97]
	NSW	(2, 2)	[7.0, 4.97]
	USW	(2, 2)	[7.0, 4.97]
(2,7)	ESW	(3, 4)	[18.0, 16.02]
	NSW	(3, 4)	[18.0, 16.02]
	USW	(4, 3)	[22.0, 12.04]
(2,10)	ESW	(4, 6)	[34.0, 33.05]
	NSW	(4, 6)	[34.0, 33.05]
	USW	(5, 5)	[40.0, 27.59]
(3,4)	ESW	(1, 1, 2)	[4.0, 3.75, 4.97]
	NSW	(1, 1, 2)	[4.0, 3.75, 4.97]
	USW	(2, 1, 1)	[7.0, 3.34, 2.45]
(3,7)	ESW	(2, 2, 3)	[13.0, 12.0, 11.85]
	NSW	(2, 2, 3)	[13.0, 12.0, 11.85]
	USW	(3, 2, 2)	[18.0, 11.19, 7.96]
(3,10)	ESW	(3, 3, 4)	[27.0, 24.86, 22.11]
	NSW	(3, 3, 4)	[27.0, 24.86, 22.11]
	USW	(4, 3, 3)	[34.0, 23.76, 16.64]
(4,4)	ESW	(1, 1, 1, 1)	[4.0, 3.74, 3.35, 2.46]
	NSW	(1, 1, 1, 1)	[4.0, 3.74, 3.35, 2.46]
	USW	(1, 1, 1, 1)	[4.0, 3.74, 3.35, 2.46]
(4,7)	ESW	(1, 2, 2, 2)	[7.0, 12.58, 11.16, 7.78]
	NSW	(1, 2, 2, 2)	[7.0, 12.58, 11.16, 7.78]
	USW	(2, 2, 2, 1)	[13.0, 11.99, 9.97, 3.9]
(4,10)	ESW	(2, 2, 2, 4)	[19.0, 18.34, 17.35, 22.15]
	NSW	(2, 2, 3, 3)	[19.0, 18.34, 23.59, 16.58]
	USW	(3, 3, 2, 2)	[27.0, 24.79, 15.4, 10.92]

Table 6: Results for PL_{ν} with $\nu = (1.1^m, 1.1^{m-1}, \dots, 1.1^1)$

(n,m)	SW	Best Policy	Best Utilities
(2,4)	ESW	(2, 2)	[7.0, 4.96]
	NSW	(2, 2)	[7.0, 4.96]
	USW	(2, 2)	[7.0, 4.96]
(2,7)	ESW	(3, 4)	[18.0, 15.9]
	NSW	(3, 4)	[18.0, 15.9]
	USW	(4, 3)	[22.0, 11.92]
(2,10)	ESW	(4, 6)	[34.0, 32.37]
	NSW	(4, 6)	[34.0, 32.37]
	USW	(5, 5)	[40.0, 26.74]
(3,4)	ESW	(1, 1, 2)	[4.0, 3.74, 5.06]
	NSW	(1, 1, 2)	[4.0, 3.74, 5.06]
	USW	(2, 1, 1)	[7.0, 3.33, 2.48]
(3,7)	ESW	(2, 2, 3)	[13.0, 11.94, 11.73]
	NSW	(2, 2, 3)	[13.0, 11.94, 11.73]
	USW	(3, 2, 2)	[18.0, 11.08, 7.9]
(3,10)	ESW	(3, 3, 4)	[27.0, 24.66, 21.45]
	NSW	(3, 3, 4)	[27.0, 24.66, 21.45]
	USW	(4, 3, 3)	[34.0, 23.31, 16.0]
(4,4)	ESW	(1, 1, 1, 1)	[4.0, 3.76, 3.33, 2.47]
	NSW	(1, 1, 1, 1)	[4.0, 3.76, 3.33, 2.47]
	USW	(1, 1, 1, 1)	[4.0, 3.76, 3.33, 2.47]
(4,7)	ESW	(1, 2, 2, 2)	[7.0, 12.55, 11.17, 8.03]
	NSW	(1, 2, 2, 2)	[7.0, 12.55, 11.17, 8.03]
	USW	(2, 2, 1, 2)	[13.0, 11.98, 5.98, 8.09]
(4,10)	ESW	(2, 2, 2, 4)	[19.0, 18.26, 17.09, 21.33]
	NSW	(2, 2, 2, 4)	[19.0, 18.26, 17.09, 21.33]
	USW	(3, 3, 2, 2)	[27.0, 24.5, 14.89, 10.35]

Table 7: Results for $M11_{\phi,\mu}$ with $\phi = 0.8$

(n,m)	SW	Best Policy	Best Utilities
(2,4)	ESW	(2, 2)	[7.0, 5.03]
	NSW	(2, 2)	[7.0, 5.03]
	USW	(2, 2)	[7.0, 5.03]
(2,7)	ESW	(3, 4)	[18.0, 15.45]
	NSW	(3, 4)	[18.0, 15.45]
	USW	(4, 3)	[22.0, 11.51]
(2,10)	ESW	(4, 6)	[34.0, 31.33]
	NSW	(4, 6)	[34.0, 31.33]
	USW	(5, 5)	[40.0, 25.79]
(3,4)	ESW	(1, 1, 2)	[4.0, 3.75, 4.98]
	NSW	(1, 1, 2)	[4.0, 3.75, 4.98]
	USW	(2, 1, 1)	[7.0, 3.32, 2.43]
(3,7)	ESW	(2, 2, 3)	[13.0, 11.86, 11.29]
	NSW	(2, 2, 3)	[13.0, 11.86, 11.29]
	USW	(3, 2, 2)	[18.0, 10.95, 7.36]
(3,10)	ESW	(3, 3, 4)	[27.0, 24.06, 20.22]
	NSW	(3, 3, 4)	[27.0, 24.06, 20.22]
	USW	(4, 3, 3)	[34.0, 22.61, 14.93]
(4,4)	ESW	(1, 1, 1, 1)	[4.0, 3.73, 3.32, 2.42]
	NSW	(1, 1, 1, 1)	[4.0, 3.73, 3.32, 2.42]
	USW	(1, 1, 1, 1)	[4.0, 3.73, 3.32, 2.42]
(4,7)	ESW	(1, 2, 2, 2)	[7.0, 12.51, 10.84, 7.68]
	NSW	(1, 2, 2, 2)	[7.0, 12.51, 10.84, 7.68]
	USW	(2, 2, 1, 2)	[13.0, 11.84, 5.77, 7.58]
(4,10)	ESW	(2, 2, 2, 4)	[19.0, 18.11, 16.69, 19.98]
	NSW	(2, 2, 2, 4)	[19.0, 18.11, 16.69, 19.98]
	USW	(3, 3, 2, 2)	[27.0, 24.02, 14.41, 9.49]

100	95	90	30	20	0	0	0	0	0	0	0
100	12	11	0	0	0	0	0	0	0	0	0
100	100	100	100	93	90	73	59	52	48	10	6
100	93	93	90	63	50	33	13	10	10	0	0
100	100	100	70	50	50	40	40	30	30	30	25
91	75	59	17	17	15	12	5	0	0	0	0
99	90	86	80	61	60	32	30	24	21	18	17
100	95	90	80	60	50	10	10	5	5	0	0
100	95	90	80	70	60	60	50	45	30	25	10
100	95	85	80	80	70	60	52	40	30	20	10
100	80	75	60	50	15	10	5	5	3	1	0
100	40	30	20	20	10	10	5	5	5	2	0
100	90	79	76	76	66	64	51	32	26	18	10
100	100	90	90	90	89	88	86	84	80	75	59
100	70	60	60	40	30	20	15	10	10	10	10
95	93	92	54	50	42	37	27	25	23	20	0
100	30	30	30	20	20	20	10	10	0	0	0
100	98	94	94	94	92	92	92	90	80	68	10
100	43	36	35	31	27	18	17	14	8	5	0
100	85	80	70	25	15	15	10	5	5	2	0
100	70	70	50	50	30	30	0	0	0	0	0
100	82	81	76	75	74	53	31	21	14	0	0
100	95	79	73	56	50	48	45	33	15	10	9
90	88	81	61	53	50	50	50	50	50	20	4
100	75	53	47	44	36	30	23	12	5	0	0
100	90	80	51	50	50	50	40	40	20	10	2
100	80	80	70	60	60	60	50	40	30	30	10

Table 2: List of the 27 scoring vectors (one per row) obtained for participants following the directives mentioning a probabilistic interpretation for values assigned to ice-cream flavours.

83	82	80	78	73	70	65	65	57	50	13	12
91	87	76	55	42	41	40	37	36	18	15	10
100	80	70	60	60	50	50	40	40	30	20	10
60	45	40	19	0	0	0	0	0	0	0	0
80	80	73	73	65	59	56	34	0	0	0	0
83	81	78	77	76	69	68	25	13	6	4	3
91	91	80	61	54	51	19	11	10	10	0	0
86	82	74	72	48	6	0	0	0	0	0	0
85	70	13	11	8	7	0	0	0	0	0	0
49	49	40	35	30	20	14	10	10	5	0	0
76	70	68	60	57	55	55	55	40	40	25	10
100	65	0	0	0	0	0	0	0	0	0	0
89	86	82	80	73	71	69	61	60	52	43	34
90	80	72	63	55	54	13	9	6	0	0	0
100	99	80	70	65	60	51	4	3	2	1	0
100	92	88	78	73	62	50	41	40	28	7	3
97	92	77	59	49	43	33	30	18	0	0	0
100	78	70	21	0	0	0	0	0	0	0	0
79	59	52	43	32	20	15	14	12	12	11	10
82	78	78	75	71	49	32	16	8	0	0	0
100	80	80	70	70	60	60	60	50	50	20	0
81	40	30	22	13	0	0	0	0	0	0	0
72	64	61	61	41	34	31	28	26	18	12	11
91	86	74	67	64	62	57	15	10	9	0	0
81	75	70	58	55	27	16	0	0	0	0	0
85	60	59	41	33	31	31	29	15	8	4	2
30	27	24	20	18	14	13	8	6	3	0	0

Table 3: List of the 27 scoring vectors (one per row) obtained for participants following the directives which did not mention a probabilistic interpretation for values assigned to ice-cream flavours.