

Near-Optimal Mechanisms for Resource Allocation Without Monetary Transfers

Moïse Blanchard

Georgia Institute of Technology
mblanchard41@gatech.edu

Patrick Jaillet

Massachusetts Institute of Technology
jaillet@mit.edu

Abstract

We study the problem in which a central planner sequentially allocates a single resource to multiple strategic agents using their utility reports at each round, but without using any monetary transfers. We consider general agent utility distributions and two standard settings: a finite horizon T and an infinite horizon with γ discounts. We provide general tools to characterize the convergence rate between the optimal mechanism for the central planner and the first-best allocation if true agent utilities were available. This heavily depends on the utility distributions, yielding rates anywhere between $1/\sqrt{T}$ and $1/T$ for the finite-horizon setting, and rates faster than $\sqrt{1-\gamma}$, including exponential rates for the infinite-horizon setting as agents are more patient $\gamma \rightarrow 1$. On the algorithmic side, we design mechanisms based on the promised-utility framework to achieve these rates and leverage structure on the utility distributions. Intuitively, the more flexibility the central planner has to reward or penalize any agent while incurring little social welfare cost, the faster the convergence rate. In particular, discrete utility distributions typically yield the slower rates $1/\sqrt{T}$ and $\sqrt{1-\gamma}$, while smooth distributions with density typically yield faster rates $1/T$ (up to logarithmic factors) and $1-\gamma$.

Contents

1	Introduction	2
1.1	Overview of the contributions	4
1.2	Related works	7
1.3	Organization of the paper	8
1.4	Notations	8
2	Problem formulation and preliminaries	9
2.1	Model definitions	9
2.2	Achievable utility sets	10
2.3	The promised utility framework	11
2.4	A convenient choice of promised utility functions	13
2.5	Characterization of scenarios without social welfare gap	14
3	A generic resource allocation mechanism	15
4	A geometric tool for proving lower bounds on the social welfare gap	18

5 Social Welfare Gap Characterizations for the discounted infinite-horizon setting	19
5.1 A universal rate $\mathcal{O}(\sqrt{1-\gamma})$	19
5.2 Sufficient conditions for faster rates $\mathcal{O}(1-\gamma)$	20
5.3 Necessary conditions for faster rates $\mathcal{O}(1-\gamma)$	23
5.4 Convergence rate bounds for general distributions	23
5.5 Mechanism trajectory and efficiency	26
6 From the discounted infinite-horizon setting to the finite-horizon setting	27
6.1 Lower bounds on the social welfare gap	27
6.2 Mechanisms for the finite-horizon setting	28
7 Examples and special cases	29
7.1 Utility distributions with density	29
7.2 Discrete utility distributions	30
8 Conclusion and future work	31
A Proofs from Section 2	36
A.1 Characterization of cases with no social welfare gap	39
B Proofs from Section 3 for the main resource allocation mechanisms	42
C Proofs from Section 4 to prove lower bounds on the social welfare gap	48
D Characterizations in the discounted infinite-horizon setting	56
D.1 A pruning procedure	57
D.2 A universal $\mathcal{O}(\sqrt{1-\gamma})$ convergence rate	60
D.3 Improved rates for smooth full-information regions $\mathcal{U}^*$	60
D.4 General rates for arbitrary utility distributions	65
D.5 Remaining proofs for discrete distributions	71
E Extension to the finite-horizon setting	76

1 Introduction

Resource allocation or auction design is one of the foundational problems in mechanism design and has been thoroughly studied across economics, operations research, and computer science. In this problem, a planner aims to allocate resources to n strategic agents while maximizing some notion of global welfare, using reports/bids from the agents on their utility for the resources. Importantly, agent reports are strategic; hence, the planner only has limited or partial information about the agents' preferences. To achieve positive results, fundamental tools in mechanism design include money transfers between the agents and the planner. The large majority of approaches heavily rely on the use of such monetary payments to align incentives between the strategic users and the central planner. For instance, the celebrated Vickrey-Clarke-Groves (VCG) mechanism [Clarke, 1971, Groves, 1973, Vickrey, 1961] ensures incentive-compatibility while implementing optimal allocations by carefully selecting payment mechanisms. In many

applications, however, money transfers may be impractical. Enforcing payments from beneficiaries to the central planner is typically undesirable when allocating public services: from designing antipoverty programs [Bardhan and Mookherjee, 2005], choosing school assignments [Ashlagi and Shi, 2016, Russell and Ebbert, 2011], allocating scientific equipment [Linden, 2022], to distributing food to food banks. Cloud computing is another application in which memory or computing resources may be scarce; hence, enforcing a scheduling policy is necessary, but monetary transfers may be unpopular/inadequate, especially when allocating internal resources [Ng, 2011].

Resource allocation without monetary transfers. Following these motivations, we consider the setting in which monetary transfers are not allowed and focus on single resource allocation for simplicity. The reports of the agents are strategic and hence form a perfect Bayesian equilibrium. In a full-information setting where the planner has direct access to all agent utilities, the optimal policy is to allocate the resource to an agent with maximum utility, also known as the *first-best* allocation. Unfortunately, prohibiting the use of money severely limits the ability to approximate the first-best allocation in light of negative results arising from Myerson-Satterthwaite’s impossibility theorem [Arrow, 2012, Myerson and Satterthwaite, 1983]. For a single allocation, whenever there are at least 3 alternative allocations, it is impossible to design a non-dictatorial incentive-compatible mechanism [Gibbard, 1973, Satterthwaite, 1975]. Indeed, without penalization from monetary transfers, agents can report as large utilities as possible for the resource; utility reports from the agents carry virtually no information. In this case, agent reports can at most signal whether agents have negative, zero, or positive utility for the resource. We note that mechanism design without monetary transfers can be successful in specific settings, e.g., if agent preferences are single-peaked, when matching users pairwise in a bipartite graph (see Schummer and Vohra [2007] for an overview of classical results), or in interdomain routing [Feigenbaum et al., 2007, Levin et al., 2008]; these exploit specific properties of the problem that are difficult to generalize.

In a repeated allocation setting, however, the central planner can use future allocations to align incentives by favoring or disfavoring agents. We consider two classical repeated settings in the literature: an infinite-horizon setting with a fixed utility time discount factor $\gamma \in [0, 1)$ and a finite-horizon T setting without discounts. To illustrate the benefit of repeated allocations, the central planner can, for instance, fix budgets for each agent that serve as virtual currency. This is routinely used in practice [Walsh, 2014], including for allocating computing resources in distributed systems [Ng, 2011] or food to food banks [Prendergast, 2022].

Characterizing the convergence to first-best allocation. From the folk theorem [Friedman, 1971, Fudenberg et al., 2009], one can approximate the optimal first-best allocation arbitrarily well as agents become more patient $\gamma \rightarrow 1$, or when the horizon grows $T \rightarrow \infty$. However, this result is not quantitative in nature and does not provide direct insights into which allocation mechanisms might be optimal.

The goal of this paper is to provide *quantitative* answers to the following set of questions:

1. What is the welfare gap between non-monetary and monetary mechanisms? Equivalently, what is the welfare gap between strategic and non-strategic agents for non-monetary mechanisms? More formally, can we quantify the *non-asymptotic* social welfare gap between the optimal first-best allocation (which can be achieved via monetary mechanisms) and

non-monetary resource allocation mechanisms for general utility distributions? Conceptually, we aim to understand how this gap depends on the utility distributions of the agents, and on the patience of agents (either through the discount factor γ or the horizon T).

2. On the algorithmic front, what are (near-)optimal allocation mechanisms without monetary transfers? Since the optimal allocation strategy may depend on the specific problem at hand, through utility distributions and horizon and agents’ patience, we aim to provide a generic non-monetary allocation mechanism strategy whose parameters can easily be adjusted for any scenario to achieve near-optimal social welfare.

From a practical perspective, understanding the main drivers (Item 1) for the social welfare gap can (1) help identify which applications are best suited for non-monetary mechanisms, or (2) inform decisions that may impact utility distributions. As a preview, our results show that the social welfare depends heavily on the geometry of utility distributions and identify a main driver for achieving small social welfare gaps: favoring (approximate) utility ties between agents is most beneficial for non-monetary mechanisms. In addition, characterizing the social welfare gap has algorithmic implications (Item 2)—this will be used in the parameter choice for the generic allocation strategy.

1.1 Overview of the contributions

To answer these questions, we use a geometric approach, in a similar spirit to [Balseiro et al. \[2019\]](#). Specifically, instead of solely focusing on approaching the social-welfare maximizing utilities, one can consider the complete region of achievable utilities for all n agents by non-monetary mechanisms, which we denote $\mathcal{U}$: it is the collection of utility vectors $(U_i)_{i \in [n]}$ for which one can achieve expected per-round U_i utility for all agents i , using some non-monetary mechanism. In this geometric perspective, the gap between non-monetary and monetary mechanisms precisely corresponds to the gap between the achievable region $\mathcal{U}$ and the target region $\mathcal{U}^*$ of utilities that could be allocated if agents were non-strategic. At the conceptual level, our results show that the local smoothness of this target region $\mathcal{U}^*$ at welfare-maximizing points, precisely determines the gap between monetary and non-monetary mechanisms.

Near-optimal allocation mechanisms. On the positive side, we construct allocation mechanisms (see Algorithm 1) that can achieve any utility inside a pre-specified ball within the target region $\mathcal{U}^*$ which has sufficient margin $\delta > 0$ to its boundary. This mechanism builds upon the so-called *promised utility* framework [[Abreu et al., 1990](#), [Spear and Srivastava, 1987](#), [Thomas and Worrall, 1990](#)] and tools from [Balseiro et al. \[2019\]](#), which reformulate the problem in a dynamic programming form. Crucially, the larger the radius of this ball, the smaller this margin δ is required to be. Geometrically, for a target region $\mathcal{U}^*$ with smoother boundary, we can employ our mechanisms to achieve utility vectors closer to social-welfare optimal boundary. In words, our results show that non-monetary mechanisms perform best when agent utility distributions significantly overlap: situations with close-ties are particularly advantageous for a central planner since they have flexibility to choose who to allocate the resource to without incurring significant social welfare cost. In fact, when exact ties arise frequently, our results prove exponentially small social welfare gaps between monetary and non-monetary mechanisms.

We complement these positive results with corresponding negative results, which essentially show that the local smoothness of the target region $\mathcal{U}^*$ indeed characterizes the social welfare gap between non-monetary and monetary mechanisms. This effectively transforms the multi-horizon

multi-agent game-theoretic resource allocation problem into a purely geometric question of understanding the local smoothness of $\mathcal{U}^*$ (so that our mechanism Algorithm 1 can be implemented with an appropriate ball parameter).

Unified perspective on γ -discounted and horizon- T settings. Up to minor modification of the algorithms, we also show that the same tools described above can be used for both the infinite-horizon setting with discount factor, and the finite horizon T setting. This unifies the analysis and the algorithmic approach for both settings. Formally, the same results can be achieved up to logarithmic factors by setting $1 - \gamma \approx 1/T$ (see Theorems 6.1 and 6.3).

Characterizations of the non-monetary/monetary social welfare gap. Using this general methodology, we can characterize the social welfare gap between non-monetary and monetary mechanisms for general utility distributions, beyond prior results.

First, the non-monetary/monetary social welfare gap is always at most $\mathcal{O}(\sqrt{1 - \gamma})$ (resp. $\mathcal{O}(1/\sqrt{T})$) for the γ -discounted infinite-horizon (resp. finite-horizon T) setting (Theorem 5.1). In particular, this holds for any number of agents and any utility distributions with bounded support. Beyond the universal asymptotic convergence guarantee from the folk theorem [Friedman, 1971, Fudenberg et al., 2009], this shows a universal convergence rate upper bound.

While this universal upper bound can be tight¹, the convergence rate can be faster in general, and depends on the local geometry of $\mathcal{U}^*$ as discussed above. Specifically, while the distributional assumptions from Balseiro et al. [2019] prohibited faster rates than $1 - \gamma$, we give non-trivial examples of exponential convergence rates as fast as $(1 - \gamma)e^{-c/\sqrt{1 - \gamma}}$ (see Theorem 7.4 which essentially corresponds to the case when any agent $i \in [n]$ can have ties for the first-best allocation with non-zero probability). From a practical standpoint, this shows that when ties between several agents are possible, our mechanisms which leverage this additional geometrical structure can achieve significantly (exponentially) smaller social welfare gap. For the finite-horizon case, faster rates than $1/\sqrt{T}$ are possible but they cannot converge faster than $1/T$ (Theorem 6.2) as a result of the boundary at time T . In particular, the exponential rates from the infinite-horizon setting do not carry to the finite-horizon case. To the best of our knowledge, this is the first example showing exponential gaps between the discounted and finite-horizon settings for non-monetary resource allocation.

Our results provide general upper and lower bounds on the optimal social welfare gap for arbitrary distributions (Theorem 5.8). While the bounds do not match in general, both share the same form and give insights into (1) the dependency of the convergence rate in terms of the utility distributions and (2) corresponding optimal allocation mechanisms. Intuitively, the main driver is whether the central planner has enough flexibility to reward or penalize any agent in future rounds, without incurring large social welfare costs. Accordingly, our mechanisms use this flexibility to align incentives: the larger the flexibility, the closer one can approach the first-best allocation.

Special cases and examples. In particular, we can give necessary (Theorem 5.5) and sufficient (Theorem 5.3) conditions to reach the faster rate $1 - \gamma$ from Balseiro et al. [2019] in the infinite-horizon setting, significantly generalizing their results (the necessary and sufficient conditions are both satisfied under their distributional assumptions). In the finite-horizon setting,

¹this is typically the case for discrete utility distributions whose boundary of $\mathcal{U}^*$ is piece-wise affine (non-smooth), see Theorem 7.3.

Setting	General distributions (Universal bounds)		Density bounded above & below		Discrete (Generic case)		Intermediary example of Fig. 6		Significant ties between agents (see Thm 7.4)	
	γ -disc.	T -hor.	γ -disc.	T -hor.	γ -disc.	T -hor.	γ -disc.	T -hor.	γ -disc.	T -hor.
SWG	$\mathcal{O}(\sqrt{1-\gamma})$	$\mathcal{O}\left(\frac{1}{\sqrt{T}}\right)$ and $\Omega\left(\frac{1}{T}\right)$	$\Theta(1-\gamma)$	$\tilde{\Theta}\left(\frac{1}{T}\right)$	$\Theta(\sqrt{1-\gamma})$	$\Theta\left(\frac{1}{\sqrt{T}}\right)$	$\Theta((1-\gamma)^{\frac{3}{4}})$	$\tilde{\Theta}\left(\frac{1}{T^{\frac{3}{4}}}\right)$	$\mathcal{O}\left((1-\gamma)e^{-\frac{c}{\sqrt{1-\gamma}}}\right)$	$\tilde{\Theta}\left(\frac{1}{T}\right)$
Ref.	Thm 5.1	+Lem 6.2&6.3	Cor 7.1	Cor 7.2	Thm 7.3	Cor 7.5	Prop 5.6	+Lem 6.1&6.3	Thm 7.4	+Lem 6.2&6.3
Prior results	—	$\mathcal{O}(1/T)$ for $\tilde{\mathcal{O}}(1/\sqrt{T})$ - approx. IC	$\Theta(1-\gamma)$ for $n=2$	—	—	—	—	—	—	—
Ref.		Gorokh et al. [2021]	Balseiro et al. [2019]							

Table 1: Examples of social welfare gap (SWG) for various utility distribution classes for γ -discounted and finite T -horizon settings. The first column gives universal bounds—always valid—while others are consequences of our general distribution-dependent characterization (Theorem 5.8). This includes distributions with well-behaved densities, discrete distributions without significant ties (see Theorems 7.3 and 7.4), an example with intermediary convergence rates, and scenarios with significant ties.

this corresponds to the fastest rate $1/T$ up to logarithmic factors. At the high-level, we show that these faster rates can be achieved if any agent $i \in [n]$ shares some common utility support with some other agent $j \neq i$: the mechanism can easily then implement future penalties (resp. rewards) on agent i by allocating to agent j (resp. i) when their utilities are sufficiently close in following rounds. Using this perspective, we introduce a sufficient condition for fast convergence rates in terms of a graph structure between agents, which can be efficiently tested given their utility distributions.

The previous example covers the case of smooth distributions (Theorems 7.1 and 7.2). On the other extreme, we also take the example of discrete utility distributions (Theorems 7.3 and 7.5). In this case, the optimal convergence rate is usually $\sqrt{1-\gamma}$ (resp. $1/\sqrt{T}$) for the infinite-horizon (resp. finite-horizon) setting.

In comparison, Gorokh et al. [2021] showed that for discrete distributions, one can reach a convergence rate $1/T$ to the first-best allocation with artificial currencies, but in a weaker form of Bayesian equilibrium in which incentive-compatibility constraints are only satisfied up to $\sqrt{\log T/T}$ terms. Our results show that the optimal rate at the perfect Bayesian equilibrium is, in general, exactly $\Theta(1/\sqrt{T})$, when incentive-compatibility constraints are perfectly enforced. Intuitively, the incentive-compatibility gap from Gorokh et al. [2021] can be traded for the social welfare gap from the central planner. Unlike Gorokh et al. [2021], our upper bound holds for arbitrary distributions.

A summary of these social welfare characterizations for various utility distribution classes is given in Table 1: (i) without any assumptions on the utility distributions one can always achieve $\mathcal{O}(\sqrt{1-\gamma})$ convergence rate, (ii) this is tight for discrete distributions in general, (iii) but may be improved to fast rates $\Theta(1-\gamma)$, for instance when distributions have density bounded above and below. (iv) Intermediate situations are also possible and we provide such an example in Fig. 6, with convergence rate $\Theta((1-\gamma)^{3/4})$. Finally, (v) when there are significant ties between agents, exponential rates are possible in the discounted setting; however these do not translate to the finite-horizon T setting.

1.2 Related works

Single-time allocation without money. We first discuss the literature on one-shot resource allocation without money. When there are two resources and agents share the same utility distribution, [Miralles \[2012\]](#) characterizes optimal mechanisms and shows that these can be constructed as combinations of lotteries and auction-like strategies. For several resources and additive utilities for the agents, [Cole et al. \[2013\]](#), [Guo and Conitzer \[2010\]](#), [Han et al. \[2011\]](#) give characterizations on the competitive ratio of strategy-proof mechanisms compared to the first-best allocation, without making prior assumptions on the agent utility distributions. Slightly different from our setting without money, the line of work including [Chakravarty and Kaplan \[2013\]](#), [Condorelli \[2012\]](#), [Hartline and Roughgarden \[2008\]](#), [Hoppe et al. \[2009\]](#) studies the setting in which the central planner can impose payments from the agents but cannot redistribute these between agents. This setting is often referred to *money burning*: any utility signal in the form of payment to the central planner is lost. As discussed above, impossibility results due to Myerson-Satterthwaite’s impossibility theorem generally prohibit the central planner from converging to the first-best allocation for a single allocation, while our focus is on designing mechanisms that converge to this optimal allocation as fast as possible.

Repeated allocation without money. Closest to our work, we now discuss the literature on the repeated allocation setting without transfers. While the folk theorem [Fudenberg et al. \[2009\]](#) guarantees the existence of a mechanism that asymptotically converges to the first-best allocation, [Guo et al. \[2009\]](#) explicitly provides a mechanism that achieves at least 75% of the optimal welfare in the infinite-horizon setting, provided agents are sufficiently patient. On the other hand, their mechanism may not converge to the optimal first-best allocation as $\gamma \rightarrow 1$. [Guo and Hörner \[2015\]](#) characterize the optimal allocation in the stylized setting in which the central planner has a fixed cost and decides at each iteration whether to allocate the resource to a single agent whose utilities evolve according to a two-state Markov chain.

We previously introduced budget-based mechanisms, also called *artificial currency* or *script* strategies [Friedman et al. \[2006\]](#), [Johnson et al. \[2014\]](#), [Kash et al. \[2007, 2015\]](#), in which each agent is given an initial budget that can be used as currency to bid on resources in each round. These have the advantage of being very simple and interpretable, and their implementation in several real-world applications has motivated the study of their theoretical performance [\[Budish, 2011, Gorokh et al., 2021, Jackson and Sonnenschein, 2007\]](#). For instance, [Gorokh et al. \[2021\]](#) shows that in the finite-horizon T setting and for discrete utility distributions their convergence to the first-best allocation decreases as $1/T$. However, their results hold under a weaker notion of approximate Nash equilibrium in which incentive-compatibility constraints are only satisfied approximately up to $\sqrt{\log T/T}$ terms. In a similar spirit to artificial currencies, [Elokda et al. \[2022\]](#) recently proposed the use of karma mechanisms and discussed its applications in ride-hailing platforms, in which agents can receive an additional budget reward for accepting to give the resource to other agents.

The work of [Balseiro et al. \[2019\]](#) is closest to the present paper. They consider the discounted infinite-horizon setting where agent utility distributions admit densities bounded above and below by some fixed constant. With these distributional assumptions and for two agents, they provide allocation mechanisms that converge to the first-best allocation as $\Theta(1 - \gamma)$ for $\gamma \rightarrow 1$ and show that this is the best possible among all mechanisms up to constants. In comparison, they show that budget-based mechanisms have a convergence rate of at most $\Omega(\sqrt{1 - \gamma})$; hence, budget-based strategies are suboptimal under these utility distribution assumptions. For $n \geq 3$

agents, they prove a $\mathcal{O}((1-\gamma)^{1/(n+4)})$ convergence rate for their mechanism, leaving a significant gap compared to their $\Omega(1-\gamma)$ lower bound.

Compared to these previous works, we give tools to prove convergence rates for *general* utility distributions. As an example, with the specific assumptions from Balseiro et al. [2019], we provide an allocation mechanism that achieves the optimal convergence rate $\Theta(1-\gamma)$ for any number of agents n . More precisely, the upper (resp. lower) bound holds whenever the utility distributions have densities lower (resp. upper) bounded by a constant. Our work also aims to reconcile the analysis for the finite-horizon and infinite-horizon settings by using the same techniques for both.

Last, we make some comments about potential extensions of our current approach. In many operational scenarios, the utility distributions of agents may evolve over time following a periodic structure—say every p rounds—for instance, due to seasonality in food banks. We expect that our current approach can be adapted to this non-i.i.d. setting by grouping p consecutive times for the allocation function—effectively replacing p distributions with an appropriate mixture on each length- p group. Beyond seasonality, utility distributions may also be dependent on the past history of allocations. We leave for further work whether our results can have implications for such adaptive utility distributions.

1.3 Organization of the paper

We define the model and the framework used for our constructions in Section 2. In Section 3, we describe our general resource allocation mechanism and then present our main geometric tool to derive social welfare lower bound gaps in Section 4. Using these two geometric tools, we derive convergence rates for social welfare in the infinite-horizon setting in Section 5. In Section 6, we show how the tools developed for the infinite-horizon setting can be used for the finite-horizon setting as well, with only minor modifications. We provide examples of the results that can be obtained with our characterizations for utility distributions of interest in Section 7. Lastly, we conclude in Section 8.

1.4 Notations

For $n \geq 1$, we denote $[n] := \{1, \dots, n\}$. We also write $\mathbb{R}_+ := [0, \infty)$ and $\mathbf{1} = (1, \dots, 1) \in \mathbb{R}^n$. By default, we use $\|\cdot\|$ for the Euclidean norm $\|\cdot\|_2$, otherwise, we will always specify the norm within the paper. For any $\mathbf{x} \in \mathbb{R}^n$ and $r > 0$, we denote $B(\mathbf{x}, r) = \{\mathbf{y} \in \mathbb{R}^n : \|\mathbf{y} - \mathbf{x}\| \leq r\}$ and $B^\circ(\mathbf{x}, r) = \{\mathbf{y} \in \mathbb{R}^n : \|\mathbf{y} - \mathbf{x}\| < r\}$, the closed and open balls centered at $\mathbf{x}$ respectively. We let $S_{n-1} = \{\mathbf{y} \in \mathbb{R}^n : \|\mathbf{y}\| = 1\}$ be the unit n -dimensional sphere. For a compact $C \subset \mathbb{R}^n$ and any $\mathbf{x} \in \mathbb{R}^n$, we define $d(\mathbf{x}, C) = \max_{\mathbf{y} \in C} \|\mathbf{y} - \mathbf{x}\|$, where the maximum is attained since C is compact. Next, for any subset $S \subseteq [n]$ and $\mathbf{x} \in \mathbb{R}^n$, we denote $\mathbf{x}_S := (x_i)_{i \in S}$. For $i \in [n]$ we also denote $\mathbf{x}_{-i} = (x_j)_{j \neq i}$. We denote $\Delta_n := \{\mathbf{x} \in \mathbb{R}_+^n, \sum_{i \in [n]} x_i \leq 1\}$ the probability simplex (augmented with a do-nothing option). We denote by $\mathbf{a} \odot \mathbf{b} = (a_i b_i)_i$ the component-wise product of two vectors $\mathbf{a}$ and $\mathbf{b}$.

We use standard Landau notations. For two functions f , and g , we write $f(x) = \mathcal{O}(g(x))$ if there exists a constant C such that $f(x) \leq Cg(x)$ for all x in the support of f . We may specify $f(x) \underset{x \rightarrow a}{=} \mathcal{O}(g(x))$ when the inequality $f(x) \leq Cg(x)$ only holds in a neighborhood of a within the support of f . We write $f(x) = \Omega(g(x))$ when $g(x) = \mathcal{O}(f(x))$, and $f(x) = \Theta(g(x))$ if both $f(x) = \mathcal{O}(g(x))$ and $f(x) = \Omega(g(x))$ hold. We write $f(x) \underset{x \rightarrow a}{=} \omega(g(x))$ if $0 \leq h(x)g(x) \leq f(x)$

where $h(x) \rightarrow \infty$ as $x \rightarrow a$. Last, the notations $\tilde{\mathcal{O}}(\cdot)$ and $\tilde{\Theta}(\cdot)$ hide logarithmic factors in the horizon T .

2 Problem formulation and preliminaries

We first detail the model and give useful preliminaries about the promised utility framework. Except for characterizations in Section 2.5, most statements in this section can be considered standard in the literature.

2.1 Model definitions

We consider the sequential problem in which a central planner repeatedly allocates a single resource between n agents. At every round $t \geq 1$, each strategic agent $i \in [n]$ observes their private realized utility $u_i(t)$, then reports a value $v_i(t)$ to the central planner only, and last, having observed the reported values $\mathbf{v}(t) = (v_i(t))_{i \in [n]}$ the central planner allocates the single resource to one or none of the agents, which we represent by the index $\hat{i}(t) \in \{0, \dots, n\}$ where $\hat{i}(t) = 0$ signifies that the resource was not allocated at time t . Alternatively, the central planner selects a probabilistic allocation $\mathbf{p}(t) = (p_i(t))_{i \in [n]} \in \Delta_n = \{\mathbf{y} \geq \mathbf{0} : \sum_{i \in [n]} y_i \leq 1\}$. Last, at the end of the round, the reported values $\mathbf{v}(t)$ and allocation $\mathbf{p}(t)$ are made public to all agents.

Crucially, we suppose that the true utilities are independent between agents and independent and identically distributed (i.i.d.) for each agent. That is, for any $i \in [n]$, we have $(u_i(t))_{t \geq 1} \stackrel{iid}{\sim} \mathcal{D}_i$. We denote by $\bar{F}_i$ the complementary cumulative function of $\mathcal{D}_i$, that is, $\bar{F}_i(x) = \mathbb{P}_{u \sim \mathcal{D}_i}[u \geq x]$. We assume that the fixed distributions $\mathcal{D}_1, \dots, \mathcal{D}_n$ have bounded support in $[0, \bar{v}]$ for some fixed value $\bar{v} > 0$, and are publicly known to the central planner and to all agents $i \in [n]$. Last, we consider the case of myopic agents, that is, at time t , agent i does not have access to their realized utilities at future times $u_i(t')$ for $t' > t$. Formally, write the history available at time t as $\mathbf{h}(t) = (v_j(t'), p_j(t'), j \in [n], t' < t)$. At time t , the central planner strategy is given by a function $S_t : [0, \bar{v}]^{n(t-1)} \times \Delta_n^{t-1} \times [0, \bar{v}]^n \rightarrow \Delta_n$ and the allocation is given by $\mathbf{p}(t) = S_t(\mathbf{h}(t), \mathbf{v}(t))$, while the strategy of agent $i \in [n]$ is given by a function $\sigma_{i,t} : [0, \bar{v}]^{n(t-1)} \times \Delta_n^{t-1} \times [0, \bar{v}] \rightarrow [0, \bar{v}]$ and their report is given by $v_i(t) = \sigma_{i,t}(\mathbf{h}(t), u_i(t))$. For convenience, throughout the paper, $\mathbf{u}$ represents a random utility vector with independent coordinates with $u_i \sim \mathcal{D}_i$.

The goal of the central planner is to maximize the social welfare or equivalently, to minimize the regret compared to the optimal allocation in hindsight $i^*(t) \in \arg \max_{i \in [n]} u_i(t)$. We consider two different settings.

1. In the *γ -discounted infinite-horizon* setting, the central planner and all agents share the same discount factor $\gamma \in [0, 1)$. That is, the central planner aims to minimize the expected regret

$$(1 - \gamma) \mathbb{E} \left[\sum_{t \geq 1} \gamma^{t-1} \left(\max_{i \in [n]} u_i(t) - u_{\hat{i}(t)}(t) \right) \right].$$

while agent $i \in [n]$ aims to maximize their total discounted utility $(1 - \gamma) \mathbb{E} \left[\sum_{t \geq 1} \gamma^{t-1} u_i(t) p_i(t) \right]$.

2. In the *finite horizon T* setting, the central planner aims to minimize the average regret

$$\frac{1}{T} \mathbb{E} \left[\sum_{t=1}^T \max_{i \in [n]} u_i(t) - u_{\hat{i}(t)}(t) \right],$$

while agent $i \in [n]$ aims to maximize their total utility $\frac{1}{T} \mathbb{E} \left[\sum_{t \in [T]} u_i(t) p_i(t) \right]$.

The reports of the strategic agent form a perfect Bayesian equilibrium, as detailed below. Using the same notations as in Balseiro et al. [2019], we let $V_{i,t}$ be the expected remaining utility of agent i starting from time t , given its own utility $u_i(t)$ and the history $\mathbf{h}(t)$, with the allocation strategy $\mathbf{S} = (S_{t'})_{t' \geq 1}$, all other agents employing their strategy $\boldsymbol{\sigma}_{-i} = (\sigma_{j,t'})_{j \neq i, t' \geq 1}$ and agent i using strategy $\tilde{\sigma}_i = (\tilde{\sigma}_{i,t'})_{t' \geq 1}$. Formally, we define $V_{i,t}(\mathbf{S}, \boldsymbol{\sigma}_{-i}, \tilde{\sigma}_i \mid u_i(t), \mathbf{h}(t))$ as

$$(1 - \gamma) \mathbb{E} \left[\sum_{t' \geq t} \gamma^{t'-t} u_i(t') p_i(t') \mid u_i(t), \mathbf{h}(t) \right]$$

for the γ -discounted setting, and

$$\frac{1}{T - t + 1} \mathbb{E} \left[\sum_{t'=t}^T u_i(t') p_i(t') \mid u_i(t), \mathbf{h}(t) \right]$$

for the finite horizon T setting. In both definitions, $\mathbf{p}(t')$ is induced by the central planner strategy $S_{t'}$ for $t' \geq t$ starting from the history $\mathbf{h}(t)$, using the reports of agents $j \neq i$ according to σ_j and the reports of agent i according to $\tilde{\sigma}_i$. The perfect Bayesian equilibrium constraints then write

$$\begin{aligned} \forall i, t, u_i(t), \mathbf{h}(t), \tilde{\sigma}_i, \quad & V_{i,t}(\mathbf{S}, \boldsymbol{\sigma}_{-i}, \sigma_i \mid u_i(t), \mathbf{h}(t)) \\ & \geq V_{i,t}(\mathbf{S}, \boldsymbol{\sigma}_{-i}, \tilde{\sigma}_i \mid u_i(t), \mathbf{h}(t)). \end{aligned} \quad (\text{PBE})$$

2.2 Achievable utility sets

As an alternative reformulation, letting $\mathbf{U} = (U_i)_{i \in [n]}$ be the vector of realized utilities where U_i is the total expected utility gained by agent i , the goal is to maximize the total utility $\mathbf{1}^\top \mathbf{U}$ among realizable utility vectors $\mathbf{U}$. For our purposes we will consider the more general problem in which agents can have weights $\boldsymbol{\alpha} \in \mathbb{R}_+^n \setminus \{\mathbf{0}\}$ and the goal of the central planner becomes maximizing the total utility $\boldsymbol{\alpha}^\top \mathbf{U}$. Equivalently, letting $V_i(\mathbf{S}, \boldsymbol{\sigma}) := \mathbb{E}_{u_i(1)}[V_{i,1}(\mathbf{S}, \boldsymbol{\sigma} \mid u_i(1))]$, we aim to characterize the region of achievable utilities:

$$\{\mathbf{U} = (V_i(\mathbf{S}, \boldsymbol{\sigma}))_{i \in [n]} \in \mathbb{R}_+^n : \mathbf{S}, \boldsymbol{\sigma} \text{ satisfy } (\text{PBE})\},$$

which we denote $\mathcal{U}_\gamma$ for the γ -discounted setting, and $\mathcal{V}_T$ for the finite horizon T setting.

Full-information set. We first characterize the optimal allocation baseline, which corresponds to the *full-information* setting in which the central planner has direct access to all agent utilities $(u_i(t))_{i \in [n]}$ before deciding its allocation. This also corresponds to the setting in which agents are non-strategic. The set of achievable utilities in this full-information setting is

$$\mathcal{U}^* = \{(\mathbb{E}[u_i p_i(\mathbf{u})])_{i \in [n]}, \mathbf{p} : [0, \bar{v}]^n \rightarrow \Delta_n\}.$$

We can also easily characterize the Pareto-optimal points in this set of achievable utilities. This boundary is parametrized by weights $\boldsymbol{\alpha} \in \mathbb{R}_+^n \setminus \{\mathbf{0}\}$ with the corresponding maximizing utility vectors $\arg \max_{\mathbf{U} \in \mathcal{U}^*} \boldsymbol{\alpha}^\top \mathbf{U} = \mathbb{E}[\max_{i \in [n]} \alpha_i u_i]$. Such $\boldsymbol{\alpha}$ -optimal vectors exactly correspond to

first-best allocation functions that always allocate the resource to $i \in \arg \max_{j \in [n]} \alpha_j u_j$. By default, we say that $\mathbf{x} \in \mathcal{U}^*$ is optimal if it is $\mathbf{1}$ -optimal.

Recall that the central planner aims to maximize the social welfare, corresponding to equalitarian weights $\alpha = \mathbf{1}$. Hence, the optimal regret in the γ -discounted setting, which corresponds to the social welfare gap between strategic and non-strategic agents is

$$\text{SWG}(\gamma) := \max_{\mathbf{U} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{U} - \max_{\mathbf{U} \in \mathcal{U}_\gamma} \mathbf{1}^\top \mathbf{U}. \quad (1)$$

By abuse of notation, we denote by $\text{SWG}(T)$ the social welfare gap in the finite-horizon- T setting, that is, replacing $\mathcal{U}_\gamma$ with $\mathcal{V}_T$ in Eq. (1).

Remark 2.1. *While the central planner mainly focuses on equalitarian weights $\alpha = \mathbf{1}$, it turns out that the geometry of the boundary of $\mathcal{U}^*$ —or equivalently, maximizing utility vectors for weights $\alpha \neq \mathbf{1}$ —is crucial to characterize the optimal social welfare gap and construct corresponding mechanisms.*

No-information set. We next define the set of utilities that can be achieved *without reports*, that is without information on the utility realizations. This corresponds to a setting in which the central planner selects a fixed allocation $\mathbf{p} \in \Delta_n$, which gives the following frontier

$$\mathcal{U}^{NI} = \{(p_i \mathbb{E}_{u_i \sim \mathcal{D}_i}[u_i])_{i \in [n]}, \mathbf{p} \in \Delta_n\}.$$

For intuition, if $\mathbb{P}(u_i = 0) = 0$ for all $i \in [n]$, we have $\mathcal{U}^{NI} = \mathcal{U}_0 = \mathcal{V}_1$. We will see, however, that if this does not hold we may have $\mathcal{U}^{NI} \subsetneq \mathcal{U}_0 = \mathcal{V}_1$, see Theorem 2.4, scenario 3(b).² Using mixed allocation strategies, we can easily check the following.

Fact 2.1. *The sets $\mathcal{U}^*$, $\mathcal{U}^{NI}$, $\mathcal{U}_\gamma$, and $\mathcal{V}_T$ are convex sets of $\mathbb{R}_+^n$ for all $\gamma \in [0, 1)$ and $T \geq 1$.*

2.3 The promised utility framework

Before giving our results, we present a convenient reformulation of the sequential allocation problem in terms of dynamic programming, referred to as the *promised utility framework* [Abreu et al., 1990, Balseiro et al., 2019, Spear and Srivastava, 1987, Thomas and Worrall, 1990]. We give below an introduction and summary of the framework for the γ -discounted infinite-horizon setting as detailed in Balseiro et al. [2019].

We focus on the γ -discounted infinite-horizon setting for now. In its direct form, maximizing social welfare requires solving an infinite-horizon perfect Bayes equilibrium. However, the symmetry of the problem in time allows to drastically simplify the formulation. In light of the revelation principle, we can without loss of generality suppose that the mechanism is incentive-compatible hence agents report truthfully. Say we want to design an allocation strategy to ensure that the final utility vector is $\mathbf{U}$. The first decision is to decide on the allocation function for the first time step $t = 1$ which we denote $\mathbf{p}(\cdot \mid \mathbf{U}) : [0, \bar{v}]^n \rightarrow \Delta_n$. It takes as input the agent reports $\mathbf{v}(t = 1) = \mathbf{u}(1)$ and outputs the allocation distribution $\mathbf{p}(1)$. We then denote by $\mathbf{W}(\cdot \mid \mathbf{U}) : [0, \bar{v}]^n \rightarrow \mathbb{R}^n$ the remaining utility to be gathered by the agents at future times $t \geq 2$, that is,

$$W_i(\mathbf{v} \mid \mathbf{U}) = (1 - \gamma) \mathbb{E} \left[\sum_{t \geq 2} \gamma^{t-2} u_i \mathbb{1}_{\hat{u}_t = i} \mid \mathbf{u}(1) = \mathbf{v} \right].$$

²Specifically, the optimal vector $\mathbf{U}$ constructed therein is outside $\mathcal{U}^{NI} = \text{Conv}(\mathbf{0}; \{\mathbb{E}[u_i] \mathbf{e}_i, i \in [n]\})$.

Note that the history also includes the allocation $\mathbf{p}(1)$ but this is completely subsumed by the knowledge of $\mathbf{u}(1) = \mathbf{v}$ via $\mathbf{p}(1) = \mathbf{p}(\mathbf{v} \mid \mathbf{U})$. The problem at time $t = 2$ is now perfectly equivalent to the original problem except that now the central planner aims to ensure that the utility vector to realize is $\mathbf{W}(\mathbf{u}(1) \mid \mathbf{U})$.

This motivates focusing only on the allocation function and the promised utility function: to achieve a region $\mathcal{R} \subseteq \mathcal{U}^*$, it suffices to specify two functions $\mathbf{p}(\cdot \mid \mathbf{U}) : [0, \bar{v}]^n \rightarrow \Delta_n$ and $\mathbf{W}(\cdot \mid \mathbf{U}) : [0, \bar{v}]^n \rightarrow \mathbb{R}^n$ for all $\mathbf{U} \in \mathcal{R}$, where the later is viewed as a utility promise to the agents. These should satisfy the following constraints. The *promise-keeping* constraint checks that the target utility $\mathbf{U}$ is met: for all $i \in [n]$ and $\mathbf{U} \in \mathcal{R}$,

$$U_i = \mathbb{E}[(1 - \gamma)u_i p_i(\mathbf{u} \mid \mathbf{U}) + \gamma W_i(\mathbf{u} \mid \mathbf{U})]. \quad (2)$$

Second, the central planner should ensure that the promised utilities can be fulfilled for any possible reports of the agent. That is,

$$\forall \mathbf{U} \in \mathcal{R}, \forall \mathbf{v} \in [0, \bar{v}]^n, \quad \mathbf{W}(\mathbf{v} \mid \mathbf{U}) \in \mathcal{R}. \quad (3)$$

The last constraint is to ensure that (PBE) is satisfied, which can be simplified since we supposed that the mechanism is incentive-compatible. Having defined the interim allocation $P_i(v \mid \mathbf{U}) := \mathbb{E}_{\mathbf{u}_{-i}}[p_i(v, \mathbf{u}_{-i} \mid \mathbf{U})]$ and the interim promised utility $W_i(v \mid \mathbf{U}) = \mathbb{E}_{\mathbf{u}_{-i}}[W_i(v, \mathbf{u}_{-i} \mid \mathbf{U})]$ for all $i \in [n]$, the perfect Bayesian incentive-compatibility equations write for all $\mathbf{U} \in \mathcal{R}$, $i \in [n]$, and $u \in [0, \bar{v}]$,

$$u \in \arg \max_{v \in [0, \bar{v}]} (1 - \gamma)u P_i(v \mid \mathbf{U}) + \gamma W_i(v \mid \mathbf{U}). \quad (4)$$

Fact 2.2. *A region $\mathcal{R}' \subseteq \mathcal{U}^*$ can be achieved in the infinite-horizon γ -discounted setting if and only there exists a superset $\mathcal{R} \supseteq \mathcal{R}'$, functions $\mathbf{p}(\cdot \mid \mathbf{U})$ and $\mathbf{W}(\cdot \mid \mathbf{U})$ for all $\mathbf{U} \in \mathcal{R}$ satisfying Eqs. (2) to (4).*

When $\gamma > 0$, we can further characterize the form of the allocation and promised utility functions following a standard argument from Myerson [1981]. Given the incentive-compatibility constraint Eq. (4), without loss of generality, we can suppose that the functions $P_i(\cdot \mid \mathbf{U})$ are non-decreasing. For convenience, we say that a choice of allocation function $\mathbf{p}(\cdot \mid \mathbf{U})$ is non-decreasing if the corresponding interim allocation function $P_i(\cdot \mid \mathbf{U})$ is non-decreasing for all $i \in [n]$. Then, Eq. (4) implies

$$\begin{aligned} & W_i(v_i \mid \mathbf{U}) - W_i(0 \mid \mathbf{U}) \\ &= \frac{1 - \gamma}{\gamma} \left(\int_0^{v_i} P_i(v \mid \mathbf{U}) dv - P_i(v_i \mid \mathbf{U}) v_i \right). \end{aligned} \quad (5)$$

Combining this with the constraint from Eq. (2), we obtain for all $\mathbf{U} \in \mathcal{R}$, $i \in [n]$, and $v_i \in [0, \bar{v}]$,

$$\begin{aligned} W_i(v_i \mid \mathbf{U}) &= \frac{1}{\gamma} \left[U_i + (1 - \gamma) \left(\int_0^{v_i} P_i(v \mid \mathbf{U}) dv \right. \right. \\ &\quad \left. \left. - P_i(v_i \mid \mathbf{U}) v_i - \int_0^{\bar{v}} \bar{F}_i(v) P_i(v \mid \mathbf{U}) dv \right) \right]. \end{aligned} \quad (6)$$

This formula can be found in Balseiro et al. [2019] as Eq (10). This equation effectively replaces the promise-keeping constraint Eq. (2) and the incentive-compatibility constraint Eq. (4). Further, we can check that $W_i(v_i \mid \mathbf{U})$ is non-increasing in v_i by re-writing $\int_0^{v_i} P_i(v \mid \mathbf{U}) dv - P_i(v_i \mid \mathbf{U}) v_i$

$\mathbf{U})v_i = \int_0^{v_i} (P_i(v \mid \mathbf{U}) - P_i(v_i \mid \mathbf{U}))dv$. The promised utility function only needs to satisfy the following validity equation for the interim promise: for all $\mathbf{U} \in \mathcal{R}$, $i \in [n]$, and $v \in [0, \bar{v}]$,

$$W_i(v \mid \mathbf{U}) = \mathbb{E}_{\mathbf{u}_{-i}}[W_i(v, \mathbf{u}_{-i} \mid \mathbf{U})]. \quad (7)$$

As a summary, we have the following characterization.

Fact 2.3. *Let $\gamma \in (0, 1)$. A region $\mathcal{R}' \subseteq \mathcal{U}^*$ can be achieved in the infinite-horizon γ -discounted setting if and only if there exists a superset $\mathcal{R} \supseteq \mathcal{R}'$, functions $\mathbf{p}(\cdot \mid \mathbf{U})$ and $\mathbf{W}(\cdot \mid \mathbf{U})$ for all $\mathbf{U} \in \mathcal{R}$ satisfying Eqs. (3) and (7), where the interim future promise is defined as in Eq. (6).*

2.4 A convenient choice of promised utility functions

Given an allocation function $\mathbf{p}(\cdot \mid \mathbf{U})$, the remaining choice of the mechanism design is in the form of the promised utility function $\mathbf{W}(\cdot \mid \mathbf{U})$ so that it satisfies Eqs. (3) and (7). We introduce a specific choice of form for these promised utility functions that will be useful for our mechanism design constructions.

The specific construction is somewhat more general than the present sequential allocation setting. Consider the following problem: given n independent random variables $\tilde{Z}_1, \dots, \tilde{Z}_n$ and a direction $\boldsymbol{\alpha} \in \mathbb{R}^n$, we aim to construct a coupling $(Z_1, \dots, Z_n)$ such that (1) for all $i \in [n]$, $\mathbb{E}[Z_i \mid \tilde{Z}_i] = \tilde{Z}_i$, and (2) that is as far along the direction $\boldsymbol{\alpha}$ as possible in the worst-case sense. That is, we aim to solve the problem

$$\max_{\mathbf{Z}} \min_{\omega \in \Omega} \boldsymbol{\alpha}^\top \mathbf{Z}(\omega), \quad \text{s.t. } \forall i \in [n], \mathbb{E}[Z_i \mid \tilde{Z}_i] = \tilde{Z}_i. \quad (8)$$

Taking the expectation, any coupling $\mathbf{Z}$ satisfying the marginal constraints must satisfy

$$\min_{\omega \in \Omega} \boldsymbol{\alpha}^\top \mathbf{Z}(\omega) \leq \mathbb{E}[\boldsymbol{\alpha}^\top \mathbf{Z}] = \boldsymbol{\alpha}^\top \mathbb{E}[\tilde{\mathbf{Z}}]. \quad (9)$$

This value can also be reached for instance using the following coupling from [Balseiro et al. \[2019\]](#). For $i \in [n]$ such that $\alpha_i = 0$, we can let $Z_i = \tilde{Z}_i$. Otherwise, we define

$$Z_i = \tilde{Z}_i - \frac{1}{n-1} \sum_{j \neq i} \frac{\alpha_j}{\alpha_i} (\tilde{Z}_j - \mathbb{E}[\tilde{Z}_j]). \quad (10)$$

We can easily check that this coupling reaches the upper bound from Eq. (9) since the coordinates of $\tilde{\mathbf{Z}}$ are independent. This coupling forces the vector $\mathbf{Z}$ to lie in the hyperplane $\{\mathbf{x} : \boldsymbol{\alpha}^\top (\mathbf{x} - \mathbb{E}[\tilde{\mathbf{Z}}]) = 0\}$.

[Balseiro et al. \[2019\]](#) used this construction to specify the future promise functions: for any $\mathbf{U}$, letting $Z_i := W_i(v_i \mid \mathbf{U})$, they use the promise function $\mathbf{W}(\mathbf{v} \mid \mathbf{U}; \boldsymbol{\alpha}(\mathbf{U}))$ defined as the previous coupling for some pre-specified vector $\boldsymbol{\alpha}(\mathbf{U}) \in \mathbb{R}^n \setminus \{\mathbf{0}\}$. For our purposes, we will use a slightly different coupling. The term α_i in the denominator leads to instabilities that we can mitigate by using the coupling from Eq. (10) only for the variables i with largest value of $|\alpha_i|$. Let i_1 and i_2 be the indices of the largest and second-largest components of $\boldsymbol{\alpha}$ respectively. We pose

$$Z_i = \begin{cases} \tilde{Z}_{i_1} - \sum_{i \neq i_1} \frac{\alpha_i}{\alpha_{i_1}} (\tilde{Z}_i - \mathbb{E}[\tilde{Z}_i]) & i = i_1, \\ \tilde{Z}_{i_2} - \frac{\alpha_{i_1}}{\alpha_{i_2}} (\tilde{Z}_{i_1} - \mathbb{E}[\tilde{Z}_{i_1}]) & i = i_2, \\ \tilde{Z}_i & i \notin \{i_1, i_2\}. \end{cases} \quad (11)$$

Both couplings defined in Eqs. (10) and (11) are optimal for the previous optimization problem.

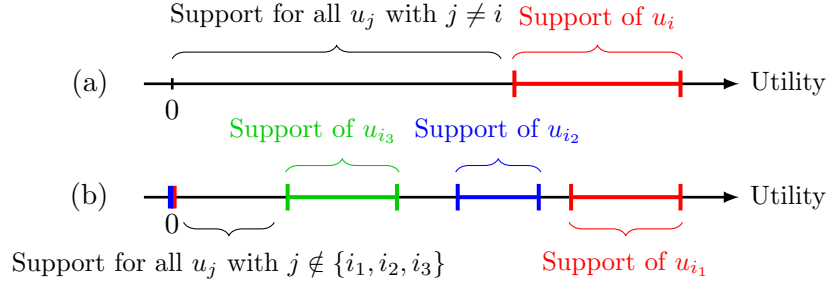


Figure 1: Illustration of the two scenarios (a) and (b) from Theorem 2.4 which correspond to cases in which there is already an optimal allocation strategy for $\gamma = 0$ (resp. $T = 1$) in the infinite-horizon (resp. finite-horizon) setting. Figure (a) corresponds to the scenario in which agent i dominates all other agents. Figure (b) corresponds to the scenario in which agent i_1 and i_2 have non-zero utility mass at 0 and the support of u_i for $i \in \{i_1, i_2, i_3\}$ follow a hierarchy.

Lemma 2.2. *For any $\alpha \in \mathbb{R}^n$, the couplings defined in Eqs. (10) and (11) are solutions of the problem in Eq. (8) with value $\alpha^\top (\mathbb{E}[\tilde{Z}])_{i \in [n]}$.*

With this construction at hand, to specify an allocation strategy for the central planner, we may only specify the allocation function $\mathbf{p}(\cdot \mid \mathbf{U})$ and the directions $\alpha(\mathbf{U})$. By defining $\mathbf{W}(\cdot \mid \mathbf{U})$ via Eq. (11), the constraint Eq. (7) is automatically satisfied. The only constraint that remains to be checked to have a valid allocation is Eq. (3), which asks that the future promises stay within the constructed region.

Remark 2.3. *We are now ready to state our main results. For ease of exposition, we state all our main results in the following Sections 3 to 7 for equalitarian weights $\alpha = \mathbf{1}$. Note that these can easily be generalized to arbitrary weights $\alpha \in \mathbb{R}_+^n \setminus \{\mathbf{0}\}$ by replacing the valuation v_i of an agent $i \in [n]$ by $\alpha_i v_i$.*

2.5 Characterization of scenarios without social welfare gap

For technical purposes, we need to treat separately trivial cases when the regions $\mathcal{U}_0 = \mathcal{V}_1$ already contain an optimal utility vector. In words, these corresponds to cases when the optimal social welfare can be achieved with *memoryless* mechanisms which allocate based solely on reports at the current iteration. In the following, we show that this may only happen in two scenarios (see Fig. 1 for an illustration). The first one corresponds to the case when one agent's utilities dominates all others, while the second one involves some slightly more complex hierarchy between agents. In both cases, we can give simple optimal mechanisms:

Lemma 2.4. *First, we have $\mathcal{U}_0 = \mathcal{V}_1$. Next, the following are equivalent.*

1. *There exists an optimal vector in $\mathcal{U}_0$ (SWG(0) = 0). In particular, $\mathcal{U}_\gamma$ and $\mathcal{V}_T$ always have an optimal vector for all $\gamma \in [0, 1)$ and $T \geq 1$.*
2. *For all $i \in [n]$, either*
 - (a) *there exists $j \neq i$ such that $\mathbb{P}(u_j \geq u_i) = 1$,*
 - (b) *or there exists $[m_i, M_i]$ such that $\mathbb{P}(u_i \in [m_i, M_i] \cup \{0\}) = 1$ but $\mathbb{P}(u_j \in (m_i, M_i)) = 0$ for all $j \neq i$.*

3. Either one of these scenarios holds

- (a) There exists $i \in [n]$ such that for all $j \neq i$, $\mathbb{P}(u_i \geq u_j) = 1$. In which case, $\mathbb{E}[u_i]e_i \in \mathcal{U}_0$ is optimal, and is reached by the mechanism that always allocates to agent i .
- (b) There exists a subset of agents $S = \{i_1, \dots, i_k\} \subseteq \{i \in [n] : \alpha_i > 0\}$ and a sequence $m_0 = \infty > m_1 \geq m_2 \geq \dots m_{k-1} > 0$ for which: (1) If $l < k$, $\mathcal{D}_{i_l}$ is supported on $\{0\} \cup [m_l, m_{l-1}]$ and $\mathbb{P}_{u \sim \mathcal{D}_{i_l}}(u = 0) > 0$. (2) $\mathcal{D}_{i_k}$ is supported on $[0, m_{k-1}]$. (3) For $i \notin S$, we have $\mathbb{P}(u_i \leq u_{i_k}) = 1$.
In this case, $\mathbf{U} = \sum_{l \in [k]} \mathbb{P}(u_{i_{l'}} = 0, l' < l) \mathbb{E}[u_{i_l}]e_{i_l} \in \mathcal{U}_0$ is optimal, and is reached by the following allocation mechanism for reported utilities $\mathbf{v} \in [0, \bar{v}]^n$: $\mathbf{p}(\mathbf{v}) = \mathbf{e}_{i_{l(\mathbf{v})}}$ where $l(\mathbf{v}) = \min\{l : v_{i_l} > 0\} \cup \{k\}$.

The proof is given in Section A.1. In light of the second characterization of Theorem 2.4, and even when $\mathcal{U}_0$ does not contain an optimal vector, we note for agents $i \in [n]$ satisfying either condition 2(a) or 2(b), ensuring incentive-compatibility while optimizing for social welfare is somewhat straightforward (see Theorem A.1 for more details). Indeed, if condition 2(a) holds, then it is possible to ignore agent i , never allocating to them, and still yield an optimal utility vector by allocating to j instead. On the other hand, if condition 2(b) holds, the reports of agent i are also not really necessary to reach an optimal vector. Knowing the utilities u_j for the other agents $j \neq i$ and whether $u_i > 0$ is sufficient: if this is the case, we can allocate to agent i when $\max_{j \neq i} u_j < M_i$ and to an agent in $\arg \max_{j \neq i} u_j$ otherwise.

This motivates the definition of the complementary set of agents, for which more work is needed to ensure incentive-compatibility: $\tilde{I} := \{i \in [n] : i \text{ does not satisfy 2(a) nor 2(b)}\}$. As it turns out, ties between different agents are especially convenient for the central planner since it can freely decide how to break ties without incurring social welfare cost. It will therefore be useful to define the set of agents I augmented with these convenient agents, on which it suffices to focus.

$$I := \tilde{I} \cup \{i \in [n] : \mathbb{P}(u_i = \max_{j \neq i} u_j > 0) > 0\}. \quad (12)$$

3 A generic resource allocation mechanism

We are now ready to detail our generic resource allocation mechanism for the γ -discounted setting. The extension to the finite horizon- T setting is discussed in Section 6. The mechanism aims to achieve any utility vector within a pre-specified ball $B(\mathbf{x}, r)$ which has sufficiently large margin within the full-information region $\mathcal{U}^*$, that is, $B(\mathbf{x}, r + \delta) \subseteq \mathcal{U}^*$ for some margin parameter $\delta > 0$. We refer to Fig. 2 for an illustration. The corresponding *Promised Utility Mechanism* (PUM) requires an appropriate margin to successfully achieve utility vectors in the target ball, as detailed below.

Lemma 3.1. Fix $\gamma \in (0, 1)$. Let $\mathbf{x} \in \mathcal{U}^*$ and $r, \delta > 0$ such that $B(\mathbf{x}, r + \delta) \subseteq \mathcal{U}^*$. If

$$\frac{r\gamma}{1-\gamma} \geq \delta \geq C \frac{1-\gamma}{\gamma r}, \quad (13)$$

then the mechanism $\text{PUM}(\mathbf{x}, r, \delta)$ (see Algorithm 1) does not fail and achieves any utility vector within $B(\mathbf{x}, r)$. In particular, when this holds we have $B(\mathbf{x}, r) \subseteq \mathcal{U}_\gamma$. Here,

$$C = 4n^2 \bar{v}^2 \left(1 + \frac{\max_{i \in [n]} \mathbb{E}[u_i]}{\min_{i \in [n]} (\mathbb{E}[u_i] - x_i - r)} \right)^2.$$

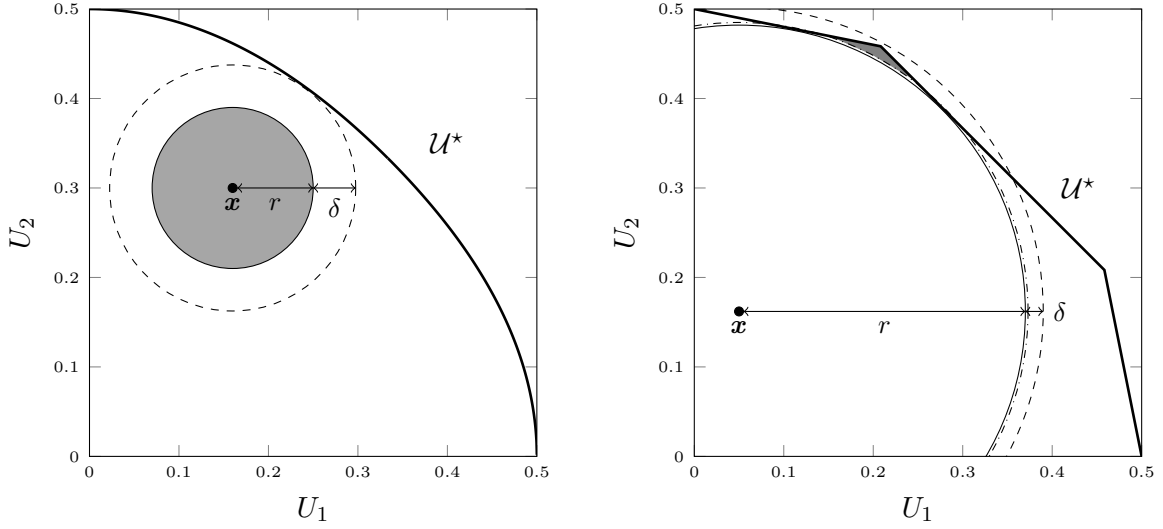


Figure 2: Illustration of Theorem 3.1 (left) and Theorem 4.1 (right) with 2 agents. (1) On the left, if the gray ball $B(\mathbf{x}, r)$ is within the full-information utility set $\mathcal{U}^*$ with a margin δ satisfying Eq (13), then $\text{PUM}(\mathbf{x}, r, \delta)$ is successful and $B(\mathbf{x}, r)$ is contained in $\mathcal{U}_\gamma$. The region $\mathcal{U}^*$ here corresponds to agents having uniform utility distributions on $[0, 1]$. (2) On the right, the ball $B(\mathbf{x}, r)$ separates the region $\mathcal{U}^*$ in two connected components, the one $\mathcal{C}$ closest to the boundary being included within a margin δ . If δ satisfies the constraints from Theorem 4.1, then the gray region is outside $\mathcal{U}_\gamma$ (note that the region is slightly smaller than the connected component $\mathcal{C}$ (corresponding to the dash-dotted ball)). The region $\mathcal{U}^*$ here corresponds to agents having uniform utility distributions on $\{1/6, 5/6\}$.

As remark, we aim to use this mechanism for balls $B(\mathbf{x}, r)$ that are close to the social-welfare-maximizing boundary of $\mathcal{U}^*$. Hence, the margin δ to this boundary intuitively corresponds to the suboptimality incurred by PUM compared to monetary mechanisms. Effectively, this reduces the multi-horizon resource allocation problem into a purely geometric one: finding the ball satisfying the margin condition Eq. (13) which is closest to the social-welfare-maximizing boundary of $\mathcal{U}^*$. Importantly, the required margin is smaller for larger balls—with larger radius r —which hints to the fact that smoothed boundaries for $\mathcal{U}^*$ will allow for smaller social-welfare gaps. This intuition will be formalized in the following sections.

We now describe the PUM resource allocation strategy. From the promised utility framework, one only needs to specify appropriate allocation and promised utility functions $\mathbf{p}(\cdot \mid \mathbf{U})$ and $\mathbf{W}(\cdot \mid \mathbf{U})$ for $\mathbf{U} \in B(\mathbf{x}, r)$. Ideally, for a point on the boundary $\mathbf{U} = \mathbf{x} + r\mathbf{y}$ with $\mathbf{y} \in S_{n-1}$, we aim to use the allocation function realizing the utility vector $\mathbf{x} + (r + \delta)\mathbf{y}$ in the full-information case and construct the promised utility function using the form described in Section 2.4 with $\boldsymbol{\alpha}(\mathbf{U}) = \mathbf{x} - \mathbf{U}$. The intuition is that this choice of allocation and promised utility function pushes the promised utility vector $\mathbf{W}(\mathbf{u} \mid \mathbf{U})$ towards the center of the ball. In practice, as per Eq. (11), we need to lower bound the second-largest component of $\boldsymbol{\alpha}(\mathbf{U})$, which will require considering a larger region $\mathcal{R} \supseteq B(\mathbf{x}, r)$ on which the allocation and promised utility functions are defined.

Input: Ball $B(\mathbf{x}, r)$, margin δ with $B(\mathbf{x}, r + \delta) \subseteq \mathcal{U}^*$, target utility $\mathbf{U}_0 \in B(\mathbf{x}, r)$

```

1 Initialize  $\mathbf{U} \leftarrow \mathbf{U}_0$ 
2 for  $t \geq 1$  do
3   Decompose  $\mathbf{U} = \mathbf{s} \odot \left( \sum_{i \in [n]} q_i \mathbb{E}[u_i] \mathbf{e}_i + q_0(\mathbf{x} + r\mathbf{y}) \right)$  where  $\mathbf{q} \geq \mathbf{0}, \sum_{i=0}^n q_i \leq 1$ ,
      $\mathbf{y} \in S_{n-1}, \mathbf{s} \in [0, 1]^n$ .
4   Fix an allocation function  $\tilde{\mathbf{p}}(\cdot)$  realizing utility  $\mathbf{x} + (r + \delta)\mathbf{y}$  in the full-information case
5   Compute interim functions  $\tilde{P}_i(\cdot) := \mathbb{E}_{\mathbf{u}_{-i}}[\tilde{p}_i(\cdot, \mathbf{u}_{-i})]$  and  $\tilde{W}_i(\cdot)$  according to Eq. (6) for
      $i \in [n]$ 
6   Observe reports  $\mathbf{v} = (v_i)_{i \in [n]}$  at round  $t$  and select allocation  $\mathbf{s} \odot \left( \sum_{i \in [n]} q_i \mathbf{e}_i + q_0 \tilde{\mathbf{p}}(\mathbf{v}) \right)$ 
7   Define  $\tilde{\mathbf{W}}$  using Eq. (11) for  $\boldsymbol{\alpha} = -r\mathbf{y}$ , and  $\tilde{Z}_i = \tilde{W}_i(v_i)$  for  $i \in [n]$ 
8   Update  $\mathbf{U} \leftarrow \mathbf{s} \odot \left( \sum_{i \in [n]} q_i \mathbb{E}[u_i] \mathbf{e}_i + q_0 \tilde{\mathbf{W}} \right)$ 
9   if  $\mathbf{U} \notin \mathcal{R}$  as defined in Eq. (14) then Mechanism fails. break;

```

Algorithm 1: Promised Utility Mechanism PUM($\mathbf{x}, r, \delta$) for the γ -discounted setting

Formal construction of the mechanism. Precisely, we construct these functions on the region

$$\mathcal{R} := \{\mathbf{y} : \mathbf{0} \leq \mathbf{y} \leq \mathbf{z}, \mathbf{z} \in \text{Conv}(\mathcal{U}^{NI}, B(\mathbf{x}, r))\}. \quad (14)$$

To give some intuition about this utility set, we have $\text{Conv}(\mathcal{U}^{NI}, B(\mathbf{x}, r)) = \text{Conv}(\mathbf{0}; \{\mathbb{E}[u_i] \mathbf{e}_i, i \in [n]\}; B(\mathbf{x}, r))$. This is the convex hull of the target ball with the trivial allocations that always allocate the resource to a single agent $i \in [n]$. Additionally, we extend this region in all directions $-\mathbf{e}_i$ for $i \in [n]$ so long as points stay within the positive orthant $\mathbb{R}_+^n$. For any $\mathbf{U} \in \mathcal{U}^*$, we fix $p(\cdot; \mathbf{U})$ an allocation function which reaches the utility vector $\mathbf{U}$ when having access to the full information of agent's utilities.

We start by specifying the allocation strategy for extreme points $\mathbf{U}$ of $\mathcal{R}$ that still belong to the ball $B(\mathbf{x}, r)$, for which we can use the strategy described above. Writing $\mathbf{U} = \mathbf{x} + r\mathbf{y}$ where $\mathbf{y} \in S_{n-1}$ is a unit vector, we define the allocation function via $p(\cdot \mid \mathbf{U}) := p(\cdot; \mathbf{x} + (r + \delta)\mathbf{y})$ and $\boldsymbol{\alpha}(\mathbf{U}) = \mathbf{x} - \mathbf{U} = -r\mathbf{y}$. We then extend the definition of the allocation function to the whole domain $\mathcal{R}$. To do so, note that any vector $\mathbf{U} \in \mathcal{R}$ can be characterized by $\mathbf{0} \leq \mathbf{U} \leq \mathbf{y}$ where $\mathbf{y} \in \text{Conv}(\mathcal{U}^{NI}, B(\mathbf{x}, r))$ and is also an extreme point of $\mathcal{R}$. In particular, $\mathbf{y}$ can be written as the convex combination of $\mathbf{0}$, vectors $\mathbb{E}[u_i] \mathbf{e}_i$ and some extreme point $\tilde{\mathbf{y}}$ of $\mathcal{R}$ within the ball $B(\mathbf{x}, r)$ (we only need one point from $B(\mathbf{x}, r)$ at most because the ball is strictly convex). Such a decomposition can be obtained via standard Carathéodory constructions. As a summary, we obtain

$$\mathbf{y} = \sum_{i \in [n]} q_i \mathbb{E}[u_i] \mathbf{e}_i + q_0 \tilde{\mathbf{y}}, \quad \text{where } \mathbf{q} \geq \mathbf{0}, \sum_{i \leq n} q_i \leq 1,$$

and $U_i = s_i y_i$ where $s_i \in [0, 1]$ for all $i \in [n]$. We then define the allocation and promised utility function via

$$p_i(\cdot \mid \mathbf{U}) = s_i \left(\sum_{i \in [n]} q_i \mathbf{e}_i + q_0 p_i(\cdot \mid \tilde{\mathbf{y}}) \right), \quad i \in [n].$$

$$W_i(\cdot \mid \mathbf{U}) = s_i \left(\sum_{i \in [n]} q_i \mathbb{E}[u_i] \mathbf{e}_i + q_0 W_i(\cdot \mid \tilde{\mathbf{y}}) \right), \quad i \in [n].$$

Note in particular that for all the extreme points $\mathbf{U} = \mathbb{E}[u_i]\mathbf{e}_i$, the allocation function $\mathbf{p}(\cdot \mid \mathbf{U}) = \mathbf{e}_i$ always gives the resource to agent i , as should be expected. The complete PUM strategy is summarized in Algorithm 1.

Remarks on Theorem 3.1. As an important remark, the term C from Theorem 3.1 should be viewed as a constant. Indeed, since we aim to bound the social welfare gap $\text{SWG}(\gamma)$ for non-monetary mechanisms, we can use Theorem 3.1 with small balls locally close to an optimal utility vector $\mathbf{U} \in \mathcal{U}^*$. For sufficiently small balls and close to this optimum point, the term C becomes

$$C \approx 4n^2\bar{v}^2 \left(1 + \frac{\max_{i \in [n]} \mathbb{E}[u_i]}{\min_{i \in [n]} (\mathbb{E}[u_i] - U_i)} \right)^2,$$

which is independent from γ . Further, the term C from Theorem 3.1 can be greatly simplified if one has additional structure in the problem, as exemplified by the following assumption.

Assumption 3.2. *The utility distributions of the agents $i \in [n]$ have support in $[\underline{v}, \bar{v}]$, where $\bar{v} \geq \underline{v} > 0$.*

Lemma 3.3. *Under Theorem 3.2, we can choose $C = \bar{v}^2(2n + \bar{v}/\underline{v})^2$ in the statement of Theorem 3.1.*

The sufficient condition from Theorem 3.1 can be generalized by splitting agents into groups and applying the resource allocation mechanism to each group separately. For simplicity, we defer the complete statement of this lower bound tool on the achievable region $\mathcal{U}_\gamma$ to the Appendix (Theorem B.2).

4 A geometric tool for proving lower bounds on the social welfare gap

In this section, we present our main tool for proving social welfare gaps between non-monetary and monetary mechanisms, for the γ -discounted infinite-horizon setting. Equivalently, we aim to show that some regions in the full-information region $\mathcal{U}^*$ cannot be achieved by non-monetary mechanisms.

As a reminder, Theorem 3.1 showed that any ball $B(\mathbf{x}, r)$ with sufficient margin δ within the full-information region $\mathcal{U}^*$ is achievable by non-monetary mechanisms. Intuitively, we show a corresponding negative counterpart: if the ball $B(\mathbf{x}, r)$ is too close to the boundary of $\mathcal{U}^*$, then the region between the ball and this boundary cannot be achieved. More formally, we show that if the ball $B(\mathbf{x}, r)$ is so close to the social-welfare-maximizing utility vectors of $\mathcal{U}^*$ that it cuts the region $\mathcal{U}^*$ in (at least) 2 components, under appropriate margin conditions, the social-welfare-maximizing component cannot be achieved (see the right figure of Fig. 2 for an illustration). The corresponding statement will therefore be *local*, that is, we focus on the region of $\mathcal{U}^*$ that is close to being optimal for the social welfare: $\{\mathbf{y} : \mathbf{1}^\top \mathbf{y} \geq \max_{\mathbf{z} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{z} - c_1\}$ for some $c_1 > 0$:

Lemma 4.1. *Suppose that $\mathcal{U}_0$ does not already contain an optimal vector (see Theorem 2.4 for a characterization). There exist two constants $c_1, c_2 > 0$ such that the following holds. For*

$\gamma \in (0, 1)$ and $r > 0$, suppose that $\mathcal{U}^* \setminus B^\circ(\mathbf{x}, r)$ has a connected component $\mathcal{C} \subseteq \{\mathbf{y} : \mathbf{1}^\top \mathbf{y} \geq \max_{\mathbf{z} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{z} - c_1\}$. If

$$\mathcal{C} \subseteq B(\mathbf{x}, r + \delta), \quad \delta = \min \left(c_2 \frac{1 - \gamma}{\gamma r}, r \right),$$

then the region $\tilde{\mathcal{C}} := \mathcal{C} \setminus B^\circ(\mathbf{x}, r + \frac{1 - \gamma}{\gamma} \delta)$ is not achievable by non-monetary mechanisms: $\mathcal{U}_\gamma \cap \tilde{\mathcal{C}} = \emptyset$. In particular, for any $\mathbf{U} \in \mathcal{C}$, writing $\mathbf{U} = \mathbf{x} + \|\mathbf{U} - \mathbf{x}\| \mathbf{d}$,

$$\max_{\mathbf{y} \in \mathcal{U}_\gamma} \mathbf{d}^\top (\mathbf{y} - \mathbf{x}) < r + \frac{1 - \gamma}{\gamma} \delta.$$

We make a few remarks about this result. First, the margin δ required ($\delta = \mathcal{O}(\frac{1 - \gamma}{\gamma r})$) has exactly the same form as the margin required for our positive results in Theorem 3.1. This will be essential to show that our mechanisms achieve near-optimal social welfare. Second, the condition $\mathcal{C} \subseteq \{\mathbf{y} : \mathbf{1}^\top \mathbf{y} \geq \max_{\mathbf{z} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{z} - c_1\}$, which localizes the region of interest, is mild since the constant $c_1 > 0$ is independent of γ (especially since it is possible to reach points that are c_1 -close to optimal for sufficiently small γ from the folk theorem). Last, ideally, Fig. 2 should prove that the full component $\mathcal{C}$ is not achievable by non-monetary mechanisms, however, for technical reasons, we prove this result for a slightly smaller region $\tilde{\mathcal{C}}$. Fortunately, this will not affect our near-optimality guarantees.

As for Theorem 3.1, the statement can be refined by considering a subset of agents. In particular, it is useful for future proofs to consider the setting in which a single agent $i \in [n]$ competes for the allocation against all other agents $[n] \setminus \{i\}$. The precise statement is deferred to the Appendix (Theorem C.3).

5 Social Welfare Gap Characterizations for the discounted infinite-horizon setting

Using the tools developed in the previous section Theorems 3.1 and 4.1 (and generalizations Theorems B.2 and C.3), we now derive our main characterizations of the social welfare gap for the discounted infinite-horizon setting.

5.1 A universal rate $\mathcal{O}(\sqrt{1 - \gamma})$

As a first consequence of Theorem 3.1, we show that for appropriately chosen parameters, the mechanism PUM always achieves social welfare gap at most $\mathcal{O}(\sqrt{1 - \gamma})$. As a remark, the folk theorem only ensures asymptotic decay of this social welfare gap. In comparison, to the best of our knowledge, this is the first result which achieves a quantitative universal convergence rate to first-best allocation as agents are more patient, irrespective of utility distributions or the number of agents.³

Theorem 5.1. *The optimal social welfare regret always satisfies $\text{SWG}(\gamma) = \mathcal{O}(\sqrt{1 - \gamma})$.*

³Gorokh et al. [2021] also provide general bounds for the finite T -horizon setting but for the weaker notion of approximate PBE, and therefore are not comparable to our perfect PBE guarantees.

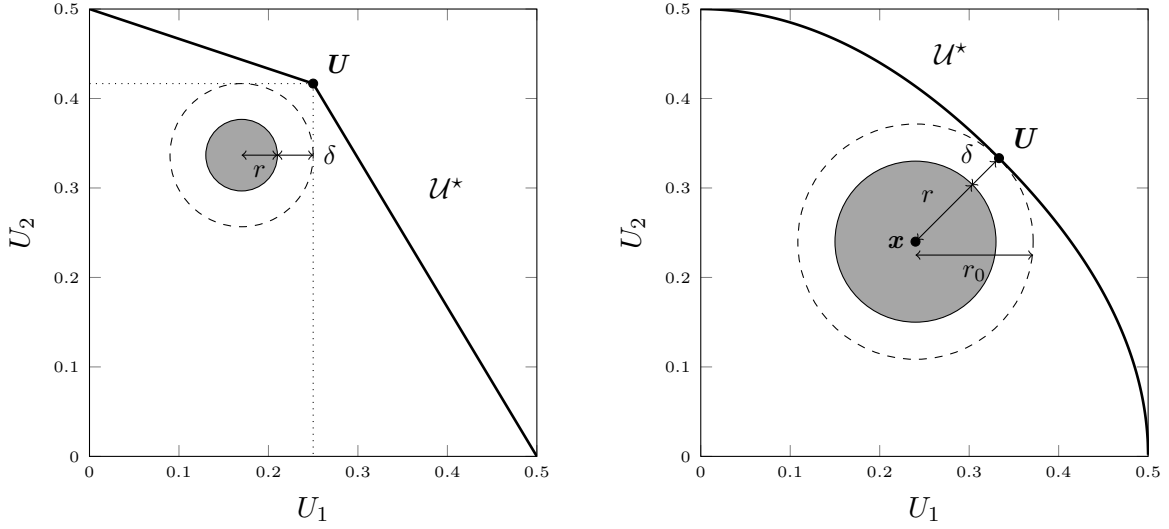


Figure 3: Visual proof of Theorem 5.1 (left) and Theorem 5.2 (right) given an optimal vector $\mathbf{U}$ with $U_i > 0$ for all $i \in [n]$. (1) On the left, the dotted region satisfies $\prod_{i \in [n]} [0, U_i] \subseteq \mathcal{U}^*$. We can then fit a ball $B(\mathbf{x}, r)$ as in the figure with $r = \delta = \sqrt{C \frac{1-\gamma}{\gamma}}$, satisfying Eq (13). Hence $B(\mathbf{x}, r) \subseteq \mathcal{U}_\gamma$ and as a result $d(\mathbf{U}, \mathcal{U}_\gamma) = \mathcal{O}(\sqrt{1-\gamma})$. (2) On the right, if $B(\mathbf{x}, r_0) \subseteq \mathcal{U}^*$ is tangent at $\mathbf{U}$, we can choose r, δ satisfying $r + \delta = r_0$ and $\delta = C \frac{1-\gamma}{\gamma r}$. These satisfy Eq (13). Note that $\delta \sim C \frac{1-\gamma}{r_0}$ as $\gamma \rightarrow 1$, which gives $d(\mathbf{U}, \mathcal{U}_\gamma) \leq \delta = \mathcal{O}(1-\gamma)$.

The proof of this result proceeds by showing that any vector $\mathbf{U} \in \mathcal{U}^*$ is at distance at most $\mathcal{O}(\sqrt{1-\gamma})$ from $\mathcal{U}_\gamma$. In the standard case where $\mathbf{U} \in \mathcal{U}^*$ satisfies $0 < U_i < \mathbb{E}[u_i]$ for all $i \in [n]$, Theorem 5.1 is a simple consequence of Theorem 3.1 of which Fig. 3 gives a visual proof. We first note that $\prod_{i \in [n]} [0, U_i] \subseteq \mathcal{U}^*$. To prove the desired result, it then suffices to construct a ball $B(\mathbf{x}, r)$ within this region and that satisfies the constraints Eq (13) with $r = \delta = \sqrt{C \frac{1-\gamma}{\gamma}}$ as depicted in the left figure of Fig. 3.

While this already shows that non-monetary mechanisms can have good non-asymptotic performance in general, the convergence rate of the social welfare gap may be significantly faster when the boundary of $\mathcal{U}^*$ is smooth. This is also captured by Theorem 3.1: the smoother the local boundary of $\mathcal{U}^*$, the larger the radius r of the ball that can be fit close to that boundary, and the smaller the margin $\delta \approx (1-\gamma)/r$ that is needed. In particular, in ideal cases where the ball has radius r independent of γ , Theorem 3.1 shows that the mechanism PUM yields smaller social welfare gap $\mathcal{O}(1-\gamma)$. We detail these scenarios in the following section.

5.2 Sufficient conditions for faster rates $\mathcal{O}(1-\gamma)$

We now discuss situations when fast social welfare gap convergence rates $\mathcal{O}(1-\gamma)$ can be achieved. Roughly speaking, these correspond to the ideal convergence rate: these are precisely the fastest rates achieved in Balseiro et al. [2019] for the case of $n = 2$ agents and well-behaved smooth utility distributions. Additionally, these will correspond to the fastest convergence rate that may carry over to the finite-horizon setting (see Section 5 and Theorem 6.2). From a practical perspective, understanding conditions for fast convergence rates is helpful to understand

which distributional properties are most desirable for the implementation of non-monetary mechanisms. As a warm-up, we start with a simple geometric condition. Intuitively, if the boundary of $\mathcal{U}^*$ is locally smooth at the neighborhood of optimal utility vectors, then one can achieve the faster rate:

Theorem 5.2. *Suppose that there exists an optimal vector $\mathbf{U} \in \mathcal{U}^*$ such that a ball tangent to $\mathcal{U}^*$ of radius $r_0 > 0$ is contained within $\mathcal{U}^*$, i.e., taking $\mathbf{x} = \mathbf{U} - r_0 \mathbf{1}$ we have $B(\mathbf{x}, r_0) \subseteq \mathcal{U}^*$. Then, $\text{SWG}(\gamma) = \mathcal{O}(1 - \gamma)$.*

For instance, Theorem 5.2 implies that whenever the boundary of $\mathcal{U}^*$ is twice-differentiable locally around an optimal point $\mathbf{U}$ then one can fit a ball tangent to $\mathbf{U}$ within the convex set $\mathcal{U}^*$ (see Theorem D.3 for a quick proof or Barb [2009, Theorem 1.0.9] for a stronger statement). This result is obtained by applying Theorem 3.1 with the parameters $\mathbf{x}$, and $r, \delta > 0$ such that $\delta = C \frac{1-\gamma}{\gamma r}$ (to satisfy Eq. (13)) and $r + \delta = r_0$. This gives a social welfare gap of $\delta = \Theta(1 - \gamma)$. This construction is depicted in the right figure of Fig. 3. While Theorem 5.2 gives a purely geometric condition, it can be translated into explicit conditions on the utility distributions: intuitively Theorem 5.2 requires that any two agents may simultaneously have near-optimal utilities with non-negligible probability. This condition can however be significantly generalized, as detailed below.

Our main characterization for fast convergence rates $\mathcal{O}(1 - \gamma)$ converts this infinite-horizon continuous incentive-compatibility problem into a purely *combinatorial* test: whether a certain directed graph $\mathcal{G}$ admits a partition into *strongly connected components*. We start by giving the main intuitions about this result then formally state it in Theorem 5.3. Intuitively, the graph $\mathcal{G}$ contains an edge $i \rightarrow j$ when the central planner can use agent i to enforce incentive-compatibility on agent j through added rewards/penalties. This flexibility typically arises when i and j share similar utilities: the central planner can then freely choose which agent to allocate to, without incurring significant social welfare cost. On each strongly-connected component of $\mathcal{G}$, we then show that the central planner can enforce incentive-compatibility at a low social welfare price. In turn, if $\mathcal{G}$ admits a partition into strongly connected components, the central planner can enforce incentive-compatibility for all agents simultaneously and achieve fast convergence $\mathcal{O}(1 - \gamma)$.

We now give the formal construction for the graph $\mathcal{G}$. As discussed in Section 2.5, it suffices to focus on the agents in I as defined in Eq. (12). For any pairs of agents $i, j \in I$, we introduce the function

$$f(\eta; i, j) := \mathbb{E} \left[u_i \mathbb{1}_{u_j = \max_{l \in [n]} u_l} \mathbb{1}_{u_j \in [u_i, (1+\eta)u_i]} \right].$$

This function quantifies to what extent (1) agent j has the highest utility among all agents and (2) agent i has near-highest utility, simultaneously. To give some intuitions, notice that under this event, the central planner has the flexibility to allocate the resource either to agent i or j at a small social welfare cost proportional to η . On the other hand, because of the incentive-compatibility constraints Eq. (4), the central planner needs to be able to penalize agent j on future rounds if they reported a high utility $v_j(t)$ at time t , and similarly, reward agent j on future rounds if they reported a low utility $v_j(t)$ (notice that the interim promised utility $W_j(v_j(t) | \mathbf{U})$ is non-increasing in $v_j(t)$ from Eq. (6)). Therefore, the higher the value of $f(\eta; i, j)$, the easier it is for the central planner to fulfill its future promises/penalties to agent j , leveraging agent i .

We then consider the following graph $\mathcal{G}$ on I such that for $i \neq j \in I$,

$$i \rightarrow j \iff f(\eta; i, j) \underset{\eta \rightarrow 0}{=} \Omega(\eta). \quad (15)$$

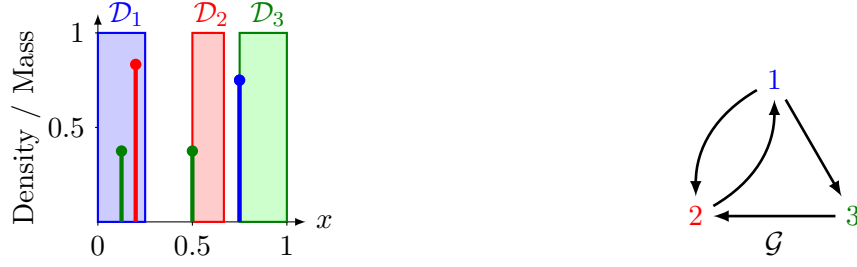


Figure 4: Utility distributions with $n = 3$ agents (left) and corresponding graph $\mathcal{G}$ (right). (1) On the left, utility distributions are a mixture of atoms—represented with bullets—and uniform distributions on intervals—represented in rectangles. E.g., $\mathcal{D}_1$ has an atom on $x = 3/4$ of mass $3/4$ and a density of 1 on $[0, 1/4]$. Formally, $\mathcal{D}_1 = \frac{1}{4}\text{Unif}([0, \frac{1}{4}]) + \frac{3}{4}\delta_{3/4}$, $\mathcal{D}_2 = \frac{1}{6}\text{Unif}([\frac{1}{2}, \frac{2}{3}]) + \frac{5}{6}\delta_{1/5}$, and $\mathcal{D}_3 = \frac{1}{4}\text{Unif}([\frac{3}{4}, 1]) + \frac{3}{8}\delta_{1/8} + \frac{3}{8}\delta_{1/2}$, where $\text{Unif}(I)$ and δ_x denotes the uniform over an interval I and the Dirac at x respectively. (2) On the right, the graph $\mathcal{G}$ as defined in Eq. (15). Intuitively, $i \rightarrow j$ when the central planner may select agent i while agent j had highest utility, without incurring large social welfare cost. Since $\{1, 2, 3\}$ is strongly-connected, Theorem 5.3 implies $\text{SWG}(\gamma) = \mathcal{O}(1 - \gamma)$.

Fig. 4 gives an illustration of $\mathcal{G}$ for concrete utility distributions. From the previous discussion, $i \rightarrow j$ may be interpreted as follows: the central planner can use agent i as a lever to flexibly ensure incentive-compatibility for agent j . In turn, we can show that fast convergence rates can be achieved if $\mathcal{G}$ is sufficiently connected.

Theorem 5.3. *Suppose that there exists a partition of agents $I = I_1 \sqcup I_2 \sqcup \dots \sqcup I_q$ such that for all $s \in [q]$, $|I_s| \geq 2$ and I_s is strongly connected in $\mathcal{G}$. Then, $\text{SWG}(\gamma) = \mathcal{O}(1 - \gamma)$.*

The proof of Theorem 5.3 uses the more involved version Theorem B.2 of Theorem 3.1 which allows to group agents according to the suitable partition. At the high-level, within each strongly-connected component I_i for $i \in [q]$, a central planner has enough flexibility to ensure incentive-compatibility for agents in I_i leveraging other agents from I_i . In particular, this requires each component I_i to have at least 2 elements, as stated in Theorem 5.3. Note that when $|I| \geq 5$, this strongly connected condition is stronger than simply assuming that for all $i \in I$, there exists $j, j' \in I$ such that $i \rightarrow j$ and $j' \rightarrow i$, as can be seen in Fig. 5.

The conditions in Theorem 5.3 hold under many standard distributional assumptions, as shown below.

Corollary 5.4. *Denote by m the infimum of the support of $\max_{i \in [n]} u_i$. Suppose that for any $i \in I$, there exists $j \in I$ with $j \neq i$ such that either*

- *there exists $m_i > 0$ such that $\mathbb{P}(u_i = m_i), \mathbb{P}(u_j = m_i) > 0$, and $\mathbb{P}(u_k \leq m_i) > 0$ for all $k \notin \{i, j\}$,*
- *or there exist $b_i > a_i > m$ such that both u_i and u_j have an absolutely-continuous part that has density bounded away from zero on $[a_i, b_i]$ by some constant.*

Then, $\text{SWG}(\gamma) = \mathcal{O}(1 - \gamma)$.

As a comparison, the main result from [Balseiro et al., 2019] implies that if the distributions $\mathcal{D}_1, \dots, \mathcal{D}_n$ have a density on $[0, \bar{v}]$ that is bounded above and below by constants, one can

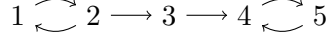


Figure 5: Example of graph for which each node has an incoming and outgoing edge but does not admit a partition into strongly connected components. For any partition of the nodes $\{1, \dots, 5\}$ into sets of size at least 2, node 3 is not strongly connected within its component.

achieve a $\mathcal{O}(1 - \gamma)$ convergence rate. In particular, Theorem 5.4 shows that this holds even if one only has densities bounded below.

5.3 Necessary conditions for faster rates $\mathcal{O}(1 - \gamma)$

We next also derive necessary conditions for having the faster rates $\mathcal{O}(1 - \gamma)$. Intuitively, we show that if $\mathcal{G}$ has an isolated node i.e., there exists $i \in I$ with no inner edges and no outer edges, then $\mathcal{O}(1 - \gamma)$ rates are impossible. Formally, we introduce the following function

$$g(\eta) := \min_{i \in I} \sum_{j \in I \setminus \{i\}} (f(\eta; i, j) + f(\eta; j, i)),$$

which quantifies the maximum ability of the central planner to fulfill its future promises/penalties for the worst agent $i \in I$ (leveraging any other agent). In particular, if this bound decays as $g(\eta) =_{\eta \rightarrow 0} o(\eta)$ then $\mathcal{G}$ has an isolated node.⁴ We then show that the condition $g(\eta) =_{\eta \rightarrow 0} o(\eta)$ is necessary for achieving fast rates for social welfare. This formalizes the intuition that if no agent can be leveraged to flexibly fulfill the future promise of some agent i , then achieving the fast rate $\mathcal{O}(1 - \gamma)$ is impossible.

Theorem 5.5. *If $g(\eta) =_{\eta \rightarrow 0} o(\eta)$, then $\text{SWG}(\gamma) = \omega(1 - \gamma)$.*

While this necessary condition does not exactly match the sufficient conditions from Theorem 5.3, these are written in similar format: enough flexibility to penalize/reward each agent, as quantified by the function f , is necessary to achieve faster rates.

5.4 Convergence rate bounds for general distributions

We are now ready to give our most general characterizations. While Theorem 5.3 and Theorem 5.5 focused on characterizing conditions for reaching the faster rate $\mathcal{O}(1 - \gamma)$, the methodology can be used to characterize any rate between $\sqrt{1 - \gamma}$ and $1 - \gamma$ more generally.

Recall that the function $f(\eta; i, j)$ measures the central planner's flexibility to meet future promises/penalties for agent j by leveraging agent i . We will show that a larger value of this function leads to a faster convergence rate of the social welfare. Specifically, when characterizing fast rates $\mathcal{O}(1 - \gamma)$, we focused on lower bounds on f of the form $f(\eta; i, j) = \Omega(\eta)$. In comparison, when weaker lower bounds are available, weaker convergence rates are still achievable.

Warm-up example with 2 agents. Before diving into the general multi-agent case, we start with a simpler example that illustrates important intuitions. We consider the 2-agent scenario with $I = \{1, 2\}$, and assume for simplicity that each agent gives similar amount of flexibility

⁴Note that converse also holds whenever f is non-pathological, e.g. if for all $i \neq j \in I$, $f(\eta; i, j) =_{\eta \rightarrow 0} \Theta(\eta^{\alpha_{i,j}})$ for some $\alpha_{i,j} \geq 0$.

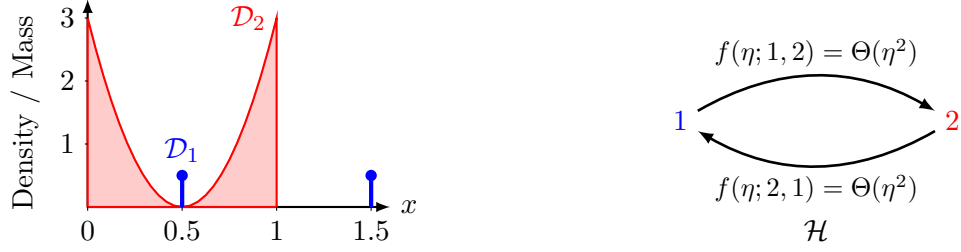


Figure 6: Utility distributions with $n = 2$ agents (left) and corresponding graph $\mathcal{H}$ (right). (1) On the left, $\mathcal{D}_1$ is uniform on $\{1/2, 3/2\}$, and $\mathcal{D}_2$ has density $3(1 - 2x)^2$ on $[0, 1]$. (2) By construction, the central planner incurs cost $\leq \eta$ when allocating to agent 1 while agent 2 had higher utility only with probability $\Theta(\eta^2)$, and vice versa: $f(\eta; 1, 2), f(\eta; 2, 1) = \Theta(\eta^2)$. In turn, this implies $f(\eta; \{1, 2\}), g(\eta) = \Theta(\eta^2)$. Hence, Theorem 5.8 (or Prop. 5.6) implies $\text{SWG}(\gamma) = \Theta((1 - \gamma)^{3/4})$.

for the central planner to enforce incentive-compatibility on the other agent in the following manner:

$$f(\eta; 1, 2), f(\eta; 2, 1) = \Theta(\eta^p), \quad \eta \in [0, 1],$$

for some parameter $p \geq 1$. Here, larger flexibility $f(\eta; 1, 2), f(\eta; 2, 1)$ correspond to smaller values of p . Note that in this scenario, the quantity $g(\eta) = \Theta(\eta^p)$ also coincides with $f(\eta; 1, 2)$ and $f(\eta; 2, 1)$ up to constant factors. We refer to Fig. 6 for a concrete example where $p = 2$.

The following result derives upper and lower bounds on convergence rates for this scenario. Specifically, it gives social welfare bounds between $\Theta(1 - \gamma)$ and $\Theta(\sqrt{1 - \gamma})$ depending on the flexibility of the central planner to enforce incentive-compatibility, as quantified by $f(\eta; 1, 2), f(\eta; 2, 1)$ through the parameter p . If the planner has larger flexibility to implement rewards/penalties—for smaller values of p —the convergence rate is faster.

Proposition 5.6. *Fix $p \geq 1$. Suppose that $I = \{1, 2\}$ and $f(\eta; 1, 2), f(\eta; 2, 1) = \Theta(\eta^p)$ for all $\eta \in [0, 1]$. Then, $\text{SWG}(\gamma) = \Theta((1 - \gamma)^{(1 + \frac{1}{p})/2})$.*

This gives tight social welfare gap characterizations in the case when the flexibility functions $f(\eta; 1, 2), f(\eta; 2, 1)$ are polynomial, interpolating between the fast rate $O(1 - \gamma)$ (when $p = 1$) and the slow rate $O(\sqrt{1 - \gamma})$ when the flexibility functions are smaller ($p \rightarrow \infty$).

General setting. Going from the 2-agent case to the general multi-agent setting requires a few notations. In particular, we first introduce a way to aggregate the flexibility functions $f(\eta; i, j)$ into a single quantity governing the social welfare convergence rate. Formally, let $\mathcal{H}$ be the graph on I with capacities $f(\eta; i, j)$ for any pair $i \neq j \in I$. We denote by $\mathcal{H}_{|S}$ this graph restricted to a subset of agents $S \subseteq I$. Now fix any partition $I = I_1 \sqcup I_2 \sqcup \dots \sqcup I_q$ with $|I_s| \geq 2$ for all $s \in [q]$, which we denote $\mathcal{P}$. For any $\eta \geq 0$, we define

$$f(\eta; \mathcal{P}) := \min_{s \in [q]} \min_{\substack{i, j \in I_s \\ i \neq j}} \left\{ \text{max-flow from } i \text{ to } j \text{ in } \mathcal{H}_{|I_s} \right\}.$$

We refer to Fig. 6 for a visualization of the graph $\mathcal{H}$ and how to derive $f(\eta; \mathcal{P})$ for a concrete example of utility distributions. Importantly, $f(\eta; \mathcal{P})$ measures the central planner's ability to meet future promises/penalties using the partition $\mathcal{P}$ to group agents, that is, only leveraging

agents from the same group to meet their incentive-compatibility constraints. Hence, we will show that the higher this quantity, the faster the convergence rate for the social welfare. In comparison, in the 2-player case, the only possible partition is to group all agents together $I_1 = \{1, 2\}$; accordingly we have $f(\eta; \mathcal{P}) = \min(f(\eta; 1, 2), f(\eta; 2, 1))$.

Remark 5.7. *In the definition of $f(\eta; \mathcal{P})$, the constraint on the max-flow between any two elements of a group I_s is analogous to the constraint that I_s was strongly connected within $\mathcal{G}$ in Theorem 5.3. Precisely, we can check that if a partition $\mathcal{P}$ satisfies the constraints for Theorem 5.3 then $f(\eta; \mathcal{P}) =_{\eta \rightarrow 0} \Omega(\eta)$, since there is a path between any two nodes $i \neq j \in I_s$ within I_s using only edges with capacities $\Omega(\eta)$.*

As in Theorem 5.3, the choice of the partition $\mathcal{P}$ allows for further flexibility for the central planner. Nevertheless, we always have $f(\eta; \mathcal{P}) \leq g(\eta)$ since the max-flow from or to some node is at most the total capacity of its incident edges. Hence, $g(\eta)$ intuitively upper bounds the ability of the central planner to ensure incentive-compatibility for agent i , leading to lower bounds on convergence rates.

We are now ready to derive general upper and lower bounds on convergence rates using $f(\eta; \mathcal{P})$ and $g(\eta)$. As in the 2-agent scenario, the following result gives social welfare upper bounds between $\Theta(1 - \gamma)$ and $\Theta(\sqrt{1 - \gamma})$ depending on the flexibility of the central planner to enforce incentive-compatibility as quantified by $f(\eta; \mathcal{P})$: the larger the function $f(\eta; \mathcal{P})$ —hence the more flexibility the planner has to implement rewards/penalties—the faster the resulting convergence. These upper bounds are complemented by corresponding lower bounds depending on $g(\eta)$. While upper and lower bounds do not match in general, they share a similar form and can be tight—e.g., see Theorem 5.6 and Section 7. We refer to Fig. 6 for a visualization of an explicit example when these yields tight rates strictly between $\Theta(1 - \gamma)$ and $\Theta(\sqrt{1 - \gamma})$.

Theorem 5.8. *There exist $C, C_\eta > 0$ (see Eq. (60) for their definition) such that the following holds. Fix any partition $I = I_1 \sqcup I_2 \sqcup \dots \sqcup I_q$ with $|I_s| \geq 2$ for all $s \in [q]$, denoted $\mathcal{P}$. Then, for any $\gamma \in [1/2, 1)$,*

$$\begin{aligned} \text{SWG}(\gamma) &\leq C(1 - \gamma) + C\sqrt{1 - \gamma} \cdot \inf \{1\} \cup \\ &\left\{ \eta \in (0, 1] : \forall \eta' \in [\eta, 1], f(\eta'; \mathcal{P}) \geq C_\eta \sqrt{1 - \gamma} \frac{\eta'}{\eta} \right\}. \end{aligned} \quad (16)$$

On the other hand, there exist $c_\eta, \eta_0, c > 0$ depending only on $\mathcal{D}_1, \dots, \mathcal{D}_n$ such that for any $\gamma \in [9/10, 1)$,

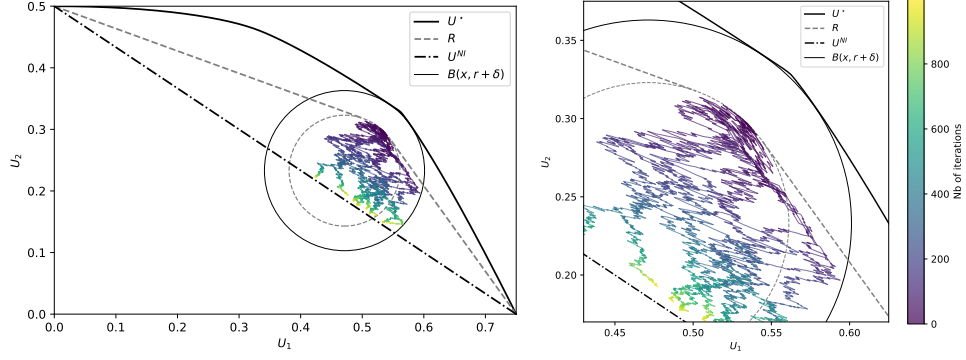
$$\begin{aligned} \text{SWG}(\gamma) &\geq c\sqrt{1 - \gamma} \cdot \\ &\sup \{0\} \cup \{\eta \in [0, \eta_0] : g(\eta) \leq c_\eta \sqrt{1 - \gamma}\}. \end{aligned} \quad (17)$$

This general social welfare characterization recovers all previous results: the universal convergence rate $\mathcal{O}(\sqrt{1 - \gamma})$ from Theorem 5.1, the sufficient conditions from Theorem 5.3 under which rates $\mathcal{O}(1 - \gamma)$ can be achieved,⁵ the corresponding necessary conditions from Eq. (17),⁶ and the 2-agent scenario Theorem 5.6.

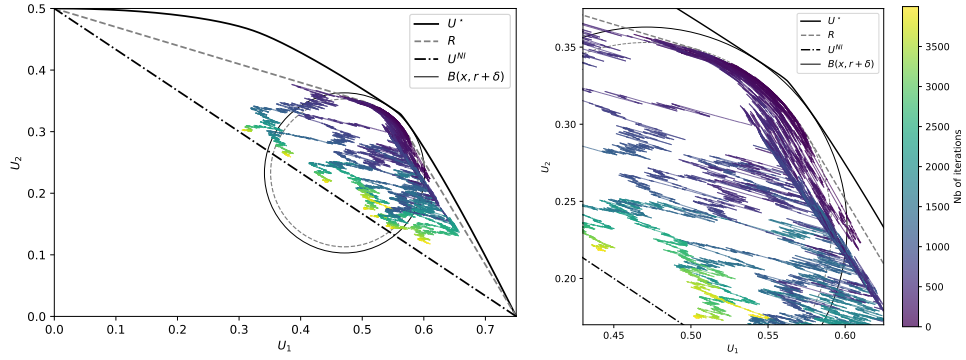
We now briefly discuss the main difference between the upper and lower bound from Theorem 5.8. The lower bound Eq. (17) corresponds to a 2-agent subproblem where one agent

⁵From Theorem 5.7, whenever $\mathcal{P}$ satisfies the constraints from Theorem 5.3 we have $f(\eta; \mathcal{P}) = \Omega(\eta)$.

⁶Theorem 5.8 shows that $g(\eta) =_{\eta \rightarrow 0} \Omega(\eta)$ is necessary for achieving rates $\mathcal{O}(1 - \gamma)$.



(a) $\delta = 0.04$ and 1000 iterations



(b) $\delta = 0.01$ and 4000 iterations

Figure 7: Sample trajectories of promised utilities for $n = 2$ agents and $\gamma = 0.99$. The full-information region $\mathcal{U}^*$ (resp. no-information region $\mathcal{U}^{NI}$) is drawn in bold solid (resp. dashed) black line. The ball $B(\mathbf{x}, r + \delta) \subseteq \mathcal{U}^*$ is drawn in regular black line. The region $\mathcal{R}$ (defined in Eq. (14)) on which the mechanism is implemented is drawn in dashed gray line. Figs. 7a and 7b each show 10 sample trajectories for a conservative margin $\delta = 0.04$ and a tighter margin $\delta = 0.01$ respectively. Plots on the right show a zoomed view of the same sample trajectories as their left plot. Crucially, with a tighter margin, trajectories in Figs. 7a and 7b linger along the boundary, where social welfare is the largest.

$i \in I$ plays against the best utility among all other agents $\max_{j \neq i} u_j$; while the upper bound Eq. (16) uses the full partition $\mathcal{P}$, hence considers more complex interactions between groups of agents. Therefore, the bounds intuitively coincide whenever there is no asymmetric advantage in clustering—see the 2-agent scenario from Theorem 5.6—or under specific partition-symmetric structure of agents.

5.5 Mechanism trajectory and efficiency

To illustrate our proposed mechanism for reaching the convergence rates from Theorem 5.8, we consider the same $n = 2$ agent setting as in Fig. 6, in which agent 1 has a discrete utility distribution and agent 2 has a uniformly-continuous utility distribution. Given a ball $B(\mathbf{x}, r)$ and an adequate margin δ such that $B(\mathbf{x}, r + \delta)$ remains within the full-information region $\mathcal{U}^*$, the mechanism implements the strategy described in Section 3 to fulfill promises within this ball

$B(\mathbf{x}, r)$.

In Fig. 7 we show sample trajectories of promised utilities under this mechanism. These trajectories follow the same general behavior: starting from the initial promised utility on the boundary of $B(\mathbf{x}, r)$, they slowly drift away from this boundary and asymptotically converge towards the boundary of the no-information region $\mathcal{U}^{NI}$ where promised utilities are straightforward to fulfill. Specifically, Fig. 7 shows trajectories for the same choice of ball $B(\mathbf{x}, r + \delta) \subseteq \mathcal{U}^*$ but different margins δ . Notably, if the mechanism converges for smaller values of this margin, its social welfare is increased since it may reach utility vectors closer to the boundary of $\mathcal{U}^*$. The margin δ for which our proofs ensure convergence are given by the constraints within Theorem 3.1. While the order of magnitude from these constraints cannot be improved in general given our proposed lower bounds, the specific constant C therein may be overly conservative. In Fig. 7a and Fig. 7b we show the convergence of the mechanism for a conservative margin and a tighter margin respectively. For conservative margins δ , the promised utilities easily drift towards $\mathcal{U}^{NI}$. In comparison, for tighter margins as in Fig. 7b, sample trajectories lingers close to the boundary of the region $B(\mathbf{x}, r)$ (first 500 iterations), where the algorithm achieves the highest social welfare efficiency.⁷

As a remark, promised utilities are not guaranteed to remain within the ball $B(\mathbf{x}, r)$. Indeed, we recall that the mechanism PUM from Section 3 instead considers an enlarged region $\mathcal{R}$ (see Eq. (14)).

6 From the discounted infinite-horizon setting to the finite-horizon setting

Up to minor changes, the mechanisms as well as the tools designed to characterize the social welfare gap in the discounted setting also carry to the finite-horizon T setting. At the conceptual level, our main finding is that the difference between the infinite-horizon discounted and finite-horizon settings is essentially *only* that in the finite-horizon case, the unavoidable last round T allocation alone forces $\Omega(1/T)$ social welfare gap (see Theorem 6.2). This is because on the last round, the central planner can no longer enforce future rewards/penalties: this corresponds to a one-shot resource allocation problem without money for which incentive-compatibility is impossible in general (see Myerson-Satterthwaite’s impossibility theorem). In comparison, in the infinite-horizon setting, much faster rates including exponential ones are possible (see Theorems 7.3 and 7.4): these are precluded by the last round allocation in the finite-horizon case.

A simple observation is that the utility regions $\mathcal{V}_T$ are still linked by a dynamic programming formulation. In turn, the promised utility framework still applies up to modifying the discount factor to be horizon-dependent $\gamma(T) = 1 - 1/T$. Naturally, for longer horizons, agents are more patient: $\gamma(T)$ is increasing in the horizon T . This is formalized by adapting Facts 2.2 and 2.3 to the finite-horizon case, which we defer to the Appendix (see E.1).

6.1 Lower bounds on the social welfare gap

We can now compare the regions $\mathcal{V}_T$ for $T \geq 1$ to the region $\mathcal{U}_{\gamma(T)}$.

Lemma 6.1. *For any $T \geq 1$, we have $\mathcal{V}_T \subseteq \mathcal{U}_{\gamma(T)}$.*

⁷This empirical behavior confirms the intuition from Theorem 3.1: the curvature of regions on which promised utility trajectories can be confined depends on the available margin δ . For smaller margins, smoother region with larger curvature radius r are required.

As a direct consequence of Theorem 6.1, all techniques developed earlier to give upper bounds on $\mathcal{U}_{\gamma(T)}$ also apply to $\mathcal{V}_T$. In particular, Theorem 4.1 (and Theorem C.3), Theorem 5.5, and the second claim Eq. (17) from Theorem 5.8 all hold by changing $\mathcal{U}_\gamma$ to $\mathcal{V}_T$ and γ to $\gamma(T) = 1 - 1/T$. In addition, we can easily prove that in the finite-horizon T setting, one cannot reach faster rates than $\mathcal{O}(1 - \gamma(T)) = \mathcal{O}(1/T)$. Intuitively, this is because on the last allocation iteration, the reports of the agents carry no information, hence the central planner can only reach a utility vector within $\mathcal{U}_0$ for this iteration, which incurs a non-zero (constant) regret.

Lemma 6.2. *Suppose that $\mathcal{U}_0$ does not already contain an optimal vector (see Theorem 2.4 for a characterization). Then, there exists $C > 0$ such that for all $T \geq 1$, we have $\text{SWG}(T) \geq \frac{C}{T}$.*

6.2 Mechanisms for the finite-horizon setting

We now show that the positive results from the γ -discounted setting can be extended to the finite-horizon case, yielding the same social welfare gap guarantees up to logarithmic factors. The main argument is the following, which shows that utility vectors within a ball with almost the same margin as in Theorem 3.1 within $\mathcal{U}^*$ are still achievable with an adjusted variant of the mechanism introduced in Section 3.

Lemma 6.3. *There exist $c_0 \geq 1$ and $r_0 > 0$ (depending only on the utility distributions) such that the following holds. Fix $T \geq 2$, $\mathbf{x} \in \mathcal{U}^*$, and $r, \delta > 0$ such that $B(\mathbf{x}, r + \delta) \subseteq \mathcal{U}^*$. If*

$$r_0 \geq r \geq \delta \geq c_0 C \frac{1 - \gamma(T)}{\gamma(T)r} \left(1 + \ln \frac{r}{\delta}\right), \quad (18)$$

where C is as defined in Theorem 3.1, then a variant of PUM($\mathbf{x}, r, \delta$) can achieve any utility vector in $B(\mathbf{x}, r)$.

In particular, we can simplify the statement as follows. Suppose that the assumptions of Theorem 3.1 are satisfied with $\gamma = \tilde{\gamma}(T) = 1 - \ln T/T$ and the constant $C' = c_0 C$, as well as $r_0 \geq r \geq \delta$. Then, $B(\mathbf{x}, r) \subseteq \mathcal{V}_T$.

As a result, all guarantees on the infinite-horizon mechanism—Theorems 5.1 and 5.3, Theorem 5.4, and the first claim Eq. (16) from Theorem 5.8—hold for the finite-horizon case by changing $\mathcal{U}_\gamma$ to $\mathcal{V}_T$ and γ to $1 - \ln T/T$.

For the sake of conciseness, we give the main structure of the variant mechanism here and defer the complete description to the Appendix (see Algorithm 2). It constructs a ball $B_t := B(\mathbf{x}^{(t)}, r_t) \subseteq \mathcal{V}_t$ with an appropriate margin δ_t within $\mathcal{U}^*$ and ensures that when t iterations remain, it only needs to fulfill a utility vector within this ball B_t . When only 1 iteration remains, the ball must start within the region $B_1 \subseteq \mathcal{U}^{NI}$; however, as the horizon $t \in [T]$ grows, the ball slowly converges to the target $B_T \approx B(\mathbf{x}, r)$. The specific trajectory depends on the parameters r and δ and is illustrated in Fig. 8. At the high-level, it proceeds in two phases: (1) first decrease the radius of the ball and the margin together, $\delta_t \approx r_t = \Theta(1/\sqrt{t})$, until the target radius r is reached, (2) then decrease the margin via $\delta_t = \Theta(1/(rt))$ while keeping the radius constant. This trajectory of the achievable balls within $\mathcal{V}_t$ for $t \in [T]$ illustrates the evolution of the strategy of the central planner throughout the sequential allocation problem.

In comparison to Theorem 6.1, in general Theorem 6.3 essentially requires a sub-optimal choice $\gamma = 1 - \ln T/T$; ideally, this should hold for $\gamma = \gamma(T)$ to match the social welfare gap lower bounds from Theorem 6.1. We note that in some cases, Theorem 6.3 only requires $\gamma = 1 - \mathcal{O}(1/T)$, yielding tight results compared to the γ -discounted setting (the extra factor

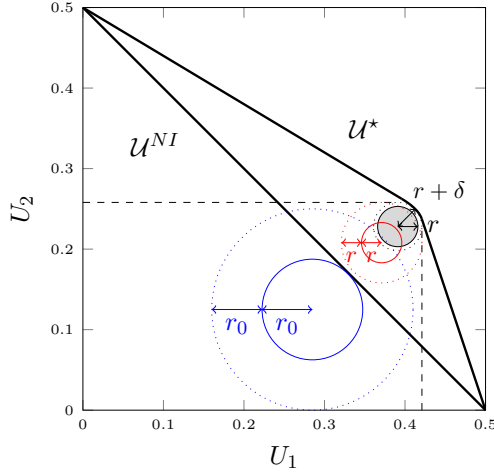


Figure 8: Illustration of the mechanism to approach the gray radius- r ball in the finite-horizon T setting (see Theorem 6.3). For all $t \in [T]$ we construct a ball within $\mathcal{V}_t$. The mechanism has two phases. For $t = 1$, we can take the blue radius- r_0 ball within $\mathcal{U}^{NI} \subseteq \mathcal{V}_1$. We first start by reducing the radius to r until we reach the red ball. As t grows, the radius reduces as $r_t = \sqrt{C \frac{1-\gamma(t)}{\gamma(t)}} = \sqrt{\frac{C}{t-1}}$, which requires conserving a margin $\delta_t = r_t$ depicted in the blue and red dotted circles. The second phase is to decrease the margin from $\delta_t = r$ (as shown in the red dotted circle) towards the desired margin δ , all while staying within the dashed region which belongs to $\mathcal{U}^*$. In this phase, the margin decreases as $\delta_t = C \frac{1-\gamma(t)}{\gamma(t)r}$.

$\ln \frac{r}{\delta}$ vanishes when $\delta = \Theta(r)$, see Theorem 7.5 for a concrete example). In other cases, however, this $\ln T$ term may not vanish, typically when utility distributions are somewhat smooth. While we believe that this logarithmic gap is essentially necessary, we leave this question open. The core of the problem is captured in the following simple scenario.

Open Question 1. Consider the $n = 2$ agent finite-horizon T resource allocation setting where the utility distributions are uniform $\mathcal{U}([0, 1])$ for both agents. Characterize the behavior (in T) of $\text{SWG}(T)$.

Our results show that the social welfare gap is between $\Omega(1/T)$ and $\mathcal{O}(\ln T/T)$. We conjecture that the answer is $\ln T/T$, which would imply that there is a separation between the infinite-horizon and finite-horizon settings beyond the fact that faster rates than $1/T$ are not achievable in the finite-horizon setting.

7 Examples and special cases

In this section, we apply the general methods described above to simple utility distributions of interest.

7.1 Utility distributions with density

We start with smooth utility distributions. Balseiro et al. [2019] shows that if the utility distributions have density upper and lower bounded by a constant on the domain $[0, \bar{v}]$ then the

convergence rate is $\Theta(1 - \gamma)$. Using the previous results, we can give further details about this convergence: in Theorem 5.4 we showed that having densities bounded below is sufficient to have a convergence rate $\mathcal{O}(1 - \gamma)$, and Theorem 5.8 implies that having densities bounded above is sufficient to prove that the rate of convergence is $\Omega(1 - \gamma)$.

Corollary 7.1. *Suppose that the utility distributions $\mathcal{D}_1, \dots, \mathcal{D}_n$ admit a density upper bounded by some constant on $[0, \bar{v}]$ and $\mathcal{U}_0$ does not already include an optimal vector. Then, $\text{SWG}(\gamma) = \Omega(1 - \gamma)$.*

If they admit a density lower bounded by some constant on $[0, \bar{v}]$, then $\text{SWG}(\gamma) = \mathcal{O}(1 - \gamma)$.

For the finite-horizon setting, we have the following.

Corollary 7.2. *Suppose that the utility distributions $\mathcal{D}_1, \dots, \mathcal{D}_n$ admit a density upper bounded by some constant on $[0, \bar{v}]$ and $\mathcal{U}_0$ does not already include an optimal vector. Then, for any $T \geq 1$, we have $\frac{c}{T} \leq \text{SWG}(T) \leq \frac{C \ln(T+1)}{T}$ for some constants $c, C > 0$.*

7.2 Discrete utility distributions

We next turn to the other extreme case of discrete utility distributions. In this case, Theorem 5.8 implies that the social welfare gap is typically $\Omega(\sqrt{1 - \gamma})$, which can be reached by Theorem 5.1. In some cases, however, when significant ties between agents arise, the convergence can be much faster (exponential) as detailed below. Geometrically, this corresponds to the case when the boundary of $\mathcal{U}^*$ is perfectly flat at social-welfare-maximizing utility vectors.

Theorem 7.3. *Suppose that all utility distributions $\mathcal{D}_1, \dots, \mathcal{D}_n$ are discrete. Suppose that $\mathcal{U}_0$ does not already contain an optimal vector. If there exists $i \in I$ such that for all $j \neq i$, $\mathbb{P}(u_i = u_j = \max_{l \in [n]} u_l > 0) = 0$, then $\text{SWG}(\gamma) = \Theta(\sqrt{1 - \gamma})$. Otherwise, there is a constant $c > 0$ such that $\text{SWG}(\gamma) = \mathcal{O}\left((1 - \gamma)e^{-c/\sqrt{1 - \gamma}}\right)$.*

The second case of Theorem 7.3 corresponds to the case when with non-zero probability there are (significant) ties between different agents. Ties are especially advantageous for the central planner because it can therefore decide which agent to allocate to without incurring any social welfare cost. As a result, exponential rates may also arise beyond purely discrete distributions, as detailed below.

Theorem 7.4. *If for all $i \in I$ there exists $j \neq i$ with $\mathbb{P}(u_i = u_j = \max_{l \in [n]} u_l > 0) > 0$, then there exists a constant $c > 0$ that only depends only on the utility distributions such that $\text{SWG}(\gamma) = \mathcal{O}\left((1 - \gamma)e^{-c/\sqrt{1 - \gamma}}\right)$.*

Achieving these faster rates requires an additional tool compared to Theorem 3.1. Indeed, applying directly this result would only prove a convergence rate $\mathcal{O}(1 - \gamma)$: this is the fastest rate that can be proved with Theorem 3.1 in its direct form. In practice, however, these can be combined with *gluing* arguments to yield faster rates when the frontier of $\mathcal{U}^*$ is particularly flat. Roughly speaking, as per Theorem 3.1, achieving exponentially small social welfare gap requires using balls with larger radius, which is prohibited by the diameter of $\mathcal{U}^*$. However, this issue may be solved by truncating this large-radius ball and gluing it to smaller-radius balls whose achievability can directly be guaranteed by Theorem 3.1. We give an illustration of this gluing strategy in Fig. 9. Repeatedly applying this process in turns yields exponentially small social welfare gaps.

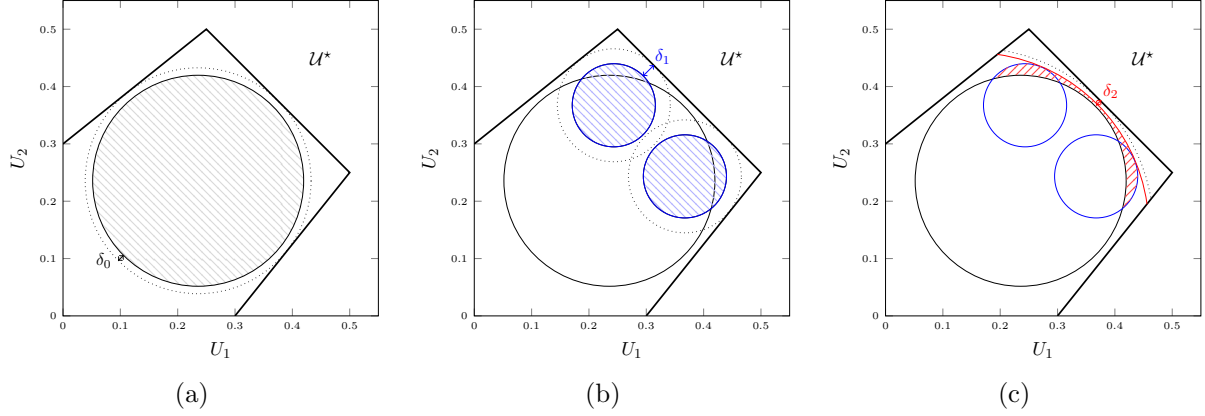


Figure 9: Illustration of how the *gluing* strategy achieves faster convergence rates when the frontier of the full-information region $\mathcal{U}^*$ is close to flat. The figure depicts a region $\mathcal{U}^*$ with an exaggerated piece-wise affine boundary for illustration purposes. In Fig. 9a, the dashed gray region corresponds to the region that can be shown to be achievable in $\mathcal{U}_\gamma$ using Theorem 3.1 directly, which is limited by the maximum-radius ball tangent to the region of optimal utility vectors that can be fit within $\mathcal{U}^*$ (represented with a dotted circle). The necessary margin δ_0 is inversely proportional to the radius of the corresponding black ball (Eq. (13)). The gluing strategy artificially increases this radius as follows: In Fig. 9b we consider the dashed blue balls which can be proved to be within $\mathcal{U}_\gamma$ directly using Theorem 3.1. Using the blue regions as anchor, we then glue the red arc in Fig. 9c, resulting in a smaller margin δ_2 . This results in the extra red-dashed region which we prove is within $\mathcal{U}_\gamma$. The gluing strategy can be used with an arbitrary number of layers to obtain exponentially small social welfare gap.

By Theorem 6.2, such exponential rates cannot be extended to the finite-horizon setting. In this case, we obtain the following social welfare gaps.

Corollary 7.5. *Suppose that all utility distributions $\mathcal{D}_1, \dots, \mathcal{D}_n$ are discrete and that $\mathcal{U}_0$ does not already contain an optimal vector. If there exists $i \in I$ such that for all $j \neq i$, $\mathbb{P}(u_i = u_j = \max_{l \in [n]} u_l > 0) = 0$, then $\text{SWG}(T) = \Theta(\frac{1}{\sqrt{T}})$. Otherwise, $\frac{c}{T} \leq \text{SWG}(T) \leq \frac{C \ln(T+1)}{T}$ for $T \geq 1$ for constants $c, C > 0$.*

8 Conclusion and future work

We consider a sequential resource allocation problem without monetary transfers, for any number of agents and general utility distributions and aim to understand the social welfare gap between monetary and non-monetary mechanisms. Our results uncover a rather surprising richness of the social welfare landscape, depending on geometrical properties of the utility distributions of agents—essentially the smoothness of the Pareto-frontier of the full-information achievable utility set. At the high-level, our results show that scenarios for which approximate ties between utilities of agents are frequent are most favorable for non-monetary mechanisms; this intuition is further quantified by corresponding social welfare gap characterizations. On the algorithmic side, we propose unified resource allocation mechanisms (PUM) based on the promised utility framework. Conveniently, while the social welfare gap depends heavily on utility distributions, it suffices to use this generic mechanism for appropriately distribution-dependent parameters

to achieve near-optimal social welfare gap guarantees among non-monetary mechanisms. Additionally, we unify the analysis of the infinite-horizon γ -discounted and finite T -horizon settings, which formalizes the links between both settings using analogous resource allocation mechanisms.

These general characterizations uncover new quantitative insights on the social welfare gap for non-monetary mechanisms. First, our mechanisms can always achieve a social welfare gap $\mathcal{O}(\sqrt{1-\gamma})$ and $\mathcal{O}(1/\sqrt{T})$ for the infinite-horizon and finite-horizon settings respectively. At the high-level this shows some quantitative efficiency of non-monetary mechanisms irrespective of the number of agents or their utility distributions; in comparison, the folk theorem ensures asymptotic convergence without specific rates. While this bound is typically tight for discrete distributions, considerably smaller social welfare gaps are possible depending on the geometry of the utility distributions: as small as $\tilde{\mathcal{O}}(1/T)$ for finite T -horizon, and even exponentially small in the infinite-horizon setting. Our quantitative bounds show that the main driver for the convergence rate is whether the central planner has enough flexibility to penalize or reward agents in future rounds, where this flexibility corresponds to scenarios when agents have similar utility for the resource.

While we focused on the standard resource allocation setting, other models without monetary transfers could also be studied. In particular, one could consider the setting in which the central planner also has a utility for the resource but still aims to maximize social welfare. This is equivalent to the mixed strategic/non-strategic setting in which a subset of the agents is strategic and at least one of the agents is not: their utilities are directly revealed to the central planner. As it turns out, our results also directly apply to this allocation problem. All lower bounds on the achievable utility region can still be achieved in this setting since the central planner has more information and hence can implement the same mechanisms. On the other hand, as emphasized in the proofs, all of our upper bounds on the achievable utility regions are derived in the setting for which a single agent $i \in [n]$ is strategic and all other agents $[n] \setminus \{i\}$ are not: the only incentive-compatibility constraint used is that for agent i . Because this is a simpler problem for the central planner, our upper bounds on the achievable utility region (showing that some regions cannot be achieved) still hold. In summary, all our characterizations are still valid for the mixed strategic/non-strategic allocation model.

Among other potential extensions, one could consider multiple resources, constructing mechanisms that do not rely on lotteries (in our constructions, the allocation at each iteration is a lottery on the simplex Δ_n), or considering correlations between the utilities of the agents for the resource. Our mechanisms also heavily rely on the assumption that the central planner has perfect knowledge of the geometry of the full-information utility set close to the optimal welfare boundary, determined by the agent utility distributions. Note that only the local geometry of the welfare-maximizing boundary is useful for our constructions. From our analysis, such perfect knowledge seems necessary to design near-optimal mechanisms given the intricate dependency of the optimal convergence rates on the boundary geometry. A natural question is to understand the trade-offs between maximizing social welfare efficiency and requiring minimal information from agent utilities for the central planner. As an initial remark, in light of the proof of Theorem 5.1, the universal rates $\mathcal{O}(\sqrt{1-\gamma})$ and $1/\sqrt{T}$ only require the knowledge of a single near-optimal utility vector in the full-information setting. We recall that optimal allocations in the full-information setting simply correspond to first-best allocations. Hence, the only information that is required is an estimate of the utility gathered by each agent under a first-best allocation. This is similar in flavor to mechanisms from artificial currencies which

require an estimate of the budget to allocate to each agent.

Another interesting research direction is to tighten the link between the infinite-horizon and finite-horizon settings. Our results show that there can be a clear separation between the two, since exponential convergence rates are possible for the infinite-horizon setting but not for the finite-horizon case: the gluing strategy we propose does not carry to the finite-horizon setting. Except for these exponential rates, however, our analysis shows that the main tools for the infinite-horizon setting can be used for the finite-horizon case up to logarithmic factors $\ln T$ by taking $\gamma(T) = 1 - 1/T$. While we give some examples in which there is no logarithmic gap between these two settings, it is unclear whether such a gap is necessary in general. We conjecture that this is the case, and formalize in Open Question 1 the simplest case for which such a gap may exist.

Acknowledgments.

This work is being partly funded by AFOSR grant FA9550-23-1-0182.

References

- D. Abreu, D. Pearce, and E. Stacchetti. Toward a theory of discounted repeated games with imperfect monitoring. *Econometrica: Journal of the Econometric Society*, pages 1041–1063, 1990.
- K. J. Arrow. *Social choice and individual values*, volume 12. Yale university press, 2012.
- I. Ashlagi and P. Shi. Optimal allocation without money: An engineering approach. *Management Science*, 62(4):1078–1097, 2016.
- S. Balseiro, H. Gurkan, and P. Sun. Multiagent mechanism design without money. *Operations Research*, 67(5):1417–1436, 2019.
- S. Barb. *Topics in geometric analysis with applications to partial differential equations*. PhD thesis, University of Missouri–Columbia, 2009.
- P. Bardhan and D. Mookherjee. Decentralizing antipoverty program delivery in developing countries. *Journal of public economics*, 89(4):675–704, 2005.
- E. Budish. The combinatorial assignment problem: Approximate competitive equilibrium from equal incomes. *Journal of Political Economy*, 119(6):1061–1103, 2011.
- S. Chakravarty and T. R. Kaplan. Optimal allocation without transfer payments. *Games and Economic Behavior*, 77(1):1–20, 2013.
- E. Clarke. Multipart pricing of public goods. *Public Choice*, 11:17–33, 1971.
- R. Cole, V. Gkatzelis, and G. Goel. Positive results for mechanism design without money. In *Proceedings of the 2013 international conference on Autonomous agents and multi-agent systems*, pages 1165–1166, 2013.
- D. Condorelli. What money cannot buy: Efficient mechanism design with costly signals. *Games and Economic Behavior*, 75(2):613–624, 2012.

- E. Elokda, S. Bolognani, A. Censi, F. Dörfler, and E. Frazzoli. A self-contained karma economy for the dynamic allocation of common resources. *arXiv preprint arXiv:2207.00495*, 2022.
- J. Feigenbaum, M. Schapira, and S. Shenker. Distributed algorithmic mechanism design. *Algorithmic Game Theory*, 14:363–384, 2007.
- E. Friedman, J. Halpern, and I. Kash. Efficiency and Nash equilibria in a scrip system for P2P networks. In *Proceedings of the 7th ACM Conference on Electronic Commerce*, pages 140–149, 2006.
- J. W. Friedman. A non-cooperative equilibrium for supergames. *The Review of Economic Studies*, 38(1):1–12, 1971.
- D. Fudenberg, D. Levine, and E. Maskin. The folk theorem with imperfect public information. In *A Long-Run Collaboration On Long-Run Games*, pages 231–273. World Scientific, 2009.
- A. Gibbard. Manipulation of voting schemes: A general result. *Econometrica*, 41(4):587–601, 1973.
- A. Gorokh, S. Banerjee, and K. Iyer. From monetary to nonmonetary mechanism design via artificial currencies. *Mathematics of Operations Research*, 46(3):835–855, 2021.
- T. Groves. Incentives in teams. *Econometrica*, 41(4):617–631, 1973.
- M. Guo and V. Conitzer. Strategy-proof allocation of multiple items between two agents without payments or priors. In *AAMAS*, pages 881–888, 2010.
- M. Guo, V. Conitzer, and D. M. Reeves. Competitive repeated allocation without payments. In *International Workshop on Internet and Network Economics*, pages 244–255. Springer, 2009.
- Y. Guo and J. Hörner. Dynamic mechanisms without money. Technical report, IHS Economics Series, 2015.
- L. Han, C. Su, L. Tang, and H. Zhang. On strategy-proof allocation without payments or priors. In *International Workshop on Internet and Network Economics*, pages 182–193. Springer, 2011.
- J. D. Hartline and T. Roughgarden. Optimal mechanism design and money burning. In *Proceedings of the fortieth annual ACM symposium on Theory of computing*, pages 75–84, 2008.
- H. C. Hoppe, B. Moldovanu, and A. Sela. The theory of assortative matching based on costly signals. *The Review of Economic Studies*, 76(1):253–281, 2009.
- M. O. Jackson and H. F. Sonnenschein. Overcoming incentive constraints by linking decisions 1. *Econometrica*, 75(1):241–257, 2007.
- K. Johnson, D. Simchi-Levi, and P. Sun. Analyzing scrip systems. *Operations Research*, 62(3):524–534, 2014.
- S. Kakade and J. Langford. Approximately optimal approximate reinforcement learning. In *Proceedings of the nineteenth international conference on machine learning*, pages 267–274, 2002.

- I. A. Kash, E. J. Friedman, and J. Y. Halpern. Optimizing scrip systems: Efficiency, crashes, hoarders, and altruists. In *Proceedings of the 8th ACM conference on Electronic commerce*, pages 305–315, 2007.
- I. A. Kash, E. J. Friedman, and J. Y. Halpern. An equilibrium analysis of scrip systems. *ACM Transactions on Economics and Computation (TEAC)*, 3(3):1–32, 2015.
- H. Levin, M. Schapira, and A. Zohar. Interdomain routing and games. In *Proceedings of the fortieth annual ACM symposium on Theory of computing*, pages 57–66, 2008.
- C. Linden, 2022. NRO internal document.
- A. Miralles. Cardinal bayesian allocation mechanisms without transfers. *Journal of Economic Theory*, 147(1):179–206, 2012.
- R. B. Myerson. Optimal auction design. *Mathematics of operations research*, 6(1):58–73, 1981.
- R. B. Myerson and M. A. Satterthwaite. Efficient mechanisms for bilateral trading. *Journal of economic theory*, 29(2):265–281, 1983.
- C. Ng. *Online Mechanism and Virtual Currency Design for Distributed Systems*. Harvard University, 2011.
- C. Prendergast. The allocation of food to food banks. *Journal of Political Economy*, 130(8):1993–2017, 2022.
- J. Russell and S. Ebbert. The high price of school assignment. *Boston Globe*, 12, 2011.
- N. Satterthwaite. Strategy-proofness and Arrow’s conditions: Existence and correspondence theorems for voting procedures and social welfare functions. *Journal of Economic Theory*, 10(2):187–217, 1975.
- J. Schummer and R. Vohra. Mechanism design without money. In *Algorithmic Game Theory*, chapter 10, pages 243–266. Cambridge University Press, 2007.
- S. E. Spear and S. Srivastava. On repeated moral hazard with discounting. *The Review of Economic Studies*, 54(4):599–617, 1987.
- J. Thomas and T. Worrall. Income fluctuation and asymmetric information: An example of a repeated principal-agent problem. *Journal of Economic Theory*, 51(2):367–390, 1990.
- W. Vickrey. Counterspeculation, auctions, and competitive sealed tenders. *The Journal of Finance*, 16(1):8–37, 1961.
- T. Walsh. Allocation in practice. In *Joint German/Austrian Conference on Artificial Intelligence (Künstliche Intelligenz)*, pages 13–24. Springer, 2014.

Appendix

In this Appendix, we give the proofs and further details on the results of the main paper. First, in Section A, we prove the preliminary results from Section 2 on the promised utility framework, including the proof of Theorem 2.4 which characterizes situations when there is no social welfare gap between monetary and non-monetary mechanisms. Next, we prove the guarantees from Section 3 in Section B on our proposed allocation resource mechanism, together with a strengthened version of this guarantee Theorem B.2. In Section C, we next prove our main geometric tool to prove lower bounds on the social welfare gap from Section 4, together with a strengthened version Theorem C.3. Next, we derive our main characterizations for the discounted infinite-horizon setting from Section 5 in Section D, including the universal $\mathcal{O}(\sqrt{1-\gamma})$ rate from Theorem 5.1, the characterizations for faster rates $\mathcal{O}(1-\gamma)$ from Theorems 5.3 and 5.5 and the full characterization from Theorem 5.8. As a consequence of the previous results, we then prove the results on specific cases from Section 7. In particular, we detail when and how exponential convergence rates can be obtained via the gluing strategy for the infinite-horizon setting. Last, we prove our results on the finite-horizon setting in Section E. In particular, we prove results from Section 6 showing how tools for the infinite-horizon setting can be used for the finite horizon. As a consequence of these derivations, we prove the results for the special cases in Section 7 in the finite-horizon setting.

Technical notations. For convenience, we introduce a few useful notations. For any $S \subseteq [n]$, denote by $P_S : \mathbb{R}^n \rightarrow \mathbb{R}^S$ the projection onto the coordinates of S . For $\mathbf{y} \in \mathbb{R}^S$, and $r > 0$, we also denote by $B_S(\mathbf{y}, r)$ the ball $B(\mathbf{y}, r)$ in $\mathbb{R}^S$. Next, we denote $Z = \max_{i \in [n]} u_i$, $Z_{-i} = \max_{j \neq i} u_j$ for $i \in [n]$. Last, for $S \subseteq [n]$ we denote $Z_S = \max_{i \in S} u_i$ and $Z_{-S} = \max_{i \notin S} u_i$.

A Proofs from Section 2

The results Facts 2.1, 2.2, and 2.3, and Theorem 2.2 can be considered as known in the literature, e.g. see Myerson [1981] and Balseiro et al. [2019] (Proposition 2.1, Proposition 2.2); these are given for completeness. First, we check that the utility regions are convex.

Proof of Fact 2.1 Let $\theta \in [0, 1]$. For any resource allocation setting, consider two realizable utilities $\mathbf{U}_1$ and $\mathbf{U}_2$. From the revelation principle there exist two incentive-compatible allocation strategies $\mathbf{S}_1$ and $\mathbf{S}_2$ realizing utilities $\mathbf{U}_1$ and $\mathbf{U}_2$ respectively. Consider the allocation strategy $\mathbf{S}$ which first samples a Bernoulli $B \sim \mathcal{B}(\theta)$, then implements $\mathbf{S}_1$ if $B = 1$ and $\mathbf{S}_2$ otherwise, for all times. By linearity, under truthful reporting, this allocation realizes the utilities $\theta \mathbf{U}_1 + (1 - \theta) \mathbf{U}_2$. It only remains to check that truthful reporting is still a PBE for $\mathbf{S}$, which is a consequence of the linearity of the value function V_i as detailed below. Fix any agent i , time t , utility $u_t(i)$, and history $\mathbf{h}(t)$. For any reporting strategy $\sigma_{i,t}$ for time t , we denote by $\sigma_{i,t} \circ \mathbf{truth}_i$ the concatenated strategy for agent i which uses strategy $\sigma_{i,t}$ at time t and truthful reporting for following rounds. We have

$$\begin{aligned}
& V_{i,t}(\mathbf{S}, \sigma_{i,t} \circ \mathbf{truth}_i, \mathbf{truth}_{-i} \mid u_t(i), \mathbf{h}(t)) \\
&= \mathbb{P}[B = 1 \mid u_t(i), \mathbf{h}(t)] \cdot V_{i,t}(\mathbf{S}_1, \sigma_{i,t} \circ \mathbf{truth}_i, \mathbf{truth}_{-i} \mid u_t(i), \mathbf{h}(t)) \\
&\quad + \mathbb{P}[B = 0 \mid u_t(i), \mathbf{h}(t)] \cdot V_{i,t}(\mathbf{S}_0, \sigma_{i,t} \circ \mathbf{truth}_i, \mathbf{truth}_{-i} \mid u_t(i), \mathbf{h}(t)) \\
&\geq \mathbb{P}[B = 1 \mid u_t(i), \mathbf{h}(t)] \cdot V_{i,t}(\mathbf{S}_1, \mathbf{truth}_i, \mathbf{truth}_{-i} \mid u_t(i), \mathbf{h}(t)) \\
&\quad + \mathbb{P}[B = 0 \mid u_t(i), \mathbf{h}(t)] \cdot V_{i,t}(\mathbf{S}_0, \mathbf{truth}_i, \mathbf{truth}_{-i} \mid u_t(i), \mathbf{h}(t)) \\
&= V_{i,t}(\mathbf{S}, \mathbf{truth}_i, \mathbf{truth}_{-i} \mid u_t(i), \mathbf{h}(t)),
\end{aligned}$$

since both $\mathbf{S}_1$ and $\mathbf{S}_2$ are incentive-compatible. In words, it is not advantageous to deviate from truthful reporting at a single round t from any state. Then, truthful reporting is a PBE by backwards induction (or e.g., using the performance difference lemma [Kakade and Langford, 2002]). ■

We next prove Fact 2.2 which gives necessary and sufficient conditions for a utility region $\mathcal{R}$ to be achievable in the γ -discounted setting.

Proof of Fact 2.2 We start with showing that the condition is sufficient to show $\mathcal{R}' \subseteq \mathcal{R} \subseteq \mathcal{U}_\gamma$. To do so, we use the allocation function $\mathbf{p}(\cdot | \mathbf{U})$ for $\mathbf{U} \in \mathcal{R}$ as the strategy at all times. Precisely, fix an arbitrary vector $\mathbf{U}_0 \in \mathcal{R}$. Suppose we have constructed $\mathbf{U}_{t-1}$ for $t \geq 1$, recalling the notation $\mathbf{v}(t)$ for the agent reports at time t , we pose

$$\mathbf{p}(t) := \mathbf{p}(\mathbf{v}(t) | \mathbf{U}_{t-1}) \quad \text{and} \quad \mathbf{U}_t := \mathbf{W}(\mathbf{v}(t) | \mathbf{U}_{t-1}).$$

It is well defined since by induction $\mathbf{U}_t \in \mathcal{R}$ for all $t \geq 0$ from Eq (3). We denote by $\mathbf{S}$ the constructed strategy. We now check that truthful reporting for all agents $\sigma = \mathbf{truth}$ forms a perfect Bayesian equilibrium.

To do so, we check that for all $i \in [n]$, any time $t \geq 0$, the expected remaining utility of agent i by reporting v at time t is exactly $W_i(v | \mathbf{U}_{t-1})$. Indeed, for any strategy σ_i for agent i , we have

$$\begin{aligned} & V_{i,t}(\mathbf{S}, \mathbf{truth}_{-i}, \sigma_i | u_i(t), \mathbf{h}(t)) \\ &= \sum_{t' \geq t} (1 - \gamma) \mathbb{E} \left[\gamma^{t'-t} u_i(t') \mathbb{E}_{\mathbf{u}_{-i}(t')} [p_i(\mathbf{u}_{-i}(t'), v_i(t') | \mathbf{U}_{t'-1})] | u_i(t), \mathbf{h}(t) \right] \\ &= \sum_{t' \geq t} (1 - \gamma) \mathbb{E} \left[\gamma^{t'-t} u_i(t') P_i(v_i(t') | \mathbf{U}_{t'-1}) | u_i(t), \mathbf{U}_{t-1} \right] \\ &\stackrel{(i)}{\leq} \sum_{t' \geq t} \gamma^{t'-t} \mathbb{E} [(1 - \gamma) u_i(t') P_i(u_i(t') | \mathbf{U}_{t'-1}) + \gamma W_i(u_i(t') | \mathbf{U}_{t'-1}) \\ &\quad - \gamma W_i(v_i(t') | \mathbf{U}_{t'-1}) | u_i(t), \mathbf{U}_{t-1}] \\ &\stackrel{(ii)}{=} (1 - \gamma) u_i(t) P_i(u_i(t) | \mathbf{U}_{t-1}) + \gamma (W_i(u_i(t) | \mathbf{U}_{t-1}) - W_i(v_i(t) | \mathbf{U}_{t-1})) \\ &\quad + \sum_{t' > t} \gamma^{t'-t} \mathbb{E} [U_{t'-1,i} - \gamma W_i(v_i(t') | \mathbf{U}_{t'-1}) | u_i(t), \mathbf{U}_{t-1}] \end{aligned}$$

In (i) we applied the incentive-compatibility equation Eq. (4), and in (ii) we took the expectation over $u_i(t') \sim \mathcal{D}_i$ within each expectation and used Eq. (2). Now note that by construction, for $t' \geq 1$ we have

$$W_i(v_i(t') | \mathbf{U}_{t'-1}) = \mathbb{E}_{\mathbf{u}_{-i}(t')} [W_i(\mathbf{u}_{-i}(t'), v_i(t') | \mathbf{U}_{t'-1})] = \mathbb{E}_{\mathbf{u}_{-i}(t')} [U_{t',i}].$$

As a result, with a telescoping argument, we obtain

$$V_{i,t}(\mathbf{S}, \mathbf{truth}_{-i}, \sigma_i | u_i(t), \mathbf{h}(t)) \leq (1 - \gamma) u_i(t) P_i(u_i(t) | \mathbf{U}_{t-1}) + \gamma (W_i(u_i(t) | \mathbf{U}_{t-1}),$$

where the only inequality used is that of (i), which is an equality for $\sigma_i = \mathbf{truth}$, that is when agent i reports truthfully. This exactly shows that Eq. (PBE) holds. We can now check that

under truthful reporting,

$$\begin{aligned} V_i(\mathbf{S}, \mathbf{truth}) &= \mathbb{E}[V_i, 1(\mathbf{S}, \mathbf{truth} \mid u_i(1))] \\ &\stackrel{(i)}{=} \mathbb{E}[(1 - \gamma)u_i(1)P_i(u_i(1) \mid \mathbf{U}_0) + \gamma(W_i(u_i(1) \mid \mathbf{U}_0))] \stackrel{(ii)}{=} U_{0,i}, \end{aligned}$$

where in (i) we used the previous computations and in (ii) we used Eq. (2). Therefore we checked that $\mathbf{U}_0 \in \mathcal{U}_\gamma$ for any arbitrary $\mathbf{U}_0 \in \mathcal{R}$.

We now prove that it is necessary. Let $\mathcal{R}$ be an achievable region, that is $\mathcal{R} \subseteq \mathcal{U}_\gamma$. By definition of $\mathcal{U}_\gamma$, for all $\mathbf{U} \in \mathcal{U}_\gamma$ there is a strategy $\mathbf{S}(\mathbf{U})$ that realizes $\mathbf{U}$. By the revelation principle, without loss of generality, truthful reporting is a perfect Bayesian equilibrium. We define $\mathbf{p}(\mathbf{v} \mid \mathbf{U}) := \mathbf{p}(1)$ where $\mathbf{p}(1)$ is the first allocation of the strategy $\mathbf{S}(\mathbf{U})$ having received reports $\mathbf{v}$. Let $\tilde{u}_i(t) \stackrel{i.i.d.}{\sim} \mathcal{D}_i$ be independent sequences for all $i \in [n]$. We then let

$$\mathbf{W}(\mathbf{v} \mid \mathbf{U}) := (1 - \gamma)\mathbb{E} \left[\sum_{t \geq 2} \gamma^{t-2} \tilde{u}_i(t) p_i(t) \mid \mathbf{v}(1) = \mathbf{v} \right], \quad \mathbf{v} \in [0, \bar{v}]^n.$$

That is, $\mathbf{W}(\mathbf{v} \mid \mathbf{U})$ is the expected remaining utility knowing that the first reports were $\mathbf{v}$. By construction, since $\mathbf{S}(\mathbf{U})$ realizes the utilities $\mathbf{U}$, Eq. (2) is satisfied. Next, note that for any first reports $\mathbf{v}(1) = \mathbf{v}$, the strategy $\mathbf{S}(\mathbf{U})$ starting from time 2 exactly realizes the utility vector $\mathbf{W}(\mathbf{v} \mid \mathbf{U})$ (the allocation problem is time-invariant): Eq. (3) holds. Last, we assumed that truthful reporting satisfies Eq. (PBE). Consider any strategy σ_i which reports an arbitrary value $v_i(1) = v(u_i(1))$ but truthfully reports at all times $t \geq 2$. Taking the expectation of Eq. (PBE) for strategy σ_i and $t = 1$, over the realization of the first allocation $\mathbf{p}(1)$ shows that Eq. (4) holds. ■

We now turn to the second characterization from Fact 2.3, which characterizes the form of the interim utility promise, following an argument from Myerson [1981].

Proof of Fact 2.3 Fix $\gamma \in (0, 1)$. Given Fact 2.2, it suffices to show that Eqs. (2) to (4) are equivalent to Eqs. (3) and (7).

We start by supposing that Eqs. (2) to (4) hold. We only need to show Eq. (7). Fix $\mathbf{U} \in \mathcal{R}$. For readability we omit $\mid \mathbf{U}$ from all functions below. First, for any $u, v \in [0, \bar{v}]$ with $u \leq v$, applying Eq. (4) to u and v implies that

$$\frac{1 - \gamma}{\gamma}(P_i(u) - P_i(v))v \leq W_i(v) - W_i(u) \leq \frac{1 - \gamma}{\gamma}(P_i(u) - P_i(v))u.$$

Note that the function $P_i(\cdot)$ takes values in $[0, 1]$ since $\mathbf{p}(\cdot)$ takes values in Δ_n . Also, the previous equation implies in particular that $P_i(\cdot)$ is non-decreasing and that $W_i(\cdot)$ is constant on any interval on which $P_i(\cdot)$ is constant. We can therefore sum the previous equations for $u = u_i$ and $v = u_{i+1}$ for regular partitions $0 \leq u_1, \dots, u_k = w$ for a fixed $w \in [0, \bar{v}]$ interpreting the sums as Riemann integrals (possible because P_i is non-decreasing hence well-behaved) which implies

$$W_i(w) - W_i(0) = -\frac{1 - \gamma}{\gamma} \int_{P_i(0)}^{P_i(w)} P_i^{<-1>}(y) dy = -\frac{1 - \gamma}{\gamma} \left(w P_i(w) - \int_0^w P_i(x) dx \right) \quad (19)$$

This proves Eq. (5). We next use Eq. (2) to compute

$$\begin{aligned}
U_i &= \mathbb{E}_{u_i}[(1 - \gamma)u_i P_i(u_i) + \gamma W_i(u_i)] \\
&= \gamma W_i(0) + (1 - \gamma) \mathbb{E}_{u_i} \left[\int_0^{u_i} P_i(x) dx \right] = \gamma W_i(0) + (1 - \gamma) \mathbb{E}_{u_i} \left[\int_0^{\bar{v}} P_i(x) \mathbb{1}_{u_i \geq x} dx \right] \\
&= \gamma W_i(0) + (1 - \gamma) \int_0^{\bar{v}} \bar{F}_i(x) P_i(x) dx.
\end{aligned}$$

Plugging this formula for $W_i(0)$ into Eq. (19) gives the desired formula for the interim allocation Eq. (6). Eq. (7) holds by definition of the interim promise $W_i(\cdot \mid \mathbf{U})$.

Now suppose that Eqs. (3) and (7) hold. Using the previous computations, we note that Eq. (2) is satisfied. It only remains to check the incentive-compatibility constraint Eq. (4). Having fixed $\mathbf{U} \in \mathcal{R}$, for any $i \in [n]$ and $u, v \in [0, \bar{v}]$, we have

$$\begin{aligned}
(1 - \gamma)u P_i(v) + \gamma W_i(v) &= U_i + (1 - \gamma) \left(P_i(v)(u - v) + \int_0^v P_i(x) dx - \int_0^{\bar{v}} \bar{F}_i(x) P_i(x) dx \right) \\
&= U_i + (1 - \gamma) \left(\int_0^u P_i(x) dx - \int_0^{\bar{v}} \bar{F}_i(x) P_i(x) dx \right) + (1 - \gamma) \int_u^v (P_i(x) - P_i(v)) dx.
\end{aligned}$$

Only the last term depends on v , hence the quantity is maximized for $v = u$ because $P_i(\cdot)$ is non-decreasing. This proves Eq. (4) and ends the proof. $\blacksquare$

Last, we prove Theorem 2.2 which shows that the couplings from Eqs. (10) and (11) are optimal.

Proof of Theorem 2.2 From Eq. (9), the optimal value of Eq. (8) is upper bounded by $\alpha^\top \mathbb{E}[\tilde{\mathbf{Z}}]$. For both couplings in Eqs. (10) and (11) we can easily check that for all realization ω , $\mathbf{Z} \in \{\mathbf{y} : \alpha^\top \mathbf{y} = \alpha^\top \mathbb{E}[\tilde{\mathbf{Z}}]\}$. Next, by independence of the variables $\tilde{Z}_i$ for $i \in [n]$, we can also prove that in both cases for all $i \in [n]$, $\mathbb{E}[Z_i \mid \tilde{Z}_i] = \tilde{Z}_i$. Hence these couplings satisfy the constraints and achieve the value $\alpha^\top \mathbb{E}[\tilde{\mathbf{Z}}]$ which ends the proof. $\blacksquare$

A.1 Characterization of cases with no social welfare gap

We now turn to the main result in this section, which is the Theorem 2.4. This gives a characterization of cases when $\mathcal{U}_0 = \mathcal{V}_1$ already contains an optimal utility vector.

Proof of Theorem 2.4 First note that the perfect Bayesian equations Eq. (PBE) characterizing $\mathcal{V}_1$ for $t = T = 1$ are identical to those for $\mathcal{U}_0$ for $t = 1$, which is the only time step that matters for the total utility $\mathbf{V}(\mathbf{S}, \boldsymbol{\sigma})$ since $\gamma = 0$. Hence we directly have $\mathcal{U}_0 = \mathcal{V}_1$. We now turn to the main claim of the result.

Proving 1. $\Rightarrow$ 2. We introduce the preorder $\succeq$ such that $i \succeq j$ if and only if $\mathbb{P}(u_i \geq u_j) = 1$ (it is reflexive and transitive, but not antisymmetric). Let $T = \{i \in [n], \nexists j \neq i, j \succeq i\}$, which gives

$$\forall i \in T, \forall j \neq i, \quad \mathbb{P}(u_i > u_j) > 0. \quad (20)$$

This implies in particular that $\mathbb{P}(u_i = 0) < 1$. For every $i \in T$, recall the notation $\text{Supp}(\mathcal{D}_i)$ for the support of $\mathcal{D}_i$. We define

$$m_i := \inf(\text{Supp}(\mathcal{D}_i) \setminus \{0\}) \quad \text{and} \quad M_i := \sup(\text{Supp}(\mathcal{D}_i) \setminus \{0\}). \quad (21)$$

Note that these are well-defined since $\text{Supp}(\mathcal{D}_i) \setminus \{0\} \neq \emptyset$, otherwise $\mathbb{P}(u_i = 0) = 1$. In particular, $M_i > 0$. By assumption, there exists an optimal vector $\mathbf{U}_0 \in \mathcal{U}_0$. Hence, by the revelation principle, there is an allocation function $\mathbf{p}(\cdot)$ that reaches $\mathbf{U}_0$ and is incentive-compatible as per Eq. (4). Now fix $i \in T$. Because $\mathbf{U}_0$ is an optimal vector, with probability one on $\mathbf{u}$, we have $p_i(\mathbf{u}) > 0$ only if $i \in \arg \max_{j \in [n]} u_j$. We also recall that $P_i(\cdot)$ is non-decreasing by hypothesis. As a result, for all $u \in \text{Supp}(\mathcal{D}_i)$ we have

$$\mathbb{P}(Z_{-i} < u) \leq P_i(u) \leq \mathbb{P}(Z_{-i} \leq u). \quad (22)$$

In particular, this holds for $u \in \{m_i, M_i\}$. In our setting, because $\gamma = 0$, the incentive-compatibility constraint implies that for all $u \in (0, \bar{v}]$, $u \in \arg \max_{v \in [0, \bar{v}]} u P_i(v)$ where $P_i(v_i) = \mathbb{E}[p_i(v_i, \mathbf{u}_{-i})]$. This implies that P_i is constant on $(0, \bar{v}]$ for all $i \in [n]$. Together with Eq. (22), this shows that

$$\mathbb{P}(Z_{-i} < M_i) \leq P_i(M_i) = \lim_{x \rightarrow m_i^+} P_i(x) \leq \lim_{x \rightarrow m_i^+} P_i(Z_{-i} \leq x) = \mathbb{P}(Z_{-i} \leq m_i). \quad (23)$$

Hence, either $m_i = M_i$, or $\mathbb{P}(Z_{-i} \leq m_i) = \mathbb{P}(Z_{-i} < M_i)$. In both cases, we obtained $\mathbb{P}(Z_{-i} \in (m_i, M_i)) = 0$, while $\text{Supp}(\mathcal{D}_i) \subseteq [m_i, M_i] \cup \{0\}$ hence $\mathbb{P}(u_i \in [m_i, M_i] \cup \{0\}) = 1$. For any $j \neq i$,

$$\begin{aligned} 0 &= \mathbb{P}(Z_{-i} \in (m_i, M_i)) \geq \mathbb{P}(u_j \in (m_i, M_i) \text{ and } \forall k \notin \{i, j\}, u_k < M_i) \\ &= \mathbb{P}(u_j \in (m_i, M_i)) \prod_{k \notin \{i, j\}} \mathbb{P}(u_k < M_i). \end{aligned}$$

However, Eq. (20) implies that for all $k \neq i$, we have $\mathbb{P}(u_k < M_i) > 0$. Together with the previous equation this gives $\mathbb{P}(u_j \in (m_i, M_i)) = 0$, which holds for all $j \neq i$.

Proving 2. $\Rightarrow$ 3. We suppose that the second statement of Theorem 2.4 holds. Further, suppose that scenario 3(a) does not hold, that is, $([n], \succeq)$ does not have a maximum element. We now construct the set M as the set of maximal elements, that is

$$M = \{i \in [n] : j \succeq i \Rightarrow i \succeq j\}.$$

On M , the preorder $\succeq$ is symmetric, hence is an equivalence relation and induces equivalence classes. We define S as a set containing exactly one representative for each equivalence class of $\succeq$ on M . Note that for any $i \neq j \in S$ we have neither $i \succeq j$ nor $j \succeq i$. By assumption, we also have $|S| \geq 2$. For any $i \in M$, we have $\mathbb{P}(u_i = 0) < 1$, otherwise for all $j \in M$, $j \succeq i$ and there would be only one equivalence class. Let T be the union of equivalence classes that are singletons. Note that T coincides with the definition from the proof that 1. $\Rightarrow$ 2., hence elements of T do not satisfy assumption 2(a) as seen in Eq. (20). Hence they must satisfy assumption 2(b): defining m_i, M_i as in Eq. (21),

$$\mathbb{P}(u_i \in [m_i, M_i] \cup \{0\}) = 1 \quad \text{and} \quad \mathbb{P}(u_j \in (m_i, M_i)) = 0, \quad j \neq i. \quad (24)$$

For $i \in M \setminus T$, there exists $j \in M$ with $j \neq i$ in its equivalence class. Then, $i \succeq j$ and $j \succeq i$, which implies that the distributions $\mathcal{D}_i$ and $\mathcal{D}_j$ are identical and are equal to a Dirac. We can define $m_i = M_i > 0$ as the value such that $\mathbb{P}(u_i = m_i) = 1$. Note that with these values, Eqs. (21) and (24) hold for all $i \in M$. Note that $M \setminus T$ only contains one equivalence class, otherwise we would have $i, j \in M \setminus T$ with say $m_i < m_j$ which implies that i is not maximal. In summary, $|S \setminus T| \leq 1$.

We next prove that for all $i, j \in S$, $M_j \notin (m_i, M_i)$. Otherwise, by definition of M_j in Eq. (21) for all $\epsilon > 0$ we have $\mathbb{P}(u_j \in (m_i, M_j)) > 0$ contradicting Eq. (24). Similarly, we have $m_j \notin (m_i, M_i)$. This implies that all intervals (m_i, M_i) for $i \in S$ are disjoint. We order them by increasing order $S = \{i_1, \dots, i_k\}$ where $M_{i_1} \geq m_{i_1} \geq M_{i_2} \geq m_{i_2} \geq \dots \geq M_{i_k} \geq m_{i_k}$.

For any $i \in T$ and $j \in S \setminus T$, we must have $m_j \leq m_i$, otherwise $m_j \geq M_i$ which implies $j \succeq i$ since $\mathbb{P}(u_j = m_j) = 1$. This shows that the element $j \in S \setminus T$ if it exists, appears last $j = i_k$. For any $l \in [k-1]$, we have $i_l \in T$ so $\text{Supp}(\mathcal{D}_{i_l}) \subseteq [m_{i_l}, M_{i_l}] \cup \{0\}$. Then, we must have $\mathbb{P}(u_{i_l} = 0) > 0$, otherwise we would have $i_l \succeq i_{l+1}$. We also note that $m_{i_{k-1}} \geq M_{i_k} > 0$. By construction of S , for any $i \notin S$ there must exist $j \in S$ such that $j \succeq i$. If $j = i_l$ where $l \in [k-1]$ we must have $\mathbb{P}(u_i = 0) = 1$ since we just proved that $\mathbb{P}(u_{i_l} = 0) > 0$. Hence, in all cases $i_k \succeq i$.

Up to renaming m_{i_l} by m_l , this ends the proof of the characterization 3(b). It is straightforward to check that the proposed allocation function $\mathbf{p}(\cdot)$ is incentive-compatible and reaches the utility vector

$$\mathbf{U} = \sum_{l \in [k]} \mathbb{P}(u_{i_l} = 0, l' < l) \mathbb{E}[u_{i_l}] \mathbf{e}_i.$$

In particular, $\mathbf{U} \in \mathcal{U}_0$. We can easily check that it always allocates to the agent within $\{i_l, l \in [k]\}$ with maximum value, which is sufficient since $i_k \succeq i$ for all $i \notin S$. This ends the proof of 2. $\Rightarrow$ 3.

Proving 3. $\Rightarrow$ 1. This is immediate because the vector $\mathbf{U} \in \mathcal{U}_0$ is optimal. ■

When $\mathcal{U}_0$ does not contain an optimal vector, the characterization from Theorem 2.4 does not hold: there exist agents $i \in [n]$ that do not satisfy assumptions 2(a) nor 2(b). However, we can still give a simple characterization of the distributions $\mathcal{D}_i$ for indices $i \in [n]$ that satisfy one of these conditions, as detailed in the following lemma. We recall that $\tilde{I}$ denotes the set of agents that do not satisfy 2(a) nor 2(b).

Lemma A.1. *Let $J = \{i \in [n] \setminus \tilde{I} : \forall j \neq i, \mathbb{P}(u_i > u_j) > 0\}$ and $K = \{i \in [n] \setminus (\tilde{I} \cup J) : \forall j \neq i, \mathbb{P}(u_i \geq u_j) > 0\}$. There is an ordering of $J = \{j_1, \dots, j_{|J|}\}$ and values $0 \leq m \leq m_{j_1} \leq M_{j_1} \leq m_{j_2} \leq M_{j_2} \leq \dots \leq m_{j_{|J|}} \leq M_{j_{|J|}}$ such that*

- For $i \notin \tilde{I} \cup J \cup K$, $\mathbb{P}(u_i = \max_{j \in [n]} u_j) = 0$.
- For $k \in K$, $\mathbb{P}(u_k \leq m) = 1$ and $\mathbb{P}(u_k = m) > 0$. Further for each $k \in K$ there exists $j(k) \neq k$ such that $\mathbb{P}(u_{j(k)} \geq m) = 1$ and $\mathbb{P}(u_{j(k)} > m) > 0$.
- For $j \in J$, $\mathbb{P}(u_j \in [m_j, M_j] \cup \{0\}) = 1$ and $\mathbb{P}(u_i \in (m_i, M_i)) = 0$ for all $i \neq j$. Also, for $j \in J \setminus \{j_1\}$, $\mathbb{P}(u_j = 0) > 0$.
- Either $\mathbb{P}(u_{j_1} = 0) > 0$, $K = \emptyset$, or: $m_{j_1} = m$ and $\mathbb{P}(u_{j_1} = m) > 0$.

Proof By construction of K , if $i \notin \tilde{I} \cup J \cup K$ there exists $j \in [n]$ such that $\mathbb{P}(u_j > u_i) = 1$. This proves the first claim. We now suppose that $K \neq \emptyset$. For $k \in K$ since $k \notin \tilde{I} \cup J$, there exists $j(k) \in [n]$ with $j(k) \neq k$ such that $\mathbb{P}(u_k > u_{j(k)}) = 0$. Using the definition of K we obtain $\mathbb{P}(u_{j(k)} = u_k) > 0$. Since u_k and $u_{j(k)}$ are independent, the two previous equations imply that there exists a value m_k such that $\mathbb{P}(u_k \leq m_k) = 1 = \mathbb{P}(u_{j(k)} \geq m_k)$ and $\mathbb{P}(u_k = m_k), \mathbb{P}(u_{j(k)} = m_k) > 0$. Further, if K contains two distinct elements $k \neq k'$, then by definition of K we have $0 < \mathbb{P}(u_{k'} \geq u_{j(k)}) \leq \mathbb{P}(u_{k'} \geq m_k)$. As a result, we have $m_{k'} \geq m_k$.

By symmetry, we obtain $m_k = m_{k'}$. In other words, there is a single value m such that for all $k \in K$:

$$\mathbb{P}(u_k \leq m) = 1 \quad \text{and} \quad \mathbb{P}(u_k = m) > 0,$$

and if $K \neq \emptyset$ there exists $l_K \in [n]$ such that $\mathbb{P}(u_{l_K} \geq m) = 1$ and $\mathbb{P}(u_{l_K} = m) > 0$.

Next, we turn to elements of J . For $j \in J$, let $m_j := \inf \text{Supp}(\mathcal{D}_j) \setminus \{0\}$ and $M_j := \sup \text{Supp}(\mathcal{D}_j)$. By definition of J , we know that for all $i \neq j$, $\mathbb{P}(u_i \in (m_j, M_j)) = 0$. Using the same arguments as in Theorem 2.4 we show that the intervals (m_j, M_j) for $i \in J$ are all disjoint, which ensures that there is an ordering $J = \{j_1, \dots, j_{|J|}\}$ satisfying the desired inequalities. The only remaining piece is to check that $m_{j_1} \geq m$, provided that $K \neq \emptyset$. Indeed, since $j_1 \in J$ and $u_k = m$ with non-zero probability for $k \in K$, we have $m \notin (m_{j_1}, M_{j_1})$. By definition of J , we also have $0 < \mathbb{P}(u_{j_1} > u_{l_K}) \leq \mathbb{P}(u_{j_1} > m)$. Hence, $M_{j_1} > m$ which implies $M_{j_1} \geq m_{j_1} \geq m$.

By definition of J , for $l \in [|J| - 1]$ one has $\mathbb{P}(u_{j_l} > u_{j_{l+1}}) > 0$. From the previous characterization of the supports of u_j for $j \in J$, this implies that $\mathbb{P}(u_{j_{l+1}} = 0) > 0$. This proves the third claim.

Suppose that $K \neq \emptyset$ and $J \neq \emptyset$. Then for $k \in K$ by definition of K we also have $0 < \mathbb{P}(u_k \geq u_{j_1}) = \mathbb{P}(u_{j_1} \leq m)$. This implies either $\mathbb{P}(u_{j_1} = 0) > 0$ or $m_{j_1} = m$ and $\mathbb{P}(u_{j_1} = m) > 0$. ■

B Proofs from Section 3 for the main resource allocation mechanisms

We now prove Theorem 3.1 which gives conditions for a ball to be achievable within $\mathcal{U}_\gamma$ by the mechanisms PUM.

Proof of Theorem 3.1 Fix parameters $\gamma, \mathbf{x}, r, \delta$ satisfying the assumptions. To show that $B(\mathbf{x}, r) \subseteq \mathcal{U}_\gamma$, we construct an online allocation strategy that reaches these utilities. In fact, we give an allocation strategy to reach utilities

$$\mathcal{R} := \{\mathbf{y} : \mathbf{0} \leq \mathbf{y} \leq \mathbf{z}, \mathbf{z} \in \text{Conv}(\mathcal{U}^{NI}, B(\mathbf{x}, r))\}.$$

Definition of the allocation function. We briefly recall the definition of the allocation function as defined in the main document. For convenience for any $\mathbf{U} \in \mathcal{U}^*$, we fix $p(\cdot; \mathbf{U})$ an allocation function which reaches the utility vector $\mathbf{U}$ when having access to the full information of agent's utilities.

We start by specifying the allocation strategy for extreme points $\mathbf{U}$ of $\mathcal{R}$ that still belong to the ball $B(\mathbf{x}, r)$. Following the promised utility definition above, we only need to specify an allocation function $p(\cdot | \mathbf{U})$ and a direction $\boldsymbol{\alpha}(\mathbf{U})$ for these utility vectors $\mathbf{U}$. In particular, these can be written as $\mathbf{U} = \mathbf{x} + r\mathbf{y}$ where $\mathbf{y} \in S_{n-1}$ is a unit vector. We define the allocation function via

$$p(\cdot | \mathbf{U}) := p(\cdot; \mathbf{x} + (r + \delta)\mathbf{y}) \quad \text{and} \quad \boldsymbol{\alpha}(\mathbf{U}) = \mathbf{x} - \mathbf{U} = -r\mathbf{y}.$$

We now extend the definition of the allocation function to the whole domain $\mathcal{R}$. To do so, note that any vector $\mathbf{U} \in \mathcal{R}$ can be characterized by $\mathbf{0} \leq \mathbf{U} \leq \mathbf{y}$ where $\mathbf{y} \in \text{Conv}(\mathcal{U}^{NI}, B(\mathbf{x}, r))$ and is also an extreme point of $\mathcal{R}$. In particular, $\mathbf{y}$ can be written as the convex combination of $\mathbf{0}$, vectors $\mathbb{E}[u_i]e_i$ and some extreme point $\tilde{\mathbf{y}}$ of $\mathcal{R}$ within the ball $B(\mathbf{x}, r)$ (we only need one point from $B(\mathbf{x}, r)$ at most because the ball is strictly convex). Such a decomposition can be

obtained via standard Carathéodory constructions. As a summary, we obtain

$$\mathbf{y} = \sum_{i \in [n]} q_i \mathbb{E}[u_i] \mathbf{e}_i + q_0 \tilde{\mathbf{y}}, \quad \text{where } \mathbf{q} \geq \mathbf{0}, \sum_{i \leq n} q_i \leq 1,$$

and $U_i = s_i y_i$ where $s_i \in [0, 1]$ for $i \in [n]$. We define the allocation and promised utility function by

$$\mathbf{p}(\cdot | \mathbf{y}) = \sum_{i \in [n]} q_i \mathbf{e}_i + q_0 \mathbf{p}(\cdot | \tilde{\mathbf{y}}) \quad \text{and} \quad p_i(\cdot | \mathbf{U}) = s_i p_i(\cdot | \mathbf{y}), \quad i \in [n]. \quad (25)$$

$$\mathbf{W}(\cdot | \mathbf{y}) = \sum_{i \in [n]} q_i \mathbb{E}[u_i] \mathbf{e}_i + q_0 \mathbf{W}(\cdot | \tilde{\mathbf{y}}) \quad \text{and} \quad W_i(\cdot | \mathbf{U}) = s_i W_i(\cdot | \mathbf{y}), \quad i \in [n]. \quad (26)$$

Checking that this is a valid allocation strategy for region $\mathcal{R}$. Using the promised utility framework, to show that this is a valid strategy for utilities within $\mathcal{R}$, we only need to show that Eqs. (3) and (7) are satisfied. That is, we check that the future promised utilities remain within $\mathcal{R}$ and their expectation with respect to other agents' utility matches the interim future promise defined as in Eq. (6). We start by proving that this is the case for the extreme points $\mathbf{U}$ of $\mathcal{R}$ in the ball $B(\mathbf{x}, r)$. For these utility vectors, because we used the choice of allocation strategy from Eq. (11) we only need to check Eq. (3). We prove the stronger statement that these stay within the ball $B(\mathbf{x}, r)$.

As above, we write $\mathbf{U} = \mathbf{x} + r\mathbf{y}$ with $\mathbf{y} \in S_{n-1}$. Note that $\boldsymbol{\alpha}(\mathbf{U})$ is the normal vector to the boundary of $\mathcal{R}$ at $\mathbf{U}$. Further, $\boldsymbol{\alpha}(\mathbf{U}) \in \mathbb{R}_+^n$. Indeed, fix $i \in [n]$ and construct $\mathbf{U}^{(i)}$ such that $U_i^{(i)} = 0$ and $U_j^{(i)} = U_j$ for all $j \neq i$. By construction of $\mathcal{R}$ using the rectangles we still have $\mathbf{U}^{(i)} \in \mathcal{R}$ so by convexity of $\mathcal{R}$, $0 \leq \boldsymbol{\alpha}(\mathbf{U})^\top (\mathbf{U} - \mathbf{U}^{(i)}) = \alpha_i(\mathbf{U}) U_i$. Now because $B(\mathbf{x}, r + \delta) \subseteq \mathcal{U}^* \subseteq \mathbb{R}_+^n$, we must have $U_i > 0$. This proves that $\alpha_i(\mathbf{U}) \geq 0$.

Intuitively, the choice of $\boldsymbol{\alpha}(\mathbf{U}) = \mathbf{x} - \mathbf{U}$ pushes the promised utilities towards the center of the ball $B(\mathbf{x}, r)$. Further, we can also characterize the expectation of the promised utility vector which we denote $\bar{\mathbf{W}} := \mathbb{E}[\mathbf{W}(v_i | \mathbf{U})]$. By the promise keeping equality Eq. (2),

$$\gamma \bar{\mathbf{W}} = \gamma \mathbb{E}[\mathbf{W}(\mathbf{v} | \mathbf{U})] = \mathbf{U} - (1 - \gamma)(\mathbb{E}[u_i p_i(\mathbf{v} | \mathbf{U})])_{i \in [n]}.$$

Now because the incentive-compatibility constraint is enforced and by definition of the allocation function, we have $\mathbb{E}[u_i p_i(\mathbf{v} | \mathbf{U})] = \mathbb{E}[u_i p_i(\mathbf{u}; \mathbf{x} + (r + \delta)\mathbf{y})] = x_i + (r + \delta)y_i$ for all $i \in [n]$. Combining this with the previous equation gives,

$$\bar{\mathbf{W}} = \mathbf{U} - \frac{1 - \gamma}{\gamma} \delta \mathbf{y} = \mathbf{x} + \left(r - \frac{1 - \gamma}{\gamma} \delta\right) \mathbf{y}. \quad (27)$$

Hence, the term in δ from Eq. (27) also pushes future promise vectors towards the ball center.

Using the definition of the promised utilities given the interim promise utility from Eq. (6), for any $i \in [n]$ since $W_i(v_i | \mathbf{U})$ is non-increasing in v_i ,

$$|W_i(v_i | \mathbf{U}) - \bar{W}_i| \leq W_i(0 | \mathbf{U}) - W_i(\bar{v} | \mathbf{U}) = \frac{1 - \gamma}{\gamma} \int_0^{\bar{v}} (P_i(\bar{v} | \mathbf{U}) - P_i(v | \mathbf{U})) dv \leq \frac{1 - \gamma}{\gamma} \bar{v}.$$

Now let i_1, i_2 be the indices of largest and second-largest value of $|\alpha_i(\mathbf{U})|$ respectively. By the coupling formula from Eq. (11) and because $\boldsymbol{\alpha}(\mathbf{U}) \geq \mathbf{0}$, we have

$$\|\mathbf{W}(\mathbf{v} | \mathbf{U}) - \bar{\mathbf{W}}\|_1 \leq \frac{1 - \gamma}{\gamma} \bar{v} \left(n + \frac{\sum_{j \neq i_1} \alpha_j(\mathbf{U})}{\alpha_{i_1}(\mathbf{U})} + \frac{\alpha_{i_1}(\mathbf{U})}{\alpha_{i_2}(\mathbf{U})} \right) \leq \frac{1 - \gamma}{\gamma} \bar{v} \left(2n + \frac{\alpha_{i_1}(\mathbf{U})}{\alpha_{i_2}(\mathbf{U})} \right). \quad (28)$$

We now give an upper bound on $\alpha_{i_1}(\mathbf{U})/\alpha_{i_2}(\mathbf{U})$. We recall that the region $\mathcal{R}$ contains all extreme points $\mathbb{E}[u_i]\mathbf{e}_i$ for $i \in [n]$. Hence, since $\boldsymbol{\alpha}(\mathbf{U})$ is a supporting hyperplane of $\mathcal{R}$ at $\mathbf{U}$, we have,

$$0 \leq \boldsymbol{\alpha}(\mathbf{U})^\top (\mathbf{U} - \mathbb{E}[u_{i_1}]\mathbf{e}_{i_1}) \leq \alpha_{i_1}(\mathbf{U})(U_{i_1} - \mathbb{E}[u_{i_1}]) + \alpha_{i_2}(\mathbf{U}) \sum_{j \neq i_1} U_j$$

In the second inequality, we used the fact that for all $j \neq i_1$, we have $0 \leq \alpha_j(\mathbf{U}) \leq \alpha_{i_2}(\mathbf{U})$. Because we can write $\mathbf{U} = \mathbf{x} + r\mathbf{y}$ for $\mathbf{y} \in S_{n-1}$, we have $U_{i_1} \leq x_{i_1} + r$. Similarly, for any $i \in [n]$, we have $U_i \leq x_i + r \leq 2x_i$, where in the last inequality, we used the fact that $\mathbf{x} - r\mathbf{e}_i \in \mathcal{R} \subseteq \mathbb{R}_+^n$. Then,

$$\alpha_{i_1}(\mathbf{U})(\mathbb{E}[u_{i_1}] - x_{i_1} - r) \leq 2\alpha_{i_2}(\mathbf{U}) \sum_{j \neq i_1} x_j \leq 2n\alpha_{i_2}(\mathbf{U}) \max_{i \in [n]} x_i,$$

which implies

$$\frac{\alpha_{i_1}(\mathbf{U})}{\alpha_{i_2}(\mathbf{U})} \leq \frac{2n \max_{i \in [n]} x_i}{\min_{i \in [n]} (\mathbb{E}[u_i] - x_i - r)}. \quad (29)$$

Putting this together with Eq. (28) gives

$$\|\mathbf{W}(\mathbf{v} \mid \mathbf{U}) - \mathbf{W}\|_2 \leq \|\mathbf{W}(\mathbf{v} \mid \mathbf{U}) - \mathbf{W}\|_1 \leq \frac{1-\gamma}{\gamma} \cdot 2n\bar{v} \left(1 + \frac{\max_{i \in [n]} \mathbb{E}[u_i]}{\min_{i \in [n]} (\mathbb{E}[u_i] - x_i - r)} \right).$$

For convenience, we will write $V := 2n\bar{v} \left(1 + \frac{\max_{i \in [n]} \mathbb{E}[u_i]}{\min_{i \in [n]} (\mathbb{E}[u_i] - x_i - r)} \right)$. We recall that by construction of the promised utilities, Theorem 2.2 shows that $\mathbf{W}(\mathbf{v} \mid \mathbf{U})$ lies in the hyperplane of equation $\boldsymbol{\alpha}(\mathbf{U})^\top (\mathbf{x} - \mathbf{W}) = 0$, that is $\mathbf{y}^\top (\mathbf{x} - \mathbf{W}) = 0$. Given Eq. (27), to show that for all $\mathbf{v}$, we have $\mathbf{W}(\mathbf{v} \mid \mathbf{U}) \in B(\mathbf{x}, r)$, it suffices to show that

$$\left(r - \frac{1-\gamma}{\gamma} \delta, \sup_{\mathbf{v}} \|\mathbf{W}(\mathbf{v} \mid \mathbf{U}) - \mathbf{W}\|_2 \right) \in B(0, r).$$

To prove this, using $\delta \leq r\gamma/(1-\gamma)$ we compute

$$\begin{aligned} r^2 - \left(r - \frac{1-\gamma}{\gamma} \delta \right)^2 - \left(\frac{1-\gamma}{\gamma} V \right)^2 &= \frac{1-\gamma}{\gamma^2} [2\gamma\delta r - (1-\gamma)\delta^2 - (1-\gamma)V^2] \\ &\geq \frac{1-\gamma}{\gamma^2} [\gamma\delta r - (1-\gamma)V^2] \geq 0. \end{aligned}$$

This ends the proof that whenever $\mathbf{U}$ is an extreme point of $\mathcal{R}$ within the ball $B(\mathbf{x}, r)$, the future promised utilities $\mathbf{W}(\mathbf{v} \mid \mathbf{U})$ stay within the ball $B(\mathbf{x}, r) \subseteq \mathcal{R}$.

It remains to check that Eqs. (3) and (7) hold for all other utility vectors $\mathbf{U} \in \mathcal{R}$. As in the definition of the allocation function, we write $U_i = s_i y_i$ where $s_i \in [0, 1]$ for $i \in [n]$, and $\mathbf{y} = \sum_{i \in [n]} q_i \mathbb{E}[u_i] \mathbf{e}_i + q_0 \tilde{\mathbf{y}}$ where $\mathbf{q} \geq \mathbf{0}$, $\sum_{i \leq n} q_i \leq 1$, and $\tilde{\mathbf{y}}$ is an extreme point of $\mathcal{R}$ within the ball $B(\mathbf{x}, r)$. In the definition of the allocation function from Eq. (25), the only component that is non-constant for $\mathbf{v} \in [0, \bar{v}]^n$ is the contribution of $\mathbf{p}(\cdot \mid \tilde{\mathbf{y}})$. By linearity, Eq. (7) is then a consequence of the definitions Eqs. (25) and (26). We turn to Eq. (3). In the previous part, we showed that $\mathbf{W}(\cdot \mid \tilde{\mathbf{y}}) \in B(\mathbf{x}, r)$ for all possible reports. As a result, by linearity this directly shows that the future promised utilities stay in the the corresponding ellipsoid $\mathcal{E}(\mathbf{U})$ defined below:

$$\forall \mathbf{v} \in [0, \bar{v}]^n, \quad \mathbf{W}(\mathbf{v} \mid \mathbf{U}) \in \left\{ (s_i(q_i \mathbb{E}[u_i] + q_0 x_i))_{i \in [n]} + \mathbf{z} : q_0^2 \sum_{i \in [n]} (s_i z_i)^2 \leq r^2 \right\} := \mathcal{E}(\mathbf{U}).$$

We can easily show that $\mathcal{E}(\mathbf{U}) \subseteq \mathcal{R}$ by construction of the region $\mathcal{R}$. This ends the proof. $\blacksquare$

In Theorem 3.3, we simplified the form of the constant C from Theorem 3.1 under Theorem 3.2 which asks that all valuations are lower bounded by some fixed value $\underline{v} > 0$. To give some intuitions, the term C in Theorem 3.1 essentially measures how close the tentative ball $B(\mathbf{x}, r)$ is to the limiting hyperplanes $\{\mathbf{x} \in \mathbb{R}^n : x_i = \mathbb{E}[u_i]\}$ for $i \in [n]$. Under Theorem 3.2, the following result gives geometrical properties of the optimal set $\mathcal{U}^*$, which intuitively imply that points in the optimal set $\mathcal{U}^*$ are sufficiently far from the hyperplanes $\{\mathbf{x} \in \mathbb{R}^n : x_i = \mathbb{E}[u_i]\}$ for $i \in [n]$.

Lemma B.1. *Under Theorem 3.2, we have*

$$\mathcal{U}^* \subseteq \mathbb{R}_+^n \cap \bigcap_{i \in [n]} \left\{ \mathbf{x} : \underline{v} \sum_{j \neq i} x_j \leq \bar{v}(\mathbb{E}[u_i] - x_i) \right\}.$$

Proof We already characterized the extreme points of $\mathcal{U}^*$ as $\mathbf{0}$ together with the utilities reached by optimal allocation functions $\mathbf{p}(\cdot)$ for α -weighted objectives for $\alpha \in \mathbb{R}_+^n \setminus \{\mathbf{0}\}$.

Now for $i \in [n]$, consider the weights α such that $\alpha_i = \bar{v}$, and $\alpha_j = \underline{v}$ for all $j \neq i$. From Theorem 3.2 the valuations always lie in $[\underline{v}, \bar{v}]$, hence, the allocation function that always allocates the resource to agent i is α -optimal. This allocation reaches the utility vector $\mathbb{E}[u_i]e_i$. Because this allocation is α -optimal, this proves the constraint $\mathcal{U}^* \subseteq \{\mathbf{x} : \alpha^\top \mathbf{x} \leq \bar{v}\mathbb{E}[u_i]\}$. $\blacksquare$

We are now ready to give the proof of Theorem 3.3.

Proof of Theorem 3.3 Going back to the proof of Theorem 3.1, we aim to improve the upper bound of $\alpha_{i_1}(\mathbf{U})/\alpha_{i_2}(\mathbf{U})$ from Eq. (29) using Theorem B.1. Up to rescaling, we suppose without loss of generality that $\alpha_{i_1}(\mathbf{U}) = \bar{v}$. Now suppose by contradiction that $\alpha_{i_2}(\mathbf{U}) < \underline{v}$. In particular, for all $j \neq i_1$, we have $\alpha_j(\mathbf{U}) \leq \alpha_{i_2}(\mathbf{U}) < \underline{v}$. Further, Theorem B.1, since $\mathbf{U} \in \mathcal{R} \subseteq \mathcal{U}^*$, we have

$$\underline{v} \sum_{j \neq i_1} U_j \leq \bar{v}(\mathbb{E}[u_{i_1}] - U_{i_1})$$

Combining the previous remarks gives,

$$\alpha(\mathbf{U})^\top (\mathbf{U} - \mathbb{E}[u_{i_1}]e_{i_1}) \leq (\alpha_{i_2} - \underline{v})U_{i_2} + \underline{v} \sum_{j \neq i_1} U_j - \bar{v}(\mathbb{E}[u_{i_1}] - U_{i_1}) \leq (\alpha_{i_2} - \underline{v})U_{i_2}.$$

We briefly argue that $U_{i_2} > 0$. We recall that $\mathbf{U} = \mathbf{x} + r\mathbf{y}$ for some element $\mathbf{y} \in S_{n-1}$. Further, because $\alpha(\mathbf{U}) \in \mathbb{R}_+^n$, we must also have $\mathbf{y} \in \mathbb{R}_+^n$ otherwise $\mathbf{U}$ wouldn't be α -optimal within $\mathcal{R}$. Because $B(x, r + \delta) \subseteq \mathbb{R}_+^n$ and $r + \delta > 0$, we have $U_{i_2} \geq x_{i_2} > 0$. This proves that $\alpha(\mathbf{U})^\top (\mathbf{U} - \mathbb{E}[u_{i_1}]e_{i_1}) < 0$. This contradicts the fact that $\alpha(\mathbf{U})$ is the normal of a supporting hyperplane at $\mathbf{U}$ of $\mathcal{R}$, which contains in particular $\mathbb{E}[u_{i_1}]e_{i_1}$. In conclusion, $\alpha_{i_1}(\mathbf{U})/\alpha_{i_2}(\mathbf{U}) \leq \bar{v}/\underline{v}$. We can now use the same proof as for Theorem 3.1 with $V := \bar{v}(2n + \bar{v}/\underline{v})$. $\blacksquare$

To understand the optimal policies for some social welfare, it may be useful to focus locally on this direction $\alpha = \mathbf{1}$. In particular, it may be possible to ignore some irrelevant agents in $[n]$. This leads to a stronger local version than the statement given in Theorem 3.1.

Lemma B.2. Fix $\gamma \in [0, 1)$. Let $\mathbf{U}$ an optimal vector and $\mathbf{x} \in \mathcal{U}^*$ such that for all $i \notin I$, we have $x_i = U_i$. Let $r, \delta > 0$ and suppose that there is a partition $I = I_1 \sqcup \dots \sqcup I_q$ with $|I_s| \geq 2$ for all $s \in [q]$ such that

$$\forall s \neq s' \in [q], \quad \mathbb{P}(Z_{I_s} = Z_{I_{s'}} = Z > 0) = 0, \quad (30)$$

and for all $s \in [q]$,

$$B_{I_s}((x_i)_{i \in I_s}, r + \delta) \subseteq P_{I_s}(\mathcal{U}^* \cap \{\mathbf{y} : \forall i \notin I_s, y_i = U_i\}). \quad (31)$$

Then, if Eq. (13) holds with $C := 4n^2\bar{v}^2 \left(1 + \frac{\max_{i \in [n]} \mathbb{E}[u_i]}{\min_{s \in [q], i \in I_s} (\mathbb{E}[u_i] \mathbb{P}(Z_{I_s} = Z > 0) - x_i - r)}\right)^2$, we have

$$B(\mathbf{x}, r) \cap \{\mathbf{y} : \forall i \notin I, y_i = x_i\} \subseteq \mathbf{x} + \{0\}^{[n] \setminus I} \otimes \bigotimes_{s \in [q]} B_{I_s}(\mathbf{0}, r) \subseteq \mathcal{U}_\gamma.$$

As a remark, note that while Theorem 2.4 essentially shows that only agents in $\tilde{I}$ need to be carefully treated (all others can be treated perfectly even with $\gamma = 0$), Theorem B.2 takes all I agents into consideration. Precisely, we added agents $i \in [n]$ such that $\mathbb{P}(u_i = Z_{-i} > 0) > 0$. Adding these agents can only help to satisfy the constraints from Theorem B.2. Indeed, by construction there is some other agent $j \neq i$ with $j \in I$ such that $\mathbb{P}(u_i = u_j = Z) > 0$ hence we can group i and j in the same set of the partition $I = I_1 \sqcup \dots \sqcup I_q$. As a result, the set of optimal vectors of $\mathcal{U}^*$ is perfectly flat along the direction $(\mathbb{1}_{k=j} - \mathbb{1}_{k=i})_{k \in [n]}$, hence we can always fit a ball to satisfy Eq (31). In summary, taking agents $I \setminus \tilde{I}$ into account within Theorem B.2 was without loss. On the other hand, these $I \setminus \tilde{I}$ agents bring extra flexibility that can be beneficial for agents $\tilde{I}$ to satisfy Eq (31). This flexibility will be useful in Section D.5 for discrete utility distributions.

Proof of Theorem B.2 We define J and K as in Theorem A.1. Throughout, we use the characterizations in this lemma. First, note that if $k \in K$ and $m > 0$, then $\mathbb{P}(u_k = Z_{-k} > 0) \geq \mathbb{P}(u_k = u_{j(k)} = m) \mathbb{P}(\forall j \notin \{k, j(k)\}, u_j \leq m) > 0$, where we bounded the second term using the definition of $k \in K$. Hence, this implies $k \in I$. Therefore, for any $i \notin I$ we have either (1) $i \notin \tilde{I} \cup J \cup K$ or (2) $i \in K$ and $m = 0$, or (3) $i \in J$. In all cases, either $x_i = U_i = 0$ or $i \in J$.

Then, for any $\mathbf{V} \in \{\mathbf{y} : \forall i \notin I, y_i = x_i = U_i\}$, and any allocation $\mathbf{p}(\cdot)$ that realizes $\mathbf{V}$ in the full-information setting, we have

$$\forall i \notin I \cup J, \quad p_i(\mathbf{u}) = 0 \text{ (a.s.)} \quad (32)$$

$$\forall i \in J \setminus I, \quad p_i(\mathbf{u}) = \mathbb{1}_{\forall j \neq i, u_j \leq m_i} \mathbb{1}_{u_i > 0} \text{ (a.s.)}. \quad (33)$$

This is because by assumption, $\mathbf{U}$ is an optimal vector. For convenience, we pose for $i \notin I$,

$$p_i^{(0)}(\mathbf{u}) = \begin{cases} 0 & i \notin I \cup J \\ \mathbb{1}_{\forall j \neq i, u_j \leq m_i} \mathbb{1}_{u_i > 0} & i \in J \setminus I. \end{cases}$$

It will be useful later to view $[n] \setminus I$ as part of the partition as well hence we pose $I_0 := [n] \setminus I$. We can easily check that Eq. (30) is also satisfied for all $r \neq r' \in \{0, \dots, q\}$ since

$$\mathbb{P}(Z_{I_0} = Z_I > 0) \leq \sum_{i \in I_0} \mathbb{P}(u_i = Z > 0) = 0.$$

We first focus on coordinates within I_s for some fixed $s \in [q]$. Let $\mathbf{V} \in \mathcal{U}^* \cap \{\mathbf{y} : \forall i \notin I_s, y_i = U_i\}$ and $\mathbf{p}(\cdot; \mathbf{V})$ an allocation that realizes $\mathbf{V}$ in the full information setting. Because $\mathbf{U}$ is an optimal vector, we have

$$\mathbb{E} [Z_{-I_s} \mathbb{1}_{Z_{-I_s} > Z_{I_s}}] \leq \sum_{i \notin I_s} U_i \leq \mathbb{E} [Z_{-I_s} \mathbb{1}_{Z_{-I_s} \geq Z_{I_s}}]$$

Because of Eq. (30), we have $\mathbb{P}(Z_{-I_s} = Z_{I_s} > 0) \leq \sum_{s' \neq s} \mathbb{P}(Z_{I_{s'}} = Z_{I_s} = Z > 0) = 0$. Then,

$$\sum_{i \notin I_s} V_i = \sum_{i \notin I_s} U_i = \mathbb{E} [Z_{-I_s} \mathbb{1}_{Z_{-I_s} \geq Z_{I_s}}].$$

As a result, on the event $\mathcal{E}_s := \{Z_{-I_s} \geq Z_{I_s}\}$, without loss of generality, we can suppose that $\mathbf{p}$ allocates to agents in $[n] \setminus I_s$:

$$\forall i \in I_s, \quad \mathbb{P}(p_i(\mathbf{u}) > 0, \mathcal{E}_s) = 0. \quad (34)$$

In the above, we used the fact that if $Z_{-I_s} = Z_{I_s} = 0$, we have $u_i = 0$ for all $i \in I_s$. Thus allocating to agents in I_s in that case is useless. We suppose that Eq. (34) will always be satisfied by the allocations of the form $\mathbf{p}(\cdot; \mathbf{V})$ from now. Note that $1 - \mathbb{P}(\mathcal{E}_s) = \mathbb{P}(Z_{I_s} > Z_{-I_s}) = \mathbb{P}(Z_{I_s} \geq Z_{-I_s}, Z_{I_s} > 0) = \mathbb{P}(Z_{I_s} = Z > 0) > 0$. In the last inequality, we used $\emptyset \subsetneq I_s \subseteq I$.

Eq. (34) is also a characterization in the following sense. Consider the resource allocation problem with only agents I_s but such that there is no allocation on the event $\mathcal{E}_s$. Denote by $\mathcal{U}_s^* \in \mathbb{R}_+^{I_s}$ the utility region that can be achieved in this setting. Then,

$$\mathcal{U}_s^* = P_{I_s}(\mathcal{U}^* \cap \{\mathbf{y} : \forall i \notin I_s, y_i = U_i\}). \quad (35)$$

Indeed, Eq. (34) shows the inclusion $\supseteq$ and for the inclusion $\subseteq$, any allocation from the game with agents I_s can be completed on $\mathcal{E}_s$ by allocating the resource to any agent in $\arg \max_{i \notin I_s} u_i$. In particular, by assumption, we have $B_{I_s}((x_i)_{i \in I_s}, r + \delta) \subseteq \mathcal{U}_s^*$. In this setting with I_s agents, we can also define the region $\mathcal{U}_s^{NI}$ that can be achieved without reports, that is $\mathcal{U}_s^{NI} = \text{Conv}(\mathbb{E}[u_i](1 - \mathbb{P}(\mathcal{E}_s))\mathbf{e}_i, i \in I_s)$. The only difference with a standard resource allocation problem between agents I_s is the constraint on $\mathcal{E}_s$. However, the proof of Theorem 3.1 still holds in this setting as specified below. We construct the utility region

$$\mathcal{R}_s := \{\mathbf{y} \in \mathbb{R}^{I_s} : \mathbf{0} \leq \mathbf{y} \leq \mathbf{z}, \mathbf{z} \in \text{Conv}(\mathcal{U}_s^{NI}, B_{I_s}((x_i)_{i \in I_s}, r))\}.$$

We then specify an allocation strategy for extreme points $\mathbf{U}$ of $\mathcal{R}_s$ that still belong to the ball $B(\mathbf{x}, r)$. We write $\mathbf{U} = \mathbf{x} + r\mathbf{y}$ where $\|\mathbf{y}\| = 1$ and pose

$$\mathbf{p}^{(s)}(\cdot \mid \mathbf{U}) := \mathbf{p}^{(s)}(\cdot; \mathbf{x} + (r + \delta)\mathbf{y}) \quad \text{and} \quad \boldsymbol{\alpha}(\mathbf{U}) = \mathbf{x} - \mathbf{U},$$

where $\mathbf{p}^{(s)}(\cdot; \mathbf{z})$ is an allocation function that realizes $\mathbf{z} \in \mathcal{U}_s^*$ in the full information setting. We then extend the definition of $\mathbf{p}^{(s)}(\cdot \mid \mathbf{U})$ and $\mathbf{W}^{(s)}(\cdot \mid \mathbf{U})$ to the complete region $\mathcal{R}_s$ as in Theorem 3.1 in Eqs. (25) and (26). The proof of Theorem 3.1 shows that if $\frac{r\gamma}{1-\gamma} \geq \delta \geq C_s \frac{1-\gamma}{\gamma r}$, where $C_s = 4|I_s|^2 \bar{v}^2 \left(1 + \frac{\max_{i \in I_s} \mathbb{E}[u_i]}{\min_{i \in I_s} (\mathbb{E}[u_i](1 - \mathbb{P}(\mathcal{E}_s)) - x_i - r)}\right)^2$, then the allocation and promised utility functions constructed are valid on $\mathcal{R}_s$ (they satisfy Eqs. (3) and (7)). We can check that this is the case since with the choice of parameters, $C \geq \max_{s \in [q]} C_s$.

We now pose $\mathcal{R}_0 := (x_i)_{i \in I_0}$ and construct an allocation mechanism to reach the utilities

$$\mathcal{R} := \bigotimes_{0 \leq s \leq q} \mathcal{R}_s \supseteq \mathbf{x} + \{0\}^{I_0} \otimes \bigotimes_{s \in [q]} B_{I_s}(\mathbf{0}, r),$$

by treating each cluster I_s of agents separately. We pose for all $\mathbf{U} \in \mathcal{R}$, $i \in [n]$ and $\mathbf{u} \in [0, \bar{v}]^n$,

$$\begin{aligned} \mathbf{p}(\mathbf{u} \mid \mathbf{U}) &:= \begin{cases} p_i^{(s)}(\mathbf{u}_{I_s} \mid \mathbf{U}_{I_s}) \frac{\mathbb{1}(Z_{I_s} > Z_{-I_s})}{1 - \mathbb{P}(\mathcal{E}_s)} & i \in I_s, s \in [q] \\ p_i^{(0)}(\mathbf{u}) & i \in I_0. \end{cases} \\ \mathbf{W}(\mathbf{u} \mid \mathbf{U}) &:= \begin{cases} W_i^{(s)}(\mathbf{u}_{I_s} \mid \mathbf{U}_{I_s}) & i \in I_s, s \in [q] \\ x_i & i \in I_0. \end{cases} \end{aligned}$$

By design of the events $\mathcal{E}_s$, we still have $\mathbf{p}(\mathbf{u} \mid \mathbf{U}) \in \Delta_n$ (there are no collisions between different partitions). By construction, Eq. (3) holds and Eq. (7) is satisfied for all $i \in I$. It only remains to check that Eq. (7) holds for $i \in I_0 = [n] \setminus I$, which is equivalent to checking both Eq. (2) and the incentive-compatibility constraint Eq. (4). Eq. (2) is directly satisfied by Eqs. (32) and (33). For the incentive-compatibility, note that for $i \notin I$ the definition of $p_i^{(0)}(\mathbf{u})$ only uses the information $\mathbb{1}_{u_i > 0}$ about u_i . Hence the interim promise $P_i^{(0)}(u_i)$ is constant on $(0, \bar{v}]$. As a result, the allocation is also incentive-compatible for agent i (this is always true if $u_i = 0$).

Altogether, this shows that $\mathcal{R} \subseteq \mathcal{U}_\gamma$, which ends the proof. $\blacksquare$

C Proofs from Section 4 to prove lower bounds on the social welfare gap

We now prove Theorem 4.1 that gives tools to prove upper bounds on the achievable region $\mathcal{U}_\gamma$. To prove this result, we first need to derive some properties of valid mechanisms for the central planner.

Using the promised utility framework, we suppose that we are given a utility region $\mathcal{R}$, allocation functions $\mathbf{p}(\cdot \mid \mathbf{U})$ and promised utility functions $\mathbf{W}(\cdot \mid \mathbf{U})$ satisfying Eqs. (3) and (7). We start by showing that to achieve the utilities from a point $\mathbf{U}$ in the boundary of $\mathcal{U}^*$ in the full information setting, the allocation function needs to discriminate enough between low and high agent utilities. In practice, this will be helpful to derive a lower bound on the range of interim future promises if one aims to reach the boundary of $\mathcal{U}^*$. We recall the definition of $\tilde{I}$ as the set of indices that do not satisfy conditions 2(a) nor 2(b) from Theorem 2.4.

Lemma C.1. *Let $i \in \tilde{I}$. Define $q_i := \min\{v \in [0, \bar{v}] : \mathbb{P}(u_i \leq v) \geq 1/2\}$ the median of $\mathcal{D}_i$. Then, there exists $\eta_i > 0$ such that the following holds. Fix any allocation function $\mathbf{p} : [0, \bar{v}]^n \rightarrow \Delta_n$ that is non-decreasing (that is, for all $i \in [n]$, $P_i(u_i) = \mathbb{E}_{\mathbf{u}_{-i}}[p_i(\mathbf{u})]$ is non-decreasing) and such that $(\mathbb{E}[u_i p_i(\mathbf{u})])_{i \in [n]}$ is an optimal utility vector. Then,*

$$\mathbb{E}_{u_i \sim \mathcal{D}_i}[|P_i(u_i) - P_i(q_i)| \min(u_i, q_i)] \geq \eta_i.$$

Proof Fix an index $i \in \tilde{I}$ satisfying the assumptions. We decompose $[0, \bar{v}] \setminus \text{Supp}(\mathcal{D}_i)$ as a union of disjoint intervals $(a_0, b_0 = \bar{v}) \cup [a_1 = 0, b_1) \cup \bigcup_{k \geq 2} (a_k, b_k)$. We note that $u_i \sim \mathcal{D}_i$ is not a constant random variable otherwise posing $m_i = M_i$ would satisfy 2(b), where $u_i = m_i$ almost surely, contradicting $i \in \tilde{I}$. This implies that $b_1 < a_0$.

Now fix an optimal allocation p for the social welfare. Because it is optimal, with probability one on $\mathbf{u}$, we have $p_i(\mathbf{u}) > 0$ only if $i \in \arg \max_{j \in [n]} u_j$. We also recall that $P_i(\cdot)$ is non-decreasing by hypothesis. As a result, for all $u \in \text{Supp}(\mathcal{D}_i) \setminus \{a_k, b_k, k \geq 0\}$ we have

$$\mathbb{P}(Z_{-i} < u) \leq P_i(u) \leq \mathbb{P}(Z_{-i} \leq u). \quad (36)$$

Next, because P_i is non-decreasing, we have $P_i(\bar{v}) \geq \mathbb{P}(Z_{-i} < a_0)$ and

$$\begin{cases} \forall u \in [a_k, b_k], & \mathbb{P}(Z_{-i} < a_k) \leq P_i(u) \leq \mathbb{P}(Z_{-i} \leq b_k), & k \geq 1 \\ \forall u \in [a_0, \bar{v}], & \mathbb{P}(Z_{-i} < a_0) \leq P_i(u) \leq P_i(\bar{v}) \end{cases} \quad (37)$$

We now consider the problem of minimizing the desired quantity

$$\begin{aligned} \mathbb{E}_{u_i \sim \mathcal{D}_i} [|P_i(u_i) - P_i(q_i)| \min(u_i, q_i)] \\ = \mathbb{E}_{u_i} [q_i P_i(u_i) \mathbb{1}_{u_i > q_i}] - \mathbb{E}_{u_i} [u_i P_i(u_i) \mathbb{1}_{u_i < q_i}] + (\mathbb{E}[u_i \mathbb{1}_{u_i < q_i}] - q_i \mathbb{P}(u_i > q_i)) P_i(q_i), \end{aligned}$$

under the constraints Eqs. (36) and (37). We can directly check that the optimum is achieved by setting $P_i^*(u)$ equal to the upper bounds from Eqs. (36) and (37) if $u < q_i$, and set $P_i^*(u)$ equal to their lower bounds if $u > q_i$. For q_i , we have $P_i^*(q_i) = \mathbb{P}(Z_{-i} < q_i)$ if $\mathbb{E}[u_i \mathbb{1}_{u_i < q_i}] \leq q_i \mathbb{P}(u_i > q_i)$ and $P_i^*(q_i) = \mathbb{P}(Z_{-i} \leq q_i)$ otherwise. We denote $\eta_i := \mathbb{E}_{u_i \sim \mathcal{D}_i} [|P_i(u_i) - P_i(q_i)| \min(u_i, q_i)]$ the corresponding value of P_i^* . As a result, we obtained that for any optimal allocation function,

$$\mathbb{E}[|P_i(u_i) - P_i(q_i)| \min(u_i, q_i)] \geq \eta_i.$$

Note that to prove the desired result, it would suffice to show that $\eta_i > 0$. The main point is that this lower bound η_i is tight since the interim function P_i^* can be achieved by an optimal allocation function that either allocates to the agent with index $\arg \max_{j \neq i} u_j$ or i (only if $u_i = Z_{-i}$). Hence, to end the proof it suffices to show that for any allocation function p that is optimal for the social welfare, one has $\mathbb{E}[|P_i(u_i) - P_i(q_i)| \min(u_i, q_i)] > 0$. Suppose this is not the case by contradiction, hence $P_i(u) = P_i(q_i)$ for $\mathcal{D}_i$ -almost all $u \in (0, \bar{v}]$. We recall that P_i is non-decreasing. Therefore, letting $m_i = \inf \text{Supp}(u_i) \setminus \{0\}$ and $M_i = \sup \text{Supp}(u_i)$, if $M_i > q_i$ we obtain that P_i is constant on the interval $[q_i, M_i]$. Similarly, if $m_i < q_i$ then P_i is constant on $(m_i, q_i]$. In all cases, we showed that P_i is constant on (m_i, M_i) . Hence, Eq. (36) implies that

$$\mathbb{P}(Z_{-i} \leq m_i) \geq \lim_{x \rightarrow m_i^+} P_i(x) = \lim_{x \rightarrow M_i^-} P_i(x) \geq \mathbb{P}(Z_{-i} < M_i).$$

Because $m_i \leq M_i$, the previous equation implies $\mathbb{P}(Z_{-i} \in (m_i, M_i)) = 0$. On the other hand, by construction we have $\mathbb{P}(u_i \in \{0\} \cup [m_i, M_i]) = 1$. Since condition 2(a) from Theorem 2.4 does not hold because $i \notin \tilde{I}$, for all $j \neq i$, we have

$$\mathbb{P}(Z_{-i} < M_i) = \prod_{j \neq i} \mathbb{P}(u_j < M_i) \geq \prod_{j \neq i} \mathbb{P}(u_j < u_i) > 0.$$

Combining this with the previous equation implies $\mathbb{P}(Z_{-i} \leq m_i) > 0$. We now fix $j \neq i$. We have

$$0 = \mathbb{P}(Z_{-i} \in (m_i, M_i)) \geq \mathbb{P}(Z_{-\{i,j\}} \leq m_i) \mathbb{P}(u_j \in (m_i, M_i)) \geq \mathbb{P}(Z_{-i} \leq m_i) \mathbb{P}(u_j \in (m_i, M_i)).$$

Therefore, $\mathbb{P}(u_j \in (m_i, M_i)) = 0$. This holds for all $j \neq i$, hence condition 2(b) of Theorem 2.4 holds which contradicts the definition of $i \in \tilde{I}$. This ends the proof that $\eta_i > 0$ and the desired result. $\blacksquare$

We note that having $i \in \tilde{I}$ is necessary to obtain such a constant $\eta_i > 0$. Indeed, suppose that condition 2(a) from Theorem 2.4 holds. In that case, there are optimal allocations that never allocate the resource to agent i (since there is an agent j that always has at least the same utility), which results in having $P_i(u) = 0$ for all $u \in [0, \bar{v}]$. On the other hand, if condition 2(b) holds, there is an interval $[m_i, M_i]$ such that $\mathbb{P}(u_i \in [m_i, M_i] \cup \{0\}) = 1$ but $\mathbb{P}(u_j \in (m_i, M_i)) = 0$ for all $j \neq i$. Then, taking into consideration the reports of agent i is not really necessary: we can safely allocate the resource to agent i if and only if $Z_{-i} \leq m_i$ and $u_i > 0$ (if $u_i = 0$, then truthful reporting is always an equilibrium for agent i). This optimal allocation then yields a constant interim allocation probability $P_i(u) = \mathbb{P}(Z_{-i} \leq m_i)$ for all $u \in (0, \bar{v}]$.

We now use Theorem C.1 to derive a the lower bound on the range of the interim future promise W_i even when the realized utility is not on the boundary of $\mathcal{U}^*$, but somewhat close to it. Having the realized utility vector $\mathbf{U}$ close to the boundary is necessary, otherwise, for instance if $\mathbf{U} \in \mathcal{U}^{NI}$, we can use a constant allocation function $\mathbf{p}(\cdot)$ so that $W_i(v)$ is constant for $v \in [0, \bar{v}]$ for all $i \in [n]$.

Lemma C.2. *Fix $\gamma \in (0, 1]$ and an index $i \in \tilde{I}$. Define $q_i \in [0, \bar{v}]$ as in Theorem C.1. Then, there exists $\delta_i, \eta_i > 0$ (η_i is the same as in Theorem C.1) such that the following holds. For any allocation function $\mathbf{p} : [0, \bar{v}]^n \rightarrow \Delta_n$ that is non-decreasing (that is, for all $i \in [n]$ $P_i : u_i \mapsto \mathbb{E}_{\mathbf{u}_{-i}}[p_i(\mathbf{u})]$ is non-decreasing) and such that*

$$\max_{\mathbf{x} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{x} - \mathbf{1}^\top (\mathbb{E}[u_i p_i(\mathbf{u})])_{i \in [n]} \leq \delta_i,$$

the interim promised utility function $W_i(\cdot)$ for $i \in [n]$ constructed via Eq. (6) using the interim allocation functions P_i is non-increasing and satisfies

$$\mathbb{E}_{u_i \sim \mathcal{D}_i} |W_i(u_i) - W_i(q_i)| \geq \frac{1 - \gamma}{2\gamma} \eta_i.$$

Proof The fact that the interim promise function $W_i(\cdot \mid \mathbf{U})$ is non-increasing is direct from the definition in Eq. (6). We omit the terms $\mid \mathbf{U}$ in the rest of the proof for readability. From the same definition, for any $0 \leq a \leq b \leq \bar{v}$,

$$\begin{aligned} W_i(a) - W_i(b) &= \frac{1 - \gamma}{\gamma} \left(\int_0^b (P_i(b) - P_i(v)) dv - \int_0^a (P_i(a) - P_i(v)) dv \right) \\ &= \frac{1 - \gamma}{\gamma} \left((P_i(b) - P_i(a))a + \int_a^b (P_i(b) - P_i(v)) dv \right) \geq \frac{1 - \gamma}{\gamma} (P_i(b) - P_i(a))a. \end{aligned}$$

As a result,

$$\mathbb{E}|W_i(u_i) - W_i(q_i)| = \mathbb{E}[(W_i(q_i) - W_i(u_i))\mathbb{1}_{u_i \geq q_i}] + \mathbb{E}[(W_i(u_i) - W_i(q_i))\mathbb{1}_{u_i \leq q_i}] \quad (38)$$

$$\geq \frac{1 - \gamma}{\gamma} \mathbb{E}[|P_i(u_i) - P_i(q_i)| \min(u_i, q_i)]. \quad (39)$$

Hence, we aim to lower bound the expectation on the right-hand side. First, recall that $\max_{\mathbf{x} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{x} = \mathbb{E}[Z] = \mathbb{E} \max(u_i, Z_{-i})$, and that the maximum is attained for any first-best allocation function. Next, for $u \in [0, \bar{v}]$, we have

$$\mathbf{1}^\top (\mathbb{E}[u_i p_i(\mathbf{u})])_{i \in [n]} \leq \mathbb{E}[u_i p_i(\mathbf{u}) + Z_{-i}(1 - p_i(\mathbf{u}))] = \mathbb{E}[u_i P_i(u_i) + Z_{-i}(1 - p_i(\mathbf{u}))].$$

Now for any value $u \in [0, \bar{v}]$, let $z_i(u) = \inf\{v : \mathbb{P}(Z_{-i} \geq v) \leq 1 - P_i(u)\}$. In particular, since $P_i(\cdot)$ is non-decreasing, so is $z_i(\cdot)$. We can check that

$$\mathbb{E}_{\mathbf{u}_{-i}}[Z_{-i}(1 - p_i(\mathbf{u}))] \leq \mathbb{E}_{\mathbf{u}_{-i}}[Z_{-i}\mathbb{1}_{Z_{-i} > z_i(u_i)}] + z_i(u_i) [\mathbb{P}(Z_{-i} \geq z_i(u_i)) - (1 - P_i(u_i))].$$

Because of the second term in the right-hand side, we define a variable $B_i \sim \mathcal{B}(\mathbb{P}(Z_{-i} \geq z_i(u_i)) - (1 - P_i(u_i)))$ which is independent from $\mathbf{u}_{-i}$. Putting the previous equations together yields

$$\begin{aligned} \max_{\mathbf{x} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{x} - \mathbf{1}^\top (\mathbb{E}[u_i p_i(\mathbf{u})])_{i \in [n]} &\geq \\ \mathbb{E}[\max(u_i, Z_{-i}) - Z_{-i}\mathbb{1}_{Z_{-i} > z_i(u_i)} - u_i\mathbb{1}_{Z_{-i} < z_i(u_i)} - \mathbb{1}_{Z_{-i} = z_i(u_i)}(Z_{-i}B_i + u_i(1 - B_i))] &=: A(P_i). \end{aligned} \quad (40)$$

This inequality is tight in the following sense. For any non-decreasing function P_i , there is a non-decreasing allocation function $\mathbf{p}(\cdot)$ with interim allocation function P_i for agent i , and

$$\max_{\mathbf{x} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{x} - \mathbf{1}^\top (\mathbb{E}[u_i p_i(\mathbf{u})])_{i \in [n]} = A(P_i).$$

This allocation can directly be obtained from the definition of $A(P_i)$: when $Z_{-i} < z_i(u_i)$ we allocate to agent i , when $Z_{-i} > z_i(u_i)$ we allocate to some agent in $\arg \max_{j \neq i} u_j$, and if $Z_{-i} = z_i(u_i)$ we allocate to one of these two cases depending on the value of B_i . In summary, for any $\delta \geq 0$,

$$\begin{aligned} \left\{ P_i(\cdot) : \exists \text{ non-decreasing } \mathbf{p}(\cdot), \max_{\mathbf{x} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{x} - \mathbf{1}^\top (\mathbb{E}[u_i p_i(\mathbf{u})])_{i \in [n]} \leq \delta, P_i(\cdot) = \mathbb{E}_{\mathbf{u}_{-i}}[p_i(\cdot, \mathbf{u}_{-i})] \right\} \\ = \{ \text{non-decreasing } P_i : [0, \bar{v}] \rightarrow [0, 1], A(P_i) \leq \delta \}. \end{aligned} \quad (41)$$

As a result, we now only focus on non-decreasing interim allocation functions P_i satisfying $A(P_i) \leq \delta$. The set of non-decreasing functions on $[0, \bar{v}] \rightarrow [0, 1]$ is closed for the product measure (point-wise convergence) and so is the set of functions satisfying the constraint $A(P_i) \leq \eta$ by the dominated convergence theorem. Further, Tychonov's theorem implies that the set of functions $[0, \bar{v}] \rightarrow [0, 1]$ is compact for the product measure. This therefore implies that the following optimization problem

$$\inf_{\substack{\text{non-decreasing } P_i : [0, \bar{v}] \rightarrow [0, 1] \\ A(P_i) \leq \delta}} \mathbb{E}_{u_i \sim \mathcal{D}_i} [|P_i(u_i) - P_i(q_i)| \min(u_i, q_i)] =: B(\delta)$$

achieves its minimum, and further that $B(\delta) \rightarrow B(0)$ as $\delta \rightarrow 0$. Note that $\delta = 0$ corresponds to the case when the allocation is exactly optimal for the social welfare. This is the case that was covered by Theorem C.1, which shows that $B(0) \geq \eta_i$ for some constant $\eta_i > 0$ (that only depends on i , and the distributions $\mathcal{D}_1, \dots, \mathcal{D}_n$). In particular, there exists $\delta_i > 0$ such that $B(\delta_i) \geq \eta_i/2$. Through Eq. (41), this exactly shows that for any non-decreasing allocation function p satisfying the constraint $\mathbf{1}^\top \mathbf{x} - \mathbf{1}^\top (\mathbb{E}[u_i p_i(\mathbf{u})])_{i \in [n]} \leq \delta_i$, we have

$$\mathbb{E} [|P_i(u_i) - P_i(q_i)| \min(u_i, q_i)] \geq \frac{\eta_i}{2}.$$

Together with Eq. (39) this ends the proof. ■

Theorem C.2 gives a lower bound on the range required for the interim future promised utilities. In fact it shows a stronger statement about their dispersion which is necessary for

our proofs. We now make use of these range lower bounds to show that some regions of $\mathcal{U}^*$ cannot be achieved by $\mathcal{U}_\gamma$. Intuitively, because future promised utilities need to remain within the achievable region $\mathcal{U}_\gamma$ (see Eq. (3)) if interim promised utilities are somewhat scattered, this gives a local upper bound on the maximum curvature of $\mathcal{U}_\gamma$.

Proof of Theorem 4.1 We start by specifying c_1 and c_2 . By the second characterization of Theorem 2.4, we have $\tilde{I} \neq \emptyset$, otherwise, $\mathcal{U}_0$ would already contain an optimal vector. Fix such an index $i \in \tilde{I}$. Let $\delta_i, \eta_i > 0$ be the constants for which Theorem C.2 holds. We pose $c_1 := \delta_i$ and $c_2 := \eta_i^2/16$.

Next, fix some parameters $\gamma, \mathbf{x}, r$ and δ satisfying the assumptions of the lemma. Since $\mathcal{U}^* \setminus B^\circ(\mathbf{x}, r)$ is compact, so is the component $\mathcal{C}$. By contradiction, we suppose that $\mathcal{U}_\gamma \cap (\mathcal{C} \setminus B^\circ(\mathbf{x}, r + \frac{1-\gamma}{\gamma}\delta)) \neq \emptyset$. We consider a solution $\mathbf{U}_0$ to the following optimization problem

$$\max_{\mathbf{y} \in \mathcal{C} \cap \mathcal{U}_\gamma} \|\mathbf{y} - \mathbf{x}\|. \quad (42)$$

The maximum is attained because both $\mathcal{C}$ and $\mathcal{U}_\gamma$ are compact. Further, by assumption,

$$\|\mathbf{U}_0 - \mathbf{x}\| \geq r + \frac{1-\gamma}{\gamma}\delta.$$

Because $\mathbf{U}_0 \in \mathcal{U}_\gamma$, there exists a mechanism for the central planner that reaches this utility. Using the promised utility framework, we fix a non-decreasing allocation function $\mathbf{p}(\cdot | \mathbf{U})$ and a promised utility function $\mathbf{W}(\cdot | \mathbf{U})$ for all $\mathbf{U} \in \mathcal{U}_\gamma$ that satisfy Eq. (3) (promises stay within $\mathcal{U}_\gamma$) and Eq. (7) (interim promises $W_i(v_i | \mathbf{U})$ as defined in Eq. (6)).

Step 1. Let $\mathbf{U}_1 := (\mathbb{E}[u_i p_i(\mathbf{v} | \mathbf{U}_0)])_{i \in [n]} \in \mathcal{U}^*$. As a first step, we aim to show that

$$\mathbf{1}^\top \mathbf{U}_1 \geq \max_{\mathbf{z} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{z} - \delta_i. \quad (43)$$

By Eq. (3) and the convexity of $\mathcal{U}_\gamma$, we have

$$\mathbf{U}_2 := \mathbb{E}[\mathbf{W}(\mathbf{u} | \mathbf{U}_0)] = \mathbf{U}_0 + \frac{1-\gamma}{\gamma}(\mathbf{U}_0 - \mathbf{U}_1) \in \mathcal{U}_\gamma. \quad (44)$$

In the equality, we used the form of the interim promises Eq. (6), or equivalently, Eq. (2). We next write $\mathbf{U}_0 = \mathbf{x} + r_0 \mathbf{y}_0$ where $r_0 \geq r + \frac{1-\gamma}{\gamma}\delta$ and $\mathbf{y}_0 \in S_{n-1}$. We argue that it suffices to show that $\mathbf{y}_0^\top (\mathbf{U}_1 - \mathbf{U}_0) \geq 0$ to obtain Eq. (43). Indeed, if this holds, then for any $\mathbf{U} \in [\mathbf{U}_0, \mathbf{U}_1]$, one has

$$\|\mathbf{U} - \mathbf{x}\| \geq \mathbf{y}_0^\top (\mathbf{U} - \mathbf{x}) = \mathbf{y}_0^\top (\mathbf{U} - \mathbf{U}_0) + r_0 \geq r_0 \geq r.$$

Because $\mathcal{U}^*$ is convex, $[\mathbf{U}_0, \mathbf{U}_1] \subseteq \mathcal{U}^*$. Combining this with the previous equation shows that $[\mathbf{U}_0, \mathbf{U}_1] \in \mathcal{U}^* \setminus B^\circ(\mathbf{x}, r)$. Hence $\mathbf{U}_1$ belongs to the same component $\mathcal{C}$ as $\mathbf{U}_0$. By assumption, $\mathcal{C} \subseteq \{\mathbf{y} : \mathbf{1}^\top \mathbf{y} \geq \max_{\mathbf{z} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{z} - c_1\}$ and we posed $c_1 = \delta_i$. This proves Eq. (43). We now show

$$\mathbf{y}_0^\top (\mathbf{U}_1 - \mathbf{U}_0) \geq 0. \quad (45)$$

By contradiction suppose that this does not hold, then Eq. (44) implies that $\mathbf{y}_0^\top (\mathbf{U}_2 - \mathbf{U}_0) > 0$. The same arguments as above then show that for any $\mathbf{U} \in [\mathbf{U}_0, \mathbf{U}_2]$,

$$\|\mathbf{U} - \mathbf{x}\| \geq \mathbf{y}_0^\top (\mathbf{U} - \mathbf{U}_0) + r_0 > r_0.$$

The previous arguments further show that $\mathbf{U}_2 \in \mathcal{C}$. Recall from Eq. (44) that $\mathbf{U}_2 \in \mathcal{U}_\gamma$. Hence, the previous inequality contradicts the definition of $\mathbf{U}_0$ as a maximizer of the problem in Eq. (42).

Step 2. We are now ready to apply Theorem C.2 to the allocation $\mathbf{p}(\cdot \mid \mathbf{U}_0)$ since the resulting utility vector $\mathbf{U}_1$ satisfies the necessary conditions as checked in Step 1. We obtain

$$\mathbb{E}_{u_i \sim \mathcal{D}_i} |W_i(u_i \mid \mathbf{U}_0) - W_i(q_i \mid \mathbf{U}_0)| \geq \frac{1-\gamma}{2\gamma} \eta_i, \quad (46)$$

where $q_i := \min\{v \in [0, \bar{v}] : \mathbb{P}(u_i \leq v) \geq 1/2\}$. Define $p_B = \frac{1/2 - \mathbb{P}(u_i < q_i)}{\mathbb{P}(u_i = q_i)}$, with the convention $0/0 = 0$ and let $B \sim \mathcal{B}(p_B)$ be a Bernoulli variable independent from all other variables. Define

$$\begin{aligned} \mathbf{U}_+ &= 2\mathbb{E}_{\mathbf{u}}[\mathbf{W}(u_i, \mathbf{u}_{-i} \mid \mathbf{U}_0) \mathbb{1}_{u_i < q_i} + B\mathbf{W}(u_i, \mathbf{u}_{-i} \mid \mathbf{U}_0) \mathbb{1}_{u_i = q_i}] \\ \mathbf{U}_- &= 2\mathbb{E}_{\mathbf{u}}[\mathbf{W}(u_i, \mathbf{u}_{-i} \mid \mathbf{U}_0) \mathbb{1}_{u_i > q_i} + (1-B)\mathbf{W}(u_i, \mathbf{u}_{-i} \mid \mathbf{U}_0) \mathbb{1}_{u_i = q_i}]. \end{aligned}$$

From Eq. (3), by convexity of $\mathcal{U}_\gamma$, and because $\mathbb{P}(u_i < q_i) + \mathbb{P}(u_i = q_i, B = 1) = 1/2$, both utility vectors satisfy $\mathbf{U}_+, \mathbf{U}_- \in \mathcal{U}_\gamma$. The goal of this step is to show that either $\mathbf{U}_+$ or $\mathbf{U}_-$ have a better objective than $\mathbf{U}_0$ for the optimization problem Eq. (42), to reach a contradiction. We first give a few properties about $\mathbf{U}_-$ and $\mathbf{U}_+$. We note that

$$\mathbf{U}_2 = \mathbb{E}[\mathbf{W}(\mathbf{u} \mid \mathbf{U}_0)] = \frac{\mathbf{U}_+ + \mathbf{U}_-}{2}. \quad (47)$$

Next, recalling that the interim allocation function $W_i(\cdot \mid \mathbf{U}_0)$ is non-increasing, we have

$$\begin{aligned} &\|\mathbf{U}_+ - \mathbf{U}_-\| \\ &\geq (\mathbf{U}_+)_i - (\mathbf{U}_-)_i \\ &= (\mathbf{U}_+)_i - W_i(q_i \mid \mathbf{U}_0) + W_i(q_i \mid \mathbf{U}_0) - (\mathbf{U}_-)_i \\ &= 2\mathbb{E}_{u_i}[(W_i(u_i \mid \mathbf{U}_0) - W_i(q_i \mid \mathbf{U}_0)) \mathbb{1}_{u_i > q_i}] + 2\mathbb{E}_{u_i}[(W_i(q_i \mid \mathbf{U}_0) - W_i(u_i \mid \mathbf{U}_0)) \mathbb{1}_{u_i < q_i}] \\ &= 2\mathbb{E}_{u_i} |W_i(u_i \mid \mathbf{U}_0) - W_i(q_i \mid \mathbf{U}_0)|. \end{aligned}$$

Combining with Eq. (46) yields

$$\|\mathbf{U}_+ - \mathbf{U}_-\| \geq \frac{1-\gamma}{\gamma} \eta_i, \quad (48)$$

We now lower bound $\|\mathbf{U}_2 - \mathbf{x}\|$. By Eq. (44), and recalling that $\mathbf{U}_0 = \mathbf{x} + r_0 \mathbf{y}_0$, we obtain

$$\mathbf{y}_0^\top (\mathbf{U}_2 - \mathbf{x}) = r_0 - \frac{1-\gamma}{\gamma} \mathbf{y}_0^\top (\mathbf{U}_1 - \mathbf{U}_0).$$

Within Step 1, we also showed that $\mathbf{U}_1 \in \mathcal{C} \subseteq B(\mathbf{x}, r + \delta)$, where in the last step we used one of the assumptions. Hence, $\mathbf{y}_0^\top (\mathbf{U}_1 - \mathbf{U}_0) = \mathbf{y}_0^\top (\mathbf{U}_1 - \mathbf{x}) - r_0 \leq (r + \delta) - r_0 \leq \delta$. Putting the two previous equations together yields

$$\|\mathbf{U}_2 - \mathbf{x}\| \geq \mathbf{y}_0^\top (\mathbf{U}_2 - \mathbf{x}) \geq r_0 - \frac{1-\gamma}{\gamma} \delta \geq r, \quad (49)$$

where in the last inequality we used the assumption on r_0 . We now argue that $\mathbf{U}_2 \in \mathcal{C}$. Indeed, since $\mathbf{U}_0, \mathbf{U}_2 \in \mathcal{U}^*$, by convexity $[\mathbf{U}_0, \mathbf{U}_2] \subseteq \mathcal{U}^*$. Next, for any $\mathbf{U} = (1-\theta)\mathbf{U}_2 + \theta\mathbf{U}_0$ for $\theta \in [0, 1]$,

$$\mathbf{y}_0^\top (\mathbf{U} - \mathbf{x}) = \mathbf{y}_0^\top (\mathbf{U}_2 - \mathbf{x}) + \theta \mathbf{y}_0^\top (\mathbf{U}_0 - \mathbf{U}_2) \geq r + \frac{1-\gamma}{\gamma} \theta \mathbf{y}_0^\top (\mathbf{U}_1 - \mathbf{U}_0) \geq r.$$

In the first inequality, we used Eq. (49), and in the second we used Eq. (45). As a result, using similar arguments as before, we have $[U_0, U_2] \in \mathcal{U}^* \setminus B^\circ(\mathbf{x}, r)$, hence $U_0 \in \mathcal{C}$ implies $U_2 \in \mathcal{C}$ as well. Next, by Eq. (47) we can write $U_+ = U_2 + \Delta$ and $U_- = U_2 - \Delta$ where $\Delta := (U_+ - U_-)/2$.

Suppose that $(U_2 - \mathbf{x})^\top \Delta \geq 0$. Then,

$$\begin{aligned} \|U_+ - \mathbf{x}\|^2 &= \|U_2 - \mathbf{x}\|^2 + \|\Delta\|^2 + 2(U_2 - \mathbf{x})^\top \Delta \\ &\geq \|U_2 - \mathbf{x}\|^2 + \left(\frac{1-\gamma}{2\gamma}\eta_i\right)^2 \\ &\geq r_0^2 + \frac{1-\gamma}{\gamma} \left(\frac{1-\gamma}{4\gamma}(\eta_i^2 + 4\delta^2) - 2r_0\delta\right) > r_0^2 + \frac{1-\gamma}{\gamma} \left(\frac{1-\gamma}{4\gamma}\eta_i^2 - 2r_0\delta\right) \end{aligned}$$

where in the second inequality we used Eq. (48) and in the third inequality we used Eq. (49). We now recall that by hypothesis $\delta = \min(c_2 \frac{1-\gamma}{\gamma r}, r)$. Hence, from the choice of c_2 , $2r_0\delta \leq 4r\delta \leq \frac{1-\gamma}{4\gamma}\eta_i^2$. Putting the two previous inequalities together, we obtain $\|U_+ - \mathbf{x}\| > \|U_0 - \mathbf{x}\| = r_0$. We recall that $U_2 \in \mathcal{C}$ and by assumption $(U_2 - \mathbf{x})^\top (U_+ - U_2) \geq 0$. Since we also have $U_+ \in \mathcal{U}^*$, the same arguments as in Step 1 show that $U_+ \in \mathcal{C}$. This contradicts the definition of U_0 as the maximizer of the problem in Eq. (42).

In the other case when $(U_2 - \mathbf{x})^\top \Delta \leq 0$, the same arguments show that $\|U_- - \mathbf{x}\| > \|U_0 - \mathbf{x}\|$ and $U_- \in \mathcal{C}$, reaching the same contradiction. This ends the proof of the desired result:

$$\mathcal{U}_\gamma \cap \left(\mathcal{C} \setminus B^\circ\left(\mathbf{x}, r + \frac{1-\gamma}{\gamma}\delta\right)\right) = \emptyset. \quad (50)$$

Proof of the second claim. We now prove the second claim of the lemma as a simple consequence of Eq. (50). First, note that because $\mathcal{U}_\gamma$ is compact, the maximum in $\max_{\mathbf{y} \in \mathcal{U}_\gamma} \mathbf{d}^\top (\mathbf{y} - \mathbf{x})$ is well defined. Fix $U = \mathbf{x} + \|U - \mathbf{x}\| \mathbf{d} \in \mathcal{C}$ such a maximizer. Consider any point $\mathbf{y} \in \mathcal{U}_\gamma$ such that $\mathbf{d}^\top (\mathbf{y} - \mathbf{x}) \geq \|U - \mathbf{x}\|$. Then, by convexity of $\mathcal{U}^*$, we have $[U, \mathbf{y}] \subseteq \mathcal{U}^*$. Further,

$$(U - \mathbf{x})^\top (\mathbf{y} - U) = \|U - \mathbf{x}\| \mathbf{d}^\top (\mathbf{y} - \mathbf{x}) - \|U - \mathbf{x}\|^2 \geq 0.$$

Then, the same arguments as in Step 1 show that since $U \in \mathcal{C}$, we must also have $\mathbf{y} \in \mathcal{C}$. Because we have $\mathbf{y} \in \mathcal{U}_\gamma$, Eq. (50) implies that $\mathbf{y} \in B^\circ(\mathbf{x}, r + \frac{1-\gamma}{\gamma}\delta)$. In particular,

$$\mathbf{d}^\top (\mathbf{y} - \mathbf{x}) \leq \|\mathbf{y} - \mathbf{x}\| \|\mathbf{d}\| < \left(r + \frac{1-\gamma}{\gamma}\delta\right) \|\mathbf{d}\|.$$

This ends the proof. ■

In some cases, it may be useful to design lower bounds based on a subproblem instead of the full resource allocation problem as stated in Theorem 4.1. Intuitively, we may want to focus on the impact of a single agent i , which corresponds to a 2-agent setting in which one has utility u_i and the other has utility Z_{-i} . In practice, this setup will be very useful, and we specify the corresponding lower bounds in the result below. This is a simple consequence of Theorem 4.1. For any $\mathbf{x} = (x_i, x_{-i}) \in \mathbb{R}^2$ and $r > 0$, we define the sets

$$\begin{aligned} B(\mathbf{x}, r; i) &:= \{\mathbf{y} \in \mathbb{R}^n : (y_i - x_i)^2 + (\mathbf{1}_{-i}^\top \mathbf{y}_{-i} - x_{-i})^2 \leq r^2\} \\ B(\mathbf{x}, r; i) &:= \{\mathbf{y} \in \mathbb{R}^n : (y_i - x_i)^2 + (\mathbf{1}_{-i}^\top \mathbf{y}_{-i} - x_{-i})^2 < r^2\}. \end{aligned}$$

Lemma C.3. Fix an index $i \in \tilde{I}$. There exist two constants $c_1, c_2 > 0$ such that the following holds.

For $\gamma \in (0, 1)$, $r > 0$, and $(x_i, x_{-i}) \in \mathbb{R}^2$, suppose that $\mathcal{U}^* \setminus B^\circ(\mathbf{x}, r; i)$ has a connected component $\mathcal{C} \subseteq \{\mathbf{y} : \mathbf{1}^\top \mathbf{y} \geq \max_{\mathbf{z} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{z} - c_1\}$. If

$$\mathcal{C} \subseteq B(\mathbf{x}, r + \delta; i), \quad \delta = \min \left(c_2 \frac{1 - \gamma}{\gamma r}, r \right),$$

then we have

$$\mathcal{U}_\gamma \cap \left(\mathcal{C} \setminus B^\circ \left(\mathbf{x}, r + \frac{1 - \gamma}{\gamma} \delta; i \right) \right) = \emptyset.$$

In particular, for any $\mathbf{U} \in \mathcal{C}$, writing $(U_i, \mathbf{1}_{-i}^\top \mathbf{U}_{-i}) = \mathbf{x} + \|(U_i, \mathbf{1}_{-i}^\top \mathbf{U}_{-i}) - \mathbf{x}\|(d_i, d_{-i})$,

$$\max_{\mathbf{y} \in \mathcal{U}_\gamma} d_i(y_i - x_i) + d_{-i}(\mathbf{1}_{-i}^\top \mathbf{y}_{-i} - x_{-i}) < \left(r + \frac{1 - \gamma}{\gamma} \delta \right) \|\mathbf{d}\|.$$

Proof In the proof of Theorem 4.1, we first used the characterization from Theorem 2.4 to show that there exists some index $i \in \tilde{I}$ ($\tilde{I} \neq \emptyset$). From then, the only incentive-compatibility constraint that was used within the proof was that of agent i . Precisely, it only uses Eq. (7) for agent i to use the fact that the interim promised utility is given according to Eq. (6) for which Theorem C.1 gives bounds.

As a result, the lower bounds also hold in a setting in which agents $j \neq i$ are non-strategic, that is, the central planner observes all utilities u_j for $j \notin i$. If we denote by $\tilde{\mathcal{U}}_{\gamma, i}$ the utility region that can be achieved in this setting, we clearly have $\mathcal{U}_\gamma \subseteq \tilde{\mathcal{U}}_{\gamma, i}$ since there are fewer constraints. From the perspective of maximizing the total utility, this corresponds to the setting in which there are effectively two agents, agent i with the same utility distribution $u_i \sim \mathcal{D}_i$ and an agent $-i$ that has utilities $u_{-i} := Z_{-i}$. We denote by $\mathcal{D}_{-i}$ its distribution and introduce the linear mapping

$$\Phi : \mathbf{U} \in \mathbb{R}^n \mapsto (U_i, \mathbf{1}_{-i}^\top \mathbf{U}_{-i}).$$

Let $\mathcal{U}_i^*$ and $\mathcal{U}_{\gamma, i}$ be the regions achievable with full information, and in the γ -discounted strategic setting for this 2-agent allocation problem. We first show that $\mathcal{U}_i^* = \Phi(\mathcal{U}^*)$. Since these are both convex sets, it suffices to focus on extreme points: by definition of $\mathcal{D}_{-i}$, for any $\mathbf{d} \in \mathbb{R}_+^2 \setminus \{\mathbf{0}\}$,

$$\max_{\mathbf{x} \in \mathcal{U}_{\gamma, i}} \mathbf{d}^\top \mathbf{x} = \mathbb{E} \max(d_i u_i, d_{-i} Z_{-i}) = \max_{\mathbf{x} \in \mathcal{U}^*} (d_i, d_{-i} \mathbf{1}_{-i})^\top \mathbf{x} = \max_{\mathbf{x} \in \mathcal{U}^*} \mathbf{d}^\top \Phi(\mathbf{x}).$$

Importantly, we also have

$$\Phi(\mathcal{U}_\gamma) \subseteq \Phi(\tilde{\mathcal{U}}_{\gamma, i}) = \mathcal{U}_{\gamma, i}. \quad (51)$$

The first inequality is immediate from $\mathcal{U}_\gamma \subseteq \tilde{\mathcal{U}}_{\gamma, i}$ and the equality comes from the definition of $\mathcal{D}_{-i}$ and Φ . Before applying the proof of Theorem 4.1, we only need to check that $\mathcal{D}_i$ still does not satisfy the conditions 2(a) nor 2(b) from Theorem 2.4 in the new 2-agent setting for the direction $\mathbf{d} := (1, 1)$. Using the assumption $i \in \tilde{I}$, for all $j \in [n]$ we have $\mathbb{P}(u_i > u_j) > 0$ hence since u_i and u_j are independent, there exists m_j such that $\mathbb{P}(u_i > m_j), \mathbb{P}(u_j \leq m_j) > 0$. Then,

$$\mathbb{P}(u_i > u_{-i}) \geq \mathbb{P}(u_i > \max_{j \neq i} m_j, \forall j \neq i, u_j \leq m_j) = \prod_{j \neq i} \mathbb{P}(u_j \leq m_j) \cdot \min_{j \neq i} \mathbb{P}(u_i > m_j) > 0.$$

Hence 2(a) does not hold. Now suppose by contradiction that 2(b) holds and let $[m_i, M_i]$ the guaranteed interval such that $\mathbb{P}(u_i \in [m_i, M_i] \cup \{0\}) = 1$ and

$$0 = \mathbb{P}(u_{-i} \in (m_i, M_i)) = \mathbb{P}(Z_{-i} \in (m_i, M_i)).$$

In the proof of Theorem 2.4 (within the proof of $1 \Rightarrow 2$), we showed that this implies $\mathbb{P}(u_j \in (m_i, M_i)) = 0$ for all $j \neq i$, which contradicts the hypothesis $i \in \tilde{I}$.

Therefore, we are ready to apply the proof of Theorem 4.1 for the 2-agent game for direction $\mathbf{1} = (1, 1)$ and agent i . Indeed, we can check that in the 2-agent setting, $B(\mathbf{x}, r) = \Phi(B(\mathbf{x}, r; i))$ and similarly for their open counterparts. Hence, using $\mathcal{U}_i^* = \Phi(\mathcal{U}^*)$,

$$\mathcal{U}_i^* \setminus B^\circ(\mathbf{x}, r) = \Phi(\mathcal{U}^* \setminus B^\circ(\mathbf{x}, r; i)),$$

and as a result $\Phi(\mathcal{C})$ is a connected component of $\mathcal{U}_i^* \setminus B^\circ(\mathbf{x}, r)$, which satisfies

$$\Phi(\mathcal{C}) \subseteq \Phi(\{\mathbf{y} : \mathbf{1}^\top \mathbf{y} \geq \max_{\mathbf{z} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{z} - c_1\}) = \{\mathbf{y} \in \mathbb{R}^2 : \mathbf{1}^\top \mathbf{y} \geq \max_{\mathbf{z} \in \mathcal{U}_i^*} \mathbf{1}^\top \mathbf{z} - c_1\},$$

$$\Phi(\mathcal{C}) \subseteq \Phi(B(\mathbf{x}, r + \delta; i)) = B(\mathbf{x}, r + \delta).$$

Thus, the proof of Theorem 4.1 gives

$$\Phi(\mathcal{U}_\gamma) \cap \left(\Phi(\mathcal{C}) \setminus B^\circ\left(\mathbf{x}, r + \frac{1-\gamma}{\gamma}\delta\right) \right) \subseteq \mathcal{U}_{\gamma,i} \cap \left(\Phi(\mathcal{C}) \setminus B^\circ\left(\mathbf{x}, r + \frac{1-\gamma}{\gamma}\delta\right) \right) = \emptyset,$$

where in the inclusion we used Eq. (51). This gives the desired result. The second claim is a direct translation of the second claim from Theorem 4.1 that we obtained on this 2-agent game. ■

To simplify the use of Theorems 3.1 and 4.1, we give simple equivalent computational constraints to prove that a ball belongs to $\mathcal{U}^*$, directly in terms of the distributions $\mathcal{D}_1, \dots, \mathcal{D}_n$.

Lemma C.4. *Let $\mathbf{x}$ and $r > 0$. $B(\mathbf{x}, r) \subseteq \mathcal{U}^*$ if and only if $\min_{i \in [n]} x_i \geq r$ and*

$$\forall \boldsymbol{\alpha} \in \mathbb{R}_+^n \setminus \{\mathbf{0}\}, \quad \mathbb{E} \left[\max_{i \in [n]} \alpha_i u_i \right] \geq \boldsymbol{\alpha}^\top \mathbf{x} + \|\boldsymbol{\alpha}\| r.$$

Proof Any convex set C is the intersection of its supporting hyperplanes. Hence,

$$\begin{aligned} B(\mathbf{x}, r) \subseteq \mathcal{U}^* &\iff \forall \boldsymbol{\alpha} \neq \mathbf{0}, \quad \max_{\mathbf{y} \in \mathcal{U}^*} \boldsymbol{\alpha}^\top \mathbf{y} \geq \max_{\mathbf{y} \in B(\mathbf{x}, r)} \boldsymbol{\alpha}^\top \mathbf{y} \\ &\iff B(\mathbf{x}, r) \subseteq \mathbb{R}_+^n \quad \text{and} \quad \forall \boldsymbol{\alpha} \in \mathbb{R}_+^n \setminus \{\mathbf{0}\}, \quad \mathbb{E} \left[\max_{i \in [n]} \alpha_i u_i \right] \geq \boldsymbol{\alpha}^\top \mathbf{x} + \|\boldsymbol{\alpha}\| r. \end{aligned}$$

This ends the proof. ■

D Characterizations in the discounted infinite-horizon setting

In this section, we prove our main results concerning lower bounds for the reachable region $\mathcal{U}_\gamma$ compared to the full-information reachable set $\mathcal{U}^*$. Before doing so, we need to introduce a pruning procedure that will help “pre-condition” the allocation problem.

D.1 A pruning procedure

For a given vector $\mathbf{U} \in \mathcal{U}^*$, we consider the following recursive pruning process. First let $I(\mathbf{U}) = \{i \in [n] : U_i > 0\}$ and initialize $\mathbf{U}^{(0)} = \mathbf{U}$. For any $k \geq 0$, if there exists $i \in I(\mathbf{U})$ with $U_i^{(k)} \geq \mathbb{E}[u_i]$, we let $i_{k+1} := i$. If $\mathbb{P}(u_i = 0) = 0$, pose $\mathbf{U}^{(k+1)} := \mathbf{0}$, otherwise

$$\mathbf{U}^{(k+1)} := \frac{\mathbf{U}^{(k)} - \mathbb{E}[u_{i_k}] \mathbf{e}_{i_k}}{\mathbb{P}(u_{i_k} = 0)}.$$

If for all $i \in I(\mathbf{U})$, $U_i < \mathbb{E}[u_i]$, the process stops. We denote by K the index of the last constructed index i_K and denote $\tilde{\mathbf{U}} := \mathbf{U}^{(K+1)}$ the final constructed vector. We pose $J(\mathbf{U}) = \{i \in I(\mathbf{U}) : \tilde{U}_i > 0\}$.

Properties of this pruning process are given below in Theorem D.1. The last claim shows that the pruning procedure preserves the regions $\mathcal{U}_\gamma$ while deleting agents. Essentially, this shows that without loss of generality, we can assume that $0 < U_i < \mathbb{E}[u_i]$ for all $i \in [n]$ (those are properties satisfied by $(\tilde{U}_j)_{j \in J(\mathbf{U})}$), which will be useful to derive general lower bounds.

Lemma D.1. *Let $\mathbf{U} \in \mathcal{U}^*$ and construct $i_1, \dots, i_K$, $\tilde{\mathbf{U}}$ and $J(\mathbf{U})$ according to the pruning process above.*

1. *For all $k \in [K]$, $\mathbf{U}^{(k)} \in \mathcal{U}^*$. In particular, $U_i^{(k)} \leq \mathbb{E}[u_i]$ for all $i \in I(\mathbf{U})$. Also, for $k \in [K-1]$, $\mathbb{P}(u_{i_k} = 0) > 0$.*

2.

$$\forall k \in [K], U_{i_k} = \mathbb{E}[u_{i_k}] \prod_{l < k} \mathbb{P}(u_{i_l} = 0) \quad \text{and} \quad \mathbf{U} = (U_i \mathbb{1}_{i \notin J(\mathbf{U})})_{i \in [n]} + \prod_{k \in [K]} \mathbb{P}(u_{i_k} = 0) \cdot \tilde{\mathbf{U}}.$$

3. *$\tilde{\mathbf{U}} \in \mathcal{U}^*$. Equivalently, $(\tilde{U}_j)_{j \in J(\mathbf{U})}$ belongs to the achievable set $\tilde{\mathcal{U}}^*$ for the full-information game in which only agents $J(\mathbf{U})$ are present.*
4. *Fix $\gamma \in [0, 1)$ and denote by $\tilde{\mathcal{U}}_\gamma$ the achievable set for the γ -discounted game in which only agents $J(\mathbf{U})$ are present. Then, $\mathbf{U} \in \mathcal{U}_\gamma$ if and only if $(\tilde{U}_j)_{j \in J(\mathbf{U})} \in \tilde{\mathcal{U}}_\gamma$.*

Proof Suppose that for some $k \in \{0, 1, \dots, K\}$ we have $\mathbf{U}^{(k)} \in \mathcal{U}^*$. By construction, $U_{i_k}^{(k)} \geq \mathbb{E}[u_{i_k}]$. However, since $\mathbf{U}^{(k)} \in \mathcal{U}^*$ we also have $U_{i_k}^{(k)} \leq \mathbb{E}[u_{i_k}]$. This proves $U_{i_k}^{(k)} = \mathbb{E}[u_{i_k}]$. Also, $i_k \in I(\mathbf{U})$ hence $\mathbb{E}[u_{i_k}] \geq U_{i_k} > 0$. As a result, $\mathbb{P}(u_{i_k} > 0) > 0$ and any allocation function $\mathbf{p}(\cdot)$ that realizes $\mathbf{U}^{(k)}$ allocates to agent i_k whenever $u_{i_k} > 0$, $\mathcal{D}_{i_k}$ -almost surely. This implies that for any $i \neq i_k$,

$$U_i^{(k)} = \mathbb{E}[u_i p_i(\mathbf{u})] = \mathbb{E}[u_i p_i(\mathbf{u}) \mathbb{1}_{u_{i_k}=0}] = \mathbb{P}(u_{i_k} = 0) \mathbb{E}[u_i p_i(0, \mathbf{u}_{-i_k})].$$

If $\mathbb{P}(u_{i_k} = 0) = 0$, then we had $\mathbf{U}^{(k)} = \mathbb{E}[u_{i_k}] \mathbf{e}_{i_k}$ and the process ends with $k = K$. Otherwise, the previous equation already shows that the allocation in which we force $u_{i_k} = 0$, that is

$$\mathbf{p}' : \mathbf{u} \in [0, \bar{v}]^n \mapsto (p_i(0, \mathbf{u}_{-i_k}) \mathbb{1}_{i \neq i_k})_{i \in [n]}, \tag{52}$$

realizes $\mathbf{U}^{(k+1)}$. Indeed, it only suffices to check $\mathbb{E}[u_{i_k} p'_i(\mathbf{u})] = 0 = U_{i_k}^{(k+1)}$. As a result, $\mathbf{U}^{(k+1)} \in \mathcal{U}^*$. This proves the first claim of the lemma as well as $\tilde{\mathbf{U}} = \mathbf{U}^{(K+1)} \in \mathcal{U}^*$, which is the third claim. The above arguments also show $U_{i_k}^{(k)} = \mathbb{E}[u_{i_k}]$ for all $k \in [K]$, and that $U_i^{(k)} = 0$ for all $i \notin I(\mathbf{U})$ and $i \in \{i_1, \dots, i_{k-1}\}$. We can then prove the second claim by simple induction.

Proving that $\mathbf{U} \in \mathcal{U}_\gamma \Rightarrow \tilde{\mathbf{U}} \in \mathcal{U}_\gamma$. To prove the last claim, further suppose that $\mathbf{U}^{(k)} \in \mathcal{U}_\gamma$ for $\gamma \in [0, 1)$ and some $k \in \{0, 1, \dots, K\}$. If $\mathbb{P}(u_{i_k} = 0) = 0$, we know that $\mathbf{U}^{(k+1)} = \mathbf{0} \in \mathcal{U}_0 \subseteq \mathcal{U}_\gamma$. From now, we suppose that

$$P_k := \mathbb{P}(u_{i_k} = 0) > 0.$$

Let $\mathbf{p}(\cdot \mid \mathbf{V})$ and $\mathbf{W}(\cdot \mid \mathbf{V})$ be valid incentive-compatible allocation and promised utility functions for all $\mathbf{V} \in \mathcal{U}_\gamma$ (that is, satisfying Eqs. (2) to (4)). Let $I_k = \{i \in [n], U_i^{(k)} > 0\}$. We denote $\mathcal{R}_k = \mathcal{U}_\gamma \cap \{\mathbf{V} : V_{i_k} = \mathbb{E}[u_{i_k}], \forall i \notin I_k, V_i = 0\}$. In particular, $\mathbf{U}^{(k)} \in \mathcal{R}_k$. For any $\mathbf{V} \in \mathcal{R}_k$,

$$\begin{aligned} \mathbb{E}[u_{i_k}] &= V_{i_k} = (1 - \gamma)\mathbb{E}[u_{i_k} p_{i_k}(\mathbf{u} \mid \mathbf{V})] + \gamma\mathbb{E}[W_{i_k}(\mathbf{u} \mid \mathbf{V})] \\ &\leq (1 - \gamma)\mathbb{E}[u_{i_k}] + \gamma\mathbb{E}[u_{i_k}] = \mathbb{E}[u_{i_k}]. \end{aligned}$$

In the inequality, we used the fact that $\mathbf{W}(\mathbf{v} \mid \mathbf{U}^{(k)}) \in \mathcal{U}_\gamma \subseteq \mathcal{U}^*$, hence $W_{i_k}(\mathbf{v} \mid \mathbf{U}^{(k)}) \leq \mathbb{E}[u_{i_k}]$. The above computations show that before, $\mathcal{D}_{i_k}$ -almost surely, $\mathbf{p}(\cdot \mid \mathbf{U}^{(k)})$ allocates to i_k whenever $u_{i_k} > 0$. Also, almost-surely $\mathbf{W}(\mathbf{u} \mid \mathbf{U}^{(k)}) \in \mathcal{R}_k$ as well. Without loss of generality, we can assume that this the case deterministically for all $\mathbf{u} \in [0, \bar{v}]^n$ (up to modification of the promised utility function on zero-measure regions). We now define

$$\mathcal{R}'_k := \left\{ \frac{\mathbf{V} - \mathbb{E}[u_{i_k}] \mathbf{e}_{i_k}}{P_k}, \mathbf{V} \in \mathcal{R}_k \right\}.$$

We construct the following allocation and promised utility functions similarly as in Eq. (52),

$$\begin{aligned} \mathbf{p}'(\mathbf{v} \mid \mathbf{V}') &:= (p_i(0, \mathbf{v}_{-i_k} \mid \mathbb{E}[u_{i_k}] \mathbf{e}_{i_k} + P_k \mathbf{V}') \mathbb{1}_{i \neq i_k})_{i \in [n]}, \\ \mathbf{W}'(\mathbf{v} \mid \mathbf{V}') &:= \frac{1}{P_k} (\mathbb{E}_{u_{i_k} \sim \mathcal{D}_{i_k}} [W_i(u_{i_k}, \mathbf{v}_{-i_k} \mid \mathbb{E}[u_{i_k}] \mathbf{e}_{i_k} + P_k \mathbf{V}')] \mathbb{1}_{i \neq i_k})_{i \in [n]}, \end{aligned}$$

for all $\mathbf{V}' \in \mathcal{R}'_k$ and $\mathbf{v} \in [0, \bar{v}]^n$. We can check that

$$\mathbf{W}'(\mathbf{v} \mid \mathbf{V}') = \frac{\mathbf{V} - \mathbb{E}[u_{i_k}] \mathbf{e}_{i_k}}{P_k}, \quad \text{where } \mathbf{V} = \mathbb{E}_{u_{i_k}} [\mathbf{W}(u_{i_k}, \mathbf{v}_{-i_k} \mid \mathbb{E}[u_{i_k}] \mathbf{e}_{i_k} + P_k \mathbf{V}')] \in \mathcal{R}_k.$$

Here we used the fact that $\mathcal{R}_k$ is convex since $\mathcal{U}_\gamma$ is convex. Thus, Eq. (3) is satisfied, that is all promised utilities stay within $\mathcal{R}'_k$. Then, the same arguments as for Eq. (52) show that for any $\mathbf{V}' \in \mathcal{R}'_k$, letting $\tilde{\mathbf{V}} := \mathbb{E}[u_{i_k}] \mathbf{e}_{i_k} + P_k \mathbf{V}'$ we have

$$(\mathbb{E}[u_i p'_i(\mathbf{u} \mid \mathbf{V}')])_{i \in [n]} = \frac{(\mathbb{E}[u_i p_i(\mathbf{u} \mid \tilde{\mathbf{V}})])_{i \in [n]} - \mathbb{E}[u_{i_k}] \mathbf{e}_{i_k}}{P_k}.$$

On the other hand,

$$\mathbb{E}[\mathbf{W}'(\mathbf{u} \mid \mathbf{V}')] = \frac{\mathbb{E}_{\mathbf{u}} \mathbb{E}_{u'_{i_k}} [\mathbf{W}(u'_{i_k}, \mathbf{u}_{-i_k} \mid \tilde{\mathbf{V}})] - \mathbb{E}[u_{i_k}] \mathbf{e}_{i_k}}{P_k} = \frac{\mathbb{E}_{\mathbf{u}} [\mathbf{W}(\mathbf{u} \mid \tilde{\mathbf{V}})] - \mathbb{E}[u_{i_k}] \mathbf{e}_{i_k}}{P_k}$$

Combining the two previous equations we obtain

$$\begin{aligned} (1 - \gamma)(\mathbb{E}[u_i p'_i(\mathbf{u} \mid \mathbf{V}')])_{i \in [n]} + \gamma \mathbb{E}[\mathbf{W}'(\mathbf{u} \mid \mathbf{V}')] \\ &= \frac{(1 - \gamma)(\mathbb{E}[u_i p_i(\mathbf{u} \mid \tilde{\mathbf{V}})])_{i \in [n]} + \gamma \mathbb{E}[\mathbf{W}(\mathbf{u} \mid \tilde{\mathbf{V}})] - \mathbb{E}[u_{i_k}] \mathbf{e}_{i_k}}{P_k} \\ &= \frac{\tilde{\mathbf{V}} - \mathbb{E}[u_{i_k}] \mathbf{e}_{i_k}}{P_k} = \mathbf{V}', \end{aligned}$$

where in the second equality we used the fact that $\mathbf{p}$ and $\mathbf{W}$ satisfy Eq. (2). The previous equation shows that $\mathbf{p}'$ and $\mathbf{W}'$ also satisfy Eq. (2). Last, for any $i \neq i_k$ and $u, v \in [0, \bar{v}]$ we have

$$(1 - \gamma)uP'_i(v | \mathbf{V}') + \gamma W'_i(v | \mathbf{V}') = \frac{(1 - \gamma)uP_i(v | \tilde{\mathbf{V}}) + \gamma W_i(v | \tilde{\mathbf{V}})}{P_k}.$$

Hence, because $\mathbf{p}, \mathbf{W}$ is incentive-compatible (Eq. (4)), so is $\mathbf{p}', \mathbf{W}'$ (we recall that for i_k , the utility of agent i_k is always 0). This ends the proof that $\mathbf{U}^{(k+1)} \in \mathcal{R}'_k \subseteq \mathcal{U}_\gamma$. The induction then shows that $\tilde{\mathbf{U}} \in \mathcal{U}_\gamma$, and hence $(\tilde{U}_j)_{j \in J(\mathbf{U})} \in \tilde{\mathcal{U}}_\gamma$.

Proving that $\tilde{\mathbf{U}} \in \mathcal{U}_\gamma \Rightarrow \mathbf{U} \in \mathcal{U}_\gamma$. Now suppose that $(\tilde{U}_j)_{j \in J(\mathbf{U})} \in \tilde{\mathcal{U}}_\gamma$, which implies in particular $\tilde{\mathbf{U}} \in \mathcal{U}_\gamma$. Let $\mathcal{R}' = \mathcal{U}_\gamma \cap \{\mathbf{V} : \forall i \notin J(\mathbf{U}), V_i = 0\}$. Let $\mathbf{p}'$ and $\mathbf{W}'$ be valid incentive-compatible allocation and promised utility functions for $\mathbf{V} \in \mathcal{U}_\gamma$. We can easily check that for any $\mathbf{V}' \in \mathcal{R}'$, then almost surely $W'_i(\mathbf{u} | \mathbf{V}') = 0$ for all $i \notin J(\mathbf{U})$, which implies $\mathbf{W}'(\mathbf{u} | \mathbf{V}') \in \mathcal{R}'$. Without loss of generality, we suppose this is the case (up to modifying $\mathbf{W}$ on zero-measure regions).

We next define $P_k := \prod_{k \in [K]} \mathbb{P}(u_{i_k} = 0)$ and $\mathbf{V}_0 := \sum_{k \in [K]} \mathbb{E}_{u_{i_k}} \prod_{l < k} \mathbb{P}(u_{i_l} = 0) \mathbf{e}_{i_l}$, we pose

$$\mathcal{R} := \{\mathbf{V}_0 + P_k \mathbf{V}', \mathbf{V}' \in \mathcal{R}'\}.$$

We consider the allocation $\mathbf{p}'$ which always tries to allocate to agents $i_1, i_2, \dots, i_K$ with this priority order whenever they have non-zero utility, then resorts to using the allocation $\mathbf{p}$. Formally, for all $\mathbf{v} \in [0, \bar{v}]^n$ and $\mathbf{V} \in \mathcal{R}$, letting $\mathbf{V}' = (\mathbf{V} - \mathbf{V}_0)/P_k \in \mathcal{R}'$ we pose

$$\begin{aligned} \mathbf{p}(\mathbf{v} | \mathbf{V}) &:= \begin{cases} \mathbf{e}_{i_l} & \text{if } \max_{k \in [K]} v_{i_k}^{(t)} > 0, \quad l = \min\{k \in [K] : v_{i_l}^{(t)} > 0\} \\ \mathbf{p}'(\mathbf{v} | \mathbf{V}') & \text{otherwise,} \end{cases} \\ \mathbf{W}(\mathbf{v} | \mathbf{V}) &:= \mathbf{V}_0 + P_k \mathbf{W}'(\mathbf{v} | \mathbf{V}'). \end{aligned}$$

We now check that $\mathbf{p}, \mathbf{W}$ are valid for utility vectors in $\mathcal{R}$. By construction of $\mathbf{W}$, it directly satisfies Eq. (3). Next, for any $\mathbf{V} \in \mathcal{R}$, with the same notation $\mathbf{V}' = (\mathbf{V} - \mathbf{V}_0)/P_k \in \mathcal{R}'$ we have

$$\begin{aligned} &(1 - \gamma)(\mathbb{E}[u_i p_i(\mathbf{u} | \mathbf{V})])_{i \in [n]} + \gamma \mathbb{E}[\mathbf{W}(\mathbf{v} | \mathbf{V})] \\ &= (1 - \gamma)(\mathbf{V}_0 + P_k (\mathbb{E}[u_i p'_i(\mathbf{u} | \mathbf{V}')]_{i \in [n]})) + \gamma(\mathbf{V}_0 + P_k \mathbb{E}[\mathbf{W}'(\mathbf{v} | \mathbf{V}')] = \mathbf{V}_0 + P_k \mathbf{V}' = \mathbf{V}. \end{aligned}$$

In the second equality, we used the fact that $\mathbf{p}', \mathbf{W}'$ satisfy Eq. (2). This shows that $\mathbf{p}, \mathbf{W}$ also satisfy it. All that remains to check is the incentive-compatibility constraint Eq. (4). For any $k \in [K]$, assuming all other agents are truthful, agent i_k is allocated the good whenever $u_{i_l} = 0$ for all $l < k$ and its report satisfies $v_{i_k} > 0$. Hence we only need to check that incentive-compatibility when $u_{i_K} = 0$, which is immediate. For any $i \notin \{i_k, k \in [K]\}$ and $u, v \in [0, \bar{v}]$, we directly have

$$(1 - \gamma)uP_i(v | \mathbf{V}) + \gamma W_i(v | \mathbf{V}) = P_k ((1 - \gamma)uP'_i(v | \mathbf{V}') + \gamma W'_i(v | \mathbf{V}')),$$

hence the incentive-compatibility is guaranteed by that of $\mathbf{p}', \mathbf{W}'$. This ends the proof. $\blacksquare$

D.2 A universal $\mathcal{O}(\sqrt{1-\gamma})$ convergence rate

With the pruning procedure and Theorem 3.1, we now show that without any assumptions, the distance between the boundary of $\mathcal{U}_\gamma$ and the optimal region $\mathcal{U}_\star$ is at most $\mathcal{O}(\sqrt{1-\gamma})$. This implies Theorem 5.1.

Theorem D.2. *For any $\mathbf{U} \in \mathcal{U}^\star$, construct $i_1, \dots, i_K$, $\tilde{\mathbf{U}}$ and $J(\mathbf{U})$ according to the pruning process. If $|J(\mathbf{U})| \leq 1$, we have $\mathbf{U} \in \mathcal{U}_0$. Otherwise, there exist $\gamma_0 \in [1/2, 1)$ and $C > 0$ such that*

$$\forall \gamma \in [\gamma_0, 1), \quad d(\mathbf{U}; \mathcal{U}_\gamma) \leq C\sqrt{1-\gamma}.$$

Precisely, introducing the constant $\tilde{C} = 4|J(\mathbf{U})|^2 \bar{v}^2 \left(1 + \frac{\max_{j \in J(\mathbf{U})} \mathbb{E}[u_i]}{\min_{j \in J(\mathbf{U})} (\mathbb{E}[u_j] - \tilde{U}_j)}\right)^2 < \infty$, we have $C = 2\tilde{C}\sqrt{|J(\mathbf{U})|} \prod_{k \in [K]} \mathbb{P}(u_{i_k} = 0)$ and $\gamma_0 = \max\left(1/2, 1 - \min_{j \in J(\mathbf{U})} \tilde{U}_i^2 / (32\tilde{C})\right) < 1$.

Proof Fix $\mathbf{U} \in \mathcal{U}^\star$ and define $i_1, \dots, i_K$, $\tilde{\mathbf{U}}$, and $J(\mathbf{U})$ according to the pruning process. If $|J(\mathbf{U})| \leq 1$, it is clear that $(\tilde{U}_j)_{j \in J(\mathbf{U})} \in \tilde{\mathcal{U}}_0$, and Theorem D.1 ensures that $\mathbf{U} \in \mathcal{U}_0$. We suppose that this is not the case from now. By construction of the pruning process, for all $j \in J(\mathbf{U})$, we have $\tilde{U}_j \in (0, \mathbb{E}[u_j])$, and $\tilde{U}_i = 0$ for all $i \notin J(\mathbf{U})$. We therefore consider the reduced problem in which only agents in $J(\mathbf{U})$ are present. By Theorem D.1, $\tilde{\mathbf{U}} \in \mathcal{U}^\star$, hence the utility vector $(\tilde{U}_j)_{j \in J(\mathbf{U})}$ can still be reached in this simplified setting with full information. We denote by $\tilde{\mathcal{U}}^\star = \{(x_j)_{j \in J(\mathbf{U})} : \mathbf{x} \in \mathcal{U}^\star\}$ the reachable utilities in this setting. By slight abuse of notation, we will still denote by $\tilde{\mathbf{U}}$ the vector $(\tilde{U}_j)_{j \in J(\mathbf{U})}$.

Note that $\mathcal{R} := \prod_{j \in J(\mathbf{U})} [0, \tilde{U}_j] \subseteq \tilde{\mathcal{U}}^\star$. Let $\tilde{C} = 4|J(\mathbf{U})|^2 \bar{v}^2 \left(1 + \frac{\max_{j \in J(\mathbf{U})} \mathbb{E}[u_i]}{\min_{j \in J(\mathbf{U})} (\mathbb{E}[u_j] - \tilde{U}_j)}\right)^2$ and pose $r := \delta := \sqrt{\tilde{C} \frac{1-\gamma}{\gamma}}$. Next, using $\gamma \geq \gamma_0$, we have

$$2(r + \delta) \leq 4\sqrt{2\tilde{C}(1-\gamma)} \leq \min_{j \in J(\mathbf{U})} \tilde{U}_i.$$

As a result, with $\mathbf{x} = \tilde{\mathbf{U}} - (r + \delta)\mathbf{1}$, we have $B(\mathbf{x}, r + \delta) \subseteq \mathcal{R} \subseteq \tilde{\mathcal{U}}^\star$. Since $\gamma \geq 1/2$, the choice of parameters $\mathbf{x}, r, \delta$ satisfies the assumptions of Theorem 3.1. Thus, $B(\mathbf{x}, r) \subseteq \tilde{\mathcal{U}}_\gamma$. In turn, Theorem D.1 implies that with $P := \prod_{k \in [K]} \mathbb{P}(u_{i_k} = 0)$, we have

$$(U_i \mathbb{1}_{i \notin J(\mathbf{U})})_{i \in [n]} + P \cdot B_{J(\mathbf{U})}(\mathbf{x}, r) \subseteq \mathcal{U}^\star,$$

where $B_{J(\mathbf{U})}(\mathbf{x}, r) = \{\mathbf{y} : (y_j)_{j \in J(\mathbf{U})} \in B(\mathbf{x}, r), \forall i \notin J(\mathbf{U}), y_i = 0\}$. Hence,

$$d(\mathbf{U}, \mathcal{U}^\star) \leq P \cdot d(\tilde{\mathbf{U}}, \tilde{\mathcal{U}}^\star) \leq P \cdot d(\tilde{\mathbf{U}}, B(\mathbf{x}, r)) = P(\sqrt{|J(\mathbf{U})|}(r + \delta) - r) \leq 2P\sqrt{|J(\mathbf{U})|}r.$$

This ends the proof of the theorem. ■

D.3 Improved rates for smooth full-information regions $\mathcal{U}^\star$

We first prove the following lemma.

Lemma D.3. *If $\mathcal{U}^\star$ is twice-differentiable locally around an optimal point $\mathbf{U}$, then it contains a ball tangent to $\mathbf{U}$ of radius $\mathcal{O}(1/\|H(\mathbf{U})\|)$ where $\|H(\mathbf{U})\|$ is the operator norm of the directional Hessian $H(\mathbf{U})$ of the boundary at $\mathbf{U}$.*

Proof Up to restricting to a subset of agents, suppose without loss of generality that $U_i > 0$ for all $i \in [n]$. In particular, $\mathcal{U}^*$ has non-empty interior. Consider the Taylor expansion of the boundary of $\mathcal{U}^*$ at $\mathbf{U}$ in coordinates given by the tangent hyperplane at $\mathbf{U}$. Then, the boundary of $\mathcal{U}^*$ is locally upper bounded by $\frac{1}{2}\mathbf{x}^\top H(\mathbf{U})\mathbf{x} + o(\|\mathbf{x}\|^2) \leq \|H(\mathbf{U})\| \cdot \|\mathbf{x}\|^2$ where $H(\mathbf{U})$ is the directional Hessian at $\mathbf{U}$ and $\|H(\mathbf{U})\|$ is its operator norm. Recalling that $\mathcal{U}^*$ is convex and has non-empty interior, this shows that $\mathcal{U}^*$ contains a ball tangent at $\mathbf{U}$ of radius $\mathcal{O}(1/\|H(\mathbf{U})\|)$. ■

Next, we prove the geometric result Theorem 5.2, which is a direct application of Theorem 3.1.

Proof of Theorem 5.2 Fix $B(\mathbf{x}, r)$ a ball tangent at $\mathbf{U} = \mathbf{x} + r\mathbf{1}$ to $\mathcal{U}^*$. Let $C > 0$ be the corresponding constant in Eq. (13) from Theorem 3.1. The constraint Eq. (13) is satisfied for $r(\gamma) = 1 - \delta(\gamma)$ and $\delta(\gamma) = 2C \frac{1-\gamma}{\gamma r}$ for $1-\gamma$ sufficiently small. Therefore, Theorem 3.1 implies that $B(\mathbf{x}, r(\gamma)) \subseteq \mathcal{U}_\gamma$. We conclude by noting $\text{SWG}(\gamma) \leq \sqrt{n} \cdot d(\mathbf{U}, B(\mathbf{x}, r(\gamma))) = \sqrt{n}\delta(\gamma) = \mathcal{O}(1-\gamma)$. ■

We next prove Theorem 5.3, which gives more flexible and combinatorial conditions under which faster rates $\mathcal{O}(1-\gamma)$ can be achieved. The most complete statement of this result is given below (and in particular implies Theorem 5.3, see condition (SC3) below). We recall the notation $P_S : \mathbb{R}^n \rightarrow \mathbb{R}^S$ for the projection on coordinates S .

Theorem D.4. *Suppose that there exists a partition $I = I_1 \sqcup I_2 \sqcup \dots \sqcup I_q$ with $|I_s| \geq 2$ for all $s \in [q]$, an optimal vector $\mathbf{U} \in \mathcal{U}^*$ and $r > 0$, such that (SC1) Eq. (30) holds and taking $\mathbf{x} := (U_i - r \sum_{s \in [q]} \frac{1}{\|\mathbf{1}_{I_s}\|} \mathbb{1}_{i \in I_s})_{i \in [n]}$, we have for all $s \in [q]$,*

$$B_{I_s}(\mathbf{x}_{I_s}, r) \subset P_{I_s}(\mathcal{U}^* \cap \{\mathbf{y} : \forall i \notin I_s, y_i = U_i\}).$$

Then $\text{SWG}(\gamma) = \mathcal{O}(1-\gamma)$.

Condition (SC1) is equivalent to the following condition (SC2): Eq. (30) holds, and for all $s \in [q]$ and $\emptyset \subsetneq B \subsetneq I_s$,

$$\mathbb{E} \left[Z_B \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} = Z \leq (1+\eta)Z_B} \right] \xrightarrow{\eta \rightarrow 0} \Omega(\eta).$$

Last, (SC1) or (SC2) are implied by the following condition (SC3): there exists a partition $I = I_1 \sqcup I_2 \sqcup \dots \sqcup I_q$ such that for all $s \in [q]$, $|I_s| \geq 2$ and I_s is strongly connected in $\mathcal{G}$.

Proof of Theorem D.4 For $s \in [q]$, we define the function $\phi : \mathbb{R}^{I_s} \rightarrow \mathbb{R}$ via

$$\phi^{(s)}(\boldsymbol{\beta}) := \max_{\mathbf{y} \in P_{I_s}(\mathcal{U}^* \cap \{\mathbf{y} : \forall i \notin I_s, y_i = U_i\})} \boldsymbol{\beta}^\top \mathbf{y} = \mathbb{E} \left[\mathbb{1}_{Z_{I_s} = Z > 0} \cdot \max_{i \in I_s} \beta_i u_i \right], \quad \boldsymbol{\beta} \in \mathbb{R}^{I_s}. \quad (53)$$

In the equality, we use the characterizations from Theorem B.2. We also recall that by Eq. (30), we have $\mathbb{P}(Z_{I_s} = Z_{[n] \setminus I_s} > 0) = 0$. Note that ϕ is convex as the maximum of linear functions.

Proving that the (SC1) is sufficient. Suppose that (SC1) holds and fix parameters satisfying the assumptions. Theorem B.2 implies that for $1-\gamma$ sufficiently small and $\delta = 2C(1-\gamma)/(\gamma r)$ for some constant $C > 0$, $B(\mathbf{x}, r-\delta) \cap \{\mathbf{y} : \forall i \notin I, y_i = U_i\} \subseteq \mathcal{U}_\gamma$. In turn, we obtain the desired result $\text{SWG} \leq \sqrt{n} \cdot d(\mathbf{U}, \mathcal{U}_\gamma) \leq \sqrt{n}\delta = \mathcal{O}(1-\gamma)$.

Proving that (SC3) $\Rightarrow$ (SC2) $\Rightarrow$ (SC1). Suppose that (SC3) holds and denote $I = J_1 \sqcup J_2 \sqcup \dots \sqcup J_{\tilde{q}}$ the guaranteed partition. To prove (SC2) and (SC1) we will use a slightly different partition by merging some sets. Precisely, consider the undirected graph $\mathcal{F}$ on $[\tilde{q}]$ such that $s \neq s' \in [q]$ are connected if and only if $\mathbb{P}(Z_{J_s} = Z_{J_{s'}} = Z > 0) > 0$. We let q be the number of connected components $C_1, \dots, C_q$ of $\mathcal{F}$. For $s \in [q]$ we let $I_s := \bigcup_{t \in C_s} J_t$, that is we merge the partition according to the graph $\mathcal{F}$. Note that if $s \neq s'$ are connected within $\mathcal{F}$ then

$$0 < \mathbb{P}(Z_{J_s} = Z_{J_{s'}} = Z > 0) \leq \sum_{i \in J_s, j \in J_{s'}} \mathbb{P}(u_i = u_j = Z > 0).$$

Hence, there exists $i \in J_s$ and $j \in J_{s'}$ such that $\mathbb{P}(u_i = u_j = Z > 0) > 0$. This implies that $i \rightarrow j$ and $j \rightarrow i$ in the graph $\mathcal{G}$, and hence J_s is strongly connected to $J_{s'}$ via this double edge. This proves that the sets of the constructed partition $I = I_1 \sqcup I_2 \sqcup \dots \sqcup I_q$ are still strongly connected within $\mathcal{G}$.

We can directly check that Eq. (30) is satisfied by construction of $I_1, \dots, I_q$: for any $s \neq s' \in [q]$

$$\mathbb{P}(Z_{I_s} = Z_{I_{s'}} = Z > 0) \leq \sum_{t \in C_s, t' \in C_{s'}} \mathbb{P}(Z_{J_t} = Z_{J_{t'}} = Z > 0) = 0.$$

In the last equality we use the fact that C_s and $C_{s'}$ are disconnected in $\mathcal{F}$.

We next define an allocation function $\mathbf{p}$ such that $\mathbf{p}(\mathbf{u})$ is the uniform distribution on $\arg \max_{j \in [n]} u_j$. We then define $\mathbf{U} = (\mathbb{E}[u_i p_i(\mathbf{u})])_{i \in [n]}$ as the utility vector realized by $\mathbf{p}$. Note that by construction, $\mathbf{U}$ is optimal and we have $U_i > 0$ for all $i \in I$. In particular, for all $s \in [q]$ we have $\phi^{(s)}(\mathbf{1}_{I_s}) = \mathbf{1}_{I_s}^\top \mathbf{U}_{I_s}$.

Suppose that there exists $c_1, c_2 > 0$ such that for any $s \in [q]$ and $\boldsymbol{\beta} \in \mathcal{H}_s := \{\boldsymbol{\delta} \in \mathbb{R}^{I_s} : \mathbf{1}_{I_s}^\top \boldsymbol{\delta} = 0\}$, with $\|\boldsymbol{\beta}\| \leq c_1$, we have

$$\phi^{(s)}(\mathbf{1}_{I_s} + \boldsymbol{\beta}) - \phi^{(s)}(\mathbf{1}_{I_s}) - \boldsymbol{\beta}^\top \mathbf{U}_{I_s} \geq c_2 \|\boldsymbol{\beta}\|^2. \quad (54)$$

Because the function $\boldsymbol{\beta} \mapsto \phi^{(s)}(\mathbf{1}_{I_s} + \boldsymbol{\beta}) - \phi^{(s)}(\mathbf{1}_{I_s}) - \boldsymbol{\beta}^\top \mathbf{U}_{I_s}$ is convex in $\boldsymbol{\beta}$ and has value 0 for $\boldsymbol{\beta} = \mathbf{0}$, this implies that for all $\boldsymbol{\beta} \in \mathcal{H}_s$,

$$\phi^{(s)}(\mathbf{1}_{I_s} + \boldsymbol{\beta}) - \phi^{(s)}(\mathbf{1}_{I_s}) - \boldsymbol{\beta}^\top \mathbf{U}_{I_s} \geq c_2 \min(\|\boldsymbol{\beta}\|^2, c_1 \|\boldsymbol{\beta}\|) \geq c_3 \min(\|\boldsymbol{\beta}\|^2, \|\mathbf{1}_{I_s}\| \|\boldsymbol{\beta}\|), \quad (55)$$

for $c_3 = \min(c_2, c_1 c_2 / \|\mathbf{1}_{I_s}\|)$. Let $r = \min(2c_3 \min_{s \in [q]} \|\mathbf{1}_{I_s}\|, \min_{i \in I} U_i)$. We note that $\mathbf{x}_{I_s} = \mathbf{U}_{I_s} - r \frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|}$. Since $\boldsymbol{\beta}^\top \mathbf{1}_{I_s} = 0$, we have

$$\begin{aligned} (\mathbf{1}_{I_s} + \boldsymbol{\beta})^\top \mathbf{x}_{I_s} + \|\mathbf{1}_{I_s} + \boldsymbol{\beta}\| r &= (\mathbf{1}_{I_s} + \boldsymbol{\beta})^\top \mathbf{U}_{I_s} + r(\|\mathbf{1}_{I_s} + \boldsymbol{\beta}\| - \|\mathbf{1}_{I_s}\|) \\ &= \phi^{(s)}(\mathbf{1}_{I_s}) + \boldsymbol{\beta}^\top \mathbf{U}_{I_s} + r \|\mathbf{1}_{I_s}\| \left(\sqrt{1 + (\|\boldsymbol{\beta}\| / \|\mathbf{1}_{I_s}\|)^2} - 1 \right) \\ &\leq \phi^{(s)}(\mathbf{1}_{I_s}) + \boldsymbol{\beta}^\top \mathbf{U}_{I_s} + \frac{2r}{\|\mathbf{1}_{I_s}\|} \min(\|\boldsymbol{\beta}\|^2, \|\mathbf{1}_{I_s}\| \|\boldsymbol{\beta}\|) \leq \phi^{(s)}(\mathbf{1}_{I_s} + \boldsymbol{\beta}). \end{aligned}$$

In the last inequality, we used Eq. (55). As a result, Theorem C.4 shows that we have $B_{I_s}(\mathbf{x}_{I_s}, r) \subseteq P_{I_s}(\mathcal{U}^* \cap \{\mathbf{y} : \forall i \notin I_s, y_i = U_i\})$. We can also easily check that it contains $\mathbf{U}$. Hence, this proves (SC1).

We now show that Eq. (54) holds. We proceed independently for each $s \in [q]$. Fix $\boldsymbol{\beta} \in \mathcal{H}_s \setminus \{\mathbf{0}\}$. We first show that there is a subset $\emptyset \subsetneq B \subsetneq I_s$ such that

$$\min_{i \in B} \beta_i - \max_{i \in I_s \setminus B} \beta_i \geq \frac{\|\boldsymbol{\beta}\|}{2n^2}. \quad (56)$$

Let $\epsilon = \frac{1}{2n^2} \|\beta\|$. We consider the partition of $\mathbb{R}$ into $(-\infty, 0]$, $(n\epsilon, \infty)$ and n length- ϵ intervals partitioning $[0, n\epsilon)$. To prove Eq. (56) it suffices to show that there is one of these intervals of length ϵ that does not contain an element of the form β_i for $i \in I_s$, but has elements below and above it. Suppose that this is not the case. Then, $\{\beta_i \geq 0, i \in I_s\}$ occupy consecutive intervals of the partition. Since $\mathbf{1}_{I_s}^\top \beta = 0$, β has both positive and negative entries, this shows that the consecutive intervals occupied by $\{\beta_i, i \in I_s\}$ must start at $[0, \epsilon)$. There are at most $|I_s| - 1 \leq n - 1$ such intervals (one point must belong to $(-\infty, 0)$), hence we proved that all elements from $\{\beta_i \geq 0, i \in I_s\}$ lie consecutively within the n intervals within $[0, (n-1)\epsilon)$. Then, since $\mathbf{1}_{I_s}^\top \beta = 0$, we have

$$\sum_{i \in I_s: \beta_i < 0} |\beta_i| = \sum_{i \in I_s: \beta_i \geq 0} \beta_i \leq |I_s|(n-1)\epsilon < n^2\epsilon.$$

Therefore, we obtain the following contradiction

$$\|\beta\| \leq \|\beta\|_1 \leq \sum_{i \in I_s} |\beta_i| < 2n^2\epsilon = \|\beta\|.$$

In the rest of the proof, we fix a subset $\emptyset \subsetneq B \subsetneq I_s$ satisfying Eq. (56). We let $\gamma_1 := \min_{i \in B} \beta_i$ and $\gamma_2 := \max_{i \in I_s \setminus B} \beta_i$. By assumption, we have $\gamma_1 \geq \gamma_2 + \epsilon$. Provided that $\|\beta\| \leq 1/2$, we have $|\gamma_2| \leq 1/2$. We suppose that this is the case from now. Recall that $\mathbf{p}$ is the allocation that realizes $\mathbf{U}$ in the full information setting. We also introduce the notation $\hat{i}$ for the agent that effectively receives the allocation (that is, $\hat{i} \sim \mathbf{p}(\mathbf{u})$). Using Eq. (53), we have

$$\phi^{(s)}(\mathbf{1}_{I_s} + \beta) - \phi^{(s)}(\mathbf{1}_{I_s}) - \beta^\top \mathbf{U}_{I_s} \tag{57}$$

$$\begin{aligned} &= \mathbb{E} \left[\left(\max_{i \in I_s} (1 + \beta_i) u_i - (1 + \beta_{\hat{i}}) u_{\hat{i}} \mathbb{1}_{\hat{i} \in I_s} \right) \mathbb{1}_{Z=Z_{I_s}} \right] \\ &\geq \mathbb{E} \left[\left(\max_{i \in I_s} (1 + \beta_i) u_i - (1 + \beta_{\hat{i}}) u_{\hat{i}} \mathbb{1}_{\hat{i} \in I_s} \right) \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} = Z \leq \frac{1+\gamma_1}{1+\gamma_2} Z_B} \right] \\ &\stackrel{(i)}{\geq} \frac{1}{n} \mathbb{E} \left[((1 + \gamma_1) Z_B - (1 + \gamma_2) Z_{I_s \setminus B}) \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} = Z \leq \frac{1+\gamma_1}{1+\gamma_2} Z_B} \right] \\ &\stackrel{(ii)}{\geq} \frac{1 + \gamma_2}{n} \mathbb{E} \left[((1 + \epsilon/2) Z_B - Z_{I_s \setminus B}) \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} = Z \leq (1+\epsilon/4) Z_B} \right] \\ &\stackrel{(iii)}{\geq} \frac{\epsilon}{8n} \mathbb{E} \left[Z_B \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} = Z \leq (1+\epsilon/4) Z_B} \right]. \end{aligned} \tag{58}$$

In (i), the additional factor $1/n$ comes from the fact that if we have a tie $Z = Z_{I_s \setminus B}$, there is at least a probability $1/n$ that $\hat{i} \in I_s \setminus B$ by definition of $\mathbf{p}$. In (ii) and (iii) we used $|\gamma_2| \leq 1/2$ so that $(1 + \gamma_1) \geq (1 + \gamma_2)(1 + \epsilon/2)$ and $1 + \gamma_2 \geq 1/2$. We compute for any $\eta \geq 0$,

$$\begin{aligned} \mathbb{E} \left[Z_B \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} = Z \leq (1+\eta) Z_B} \right] &\geq \max_{j \in I_s \setminus B} \mathbb{E} \left[Z_B \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} \leq (1+\eta) Z_B} \mathbb{1}_{u_j = Z_{I_s \setminus B} = Z} \right] \\ &\geq \max_{i \in B, j \in I_s \setminus B} \mathbb{E} \left[u_i \mathbb{1}_{u_j = Z} \mathbb{1}_{u_j \in [u_i, (1+\eta) u_i]} \right]. \end{aligned}$$

By assumption, because $\mathcal{G}$ is strongly connected, there exists an outgoing edge from B . This implies that the right-hand side of the previous equation is $\Omega(\eta)$. This proves the condition

(SC2). Because there are only a finite number of subsets $\emptyset \subsetneq B \subsetneq I_s$ and $s \in [q]$, this shows that there exists $c_3, c_4 > 0$ such that for all $s \in [q]$ and $\emptyset \subsetneq B \subsetneq I_s$,

$$\mathbb{E} \left[Z_B \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} = Z \leq (1+\eta)Z_B} \right] \geq c_3 \eta, \quad \eta \in [0, c_4].$$

Together with the previous estimates, this shows that for $s \in [q]$, and any $\beta \in \mathcal{H}_s$ satisfying $\|\beta\| \leq \min(1/2, c_4 n^{3/2})$, we have $\epsilon \leq c_4$ hence we can prove Eq. (54) as follows

$$\phi^{(s)}(\mathbf{1}_{I_s} + \beta) - \phi^{(s)}(\mathbf{1}_{I_s}) - \beta^\top \mathbf{U}_{I_s} \geq \frac{c_3}{32n} \epsilon^2 = \frac{c_3}{32n^4} \|\beta\|^2.$$

Proving that (SC1) $\Rightarrow$ (SC2). Suppose that we are given $\mathbf{U}$ and a radius $r > 0$ satisfying (SC1). We use similar notations as before: let $\mathbf{p}(\cdot)$ be an allocation function that realizes $\mathbf{U}$ and denote by $\hat{i} \sim \mathbf{p}(\mathbf{u})$ the agent that receives the resource with this allocation. By Theorem C.4, and using the same arguments as before, for $s \in [q]$ and every $\beta \in \mathcal{H}_s$, we have

$$\phi^{(s)}(\mathbf{1}_{I_s} + \beta) \geq \phi^{(I_s)}(\mathbf{1}_{I_s}) + \beta^\top \mathbf{U}_{I_s} + r \|\mathbf{1}_{I_s}\| \left(\sqrt{1 + (\|\beta\|/\|\mathbf{1}_{I_s}\|)^2} - 1 \right).$$

Using the inequality $\sqrt{1+x} - x \geq (\sqrt{2} - 1)x$ for $x \in [0, 1]$, this shows that for $\beta \in \mathcal{H}_s$ with $\|\beta\| \leq \|\mathbf{1}_{I_s}\|$,

$$\phi^{(s)}(\mathbf{1}_{I_s} + \beta) - \phi^{(s)}(\mathbf{1}_{I_s}) - \beta^\top \mathbf{U}_{I_s} \geq \frac{(\sqrt{2} - 1)r}{\|\mathbf{1}_{I_s}\|} \|\beta\|^2.$$

Next fix any $\emptyset \subsetneq B \subsetneq I_s$. We consider the value $\beta_B = (\eta \mathbb{1}_{i \in B} - \eta' \mathbb{1}_{i \in I_s \setminus B})_{i \in I_s}$ where $\eta \geq 0$ and we posed $\eta' := \eta|B|/|I_s \setminus B|$, so that $\mathbf{1}_{I_s}^\top \beta = 0$. For sufficiently small $\eta \geq 0$, we have $\|\beta\| = \eta\sqrt{2|B|} \leq \|\mathbf{1}_{I_s}\|$. Then,

$$\begin{aligned} \phi^{(s)}(\mathbf{1} + \beta) - \phi^{(s)}(\mathbf{1}) - \beta^\top \mathbf{U} &= \mathbb{E} \left[\left(\max_{i \in I_s} (1 + \beta_i) u_i - (1 + \beta_i) u_{\hat{i}} \mathbb{1}_{\hat{i} \in I_s} \right) \mathbb{1}_{Z = Z_{I_s}} \right] \\ &\leq \mathbb{E} \left[\left(\max((1 + \eta)Z_B, (1 - \eta')Z_{I_s \setminus B}) - (1 + \eta)Z_B \mathbb{1}_{\hat{i} \in B} - (1 - \eta')Z_{I_s \setminus B} \mathbb{1}_{\hat{i} \in I_s \setminus B} \right) \mathbb{1}_{Z = Z_{I_s}} \right] \\ &\leq \mathbb{E} \left[\left((1 + \eta)Z_B - (1 - \eta')Z_{I_s \setminus B} \right) \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} = Z \leq \frac{1+\eta}{1-\eta'} Z_B} \right] \\ &\leq (\eta + \eta') \mathbb{E} \left[Z_B \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} = Z \leq \frac{1+\eta}{1-\eta'} Z_B} \right]. \end{aligned}$$

For $\eta \geq 0$ sufficiently small, we have $\frac{1+\eta}{1-\eta'} \leq 1 + 2(\eta + \eta')$. Therefore, for $\zeta \geq 0$ sufficiently small, taking $\eta \geq \zeta$ such that $\zeta = \eta + \eta'$, we obtained

$$\mathbb{E} \left[Z_B \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} = Z \leq (1+2\zeta)Z_B} \right] \geq \frac{(\sqrt{2} - 1)r|B|}{2\|\mathbf{1}_{I_s}\| (1 + |B|/|I_s \setminus B|)^2} \zeta.$$

This proves (SC2). ■

We now prove Theorem 5.4 which gives simpler conditions for reaching the fast rates $\mathcal{O}(1-\gamma)$. **Proof of Theorem 5.4** We show that condition (SC3) from Theorem 5.3 is satisfied. For any $i \in I$, let $j \in I$ be the corresponding index satisfying the assumption for i . In the first case, we obtain directly

$$\mathbb{P}(u_i = u_j = Z = m_i) \geq \mathbb{P}(u_i = m_i) \mathbb{P}(u_j = m_i) \prod_{k \notin \{i, j\}} \mathbb{P}(u_k \leq m_i) > 0.$$

Hence $m_i > 0$ implies $i \rightarrow j$ and $j \rightarrow i$. In the second case, let f_i be the lower bound on the density of the absolutely-continuous part of u_i and u_j on $[a_i, b_i]$. Since $a_i > m$, we have $\mathbb{P}(Z \leq a_i) > 0$. This shows that for all $k \in [n]$, $\mathbb{P}(u_k \leq a_i) > 0$. Then,

$$\begin{aligned} \mathbb{E} \left[u_i \mathbb{1}_{u_j=Z} \mathbb{1}_{u_j \in [u_i, (1+\eta)u_i]} \right] &\geq a_i \mathbb{P}(\forall k \notin \{i, j\}, u_k \leq a_i) \mathbb{P}(a_i \leq u_i \leq u_j \leq u_i + \eta a_i) \\ &\geq a_i \prod_{k \notin \{i, j\}} \mathbb{P}(u_k \leq a_i) \cdot f_i^2(b_i - a_i - \eta a_i) \eta a_i = \Omega(\eta). \end{aligned}$$

This shows that $i \rightarrow j$. The same arguments also show that $j \rightarrow i$.

In summary, $\mathcal{G}$ contains a subgraph $\mathcal{G}'$ in which $i \rightarrow j$ implies $j \rightarrow i$ and such that every node has an outgoing edge. We then use the partition of I given by the connected components of $\mathcal{G}'$, that all have at least 2 elements and are strongly connected. Hence (SC3) is satisfied. ■

D.4 General rates for arbitrary utility distributions

Theorems B.2 and 3.1, and Theorems C.3 and 4.1 can be used more generally to prove any convergence rates between $\sqrt{1-\gamma}$ and $1-\gamma$, as characterized in Theorem 5.8. Before doing so, we introduce an additional notation. Fix a partition $I = I_1 \sqcup I_2 \sqcup \dots \sqcup I_q$ with $|I_s| \geq 2$ for all $s \in [q]$, which we denote $\mathcal{P}$. In the main text, we introduced the quantity $f(\eta; \mathcal{P})$ to quantify the central planner's ability to meet future promises using the partition $\mathcal{P}$ to group agents. For technical purposes, we introduce another related quantity as follows:

$$\tilde{f}(\eta; \mathcal{P}) := \min_{s \in [q]} \min_{\emptyset \subsetneq B \subsetneq I_s} \mathbb{E} \left[Z_B \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} = Z \leq (1+\eta)Z_B} \right].$$

These functions are identical up to constant factors as checked below.

Lemma D.5. *Fix a partition $I = I_1 \sqcup \dots \sqcup I_q$ with $|I_s| \geq 2$ for all $s \in [q]$, which we denote $\mathcal{P}$. For any $\eta \geq 0$, we have*

$$\frac{1}{n^2} f(\eta; \mathcal{P}) \leq \tilde{f}(\eta; \mathcal{P}) \leq f(\eta; \mathcal{P}).$$

Proof Fix $\eta \geq 0$ and $s \in [q]$. Using the min-cut/max-flow duality theorem, we have

$$\begin{aligned} \min_{i \neq j \in I_s} (\text{max flow from } i \text{ to } j \text{ on } \mathcal{H}_{|I_s|}) &= \min_{i \neq j \in I_s} \min_{\{i\} \subseteq B \subseteq I_s \setminus \{j\}} \sum_{k \in B, l \in I_s \setminus B} f(\eta; k, l) \\ &= \min_{\emptyset \subsetneq B \subsetneq I_s} \sum_{i \in B, j \in I_s \setminus B} f(\eta; i, j). \end{aligned}$$

We next note that for any $\emptyset \subsetneq B \subsetneq I_s$, we have

$$\mathbb{E} \left[Z_B \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} = Z \leq (1+\eta)Z_B} \right] = \mathbb{E} \left[\max_{i \in B, j \in I_s \setminus B} u_i \mathbb{1}_{u_i = Z_B} \mathbb{1}_{u_j = Z} \mathbb{1}_{u_j \in [u_i, (1+\eta)u_i]} \right].$$

Therefore,

$$\begin{aligned} \frac{1}{n^2} \sum_{i \in B, j \in I_s \setminus B} f(\eta; i, j) &\leq \max_{i \in B, j \in I_s \setminus B} f(\eta; i, j) \leq \mathbb{E} \left[Z_B \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} = Z \leq (1+\eta)Z_B} \right] \\ &\leq \sum_{i \in B, j \in I_s \setminus B} f(\eta; i, j) \leq \sum_{i \in B, j \in I_s \setminus B} f(\eta; i, j). \end{aligned}$$

Taking the minimum over all $s \in [q]$ and $\emptyset \subsetneq B \subsetneq I_s$ then gives the desired result. $\blacksquare$

We are now ready to prove the complete characterization from Theorem 5.8

Proof of Theorem 5.8 We prove the upper and lower bound separately.

Proof of the first claim. Fix a partition $\tilde{\mathcal{P}}$ of $I = J_1 \sqcup J_2 \sqcup \dots \sqcup J_{\tilde{q}}$ with $|J_2| \geq 2$ for $s \in [\tilde{q}]$. As a first step, we construct a new partition $I = I_1 \sqcup I_2 \sqcup \dots \sqcup I_q$ which we denote $\mathcal{P}$ by merging together parts J_s and $J_{s'}$ satisfying $\mathbb{P}(Z_{J_s} = Z_{J_{s'}} = Z > 0) > 0$. This is the same procedure as within the proof of (SC3) $\implies$ (SC2) for Theorem D.4. This proof then shows that the resulting partition $\mathcal{P}$ satisfies Eq. (30). Further, if two parts J_s and $J_{s'}$ are merged, there exists $j_s \in J_s$ and $j_{s'} \in J_{s'}$ such that $\mathbb{P}(u_{j_s} = u_{j_{s'}} = Z > 0) > 0$. In particular, $f(\eta; j_s, j_{s'}) = \Omega(1)$ as $\eta \rightarrow 0$: in words, there is a link between the two merged parts with capacity is independent of η in $\mathcal{H}$. Since there are finitely-many merged parts, this implies that there is a constant $\zeta > 0$ (dependent only on the graph $\mathcal{H}$) such that for all $\eta \geq 0$,

$$f(\eta; \mathcal{P}) \geq \zeta \cdot f(\eta; \tilde{\mathcal{P}}). \quad (59)$$

We may now focus on the partition $\mathcal{P}$ only. We will later use Eq. (59) to translate the obtained guarantee to $f(\eta; \tilde{\mathcal{P}})$. Precisely, we aim to apply Theorem B.2 with the partition $\mathcal{P}$, which now satisfies Eq. (30). We follow similar arguments as in the proof of Theorem 5.3. In particular, we define the function $\phi^{(s)} : \mathbb{R}^{I_s} \rightarrow \mathbb{R}$ as in Eq. (53). We consider the allocation function $\mathbf{p}(\mathbf{u})$ that allocates uniformly on $\arg \max_{j \in [n]} u_j$ and let $\mathbf{U}$ be the utility vector realized by $\mathbf{p}$.

We next introduce the parameters used. Let

$$C_2 := 16n^2 \bar{v}^2 \left(1 + \frac{\max_{i \in [n]} \mathbb{E}[u_i]}{\min_{s \in [q], i \in I_s} (\mathbb{E}[u_i] \mathbb{P}(Z_{I_s} = Z > 0) - U_i)} \right)^2 \quad (60)$$

and $r_0 = \min_{s \in [q], i \in I_s} (\mathbb{E}[u_i] \mathbb{P}(Z_{I_s} = Z > 0) - U_i)/2$. Next, for convenience, we pose $c_0 = \frac{1}{16n^5}$ and $c' = \frac{1}{8n^2}$. Let $C_r := c' \sqrt{\frac{C_2}{2}}$, $C_\eta = 4c_0 c' C_r$, and $C_\delta = \frac{4C_r}{c'^2} = \frac{2C_2}{C_r}$. Last, we pose $f_0 = f(c'/2; \mathcal{P})$. We define

$$\eta(\gamma) := \inf \{1\} \cup \left\{ \eta \in \left[\frac{C_r}{\min(r_0, c_0 f_0/4)} \sqrt{1-\gamma}, 1 \right] : \forall \eta' \in [\eta, 1], f(\eta'; \mathcal{P}) \geq C_\eta \sqrt{1-\gamma} \frac{\eta'}{\eta} \right\}.$$

First, if $\eta(\gamma) \geq \frac{c'}{2\sqrt{2}} > 0$, by Theorem D.2 we have $\text{SWG}(\gamma) \leq C_1 \sqrt{1-\gamma}$, for the corresponding constant C_1 . Hence ensuring that $C \geq \frac{2\sqrt{2}C_1}{c'}$ gives the desired result. We now suppose that we have $\eta(\gamma) < \frac{c'}{2\sqrt{2}}$. Note that f is right-continuous, hence we have $f(\eta'; \mathcal{P}) \geq C_\eta \sqrt{1-\gamma} \frac{\eta'}{\eta(\gamma)}$ for all $\eta' \in [\eta(\gamma), 1]$. In the following, we check that the assumptions of Theorem B.2 are satisfied for $\gamma \in [1/2, 1)$ with the parameters:

$$r = C_r \frac{\sqrt{1-\gamma}}{\eta(\gamma)}, \quad \delta = C_\delta \sqrt{1-\gamma} \eta(\gamma), \quad \text{and} \quad \mathbf{x} = \left(U_i - (r + 2\delta) \frac{\mathbb{1}_{i \in I_s}}{\|\mathbf{1}_{I_s}\|} \right)_{i \in [n]}.$$

First note that by assumption we have $\eta(\delta) \leq \frac{c'}{2\sqrt{2}} = \sqrt{\frac{C_r}{2C_\delta}}$ so that $r \geq 2\delta$. In particular, since $\gamma \geq 1/2$ the constraint $\frac{r\gamma}{1-\gamma} \geq \delta$ from Eq. (13) is satisfied.

We now fix $s \in [q]$ and let $\mathcal{H}_s := \{\boldsymbol{\delta} \in \mathbb{R}^{I_s} : \mathbf{1}_{I_s}^\top \boldsymbol{\delta} = 0\}$. From the proof of Theorem D.4 (see Eq. (58)) for any $\boldsymbol{\beta} \in \mathcal{H}_s$ with $\|\boldsymbol{\beta}\| \leq 1/2$, we have

$$\begin{aligned} \phi^{(s)}(\mathbf{1}_{I_s} + \boldsymbol{\beta}) - \phi^{(s)}(\mathbf{1}_{I_s}) - \boldsymbol{\beta}^\top \mathbf{U}_{I_s} &\geq n^2 c_0 \mathbb{E} \left[Z_B \mathbb{1}_{Z_B \leq Z_{I_s \setminus B} = Z \leq (1+c'\|\boldsymbol{\beta}\|)Z_B} \right] \|\boldsymbol{\beta}\| \\ &\geq n^2 c_0 \tilde{f}(c'\|\boldsymbol{\beta}\|; \mathcal{P}) \|\boldsymbol{\beta}\| \stackrel{(i)}{\geq} c_0 f(c'\|\boldsymbol{\beta}\|; \mathcal{P}) \|\boldsymbol{\beta}\|, \end{aligned}$$

where in (i) we used Theorem D.5. Because the function $\boldsymbol{\beta} \mapsto \phi^{(s)}(\mathbf{1}_{I_s} + \boldsymbol{\beta}) - \phi^{(s)}(\mathbf{1}_{I_s}) - \boldsymbol{\beta}^\top \mathbf{U}_{I_s}$ is convex in γ and has value 0 for $\boldsymbol{\beta} = \mathbf{0}$, this implies that for all $\boldsymbol{\beta} \in \mathcal{H}_s$,

$$\phi^{(s)}(\mathbf{1}_{I_s} + \boldsymbol{\beta}) - \phi^{(s)}(\mathbf{1}_{I_s}) - \boldsymbol{\beta}^\top \mathbf{U}_{I_s} \geq c_0 \min(f(c'\|\boldsymbol{\beta}\|; \mathcal{P}), f_0) \cdot \|\boldsymbol{\beta}\|.$$

As a result, for any $\boldsymbol{\beta} \in \mathcal{H}_s$,

$$\begin{aligned} A_s(\boldsymbol{\beta}) &:= \phi^{(s)}(\mathbf{1}_{I_s} + \boldsymbol{\beta}) - (\mathbf{1}_{I_s} + \boldsymbol{\beta})^\top \mathbf{x}_{I_s} - (r + \delta) \|\mathbf{1}_{I_s} + \boldsymbol{\beta}\| \\ &\stackrel{(i)}{\geq} c_0 \min(f(c'\|\boldsymbol{\beta}\|; \mathcal{P}), f_s) \|\boldsymbol{\beta}\| - (r + \delta) (\|\mathbf{1}_{I_s} + \boldsymbol{\beta}\| - \|\mathbf{1}_{I_s}\|) + \delta \|\mathbf{1}_{I_s}\| \\ &\stackrel{(ii)}{\geq} c_0 \min(f(c'\|\boldsymbol{\beta}\|; \mathcal{P}), f_s) \|\boldsymbol{\beta}\| - \frac{4r}{\|\mathbf{1}_{I_s}\|} \min(\|\boldsymbol{\beta}\|^2, \|\mathbf{1}_{I_s}\| \|\boldsymbol{\beta}\|) + \delta \|\mathbf{1}_{I_s}\|. \end{aligned}$$

In (i) we used $\phi^{(s)}(\mathbf{1}_{I_s}) = \mathbf{1}_{I_s}^\top \mathbf{U}_{I_s}$ (see proof of Theorem D.4), $\mathbf{1}_{I_s}^\top \boldsymbol{\beta} = 0$, and the previous inequality. In (ii) we used $r \geq 2\delta$ and $\mathbf{1}_{I_s}^\top \boldsymbol{\beta} = 0$. We now distinguish between three cases. First suppose that $\|\boldsymbol{\beta}\| \leq \eta(\gamma)/c'$. Then,

$$A_s(\boldsymbol{\beta}) \geq \delta \|\mathbf{1}_{I_s}\| - \frac{4r}{c'^2 \|\mathbf{1}_{I_s}\|} \eta(\gamma)^2 = \sqrt{1-\gamma} \eta(\gamma) \left(C_\delta \|\mathbf{1}_{I_s}\| - \frac{4C_r}{c'^2 \|\mathbf{1}_{I_s}\|} \right) \geq 0.$$

Next suppose that $\eta(\gamma)/c' \leq \|\boldsymbol{\beta}\| \leq 1/2$. Note that by assumption, we have $\eta(\gamma)/c' \leq 1/2$. Therefore, since $f(\cdot; \mathcal{P})$ is non-decreasing we have $f(c'\|\boldsymbol{\beta}\|; \mathcal{P}) \leq f_0$. Further, since $c'\|\boldsymbol{\beta}\| \geq \eta(\gamma)$, we also have $f(c'\|\boldsymbol{\beta}\|) \geq C_\eta \sqrt{1-\gamma} c' \|\boldsymbol{\beta}\| / \eta(\gamma)$. Hence,

$$A_s(\boldsymbol{\beta}) \geq c f(c'\|\boldsymbol{\beta}\|) \|\boldsymbol{\beta}\| - \frac{4r}{\|\mathbf{1}_{I_s}\|} \|\boldsymbol{\beta}\|^2 \geq \left(c_0 c' C_\eta - \frac{4C_r}{\|\mathbf{1}_{I_s}\|} \right) \frac{\sqrt{1-\gamma}}{\eta(\gamma)} \|\boldsymbol{\beta}\|^2 \geq 0.$$

In the last inequality, we used the definition of C_η . Last, for $\|\boldsymbol{\beta}\| \geq 1/2$, we obtain $A_s(\boldsymbol{\beta}) \geq (c_0 f_0 - 4r) \|\boldsymbol{\beta}\|$. Now by construction of $\eta(\gamma)$, we have

$$r = C_r \frac{\sqrt{1-\gamma}}{\eta(\gamma)} \leq \min(r_0, c f_0 / 4).$$

This therefore shows that for all $\boldsymbol{\beta} \in \mathcal{H}_s$ we have $A_s(\boldsymbol{\beta}) \geq 0$. Finally, we check that the second constraint from Eq. (13) is satisfied. By construction and because $\gamma \geq 1/2$,

$$\delta = 2C_2 \frac{1-\gamma}{r} \geq C_2 \frac{1-\gamma}{\gamma r}.$$

Further, since $r \leq r_0$, we can check that

$$4n^2 \bar{v}^2 \left(1 + \frac{\max_{i \in [n]} \mathbb{E}[u_i]}{\min_{s \in [q], i \in I_s} (\mathbb{E}[u_i] \mathbb{P}(Z_{I_s} = Z > 0) - x_i - r)} \right)^2 \leq C_2.$$

Finally, we have checked that all assumptions to apply Theorem B.2 are satisfied. Therefore,

$$\mathbf{x} + \{0\}^{[n] \setminus I} \otimes \bigotimes_{s \in [q]} r \frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|} \in \mathcal{U}_\gamma.$$

As a result,

$$\text{SWG}(\gamma) = \max_{\mathbf{x} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{x} - \max_{\mathbf{x} \in \mathcal{U}_\gamma} \mathbf{1}^\top \mathbf{x} \leq 2\delta \sum_{s \in [q]} \|\mathbf{1}_{I_s}\| \leq n\sqrt{n}C_\delta \sqrt{1-\gamma}\eta(\gamma).$$

Ensuring that $C \geq n\sqrt{n}C_\delta \max\left(1, \frac{C_r}{\min(r_0, c_0 f_0/4)}\right)$ and up to multiplying C_η by a factor $1/\zeta$, the previous equation together with Eq. (59) this implies the desired result. Note that the term $C(1-\gamma)$ takes care of the case when $\eta(\gamma) = \frac{C_r}{\min(r_0, c_0 f_0/4)} \sqrt{1-\gamma}$.

Proof of the second claim. For $i \in I$, we define the functions g_i^+ and g_i^- on $\mathbb{R}_+$ via

$$g_i^+(\eta) := \mathbb{E}[u_i \mathbb{1}_{u_i \leq Z_{-i} \leq (1+\eta)u_i}] \quad \text{and} \quad g_i^-(\eta) := \mathbb{E}[u_i \mathbb{1}_{Z_{-i} < u_i \leq (1+\eta)Z_{-i}}].$$

Last, we define $g_i(\eta) = g_i^+(\eta) + g_i^-(\eta)$. Let $\tilde{c}_\eta = \min\{1\} \cup \{\mathbb{E}[u_i \mathbb{1}_{u_i = Z_{-i}}]/2 : i \in I, \mathbb{P}(u_i = Z_{-i} > 0) > 0\} > 0$. Ensuring $c_\eta \leq \tilde{c}_\eta$ implies that for any $i \in I$ such that $\mathbb{P}(u_i = Z_{-i} > 0) > 0$, we have (for any $\eta_0 > 0$)

$$\sup\{0\} \cup \{\eta \in [0, \eta_0] : g_i(\eta) \leq c_\eta \sqrt{1-\gamma}\} \leq \sup\{0\} \cup \{\eta \in [0, \eta_0] : g_i(\eta) \leq \tilde{c}_\eta\} = 0.$$

Hence, the desired inequality is trivial in that case. In the rest of the proof, we therefore focus on an index $i \in I$ with $\mathbb{P}(u_i = Z_{-i} > 0) = 0$. In particular, we necessarily have $i \in \tilde{I}$. We now introduce the function

$$\phi(\mathbf{y}) := \max_{\mathbf{x} \in \mathcal{U}^*} \mathbf{y}^\top (x_i, \mathbf{1}_{-i}^\top \mathbf{x}_{-i}) = \mathbb{E}[\max(y_i u_i, y_{-i} Z_{-i})], \quad \mathbf{y} \in \mathbb{R}^2 \setminus \{\mathbf{0}\}.$$

We fix $\mathbf{d} := (1, 1)$ and $\mathbf{f} := (1, -1)$ so that $\mathbf{d}^\top \mathbf{f} = 0$. Let $\mathbf{U}$ be an optimal vector and let $\mathbf{p}(\cdot)$ be an optimal allocation that realizes $\mathbf{U}$. In the 2-agent game it induces the utility vector $\mathbf{V} = (U_i, \mathbf{1}_{-i}^\top \mathbf{U}_{-i})$. We denote $\hat{i} \sim \mathbf{p}(\mathbf{u})$ the agent that receives the allocation. Then, for any $\eta \geq 0$, using the fact that $\phi(\mathbf{d}) = \mathbf{d}^\top \mathbf{V}$ and the same arguments as in the proof of Theorem D.4, we have

$$\begin{aligned} \phi(\mathbf{d} + \eta \mathbf{f}) - \phi(\mathbf{d}) - \eta \mathbf{f}^\perp \mathbf{V} &= \mathbb{E} \left[\max((1+\eta)u_i, (1-\eta)Z_{-i}) - (1+\eta)u_i \mathbb{1}_{\hat{i}=i} - (1-\eta)Z_{-i} \mathbb{1}_{\hat{i} \neq i} \right] \\ &= \mathbb{E} \left[((1+\eta)u_i - (1-\eta)Z_{-i}) \mathbb{1}_{\hat{i} \neq i} \mathbb{1}_{(1+\eta)u_i \geq (1-\eta)Z_{-i}} \right] \\ &\stackrel{(i)}{=} \mathbb{E} \left[((1+\eta)u_i - (1-\eta)Z_{-i}) \mathbb{1}_{u_i < Z_{-i} \leq \frac{1+\eta}{1-\eta} u_i} \right]. \end{aligned}$$

In (i) we used the fact that $\mathbb{P}(u_i = Z_{-i} > 0) = 0$ hence almost surely, $\mathbf{p}$ only allocates to i if $u_i > Z_{-i}$ or if $u_i = Z_{-i} = 0$ in which case $u_i = 0$ and the contribution is null. For $\eta \geq 0$ sufficiently small (independently of γ), we have $\frac{1+\eta}{1-\eta} \leq 1 + 3\eta$. Then,

$$\phi(\mathbf{d} + \eta \mathbf{f}) - \phi(\mathbf{d}) - \eta \mathbf{f}^\perp \mathbf{V} \leq 2\eta \mathbb{E} [u_i \mathbb{1}_{u_i < Z_{-i} \leq (1+3\eta)u_i}] = 2\eta g_i^+(3\eta).$$

In the last inequality we used the fact that $\mathbb{P}(u_i = Z_{-i} > 0) = 0$. Similarly, for $\eta \leq 0$, with $|\eta|$ sufficiently small, we have

$$\begin{aligned}\phi(\mathbf{d} + \eta \mathbf{f}) - \phi(\mathbf{d}) - \eta \mathbf{f}^\top \mathbf{V} &= \mathbb{E} \left[((1 - \eta)Z_{-i} - (1 + \eta)u_i) \mathbb{1}_{\hat{i}=i} \mathbb{1}_{(1+\eta)u_i \leq (1-\eta)Z_{-i}} \right] \\ &= \mathbb{E} \left[((1 - \eta)Z_{-i} - (1 + \eta)u_i) \mathbb{1}_{Z_{-i} < u_i \leq \frac{1-\eta}{1+\eta} Z_{-i}} \right] \\ &\leq 2|\eta| \mathbb{E} [u_i \mathbb{1}_{Z_{-i} < u_i \leq (1+3|\eta|)Z_{-i}}] = 2|\eta| g_i^-(3|\eta|).\end{aligned}$$

We next introduce $c_1, c_2 > 0$ the constants guaranteed from Theorem C.3. We now pose $c_\eta = \min\left(\tilde{c}_\eta, \frac{\sqrt{3c_2}}{48}\right)$, $c_r = 3\sqrt{3c_2}$, and fix $0 < \eta_0 \leq \min(1, \sqrt{2c_r^2/(3c_2)}, c_1/(\frac{c_2}{\sqrt{2c_r}} + \frac{c_r}{9\sqrt{2}}))$ sufficiently small such that all the previous estimates hold for $|\eta| \leq \eta_0$. Now let $\gamma \in [9/10, 1)$ and define

$$\eta'_i(\gamma) := \sup \{0\} \cup \left\{ \eta \in [0, \eta_0] : g_i(\eta) = g_i^+(\eta) + g_i^-(\eta) \leq c_\eta \sqrt{1 - \gamma} \right\}.$$

If $\eta'_i(\gamma) = 0$, we always have $\text{SWG}(\gamma) \geq c\sqrt{1 - \gamma} \cdot \eta'_i(\gamma)$ for any choice of constants. Suppose that is not the case and fix $\lambda_i(\gamma) \in (0, \eta'_i(\gamma))$. By construction, since g_i^+ and g_i^- are non-decreasing, we have $g_i(\lambda_i(\gamma)) \leq c_\eta \sqrt{1 - \gamma}$. We then pose

$$r := c_r \frac{\sqrt{1 - \gamma}}{\lambda_i(\gamma)}, \quad \xi := \frac{c_2}{2c_r} \sqrt{1 - \gamma} \lambda_i(\gamma), \quad \text{and} \quad \mathbf{x} := \mathbf{V} - (r + \xi) \frac{\mathbf{d}}{\|\mathbf{d}\|}.$$

For convenience, we introduce the function

$$\psi(\mathbf{y}) := \max_{\mathbf{z} \in B(\mathbf{x}, r)} \mathbf{y}^\top \mathbf{z} = \mathbf{y}^\top \mathbf{x} + r \|\mathbf{y}\|$$

We then compute for any η ,

$$\psi(\mathbf{d} + \eta \mathbf{f}) - \phi(\mathbf{d}) - \eta \mathbf{f}^\top \mathbf{V} = r(\|\mathbf{d} + \eta \mathbf{f}\| - \|\mathbf{d}\|) - \xi \|\mathbf{d}\| = r \|\mathbf{d}\| \left(\sqrt{1 + \eta^2} - 1 \right) - \xi \|\mathbf{d}\|$$

In the first equality, we used the fact that $\mathbf{U}$ is optimal, hence $\mathbf{d}^\top \mathbf{V} = \phi(\mathbf{d})$. Next, using $1 + x/2 \geq \sqrt{1 + x} \geq 1 + x/3$ for all $x \in [0, 1]$, we obtain that for η with $|\eta| \leq 1$,

$$\frac{1}{3} r \|\mathbf{d}\| \eta^2 - \xi \|\mathbf{d}\| \leq \psi(\mathbf{d} + \eta \mathbf{f}) - \phi(\mathbf{d}) - \eta \mathbf{f}^\top \mathbf{V} \leq \frac{1}{2} r \|\mathbf{d}\| \eta^2 - \xi \|\mathbf{d}\|.$$

Putting this with the previous equations, for any η with $|\eta| \leq \eta_0$, we obtained

$$\phi(\mathbf{d} + \eta \mathbf{f}) - \psi(\mathbf{d} + \eta \mathbf{f}) \leq |\eta| \left(2g_i(3|\eta|) - \frac{\|\mathbf{d}\|}{3} |\eta| r \right) + \xi \|\mathbf{d}\|.$$

In particular, for $\eta \in \{\pm \lambda_i(\gamma)/3\}$, we have

$$\begin{aligned}\phi(\mathbf{d} + \eta \mathbf{f}) - \psi(\mathbf{d} + \eta \mathbf{f}) &\leq \frac{\lambda_i(\gamma)}{3} \left(2c_\eta \sqrt{1 - \gamma} - \frac{c_r \|\mathbf{d}\|}{9} \sqrt{1 - \gamma} \right) + \frac{c_2}{2c_r} \|\mathbf{d}\| \lambda_i(\gamma) \sqrt{1 - \gamma} \\ &\stackrel{(i)}{=} \lambda_i(\gamma) \sqrt{1 - \gamma} \left(\frac{2c_\eta}{3} - \frac{\sqrt{c_2} \|\mathbf{d}\|}{6\sqrt{3}} \right) \stackrel{(ii)}{=} -\lambda_i(\gamma) \sqrt{1 - \gamma} \frac{c_\eta}{2} < 0.\end{aligned}$$

In (i) we used the definition of c_r and in (ii) we used the definition of c_η and $\|\mathbf{d}\| = \sqrt{2}$. As a result, $\mathbf{x} + r \frac{\mathbf{d} + \eta \mathbf{f}}{\|\mathbf{d} + \eta \mathbf{f}\|} \notin \{(u_i, \mathbf{1}_{-i}^\top \mathbf{u}_{-i}), \mathbf{u} \in \mathcal{U}^*\}$ for $\eta \in \{\pm \lambda_i(\gamma)/3\}$. We can then define $\mathcal{C}'$ as the union of all connected components of $\mathcal{U}^* \setminus B^\circ(\mathbf{x}, r; i)$ that are included within

$$\left\{ \mathbf{u} : \mathbf{d}^\top (u_i, \mathbf{1}_{-i}^\top \mathbf{u}_{-i}) > \mathbf{d}^\top \left(\mathbf{x} + r \frac{\mathbf{d} + \eta \mathbf{f}}{\|\mathbf{d} + \eta \mathbf{f}\|} \right), \eta = \frac{\lambda_i(\gamma)}{3} \right\}.$$

Note that for $|\eta|$ sufficiently small, we have

$$\mathbf{d}^\top \left(\mathbf{x} + r \frac{\mathbf{d} + \eta \mathbf{f}}{\|\mathbf{d} + \eta \mathbf{f}\|} \right) = \max_{\mathbf{z} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{z} - \xi \|\mathbf{d}\| - r \|\mathbf{d}\| \left(1 - \frac{1}{\sqrt{1 + \eta^2}} \right).$$

Now by construction of η_0 , for $|\eta| = \lambda_i(\gamma)3$, since $\lambda_i(\gamma) < \eta_i(\gamma) \leq \eta_0$, we have

$$\xi \|\mathbf{d}\| + r \|\mathbf{d}\| \left(1 - \frac{1}{\sqrt{1 + \eta^2}} \right) \leq \xi \|\mathbf{d}\| + \frac{1}{2} r \|\mathbf{d}\| \eta^2 \leq \frac{c_2}{2c_r} \|\mathbf{d}\| \lambda_i(\gamma) + \frac{c_r}{18} \|\mathbf{d}\| \lambda_i(\gamma) \leq c_1.$$

This proves that $\mathcal{C}' \subseteq \{\mathbf{u} : \mathbf{1}^\top \mathbf{u} \geq \max_{\mathbf{z} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{z} - c_1\}$. Pose $\delta = c_2 \frac{1-\gamma}{\gamma r} = \frac{2}{\gamma} \xi \geq 2\xi$. Since $\gamma \geq 9/10$,

$$\frac{\delta}{r} \leq \frac{3\xi}{r} = \frac{3c_2}{2c_r^2} \lambda_i(\gamma)^2 \leq \frac{3c_2}{2c_r^2} \eta_0^2 \leq 1.$$

That is, $\delta \leq r$. We next prove $\mathcal{C}' \subseteq B(\mathbf{x}, r + \delta; i)$. To do so, it suffices to check that

$$\phi(\mathbf{d} + \eta \mathbf{f}) \leq (\mathbf{d} + \eta \mathbf{f})^\top \mathbf{x} + (r + \delta) \|\mathbf{d} + \eta \mathbf{f}\|, \quad \eta \in \left[-\frac{\lambda_i(\gamma)}{3}, \frac{\lambda_i(\gamma)}{3} \right].$$

Using the previous computations, for any η such that $|\eta| \leq \lambda_i(\gamma)/3$, we have

$$\begin{aligned} & \phi(\mathbf{d} + \eta \mathbf{f}) - (\mathbf{d} + \eta \mathbf{f})^\top \mathbf{x} - (r + \delta) \|\mathbf{d} + \eta \mathbf{f}\| \\ & \leq 2|\eta|g_i(3|\eta|) - \frac{1}{3}(r + \delta) \|\mathbf{d}\| \eta^2 - (\delta - \xi) \|\mathbf{d}\| \\ & \stackrel{(i)}{\leq} \frac{2c_\eta \lambda_i(\gamma)}{3} \sqrt{1 - \gamma} - \xi \|\mathbf{d}\| \stackrel{(ii)}{\leq} 0, \end{aligned}$$

where in (i) we used $|\eta| \leq \lambda_i(\gamma)/3$ and $\delta \geq 2\xi$, and in (ii) we used the definition of c_η and ξ . This proves $\mathcal{C}' \subseteq B(\mathbf{x}, r + \delta; i)$. We have now checked all the assumptions for Theorem C.3 which implies

$$\mathcal{U}_\gamma \cap \left(\mathcal{C}' \cap B^\circ \left(\mathbf{x}, r + \frac{1-\gamma}{\gamma} \delta; i \right) \right) = \emptyset.$$

Next, note that by construction, $\mathbf{U} \in \mathcal{C}$ and $(U_i, mb1_{-i}^\top \mathbf{U}_{-i}) = \mathbf{x} + (r + \xi) \frac{\mathbf{d}}{\|\mathbf{d}\|}$. Hence, by Theorem C.3,

$$\max_{\mathbf{y} \in \mathcal{U}_\gamma} \mathbf{1}^\top \mathbf{y} < \mathbf{d}^\top \mathbf{x} - \left(r + \frac{1-\gamma}{\gamma} \delta \right) \|\mathbf{d}\| = \max_{\mathbf{y} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{y} - \left(\xi - \frac{1-\gamma}{\gamma} \delta \right) \|\mathbf{d}\|$$

Hence, for $1 - \gamma \leq 1/10$, we obtain

$$\text{SWG}(\gamma) = \max_{\mathbf{y} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{y} - \max_{\mathbf{y} \in \mathcal{U}_\gamma} \mathbf{1}^\top \mathbf{y} \geq \frac{\xi}{2} \|\mathbf{d}\| = \frac{c_2 \|\mathbf{d}\|}{2c_r} \sqrt{1 - \gamma} \lambda_i(\gamma).$$

Because this holds for all $\lambda_i(\gamma) \in (0, \eta'_i(\gamma))$, this also holds for $\eta'_i(\gamma)$.

In summary, we proved that for each $i \in I$, there exists constants $c_\eta(i), \eta_0(i), c(i) > 0$ depending only on the utility distributions such that for $1 - \gamma \leq 1/10$,

$$\text{SWG}(\gamma) \geq c(i) \sqrt{1 - \gamma} \cdot \sup \{0\} \cup \{\eta \in [0, \eta_0(i)] : g_i(\eta) \leq c_\eta(i) \sqrt{1 - \gamma}\}$$

We take c_η, η_0, c to be the minimum of these quantities for all $i \in I$, which shows that the above equation holds by replacing g_i with $\min_{i \in I} g_i$, and $c_\eta(i), \eta_0(i), c(i)$ with c_η, η_0, c . The last step of the proof is to relate the functions g and $\min_{i \in I} g_i$. For a given $\eta \in [0, \eta_0]$, there is $i \in I$ such that

$$\begin{aligned} g(\eta) &= \sum_{j \in I \setminus \{i\}} f(\eta; i, j) + f(\eta; j, i) \\ &\geq \mathbb{E} [u_i \mathbb{1}_{Z_{-i}=Z} \mathbb{1}_{Z_{-i} \in [u_i, (1+\eta)u_i]}] + \mathbb{E} [Z_{-i} \mathbb{1}_{u_i=Z} \mathbb{1}_{u_i \in [Z_{-i}, (1+\eta)Z_{-i}]}] \\ &\geq g_i^+(\eta) + \frac{g_i^-(\eta)}{1+\eta} \geq \frac{1}{2} \min_{i \in I} g_i(\eta), \end{aligned}$$

where in the last inequality we used $\eta_0 \leq 1/2$. Therefore, up to halving the parameter c_η , this proves the desired result. $\blacksquare$

We now derive some immediate consequences from Theorem 5.8. First, we prove Theorem 5.5 that gives necessary conditions for having rates $\mathcal{O}(1 - \gamma)$.

Proof of Theorem 5.5 If $g(\eta) \underset{\eta \rightarrow 0}{=} o(\eta)$, then $\sup \{0\} \cup \{\eta \in [0, \eta_0] : g(\eta) \leq c_\eta \sqrt{1 - \gamma}\} = \omega(\sqrt{1 - \gamma})$, which gives the desired lower bound $\omega(1 - \gamma)$ on the convergence rate. $\blacksquare$

Next, we can prove the 2-agent scenario characterization from Theorem 5.6.

Proof of Theorem 5.6 Since $I = \{1, 2\}$ and $f(\eta; 1, 2), f(\eta; 2, 1) = \Theta(\eta^p)$, we have $g(\eta) = \Theta(\eta^p)$. Also, using the partition $\mathcal{P} = \{\{1, 2\}\}$ which groups all agents together we also obtain $f(\eta; \mathcal{P}) = \Theta(\eta^p)$, where $p \geq 1$. A direct computation then shows that Theorem 5.8 implies $\text{SWG}(\eta) = \Theta((1 - \gamma)^{(1+\frac{1}{p})/2})$ as desired. $\blacksquare$

Next, we prove Theorem 7.1.

Proof of Theorem 7.1 We use Theorem 5.8. Suppose that we have density upper bounds. From the distributional assumptions, for any $i, j \in I$ we have $f(\eta; i, j) \underset{\eta \rightarrow 0}{=} \mathcal{O}(\eta)$. Hence, we also have $g(\eta) \underset{\eta \rightarrow 0}{=} \mathcal{O}(\eta)$. This implies that $g(\eta) \underset{\eta \rightarrow 0}{=} \mathcal{O}(\eta)$. Hence $\sup\{\eta \in [0, \eta_0] : g(\eta) \leq c_\eta \sqrt{1 - \gamma}\} = \Omega(\sqrt{1 - \gamma})$ as $\gamma \rightarrow 1$ and Theorem 5.8 provides the desired first claim. The second claim is taken directly from Theorem 5.4. $\blacksquare$

D.5 Remaining proofs for discrete distributions

In this section, we give the remaining proof of Theorem 7.3 from Section 7. The tools developed until now already imply the following.

Corollary D.6. *Suppose that all $\mathcal{D}_1, \dots, \mathcal{D}_n$ are discrete and that $\mathcal{U}_0$ does not already contain an optimal vector (see Theorem 2.4 for a characterization). If there exists $i \in I$ such that for all $j \neq i$, $\mathbb{P}(u_i = u_j = Z > 0) = 0$, then $\text{SWG}(\gamma) = \Theta(\sqrt{1 - \gamma})$. Otherwise, $\text{SWG}(\gamma) = \mathcal{O}(1 - \gamma)$.*

Proof We let $m := \max_{i \in [n]} \min(\text{Supp}(\mathcal{D}_i))$ be the minimum of the support of Z . We start with the first case. The assumption implies that for some $i \in I$, $\mathbb{P}(u_i = Z_{-i} > 0) = 0$. As a result, there exists $\eta_1 > 0$ such that

$$\mathbb{P}(u_i > 0, Z_{-i} \in [u_i/(1 + \eta_1), (1 + \eta_1)u_i]) = 0.$$

This implies that for all $\eta \in [0, \eta_1]$, we have $g_i(\eta) = 0$. Then, Theorem 5.8 implies

$$\text{SWG}(\gamma) \geq c\sqrt{1 - \gamma} \min(\eta_0, \eta_1).$$

The other inequality is already guaranteed by Theorem D.2, which ends the first claim.

We now consider the second alternative. In this case, the first bullet point from Theorem 5.4 is satisfied for all $i \in I$. This implies the second claim of the corollary. $\blacksquare$

We now use a gluing method to improve the convergence rate in the second case of Theorem D.6. In this case, the frontier of $\mathcal{U}^*$ at optimal points is flat, which will allow to fit local caps from larger radius balls close to this flat surface. Here, the form of the resulting frontier can be explicitly constructed, which allows for a simpler analysis. The intuition is however the same as that from the gluing approach described in Section 7 and gives the same convergence rates up to constants (potentially in exponents). The proof of Theorem 7.4 is detailed below. In particular, it implies the second claim of Theorem 7.3.

Proof of Theorem 7.4 Suppose that the distributions satisfy the assumptions that for all $i \in I$, there exists $j \neq i$ with $\mathbb{P}(u_i = u_j = Z > 0) > 0$. We denote the set of optimal utility vectors by

$$\mathcal{U}^*(\mathbf{1}) := \left\{ \mathbf{x} \in \mathcal{U}^*, \mathbf{1}^\top \mathbf{x} = \max_{\mathbf{y} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{y} \right\}.$$

We next construct the (undirected) graph $\mathcal{G}$ on I such that for any $i \neq j \in I$, i is connected to j if and only if $\mathbb{P}(u_i = u_j = Z > 0) > 0$. By assumption, for all $i \in I$ there exists $j \neq i$ such that $\mathbb{P}(u_i = u_j = Z > 0) > 0$. In particular, this implies $j \in I$. This shows that $\mathcal{G}$ does not have any isolated node. We then consider the partition $I = I_1 \sqcup \dots \sqcup I_q$ by connected components of $\mathcal{G}$. Since there is no isolated node, we have $|I_s| \geq 2$ for all $s \in [q]$. By construction, since two components are not connected, Eq. (30) is satisfied. Following the same arguments as in Theorem B.2, it suffices to focus on each set I_s independently for all $s \in [q]$. Precisely, we focus on the game where only agents I_s are present but the central planner cannot allocate on the event $\mathcal{E}_s := \{Z_{-I_s} \geq Z_{I_s}\}$. Suppose we showed that the region $\mathcal{R}_s \subseteq \mathbb{R}_+^{I_s}$ is achievable in this setting for all $s \in [q]$, then the proof of Theorem 4.1 implies that for any optimal vector $\mathbf{U}$, the region

$$\mathcal{R} := \mathbf{U}_{[n] \setminus I} \otimes \bigotimes_{s \in [q]} \mathcal{R}_s \quad (61)$$

is achievable in the original setting with $[n]$ agents: $\mathcal{R} \subseteq \mathcal{U}_\gamma$.

We therefore focus on a single connected component I_s for $s \in [n]$ from now. We fix any optimal vector $\mathbf{U}$ and let $\mathcal{H}_s = \{\mathbf{y} : \forall i \notin I_s, y_i = U_i\}$. For any $\mathbf{x} \in \mathbb{R}^{I_s}$ and $r > 0$, we also denote $B_{I_s}(\mathbf{x}, r)$ the corresponding ball on the space $\mathbb{R}^{I_s}$. We first show that $\mathcal{U}^*(\mathbf{1}) \cap \mathcal{H}_s$ has (full) dimension $|I_s| - 1$. Note that since I_s is connected, there exist values $m_1 > \dots > m_T > 0$ such that letting

$$J_t := \{i \in I_s : \mathbb{P}(u_i = m_t = Z) > 0\}, \quad t \in [T],$$

we have (1) $|J_t| \geq 2$ for $t \in [T]$ and (2) no set J_t is disjoint from the others, that is $J_t \cap \bigcup_{t' \neq t} J_{t'} \neq \emptyset$ for all $t \in [T]$. We now define the event $\mathcal{F}_t = \{\forall i \in J_t, u_i = m_t = Z\}$. By construction, $\mathbb{P}(\mathcal{F}_t) > 0$ and all events $\mathcal{F}_1, \dots, \mathcal{F}_T$ are disjoint. On each event $\mathcal{F}_t$, an optimal allocation can choose to allocate to any agent in J_t . Let $\mathbf{p}^{(0)}(\cdot)$ an optimal allocation such that $\mathbb{E}[u_i p_i^{(0)}(\mathbf{u})] = U_i$ for all $i \notin I_s$ and that for any $t \in [T]$, allocates uniformly among agents J_t on the event $\mathcal{F}_t$. Denote

$\mathbf{U}^{(0)}$ the utility vector realized by $\mathbf{p}^{(0)}$. The previous arguments imply

$$\mathcal{S} := \mathbf{U}^{(0)} + \sum_{t \in [T]} \mathbb{P}(\mathcal{F}_t) \left\{ \left(m_t \left(q_i - \frac{1}{|J_t|} \right) \mathbb{1}_{i \in J_t} \right)_{i \in [n]}, \mathbf{q} \in \Delta_{J_t} \right\} \subseteq \mathcal{U}^*(\mathbf{1}),$$

where the sum between sets are Minkowski sums. Here each sum corresponds to the freedom that an optimal allocation has on the event $\mathcal{F}_t$ to use any allocation within Δ_{J_t} . From the assumptions (1) and (2) on the sets $J_1, \dots, J_T$, with $\tilde{r}_0 = \min_{t \in [T]} \mathbb{P}(\mathcal{F}_t) m_t / n$, we have

$$\mathbf{U}^{(0)} + \left\{ (y_i \mathbb{1}_{i \in I_s})_{i \in [n]}, \mathbf{y} \in B_{I_s}(0, \tilde{r}_0), \mathbf{1}_{I_s}^\top \mathbf{y} = 0 \right\} \subseteq \mathcal{U}^*(\mathbf{1}) \cap \mathcal{H}_s. \quad (62)$$

We recall that for any $\mathbf{x} \in \mathcal{U}^*$, any $\mathbf{0} \leq \mathbf{y} \leq \mathbf{x}$ satisfies $\mathbf{y} \in \mathcal{U}^*$. This also holds in the setting where only agents I_s are present and there is no allocation on the event $\mathcal{E}_s$. We denote by $\mathcal{U}_s^*$ the corresponding achievable region. In the proof of Theorem B.2, we showed (Eq. (35)) that $\mathcal{U}_s^* = P_{I_s}(\mathcal{U}^* \cap \mathcal{H}_s)$, where $P_{I_s} : \mathbb{R}^n \rightarrow \mathbb{R}^{I_s}$ is the projection on the I_s coordinates. For convenience, we also write $\mathbf{U}_s := \mathbf{U}_{I_s}^{(0)}$. As a result, from Eq. (62) we can check that with $r_0 = \tilde{r}_0 / (8\|\mathbf{1}\|)$, for any $\mathbf{z} \in B_{I_s}(\mathbf{U}_s, r_0) \cap \{\mathbf{y} : \mathbf{1}_{I_s}^\top (\mathbf{y} - \mathbf{U}_s) = 0\} := S_s$,

$$B_{I_s} \left(\mathbf{z} - r_0 \frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|}, r_0 \right) \subseteq P_{I_s}(\mathcal{U}^* \cap \mathcal{H}_s) = \mathcal{U}_s^*.$$

Note that for all $i \in I_s$, we have $\mathbb{E}[u_i] \mathbb{P}(Z_{I_s} = Z > 0) - (\mathbf{z} - r_0 \frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|})_i - r_0 \geq \mathbb{P}(\mathcal{F}_t) m_t (1 - 1/|J_t|) - 2r_0 - r_0 \geq 8r_0 - 3r_0 \geq 2r_0$, where t is such that $i \in J_t$. Then, with

$$C = 4n^2 \bar{v}^2 \left(1 + \frac{\max_{i \in [n]} \mathbb{E}[u_i] + r_0}{r_0} \right)^2$$

and $\delta = 4C \frac{1-\gamma}{r_0}$, provided that $\delta \leq r_0/2$ then Eq. (13) is satisfied with $r = r_0 - \delta$ and δ :

$$\frac{r\gamma}{1-\gamma} \geq r \geq \frac{r_0}{2} \geq \delta = 4C \frac{1-\gamma}{r_0} \geq C \frac{1-\gamma}{\gamma r}.$$

Here, we used $\gamma \geq 1/2$ and $r_0 - \delta \geq r_0/2$. Let $C_1 := 5n\bar{v}$. The inequality $\delta \leq r_0/4$ holds whenever $\gamma \geq \gamma_1 := \max(1 - \min(\frac{r_0^2}{16C}, (\frac{r_0}{8C\|\mathbf{1}\|})^2, \frac{r_0}{4C_1}), 1/2)$. Then, for $\gamma \in [\gamma_1, 1)$, the proof of Theorem B.2 shows that

$$B_{I_s} \left(\mathbf{z} - r_0 \frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|}, r_0 - \delta \right) \subseteq \mathcal{U}_{s,\gamma}, \quad (63)$$

where $\mathcal{U}_{s,\gamma}$ is the achievable region for the setting with only agents I_s and in which the central planner cannot allocate on the event $\mathcal{E}_s$, with discount factor γ . We define the following function $f : \mathbb{R} \rightarrow \mathbb{R}$,

$$f(r) := \delta \frac{\cosh\left(\frac{c_f r}{\sqrt{1-\gamma}}\right)}{\cosh\left(\frac{c_f r_0}{\sqrt{1-\gamma}}\right)} = 4C \frac{1-\gamma}{r_0} \frac{\cosh\left(\frac{c_f r}{\sqrt{1-\gamma}}\right)}{\cosh\left(\frac{c_f r_0}{\sqrt{1-\gamma}}\right)}, \quad r \geq 0.$$

where $c_f = \min(\frac{1}{4C_1}, 1)$. We then pose $\gamma_0 = \max(\gamma_1, 1 - \min(r_0^2 c_f^2, (\frac{c_f C}{C_1})^2))$ and consider $\gamma \in [\gamma_0, 1)$. We focus on the following region

$$\mathcal{R}_s = \left\{ \mathbf{y} \in \mathbb{R}^{I_s} : \mathbf{0} \leq \mathbf{y} \leq \mathbf{z} - f(r) \frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|}, \|\mathbf{z} - \mathbf{U}_{I_s}\| = r, \mathbf{z} \in S_s \right\}$$

and now construct an allocation and promised utility function on $\mathcal{R}_s$. For any $\mathbf{z} \in S_s$, we denote by $\mathbf{p}(\cdot; \mathbf{z})$ an allocation that realizes $\mathbf{z}$ in the full information setting (within the game when only agents I_s are present). We start by specifying the strategy on the boundary points of $\mathcal{R}_s$. For $\mathbf{z} \in S_s$, let $r = \|\mathbf{z} - \mathbf{U}_s\|$ and $\mathbf{U} = \mathbf{z} - f(r)\mathbf{1}_{I_s}/\|\mathbf{1}_{I_s}\|$. We pose

$$\mathbf{p}(\cdot | \mathbf{U}) := \mathbf{p}(\cdot; \mathbf{z}) \quad \text{and} \quad \boldsymbol{\alpha}(\mathbf{U}) := -\frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|} - f'(r)\frac{\mathbf{z} - \mathbf{U}_s}{\|\mathbf{z} - \mathbf{U}_s\|}$$

Note that $\boldsymbol{\alpha}(\mathbf{U})$ was selected (up to the sign) as the normal to the frontier manifold $\mathcal{M}_s := \{\mathbf{z}' - f(\|\mathbf{z}' - \mathbf{U}_s\|)\mathbf{1}_{I_s}/\|\mathbf{1}_{I_s}\|, \mathbf{z}' \in S_s\}$. We then extend the definition of the allocation and promised utility functions to all $\mathbf{U} \in \mathcal{R}_s$ as follows. Let $\mathbf{V} \in \mathcal{M}_s$ such that $\mathbf{0} \leq \mathbf{U} \leq \mathbf{V}$. We let

$$p_i(\cdot | \mathbf{U}) := \frac{U_i}{V_i}p_i(\cdot | \mathbf{V}) \quad \text{and} \quad W_i(\cdot | \mathbf{U}) := \frac{U_i}{V_i}W_i(\cdot | \mathbf{V}), \quad i \in I_s,$$

with the convention $0/0 = 0$ when $U_i = V_i = 0$. To prove that $\mathcal{R}_s \subseteq \mathcal{U}_{s,\gamma}$ it suffices to show that for all $\mathbf{U} \in \mathcal{R}_s$ and $\mathbf{v} \in [0, \bar{v}]^{I_s}$,

$$\mathbf{W}(\mathbf{v} | \mathbf{U}) \in \mathcal{R}_s \cup \bigcup_{\mathbf{z} \in S_s} \left\{ \mathbf{y} : \mathbf{0} \leq \mathbf{y} \leq \mathbf{x}, \mathbf{x} \in B_{I_s} \left(\mathbf{z} - r_0 \frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|}, r_0 - \delta \right) \right\}, \quad (64)$$

and Eq. (7) is satisfied for all $\mathbf{U} \in \mathcal{R}_s$. Here, we used Eq. (63) that already guarantees that the second term in the right-hand side of Eq. (64) is contained within $\mathcal{U}_\gamma$. By the linearity of the definitions of the allocation and promised utility function, it suffices to check that for all boundary utility vectors $\mathbf{U} \in \mathcal{M}_s$, Eq. (7) is satisfied (this is guaranteed by Eq. (11)) and to check that the following holds:

$$\mathbf{W}(\mathbf{v} | \mathbf{U}) \in \mathcal{R}_s \cup \bigcup_{\mathbf{z} \in S_s} B_{I_s} \left(\mathbf{z} - r_0 \frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|}, r_0 - \delta \right), \quad \mathbf{U} \in \mathcal{M}_s, \mathbf{v} \in [0, \bar{v}]^{I_s}, \quad (65)$$

Fix such a vector $\mathbf{U} \in \mathcal{M}_s$ and write $\mathbf{U} = \mathbf{z} - f(r)\frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|}$ where $r = \|\mathbf{z} - \mathbf{U}_s\|$. By the promise-keeping equality Eq. (2) (which is a consequence of the construction),

$$\bar{\mathbf{W}} := \mathbb{E}[\mathbf{W}(\mathbf{u} | \mathbf{U})] = \mathbf{U} - \frac{1-\gamma}{\gamma}f(r)\frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|}. \quad (66)$$

By construction, we also have $\boldsymbol{\alpha}(\mathbf{U})^\top (\mathbf{W}(\mathbf{v} | \mathbf{U}) - \bar{\mathbf{W}}) = 0$ for $\mathbf{v} \in [0, \bar{v}]^{I_s}$. Last, the proof of Theorem 3.1 (Eq. (28)) shows that

$$\|\mathbf{W}(\mathbf{v} | \mathbf{U}) - \bar{\mathbf{W}}\| \leq \frac{1-\gamma}{\gamma}\bar{v} \left(2n + \frac{\alpha_{i_1}(\mathbf{U})}{\alpha_{i_2}(\mathbf{U})} \right), \quad \mathbf{v} \in [0, \bar{v}]^{I_s},$$

where $\alpha_{i_1}(\mathbf{U})$ and $\alpha_{i_2}(\mathbf{U})$ are the first and second largest components in absolute value of $\boldsymbol{\alpha}(\mathbf{U})$. Note that from $\gamma \geq \gamma_1$ and the definition of γ_1 , we have

$$\left\| \boldsymbol{\alpha}(\mathbf{U}) - \frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|} \right\| = |f'(r)| = c_f \frac{\delta}{\sqrt{1-\gamma}} \frac{|\sinh(c_f r / \sqrt{1-\gamma})|}{\cosh(c_f r_0 / \sqrt{1-\gamma})} \leq \frac{1}{2\|\mathbf{1}\|}.$$

As a result, this shows that for all $i \in I_s$ we have $\alpha_i(\mathbf{U})\|\mathbf{1}_{I_s}\| \in [1/2, 3/2]$. Therefore, we obtained

$$\|\mathbf{W}(\mathbf{v} | \mathbf{U}) - \bar{\mathbf{W}}\| \leq \frac{1-\gamma}{\gamma}\bar{v}(2n + 3\|\mathbf{1}\|) \leq 2C_1(1-\gamma), \quad \mathbf{v} \in [0, \bar{v}]^{I_s}. \quad (67)$$

Now let $d(\mathbf{v}) := \|P_{\mathbf{1}_{I_s}^\perp}(\mathbf{W}(\mathbf{v} | \mathbf{U}) - \mathbf{U}_s)\| - \|P_{\mathbf{1}_{I_s}^\perp}(\bar{\mathbf{W}} - \mathbf{U}_s)\|$, where $P_{\mathbf{1}_{I_s}^\perp}$ denotes the projection onto the orthogonal space to $\mathbf{1}_{I_s}\mathbb{R}$. Because $\boldsymbol{\alpha}(\mathbf{U})^\top(\mathbf{W}(\mathbf{v} | \mathbf{U}) - \bar{\mathbf{W}}) = 0$, we can write

$$\mathbf{W}(\mathbf{v} | \mathbf{U}) - \bar{\mathbf{W}} = d(\mathbf{v}) \frac{\mathbf{z} - \mathbf{U}_s}{\|\mathbf{z} - \mathbf{U}_s\|} + d(\mathbf{v})f'(r) \frac{(-\mathbf{1}_{I_s})}{\|\mathbf{1}_{I_s}\|} + \mathbf{W}_0(\mathbf{v}),$$

where $\mathbf{W}_0(\mathbf{v}) \in \text{Span}(\mathbf{1}_{I_s}, \mathbf{z} - \mathbf{U}_s)^\perp$. Using Eq. (66) we have

$$\begin{aligned} \frac{-\mathbf{1}_{I_s}^\top}{\|\mathbf{1}_{I_s}\|}(\mathbf{W}(\mathbf{v} | \mathbf{U}) - \mathbf{U}_s) &= f(r) + f(r) \frac{1-\gamma}{\gamma} + \frac{-\mathbf{1}_{I_s}^\top}{\|\mathbf{1}_{I_s}\|}(\mathbf{W}(\mathbf{v} | \mathbf{U}) - \bar{\mathbf{W}}) \\ &= f(r) + f(r) \frac{1-\gamma}{\gamma} + d(\mathbf{v})f'(r). \end{aligned}$$

Using Taylor expansion's theorem, there exists $\tilde{r} \in [\min(r, r + d(\mathbf{v})), \max(r, r + d(\mathbf{v}))]$ such that

$$\begin{aligned} f(r) \frac{1-\gamma}{\gamma} &\geq \frac{-\mathbf{1}_{I_s}^\top}{\|\mathbf{1}_{I_s}\|}(\mathbf{W}(\mathbf{v} | \mathbf{U}) - \mathbf{U}_s) - f(r + d(\mathbf{v})) = f(r) \frac{1-\gamma}{\gamma} - \frac{f''(\tilde{r})}{2} d(\mathbf{v})^2 \\ &\stackrel{(i)}{=} f(r)(1-\gamma) - \frac{c_f^2 f(\tilde{r})}{2(1-\gamma)} d(\mathbf{v})^2 \\ &\stackrel{(ii)}{\geq} (1-\gamma) (f(r) - 2c_f^2 C_1^2 f(\tilde{r})) \\ &\stackrel{(iii)}{\geq} (1-\gamma) \left(f(r) - \frac{f(r + |d(\mathbf{v})|)}{2} \right) \end{aligned}$$

In (i) we used the definition the fact that $\cosh''(x) = \cosh(x)$ and in (ii) we used $|d(\mathbf{v})| \leq \|\mathbf{W}(\mathbf{v} | \mathbf{U}) - \bar{\mathbf{W}}\|$ together with Eq. (67). Last, in (iii) we used the definition of c_f together with the fact that $f(\tilde{r}) \leq f(r + |d(\mathbf{v})|)$. Now note that

$$f(r + |d(\mathbf{v})|) \leq f(r) e^{\frac{c_f |d(\mathbf{v})|}{\sqrt{1-\gamma}}} \leq f(r) e^{2c_f C_1 \sqrt{1-\gamma}} \leq 2f(r). \quad (68)$$

In summary, since $\gamma \geq 1/2$, we showed that

$$f(r + d(\mathbf{v})) \leq \frac{-\mathbf{1}_{I_s}^\top}{\|\mathbf{1}_{I_s}\|}(\mathbf{W}(\mathbf{v} | \mathbf{U}) - \mathbf{U}_s) \leq f(r + d(\mathbf{v})) + 2f(r)(1-\gamma). \quad (69)$$

First suppose that $|r + d(\mathbf{v})| \leq r_0$. We let $\mathbf{z}(\mathbf{v}) := \mathbf{U}_s + P_{\mathbf{1}_{I_s}^\perp}(\mathbf{W}(\mathbf{v} | \mathbf{U}) - \mathbf{U}_s) \in S_s$. The previous equation implies that $\mathbf{W}(\mathbf{v} | \mathbf{U}) \leq \mathbf{z}(\mathbf{v}) - f(\|\mathbf{z}(\mathbf{v}) - \mathbf{U}_s\|) \frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|}$. Also, $2f(r)(1-\gamma) \leq f(r_0) = \delta$, hence, we also obtain $\mathbf{0} \leq \mathbf{W}(\mathbf{v} | \mathbf{U})$. Together with the previous inequalities, this gives $\mathbf{W}(\mathbf{v} | \mathbf{U}) \in \mathcal{R}_s$.

We now treat the remaining case when $|r + d(\mathbf{v})| \geq r_0$. We first show $d(\mathbf{v}) \geq 0$. Otherwise, we have $2C_1(1-\gamma) \geq d(\mathbf{v}) \geq r_0$, which is absurd since $\gamma \geq \gamma_1$. For this boundary case, we let

$$\mathbf{z}(\mathbf{v}) := \mathbf{U}_s + r_0 \frac{P_{\mathbf{1}_{I_s}^\perp}(\mathbf{W}(\mathbf{v} | \mathbf{U}) - \mathbf{U}_s)}{\|P_{\mathbf{1}_{I_s}^\perp}(\mathbf{W}(\mathbf{v} | \mathbf{U}) - \mathbf{U}_s)\|},$$

and aim to show that

$$\mathbf{W}(\mathbf{v}, \mathbf{U}) \in B_{I_s} \left(\mathbf{z}(\mathbf{v}) - r_0 \frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|}, r_0 - \delta \right). \quad (70)$$

Note that $f(r_0) = \delta$ and by Eq. (69) there exists $\tilde{f} \in [0, 2f(r_0)(1 - \gamma)] \subseteq [0, f(r_0)]$ such that

$$\begin{aligned}
R(\mathbf{v}) &:= \left\| \mathbf{W}(\mathbf{v}, \mathbf{U}) - \left(\mathbf{z}(\mathbf{v}) - r_0 \frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|} \right) \right\|^2 = (r + d(\mathbf{v}) - r_0)^2 + (f(r + d(\mathbf{v})) + \tilde{f} - r_0)^2 \\
&\stackrel{(i)}{\leq} (r + d(\mathbf{v}) - r_0)^2 + (r_0 - f(r + d(\mathbf{v})))^2 \\
&= (r_0 - \delta)^2 + (r + d(\mathbf{v}) - r_0)^2 - (2r_0 + f(r_0) - f(r + d(\mathbf{v}))) (f(r + d(\mathbf{v})) - f(r_0)) \\
&\stackrel{(ii)}{\leq} (r_0 - \delta)^2 + (r + d(\mathbf{v}) - r_0)^2 - r_0(f(r + d(\mathbf{v})) - f(r_0)).
\end{aligned}$$

In (i) we used Eq. (68) to have $\delta = f(r_0) \leq f(r + d(\mathbf{v})) \leq 2f(r) \leq 2\delta$ which implies in particular $r_0 - f(r + d(\mathbf{v})) - \tilde{f} \geq r_0 - 3\delta \geq r_0/4$ since $\delta \leq r_0/4$, and in (ii) we again used $f(r + d(\mathbf{v})) \leq 2\delta$ and $r_0 \geq \delta$. Now because f is convex, we have $f(r + d(\mathbf{v})) - f(r_0) \geq f'(r_0)(r + d(\mathbf{v}) - r_0)$. Further, by definition of $\gamma \geq \gamma_0$, we have $\frac{c_f r_0}{\sqrt{1-\gamma}} \geq 1$. As a result, we obtain $f'(r_0) \geq \frac{c_f \delta}{\sqrt{1-\gamma}} \cdot \frac{e-1/e}{e+1/e} \geq \frac{c_f \delta}{2\sqrt{1-\gamma}}$. Altogether, this implies

$$\begin{aligned}
R(\mathbf{v}) - (r_0 - \delta)^2 &\leq (r + d(\mathbf{v}) - r_0)(r + d(\mathbf{v}) - r_0 - r_0 f'(r_0)) \\
&\stackrel{(i)}{\leq} (r + d(\mathbf{v}) - r_0)(2C_1(1 - \gamma) - 2c_f C \sqrt{1 - \gamma}) \stackrel{(ii)}{\leq} 0
\end{aligned}$$

where in (i) we used $r + d(\mathbf{v}) - r_0 \leq d(\mathbf{v}) \leq 2C_1(1 - \gamma)$ and in (ii) we used the definition of γ_0 . This ends the proof of Eq. (70).

In both cases, we proved that Eq. (65) holds. Therefore, we have $\mathcal{R}_s \subseteq \mathcal{U}_{s, \gamma}$. This holds for all $s \in [q]$, hence the tensored region $\mathcal{R}$ defined in Eq. (61) is achievable in the original setting with $[n]$ agents: $\mathcal{R} \subseteq \mathcal{U}_\gamma$. In particular, with $\mathbf{U}_1 := \mathbf{U}_{[n] \setminus I} \otimes \bigotimes_{s \in [q]} \left(\mathbf{U}_s - f(0) \frac{\mathbf{1}_{I_s}}{\|\mathbf{1}_{I_s}\|} \right)$, we obtain

$$\text{SWG}(\gamma) = \max_{\mathbf{x} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{x} - \max_{\mathbf{x} \in \mathcal{U}_\gamma} \mathbf{1}^\top \mathbf{x} \leq f(0) \sum_{s \in [q]} \|\mathbf{1}_{I_s}\| \leq \frac{4n \|\mathbf{1}\| (1 - \gamma)}{r_0 \cosh\left(\frac{c_f r_0}{\sqrt{1-\gamma}}\right)} = \mathcal{O}\left((1 - \gamma) e^{-c_f r_0 / \sqrt{1-\gamma}}\right),$$

which ends the proof of the theorem. ■

E Extension to the finite-horizon setting

In this section, we show how the techniques developed in previous sections for the γ -discounted infinite-horizon setting can be used for the finite-horizon setting.

We start by proving the following finite-horizon counterpart to Fact 2.2 and 2.3 which allows to use the promised utility framework for the finite- T horizon setting also.

Fact E.1. *Fix $\mathcal{R} \subseteq \mathcal{U}^*$ and $T \geq 2$. Then $\mathcal{R} \subseteq \mathcal{V}_T$ if and only if there exists functions $\mathbf{p}(\cdot \mid \mathbf{U})$ and $\mathbf{W}(\cdot \mid \mathbf{U})$ for all $\mathbf{U} \in \mathcal{R}$ satisfying Eqs. (2) and (4) for $\gamma = \gamma(T)$, and instead of the original constraint Eq. (3) they satisfy*

$$\forall \mathbf{U} \in \mathcal{R}, \forall \mathbf{v} \in [0, \bar{v}]^n, \quad \mathbf{W}(\mathbf{v} \mid \mathbf{U}) \in \mathcal{V}_{T-1}. \quad (71)$$

Equivalently, $\mathcal{R} \subseteq \mathcal{V}_T$ if and only if the promised utility function satisfies Eq. (7) where the interim promised utility was defined via Eq. (6) with $\gamma = \gamma(T)$, and satisfies Eq. (71).

Proof Suppose that Eqs. (2) and (4) for $\gamma = \gamma(T)$, and Eq. (71) are satisfied. Then, for any $\mathbf{U} \in \mathcal{R}$, we can use the allocation function $p(\mathbf{v}(1) \mid \mathbf{U})$ for the first time $t = 1$ and reports $\mathbf{v}(1)$, then realize the utility vector $\mathbf{W}(\mathbf{v}(1) \mid \mathbf{U})$ since it belongs to the achievable region $\mathcal{V}_{T-1}$ by Eq. (71). This strategy is incentive-compatible at least for time $t = 1$ hence from Eq. (2), it realizes the vector $\mathbf{U}$, that is, $\mathbf{U} \in \mathcal{V}_T$. On the other hand, if $\mathbf{U}$ can be realized, fix an incentive-compatible strategy to realize $\mathbf{U}$. Let $p(\cdot \mid \mathbf{U})$ be the allocation function for the first time $t = 1$ and define $\mathbf{W}(\mathbf{v} \mid \mathbf{U}) = \frac{1}{T} \mathbb{E} \left[\sum_{t \geq 2} \tilde{u}_i(t) p_i(t) \mid \mathbf{v}(1) = \mathbf{v} \right]$ the remaining utility for times in $[2, T]$. We can easily check that these satisfy all the desired equations.

The second claim can be proved with the same arguments as in the proof of Fact 2.3. ■

We next prove Theorem 6.1 which shows that for $T \geq 2$, we have $\mathcal{V}_T \subseteq \mathcal{U}_{\gamma(T)}$, where $\gamma(T) = 1 - 1/T$.

Proof of Theorem 6.1 For $t \geq 2$, we let $\mathbf{p}^{(t)}(\cdot \mid \mathbf{U})$ and $\mathbf{W}^{(t)}(\cdot \mid \mathbf{U})$ be allocation and promised utility functions for $\mathbf{U} \in \mathcal{V}_t$ satisfying Eqs. (2) and (4) for $\gamma(t)$ and Eq. (71) for t .

For $T = 1$, we have directly $\mathcal{V}_1 = \mathcal{U}_0 = \mathcal{U}_{\gamma(1)}$. For $T \geq 2$, note that $\mathbf{p}^{(T)}$ and $\mathbf{W}^{(T)}$ already satisfy almost all the equations to show that the region $\mathcal{V}_T$ belongs to $\mathcal{U}_{\gamma(T)}$. The only difference is in Eq. (71) which implies that the promised utilities stay in $\mathcal{V}_{T-1}$. Precisely, to show $\mathcal{V}_T \subseteq \mathcal{U}_{\gamma(T)}$ it suffices to show that $\mathcal{V}_{T-1} \subseteq \mathcal{V}_T$. This is also immediate if $T = 2$ since $\mathcal{V}_1 = \mathcal{U}_0 \subseteq \mathcal{V}_t$ for all $t \geq 1$. We now suppose that $T \geq 3$ and that we showed $\mathcal{V}_{T-2} \subseteq \mathcal{V}_{T-1}$.

We introduce an allocation function and promised utility function on $\mathcal{V}_{T-1}$ via

$$\mathbf{p}(\cdot \mid \mathbf{U}) := \mathbf{p}^{(T-1)}(\cdot \mid \mathbf{U}) \quad \text{and} \quad \mathbf{W}(\cdot \mid \mathbf{U}) := \frac{T-2}{T-1} \mathbf{W}^{(T-1)}(\cdot \mid \mathbf{U}) + \frac{1}{T-1} \mathbf{U}, \quad \mathbf{U} \in \mathcal{V}_{T-1}.$$

We now check that this is a valid allocation strategy. Indeed, with $\mathbf{U} \in \mathcal{V}_{T-1}$, for all $i \in [n]$,

$$\begin{aligned} (1 - \gamma(T)) \mathbb{E}[u_i p_i(\mathbf{u})] + \gamma(T) \mathbb{E}[\mathbf{W}(\mathbf{u} \mid \mathbf{U})] \\ = \frac{\mathbf{U}}{T} + \frac{T-1}{T} \left((1 - \gamma(T-1)) \mathbb{E}[u_i p_i^{(T-1)}(\mathbf{u})] + \gamma(T-1) \mathbb{E}[\mathbf{W}^{(T-1)}(\mathbf{u} \mid \mathbf{U})] \right) = \mathbf{U}. \end{aligned}$$

Hence Eq. (2) holds. Here, in the last inequality we used the fact that $\mathbf{p}^{(T-1)}$ and $\mathbf{W}^{(T-1)}$ satisfy Eq. (2) for $\gamma = \gamma(T-1)$. Similarly, for any $i \in [n]$ and $u, v \in [0, \bar{v}]$,

$$\begin{aligned} (1 - \gamma(T)) u P_i(v \mid \mathbf{U}) + \gamma(T) W_i(v \mid \mathbf{U}) \\ = \frac{U_i}{T} + \frac{T-1}{T} \left((1 - \gamma(T-1)) u P_i^{(T-1)}(v \mid \mathbf{U}) + \gamma(T-1) W_i^{(T-1)}(v \mid \mathbf{U}) \right), \end{aligned}$$

hence the incentive-compatibility Eq. (4) also holds from that of $\mathbf{p}^{(T-1)}$ and $\mathbf{W}^{(T-1)}$. Last, for any $\mathbf{u} \in [0, \bar{v}]^n$, we have $\mathbf{W}^{(T-1)}(\mathbf{u} \mid \mathbf{U}) \in \mathcal{V}_{T-2} \subseteq \mathcal{V}_{T-1}$ and $\mathbf{U} \in \mathcal{V}_{T-1}$. Because $\mathcal{V}_{T-1}$ is convex, this implies $\mathbf{W}(\mathbf{u} \mid \mathbf{U}) \in \mathcal{V}_{T-1}$. That is, Eq. (71) also holds, which ends the inductive proof that $\mathcal{V}_{T-1} \subseteq \mathcal{V}_T$. This proves the desired result. ■

As a consequence of Theorem 6.1, all techniques developed earlier to give upper bounds on $\mathcal{U}_{\gamma(T)}$ also apply to $\mathcal{V}_T$. We next show the simple universal $1/T$ lower bound for the finite-horizon setting.

Proof of Theorem 6.2 Let $\mathbf{U} \in \mathcal{V}_T$ and consider an allocation strategy that realizes $\mathbf{U}$. For any $t \in [T]$, we recall the notation $\mathbf{p}(t) \in \Delta_n$ and $\mathbf{u}(t)$ for the allocation distribution and the

agent utilities respectively at time t . The main remark is that the suboptimality of $\mathbf{U}$ is at least that of the first allocation $\mathbf{p}(1)$ and this corresponds to an allocation from $\mathcal{V}_1$:

$$\begin{aligned} \max_{\mathbf{x} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{x} - \mathbf{1}^\top \mathbf{U} &= \mathbb{E} \left[\frac{1}{T} \sum_{t \in [T]} \max_{\mathbf{x} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{x} - \mathbf{1}^\top (p_i(t) u_i(t))_{i \in [n]} \right] \\ &\geq \frac{1}{T} \mathbb{E} \left[\max_{\mathbf{x} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{x} - \mathbf{1}^\top (p_i(1) u_i(1))_{i \in [n]} \right] \stackrel{(i)}{\geq} \frac{1}{T} \left(\max_{\mathbf{x} \in \mathcal{U}^*} \mathbf{1}^\top \mathbf{x} - \max_{\mathbf{x} \in \mathcal{U}_0} \mathbf{1}^\top \mathbf{x} \right). \end{aligned}$$

In (i) we used the fact that $(\mathbb{E}[p_i(1) u_i(1)])_{i \in [n]} \in \mathcal{V}_1 = \mathcal{U}_0$. By assumption, the term in the bracket in the last inequality is nonzero, otherwise $\mathcal{U}_0$ would contain an optimal vector. ■

We next turn to the lower bound tools on $\mathcal{U}_\gamma$ and show that they still hold in the finite-horizon T setting up to logarithmic factors. We start by showing that Theorem 3.1 still holds with a variant of the mechanism PUM. The proof proceeds by giving a trajectory of achievable regions for all times in $t \in [T]$, illustrated in Fig. 8: one can approximate balls within $\mathcal{U}^*$ by constructing balls within $\mathcal{V}_t$ converging to the original ball as $t \in [T]$ grows. The corresponding mechanism FH-PUM (Finite-Horizon Promised Utility Mechanism) is described formally within the following proof and summarized in Algorithm 2.

Proof of Theorem 6.3 Let $\mathbf{x}$ and $r, \delta > 0$ satisfying the requirements. The assumptions imply that $0 < x_i \leq \mathbb{E}[u_i]$ for all $i \in [n]$. Now let $r_0 := \min_{i \in [n]} \mathbb{E}[u_i]/(6n)$ and let $\mathbf{x}^{(0)} := 3r_0 \mathbf{1}$. We can note that

$$B(\mathbf{x}^{(0)}, 3r_0) \subseteq [0, 6r_0]^n \subseteq \{\mathbf{y} : \mathbf{0} \leq \mathbf{y} \leq (\mathbb{E}[u_i]/n)_{i \in [n]}\} \subseteq \mathcal{U}_0.$$

Also, up to renaming the constant r_0 to $r_0/2$, the assumption gives $r + \delta \leq 2r \leq r_0$. We next introduce the constant $\tilde{c}_0 := (1 + \frac{\max_{i \in [n]} \mathbb{E}[u_i]}{r_0})^2$.

Step 1. We start with the case in which the assumption that Eq. (13) holds with the same constant $\tilde{C} := \tilde{c}_0 C$ where C is the same constant as in the original statement of Theorem 3.1, and we have $\delta \leq r \leq r_0/2$. We will later see that from there we can easily extend to the general case.

We define the set $\mathcal{T} = \left\{ t \in \{2, \dots, T\}, r \geq \sqrt{\frac{\tilde{C}(1-\gamma(t))}{\gamma(t)}} = \sqrt{\frac{\tilde{C}}{t-1}} \right\}$. Note that by hypothesis, we have $2r \geq r + \delta \geq r + \tilde{C} \frac{1-\gamma(T)}{\gamma(T)r} \geq 2\sqrt{\frac{\tilde{C}(1-\gamma(T))}{\gamma(T)}}$. Thus, $T \in \mathcal{T}$. Because $\frac{1-\gamma(t)}{\gamma(t)}$ is decreasing in t , there exists $T_1 \in [T]$ such that $\mathcal{T} = \{T_1, T_1 + 1, \dots, T\}$. We next let

$$T_0 = \max \left\{ t \in [T], 2\sqrt{\frac{\tilde{C}(1-\gamma(t))}{\gamma(t)}} \geq r_0 \right\},$$

which is well defined since $\gamma(1) = 0$ and hence the set always contains time $t = 1$. We note that for T_1 , we have

$$2\sqrt{\frac{\tilde{C}(1-\gamma(T_1))}{\gamma(T_1)}} < 2r \leq r_0,$$

Therefore, $T_0 < T_1$ (by construction $T_1 \geq 2$). We now specify our allocation mechanism.

Input: Ball $B(\mathbf{x}, r)$, margin δ with $B(\mathbf{x}, r + \delta) \subseteq \mathcal{U}^*$, target utility $\mathbf{U}_0 \in B(\mathbf{x}, r)$, horizon T

1 **Initialization:** $r_0 := \frac{\min_{i \in [n]} \mathbb{E}[u_i]}{6n}$, $\mathbf{x}^{(0)} := 3r_0 \mathbf{1}$, $\tilde{\delta} := \frac{\delta}{2}$, $\mathbf{z} = \mathbf{x} + \tilde{\delta} \frac{\mathbf{U}_0 - \mathbf{x}}{\|\mathbf{U}_0 - \mathbf{x}\|}$, constant $\tilde{C} := 4(1 + \frac{\max_{i \in [n]} \mathbb{E}[u_i]}{r_0})^2 C$ where C is defined as in Theorem 3.1.

2 Let $\mathbf{U} \leftarrow \mathbf{U}_0$, $T_0 = \max\{t \in [T] : 2\sqrt{\frac{\tilde{C}}{t-1}} \geq r_0\}$, and $\mathbf{x}^{(T_0)} = \mathbf{z}^{(0)}$

3 **for** $t = T_0 + 1, \dots, T$ **do**

4 **if** $r < \sqrt{\frac{\tilde{C}}{t-1}}$ **then** Define $r_t = \delta_t := \sqrt{\frac{\tilde{C}}{t-1}}$;

5 **else** Define $r_t := r$ and $\delta_t := \max(\frac{\tilde{C}}{r(t-1)}, \tilde{\delta})$;

6 Define $\mathbf{y}^{(t)} := \mathbf{z} + \frac{r_t + \delta_t - (r + \tilde{\delta})}{2r_0 - (r + \tilde{\delta})}(\mathbf{x}^{(0)} - \mathbf{z})$ and $\mathbf{x}^{(t)} := (1 - \frac{1}{t})\mathbf{x}^{(t-1)} + \frac{1}{t}\mathbf{y}^{(t)}$

7 **Main mechanism:** **for** $t = T, T-1, \dots, 1$ **do**

8 **if** $\mathbf{U} \in \mathcal{U}^{NI}$ **then** Observe new reports $\mathbf{v} = (v_i)_{i \in [n]}$ and select allocation $(\frac{v_i}{\mathbb{E}[u_i]})_{i \in [n]}$;

9 **else**

10 Decompose $\mathbf{U} = \mathbf{s} \odot \left(\sum_{i \in [n]} q_i \mathbb{E}[u_i] \mathbf{e}_i + q_0(\mathbf{x}^{(t)} + r_t \mathbf{y}) \right)$ where $\mathbf{q} \geq 0$, $\sum_{i=0}^n q_i \leq 1$, $\mathbf{y} \in S_{n-1}$, $\mathbf{s} \in [0, 1]^n$.

11 Fix an allocation function $\tilde{\mathbf{p}}(\cdot)$ realizing utility $\mathbf{y}^{(t)} + (r_t + \delta_t)\mathbf{y}$ for full-information case

12 Compute interim functions $\tilde{P}_i(\cdot) := \mathbb{E}_{\mathbf{u}_{-i}}[\tilde{p}_i(\cdot, \mathbf{u}_{-i})]$ and $\tilde{W}_i(\cdot)$ according to Eq. (6) for $i \in [n]$

13 Observe new reports $\mathbf{v} = (v_i)_{i \in [n]}$ and select allocation $\mathbf{s} \odot \left(\sum_{i \in [n]} q_i \mathbf{e}_i + q_0 \tilde{\mathbf{p}}(\mathbf{v}) \right)$

14 Define $\tilde{\mathbf{W}}$ using Eq. (11) for $\gamma = 1 - \frac{1}{t}$, $\boldsymbol{\alpha} = -r_t \mathbf{y}$, and $\tilde{Z}_i = \tilde{W}_i(v_i)$ for $i \in [n]$

15 Update $\mathbf{U} \leftarrow \mathbf{s} \odot \left(\sum_{i \in [n]} q_i \mathbb{E}[u_i] \mathbf{e}_i + q_0 \tilde{\mathbf{W}} \right)$

16 **if** $\mathbf{U} \notin \mathcal{U}^{NI}$ and $\mathbf{U} \notin \mathcal{R}_{t-1}$ as defined in Eq. (73) **then** Mechanism fails. **break**;

Algorithm 2: Promised Utility Mechanism FH-PUM($\mathbf{x}, r, \delta$) for the T -horizon setting

Definition of the allocation strategy. We use a trivial allocation for times in $[T_0]$: we will only use the fact that $B(\mathbf{x}^{(0)}, 3r_0) \subseteq \mathcal{U}_0 \subseteq \mathcal{V}_{T_0}$. For times in $(T_0, T]$, we let

$$r_t := \begin{cases} \sqrt{\frac{\tilde{C}(1-\gamma(t))}{\gamma(t)}} = \sqrt{\frac{\tilde{C}}{t-1}} & t \in \{T_0, \dots, T_1 - 1\} \\ r & t \in \{T_1, \dots, T\}. \end{cases}$$

To prove the desired result, we inductively prove that $B(\mathbf{x}^{(t)}, r_t) \subseteq \mathcal{V}_t$ for appropriately chosen parameters $\mathbf{x}^{(t)}$. Precisely, we first introduce the parameter

$$\delta_t := \begin{cases} r_t & t \in \{T_0 + 1, \dots, T_1 - 1\} \\ \max\left(\tilde{C} \frac{1-\gamma(t)}{\gamma(t)r_t}, \delta\right) & t \in \{T_1, \dots, T\}. \end{cases}$$

We pose $\mathbf{x}^{(t)} := \mathbf{x}^{(0)}$ for all $t \in [T_0]$ and for $t > T_0$ we introduce

$$\mathbf{y}^{(t)} := \mathbf{x} + \frac{r_t + \delta_t - (r + \delta)}{2r_0 - (r + \delta)}(\mathbf{x}^{(0)} - \mathbf{x}).$$

Note that by construction, $r_t + \delta_t \geq r + \delta$: this is immediate for $t \in [T_1, T]$ and for $t \in (T_0, T_1)$ since $t < T_1$ one has $r_t + \delta_t = 2r_t > 2r \geq r + \delta$. We then construct the sequence of vectors $\mathbf{x}^{(t)}$

via

$$\mathbf{x}^{(t)} := \gamma(t)\mathbf{x}^{(t-1)} + (1 - \gamma(t))\mathbf{y}^{(t)}, \quad t \in \{T_0 + 1, \dots, T\}. \quad (72)$$

To prove that $B(\mathbf{x}^{(t)}, r_t) \subseteq \mathcal{V}_t$, we construct an allocation mechanism on the region

$$\mathcal{R}_t := \{\mathbf{y} : \mathbf{0} \leq \mathbf{y} \leq \mathbf{z}, \mathbf{z} \in \text{Conv}(\mathcal{U}^{NI}, B(\mathbf{x}^{(t)}, r_t))\}. \quad (73)$$

Next, by construction, note that

$$B(\mathbf{y}^{(t)}, r_t + \delta_t) = \frac{r_t + \delta_t - (r + \delta)}{2r_0 - (r + \delta)} B(\mathbf{x}^{(0)}, 2r_0) + \frac{2r_0 - (r_t + \delta_t)}{2r_0 - (r + \delta)} B(\mathbf{x}, r + \delta). \quad (74)$$

By convexity of $\mathcal{U}^*$, since we have $B(\mathbf{x}^{(0)}, 2r_0), B(\mathbf{x}, r + \delta) \subseteq \mathcal{U}^*$ this also implies that $B(\mathbf{y}^{(t)}, r_t + \delta_t) \subseteq \mathcal{U}^*$. This is precisely the reason behind the definition of $\mathbf{y}^{(t)}$.

We now construct an allocation function $\mathbf{p}^{(t)}(\cdot \mid \mathbf{U})$ and a promised utility function $\mathbf{W}^{(t)}(\cdot \mid \mathbf{U})$ for $\mathbf{U} \in \mathcal{R}_t$ in a similar fashion to those constructed in the proof of Theorem 3.1. We first specify the allocation strategy when $\mathbf{U}$ is an extreme point of $\mathcal{R}_t$ that still belongs to $B(\mathbf{x}^{(t)}, r_t)$. We write $\mathbf{U} = \mathbf{x}^{(t)} + r_t \mathbf{y}$ where $\mathbf{y} \in S_{n-1}$ and let

$$\mathbf{p}^{(t)}(\cdot \mid \mathbf{U}) := \mathbf{p}(\cdot; \mathbf{y}^{(t)} + (r_t + \delta_t)\mathbf{y}) \quad \text{and} \quad \boldsymbol{\alpha}(\mathbf{U}) = \mathbf{x}^{(t)} - \mathbf{U} = -r_t \mathbf{y}, \quad t \in (T_0, T], \quad (75)$$

where $\mathbf{p}(\cdot; \mathbf{z})$ refers to an allocation that realizes $\mathbf{z}$ in the full information setting (we checked earlier that we always used an allocation $\mathbf{p}(\cdot; \mathbf{z})$ with $\mathbf{z} \in \mathcal{U}^*$). We then extend the definition of the allocation mechanism to the complete region $\mathcal{R}_t$ exactly as in the original proof using Eqs. (25) and (26).

Checking that the allocation strategy is valid. We now check that this is a valid allocation. Using the same linearity arguments as in the proof of Theorem 3.1, it suffices to show that for extreme points $\mathbf{U}$ of $\mathcal{R}_t$ within $B(\mathbf{x}^{(t)}, r_t)$, the promised utilities lie in $B(\mathbf{x}^{(t-1)}, r_{t-1})$ to show that for all $\mathbf{U} \in \mathcal{R}_t$, Eq. (7) is satisfied, and

$$\forall \mathbf{v} \in [0, \bar{v}]^n, \quad \mathbf{W}^{(t)}(\mathbf{v} \mid \mathbf{U}) \in \mathcal{R}_{t-1}. \quad (76)$$

Fix such an extreme point $\mathbf{U} = \mathbf{x}^{(t)} + r_t \mathbf{y}$ where $\mathbf{y} \in S_{n-1}$. First, recall that $B(\mathbf{y}^{(t)}, r_t + \delta_t) \subseteq \mathcal{U}^*$. We now check that Eq. (13) is also satisfied. Note that by construction the sequences r_t and δ_t are both non-increasing on $(T_0, T]$. This is clear by definition of T_1 for the sequence $(r_t)_t$. For δ_t , we only need to check that $\delta_{T_1-1} = r_{T_1-1} \geq \delta_{T_1}$. To do so, we use $\delta_{T_1} \leq r_{T_1} \leq r_{T_1-1}$. By Eq. (74), we have $y_i^{(t)} + r_t + \delta_t \leq \max(x_i^{(0)} + 2r_0, x_i + r + \delta)$. Now by hypothesis Eq. (13) is satisfied for the parameters $\mathbf{x}, r, \delta$ and $\gamma = \gamma(T)$, and the constant $\tilde{C}$ instead of C . Hence, $\delta \geq \tilde{C} \frac{1-\gamma(T)}{\gamma(T)r}$ which implies $\delta_T = \delta$. Hence, because δ_t is non-increasing in t , we have $\delta_t \leq \delta$ for all $t \in [T_0, T_1]$, which implies

$$y_i^{(t)} + r_t \leq \max(x_i^{(0)} + 2r_0, x_i + r) \leq \max(\mathbb{E}[u_i] - r_0, x_i + r),$$

where in the last inequality we used the fact that $B(\mathbf{x}^{(0)}, 3r_0) \subseteq \mathcal{U}_0$. The previous remarks imply that

$$\tilde{C} = \tilde{c}_0 4n^2 \bar{v}^2 \left(1 + \frac{\max_{i \in [n]} \mathbb{E}[u_i]}{\min_{i \in [n]} (\mathbb{E}[u_i] - x_i - r)} \right)^2 \geq 4n^2 \bar{v}^2 \left(1 + \frac{\max_{i \in [n]} \mathbb{E}[u_i]}{\min_{i \in [n]} (\mathbb{E}[u_i] - y_i^{(t)} - r_t)} \right)^2.$$

Last, we check that

$$\frac{\gamma(t)r_t}{1-\gamma(t)} \stackrel{(i)}{\geq} r_t \geq \delta_t \geq \tilde{C} \frac{1-\gamma(t)}{\gamma(t)r_t}.$$

where in (i) we used the fact that $t \geq T_0 + 1 \geq 2$, hence $\gamma(t) \geq 1/2$. Altogether, the previous arguments show that Eq. (13) is satisfied for the parameters $\mathbf{y}^{(t)}, r_t, \delta_t$ and $\gamma = \gamma(t)$.

Next, compared to the definition of Theorem 3.1, these would be identical if the allocation was for $\mathbf{U}' := \mathbf{y}^{(t)} + r_t \mathbf{y}$ instead of $\mathbf{U} = \mathbf{x}^{(t)} + r_t \mathbf{y}$. To clarify this comparison, we introduce the allocation function $\mathbf{p}'(\cdot | \mathbf{U}') := \mathbf{p}(\cdot; \mathbf{U})$ and denote by $\mathbf{W}'(\cdot | \mathbf{U}')$ the promised utility function resulting from this choice of allocation function and the same direction choice $\boldsymbol{\alpha}(\mathbf{U})$ as in Eq. (75). Having checked that all conditions apply, the proof of Theorem 3.1 implies that for all $\mathbf{v} \in [0, \bar{v}]^n$, $\mathbf{W}'(\mathbf{v} | \mathbf{U}') \in B(\mathbf{y}^{(t)}, r_t)$. Now by definition of the promised utility functions (see Eqs. (6) and (11)), for any $\mathbf{v} \in [0, \bar{v}]^n$,

$$\mathbf{W}^{(t)}(\mathbf{v} | \mathbf{U}) = \mathbf{W}'(\mathbf{v} | \mathbf{U}') + \frac{\mathbf{U} - \mathbf{U}'}{\gamma(t)} = \mathbf{W}'(\mathbf{v} | \mathbf{U}') - \frac{\mathbf{y}^{(t)} - \mathbf{x}^{(t)}}{\gamma(t)}.$$

As a result, using $r_t \leq r_{t-1}$, we obtain

$$\mathbf{W}^{(t)}(\mathbf{v} | \mathbf{U}) \in B\left(\mathbf{y}^{(t)} + \frac{\mathbf{x}^{(t)} - \mathbf{y}^{(t)}}{\gamma(t)}, r_t\right) = B(\mathbf{x}^{(t-1)}, r_t) \subseteq B(\mathbf{x}^{(t-1)}, r_{t-1}).$$

This ends the proof of Eq. (76) hence $\mathcal{R}_t \subseteq \mathcal{V}_t$. As a summary, we obtained the induction: if $B(\mathbf{x}^{(t-1)}, r_{t-1}) \subseteq \mathcal{V}_{t-1}$, then $B(\mathbf{x}^{(t)}, r_t) \subseteq \mathcal{V}_t$ where $\mathbf{x}^{(t)}$ is defined in Eq. (72). We recall that the initialization at $t = T_0$ is immediate from $B(\mathbf{x}^{(0)}, r_0) \subseteq \mathcal{U}_0 \subseteq \mathcal{V}_{T_0}$.

Putting the recursive construction together. Using the induction formula Eq. (72), we can check that $\mathbf{x}^{(t)} = a_t \mathbf{x}^{(0)} + b_t(\mathbf{x}^{(0)} - \mathbf{x}) + c_t \mathbf{x}$ for $t \in [T_0, T]$ where $a_{T_0} = 1$, $b_{T_0} = c_{T_0} = 0$ and

$$\begin{cases} a_t = \gamma(t)a_{t-1} \\ b_t = \gamma(t)b_{t-1} + (1-\gamma(t))\frac{r_t + \delta_t - (r + \delta)}{2r_0 - (r + \delta)} \\ c_t = \gamma(t)c_{t-1} + (1-\gamma(t)), \end{cases} \quad t \in [T_0 + 1, \dots, T_1 - 1].$$

Using the convexity inequality $\ln(1-x) \leq -x$ for $x \leq 1$, we obtain

$$0 \leq a_T = 1 - c_T = \prod_{t=T_0+1}^T \gamma(t) \leq e^{-\sum_{t=T_0+1}^T \gamma(t)} \leq \frac{T_0 + 1}{T + 1}. \quad (77)$$

The recursion for b_t requires more work. First, recall that $r + \delta \leq r_0$ and $r_t + \delta_t \geq r + \delta$. Therefore, if we define the inductive sequence $\tilde{b}_{T_0} := 0$ and

$$\tilde{b}_t = \gamma(t)\tilde{b}_{t-1} + (1-\gamma(t))\frac{r_t + \delta_t - (r + \delta)}{r_0},$$

we have $0 \leq b_t \leq \tilde{b}_t$ for all $t \in [T_0, T]$. Note that the induction can be rewritten as follows,

$$r_0 \tilde{b}_t = r_0(t-1)\tilde{b}_{t-1} + r_t + \delta_t - (r + \delta).$$

As a result, letting $T_2 = \min\{t \geq T_1, \delta_t = \delta\}$, we obtain

$$\begin{aligned} r_0 T \tilde{b}_T &= \sum_{t=T_0+1}^T (r_t + \delta_t - r - \delta) = 2 \sum_{t=T_0+1}^{T_1-1} r_t - (T_1 - T_0 - 1)(r + \delta) + \sum_{t=T_1}^{T_2-1} \delta_t - (T_2 - T_1)\delta \\ &\stackrel{(i)}{\leq} 4\sqrt{\tilde{C}(T_1 - 2)} + \frac{\tilde{C}}{r} \left(1 + \ln \frac{T_2 - 2}{T_1 - 1} \mathbb{1}_{T_2 > T_1}\right), \end{aligned}$$

where in (i) we used $r_t = \sqrt{\tilde{C}/(t-1)}$ for $t \in (T_0, T_1)$ and $\delta_t = \tilde{C}/((t-1)r)$ for $t \in [T_1, T_2)$, and used sum-integral inequalities. Now by definition of T_1 and T_2 , we have $\sqrt{\frac{\tilde{C}}{T_1-2}} > r \geq \sqrt{\frac{\tilde{C}}{T_1-1}}$ and if $T_2 > T_1$, we have $\delta < \tilde{C} \frac{1-\gamma(T_2-1)}{\gamma(T_2-1)r} = \frac{\tilde{C}}{r(T_2-2)}$. As a result, we obtained

$$0 \leq b_T \leq \frac{4\sqrt{\tilde{C}(T_1 - 2)}}{r_0 T} + \frac{\tilde{C}}{r_0 r T} \left(1 + \ln \frac{r}{\delta}\right) \leq \frac{\tilde{C}}{r_0 r T} \left(5 + \ln \frac{r}{\delta}\right). \quad (78)$$

Step 2. We now consider the original case in which $B(\mathbf{x}, r + \delta) \subseteq \mathcal{U}^*$, $\delta \leq r \leq r_0/2$, and Eq. (18) holds with the constant $c_0 = 8\tilde{c}_0 + 20\tilde{c}_0/r_0 + 4(T_0 + 1)r_0/C$. We denote by C the term in the statement of Theorem 3.1 for parameters $\mathbf{x}$ and r . Fix any $\mathbf{z} \in B(\mathbf{x}, \delta/2)$. Then, with $\tilde{\delta} = \delta/2$, we still have $B(\mathbf{z}, r + \tilde{\delta}) \subseteq \mathcal{U}^*$, and

$$r \geq \tilde{\delta} \geq \frac{c_0 C}{2} \frac{1 - \gamma(T)}{\gamma(T)r} \geq 4\tilde{c}_0 C \frac{1 - \gamma(T)}{\gamma(T)r}.$$

Note that the constant corresponding to the parameters $\mathbf{z}, r$ in the statement of Theorem 3.1 is at most $4C$ (note that $\|\mathbf{z} - \mathbf{x}\| \leq \delta/2$ and since $B(\mathbf{x}, r + \delta) \subseteq \mathcal{U}^*$, this still leaves a margin at least $\delta/2$ to the boundary of $\mathcal{U}^*$ for $B(\mathbf{z}, r)$). Therefore, the previous step shows that $B(\mathbf{z}^{(T)}, r) \subseteq \mathcal{V}_T$ where

$$\mathbf{z}^{(T)} = a_T \mathbf{x}^{(0)} + b_T (\mathbf{x}^{(0)} - \mathbf{z}) + c_T \mathbf{z}.$$

Because this holds for all $\mathbf{z} \in B(\mathbf{x}, \delta/2)$, we obtained

$$B\left(a_T \mathbf{x}^{(0)} - b_T \mathbf{1} + c_T \mathbf{x}, r + \frac{\delta}{2}\right) \subseteq \mathcal{V}_T. \quad (79)$$

Now using Eqs. (77) and (78), we have

$$\begin{aligned} \|a_T \mathbf{x}^{(0)} - b_T \mathbf{1} + c_T \mathbf{x} - \mathbf{x}\| &\leq \frac{T_0 + 1}{T + 1} \|\mathbf{x} - \mathbf{x}^{(0)}\| + \frac{\tilde{C}(5 + \ln \frac{r}{\delta})}{r_0 r T} \|\mathbf{x} - \mathbf{x}^{(0)}\| \\ &\leq \frac{(T_0 + 1)r_0^2 + \tilde{C}(5 + \ln \frac{r}{\delta})}{r_0 r T} \sqrt{n} \max_{i \in [n]} \mathbb{E}[u_i] \end{aligned}$$

As a result, recalling the definition of c_0 , we obtain

$$\|a_T \mathbf{x}^{(0)} - b_T \mathbf{1} + c_T \mathbf{x} - \mathbf{x}\| \leq \frac{c_0 C (1 + \ln \frac{r}{\delta})}{2r(T + 1)} \leq \frac{\delta}{2}.$$

Together with Eq. (79) this implies the desired result $B(\mathbf{x}, r) \subseteq \mathcal{V}_T$. This proves the first claim. The corresponding mechanism FH-PUM($\mathbf{x}, r, \delta$) is summarized in Algorithm 2.

For the second claim, we simply note that by assumption, we have $r \geq \delta \geq \frac{c_0 C}{r(T-1)}$. In particular, $\frac{r}{\delta} \leq \frac{r^2}{c_0 C}(T-1) \leq \frac{r_0^2}{c_0 C}T$. Hence, up to modifying the constant c_0 , the result holds if we replace the lower bound from Eq. (18) with $c_0 C^{\frac{1-\tilde{\gamma}(T)}{\tilde{\gamma}(T)}}$ where $\tilde{\gamma}(T) := 1 - (1 + \ln T)/T$. ■

As a remark, note that when the margin δ is of the same order as r within Theorem 6.3, then the logarithmic term $\ln \frac{r}{\delta}$ is a constant, which essentially means that we can use the same techniques to prove lower bounds on $\mathcal{V}_T$ as for $\mathcal{U}_{\gamma(T)} = \mathcal{U}_{1-1/T}$, up to worsening constants. This is significant since from Theorem 6.1, we also have $\mathcal{V}_T \subseteq \mathcal{U}_{\gamma(T)}$: in these cases, characterizing $\mathcal{V}_T$ or $\mathcal{U}_{\gamma(T)}$ with our tools is essentially equivalent. On the other hand, as per the second claim of Theorem 6.3, the extra factor $\ln \frac{r}{\delta}$ can be of order $\ln T$, which potentially induces additional logarithmic factors to translate results from the infinite-horizon setting to the finite-horizon setting.

Similarly as for Theorem 3.1, the optimized version Theorem B.2 also holds in the finite-horizon setting up to worsening the constants.

Lemma E.1. *Fix $T \geq 2$. There exist constants $c_0 \geq 8$ and $r_0 > 0$ such that the following holds. Suppose that all assumptions from Theorem B.2 are satisfied with the exception that Eq. (13) is replaced with Eq. (18). Then, using the same notations,*

$$B(\mathbf{x}, r) \cap \{\mathbf{y} : \forall i \notin I, y_i = x_i\} \subseteq \mathbf{x} + \{0\}^{[n] \setminus I} \otimes \bigotimes_{s \in [q]} B_{I_s}(\mathbf{0}, r) \subseteq \mathcal{V}_T.$$

Alternatively, suppose that all assumptions from Theorem B.2 are satisfied with $\gamma = \tilde{\gamma}(T) = 1 - \ln T/T$ and the constant $C' = c_0 C$, as well as $r_0 \geq r \geq \delta$. Then the previous conclusion also holds.

Proof The same proof as in Theorem B.2 carries to the finite-horizon setting: we can construct an allocation mechanism that treats separately agents by clusters $I_1, \dots, I_q$ as described in the proof. The only difference is that instead of applying Theorem 3.1 we apply Theorem 6.3, which required exactly the assumptions made within Theorem E.1. ■

As a conclusion, up to logarithmic factors, all lower bounds developed in the previous sections apply to the finite-horizon setting choosing $\gamma = \gamma(T) = 1 - 1/T$. We now prove Theorems 7.2 and 7.5 as examples of how the results from the γ -discounted setting can help for the finite-horizon setting.

Proof of Theorem 7.2 It suffices to prove the result for $T \geq 2$. The lower bound is directly taken from Theorem 6.2 and the upper bound is a consequence of the second claim of Theorem 7.1 together with Theorem E.1. ■

Proof of Theorem 7.5 It suffices to prove the result for $T \geq 2$. The first claim is a consequence of the first result from Theorem D.6 together with Theorem 6.1 for the lower bound and Theorem 6.3 for the upper bound. Note that there is no extra $\ln T$ factor (we can take directly $\gamma = \gamma(T) = 1 - 1/T$) in the upper bound because the proof of the convergence rate $\mathcal{O}(\sqrt{1-\gamma})$ comes from Theorem D.2, which only uses Theorem 3.1 with parameters $r = \delta$. As a result, the extra term $\ln \frac{r}{\delta}$ from Theorem 6.3 vanishes. The second claim is a consequence of the second claim of Theorem D.6 together with Theorem E.1 for the upper bound. The lower bound is directly Theorem 6.2. ■