

Efficient Mechanisms under Unawareness

KYM PRAM

Department of Economics, University of Nevada, Reno

BURKHARD C. SCHIPPER

Department of Economics, University of California, Davis

We study the design of efficient mechanisms under asymmetric awareness and information. Unawareness refers to the lack of conception rather than the lack of information. Assuming quasi-linear utilities and private values, we show that we can implement in conditional dominant strategies a social choice function that is utilitarian ex-post efficient when pooling all awareness of all agents without the need of the social planner being fully aware ex-ante. To this end, we develop novel dynamic versions of Vickrey-Clarke-Groves mechanisms in which types are revealed and subsequently elaborated at endogenous higher awareness levels. We explore how asymmetric awareness affects budget balance and participation constraints. We show that ex-ante unforeseen contingencies are no excuse for deficits. Finally, we propose a modified reverse second price auction for efficient procurement of complex incompletely specified projects.

KEYWORDS. Dynamic mechanism design, VCG mechanisms, auctions versus negotiations, unknown unknowns, complex projects.

1. INTRODUCTION

Mechanism design studies the design of institutions governing collective decisions such as markets, contracts, or political systems in the presence of asymmetric information. However, agents may not just face asymmetric information but also asymmetric awareness. Unawareness refers to the lack of conception rather than the lack of information. Agents and the designer of mechanisms may be unaware of some events and actions affecting values and costs of complex private or public projects, allocations, and outcomes and may not even realize this fact. In this paper, we extend mechanism design to unawareness. Under the assumption of quasi-linear preferences, we show how to design efficient mechanisms in the presence of asymmetric awareness. To this end, we introduce dynamic direct elaboration mechanisms in which not only are types communicated

Kym Pram: kpram@unr.edu

Burkhard C. Schipper: bcschipper@ucdavis.edu

04/05/2025. We thank Sarah Auster, Botond Koeszegi, Takashi Kunimoto, Matthias Lang, Lu Jingfeng, Moritz Meyer-ter-Vehn, Benny Moldovanu, Alessandro Pavan, Tomasz Sadzik, Klaus Schmidt, Roland Strausz, Tymon Tatur, Joel Watson, and participants in seminars at Bonn, HU Berlin, UC Davis, UC Irvine, LMU Munich, NSU Singapore, SWET 2025, SEA 2023, and SAET 2023 for useful discussions. Burkhard gratefully acknowledges financial support via ARO Contract W911NF2210282.

from agents but also awareness is raised among participants in back-and-forth communication between participants and transfers are inspired by Vickrey-Clarke-Groves (VCG) mechanisms.

Besides our theoretical motivation for extending mechanism design to unawareness, we are motivated by practical concerns about the usefulness of explicit mechanisms in reality. For instance, it has been argued that auctions are inappropriate when projects are complex, incompletely designed, and custom-made such as the procurement of new fighter jets, buildings, consulting services, and IT projects. It is claimed that auctions may stifle communication between buyers and sellers, preventing buyers and sellers to use each other's expertise when designing projects (Goldberg (1977), Bajari and Tadelis (2001), Bajari et al. (2008)). The Department of Defense (DoD), the General Service Administration (GSA), and the National Aeronautics and Space Administration (NASA) recently proposed to amend the Federal Acquisition Regulation (FAR) suggesting the prohibition of the use of reverse auctions for complex, specialized, or substantial design and construction services (Office of the Federal Register (2024)).¹ The received view is that for complex projects, negotiations outperform auctions because they can better take advantage of the expertise and know-how of contractors (Goldberg (1977), Sweet (1994), Bajari et al. (2008)). We seek an extension of mechanism design to unawareness to make mechanisms also applicable to complex and incompletely designed projects whose efficient implementation require the pooling of expertise among participants. Our proposed mechanisms combine features of common business practices such as Request for Information (RFI) and Request for Proposals (RFP)² with standard VCG mechanisms such as second price auctions. That is, we combine features of negotiations and traditional mechanism design. While traditionally in economics, negotiations have been interpreted mostly as bargaining over the surplus and thus as a substitute to mechanisms, we view negotiations more as interactively defining the surplus using the expertise of participants, which is complementary to traditional mechanisms implementing and allocating the surplus.

Extending mechanism design to unawareness has to overcome several obstacles. First, since the mechanism designer herself may be unaware of some payoff type profiles, how could she commit to outcomes for such type profiles? That is, how can she commit to a social choice function without being fully aware of the domain of the social choice function? It is well known that the revelation principle may fail if the mechanism designer cannot commit to the mechanism (e.g., Bester and Strausz (2001)).³ Yet, the

¹When analyzing the use of auctions for procurement at the DoD, Alper and Boning (2003) state that auctions work best for well-specified and off-the-shelf commodities but nevertheless remark that the Navy was able to achieve substantial savings with procurement auctions for customized and complex contracts such as CVN camels, i.e., devices for safe separation of ships and piers.

²E.g., Federal Acquisitions Regulation 15.203.

³In mechanisms under unawareness, there is another potential failure of the revelation principle. Types are always informative about the awareness of an agent because an agent cannot pretend to be a type of which she herself is unaware. In standard mechanism design, type-dependent message sets may lead to a failure of the revelation principle (e.g., Green and Laffont (1986), Bull and Watson (2007), Deneckere and Severinov (2008)). Nevertheless, in a separate note, we are able to prove a revelation principle for mechanisms under unawareness.

mechanism designer can at least commit to general properties of outcome functions like efficiency given ex-post awareness. Promising to implement an efficient outcome given what transpires in the mechanisms could be verified ex-post by a court of law. That's why we focus on efficient mechanism design under unawareness. We restrict ourselves to the economically relevant class of quasilinear preferences and the notion of utilitarian ex-post efficiency.

The second obstacle is that, as we will demonstrate in Section 2.2, static mechanisms will not suffice for efficient implementation under unawareness. We desire to implement efficiently at the *highest awareness level possible*. The problem is that no agent or the mechanism designer might be aware of everything. We seek to pool awareness of all agents and aim to implement the social choice that would be utilitarian ex-post efficient at this pooled awareness level. To achieve this, we need dynamic mechanisms that allow prior reported payoff types to be elaborated in light of awareness raised by all agents that is communicated back from the mediator to the agents.

This leads us to the third obstacle, namely how to model dynamic interactive unawareness. Recent years witnessed laying the foundation for modeling unawareness (Fagin and Halpern (1988), Heifetz et al. (2006, 2013b)). Results by Dekel et al. (1998) imply that standard type spaces preclude unawareness. As remedy, we introduce payoff type spaces that are simplified versions of unawareness type spaces (Heifetz et al. (2006, 2013b)). Payoff types are more than a value or cost. They can be complex descriptions of cost structures and factors affecting valuations involving both verbal, quantitative or pictorial descriptions that might be associated with RFI, RFP, Request for Tender (RFT), and Requests for Quotation (PFQ) in business practice. Important for modelling unawareness is that these descriptions can be ordered by how multi-dimensional, expressive, detailed, fine grained, rich etc. they are. More elaborate descriptions potentially raise awareness of factors that remain hidden or neglected in coarser descriptions of payoff types. The advantage of our abstract formalism is that it provides for a variety of complex descriptions of payoff types while allowing for tractability and generality. Recent years also saw the development of game theory with unawareness (Heifetz et al. (2013a), Meier and Schipper (2014, 2024), Schipper (2021), Halpern and Rêgo (2014), Feinberg (2021)). Our mechanisms introduced in this paper induce dynamic games with unawareness. Dynamic games with unawareness consist of a forest of trees ordered by expressiveness (i.e., the richness of payoff types in our case). An agent's information set at a history in one tree may comprise of histories in a less expressive tree precisely when this agent is unaware of some factors. After nature selected a payoff type and awareness level for each agent, agents report their perceived payoff type in the initial stage of the dynamic direct elaboration mechanism. After that, their awareness may be raised when the mediator provides feedback of the pooled awareness level from first-stage reports of all agents. At this point, agents have the opportunity to elaborate on their previously reported type at the pooled awareness level. This process of elaborations by agents and subsequent communication of pooled awareness by the mediator continues until no agent wants to elaborate any further at which point the mechanism stops and a utilitarian ex-post efficient outcome is implemented.

Implementation requires a solution concept, a fourth obstacle under unawareness. The dynamics of beliefs and awareness are highly complex, especially in light of awareness changing communication from the mediator. We avoid having to explicitly model the dynamics of beliefs by sticking to the belief-free approach of dominant strategy implementation. However, while dominant strategy implementation may be appropriate for static mechanisms, for dynamic mechanisms it would imply that agents select their entire dominant reporting strategy *ex-ante*. How can they do this if they are not even aware of all their future information sets and consequently their entire set of strategies? Inspired by conditional dominance for standard games in extensive form by [Shimoi and Watson \(1998\)](#) and extended to games with unawareness by [Meier and Schipper \(2024\)](#), we use implementation in conditional dominant strategies: Conditional on reaching an information set, the agent selects a weak dominant continuation strategy as far as he can anticipate at this information set. A side-benefit of being able to retain the belief-free approach under unawareness is that outcomes are robust to misspecifications of beliefs.

The last but crucial obstacle is how to incentivize not just the revelation of information but also awareness. Awareness is different from information. An example of payoff type information is “the cost of producing the item is at least x ” while an example of payoff type awareness is “when digging the foundation near the railway station we may or may not discover a WWII unexploded ordnance whose defusion will cause construction delays”. Note that latter statement involving “may or may not” is a tautology. As such, it does not carry information as it does not exclude any event. However, it surely raises awareness of the possibility of such events. In practice, payoff relevant statements often carry both information and awareness like “our prior experience suggests that when digging near a railways station, we are likely to find a WWII unexploded ordnance whose defusion will create a three day delay of construction at the cost of x ”. It raises awareness when the buyer has been unaware of the possibility. Moreover, it provides information in stating that such a discovery is “likely” and an estimate of the costs. The crucial point though is that when the buyer is unaware, she cannot ask for information about it in the contract tender and agents may not find it in their interest to volunteer such awareness, while when the buyer is aware, she can ask for information (and even infer information from silence; see [Milgrom \(1981\)](#), [Heifetz et al. \(2021\)](#)). Moreover, once the buyer becomes aware of the possibility, she might want to explicitly request information on their cost for such a contingency from *all* agents and thus sharing her newly won awareness among all agents.

We show that when we pair dynamic direct elaboration mechanisms with transfer schemes inspired by VCG mechanisms ([Groves \(1973\)](#), [Groves and Loeb \(1975\)](#)), then we can implement efficiently at the pooled awareness in conditional dominant strategies. It is well known that these transfers incentivize information revelation given any awareness level. However, our transfer schemes also include an additional term that incentivizes raising awareness. This term rewards raising awareness commensurate with its possible effect on social welfare. It does so only if there is a unique agent who raises awareness to the pooled awareness level. Moreover, compared to the standard VCG mechanism without unawareness, an agent has to pay its “fair” share of the awareness

raising incentives if some other unique agent raises awareness to the final pooled awareness level before she does herself (if such an agent exists).

Next, we investigate properties beyond efficiency. It is well known that standard VCG mechanisms without unawareness do not necessarily satisfy budget balance even without unawareness ([Green and Laffont \(1979\)](#)). We show that awareness pooling does not impose an additional constraint on budget balance. In particular, the characterizing condition for budget balance of our dynamic elaboration VCG mechanisms is the same as for standard VCG mechanisms without unawareness ([Holmström \(1977\)](#)). The additional incentives for raising awareness are designed to be budget neutral and deviations to lower awareness levels do not impose additional constraints on budget balance. Since budget balance is elusive in the general case (with or without unawareness), we need to look beyond budget balance. Unforeseen contingencies are often used to justify budget overruns. Thus, we are more interested in running no deficit than running a surplus. It is well-known that in standard mechanism design without unawareness, the Clarke or Pivot mechanisms ([Clarke \(1971\)](#)) satisfy no deficit. We show that we can implement a utilitarian ex-post efficient outcome under pooled awareness with no deficit in conditional dominant strategies using a dynamic direct elaboration mechanisms in which transfers are inspired by the Clarke mechanism plus the additional budget neutral term that incentivizes raising awareness. This result is important in the context of unawareness for two reasons: First, as mentioned above, often unforeseen contingencies are used as a justification of cost overruns. Our results imply that ex-ante unforeseen contingencies are not necessarily an excuse for cost overruns. Second, under unawareness the mechanism designer may not be able to anticipate how large subsidies in the mechanism could become. Agents who anticipate this may question whether the mechanism designer could really commit to such large transfers and consequently may not want to report truthfully. By showing that the dynamic elaboration Clarke mechanism satisfies no deficit, we can mitigate this concern.

Next, we investigate ex-post participation constraints. It is well-known in standard mechanism design that the Clarke mechanism satisfies participation constraints. We show that our dynamic elaboration Clarke mechanism satisfies ex-post participation constraints except that it only holds for ex-post outcomes that are ex-ante anticipated by the agent. While the agent may ex-ante be happy to participate in the mechanism, she may regret it at an interim stage because she comes to realize that she has to pay for the awareness raising incentives. This implies that when the agent is aware of her unawareness (i.e., she realizes that she might be unaware of something without being able to comprehend what she is unaware of; see [Schipper \(2024\)](#)), she might be reluctant to participate in the mechanisms ex-ante because she fears to be penalized when others raise awareness more than she herself does although she has no idea of what other agents could raise awareness of.

Finally, we focus on the special but important case of procurement under ex-ante unforeseen contingencies. We introduce the dynamic elaboration reverse second price auction that implements an utilitarian ex-post efficient outcome under pooled awareness with budget balance satisfying ex-post participation constraints of all sellers. After the revelation and elaboration stages of the mechanism conclude, the buyer pays the

second lowest bid to the seller with the lowest bid who gets the project awarded. The buyer also compensates the unique seller (which may not be the winning seller) who first raises awareness to the pooled awareness level if such a seller exists.

The paper is organized as follows: In the following section, we introduce payoff type spaces with unawareness and show via a simple example the failure of static VCG mechanisms under unawareness. This is followed by an exposition of dynamic direct elaboration mechanisms in Section 2.3 and the proof of efficiency in Section 3. In Section 4, we characterize budget balance. We also introduce the dynamic elaboration Clarke mechanisms and show no deficit. Participation constraints are discussed in Section 5 and procurement under ex-ante unforeseen contingencies is studied in Section 6. We conclude in Section 7.

1.1 *Related Literature*

Since mechanism design is closely related to contract theory, our paper contributes to the recent literature on contracting under unawareness (e.g., [Lee \(2008\)](#), [Sommer and Loch \(2009\)](#), [von Thadden and Zhao \(2012\)](#), [Auster \(2003\)](#), [Filiz-Ozbay \(2012\)](#), [Grant et al. \(2012\)](#), [Chung and Fortnow \(2016\)](#), [Auster and Pavoni \(2024\)](#), [Lei and Zhao \(2021\)](#), [Francetich and Schipper \(2024\)](#)). For instance, [Francetich and Schipper \(2024\)](#) show that screening contracts may not provide sufficient incentives to agents to reveal their awareness. Consequently, the principal may not be able to consider the full set of incentive constraints. In contrast, we show that we can go beyond incomplete contracts and reveal awareness in dynamic direct elaboration mechanisms that provide sufficient incentives for truthful reporting and raising awareness. However, we focus on efficient mechanisms rather than optimal mechanisms, affording us the opportunity of working in the belief-free paradigm of (conditional) dominant strategy implementation and allowing us to conveniently bypass the subtleties of updating beliefs under dynamic awareness. The literature on unawareness in contracting is very different from the earlier literature on indescribable contingencies in contracting. For instance, in [Maskin and Tirole \(1999\)](#), agents are fully aware of every particular payoff consequence but cannot ex-ante describe them. In contrast, under unawareness agents may not be fully aware of all payoff consequences in contracting. Closer to mechanism design, [Li and Schipper \(2024\)](#) study the seller's decision to raise bidders' awareness of characteristics before a second-price auction with entry fees. Optimal entry fees capture an additional unawareness rent due to unaware bidders misperceiving their probability of winning and the price to be paid upon winning. In contrast to our setting, the auctioneer is aware of everything ex-ante. [Piermont \(2024\)](#) studies how a decision maker can incentivize an expert to reveal novel aspects about a decision problem via an iterated revelation mechanism in which at each round the expert decides on whether or not to raise awareness of a contingency that in turn the decision maker considers when proposing a new contract that the expert can accept or reject. This bears some similarity with our dynamic direct elaboration mechanisms. Most importantly, we focus on a multi-agent setting with transferable utilities. [Piermont \(2024\)](#) shows that iterated revelation allows for efficient outcomes to emerge.

[Herweg and Schmidt \(2020\)](#) study a procurement problem with a principal and two agents who may be aware of some design flaws. Our paper differs from theirs in many

respects: First, in [Herweg and Schmidt \(2020\)](#) agents can raise awareness of realized design flaws, agent's private costs are independent of the design flaw, and fixing the design flaw requires a known common cost that is interim verified by an industry expert. In our model, agents report payoff types, can raise awareness of potential factors affecting payoffs, and individually elaborate how the payoffs change in light of new awareness. Second, [Herweg and Schmidt \(2020\)](#) construct an efficient direct mechanism under common awareness and then argue that there is an indirect mechanism that under asymmetric awareness that gives rise to the same outcomes and incentives. This leads them to conclude that there are also corresponding equilibria in these two mechanisms. Yet, an appropriate notion of equilibrium in mechanisms under asymmetric awareness should verify equilibrium behavior w.r.t what outcomes agents anticipate and how behavior is affected by changing anticipations of outcomes during the play. Our approach is more “direct”: We conduct our entire analysis of efficient conditional dominant strategy implementation in a direct mechanisms under asymmetric awareness thereby modeling every potential deviation from equilibrium behavior from the agent's point view. Finally, [Herweg and Schmidt \(2020\)](#) focus on the particular but very relevant application to procurement while we consider more generally the efficient mechanism design problem under asymmetric awareness.

We study the design of efficient mechanisms in which raising awareness by one agent can via the mediator awareness and thus valuations of other agents. That is, even though we work with private valuations, valuations become endogenously interdependent. We know from [Jehiel and Moldovanu \(2001\)](#) that it is generically impossible to implement ex-post efficient outcome functions in settings with interdependent valuations. Why are we able to nevertheless implement efficiently? The key is that the interdependence in payoffs stems from awareness, which agents can only misreport in one direction. By using transfers that are sufficiently increasing in reported awareness, we incentivize agents to report as much awareness as possible — namely, their true awareness level. A related result is given by [Krähmer and Strausz \(2024\)](#). We discuss the connection further after the results.

We make use of unawareness type spaces introduced by [Heifetz et al. \(2006, 2013b\)](#) as well as games with unawareness developed by [Heifetz et al. \(2013a\)](#), [Meier and Schipper \(2014\)](#), and [Schipper \(2021\)](#). For alternative approaches, see [Fagin and Halpern \(1988\)](#), [Halpern and Rêgo \(2014\)](#), [Feinberg \(2021\)](#), and [Board and Chung \(2021\)](#).

2. MODEL

2.1 *Payoff Types with Unawareness*

Let L be a finite lattice with order $\supseteq$. Elements of the lattice represent awareness levels. The join of the lattice is denoted by $\bar{\ell} \in L$. Define $L(\ell) := \{\ell' \in L : \ell' \trianglelefteq \ell\}$ for $\ell \in L$. This is the sublattice of L with join ℓ . The significance of $L(\ell)$ is that an agent with awareness level ℓ can only reason about awareness levels in $L(\ell)$.

Fix a nonempty finite set of agents I . For each agent $i \in I$, there is a collection of nonempty disjoint payoff type spaces $\{T_i^\ell\}_{\ell \in L}$. Let $\mathcal{T}_i := \bigcup_{\ell \in L} T_i^\ell$. A payoff type $t_i \in \mathcal{T}_i$ is more than just a value for an object. If $t_i \in T_i^\ell$, then describing the payoff type also

requires at least awareness level ℓ . We illustrate this feature with the following example.

Example 1 Consider the context of procurement. A principal may invite agents to bid on a complex project. Agents do their due diligence and identify relevant items that drive their costs. Suppose there are items $\{a, b, c\}$. Agents may be unaware of some items even after their due diligence. Agent 1 may only be aware of items $\ell_1 = \{a, b\}$ while agent 2 is only aware of items $\ell_2 = \{b, c\}$. In this case, L is isomorphic to the set of all subsets in $2^{\{a, b, c\}}$ and the natural lattice order is induced by set inclusion on $2^{\{a, b, c\}}$. Agent 1 may submit a bid given by the following table t_1 while agent 2 may submit a bid given by table t_2 .

$t_1 =$	Item	Cost
	a	23
	b	41
	Total	64

$t_2 =$	Item	Cost
	b	38
	c	29
	Total	67

As the example demonstrates, payoff types can be multi-dimensional with varying dimensions. Needless to say, the lattice approach is more general than multi-dimensional payoff type spaces. $\square$

For any agent $i \in I$ and awareness levels $\ell, k \in L$ with $k \supseteq \ell$, we require a surjective projection $r_\ell^k : T_i^k \rightarrow T_i^\ell$ such that for all $\ell, \ell', \ell'' \in L$, $\ell'' \supseteq \ell' \supseteq \ell$, $r_\ell^{\ell'}(r_{\ell'}^{\ell''}(t_i)) = r_\ell^{\ell''}(t_i)$ and $r_\ell^\ell(t_i) = t_i$. For brevity, we do not index r_ℓ^k by agents. The projection relates payoff types across awareness levels. Before we illustrate this notion, we also define extensions of payoff types to greater awareness levels. For any agent $i \in I$, awareness level $\ell \in L$, and payoff type $t_i \in T_i^\ell$, we let $t_i^\uparrow := \bigcup_{k \supseteq \ell, k \in L} (r_\ell^k)^{-1}(t_i)$. That is, $t_i^\uparrow$ is the union of inverse images at (weakly) greater awareness levels of payoff type t_i . Similarly, for any subset of payoff types in a given payoff type space, we let superscript “ $\uparrow$ ” indicate the union of inverse images in payoff type spaces corresponding to greater awareness levels. We illustrate these notions in our prior example.

Example 1 (Continuation) Continuing our prior example, consider payoff types t'_1 and t''_1 of agent 1:

	Item	Cost		Item	Cost
$t'_1 =$	a	23	$t''_1 =$	a	23
	b	41		b	41
	c	16		c	15
	Total	80		Total	79

Both payoff types are described with items in $\{a, b, c\}$. That is, both $t'_1, t''_1 \in T_1^{\{a, b, c\}}$. When payoff information on item c is stripped away from payoff type t'_1 we obtain payoff type t_1 . That is, $r_{\{a, b\}}^{\{a, b, c\}}(t'_1) = t_1$. Similarly, $r_{\{a, b\}}^{\{a, b, c\}}(t''_1) = t_1$. Thus, $t_1, t'_1, t''_1 \in t_1^\uparrow$. $\square$

For any agent $i \in I$, we define $\lambda: \mathcal{T}_i \rightarrow L$ by $\lambda(t_i) = \ell$ if $t_i \in T_i^\ell$. Again, for brevity we do not index λ by agents. The function λ indicates the awareness level that is required to describe the payoff type. For instance, in Example 1, $\lambda(t_1) = \{a, b\}$ while $\lambda(t'_1) = \{a, b, c\}$.

Our framework is general enough to capture payoff types described by all kinds of formal objects like sets of formulae in a formal language (e.g., [Heifetz et al. \(2008\)](#)), abstract sets, vectors, matrices, formal concepts, pre-sheaves etc. that may represent verbal, quantitative, or pictorial features of proposals, tenders, quotations, messages etc. in business practice. By abstracting from these particular features, we obtain a theory that works for more than one kind of formalism, ensure tractability, and focus on what is really essential for modeling awareness levels, namely the existence of an order of expressiveness of descriptions.

For each agent $i \in I$, payoff types and awareness levels are drawn consistently as follows: At awareness level $\bar{\ell}$, nature draws a payoff type $\bar{t}_i \in T_i^{\bar{\ell}}$ in the upmost payoff type space, interpreted as agent i 's true payoff type if she were aware of everything, and an awareness level $\ell_i \in L$. Consequently, the agent's perceived payoff type is $r_{\ell_i}^{\bar{\ell}}(t_i)$. That is, agent i can “miss something” but he cannot perceive the “wrong” payoff type w.r.t. what he is aware.

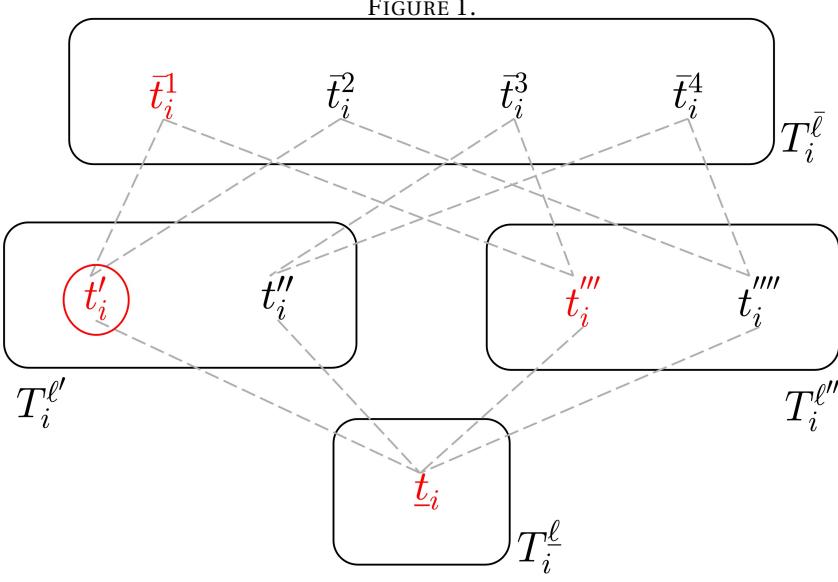
More generally, for any awareness level $\ell \in L$, nature draws agent i 's corresponding type $r_{\ell}^{\bar{\ell}}(\bar{t}_i)$ and agent i 's awareness level $\ell' = \ell_i \wedge \ell$. This means in particular, that if $\ell = \ell_i$, then it is payoff type $r_{\ell_i}^{\bar{\ell}}(t_i)$ and awareness level ℓ_i as just discussed. If $\bar{\ell} \supseteq \ell \supseteq \ell_i$, nature draws agent i 's corresponding type $r_{\ell}^{\bar{\ell}}(\bar{t}_i)$ and awareness level ℓ_i . If $\ell \in L$ with $\ell_i \supseteq \ell$, nature draws agent i 's corresponding type $r_{\ell}^{\bar{\ell}}(\bar{t}_i)$ and awareness level ℓ . Finally, if $\ell \in L$ with $\ell \not\supseteq \ell_i$, nature draws agent i 's corresponding type $r_{\ell}^{\bar{\ell}}(\bar{t}_i)$ and awareness level $\ell' = \ell_i \wedge \ell$. This specifies perceived payoff types and awareness levels consistently across the type spaces $\{T_i^{\ell}\}_{\ell \in L}$. Notice that if $t_i \in T_i^{\ell}$ for some $\ell \in L$, then the agent's perceived type is always in a type space with (weakly) less awareness than ℓ .

Figure 1 illustrates the payoff type structure. There are four awareness levels $L = \{\bar{\ell}, \ell', \ell'', \underline{\ell}\}$ with $\bar{\ell} \supset \ell' \supset \ell'' \supset \underline{\ell}$ and $\bar{\ell} \supset \ell'' \supset \underline{\ell}$. Associated with each awareness level is a payoff type space for agent i . Projections from richer type spaces to poorer type spaces are indicated by dashed lines. Suppose that the payoff type selected by nature in the upmost payoff type space $T_i^{\bar{\ell}}$ is $\bar{t}_i^1$ indicated in red. Then the corresponding payoff types selected in the other payoff type spaces follow by projections and are also indicated in red. When the awareness level of agent i selected by nature is $\ell_i = \ell'$, then t_i^1 is the payoff type perceived by agent i . This is indicated with the red circle in type space $T_i^{\ell'}$.

For any $\ell \in L$, let $\mathcal{T}^{\ell} := \times_{i \in I} T_i^{\ell}$. Moreover, let $\mathcal{T} := \times_{i \in I} \mathcal{T}_i$. Similarly, we let $\mathcal{T}_{-i}^{\ell} := \times_{j \in I \setminus \{i\}} T_j^{\ell}$ and $\mathcal{T}_{-i} := \times_{j \in I \setminus \{i\}} \mathcal{T}_j$.

When agents have the payoff type profile $\mathbf{t} = (t_i)_{i \in I} \in \mathcal{T}$, their *pooled* awareness level is $\check{\lambda}(\mathbf{t}) := \bigvee_{i \in I} \lambda(t_i) \in L$, i.e., the join of all agent's awareness levels at payoff type profile $\mathbf{t}$. Since L is a finite lattice, the join always exists in L . Note that $\check{\lambda}(\mathbf{t})$ may be a greater awareness level than any of the agent's awareness levels.

For each $\ell \in L$, there is a nonempty compact set of outcomes or allocations $X^{\ell} \subseteq X$. This formulation allows both for awareness of outcomes and awareness that some outcomes are infeasible.



Each agent $i \in I$ has an upper semi-continuous utility function $u_i : \bigcup_{\ell \in L} X^\ell \times T_i^\ell \rightarrow \mathbb{R}$. From this formulation it is clear that we focus on private payoff types.

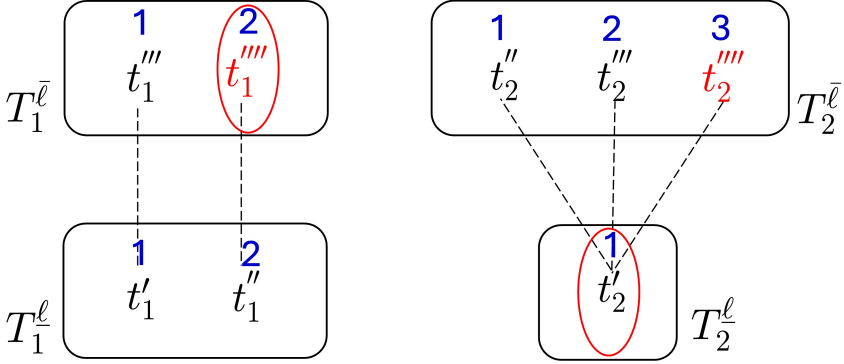
To facilitate a commonly used notion of efficiency, we assume that each agent's utility function is quasilinear. I.e., for each $i \in I$, $u_i(x, t_i) := v_i(x_0, t_i) + x_i$ for $x = (x_0, x_1, \dots, x_{|I|}) \in X^\ell := X_0^\ell \times \mathbb{R}^{|I|}$ and $t_i \in T_i^\ell$, $\ell \in L$. As usual, x_0 describes the physical properties of the outcome while $(x_i)_{i \in I}$ represents the vector of transfers made to agents.

We denote the outcome function by $f_0 : \mathcal{T} \rightarrow X_0$. That is, $f_0(t)$ is the physical outcome prescribed by the outcome function f_0 to type profile t . We require that for any $t \in \mathcal{T}$, $f_0(t) \in X_0^{\check{\lambda}(t)}$. That is, the social planner can use joint awareness to select the outcome/allocation. If the reported payoff type profile is t , then the planner's awareness is $\check{\lambda}(t)$ and hence outcomes in $X_0^{\check{\lambda}(t)}$ are selected.⁴ The formulation also makes clear that the social planner cannot necessarily describe the outcome function to agents in advance if the social planner is unaware of something herself. However, we assume that the social planner can commit to abstract properties of outcome functions such as efficiency.

We generalize utilitarian ex-post efficiency of outcome functions to unawareness by requiring it for every awareness level.

⁴Typically the planner is not considered as an agent in mechanism design and hence we do not explicitly consider the awareness that the planner may have. Yet, all of our results remain intact when the planner joins her awareness together with the awareness of all agents whenever communicating the pooled awareness level in the dynamic elaboration mechanisms introduced in the next section.

FIGURE 2. Payoff Types in Example 2



DEFINITION 1. The outcome function f_0 is *utilitarian ex-post efficient* if for all $\ell \in L$ and $\mathbf{t} = (t_i)_{i \in I} \in \mathbf{T}^\ell$,

$$\sum_{i \in I} v_i(f_0(\mathbf{t}), t_i) \geq \sum_{i \in I} v_i(x_0, t_i) \text{ for all } x_0 \in X_0^\ell. \quad (1)$$

2.2 Failure of Static VCG Mechanisms

In standard mechanism design with quasilinear utilities and private values but without asymmetric unawareness, efficient implementation is almost equivalent to the use of Vickrey-Clark-Groves (VCG) mechanisms. We demonstrate that VCG mechanisms fail to implement efficiently under unawareness, quasilinear utilities, and private values.

Example 2 Consider two agents, 1 and 2. There are two awareness levels, $\bar{\ell}$ and $\underline{\ell}$ with $\bar{\ell} \succ \underline{\ell}$. The payoff type structure is given by Figure 2. For each agent, there are two payoff type spaces, one each corresponding to the two awareness levels. The left payoff types spaces belong to agent 1 while the right ones to agent 2. For each agent's type structure, projections from the larger payoff type space to the smaller one are indicated by dashed lines. For each payoff type, the subscript indices the agent. The superscripts are arbitrary and are made in order to distinguish different payoff types of an agent.

There is a single object to be allocated. For simplicity, we assume that the value of not getting the object is always zero. The value of getting the object depends on the payoff type. It is written in blue above each payoff type, representing v_i . Consider now a payoff type profile (t_1''', t_2''') and awareness levels $\ell_1 = \bar{\ell}$ and $\ell_2 = \underline{\ell}$ for agents 1 and 2, respectively. In Figure 2, we print the true payoff types in red. Yet, since agents differ in their awareness, we also indicate their corresponding perceived payoff types by red ovals. That is, agent 1 perceives his payoff type to be equal to his actual payoff type t_1''' , while agent 2 perceives his payoff type to be t_2' because his awareness level is $\underline{\ell}$.

Consider running a standard Vickrey second price auction in which the highest bidder wins the object and pays the second highest bid. In the dominant strategy equilibrium of the auction, agents bid truthfully their own perceived payoff type. That is, agent

1 bids 2 and agent 2 bids 1. Consequently, agent 1 obtains the object and pays a price of 1. His (net) utility is 1. Note that awareness is not raised (since agents just submit bids consisting of a value for the object). Clearly, the auction is not efficient since it does not allocate the object to the bidder with the highest true value for the object, which would be bidder 2.

If bidder 1 would have raised awareness of $\bar{\ell}$, then bidder 2 would realize that her value is 3. In such a case, bidder 2 might have bid 3, obtained the object, and paid a price of 2. This leaves bidder 1 with a utility 0, which is worse than his utility of 1 in the Vickrey auction. Thus, in the second price auction, bidder 1 has no incentive to raise bidder 2's awareness.

2.3 Dynamic Direct Elaboration Mechanisms

In order to pool awareness among agents and allow agents to provide information on issues they are or became aware, we allow for communication not just *from* agents but also *to* agents. That is, we envision a mediator who receives messages about payoff types from agents, pools the awareness contained in those messages, and sends back messages to agents that potentially raise the awareness of agents to the pooled awareness level. Subsequently, agents may want to elaborate on their prior messages at least to the details of the pooled awareness level. This procedure may be repeated till no agent wants to elaborate further. To this end, we introduce a new class of dynamic mechanisms.

DEFINITION 2 (Dynamic Direct Elaboration Mechanism). The *dynamic direct elaboration mechanism* implementing outcome function f_0 is defined recursively by the following algorithm:⁵

Stage $n = 1$: Each agent $i \in I$ must report a type t_i^1 .

Stages $n > 1$: Each agent i must report a type $t_i^n \in (t_i^{n-1})^\uparrow \cap \left(T_i^{\check{\lambda}(t^{n-1})}\right)^\uparrow$.

Stop: If $t_i^{n+1} = t_i^n$ for all $i \in I$, then $f_0(t^{n+1})$ is implemented.

At the first stage, every agent reports a type. At later stages, agents can elaborate on their prior reported types. When agents report the payoff type profile t^{n-1} in stage $n - 1$, the mediator awareness level is the pooled awareness level $\check{\lambda}(t^{n-1})$. Consequently, he communicates back the pooled awareness level $\check{\lambda}(t^{n-1})$ to all agents.⁶ Importantly, all agents receive the same message from the mediator. That is, we can use public messages. At the next stage, each agent i must report a type at least at the pooled awareness level $T_i^{\check{\lambda}(t^{n-1})}$. This report must be consistent with her prior reported type which is implied by the requirement $(t_i^{n-1})^\uparrow \cap \left(T_i^{\check{\lambda}(t^{n-1})}\right)^\uparrow$. For instance, if an agent reported a type

⁵Transfers can be arbitrary at this point and will be specified later when we consider particular dynamic direct elaboration mechanisms.

⁶The current formulation presumes that the mediator, mechanism designer, or social planner has no relevant awareness herself. If she does have relevant awareness, she can incorporate it easily into the message communicated back to agents.

whose awareness is strictly below the pooled awareness of all agents' reports, then she must now report a type at least at the pooled awareness level $\check{\lambda}(t_i^{n-1})$ that is consistent with her prior reported type, i.e., a type in $(t_i^{n-1})^\uparrow$. However, it could be the case that she has previously pretended to have much less awareness so that the pooled awareness level did not yet incorporate all her awareness. Thus, we allow her to report a type with awareness strictly larger than the pooled awareness level. This motivates the requirement $t_i^n \in (t_i^{n-1})^\uparrow \cap \left(T_i^{\check{\lambda}(t^{n-1})}\right)^\uparrow$ rather than $t_i^n \in (t_i^{n-1})^\uparrow \cap T_i^{\check{\lambda}(t^{n-1})}$. If an agent did report at the pooled awareness level in the previous stage, then she is not able to change her report in the current stage. The mechanism stops once no agents revise their reports any further. Once it stops, it implements the physical outcome associated to the final reported type profile by f_0 .

Example 1 (Continued) To continue our Example 1 above, suppose that agent 1 reported in stage 1 payoff type $t_1^1 = t_1$ and agent 2 reported $t_2^1 = t_2$. Pooling awareness leads to $\check{\lambda}(t_1, t_2) = \{a, b, c\}$. At stage 2, both agents must now elaborate their prior reported type and report a type in $T_i^{\{a, b, c\}}$. For instance, agent 1 could report $t_1^2 = t_1'$. Since no further awareness can be revealed after stage 2, the mechanism must conclude after stage 3. $\square$

Note that awareness and information is transmitted both from agents to the mediator but also from the mediator to the agents. Thus, awareness of agents may change endogenously when interacting in the mechanism. Since agents report directly types and elaborate in later stages on their prior reported types, we call it a “direct elaboration” mechanism. Note that by the definition of join of the lattice of spaces, $\check{\lambda}(t^{n-1}) \supseteq \lambda(t_i^{n-1})$ for any $i \in I$.

The mechanism stops when no agent wants to further elaborate on her type. Clearly, since L is finite, the mechanism must stop at some finite stage n . Moreover, when it stops at n , then $t^n \in T^{\check{\lambda}(t^n)}$. That is, all agents must have reported twice in a row types at the *same* awareness level.

The dynamic direct mechanism induces a game with unawareness in extensive form à la [Heifetz et al. \(2013a\)](#) and [Schipper \(2021\)](#). Compared to standard games in extensive form, extensive-form games with unawareness feature a forest of game trees with exactly one tree for each level of awareness $\ell \in L$. Moreover, the information set of a player at a history in a tree associated with awareness level ℓ may be located in a tree associated with a smaller awareness level $\ell' \leq \ell$. For each awareness level $\ell \in L$, in the corresponding tree, nature moves first drawing for each agent i a payoff type in T_i^ℓ and an awareness level in $L(\ell)$. Across these trees, the draws must be consistent as outlined in the section on payoff types. That is, if payoff type $\bar{t}_i \in T_i^{\bar{\ell}}$ and awareness level ℓ_i are drawn in the tree associated with $\bar{\ell}$, then payoff type $r_i^{\bar{\ell}}(\bar{t}) \in T_i^{\ell}$ and awareness level $\ell_i \wedge \ell \in L(\ell)$ are drawn in the tree associated with ℓ . In any game tree of the forest, the move of nature is followed by histories created by the play of the dynamic direct elaboration mechanism. For any $\ell \in L$, the ℓ -partial game consists of all game trees associated with awareness levels in $L(\ell)$. The idea is that an agent with awareness level ℓ can reason about opponents having awareness associated with any level in $L(\ell)$. At any point during the play, the game perceived by the agent with an awareness level $\ell \in L$ is the ℓ -partial game.

To describe information sets, focus first on initial information sets. If nature draws payoff type $\bar{t}_i \in T_i^{\bar{\ell}}$ and awareness level ℓ_i for agent i , his initial actual information set is $h_i^1 = r_{\ell_i}^{\bar{\ell}}(\bar{t}_i)$. Here $r_{\ell_i}^{\bar{\ell}}(\bar{t}_i)$ is agent i 's initially perceived type. Notice that in this case h_i^1 is an information set in the tree associated with awareness level ℓ_i . That is, in games with unawareness, the information set of a player at a history in a tree may be an object of a “less expressive” tree that misses aspects of which the agent is unaware.

In correspondence to the draw of nature in the full game, for any $\ell \in L$ in the ℓ -partial game, nature draws payoff type $t_i = r_{\ell}^{\bar{\ell}}(\bar{t}_i) \in T_i^{\ell}$ and awareness level $\ell_i \wedge \ell \in L(\ell)$ for agent i . The initial information set of agent i in the ℓ -partial game is $h_i^1 = r_{\ell_i \wedge \ell}^{\bar{\ell}}(\bar{t}_i) = r_{\ell_i \wedge \ell}^{\ell}(t_i)$. Here $r_{\ell_i \wedge \ell}^{\ell}(t_i)$ is agent i 's initially perceived type in the ℓ -partial game. (Recall that ℓ_i is agent i 's true initial awareness level selected by nature. It could be that $\ell_i \not\leq \ell$. In such a case, the meet notation $\ell_i \wedge \ell$ becomes crucial when specifying the agent's awareness level in the ℓ -partial game as discussed when introducing the payoff type spaces.) Note that even though $t_i \in T_i^{\ell}$ and thus it is a payoff type in the tree associated with awareness level ℓ , the corresponding information set and initially perceived payoff type by agent i , $h_i^1 = r_{\ell_i \wedge \ell}^{\ell}(t_i)$, is a payoff type in the tree associated with awareness level $\ell_i \wedge \ell$ which, by definition of the meet, is necessarily weakly less than ℓ . Again, in games with unawareness the information set associated to a history in one tree may be an object of a less expressive tree.

Given we have defined all initial information sets, we turn to second stage information sets. For any $i \in I$, fix any initial information set $h_i^1 = t_i$. A successor of h_i^1 is defined by $h_i^2 = (t'_i, \mathbf{t}^1)$ with $t'_i \in \left(r_{\lambda(t_i)}^{\lambda(t_i) \vee \check{\lambda}(\mathbf{t}^1)} \right)^{-1}(t_i)$ and $\mathbf{t}^1 = (t_i^1, \mathbf{t}_{-i}^1)$ with $\lambda(t_i^1) \leq \lambda(t_i)$. That is, agent i becoming aware in information set h_i^2 of $\check{\lambda}(\mathbf{t}^1)$, joins it with her own awareness $\lambda(t_i)$, and now considers a more elaborate payoff type $t'_i \in \left(r_{\lambda(t_i)}^{\lambda(t_i) \vee \check{\lambda}(\mathbf{t}^1)} \right)^{-1}(t_i)$. Of course, when communicating at the first stage to the mediator, she could only communicate a payoff type involving awareness not greater than her perceived type at that time t_i , i.e., $\lambda(t_i^1) \leq \lambda(t_i)$.

Inductively, for any agent $i \in I$ and $n > 2$, fix a $n - 1$ stage information set $h_i^{n-1} = (t_i, (\mathbf{t}^k)_{k < n-1})$. A successor of h_i^{n-1} is defined by $h_i^n = (t'_i, (\mathbf{t}^k)_{k < n})$ with $t'_i \in \left(r_{\lambda(t_i)}^{\lambda(t_i) \vee \check{\lambda}(\mathbf{t}^{n-1})} \right)^{-1}(t_i)$, $\mathbf{t}^{n-1} = (t_i^{n-1}, \mathbf{t}_{-i}^{n-1})$ with $\lambda(t_i^{n-1}) \leq \lambda(t_i)$ and $t_j^{n-1} \in (t_j^{n-2})^\uparrow \cap (T_j^{\check{\lambda}(\mathbf{t}^{n-2})})^\uparrow$ for all $j \in I$. The interpretations of the conditions are analogous to the ones for h_i^2 . The additional condition $t_j^{n-1} \in (t_j^{n-2})^\uparrow \cap (T_j^{\check{\lambda}(\mathbf{t}^{n-2})})^\uparrow$ for all $j \in I$ just says that at the previous stage all agents must have reported a (weak) elaboration of the payoff type reported at the stage before the previous stage since this is implied by the dynamic direct elaboration mechanism. We write $h_i^{n-1} \rightsquigarrow h_i^n$ to indicate that h_i^n is a successor of h_i^{n-1} .

Since L is finite, the mechanism must stop after some finite number of stages. Denote the set of agent i 's information sets by H_i .

While we defined each information set essentially as a history, most versions of the dynamic direct elaboration mechanisms we will discuss below will not make use of all this information. Our mechanisms will only make use of the awareness embodied in the

reported payoff types profiles but not of the precise payoff type profiles reported along the way except for the last one. Yet, writing information sets with sequences of reported payoff type profiles allows us to keep a uniform notation for all mechanisms we discuss and is notationally simpler than alternatives.⁷

We can associate with each information set $h_i \in H_i$ an awareness level. When $h_i^1 = t_i$, agent i has awareness level $\lambda(t_i)$. This motivates an abuse of notation by letting $\lambda(h_i^1) := \lambda(t_i)$ defined by $h_i^1 = t_i$. At later stages $n > 1$, when $h_i^n = (t_i, (t^k)_{k < n})$ occurs, the agent has awareness level $\lambda(t_i)$. (Recall that by the definition of information sets, t_i is already the elaboration of agent i 's payoff type in light of awareness raised along the sequence of reported payoff type profiles $(t^k)_{k < n}$.) Again, we abuse notation and write $\lambda(h_i^n) := \lambda(t_i)$.

For $h_i \in H_i$, let $t_i(h_i) = t_i$ be defined by $h_i^1 = t_i$ if $h_i = h_i^1$ and by $h_i^n = (t_i, (t^k)_{k < n})$ if $h_i = h_i^n$. This is the payoff type perceived by agent i in information set h_i .

Since an agent may not necessarily be aware of everything ex-ante, she cannot necessarily envision all of her information sets and choose a strategy that prescribes an action to any of her information sets. At any information set $h_i \in H_i$, agent i 's perception of the game is the $\lambda(h_i)$ -partial game. For any $\ell \in L$, the set of information sets perceived by her in the ℓ -partial game is $H_i^\ell := \{h_i \in H_i : \lambda(h_i) \leq \ell\}$.

Having specified the representation of information sets in the game with unawareness induced by the dynamic direct elaboration mechanism, we proceed with defining strategies. As standard in game theory, a strategy of an agent assigns to each of her information sets an action. More formally, a (pure) *strategy* of agent i in the game induced by the dynamic direct elaboration mechanism is a mapping $\sigma_i : H_i \rightarrow \mathcal{T}_i$ such that

- (i) for all $h_i^1 \in H_i$, $\sigma_i(h_i^1) \in \bigcup_{\ell \in L(\lambda(h_i))} \mathcal{T}_i^\ell$,
- (ii) for $n \in \mathbb{N}$ and $h_i^n \in H_i$, $\sigma_i(h_i^n) \in (t_i^{n-1})^\uparrow \cap (T_i^{\lambda(h_i^n)})^\uparrow$, where t_i^{n-1} is agent i 's previous period's reported payoff type according to h_i^n .

Property (i) applies to all initial information sets of agent i . It says that the agent can report any type that she is aware of at her perceived true type. Property (ii) is a constraint imposed by the dynamic direct elaboration mechanism as it allows only elaborations of the previously reported type.⁸ An agent can only provide elaborations of her payoff type that she is aware of herself at the information set. She has to elaborate at least at the pooled awareness level ℓ . Yet, she is also free to elaborate at an even greater awareness level.

For any awareness level $\ell \in L$ and strategy σ_i , an ℓ -partial strategy σ_i^ℓ is the strategy σ_i restricted to information sets in H_i^ℓ . It assigns an action to each information set of agent i in any game tree of the ℓ -partial game induced by the mechanism.

We denote by Σ_i the set of agent i 's strategies and by Σ_i^ℓ the set of agent i 's ℓ -partial strategies for $\ell \in L$. We denote by $\Sigma_{-i} := \times_{j \in I \setminus \{i\}} \Sigma_j$, $\Sigma_{-i}^\ell := \times_{j \in I \setminus \{i\}} \Sigma_j^\ell$, $\Sigma := \times_{i \in I} \Sigma_i$ and $\Sigma^\ell := \times_{i \in I} \Sigma_i^\ell$ the set of strategy profiles of agent i 's opponents, the set of ℓ -partial

⁷A more accurate notation for information sets would be $h_i^n = (t_i, (t_i^k, \tilde{\lambda}(t^k))_{k < n})$.

⁸This constraint is not necessary for our results. However, it is a natural simplifying assumption on communication, reducing the set of messages over which each agent optimizes and reducing the size of game trees.

strategy profiles of agent i 's opponents, the set of strategy profiles, and the set of ℓ -partial strategy profiles, respectively.

For agent i , the *truth-telling and elaboration strategy* σ_i^* is defined by:

- For any $h_i^1 \in H_i$, $h_i^1 = t_i$ implies $\sigma_i^*(h_i^1) = t_i$.
- For any $n > 1$, $h_i^n \in H_i$, $h_i^n = (t_i, (t^k)_{k < n})$ implies $\sigma_i^*(h_i^n) = t_i$.

That is, at each information set, agent i tells her true payoff type as currently perceived at this information set. (Note that at $h_i^n = (t_i, (t^k)_{k < n})$, payoff type t_i incorporates her awareness associated with t^{n-1} .)

For any agent $i \in I$, awareness level $\ell \in L$, “initial” payoff type profile $t = (t_j)_{j \in I} \in \mathbf{T}^\ell$, “initial” profile of awareness levels $\ell = (\ell_j)_{j \in I} \in L(\ell)^{|I|}$, and profile of strategies $\sigma = (\sigma_j)_{j \in I} \in \Sigma$, we let $H_i(t, \ell, \sigma)$ denote the set of agent i 's histories reached with t , ℓ , and σ (in the ℓ -partial game). It is defined inductively as follows:

$$h_i^1 \in H_i(t, \ell, \sigma) \text{ if } h_i^1 = r_{\ell_i}^\ell(t_i)$$

$$h_i^2 \in H_i(t, \ell, \sigma) \text{ if } h_i^2 = \left(r_{\ell_i \vee \tilde{\lambda}((\sigma_j(h_j^1))_{j \in I})}^\ell(t_i), (\sigma_j(h_j^1))_{j \in I} \right) \text{ for } h_j^1 \in H_j(t, \ell, \sigma)$$

and for $n > 1$,

$$h_i^n \in H_i(t, \ell, \sigma) \text{ if } h_i^n = \left(r_{\ell_i \vee \tilde{\lambda}((\sigma_j(h_j^{n-1}))_{j \in I})}^\ell(t_i), (t^k)_{k < n-1}, (\sigma_j(h_j^{n-1}))_{j \in I} \right) \text{ for } h_j^{n-1} \in H_j(t, \ell, \sigma) \text{ and } h_i^{n-1} \rightsquigarrow h_i^n.$$

The initial information set is agent i 's perceived payoff type in the ℓ -partial game. At successive information sets, the initial payoff type is elaborated in light of the pooled awareness raised in the reports induced by the strategy profile at their prior reached information sets. The elaboration is consistent with the initial payoff type profile selected in $\mathbf{T}^\ell$. This feature implies that $H_i(t, \ell, \sigma)$ consists of a *unique* path (a unique sequence) of information sets for agent i . The requirement $h_i^{n-1} \rightsquigarrow h_i^n$ implies that the sequence $(t^k)_{k < n-1}$ in the definition of h_i^n is identical to the sequence of payoff type profiles in h_i^{n-1} .

For any agent $i \in I$, awareness level $\ell \in L$, “initial” payoff type profile $t = (t_j)_{j \in I} \in \mathbf{T}^\ell$, “initial” profile of awareness levels $\ell = (\ell_j)_{j \in I} \in L(\ell)^{|I|}$, and strategy $\sigma_i \in \Sigma_i$, we let $H_i(t, \ell, \sigma_i) := \bigcup_{\sigma_{-i} \in \Sigma_{-i}} H_i(t, \ell, \sigma_i, \sigma_{-i})$. Similarly, we let $H_i(\sigma_i) := \bigcup_{(t, \ell) \in \bigcup_{\ell \in L} \mathbf{T}^\ell \times L(\ell)^{|I|}} H_i(t, \ell, \sigma_i)$.

For every $\ell \in L$, any “initial” payoff type profile $t = (t_j)_{j \in I} \in \mathbf{T}^\ell$, “initial” profile of awareness levels $\ell = (\ell_j)_{j \in I} \in L(\ell)^{|I|}$, and profile of strategies $\sigma = (\sigma_j)_{j \in I} \in \Sigma$ give rise to a *unique* sequence of payoff type profiles $\tau(t, \ell, \sigma) = (t^1, \dots, t^{n-1}, t^n)$ (for some $n > 2$) defined by $t^k = (\sigma_j(h_j^k))_{j \in I}$ with $h_j^k \in H_j(t, \ell, \sigma)$ for $j \in I$, $k \leq n$, and $t^{n-1} = t^n$. This sequence is unique because as we observed above h_j^1 is unique for each $j \in I$ given t, ℓ , and σ , h_j^2 is unique for each $j \in I$ given t, ℓ , and σ , etc. The requirement $t^{n-1} = t^n$ means that $\tau(t, \ell, \sigma)$ consists of the entire sequence of payoff type profiles until the dynamic direct elaboration mechanisms stops. For the implementation of outcomes, we

are especially interested in the last profile of payoff types of any such sequence and denote it by $\tau^*(t, \ell, \sigma)$. The set of all sequences is denoted by $\mathfrak{T} := \{\tau(t, \ell, \sigma) : \sigma \in \Sigma, (t, \ell) \in \bigcup_{\ell \in L} (\mathbf{T}^\ell \times L(\ell)^{|I|})\}$

For any agent $i \in I$ and information set $h_i \in H_i$, define

$$\Sigma_i(h_i) := \left\{ \sigma_i \in \Sigma_i : \exists (t, \ell, \sigma_{-i}) \in \left(\bigcup_{\ell \in L} (\mathbf{T}^\ell \times L(\ell)^{|I|}) \right) \times \Sigma_{-i} (h_i \in H_i(t, \ell, (\sigma_i, \sigma_{-i}))) \right\}.$$

This is the set of agent i 's strategies that are consistent with information set h_i .

For any agent $i \in I$, let $f_i : \mathfrak{T} \rightarrow \mathbb{R}$ denote the transfer paid to agent i in the dynamic direct elaboration mechanisms. These transfers will depend on the precise versions of the direct elaboration mechanism studied below. In contrast to the outcome function f_0 we allow transfers to depend on the entire sequence of reported type profiles. The reason is that it will become important who first reported which awareness level.

We let the social choice function (i.e., the outcome function and transfers) be denoted by $f : \mathfrak{T} \rightarrow \mathbb{R}$ and defined by $f = (f_0, f_1, \dots, f_{|I|})$. That is, for any strategy profile $\sigma \in \Sigma$, $\ell \in L$, initial profiles of payoff types $t \in \mathbf{T}^\ell$, initial profiles of awareness levels $\ell \in L(\ell)^{|I|}$, $f(\tau(t, \ell, \sigma)) = (f_0(\tau(t, \ell, \sigma)), f_1(\tau(t, \ell, \sigma)), \dots, f_{|I|}(\tau(t, \ell, \sigma)))$.

DEFINITION 3. The dynamic direct elaboration mechanism truthfully implements the social choice function f in conditional dominant strategies if for all agents $i \in I$, information sets $h_i \in H_i(\sigma_i^*)$, opponents' strategy profiles $\sigma_{-i} \in \Sigma_{-i}$, initial profiles of payoff types $t \in \mathbf{T}^{\lambda(h_i)}$ (as perceived by i in h_i), and initial profiles of awareness levels $\ell \in L(\lambda(h_i))^{|I|}$ (as perceived by i in h_i) such that $h_i \in H_i(t, \ell, (\sigma_i^*, \sigma_{-i}))$,

$$u_i(f(\tau(t, \ell, (\sigma_i^*, \sigma_{-i}))), t_i(h_i)) \geq u_i(f(\tau(t, \ell, (\sigma_i, \sigma_{-i}))), t_i(h_i)) \quad (2)$$

for all $\sigma_i \in \Sigma_i(h_i)$.

Conditional dominance strengthens dominance by requiring each agent not only to select a (partial) strategy that is ex-ante dominant but also dominant conditional on each information set. This becomes important when agents cannot anticipate all information sets ex-ante and thus cannot select ex-ante a strategy for the entire game in extensive form.

DEFINITION 4. A outcome function f_0 is truthfully implemented at the pooled awareness level if for any $\bar{t} \in \mathbf{T}^{\bar{\ell}}$ and $(\ell_i)_{i \in I} \in L^{|I|}$, $f_0 \left(r_{(\bigvee_{i \in I} \ell_i)}^{\bar{\ell}}(\bar{t}) \right)$ is implemented.

Recall that $(\bar{t}, (\ell_i)_{i \in I}) \in \mathbf{T}^{\bar{\ell}} \times L^{|I|}$ represents the move of nature selecting both actual payoff types and awareness levels for all agents. Consequently, $\bigvee_{i \in I} \ell_i$ represents the pooled awareness level.

3. EFFICIENT IMPLEMENTATION

To implement efficiently at the pooled awareness level, we specify for each agent $i \in I$ the transfers f_i that incentivize both the revelation of information and raising awareness. Revelation of information is achieved via VCG transfers. Raising awareness requires an extra term in the transfer functions.

DEFINITION 5 (Dynamic Elaboration VCG Mechanism). We say that the dynamic direct elaboration mechanism implementing f is a dynamic elaboration VCG mechanism if f_0 is utilitarian ex-post efficient and transfers f_i to agent $i \in I$ are given by for any $(t^1, \dots, t^n) \in \mathfrak{T}$,

$$f_i(t^1, \dots, t^n) := \sum_{j \neq i} v_j(f_0(t^n), t_j^n) + y_i^{\check{\lambda}(t^n)}(t_{-i}^n) + a_i(t^1, \dots, t^n), \quad (3)$$

where, for each ℓ , $y_i^\ell : \mathbf{T}_{-i}^\ell \rightarrow \mathbb{R}$ is an arbitrary function and $a_i(\cdot)$ is defined by:

$$a_i(t^1, \dots, t^n) := \begin{cases} m_i(\check{\lambda}(t^n)) & \text{if } i = i^*(t^1, \dots, t^n) \\ -\frac{1}{|I|-1} m_j(\check{\lambda}(t^n)) & \text{if } j = i^*(t^1, \dots, t^n) \neq i \\ 0 & \text{otherwise} \end{cases} \quad (4)$$

where

$$i^*(t^1, \dots, t^n) := \left\{ i \in I : \exists k \leq n \left((\lambda(t_i^k) = \check{\lambda}(t^n)) \text{ and } \nexists j \neq i, k' \leq k (\lambda(t_j^{k'}) = \check{\lambda}(t^n)) \right) \right\}$$

and $m_i(\ell)$ is defined recursively, as follows: $m_i(\underline{\ell}) := 0$ and for any $\ell \triangleright \underline{\ell}$,

$$m_i(\ell) := \max_{\ell' \triangleleft \ell, t' \in \mathbf{T}^{\ell'}, t \in \mathbf{T}^\ell} \left\{ \left(m_i(\ell') + v_i(f_0(t'), t_i) + \sum_{j \neq i} v_j(f_0(t'), t_j') \right. \right. \\ \left. \left. + y_i^{\ell'}(t'_{-i}) - \sum_j v_j(f_0(t), t_j) - y_i^\ell(t_{-i}) \right), 0 \right\}. \quad (5)$$

The mechanism can be viewed as a dynamic version of the Vickrey-Clarke-Groves (VCG) mechanisms (Groves (1973), Groves and Loeb (1975)). Note that the transfers can be described without necessarily being aware of all payoff type profiles. The mechanism designer commits to implement a utilitarian ex-post efficient outcome given the agents' final reports of payoff types. Each agent is paid the total welfare of others given the final reports of payoff types plus a term that depends on only the opponents' payoff types (and the final pooled awareness level) and additionally a term incentivizing raising of awareness. This last term does not only depend on the final reported payoff type profile but on the sequence of reported payoff type profiles. It matters who reports the pooled awareness level first. Note that $i^*(t^1, \dots, t^n) = \emptyset$ if there is no agent who reports the pooled awareness level (which can happen if the individual awareness levels are incomparable) or if several agents simultaneously report the pooled awareness level first.

If i is the unique agent who first reports the pooled awareness level, then a_i receives a payment related to the cost that raising awareness could impose on her utility from the mechanism, net of the $a_i(\cdot)$ term itself.

THEOREM 1. *The dynamic elaboration VCG mechanism truthfully implements in conditionally dominant strategies a utilitarian ex-post efficient outcome under pooled awareness.*

PROOF. Using equations (3), (4), and (5) for transfers, rewrite the defining inequality (2) of conditional dominant strategy implementation as follows: For all agents $i \in I$, information sets $h_i \in H_i(\sigma_i^*)$, opponents' strategy profiles $\sigma_{-i} \in \Sigma_{-i}$, initial profiles of payoff types $t \in T^{\lambda(h_i)}$ (as perceived by i at h_i), and initial profiles of awareness levels $\ell \in L(\lambda(h_i))^{|I|}$ (as perceived by i at h_i) such that $h_i \in H_i(t, \ell, (\sigma_i^*, \sigma_{-i}))$,

$$\begin{aligned}
& v_i(f_0(\tau^*(t, \ell, (\sigma_i^*, \sigma_{-i}))), t_i(h_i)) + \sum_{j \neq i} v_j(f_0(\tau^*(t, \ell, (\sigma_i^*, \sigma_{-i}))), \tau_j^*(t, \ell, (\sigma_i^*, \sigma_{-i}))) \\
& + y_i^{\check{\lambda}(\tau^*(t, \ell, (\sigma_i^*, \sigma_{-i})))}(\tau_{-i}^*(t, \ell, (\sigma_i^*, \sigma_{-i}))) + a_i(\tau(t, \ell, (\sigma_i^*, \sigma_{-i}))) \\
& \geq v_i(f_0(\tau^*(t, \ell, (\sigma_i, \sigma_{-i}))), t_i(h_i)) + \sum_{j \neq i} v_j(f_0(\tau^*(t, \ell, (\sigma_i, \sigma_{-i}))), \tau_j^*(t, \ell, (\sigma_i, \sigma_{-i}))) \\
& + y_i^{\check{\lambda}(\tau^*(t, \ell, (\sigma_i, \sigma_{-i})))}(\tau_{-i}^*(t, \ell, (\sigma_i, \sigma_{-i}))) + a_i(\tau(t, \ell, (\sigma_i, \sigma_{-i}))) \tag{6}
\end{aligned}$$

for all $\sigma_i \in \Sigma_i(h_i)$.

We need to show that this inequality holds in the following cases:

Case 1: $h_i = h_i^n \in H_i(\sigma_i^*)$ such that $\nexists j \neq i, k \leq n$ with $\lambda(t_j^k) = \lambda(h_i^n)$. That is, at stage n , nobody raised awareness to $\lambda(h_i^n)$ in any previous stage and nobody other than i raises it at the current stage n (and i can raise it at the current stage with the truth-telling strategy).

Case 2: $h_i = h_i^n \in H_i(\sigma_i^*)$ such that $\nexists j \neq i, k < n$ with $\lambda(t_j^k) = \lambda(h_i^n)$ but $\exists j \neq i$ such that $\lambda(t_j^n) = \lambda(h_i^n)$. That is, at stage n , nobody raised awareness to $\lambda(h_i^n)$ at any previous stage but there exists agent $j \neq i$ who raises awareness to $\lambda(h_i^n)$ in the current stage as well.

Case 3: $h_i = h_i^n \in H_i(\sigma_i^*)$ such that $\exists j \in I$ (possibly $j = i$) and $k < n$ such that $\lambda(t_j^k) = \lambda(h_i^n)$. That is, at stage n , somebody has raised awareness to $\lambda(h_i)$ already in some earlier stage.

We deal with the cases one-by-one.

Case 1: We have $i = i^*(\tau(t, \ell, (\sigma_i^*, \sigma_{-i})))$ and $a_i(\tau(t, \ell, (\sigma_i^*, \sigma_{-i}))) = m_i(\lambda(h_i))$.

Case 1a - Deviations within the same awareness level: $\sigma_i(h_i) \in T_i^{\lambda(h_i)}$. By Definition 2, $\check{\lambda}(\tau^*(t, \ell, (\sigma_i^*, \sigma_{-i}))) = \check{\lambda}(\tau^*(t, \ell, (\sigma_i, \sigma_{-i}))) = \lambda(h_i)$. It follows that in this case $i = i^*(\tau(t, \ell, (\sigma_i^*, \sigma_{-i}))) = i^*(\tau(t, \ell, (\sigma_i, \sigma_{-i})))$ and $a_i(\tau(t, \ell, (\sigma_i^*, \sigma_{-i}))) = a_i(\tau(t, \ell, (\sigma_i, \sigma_{-i})))$. Thus, these terms cancel out at both sides of inequality (6). Since the remaining terms

represent utilitarian ex-post welfare as anticipated at h_i , plus a term that does not depend on i 's report within the same awareness level, inequality (6) is now implied by utilitarian ex-post efficiency of the outcome function f_0 , i.e., inequality (1).

Case 1b - Deviations to lower awareness levels: $\sigma_i(h_i) \in T_i^\ell$ for some $\ell \triangleleft \lambda(h_i)$. We can rewrite inequality (6) as

$$\begin{aligned}
& m_i(\check{\lambda}(\tau^*(\mathbf{t}, \ell, (\sigma_i^*, \sigma_{-i})))) - a_i(\tau(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i}))) \\
& \geq v_i(f_0(\tau^*(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i}))), t_i(h_i)) + \sum_{j \neq i} v_j(f_0(\tau^*(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i}))), \tau_j^*(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i}))) \\
& \quad + y_i^{\check{\lambda}(\tau^*(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i})))}(\tau_{-i}^*(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i}))) \\
& \quad - v_i(f_0(\tau^*(\mathbf{t}, \ell, (\sigma_i^*, \sigma_{-i}))), t_i(h_i)) - \sum_{j \neq i} v_j(f_0(\tau^*(\mathbf{t}, \ell, (\sigma_i^*, \sigma_{-i}))), \tau_j^*(\mathbf{t}, \ell, (\sigma_i^*, \sigma_{-i}))) \\
& \quad - y_i^{\check{\lambda}(\tau^*(\mathbf{t}, \ell, (\sigma_i^*, \sigma_{-i})))}(\tau_{-i}^*(\mathbf{t}, \ell, (\sigma_i^*, \sigma_{-i}))) \tag{7}
\end{aligned}$$

Moreover, the left hand side is greater than or equal to $m_i(\check{\lambda}(\tau^*(\mathbf{t}, \ell, (\sigma_i^*, \sigma_{-i})))) - m_i(\check{\lambda}(\tau^*(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i}))))$, since either σ_i still results in player i being the first to report the final awareness level as anticipated by i at h_i (i.e., $i = i^*(\tau(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i})))$), in which case the $a_i(\cdot)$ -term is $m_i(\check{\lambda}(\tau^*(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i}))))$, or it does not ($i \neq i^*(\tau(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i})))$), in which case the $a_i(\cdot)$ -term is non-positive. Note that $m_i(\cdot)$ is non-negative by construction. Therefore, it is sufficient to show:

$$\begin{aligned}
& \ell' \triangleleft \check{\lambda}(\tau^*(\mathbf{t}, \ell, (\sigma_i^*, \sigma_{-i}))), \quad \left(m_i(\ell') + v_i(f_0(\mathbf{t}', t'_i)) + \sum_{j \neq i} v_j(f_0(\mathbf{t}'), t'_j) + y_i^{\check{\lambda}(\mathbf{t}')}(\mathbf{t}'_{-i}) \right. \\
& \quad \left. - \sum_j v_j(f_0(\mathbf{t}''), t''_j) - y_i^{\check{\lambda}(\mathbf{t}'')}(\mathbf{t}''_{-i}) \right) - m_i(\check{\lambda}(\tau^*(\mathbf{t}, \ell, (\sigma_i^*, \sigma_{-i})))) \\
& \geq \max_{\substack{\mathbf{t}' \in T^{\check{\lambda}(\tau^*(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i})))}, \\ \mathbf{t}'' \in T^{\check{\lambda}(\tau^*(\mathbf{t}, \ell, (\sigma_i^*, \sigma_{-i})))}}} \left(v_i(f_0(\mathbf{t}'), t'_i) + \sum_{j \neq i} v_j(f_0(\mathbf{t}'), t'_j) + y_i^{\check{\lambda}(\mathbf{t}')}(\mathbf{t}'_{-i}) \right. \\
& \quad \left. - \sum_j v_j(f_0(\mathbf{t}''), t''_j) - y_i^{\check{\lambda}(\mathbf{t}'')}(\mathbf{t}''_{-i}) \right) \\
& \geq v_i(f_0(\tau^*(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i}))), t_i(h_i)) + \sum_{j \neq i} v_j(f_0(\tau^*(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i}))), \tau_j^*(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i}))) \\
& \quad + y_i^{\check{\lambda}(\tau^*(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i})))}(\tau_{-i}^*(\mathbf{t}, \ell, (\sigma_i, \sigma_{-i}))) \\
& \quad - v_i(f_0(\tau^*(\mathbf{t}, \ell, (\sigma_i^*, \sigma_{-i}))), t_i(h_i)) - \sum_{j \neq i} v_j(f_0(\tau^*(\mathbf{t}, \ell, (\sigma_i^*, \sigma_{-i}))), \tau_j^*(\mathbf{t}, \ell, (\sigma_i^*, \sigma_{-i})))
\end{aligned}$$

$$-g_i^{\check{\lambda}(\tau^*(t, \ell, (\sigma_i^*, \sigma_{-i})))}(\tau_{-i}^*(t, \ell, (\sigma_i^*, \sigma_{-i}))) \quad (8)$$

The first inequality follows from the fact that $\check{\lambda}(\tau^*(t, \ell, (\sigma_i, \sigma_{-i})))$ is in the set of awareness levels we maximize over at the l.h.s. of the inequality. The second equality follows from the fact that both $\tau^*(t, \ell, (\sigma_i^*, \sigma_{-i}))$ and $\tau^*(t, \ell, (\sigma_i, \sigma_{-i}))$ are profiles of types we maximize over at the l.h.s. of the inequality. That is, the r.h.s. is just an instance of the expression being maximized. Thus, inequality (6) holds by construction of the mechanism in this case.

Case 2a - Deviations within the same awareness level: $\sigma_i(h_i) \in T_i^{\lambda(h_i)}$. In this case, $i^*(\tau(t, \ell, (\sigma_i, \sigma_{-i}))) = i^*(\tau(t, \ell, (\sigma_i^*, \sigma_{-i}))) = \emptyset$. Thus, $a_i(\tau(t, \ell, (\sigma_i, \sigma_{-i}))) = a_i(\tau(t, \ell, (\sigma_i^*, \sigma_{-i}))) = 0$. The same argument as in Case 1a shows that truth-telling is conditionally dominant among deviations within the same awareness level.

Case 2b - Deviations to lower awareness levels: $\sigma_i(h_i) \in T_i^\ell$ for some $\ell \triangleleft \lambda(h_i)$. In this case, there exists $j \neq i$, such that $i^*(\tau(t, \ell, (\sigma_i, \sigma_{-i}))) = j$ while $i^*(\tau(t, \ell, (\sigma_i^*, \sigma_{-i}))) = \emptyset$. It follows that $a_i(\tau(t, \ell, (\sigma_i^*, \sigma_{-i}))) = 0$ while $a_i(\tau(t, \ell, (\sigma_i, \sigma_{-i}))) \leq 0$. To show inequality (6) in this case, it is therefore sufficient to show that

$$\begin{aligned} & v_i(f_0(\tau^*(t, \ell, (\sigma_i^*, \sigma_{-i}))), t_i(h_i)) + \sum_{j \neq i} v_j(f_0(\tau^*(t, \ell, (\sigma_i^*, \sigma_{-i}))), \tau_j^*(t, \ell, (\sigma_i^*, \sigma_{-i}))) \\ & + y_i^{\check{\lambda}(\tau^*(t, \ell, (\sigma_i^*, \sigma_{-i})))}(\tau_{-i}^*(t, \ell, (\sigma_i^*, \sigma_{-i}))) \\ & \geq v_i(f_0(\tau^*(t, \ell, (\sigma_i, \sigma_{-i}))), t_i(h_i)) + \sum_{j \neq i} v_j(f_0(\tau^*(t, \ell, (\sigma_i, \sigma_{-i}))), \tau_j^*(t, \ell, (\sigma_i, \sigma_{-i}))) \\ & + y_i^{\check{\lambda}(\tau^*(t, \ell, (\sigma_i, \sigma_{-i})))}(\tau_{-i}^*(t, \ell, (\sigma_i, \sigma_{-i}))) \end{aligned} \quad (9)$$

Since, in this case, some player has already raised awareness $\lambda(h_i) = \check{\lambda}(\tau^*(t, \ell, (\sigma_i^*, \sigma_{-i})))$, we have that $\check{\lambda}(\tau^*(t, \ell, (\sigma_i^*, \sigma_{-i}))) = \check{\lambda}(\tau^*(t, \ell, (\sigma_i, \sigma_{-i})))$. Therefore, inequality (9) holds by the usual argument that truth-telling is a dominant strategy in the VCG mechanism.

Case 3: In this case, the $a_i(\cdot)$ -terms on both the l.h.s. and r.h.s of inequality (6) cancel (as they have already been determined, as far as player i is aware). Moreover, player i can only report a type in $T_i^{\lambda(h_i)}$ (types expressing more awareness are infeasible at h_i and types expressing less awareness are ruled out by the requirement that the type announced in stage n must be in $(t_i^{n-1})^\uparrow \cap (T_i^{\check{\lambda}(t^{n-1})})^\uparrow$). That truth-telling is dominant now follows from the standard argument that the VCG mechanism is dominant strategy incentive compatible with constant awareness.

Finally, with respect to all three cases, recall that deviations to higher or non-comparable awareness levels are not feasible as agent i has no further awareness in h_i than $\lambda(h_i)$. Therefore, we have shown that a truth-telling and elaboration strategy is conditionally dominant in this mechanism. $\square$

The proof reveals that every agent has an incentive to raise awareness and no agent has an incentive to report a different payoff type given the awareness level. Initially agents report their awareness level and the perceived true payoff type. At a second stage,

they elaborate on their previously reported payoff type at the pooled awareness. The final stage just ends the mechanisms. Nothing is revealed in the last stage except that it becomes common knowledge that the mechanism ends. Thus, the dynamic elaboration VCG mechanism can be implemented in three stages. Each agent has even a strict incentive to be the first to raises awareness to the pooled awareness level.

PROPOSITION 1. *The utilitarian ex-post efficient outcome under pooled awareness is truth-fully implemented in conditional dominant strategies in the game induced by the dynamic elaboration VCG mechanism in at most three stages.*

4. BUDGET BALANCE

Ideally, we would like our mechanisms to satisfy further properties beyond utilitarian ex-post efficiency. We are interested in conditions for budget balance.

DEFINITION 6 (Budget Balance). We say that the dynamic direct elaboration mechanism with transfer functions $(f_i)_{i \in I}$ is ex-post budget balanced if for all $(t^1, \dots, t^n) \in \mathfrak{T}$,

$$\sum_{i \in I} f_i(t^1, \dots, t^n) = 0. \quad (10)$$

Note that the a_i -terms are budget neutral, i.e., for any $(t^1, \dots, t^n) \in \mathfrak{T}$ we have by construction, $\sum_{j \in I} a_j(t^1, \dots, t^n) = 0$. Thus, we can show that the condition for budget balance of the dynamic elaboration VCG mechanism is identical to the condition for budget balance of the standard VCG mechanism without unawareness (i.e., [Holmström \(1977\)](#)) holding for each awareness level. Awareness pooling does not impose an additional constraint on budget balance.

PROPOSITION 2. *The dynamic elaboration VCG mechanism implements the social choice function f with budget balance if and only if for every agent $i \in I$ and awareness level $\ell \in L$, there exists a function $g_i^\ell : T_{-i}^\ell \rightarrow \mathbb{R}$ such that for all $i \in I$ and $t^n = (t_i^n, t_{-i}^n) \in T^\ell$*

$$\sum_{i \in I} v_i(f_0(t^n), t_i) = \sum_{i \in I} g_i^{\check{\lambda}(t^n)}(t_{-i}^n). \quad (11)$$

PROOF. The proof is an extension of the proof in [Börger \(2015, Proposition 7.10\)](#); see also [Milgrom \(2004, pp. 53–54\)](#).

“Only if”: The proof is constructive. Using equations (3), (4), and (5) for transfers and (10) for budget balance, the dynamic elaboration VCG mechanism is budget balanced if for all $i \in I$ and $(t^1, \dots, t^n) \in \mathfrak{T}$,

$$\sum_{i \in I} \left(\sum_{j \neq i} v_j(f_0(t^n), t_j^n) + y_i^{\check{\lambda}(t^n)}(t_{-i}^n) + a_i(t^1, \dots, t^n) \right) = 0. \quad (12)$$

Since the a_i -terms are budget neutral by design, this equation is equivalent to

$$\sum_{i \in I} \left(\sum_{j \neq i} v_j(f_0(\mathbf{t}^n), t_j^n) + y_i^{\check{\lambda}(\mathbf{t}^n)}(t_{-i}^n) \right) = 0 \quad (13)$$

which implies that budget balance does not depend on the entire sequence $(\mathbf{t}^1, \dots, \mathbf{t}^n)$ of reported payoff type profiles but just on the final reported payoff type profile $\mathbf{t}^n \in \mathbf{T}^{\check{\lambda}(\mathbf{t}^n)}$. The last equation is equivalent to

$$\begin{aligned} (|I| - 1) \sum_{i \in I} v_i(f_0(\mathbf{t}^n), t_i^n) &= - \sum_{i \in I} y_i^{\check{\lambda}(\mathbf{t}^n)}(t_{-i}^n) \\ \sum_{i \in I} v_i(f_0(\mathbf{t}^n), t_i^n) &= - \sum_{i \in I} \frac{y_i^{\check{\lambda}(\mathbf{t}^n)}(t_{-i}^n)}{|I| - 1}. \end{aligned}$$

Set

$$g_i^{\check{\lambda}(\mathbf{t}^n)}(t_{-i}^n) := - \frac{y_i^{\check{\lambda}(\mathbf{t}^n)}(t_{-i}^n)}{|I| - 1},$$

obtaining equation (11). This shows necessity.

“If”: Assume equation (11) and define

$$y_i^{\check{\lambda}(\mathbf{t}^n)}(t_{-i}^n) := -(|I| - 1) g_i^{\check{\lambda}(\mathbf{t}^n)}(t_{-i}^n). \quad (14)$$

Then

$$\begin{aligned} & \sum_{i \in I} \left(\sum_{j \neq i} v_j(f_0(\mathbf{t}), t_j) - (|I| - 1) g_i^{\check{\lambda}(\mathbf{t}^n)}(t_{-i}^n) + a_i(\mathbf{t}^1, \dots, \mathbf{t}^n) \right) \\ &= \sum_{i \in I} \left(\sum_{j \neq i} v_j(f_0(\mathbf{t}), t_j) - (|I| - 1) g_i^{\check{\lambda}(\mathbf{t}^n)}(t_{-i}^n) \right) \\ &= (|I| - 1) \sum_{i \in I} v_i(f_0(\mathbf{t}), t_i) - (|I| - 1) \sum_{i \in I} g_i^{\check{\lambda}(\mathbf{t}^n)}(t_{-i}^n) = 0 \end{aligned}$$

where the second line follows from budget neutrality of a_i -terms and the last line follows from equation (11). $\square$

While the proposition demonstrates that the incentives for awareness pooling do not necessarily add additional constraints on budget balance, it does not mean that budget balance is easy to satisfy with dynamic elaboration VCG mechanisms.

4.1 No Deficit

In classical mechanism design, it is well known that the VCG mechanisms like the Groves mechanisms can not satisfy budget balance in general (Green and Laffont (1979)). We

may want to look for weaker requirements. Not being budget balanced means that the mechanisms can run a deficit or surplus. In the context of unawareness, we are interested more in running no deficit than running no surplus for two reasons: First, if the mechanisms runs a deficit larger than the mechanism designer anticipated, agents may suspect that the mechanism designer is not committed to the mechanisms and become reluctant to truth-fully report their payoff types. Second, unforeseen contingencies are often used to justify budget overruns. But are deficits really inevitable in the presence of unawareness?

Recall that we used the convention that transfers f_i denote transfers *to* agent i . The following property is sometimes also called weak budget balance (e.g., [Shoham and Leyton-Brown \(2012\)](#)).

DEFINITION 7 (No deficit). We say that the dynamic direct elaboration mechanism with transfer functions $(f_i)_{i \in I}$ satisfies no deficit if for all $(t^1, \dots, t^n) \in \mathcal{T}$,

$$\sum_{i \in I} f_i(t^1, \dots, t^n) \leq 0. \quad (15)$$

In classical mechanism design, it is well known that a Groves mechanism may run a deficit, but the Clarke (or Pivotal) mechanism, in which each agent pays the negative externality that they impose on other agents through their effect on the social choice, does not. Therefore, the idea is to design a version of dynamic elaboration VCG mechanisms with Clarke transfers and show that it does not run a deficit.

To define the mechanism, we have to specify externalities. For any agent $i \in I$, define the $(-i)$ -utilitarian ex-post efficient outcome function $f_0^{-i} : \mathcal{T} \rightarrow X_0$ by $f_0^{-i}(t) = \arg \max_{x_0 \in X_0} \sum_{j \neq i} v_j(x_0, t_j)$. That is, for each profile of payoff types $t \in \mathcal{T}$, f_0^{-i} maximizes utilitarian ex-post welfare taking into account only the value functions of agent i 's opponents. Note that different from restricted outcome functions in standard Clarke mechanism the argument of f_0^{-i} is the *full* profile of payoff types $t = (t_j)_{j \in I}$ rather than just t_{-i} . The reason is that in order to compute f_0^{-i} we need the awareness of *all* agents, not just agents $j \neq i$: Although agent i value function is not considered when evaluating the social welfare of an outcome, for agents $j \neq i$, the efficiency of this outcome is still evaluated at the pooled awareness level.

DEFINITION 8 (Dynamic Elaboration Clarke Mechanism). We say that the dynamic direct elaboration mechanism implementing f is a dynamic elaboration Clarke mechanism if it is a dynamic elaboration VCG mechanisms with, for all $t^n \in \mathcal{T}^{\check{\lambda}(t^n)}$ and $i \in I$,

$$y_i^{\check{\lambda}(t^n)}(t_{-i}^n) := - \sum_{j \neq i} v_j(f_0^{-i}(t^n), t_j^n).$$

Observe that for any $i \in I$, if $\sum_{j \neq i} v_j(f_0^{-i}(t^n), t_j^n) \geq \sum_{j \neq i} v_j(f_0(t^n), t_j^n)$ then $f_i(t^1, \dots, t^n) - a_i(t^1, \dots, t^n)$ is non-positive. Since $f_0^{-i}(t^n)$ maximizes the sum of values over $j \in I \setminus \{i\}$, we have $\sum_{j \neq i} v_j(f_0^{-i}(t^n), t_j^n) \geq \sum_{j \neq i} v_j(f_0(t^n), t_j^n)$. Together with budget neutrality of the a_i -terms, this implies now that the mechanism satisfies no deficit. Utilitarian ex-post efficiency is implied by Theorem 1.

THEOREM 2. *The dynamic elaboration Clarke mechanism truthfully implements in conditionally dominant strategies a utilitarian ex-post efficient outcome under pooled awareness with no deficit.*

We illustrate the dynamic elaboration Clarke mechanisms with two examples. In the first example, neither agent will announce the pooled awareness level as the pooled awareness level is the join that is strictly greater than any agent's awareness.

Example 1 (Continuation) Consider again Example 1. In the conditional dominant solution to the dynamic elaboration Clarke mechanism, agents report in the first stage respectively,

$t_1 =$	Item	Cost	$t_2 =$	Item	Cost
	a	23		b	38
	b	41		c	29
	Total	64		Total	67

Since $\lambda(t_1) = \{a, b\}$ and $\lambda(t_2) = \{b, c\}$, agents are made aware of the join $\check{\lambda}(t_1, t_2) = \{a, b, c\}$ and are invited to report an elaborated type in $T_i^{\{a, b, c\}}$, $i \in 1, 2$, respectively. Extending the example slightly, let agents report truthfully their elaborations, respectively,

$t'_1 =$	Item	Cost	$t'_2 =$	Item	Cost
	a	23		a	19
	b	41		b	38
	c	16		c	29
	Total	80		Total	86

To complete the example, suppose that the possible physical outcomes (at any awareness level) are that the good is produced by agent 1, the good is produced by agent 2 or the good is not produced, i.e., $X_0 = \{1, 2, \emptyset\}$. Finally, let there be a third agent, a buyer, who always values the good at 100 no matter whether it comes from agent 1 or agent 2, i.e., $v_3(1, t_3) = v_3(2, t_3) = 100$ and $v_3(\emptyset, t_3) = 0$ for any t_3 .

Then, the efficient decision at the updated type profile is for the good to be produced by agent 1 (at a total cost of 80). Since no player is the first to announce the joint awareness level, the $a_i(\cdot)$ -term is 0 for each agent. Agent 3 is pivotal in the sense that if agent 3's valuation is not considered, the good would not be produced at all. We have for all t , $f(t) = 1$, $f^{-1}(t) = f^{-2}(t) = 1$, and $f^{-3}(t) = \emptyset$. The transfers to agent 1 are $100 - 0 - (100 - 0) + 0 = 0$, to agent 2 they are $20 - (100 - 80) - 0 = 0$, and to agent 3 are $-80 - 0 + 0 = -80$. Thus, the mechanism runs a surplus of 80.⁹ $\square$

⁹Note that the mechanism does not satisfy agent 1's ex-post participation constraints. This is true of the Clarke mechanism in this context even without any unawareness. It is known that the Clarke mechanisms may not necessarily satisfy ex-post participation constraints. We will analyze participation constraints in the next section.

Note that in this example awareness did not improve perceived utilitarian ex-post welfare. Nevertheless, we are able to raise awareness to the pooled awareness level with out mechanisms.

The next example illustrates a case in which the novel a_i -terms are non-trivial.

Example 2 (Continuation) Consider again Example 2, as depicted in Figure 2. The dynamic Clarke mechanism would proceed as follows: In the first stage, agent 1 announces t_1''' (with $\lambda(t_1''') = \bar{\ell}$) and agent 2 announces t_2' (with $\lambda(t_2') = \underline{\ell}$). In the second stage, agent 1 repeats his announcement and agent 2 reports t_2'''' (with $\lambda(t_2''') = \bar{\ell}$), his elaboration of t_2' . At the third stage both agents repeat the previous announcement and the mechanism stops. Agent 2 receives the good. Since agent 1 was the unique first agent to announce the joint awareness, agent 2 pays $m_1(\bar{\ell})$ to agent 1 (in addition to making the standard Clarke payment to the mechanism, latter amounting to the payment of 2 like in the second price auction in this example). For computing $m_1(\bar{\ell})$, assume for simplicity that the good is always given to agent 1 when both agents have the same value for the good. Then

$$m_1(\bar{\ell}) = \max_{\underline{t} \in T^{\underline{\ell}}, \bar{t} \in T^{\bar{\ell}}} \left(v_1(f_0(\underline{t}), \underline{t}_1) + v_2(f_0(\underline{t}), \underline{t}_2) - v_2(f_0^{-1}(\underline{t}_2), \underline{t}_2) \right. \\ \left. - \left(v_1(f_0(\bar{t}), \bar{t}_1) + v_2(f_0(\bar{t}), \bar{t}_2) - v_2(f_0^{-1}(\bar{t}_2), \bar{t}_2) \right) \right)$$

The term in the second line is non-negative. In the maximum, it is equal to zero. This is the case when we take $\bar{t}_2 = t_2'''$ and $\bar{t}_1 = t_1''''$. In such a case, the term is $2 + 0 - 2 = 0$. The r.h. term in the first line is at most 1, which is achieved when $\underline{t} = (t_1'', t_2')$ with associated values equal to $(2, 1)$, and noting that the only other type profile at this awareness level makes the term zero. Thus, $m_1(\bar{\ell}) = 1$. Note that this is non-negative, so the restriction that $m_i(\cdot)$ must be non-negative is not binding.

The additional terms in the transfers are just the payments from a second price auction at the joint awareness level (0 from agent 1 and 2 from agent 2). Therefore, agent 1 receives 1 and agent 2 pays 2. The mechanism runs a surplus of 1. Agent 1 has no incentive to conceal her awareness, since she gets a utility of 1 no matter whether she reveals it or not. $\square$

5. PARTICIPATION CONSTRAINTS

In this section, we study voluntary participation in dynamic elaboration Clarke mechanisms. We assume that agents have an outside option yielding them a payoff of zero for sure. As for VCG mechanisms without unawareness, we focus on ex-post participation constraints. In the presence of unawareness, at the onset of the mechanism, when they might have to sign up for participation, agents do not necessary anticipate all ex-post outcomes. That's why we will distinguish between satisfying ex-post participation constraints and satisfying just *ex-ante anticipated* ex-post participation constraints.

DEFINITION 9 (ex-post participation constraints). A dynamic direct elaboration mechanism implementing the social choice function f satisfies ex-post participation constraints if for all agents $i \in I$, information sets $h_i \in H_i(\sigma^*)$, initial profiles of payoff types $\mathbf{t} \in \mathbf{T}^{\lambda(h_i)}$ (as perceived by i at h_i), and initial profiles of awareness levels $\ell \in L(\lambda(h_i))^{|I|}$ (as perceived by i at h_i) such that $h_i \in H_i(\mathbf{t}, \ell, \sigma^*)$,

$$u_i(f(\tau(\mathbf{t}, \ell, \sigma^*)), t_i(h_i)) \geq 0.$$

DEFINITION 10 (ex-ante anticipated ex-post participation constraints). The dynamic direct elaboration mechanism satisfies ex-ante anticipated ex-post participation constraints if for every agent $i \in I$ it satisfies ex-post participation constraints for all initial information sets $h_i^1 \in H_i(\sigma^*)$.

In standard mechanism design with unawareness, participation constraints are typically not satisfied by the Clarke mechanisms without further assumptions. One such assumption in the literature are non-negative valuations:

ASSUMPTION 1 (Non-negative valuations). For any $\ell \in L$, $\mathbf{t} = (t_i)_{i \in I} \in \mathbf{T}^\ell$, $i \in I$, and $x_0 \in X_0^\ell$, $v_i(x_0, t_i) \geq 0$.

This assumption would be satisfied for instance in the context of the sale of a good via an auction or in the context of public goods/projects.

PROPOSITION 3. Assumption 1 implies that the dynamic elaboration Clarke mechanism satisfies ex-ante anticipated ex-post participation constraints.

PROOF. In the conditionally dominant strategy outcome, for all agents $i \in I$, initial information sets $h_i^1 \in H_i(\sigma^*)$, initial profiles of payoff types $\mathbf{t} \in \mathbf{T}^{\lambda(h_i^1)}$ (as perceived by i at h_i^1), and initial profiles of awareness levels $\ell \in L(\lambda(h_i^1))^{|I|}$ (as perceived by i at h_i^1) such that $h_i^1 \in H_i(\mathbf{t}, \ell, \sigma^*)$,

$$\begin{aligned} & u_i(f(\tau(\mathbf{t}, \ell, \sigma^*)), t_i(h_i^1)) \\ &= v_i(f_0(\tau^*(\mathbf{t}, \ell, \sigma^*)), t_i(h_i^1)) + f_i(\tau(\mathbf{t}, \ell, \sigma^*)) \\ &= v_i(f_0(\tau^*(\mathbf{t}, \ell, \sigma^*)), t_i(h_i^1)) + \sum_{j \neq i} v_j(f_0(\tau^*(\mathbf{t}, \ell, \sigma^*)), \tau_j^*(\mathbf{t}, \ell, \sigma^*)) \\ & \quad - \sum_{j \neq i} v_j(f_0^{-i}(\tau^*(\mathbf{t}, \ell, \sigma^*)), \tau_j^*(\mathbf{t}, \ell, \sigma^*)) + a_i(\tau(\mathbf{t}, \ell, \sigma^*)) \\ &= \sum_{j \in I} v_j(f_0(\tau^*(\mathbf{t}, \ell, \sigma^*)), \tau_j^*(\mathbf{t}, \ell, \sigma^*)) - \sum_{j \neq i} v_j(f_0^{-i}(\tau^*(\mathbf{t}, \ell, \sigma^*)), \tau_j^*(\mathbf{t}, \ell, \sigma^*)) + a_i(\tau(\mathbf{t}, \ell, \sigma^*)) \\ &\geq \sum_{j \in I} v_j(f_0(\tau^*(\mathbf{t}, \ell, \sigma^*)), \tau_j^*(\mathbf{t}, \ell, \sigma^*)) - \sum_{j \neq i} v_j(f_0^{-i}(\tau^*(\mathbf{t}, \ell, \sigma^*)), \tau_j^*(\mathbf{t}, \ell, \sigma^*)) \end{aligned}$$

where equality follows from $t_i(h_i^1) = \tau_i^*(\mathbf{t}, \ell, \sigma^*)$ and the last inequality follows from $\lambda(h_i^1) = \check{\lambda}(\tau^*(\mathbf{t}, \ell, \sigma^*))$ implying $a_i(\tau(\mathbf{t}, \ell, \sigma^*)) \geq 0$. (Recall that if $\mathbf{t} \in \mathbf{T}^{\lambda(h_i^1)}$ and $\ell \in$

$L(\lambda(h_i^1))^{|I|}$, then $\tau^*(t, \ell, \sigma^*) \in T^{\lambda(h_i^1)}$. If $i = i^*(\tau(t, \ell, \sigma^*))$, then $a_i(\tau(t, \ell, \sigma^*)) \geq 0$. Otherwise, if $i^*(\tau(t, \ell, \sigma^*)) = \emptyset$ (because several players raise the pooled awareness level at the first stage), then $a_i(\tau(t, \ell, \sigma^*)) = 0$. The case $k = i^*(\tau(t, \ell, \sigma^*)) \neq i$ is ruled out by playing σ_i^* .) That is, agent i at h_i^1 possesses already the anticipated pooled awareness level and raises it in the truth-telling solution.

Note $f_0(\tau^*(t, \ell, \sigma^*))$ could pick $f_0^{-i}(\tau^*(t, \ell, \sigma^*))$ if latter is the utilitarian ex-post welfare maximizing outcome when considering the utility of all agents in I . Thus,

$$\sum_{j \in I} v_j(f_0(\tau^*(t, \ell, \sigma^*)), \tau_j^*(t, \ell, \sigma^*)) \geq \sum_{j \in I} v_j(f_0^{-i}(\tau^*(t, \ell, \sigma^*)), \tau_j^*(t, \ell, \sigma^*))$$

Moreover, by Assumption 1,

$$\begin{aligned} \sum_{j \neq i} v_j(f_0^{-i}(\tau^*(t, \ell, \sigma^*)), \tau_j^*(t, \ell, \sigma^*)) + v_i(f_0^{-i}(\tau^*(t, \ell, \sigma^*)), \tau_i^*(t, \ell, \sigma^*)) \\ \geq \sum_{j \neq i} v_j(f_0^{-i}(\tau^*(t, \ell, \sigma^*)), \tau_j^*(t, \ell, \sigma^*)) \end{aligned}$$

Thus,

$$\sum_{j \in I} v_j(f_0(\tau^*(t, \ell, \sigma^*)), \tau_j^*(t, \ell, \sigma^*)) \geq \sum_{j \neq i} v_j(f_0^{-i}(\tau^*(t, \ell, \sigma^*)), \tau_j^*(t, \ell, \sigma^*))$$

which together with the arguments above implies

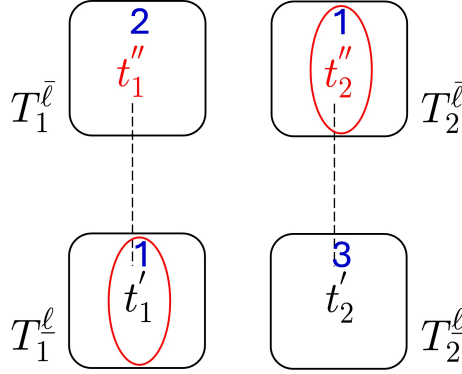
$$u_i(f(\tau(t, \ell, \sigma^*)), t_i(h_i^1)) \geq 0$$

□

While the dynamic elaboration Clarke mechanisms satisfies *ex-ante anticipate* ex-post participation constraints under Assumption 1, it does not satisfy ex-post participation constraints in general because during the play of the mechanisms agent i may become aware due to some agent raising awareness to a level strictly larger than i 's awareness level. This is illustrated in the following example.

Example 4 Similar to Example 3, there are two agents and two awareness levels $L = \{\underline{\ell}, \bar{\ell}\}$. For each agent and awareness level, the payoff type space is a singleton. Again, this will make the computation of the maximum in the definition of m_i for $i \in \{1, 2\}$ trivial. The payoff type spaces are given in Figure 3. There is a single object to be allocated. When an agent obtains the object, the valuation for the object is printed in blue above each payoff type. Otherwise, if the agent does not obtain the object, her value is zero. At awareness level $\bar{\ell}$, the utilitarian ex-post efficient outcome function must allocate the object to agent 1 while at awareness level $\underline{\ell}$ it must allocate it to agent 2. The twist is that only agent 2 has awareness level $\bar{\ell}$. In the dynamic elaboration Clarke mechanisms, initially agent 1 reports t'_1 and agent 2 reports t''_2 . The pooled awareness level is $\bar{\ell}$. At the next stage, agent 1 elaborates on her prior report by submitting report t''_1 . Agent 2 has nothing to elaborate. The following table shows the computations of the m_i , f_i , and u_i

FIGURE 3. Payoff Types in Example 3



for $i \in \{1, 2\}$ and each awareness level. The order of the terms in each sum follows the order of terms in the sum making up the definition of m_i , f_i and u_i , respectively:

	a_2	f_1	f_2	u_1	u_2
$\bar{\ell}$	$0 + 3 - 0 - 1 - 2 + 2 = 2$	$0 - 1 - 2 = -3$	$2 - 2 + 2 = 2$	$2 - 3 = -1$	$0 + 2 = 2$
$\underline{\ell}$	0	$3 - 3 + 0 = 0$	$0 - 1 + 0 = -1$	$0 + 0 = 0$	$3 - 1 = 2$

We observe that for agent 1, the ex-ante anticipated ex-post participation constraint is satisfied. At his initial awareness level $\underline{\ell}$, she does not expect to obtain the good and does not expect to pay. However, upon agent 2 raising pooled awareness to $\bar{\ell}$, agent 1 realizes that she will obtain the good but beyond the Clarke payments of 1 she also has to compensate agent 2 for raising awareness with an additional payments of 2. Thus, her ex-post utility is -1, violating her ex-post participation constraints at interim after stage 1. $\square$

The example underlines the importance of agents' ex-ante commitment to participating in the mechanisms in the face of unanticipated experiences.¹⁰

Showing that the dynamic elaboration Clarke mechanism satisfies ex-ante anticipated ex-post participation constraints relies on the fact that agents do not consider other agents to be more aware than themselves. One feature of unawareness is that if an agent is unaware of an event, then she is unaware that she is unaware of the event (so called AU introspection; see Heifetz et al. (2006)). While an agent who is unaware of an event cannot contemplate about this particular event and hence cannot contemplate about other agents' awareness of this event, she could nevertheless be aware that she might be unaware of something that other agents are aware of. Schipper (2024) extended unawareness structures of Heifetz et al. (2006, 2013b) to awareness of unawareness.¹¹ Since our payoff type space structure is a simplified version of unawareness structures,

¹⁰Typically, contracts exposing agents to unanticipated experiences such as enlisting in the military also pose draconian penalties for reneging the commitment. For instance, desertion was often punishable by death, especially during wartime.

¹¹For other approaches to awareness of unawareness in interactive settings, see Halpern and R go (2013) and Board and Chung (2021). In contrast to Schipper (2024), those approaches are not event-based.

we could also extend them to awareness of unawareness and model explicitly the state of the mind when an agent contemplates her participation constraints under awareness of unawareness. In such a model, there can be two situations, one in which the agent is unaware of her unawareness and one in which she has the same payoff type but is aware of her unawareness. While in the former, the agent would happily sign on the mechanisms, in latter the agent might be reluctant to do so because although she cannot anticipate what she is currently unaware of, she suspects that she has to pay agent i for raising awareness. Such a scenario is quite realistic as agents may have experience in prior mechanisms and are able to reason about other agents having less awareness. We leave the extension to future work. However, in the following section we develop a mechanism for procurement that satisfies participation constraints even after awareness is raised.

6. PROCUREMENT UNDER EX-ANTE UNFORESEEN CONTINGENCIES

A common mechanism used in procurement is the reverse action. Despite the similarity of the second price auction (Vickrey (1961)) to the Clarke mechanism, the reverse second price auctions is not an example of the Clarke mechanism. The reasons are twofold: First, as usual in economics textbooks, we defined f_0^{-i} by $f_0^{-i}(t) = \arg \max_{x_0 \in X_0} \sum_{j \neq i} v_j(x_0, t_j)$. With such a definition the payment received by the winning bidder uses an outcome where i supplies the good when computing the maximum welfare without i 's cost/value function. Thus, she does not pay the second smallest bid but zero, which differs from the reverse second price auction.¹² The other reason is that the buyer is special in that her payments are not computed using the Clarke transfers but she pays the second lowest bid to the winning bidder as required by the reverse second price auction. She is what has been called a “sink-agent” (e.g., Nath and Sandholm (2019)) with respect to transfers. Since procurement of complex projects under ex-ante unforeseen contingencies and incomplete specifications is an extremely relevant topic, we like to extend our setting to it. In particular, using dynamic direct elaboration mechanisms we aim to remedy a shortcoming of standard reverse second price auctions and allow the buyer to take advantage of the expertise of bidders by pooling their awareness.

Let $I = \{1, \dots, b\}$. We say that $\langle X_{0,-i}, c_i, \rangle_{i \in I \setminus \{b\}}$ is a *procurement context* if for all $i \in I \setminus \{b\}$, $X_{0,-i} = X_{0,-i}^\ell = \{\emptyset\} \cup (I \setminus \{b, i\})$ for all $\ell \in L$, and $c_i = -v_i$ is the cost function of agent $i \in I \setminus \{b\}$. This is interpreted as follows: There is a special agent, the buyer denoted by b . The other agents, $I \setminus \{b\}$, are potential suppliers/sellers. The set of outcomes $X_{0,-i} = \{1, 2, i-1, i+1, \dots, b-1, \emptyset\}$ specifies which of the sellers in $\{1, 2, i-1, i+1, \dots, b-1\}$ supplies the good or whether nobody supplies the good, $\emptyset$. This is the set of outcomes feasible without supplier i . The value $c_i(t_i)$ is seller i 's cost of supplying the good when her type is t_i . If she does not supply the good, her cost is zero. Note that in this setting, the strategy of agent i is still to report an entire payoff type,

¹²We could redefine dynamic elaboration Clarke mechanisms by letting f_0^{-i} be the maximum of the sum of valuations of agent i 's opponents only over outcomes still available without agent i . In such a case, we would need to require further assumptions in order to show no deficit; see Shoham and Leyton-Brown (2012).

while $c_i(t_i)$ would be agent i 's total cost/value associated with the reported payoff type t_i .

DEFINITION 11 (Dynamic Elaboration Reverse Second Price Auction). Given a procurement context, the dynamic direct elaboration mechanism implementing f is a dynamic elaboration reverse second price auction if f_0 assigns the project to a lowest bidder (with an arbitrary but ex-ante given tie breaking rule) and transfers f_i to seller $i \in I \setminus \{b\}$ are given for any sequence of reported payoff type profiles $(t^1, \dots, t^n) \in \mathcal{T}$ by

$$f_i(t^1, \dots, t^n) := \mathbb{I}_i(t^n) c_{(2)}(t^n) + a_i(t^1, \dots, t^n), \quad (16)$$

where for $i \in I \setminus \{b\}$ indicator function $\mathbb{I}_i$ is defined by

$$\mathbb{I}_i(t^n) := \begin{cases} 1 & \text{if } f_0(t^n) = i \\ 0 & \text{otherwise} \end{cases} \quad (17)$$

and $c_{(2)}(t^n)$ is the second smallest order statistic of $\{c_i(t_i^n)\}_{i \in I \setminus \{b\}}$. The term $a_i(\cdot)$ is defined for $i \in I \setminus \{b\}$ by

$$a_i(t^1, \dots, t^n) := \begin{cases} m_i(\check{\lambda}(t^n)) & \text{if } i = i^*(t^1, \dots, t^n) \\ 0 & \text{otherwise} \end{cases} \quad (18)$$

where as before

$$i^*(t^1, \dots, t^n) := \left\{ i \in I : \exists k \leq n \left((\lambda(t_i^k) = \check{\lambda}(t^n)) \text{ and } \nexists j \neq i, k' \leq k (\lambda(t_j^{k'}) = \check{\lambda}(t^n)) \right) \right\}$$

and $m_i(\ell)$ is defined recursively, as follows: $m_i(\underline{\ell}) := 0$ and for any $\ell \triangleright \underline{\ell}$,

$$m_i(\ell) := \max_{\ell' \triangleleft \ell, t' \in \mathbf{T}^{\ell'}, t \in \mathbf{T}^{\ell}} \left\{ (m_i(\ell') + \mathbb{I}_i(t') (c_{(2)}(t') - c_i(t'_j)) - \mathbb{I}_i(t) (c_{(2)}(t) - c_i(t_j))) \right\}, \quad (19)$$

Finally, the transfer to the buyer is

$$f_b(t^1, \dots, t^n) := -c_{(2)}(t^n) - \sum_{i \in I \setminus \{b\}} a_i(t^1, \dots, t^n). \quad (20)$$

In the dynamic elaboration reverse second price auction, there are a finite number of stages of raising awareness after which the project is awarded to the bidder with the lowest bid price. She is paid by the buyer the second lowest bid price. Any bidder, no matter whether she is awarded the contract or not, may receive a compensation for raising awareness if she is the unique agent who is the first to raise awareness to the pooled awareness level. In such a case, the agent is also compensated for the best possible profit she would lose by pretending to be less aware, which is similar to the m_i -terms of the transfers of the dynamic elaboration Clarke mechanisms. The buyer pays the second lowest bid to the lowest bidder and the compensation for raising awareness if any. Note that the unique agent raising the awareness level to the pooled awareness level may also be the buyer, in which case she would not have to pay any compensation

for raising awareness to the sellers. She also does not pay any compensation for raising awareness if there is no unique bidder who raises awareness to the pooled awareness level or if the pooled awareness level is the join that is strictly greater than awareness levels communicated by the sellers.

Note that if we assume that for every awareness level and seller there is a payoff type profile with which she does not have the lowest bid price or she could always opt out even if it would be utilitarian ex-post efficient for her to produce, then the definition of $m_i(\ell)$ of equation (19) simplifies to

$$m_i(\ell) = \max_{\ell' \triangleleft \ell, \mathbf{t}' \in \mathbf{T}^{\ell'}} (m_i(\ell') + \mathbb{I}_i(\mathbf{t}') (c_{(2)}(\mathbf{t}') - c_i(t'_i))) .$$

The dynamic elaboration reverse second price auctions satisfies the following:

THEOREM 3. *Given any procurement context, the dynamic elaboration reverse second price auction implements a utilitarian ex-post efficient outcome under pooled awareness in conditionally dominant strategies with budget balance and satisfies ex-post participation constraints of all sellers/bidders.*

The proof of utilitarian ex-post efficiency in conditional dominant strategies under pooled awareness is analogous to the proof of Theorem 1 and thus omitted. Budget balance follows immediately from the fact that the buyer makes any payments to the sellers. Satisfaction of ex-post participation constraints for bidders follow from the fact that at worst, a bidder is left with zero transfers, in which case she also does not supply the object. If she supplies the object, she must be the lowest bidder but is paid the second lowest bid, leaving her a non-negative profit. Note that not just ex-ante anticipated ex-post participation constraints are satisfied for each bidder but even the interim anticipated ex-post participation constraints are satisfied. Moreover, awareness of unawareness does not play any role. Even if a bidder is aware that other bidders may be aware of something that she herself is not, her participation constraints are satisfied.

7. DISCUSSION

One might be concerned that our mechanisms incentivize agents to raise awareness of irrelevant events in order to receive larger side payments. Formally, this will not affect any of our results, since to the extent these events are irrelevant, awareness of them will not change the physical outcome. However, we may still be concerned that this will make payments larger — in the case of the reverse auction, payments from the mechanism. This concern can be addressed by modifying the mechanism so that it ignores awareness of irrelevant events when computing the transfers. Formalizing this modified mechanism comes at the cost of some additional notation, so we relegate it to the [Online Appendix](#).

Our results appear to contradict [Jehiel and Moldovanu \(2001\)](#)'s impossibility theorem on efficient implementation under interdependent valuations. We achieve ex-post efficiency despite asymmetric awareness introducing interdependent values because

agents can only report less awareness than they have, not more (note also we assume private values *given* an awareness level). This creates a unidirectional incentive compatibility environment. [Krähmer and Strausz \(2024\)](#) show that when agents can only report lower types, incentive compatibility does not restrict implementable allocation rules. Although in their model the order on types is also connected to payoffs, while in our model there is no necessary connection between awareness and payoffs, this assumption is only needed for other results in their paper. In both cases, the key insight is that sufficient transfers can incentivize truthful reporting when truth is the maximal feasible report.

It is interesting to note that [Mezzetti \(2004\)](#) overcomes the impossibility of ex-post efficient implementation under interdependent valuations with two-stage Groves mechanisms in which agents first report benefits from outcomes and transfers are determined only after every valuation becomes transparent. Our mechanism has some similarity to his, in that transfers are determined only after payoffs are clarified and reported in a second stage. However, unlike in Mezzetti, we require a second stage of reporting to determine the allocation, and the reason we are able to elicit truthful reports is because of the “evidence”-like properties of awareness. In contrast, Mezzetti utilizes payoff information reported at an ex-post stage where payoff interdependencies are no longer relevant. Moreover, Mezzetti’s mechanism is incentive compatible in (ex-post) equilibrium rather than dominant strategies.

While this is the first paper on general mechanism design under unawareness, it is just a first step. One shortcoming is that upon becoming aware, we assume that their more elaborated payoff type becomes immediately transparent to agents. While sometimes raising awareness may be an “eye opener”, other times it may also require substantial efforts on part of the agents to acquire information on relevant events and issues they newly became aware. In a natural next step, we study information acquisition upon becoming aware using ideas from [Bergemann and Välimäki \(2002\)](#).

The next step is to ask for optimal (e.g., revenue maximizing) mechanism design under unawareness. How would a designer maximize surplus from agents with heterogeneous awareness. As standard mechanism design, we would need to go beyond belief-free mechanism design. In our dynamic context, this would be complicated by updating of beliefs in light of new information and awareness. Since in standard mechanism design, optimal mechanism design is the multiagent extension of screening problems, under unawareness the screening problem studied by [Francetich and Schipper \(2024\)](#) should be relevant. Also [Li and Schipper \(2024\)](#) study raising bidders’ awareness before second price auctions in a revenue maximizing way. Nevertheless, optimal mechanism design under unawareness remains an unexplored field.

REFERENCES

- ALPER, O. AND W.B. BONING (2003): “Using procurement auctions in the Department of Defense,” *Center for Naval Analysis*, CRM D0007515.A3. [2]
- AUSTER, S. (2003): “Asymmetric awareness and moral hazard,” *Games and Economic Behavior*, 82, 503–521. [6]

- AUSTER, S. AND N. PAVONI (2024): “Optimal delegation and information transmission under limited awareness,” *Theoretical Economics*, 19, 245–284. [6]
- BAJARI, P., R. MCMILLAN, AND S. TADELIS (2008): “Auctions versus negotiations in procurement: An empirical analysis,” *Journal of Law, Economics & Organization*, 25, 372–399. [2]
- BAJARI, P. AND S. TADELIS (2001): “Incentives versus transaction costs: A theory of procurement contracts,” *RAND Journal of Economics*, 32, 387–407. [2]
- BERGEMANN, D. AND J. VÄLIMÄKI (2002): “Information acquisition and efficient mechanisms design,” *Econometrica*, 70, 1007–1033. [33]
- BESTER, H. AND R. STRAUSS (2001): “Contracting with imperfect commitment and the revelation principle: The single agent case,” *Econometrica*, 69, 1077–1098. [2]
- BOARD, O. AND K.S. CHUNG (2021): “Object-based unawareness: Axioms,” *Journal of Mechanism and Institution Design*, 6, 1–36. [7, 29]
- BÖRGERS, T. (2015): *An introduction to the theory of mechanism design*, Oxford: Oxford University Press. [22]
- BULL, J. AND J. WATSON (2007): “Hard evidence and mechanism design,” *Games and Economic Behavior*, 58, 75–93. [2]
- CHUNG, K.S. AND L. FORTNOW (2016): “Loopholes,” *Economic Journal*, 126, 1774–1797. [6]
- CLARKE, E.H. (1971): “Multipart pricing of public goods,” *Public Choice*, 8, 19–33. [5]
- DEKEL, E., B. LIPMAN, AND A. RUSTICHINI (1998): “Standard state-space models preclude unawareness,” *Econometrica*, 66, 159–173. [3]
- DENECKERE, R. AND S. SEVERINOV (2008): “Mechanism design with partial state verifiability,” *Games and Economic Behavior*, 64, 487–513. [2]
- FAGIN, R. AND J. HALPERN (1988): “Belief, awareness, and limited reasoning,” *Artificial Intelligence*, 34, 39–76. [3, 7]
- FEINBERG, Y. (2021): “Games with unawareness,” *B.E. Journal of Theoretical Economics*, 21, 433–488. [3, 7]
- FILIZ-OZBAY, E. (2012): “Incorporating unawareness into contract theory,” *Games and Economic Behavior*, 76, 191–194. [6]
- FRANCETICH, A. AND B.C. SCHIPPER (2024): “Rationalizable screening and disclosure under unawareness,” Tech. rep., University of Washington, Bothell, and University of California, Davis. [6, 33]
- GOLDBERG, V.P. (1977): “Competitive bidding and the production of precontract information,” *Bell Journal of Economics*, 8, 250–261. [2]
- GRANT, S., J. KLINE, AND J. QUIGGIN (2012): “Differential awareness, ambiguity, and incomplete contracts: A model of contractual disputes,” *Journal of Economic Behavior and Organization*, 82, 494–504. [6]

- GREEN, J. AND J.-J. LAFFONT (1979): *Incentives in public decision-making*, Amsterdam: North-Holland. [5, 23]
- (1986): “Partially verifiable information and mechanism design,” *Review of Economic Studies*, 53, 447–456. [2]
- GROVES, T. (1973): “Incentives in teams,” *Econometrica*, 41, 617–631. [4, 18]
- GROVES, T. AND M. LOEB (1975): “Incentives for public inputs,” *Journal of Public Economics*, 4, 211–226. [4, 18]
- HALPERN, J. AND L.C. RÊGO (2013): “Reasoning about knowledge of unawareness revisited,” *Mathematical Social Sciences*, 66, 73–84. [29]
- HALPERN, J. AND L.C. RÊGO (2014): “Extensive games with possibly unaware players,” *Mathematical Social Sciences*, 70, 42–58. [3, 7]
- HEIFETZ, A., M. MEIER, AND B.C. SCHIPPER (2006): “Interactive unawareness,” *Journal of Economic Theory*, 130, 78–94. [3, 7, 29]
- (2008): “A canonical model of interactive unawareness,” *Games and Economic Behavior*, 62, 304–324. [9]
- (2013a): “Dynamic unawareness and rationalizable behavior,” *Games and Economic Behavior*, 81, 100–121. [3, 7, 13]
- (2013b): “Unawareness, beliefs, and speculative trade,” *Games and Economic Behavior*, 77, 100–121. [3, 7, 29]
- (2021): “Prudent rationalizability in generalized extensive-form games with unawareness,” *B.E. Journal of Theoretical Economics*, 21, 525–556. [4]
- HERWEG, F. AND K. SCHMIDT (2020): “Procurement with unforeseen contingencies,” *Management Science*, 66, 2194–2212. [6, 7]
- HOLMSTRÖM, B. (1977): *On incentives and control in organizations*, Dissertation: MIT. [5, 22]
- JEHIEL, P. AND B. MOLDOVANU (2001): “Efficient design with interdependent valuations,” *Econometrica*, 69, 1237–1259. [7, 32]
- KRÄHMER, D. AND R. STRAUSS (2024): “Unidirectional incentive compatibility,” Tech. rep., University of Bonn. [7, 33]
- LEE, J. (2008): “Unforeseen contingencies and renegotiation with asymmetric information,” *Economic Journal*, 118, 678–694. [6]
- LEI, H. AND X.J. ZHAO (2021): “Delegation and information disclosure with unforeseen contingencies,” *B.E. Journal of Theoretical Economics*, 21, 637–656. [6]
- LI, Y.X. AND B.C. SCHIPPER (2024): “Raising bidders’ awareness in second-price auctions,” Tech. rep., University of California, Davis. [6, 33]
- MASKIN, E. AND J. TIROLE (1999): “Unforeseen contingencies and incomplete contracts,” *Review of Economic Studies*, 66, 83–114. [6]

MEIER, M. AND B.C. SCHIPPER (2014): “Bayesian games with unawareness and unawareness perfection,” *Economic Theory*, 56, 219–249. [3, 7]

——— (2024): “Conditional dominance in games with unawareness,” Tech. rep., University of California, Davis. [3, 4]

MEZZETTI, C. (2004): “Mechanism design with interdependent valuations: Efficiency,” *Econometrica*, 72, 1617–1626. [33]

MILGROM, P. (1981): “Good news and bad news: Representation theorems and applications,” *Bell Journal of Economics*, 12, 380–391. [4]

——— (2004): *Putting auction theory to work*, Cambridge, M.A.: Cambridge University Press. [22]

NATH, S. AND T. SANDHOLM (2019): “Efficiency and budget balance in general quasi-linear domains,” *Games and Economic Behavior*, 113, 673–693. [30]

OFFICE OF THE FEDERAL REGISTER (2024): “Federal Acquisition Regulation: Prohibition on the use of reverse auctions for complex, specialized, or substantial design and construction services (FAR Case 2023–003),” *Federal Register*, 89, 70157–70160. [2]

PIERMONT, E. (2024): “Iterated revelation: How to incentive experts to complete incomplete contracts,” Tech. rep., Royal Holloway. [6]

SCHIPPER, B.C. (2021): “Discovery and equilibrium in games with unawareness,” *Journal of Economic Theory*, 198, 105365. [3, 7, 13]

——— (2024): “Interactive awareness of unawareness,” Tech. rep., University of California, Davis. [5, 29]

SHIMOJI, M. AND J. WATSON (1998): “Conditional dominance, rationalizability, and game forms,” *Journal of Economic Theory*, 83, 161–195. [4]

SHOHAM, Y. AND K. LEYTON-BROWN (2012): *Multiagent systems. Algorithmic, game-theoretic, and logical foundations*, Cambridge, M.A.: Cambridge University Press. [24, 30]

SOMMER, S.C. AND C.H. LOCH (2009): “Incentive contracts in projects with unforeseeable uncertainty,” *Production and Operations Management*, 18, 185–196. [6]

SWEET, J. (1994): *Legal aspects of architecture, engineering and the construction process*, St. Paul: West Publishing Company. [2]

VICKREY, W. (1961): “Counterspeculation, auctions, and competitive sealed tenders,” *Journal of Finance*, 16, 8–37. [30]

VON THADDEN, E.-L. AND X.J. ZHAO (2012): “Incentives for unaware agents,” *Review of Economic Studies*, 79, 1151–1174. [6]