

Optimal mechanism for the sale of a durable good^{*}

Laura Doval[†]

Vasiliki Skreta[‡]

January 20, 2020

Abstract

We show that posted prices are the optimal mechanism to sell a durable good to a privately informed buyer when the seller has limited commitment in an infinite horizon setting. We provide a methodology for mechanism design with limited commitment and transferable utility. Whereas in the case of commitment, subject to the buyer's truth-telling and participation constraints, the seller's problem is a decision problem, in the case of limited commitment, the seller's problem corresponds to an intrapersonal game, where different "incarnations" of the seller represent the different beliefs he may have about the buyer's valuation.

KEYWORDS: mechanism design, limited commitment, intrapersonal equilibrium, information design, self-generation, posted price

JEL CLASSIFICATION: D84, D86

^{*}We would like to thank Rahul Deb, Frederic Koessler, Pablo Schenone, and especially Max Stinchcombe, as well as audiences at Cowles, SITE, and Stony Brook, for thought-provoking questions and illuminating discussions. Vasiliki Skreta is grateful for generous financial support through the ERC consolidator grant 682417 "Frontiers in design." This research was supported by grants from the National Science Foundation (SES-1851744 and SES-1851729).

[†]California Institute of Technology, Pasadena, CA 91125. E-mail: ldoval@caltech.edu

[‡]University of Texas at Austin, University College London, and CEPR. E-mail: vskreta@gmail.com

1 Introduction

A classic problem in mechanism design is how to sell a good to a buyer so as to maximize the seller's revenue. The answer is surprisingly simple: out of the many ways in which a seller could achieve this feat, a posted price is optimal. This result does not depend on the length of the interaction between the buyer and the seller: even if they would interact for infinitely many periods, the revenue-maximizing mechanism is to post the same price in each period. An important assumption behind this result is that the seller can commit to the whole sequence of mechanisms that the buyer faces. The commitment assumption is important because the optimal mechanism is, in general, *time-inconsistent*. When the buyer does not buy the good at the posted price, the seller has to be able to not lower the price even if there is common knowledge of unrealized gains from trade.

However, much less is known about how to sell a durable good when the seller can commit to today's mechanism, but not to the mechanism he will offer if no sale occurs. In an infinite horizon game in which the seller may only post a price in each period, [Ausubel and Deneckere \(1989\)](#) show the seller can achieve monopoly profits provided that (i) the buyer's valuation is drawn from a continuum, with the lowest valuation being smaller than the seller's marginal cost (i.e., the *no gap* case), and (ii) we consider the limit case in which the seller is fully patient. In a finite horizon game in which the seller can choose mechanisms in each period and the buyer's valuation is drawn from a continuum, [Skreta \(2006\)](#) shows the seller's optimal mechanism under limited commitment is a sequence of posted prices. These results rely, however, on the assumption that the seller *directly* observes the buyer's choices – the rejection of the price in [Ausubel and Deneckere \(1989\)](#) and the choice of messages in [Skreta \(2006\)](#). Indeed, in a two-period model, [Denicolo and Garella \(1999\)](#) show that when the seller only observes whether trade happens or not, the seller may instead prefer to ration high valuation buyers in the first period, in order to induce a strong demand in the second period.

In this paper, we study an *infinite-horizon mechanism-selection* game between an uninformed seller, who owns one unit of a durable good, and a privately informed buyer. The buyer's valuation is binary and fully persistent. At the beginning of each period, as long as the good has not been sold, the seller offers the buyer a *mechanism*, the rules of which determine the allocation for that period. Following [Bester and Strausz \(2007\)](#) and [Doval and Skreta \(2018\)](#), the set of mechanisms the seller may choose from allows the seller to design how much he observes about the buyer's choices, and hence, design his beliefs about the buyer's type. Since

the seller's beliefs define the demand schedule he faces, it is feasible for him to implement, amongst others, the demand schedule that obtains from rationing in Denicolo and Garella (1999).

The main result of this paper is that, among all mechanisms, posted prices are optimal when the seller has limited commitment. [Theorem 1](#) characterizes the revenue-maximizing Perfect Bayesian equilibrium and a strategy profile that achieves it. The latter is such that, as long as a sale has not occurred, the seller will choose a mechanism that can be implemented as a posted price. Contrary to the case in which the seller has commitment to the entire sequence of mechanisms, the sequence of posted prices is not constant, reflecting the fact that the seller makes inferences about the buyer's valuation in the event that no trade happens.

The optimality of posted prices echoes the result for the case in which the seller has commitment, showing that the *shape* of the mechanism does not depend on whether there is commitment to short- or long-term mechanisms. Moreover, this result provides a microfoundation for the strategy space used by the literature that studies the sale of a durable good by a monopolist. Inasmuch as we are interested in understanding the best outcome for the monopolist, assuming he can post prices is without loss of generality.

The key conceptual innovation of the analysis is to translate the problem of finding the revenue-maximizing PBE in the mechanism-selection game between the seller and the buyer into an *auxiliary* problem, which only involves the seller. This auxiliary problem is an *intrapersonal game*, where the seller is the only player, and each of his "incarnations" represents the different beliefs he may have about the buyer's valuation. This step establishes a formal connection between the literatures on mechanism design with limited commitment and that on time inconsistency. Time inconsistency of optimal mechanisms under full commitment lies at the heart of the difficulties in mechanism design with limited commitment ([Laffont and Tirole \(1990\)](#)). This paper is the first to show that these two problems, mechanism design with limited commitment and intrapersonal games, are formally related and to use the solution of one to solve the other.

To understand how we arrive to this auxiliary problem, a review of the four main steps involved in the proof of [Theorem 1](#) is useful (We contrast our approach to previous work on mechanism design with limited commitment, especially [Skreta \(2006\)](#) in the related literature). First, we rely on the results in our previous work, [Doval and Skreta \(2018\)](#), to simplify the class of mechanisms the seller will offer in any equilibrium of the game and, more importantly, the buyer's

equilibrium behavior. This step reduces the search of the optimal sequence of mechanisms to those that satisfy, loosely speaking, a sequence of *participation* and *truthtelling* constraints. Second, we derive necessary conditions that a revenue-maximizing PBE assessment must satisfy. These necessary conditions are the dynamic analogue of establishing which participation and incentive compatibility constraints are binding in the case of commitment. We use these conditions to obtain a recursive program which provides an upper bound on the seller's payoff in the payoff-maximizing PBE, analogous to the maximization of the *virtual surplus* in the case of commitment. Indeed, as in the case of commitment, in this recursive program, the seller's payoff is written solely as a function of the allocation: the transfers have been removed by using the binding constraints.

The third step shows this recursive program has a unique solution. Whereas in the case of commitment the maximization of the virtual surplus is a decision problem, it corresponds to an *intrapersonal game* in the case of limited commitment. Indeed, the best response condition of the intrapersonal equilibrium captures the restrictions imposed by sequential rationality on the seller's behavior. Whenever the seller chooses an allocation that involves no trade, he internalizes how his beliefs will change and how that change will determine the choice of mechanisms in the continuation.

The output of this third step is a mapping that assigns a choice of a mechanism to each belief the seller may hold about the buyer's valuation. The final step is to return to our dynamic incomplete information game and show how to use this mapping to construct the revenue-maximizing PBE assessment. Crucial to this step is a result in our previous work that allows us to apply self-generation techniques as in [Abreu et al. \(1990\)](#); [Athey and Bagwell \(2008\)](#) to characterize the whole set of equilibrium payoffs.

Notwithstanding the conceptual contributions of the analysis, our work also makes a methodological contribution in that it provides a *recipe* for analyzing mechanism design problems with limited commitment with transferable utility. Indeed, the four steps we described above follow very closely the prototypical steps in classical mechanism design: (i) invoke the revelation principle to simplify the class of mechanisms and the agent's behavior, (ii) find the binding constraints, (iii) maximize the virtual surplus, and (iv) show the solution satisfies any constraints that may have been ignored. As will become clear from the analysis that follows, these four steps will also be necessary in other problems of mechanism design with limited commitment and transferable utility beyond the one we consider here.

Our goal is to highlight the conceptual and technical nuances that arise in an infinite-horizon *mechanism selection* game when the designer has limited commitment, and thus our focus on the case of binary valuations. This allows us to bring to the forefront the complexities that arise because of the designer’s rich action space. The extension to the case in which the buyer’s valuation is drawn from a continuum is of interest, and we plan to address it in future work. When possible, we highlight in footnotes how the results presented extend to the continuum.

Related Literature: The paper contributes mainly to four strands of literature. The first strand, similar to this paper, derives optimal mechanisms when the designer has limited commitment, but unlike this paper, these studies consider finite horizon settings (see Laffont and Tirole (1988); Skreta (2006, 2015); Deb and Said (2015); Fiocco and Strausz (2015); Beccuti and Möller (2018)). In these papers, the observability of the agent’s choices in the mechanism makes it hard to disentangle the principal’s choice of mechanism from the agent’s best response, which determines, in turn, the principal’s continuation beliefs, and hence, his continuation payoffs. In their seminal paper, to get around this difficulty, Bester and Strausz (2001) allow the agent to misrepresent her type with positive probability. Instead, Skreta (2006) circumvents the same problem by leveraging the properties of the set of incentive feasible outcomes and the finite horizon, which pins down the optimal mechanism in the final period, to study the implications of the seller’s sequential rationality constraints for the set of incentive feasible outcomes.

Our approach allows us to reduce the seller’s search for the optimum to a set of mechanisms that satisfy the truthtelling and participation constraints of the buyer, without having to impose any restrictions on the horizon of the game. This simplifies the analysis by allowing us, for the most part, to ignore the buyer as a player. This approach relies on endowing the seller with a larger set of mechanisms, which comes “at the cost” of having to rule out as optimal a richer set of possibilities. As we discuss Section 3.2, the richer set of mechanisms and no exogenous deadline on the interaction between the seller and the buyer, make the analysis differ from the case in which the horizon is finite.

The second strand, similar to this paper, studies infinite horizon problems in which the designer has limited commitment and faces an agent with one of finitely many types.¹ Given the difficulties with the revelation principle when the de-

¹Beccuti and Möller (2018) lie somewhat in between these two strands because they take limits of their finite horizon results to draw conclusions about the infinite horizon setting. However, they do not show that this limit corresponds to the seller’s revenue-maximizing equilibrium in the

signer has limited commitment (Laffont and Tirole (1988); Bester and Strausz (2001)), the works on this second strand either impose restrictions on the class of contracts that can be offered (e.g., Strulovici (2017); Gerardi and Maestri (2018)), or on the solution concept (e.g., Acharya and Ortner (2017)). This paper contributes to the literature in two ways. First, by appealing to the results in Doval and Skreta (2018), we avoid imposing restrictions on the length of the interaction, the class of mechanisms, and the solution concept. Second, whereas we exploit the specifics of our problem to show equilibrium existence, the methodology suggested by the paper can be imported into other settings of interest and, as we discuss throughout the paper, it can be particularly powerful in the case of transferable utility.

The third strand is the literature that starts with the observation in Coase (1972) that the durable-good monopolist faces a time-inconsistency problem, which in turn limits his monopoly power. The papers in the durable-good monopolist literature (Stokey (1981); Bulow (1982); Gul et al. (1986); Sobel (1991); Ortner (2017)) study price dynamics² and establish (under some conditions) Coase's conjecture.³ Related to this literature is the problem of dynamic bargaining with one-sided incomplete information, where an uninformed proposer each period makes a take-it-or-leave-it offer to a privately informed receiver.⁴ Sobel and Takahashi (1983), Fudenberg et al. (1985), and Ausubel and Deneckere (1989) characterize the equilibria of this game for the case of finite horizon, infinite horizon and gap, infinite horizon and no gap, respectively. In an analogous vein, Burguet and Sakovics (1996), McAfee and Vincent (1997), Caillaud and Mezzetti (2004), and Liu et al. (2019) study equilibrium reserve-price dynamics without commitment in different settings.⁵ The common thread in all these papers is that the seller's inability to commit reduces monopoly profits.

The last strand is the literature on games with time-inconsistent preferences, which builds on the seminal works of Strotz (1955), Peleg and Yaari (1973), and Pollak (1968). Our contribution to this literature is to show that optimal mecha-

infinite horizon game.

²Even when the seller can commit not to revise prices, price dynamics can arise if there are demand or cost changes. See, for example, Stokey (1979), Conlisk et al. (1984), and Garrett (2016).

³There is also a literature that analyzes different variations on the sale of a durable good model to understand whether Coase's conjecture still obtains. See, for instance, McAfee and Wiseman (2008); Board and Pycia (2014).

⁴More recently, Peski (2019) studies alternating bargaining games where players can offer menus.

⁵Dilmé and Li (Forthcoming) study revenue management when the seller cannot commit to revise prices.

nisms under limited commitment can be obtained as a solution to an intrapersonal game, even when both the seller and the buyer have time consistent preferences. In this way, we extend the range of applications of this literature, which has already seen applications ranging from growth and development (Bernheim et al. (2015)) to saving decisions (Harris and Laibson (2001)).

Organization The rest of the paper is organized as follows. Section 2 describes the model; Section 2.1 summarizes the results in Doval and Skreta (2018) used to simplify the analysis that follows. Section 3 formally states the main result and describes the main steps in the proof of Theorem 1. Section 3.1 derives the recursive formulation that is the basis of the intrapersonal game; Section 3.2 characterizes the unique equilibrium of the intrapersonal game; Section 3.3 uses the intrapersonal equilibrium to build a PBE assessment that delivers the revenue-maximizing PBE. Section 4 concludes.

2 Model

Primitives: Two players, a seller and a buyer, interact over infinitely many periods. The seller owns one unit of a durable good to which he attaches value 0. The buyer has private information: before her interaction with the seller starts, she observes her valuation $v \in \{v_L, v_H\} \equiv V$, with $0 \leq v_L < v_H$. Let μ_0 denote the probability that the buyer's valuation is v_H at the beginning of the game. In what follows, we denote by $\Delta(V)$ the set of distributions on V .

An allocation in period t is a pair $(q, x) \in \{0, 1\} \times \mathbb{R}$, where q indicates whether the good is traded ($q = 1$) or not ($q = 0$), and x is a payment from the buyer to the seller. The game ends the first time the good is traded.

Payoffs are as follows. If in period t , the allocation is (q, x) , the flow payoffs are $u_B(q, x, v) = vq - x$ and $u_S(q, x) = x$ for the buyer and the seller, respectively. The seller and the buyer maximize the expected discounted sum of flow payoffs. They share a common discount factor $\delta \in (0, 1)$.

Mechanisms: To introduce the timing of the game, we first define the action space of the seller in each period. In each period, the seller offers the buyer a mechanism. Following Myerson (1982); Bester and Strausz (2007); Doval and Skreta (2018), we define a mechanism as follows. A mechanism, $\mathbf{M} = \langle (M^{\mathbf{M}}, \beta^{\mathbf{M}}, S^{\mathbf{M}}), (q^{\mathbf{M}}, x^{\mathbf{M}}) \rangle$, consists of a communication device $\beta^{\mathbf{M}}$, which maps an input message $m \in M^{\mathbf{M}}$

to a distribution with finite support on $S^{\mathbf{M}}$, and an allocation rule $(q^{\mathbf{M}}, x^{\mathbf{M}})$, which maps each element s in $S^{\mathbf{M}}$ to a probability of trade, $q^{\mathbf{M}}(\mu')$, and a payment from the buyer to the seller, $x^{\mathbf{M}}(\mu')$.⁶ We endow the seller with a collection $(M_i, S_i)_{i \in \mathcal{I}}$ of input and output messages in which each M_i is finite, contains at least two elements, and each S_i contains $\Delta(M_i)$. Denote by $\mathcal{M}$ the set of all mechanisms with message sets $(M_i, S_i)_{i \in \mathcal{I}}$.⁷

Timing: If in period t , the good is yet to be traded, then

- $t. 0$ Both players observe a draw from a public randomization device $\omega \in [0, 1]$.⁸
- $t. 1$ The seller offers the buyer a mechanism $\mathbf{M} = (\langle M^{\mathbf{M}}, \beta^{\mathbf{M}}, S^{\mathbf{M}} \rangle, (q^{\mathbf{M}}, x^{\mathbf{M}}))$.
- $t. 2$ The buyer observes the mechanism and decides to participate ($p = 1$) or not ($p = 0$).
 - $t. 2.1$ If the buyer does not participate in the mechanism, the good is not sold and no payments are made; that is, the allocation is $(q, x) = (0, 0)$.
 - $t. 2.2$ If the buyer participates in the mechanism,
 - i. She chooses a report $m \in M^{\mathbf{M}}$, which is unobserved by the seller,
 - ii. An output message $s \in S^{\mathbf{M}}$ is drawn from $\beta^{\mathbf{M}}(\cdot | m)$, which, in turn, determines the probability of trade $q^{\mathbf{M}}(s)$, and the payment, $x^{\mathbf{M}}(s)$. The output message and the allocation are observed by both the seller and the buyer. If the good is not traded, the game proceeds to period $t + 1$.

The above description determines an infinite horizon dynamic game, which we denote by $G_{\mathcal{M}}^{\infty}(\mu_0)$ to remind the reader that the seller's beliefs at the beginning

⁶A priori, we could have allowed a mechanism to offer a randomization over allocations, i.e. a randomization over $\{0, 1\} \times \mathbb{R}$. However, because both players have quasilinear payoffs, the seller does not benefit from randomizing over transfers (see [Lemma II.1](#) in [Doval and Skreta \(2019\)](#) for a proof). He may, however, benefit from randomizing over whether the good is traded.

⁷We restrict the seller to choose mechanisms with input and output messages in $(M_i, S_i)_{i \in \mathcal{I}}$ to have a well-defined action space for the seller. This allows us to have a well-defined set of deviations, avoiding set-theoretic issues related to self-referential sets. The analysis in [Doval and Skreta \(2018\)](#) shows that the choice of the collection plays no further role in the analysis.

⁸Public randomization devices are standard in repeated games (see, e.g., [Fudenberg and Maskin \(1986\)](#)).

of the game are given by μ_0 . Public histories in this game are

$$h^t = (\omega_0, \mathbf{M}_0, p_0, s_0, (q_0, x_0), \dots, \omega_{t-1}, \mathbf{M}_{t-1}, p_{t-1}, s_{t-1}, (q_{t-1}, x_{t-1}), \omega_t),$$

where $p_k \in \{0, 1\}$ denotes the buyer's participation decision, with the restriction that if $p_k = 0$, then $s_k = \emptyset$, i.e., no signal is generated, and $(q_k, x_k) = (0, 0)$.⁹ Let H^t denote the set of all period t public histories; they capture what the seller knows through period t . A strategy for the seller is then $\Gamma : \cup_{t=0}^{\infty} H^t \mapsto \Delta(\mathcal{M}_C)$.

A history for the buyer consists of the *public* history of the game together with the buyer's reports into the mechanism (henceforth, the buyer history) and her private information. Formally, a buyer history is an element

$$h_B^t = (\omega_0, \mathbf{M}_0, p_0, m_0, s_0, (q_0, x_0), \dots, \omega_{t-1}, \mathbf{M}_{t-1}, p_{t-1}, m_{t-1}, s_{t-1}, (q_{t-1}, x_{t-1}), \omega_t).$$

Let H_B^t denote the set of buyer histories through period t . The buyer also knows her valuation and hence, when her valuation is v , a history through period t is an element of $\{v\} \times H_B^t$. The buyer's participation strategy is $\pi_v : \cup_{t=0}^{\infty} (H_B^t \times \mathcal{M}) \mapsto [0, 1]$. Conditional on participating in the mechanism $\mathbf{M}$, her reporting strategy is a distribution $r_v(h_B^t, \mathbf{M}, 1) \in \Delta(M^{\mathbf{M}})$ for each of her types v and each $h_B^t \in H_B^t$.

A belief for the seller at the beginning of time t , history h^t , is a distribution $\mu(h^t) \in \Delta(V \times H_B^t(h^t))$, where $H_B^t(h^t)$ is the set of buyer histories consistent with the public history, h^t .

Solution concept: We are interested in studying the Perfect Bayesian equilibrium (henceforth, PBE) payoffs of this game, where PBE is defined as follows. An assessment, $\langle \Gamma, (\pi_v, r_v)_{v \in V}, \mu \rangle$, is a PBE if the following hold:

1. Given $\mu(h^t)$, $\Gamma(h^t)$ is sequentially rational given $(\pi_v, r_v)_{v \in V}$,
2. Given $\Gamma(h^t)$, $\pi_v(h_B^t, \cdot)$, $r_v(h_B^t, \cdot, 1)$ are sequentially rational for all $h_B^t \in H_B^t$ and $v \in V$,
3. $\mu(h^t)$ is derived via Bayes' rule where possible (see [Definition 4 in Section A.1](#)).¹⁰

⁹While there is no output message when the buyer does not participate in the mechanism, we denote this by $s = \emptyset$ to keep the length of all histories the same.

¹⁰[Section A.1](#) contains the formal statement of Bayes' rule where possible. To wit, we impose the following requirement. Let h and h' be two consecutive information sets for the seller. The beliefs at h' are obtained via Bayes' rule from the beliefs at h if one of the following holds: (i) conditional

We denote by $\mathcal{E}_{\mathcal{M}}^*(\mu_0) \subseteq \mathbb{R}^3$ the set of PBE payoffs of $G_{\mathcal{M}}^\infty(\mu_0)$ and we use (u_S, u_H, u_L) to denote a generic element of $\mathcal{E}_{\mathcal{M}}^*(\mu_0)$, where u_S is the seller's payoff and u_H, u_L denote the buyer's payoff when her type is v_H, v_L , respectively.

The seller's highest equilibrium payoff in $G_{\mathcal{M}}^\infty(\mu_0)$ is of particular interest to us:

$$u_S^*(\mu_0) = \max\{u_S : (u_S, u_H, u_L) \in \mathcal{E}_{\mathcal{M}}^*(\mu_0)\}. \quad (1)$$

Theorem 1 characterizes $u_S^*(\mu_0)$ and an assessment that achieves it for all $\mu_0 \in \Delta(V)$, which we denote by $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$.

2.1 Revelation Principle

There are at least two reasons that the game $G_{\mathcal{M}}^\infty(\mu_0)$ is not a simple to analyze. First, the seller's action space is large and, a priori, it is not clear which mechanisms could be ruled out from consideration. Second, fixed a seller's strategy, and hence a sequence of mechanisms faced by the buyer, we still need to understand the buyer's best response in the game induced by the sequence of mechanisms. The main result of [Doval and Skreta \(2018\)](#) allows us to deal with these two complications by (i) pinning down a well-defined set of mechanisms that we can restrict attention to, and (ii) pinning down the buyer's behavior as a best response to the sequence of mechanisms offered by the seller.

Let $\mathcal{M}_C$ denote the set of all mechanisms where $(M, S) = (V, \Delta(V))$. Foreshadowing the result in [Lemma 1](#), we label these mechanisms *canonical*. Let $G^\infty(\mu_0)$ denote the same game in the previous section, except that in each period the seller's action space is $\mathcal{M}_C$ and let $\mathcal{E}^*(\mu_0)$ denote the set of equilibrium payoffs of $G^\infty(\mu_0)$.

In [Doval and Skreta \(2018\)](#), we show that to characterize the set of equilibrium payoffs of $G_{\mathcal{M}}^\infty(\mu_0)$, it is enough to characterize the set of equilibrium payoffs of $G^\infty(\mu_0)$. Moreover, we also show that it is without loss of generality to focus on a particular class of assessments to characterize *all* equilibrium payoffs in $\mathcal{E}^*(\mu_0)$. [Lemma 1](#) below summarizes these results for future reference:

Lemma 1 ([Doval and Skreta \(2018\)](#)). $G_{\mathcal{M}}^\infty(\mu_0)$ and $G^\infty(\mu_0)$ have the same set of equilibrium payoffs, i.e., $\mathcal{E}_{\mathcal{M}}^*(\mu_0) = \mathcal{E}^*(\mu_0)$.

on reaching h , h' is reached with positive probability under the equilibrium strategy profile, or (ii) conditional on reaching h , h' is reached from h through a deviation by the seller and the buyer playing the equilibrium strategy profile.

Moreover, let $(u_S, u_H, u_L) \in \mathcal{E}^*(\mu_0)$. Then, there exists a PBE assessment $\langle \Gamma, (\pi_v, r_v)_{v \in V}, \mu \rangle$ of $G^\infty(\mu_0)$ that achieves payoff (u_S, u_H, u_L) and satisfies the following properties:

1. For all histories h^t , the buyer participates in the mechanism offered by the seller at that history and truthfully reports her type, with probability 1,
2. For all histories h^t , if the mechanism offered by the seller at h^t outputs posterior μ' , the seller's updated equilibrium beliefs about the buyer coincide with μ' ,
3. The strategy of the buyer depends only on her private valuation and the public history.

Lemma 1 has two implications. First, it is without loss of generality to restrict the seller to offering mechanisms where the set of input messages are the set of type reports and the set of output messages is the set of beliefs the seller has about the buyer's type. Second, it is without loss of generality to restrict attention to assessments that satisfy certain properties.

Part 1 of **Lemma 1** implies the mechanisms chosen by the seller in equilibrium must satisfy a *participation* constraint and an *incentive compatibility* constraint for each buyer type and each public history. For completeness, these constraints are stated in [Section A.2](#); see [Equation \$PC_{v,h^t}\$](#) and [Equation \$IC_{v,h^t}\$](#) . As in the case of commitment to long-term mechanisms, part 1 simplifies the analysis of the buyer's behavior, by reducing it to a series of constraints. Part 2 implies we can interpret the seller's choice of a communication device as a choice of a distribution over posteriors that satisfies a Bayes' plausibility constraint (see [Equation \$BC_{\mu\(h^t\)}\$](#) in [Section A.2](#)). As a consequence, our analysis involves aspects of *information design*. Part 3 implies the set of PBE payoffs of $G^\infty(\mu_0)$ coincides with the set of *Public PBE* payoffs of $G^\infty(\mu_0)$ ([Athey and Bagwell \(2008\)](#)), allowing us to invoke self-generation techniques as in [Abreu et al. \(1990\)](#); [Athey and Bagwell \(2008\)](#) to argue the assessment we construct in [Section 3.3](#) is indeed a PBE assessment.

A corollary of **Lemma 1** is that $u_S^*(\mu_0)$ is also the highest equilibrium payoff the seller can obtain in $G^\infty(\mu_0)$. The rest of the paper focuses on studying the equilibrium payoffs of $G^\infty(\mu_0)$ and when we refer to a PBE assessment, we mean one that satisfies the conditions of **Lemma 1**.

In what follows, we abuse notation in the following way. Because valuations are binary, we can think of an element in $\Delta(V)$ (a distribution over v_L and v_H) as an element of the interval $[0, 1]$ (the probability assigned to v_H). We use the latter formulation in what follows. That is, whereas the mechanism outputs a distribution over v_L and v_H , we index this distribution by the probability of v_H .

3 Main Result

Section 3 contains the main result of the paper: the optimal mechanism for a seller of a durable good is a sequence of posted prices. To state the result, we proceed as follows. First, we define what it means for the seller to attain a certain equilibrium payoff via a sequence of posted prices, when the seller's action space consists of mechanisms (Definition 1). This is enough to informally state our main result. Second, we describe a PBE assessment, $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$, that achieves $u_S^*(\mu_0)$ and argue that, indeed, $u_S^*(\mu_0)$ can be implemented via a sequence of posted prices. The rest of this section describes the main steps to prove $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$ indeed achieves $u_S^*(\mu_0)$.

Definition 1. Let u_S denote a seller's payoff in $\mathcal{E}^*(\mu_0)$. The payoff u_S can be *implemented via a sequence of posted prices* if a PBE assessment $\langle \Gamma, (\pi_v, r_v)_{v \in V}, \mu \rangle$ exists such that in each history h^t , the seller's mechanism $\mathbf{M}$ satisfies the following: there exist $\mu'_1 \neq \mu'_2$ and a price $p(h^t)$ such that

1. $\beta^{\mathbf{M}}(\mu'_1 | v_L) = \beta^{\mathbf{M}}(\mu'_2 | v_H) = 1$,
2. $(q^{\mathbf{M}}(\mu'_1), x^{\mathbf{M}}(\mu'_1)) = (0, 0)$ and $(q^{\mathbf{M}}(\mu'_2), x^{\mathbf{M}}(\mu'_2)) = (1, p(h^t))$.

Thus, a mechanism is a posted price whenever it can be implemented as if the buyer could choose one of two options: either she buys the good at price $p(h^t)$, or she does not buy the good, in which case she pays nothing. Note that when the seller posts a price in each period, he does not use *noisy* communication devices: by observing the posterior, the seller knows which report the buyer submitted into the mechanism. However, the buyer's report *need not be truthful* in this implementation.

Theorem 1. Let $u_S^*(\mu_0)$ denote the seller's maximum revenue in a PBE. Then, $u_S^*(\mu_0)$ can be implemented via a sequence of posted prices.

Theorem 1 shows that, amongst all trading protocols, posted prices are optimal. This provides a microfoundation for the seller's strategy space in the literature that studies the sale of a durable good.

Theorem 1 echoes the result for the case in which the seller has commitment and faces one buyer, where it is also well known that a posted price is optimal. However, in the latter case, the seller would post the same price at all histories, which would fail to be sequentially rational in $G^\infty(\mu_0)$. As we show in the analysis that follows, we find that the seller's mechanism determines a price path that is

decreasing in the seller's belief that the buyer's valuation is v_H , as in the literature that analyzes the Coase conjecture.

PBE assessment that achieves $u_S^*(\mu_0)$: The proof of [Theorem 1](#) singles out a particular PBE assessment, $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$, that achieves $u_S^*(\mu_0)$, which we now define informally (the formal statement is in [Section E.2](#)). After defining it, we use it to argue that $u_S^*(\mu_0)$ can be implemented via a sequence of posted prices. In what follows, the ratio v_L/v_H is important and we denote it by $\bar{\mu}_1$. To see why $\bar{\mu}_1$ is important, recall the following fact from second-degree price discrimination: when the seller sells at a price of v_L , he leaves rents, $\Delta v \equiv v_H - v_L$, to v_H , so that when v_H has probability μ , the seller's revenue, v_L , can be written as

$$v_L = \mu(v_H - \Delta v) + (1 - \mu)v_L = \mu v_H + (1 - \mu)(v_L - \frac{\mu}{1 - \mu} \Delta v) = \mu v_H + (1 - \mu)\hat{v}_L(\mu).$$

The first equality represents revenue as the surplus extracted from each type. The second equality represents revenue as the *virtual surplus*, where the value of allocating the good to v_L is adjusted to capture the fact that when v_L is served, so is v_H , which leaves rents to v_H . $\bar{\mu}_1$ is the belief at which $\hat{v}_L(\mu) = 0$. At that belief, the seller is indifferent between serving the buyer for both of her valuations and excluding the low-valuation buyer.

We focus here and in the rest of this section on the case in which $v_L > 0$; we return to the (less interesting) case of $v_L = 0$ at the end of this section. When $v_L > 0$, we show a sequence $0 = \bar{\mu}_0 < \bar{\mu}_1 = v_L/v_H < \dots < \bar{\mu}_n < \dots$ exists such that the following holds:¹¹

1. Along the equilibrium path:

- (a) If at history h^t , the seller's beliefs, $\mu^*(h^t)$, are in $[\bar{\mu}_0, \bar{\mu}_1)$, he chooses a mechanism such that $(q^{\mathbf{M}_t^*}(\mu^*(h^t)), x^{\mathbf{M}_t^*}(\mu^*(h^t))) = (1, v_L)$ and the communication device satisfies that $\beta^{\mathbf{M}_t^*}(\mu^*(h^t)|v) = 1$ for $v \in \{v_L, v_H\}$.

¹¹Chapter 10 in [Fudenberg and Tirole \(1991\)](#) solves for the equilibrium of the price posting game for the binary type case. [Hart and Tirole \(1988\)](#) compare price paths and the seller's revenue across different sale modes (sale vs. leasing) in a finite-horizon setting with limited commitment and binary valuations. As in their model with sale, the optimal price path in our model is determined by the beliefs at which the seller is indifferent between trading with the low-valuation buyer in n versus $n + 1$ periods. However, because their model has a finite horizon and they define the buyer's value as a flow payoff from consuming the good, it is not immediate to compare the price paths in both models.

- (b) If at history h^t , the seller's beliefs, $\mu^*(h^t)$, are in $[\bar{\mu}_n, \bar{\mu}_{n+1})$ for $n \geq 1$, the seller's mechanism satisfies the following. First, it induces two posteriors, $\bar{\mu}_{n-1}$ and 1. Second, the allocation rule satisfies that $(q^{\mathbf{M}_i^*}(1), x^{\mathbf{M}_i^*}(1)) = (1, v_L + (1 - \delta^n)\Delta v)$, whereas $(q^{\mathbf{M}_i^*}(\bar{\mu}_{n-1}), x^{\mathbf{M}_i^*}(\bar{\mu}_{n-1})) = (0, 0)$. Finally, the communication device maps v_L to $\bar{\mu}_{n-1}$, whereas it maps v_H to both $\bar{\mu}_{n-1}$ and 1 with positive probability. The probabilities $\beta^{\mathbf{M}_i^*}(\bar{\mu}_{n-1}|v_H), \beta^{\mathbf{M}_i^*}(1|v_H)$ are chosen so that when the seller observes $\bar{\mu}_{n-1}$, his updated belief coincides with $\bar{\mu}_{n-1}$.
2. Off the equilibrium path, the seller's strategy coincides with the above, except that when $\mu^*(h^t) = \bar{\mu}_n$ for some $n \geq 1$, the seller may randomize between the mechanism he offers on the path of play when his belief is $\bar{\mu}_n$ and the one he offers on the path of play when his belief is $\bar{\mu}_{n-1}$.¹²
 3. At each history h^t , the buyer's best response to the seller's equilibrium offer at h^t is to participate in the mechanism and truthfully report her valuation.

Implementation via posted prices: We now argue that if the assessment described above achieves $u_S^*(\mu_0)$, then $u_S^*(\mu_0)$ can be achieved via a sequence of posted prices. Clearly, when the seller's beliefs are below $\bar{\mu}_1$, the seller's mechanism corresponds to selling the good at a price of v_L . Consider then the case in which the seller's beliefs are in $[\bar{\mu}_1, \bar{\mu}_2)$. Note that when the buyer's valuation is v_H and the realized allocation is trade, then her payoff is $v_H - v_L - (1 - \delta)\Delta v = \delta\Delta v$. On the other hand, when the buyer's valuation is v_H and the realized allocation is no trade, then the seller's beliefs next period are $\bar{\mu}_0 = 0$, so that the buyer's continuation payoff is $\delta\Delta v$. That is, the buyer with valuation v_H is indifferent between obtaining the good at price $v_L + (1 - \delta)\Delta v$ and not obtaining the good, and paying a price of v_L in the next period. Since the buyer with valuation is v_H is indifferent between these two options, she is willing to mix between buying at price $v_L + (1 - \delta)\Delta v$ and not obtaining the good. She does so in a way that the seller's belief is $\bar{\mu}_0$ when the allocation is $(0, 0)$. Since $\bar{\mu}_0 = 0$, it implies that when the seller's belief is in $[\bar{\mu}_1, \bar{\mu}_2)$, the buyer buys with probability 1 at a price of $v_L + (1 - \delta)\Delta v$. Working recursively through the equations, one can show that when the seller's prior is in $[\bar{\mu}_n, \bar{\mu}_{n+1})$, the mechanism is equivalent to posting a price of $v_L + (1 - \delta^n)\Delta v$. In this case, the low-valuation buyer chooses the $(0, 0)$

¹²The need for mixing arises for technical reasons: it ensures that the buyer's continuation payoffs when her valuation is v_H are upper-semicontinuous and, thus, guarantees that a best response exists after any deviation by the seller (see [Section E.1](#)). Indeed, we appeal to the results in [Simon and Zame \(1990\)](#) to simultaneously determine the buyer's best response and the seller's mixing.

allocation, whereas the high-valuation buyer mixes so that the seller's belief is $\bar{\mu}_{n-1}$ when the allocation is $(0, 0)$.

It is interesting to contrast the implementation under the PBE assessment $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$ described above with the implementation via posted prices. In the former, the buyer is truthful and the seller *rationes* the high valuation buyer as in [Denicolo and Garella \(1999\)](#), which slows down the rate at which the seller's beliefs fall conditional on the good not being sold. Instead, in the implementation via posted prices, the high valuation buyer misreports her type with positive probability, which, like rationing, prevents the seller from becoming pessimistic too quickly about the buyer's valuation. Since both implementations are payoff equivalent, the seller cannot do better with rationing than with posted prices.

The rest of this section describes the steps to show the assessment, $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$, achieves $u_S^*(\mu_0)$. [Section 3.1](#) derives necessary conditions that the PBE assessment, $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$, must satisfy. They are the basis for the formulation of a recursive problem, the intrapersonal game, studied in [Section 3.2](#). The main output of [Section 3.2](#) is a mapping $\gamma^* : \Delta(V) \mapsto \mathcal{M}_C$, which describes for each prior μ'_0 the seller may have about the buyer, his optimal choice of mechanism. [Section 3.3](#) discusses then how we use γ^* to construct $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$. Finally, we appeal to self-generation techniques to show $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$ is indeed a PBE assessment.

3.1 Binding constraints and virtual surplus

The main output of [Section 3.1](#) is [Problem 1](#) at the end of this section, which relies on [Proposition 1](#) below. To introduce [Proposition 1](#), we introduce the definition of incentive efficiency (see [Bester and Strausz \(2001\)](#)):¹³

Definition 2 (Incentive efficiency). Let $\langle \Gamma, (\pi_v, r_v)_{v \in V}, \mu \rangle$ be a PBE assessment and (u_S, u_H, u_L) its associated payoff. $\langle \Gamma, (\pi_v, r_v)_{v \in V}, \mu \rangle$ is incentive efficient if no $u'_S > u_S$ exists such that $(u'_S, u_H, u_L) \in \mathcal{E}^*(\mu_0)$.

In an incentive efficient PBE, the seller at the beginning of the game is earning his best payoff consistent with the buyer's equilibrium payoff.

¹³Incentive efficiency as defined in [Bester and Strausz \(2001\)](#) is related, but different from, incentive efficiency as defined in [Holmström and Myerson \(1983\)](#). While the definition in [Holmström and Myerson \(1983\)](#) would allow for Pareto improvements for both the buyer and the seller, the definition in [Bester and Strausz \(2001\)](#) only considers Pareto improvements for the seller.

Proposition 1. *If $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$ is the PBE assessment that delivers the seller's best payoff in $\mathcal{E}^*(\mu_0)$, $u_S^*(\mu_0)$, the following hold:*

1. *Without loss of generality, if the buyer rejects the seller's equilibrium choice of mechanism at history h^t , the seller assigns probability 1 to the buyer's valuation being v_H .*
2. *For every history on the path of play, h^t , $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle|_{h^t}$ is incentive efficient.*
3. *For all histories on the path of $(\Gamma^*, (\pi_v^*, r_v^*)_{v \in V})$, the buyer is indifferent between participating in the mechanism and not when her valuation is v_L .*
4. *For all histories on the path of $(\Gamma^*, (\pi_v^*, r_v^*)_{v \in V})$, the buyer is indifferent between reporting v_H and v_L when her valuation is v_H .*

Part 1 implies that it is without loss of generality to assume that the buyer's value of rejecting the seller's equilibrium choice of mechanism is 0: when the seller assigns probability 1 to v_H , the seller asks for a payment of v_H . Part 1 follows from two observations: (i) By Lemma 1, the buyer's rejection of the seller's equilibrium offer is an off the path event, so that the seller's beliefs are not pinned down by Bayes' rule, and (ii) by setting the beliefs of the seller to 1 on that event, we only relax the buyer's participation constraint without affecting the seller's incentives.¹⁴

Part 2 is a consequence of the result that it is without loss of generality to focus on Public PBE. While in general it is not true that $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle|_{h^t}$ is a PBE assessment of $G^\infty(\mu^*(h^t))$, this is indeed the case when $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$ is a Public PBE assessment (see Proposition I.1 in Doval and Skreta (2019)). Part 2 then says that along the equilibrium path the seller's continuation payoff is the highest equilibrium payoff consistent with the buyer's equilibrium payoff. Thus, while the seller may not be earning $u_S^*(\mu^*(h^t))$ for a history h^t on the equilibrium path, the payoffs starting from h^t are the highest equilibrium payoffs consistent with the buyer's payoff at h^t .

Parts 3 and 4 imply that the standard observations from second-degree price discrimination hold along the path of play of $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$, namely, that all surplus is extracted from the buyer when her valuation is v_L and the buyer's

¹⁴Part 1 of Proposition 1 is not a consequence of using PBE as opposed to sequential equilibrium, since the buyer could always tremble and reject the mechanism with a higher probability when her valuation is v_H .

payoff coincides with her information rents when her valuation is v_H . Clearly, they must hold at the initial history, because either of them failing implies the seller is “leaving money on the table.” That they hold after every history on the path of play is a result of discounting and the linearity of payoffs in transfers, which allows us to distribute the players’ payoffs across time without affecting (and sometimes even increasing) payoffs at some history.¹⁵

It follows from parts 3 and 4 that the participation constraint for v_L and the incentive compatibility constraint for v_H bind along the equilibrium path.¹⁶ Recall that when the seller has commitment to long-term mechanisms, these constraints are the ones that we use to replace the transfers out of the seller’s payoffs, so that they are expressed solely in terms of the allocation. In what follows, we show that the same can be done when analyzing the revenue-maximizing PBE when the seller can only commit to short-term mechanisms.

Let $u_S^*(\mu_0)$ denote the seller’s highest equilibrium payoff in $G^\infty(\mu_0)$ and let $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$ denote the PBE assessment that achieves it. Lemma 1 implies that we can write $u_S^*(\mu_0)$ as follows:

$$u_S^*(\mu_0) = \sum_{v \in V} \mu_0(v) \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}_0^*}(\mu'|v) [x^{\mathbf{M}_0^*}(\mu') + \delta(1 - q^{\mathbf{M}_0^*}(\mu')) U_S^*(\mathbf{M}_0^*, 1, \mu', 0, x^{\mathbf{M}_0^*})], \quad (2)$$

where $U_S^*(h^1)$ denotes the seller’s payoffs under $(\Gamma^*, (\pi_v^*, r_v^*)_{v \in V})$ from $t = 1$ onwards. (Note that $U_S^*(h^1)$ is not necessarily $u_S^*(\mu^*(h^1))$ and hence the difference in notation.) By definition, $u_S^*(\mu_0)$ is the highest payoff the seller can achieve, given the continuation strategy, from among all mechanisms $\mathbf{M} \in \mathcal{M}_C$ that satisfy the participation and incentive compatibility constraints at the initial history. Proposition 1 implies $\mathbf{M}_0^*$ satisfies the following three properties. First, the buyer’s payoff

¹⁵Parts 2-4 may sound counterintuitive from a dynamic game perspective: there are games where, to sustain high payoffs today, low continuation payoffs are needed, even on the equilibrium path. Importantly, when we prove Proposition 1, we show that when we modify the strategy profile from history h^t onwards, we do not upset the equilibrium constraints at the histories that precede h^t , thereby showing that the new assessment is also a PBE assessment.

¹⁶Although with binary types, the incentive constraint for v_H binds as a result of revenue maximization, with a continuum of types the *adjacent downward-looking* incentive constraints obtain because incentive compatibility implies the envelope representation of payoffs. In the online appendix to Doval and Skreta (2018), we show how to obtain the envelope representation of payoffs in the abstract mechanism selection game we study there; it is immediate to show that it holds in this game as well. Thus, with a continuum of types, we obtain a representation of the seller’s payoff similar to the one in this section.

is 0 when her valuation is v_L , so that the following holds:

$$\sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}_0^*}(\mu'|v_L) v_L q^{\mathbf{M}_0^*}(\mu') = \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}_0^*}(\mu'|v_L) x^{\mathbf{M}_0^*}(\mu'), \quad (3)$$

where we have already replaced the buyer's continuation value with 0. Second, when the buyer's valuation is v_H , she is indifferent between reporting v_H and v_L , so that the following equality holds:

$$\begin{aligned} & \sum_{\mu' \in \Delta(V)} (\beta^{\mathbf{M}_0^*}(\mu'|v_H) - \beta^{\mathbf{M}_0^*}(\mu'|v_L)) (v_H q^{\mathbf{M}_0^*}(\mu') + (1 - q^{\mathbf{M}_0^*}(\mu')) U_{H|L}^*(h^1)) \\ &= \sum_{\mu' \in \Delta(V)} (\beta^{\mathbf{M}_0^*}(\mu'|v_H) - \beta^{\mathbf{M}_0^*}(\mu'|v_L)) x^{\mathbf{M}_0^*}(\mu'), \end{aligned} \quad (4)$$

where the notation $U_{H|L}^*(h^1)$ signifies that the continuation value of a buyer of valuation v_H is the utility that she would obtain from the allocation that corresponds to v_L – these are her information rents. Finally, the induced distribution over posteriors averages out to μ_0 . That is, letting $\tau^{\mathbf{M}_0^*}(\mu_0, \mu') = \sum_{v \in V} \mu_0(v) \beta^{\mathbf{M}_0^*}(\mu'|v)$ denote the probability that posterior μ' is induced, we have

$$\sum_{\mu' \in \Delta(V)} \tau^{\mathbf{M}_0^*}(\mu_0, \mu') \mu' = \mu_0. \quad (5)$$

Replacing transfers out of seller's payoff: We can replace equations 3-5 in Equation 2 to obtain

$$u_S^*(\mu_0) = \sum_{\mu' \in \Delta(V)} \tau^{\mathbf{M}_0^*}(\mu_0, \mu') \left[q^{\mathbf{M}_0^*}(\mu') (\mu' v_H + (1 - \mu') \hat{v}_L(\mu_0)) + (1 - q^{\mathbf{M}_0^*}(\mu')) \times \delta \left(U_S^*(h^1) + \left(\frac{\mu'}{1 - \mu'} - \frac{\mu_0}{1 - \mu_0} \right) (1 - \mu') U_{H|L}^*(h^1) \right) \right], \quad (6)$$

where

$$\hat{v}_L(\mu_0) = v_L - \frac{\mu_0}{1 - \mu_0} \Delta v,$$

is the buyer's *virtual value* when her valuation is v_L and the seller assigns probability μ_0 to the buyer's valuation being v_H .

Equation 6 says that we can think of the seller's optimal mechanism at $t = 0$ as choosing a distribution over posteriors, $\tau^{\mathbf{M}_0^*}$, and, for each posterior he induces, a probability of trade, $q^{\mathbf{M}_0^*}$. If at a given posterior, μ' , he sells the good (i.e., $q^{\mathbf{M}_0^*}(\mu') = 1$), he obtains the expected *virtual surplus*, where the expectation is

taken with respect to μ' , but the low-valuation buyer's virtual value is evaluated at μ_0 . That the low-valuation buyer's virtual value is evaluated at μ_0 instead of at μ' is intuitive. After all, the seller assigns probability μ_0 to the buyer's valuation being v_H so that, whenever he sells to both buyer types (i.e., $\mu' > 0$), he leaves rents to v_H with probability μ_0 . If at a given posterior, μ' , he delays trade (i.e., $q^{\mathbf{M}_0^*}(\mu') = 0$), then he receives his continuation payoff. However, it is modified by the buyer's continuation rents: the seller internalizes that, whenever trade is delayed, the buyer receives continuation rents when her valuation is v_H . Note the continuation rents enter with a term that depends on μ_0 : the seller assigns probability μ_0 to v_H , and hence that is the rate at which he pays rents (this time in terms of continuation values) to the buyer.

It follows from [Proposition 1](#) that an expression similar to that of [Equation 6](#) holds for any history h^t on the path of $(\Gamma^*, (\pi_v^*, r_v^*)_{v \in V})$. That is,¹⁷

$$U_S^*(h^t) = \sum_{\mu' \in \Delta(V)} \tau^{\mathbf{M}_t^*}(\mu^*(h^t), \mu') \left[q^{\mathbf{M}_t^*}(\mu') (\mu' v_H + (1 - \mu') \hat{v}_L(\mu^*(h^t))) + (1 - q^{\mathbf{M}_t^*}(\mu')) \times \right. \\ \left. \delta \left(U_S^*(h^{t+1}) + \left(\frac{\mu'}{1 - \mu'} - \frac{\mu^*(h^t)}{1 - \mu^*(h^t)} \right) (1 - \mu') U_{H|L}^*(h^{t+1}) \right) \right], \quad (7)$$

where $\sum_{\mu' \in \Delta(V)} \tau^{\mathbf{M}_t^*}(\mu^*(h^t), \mu') \mu' = \mu^*(h^t)$.

Part 2 implies that this payoff is the best one the seller can achieve in $\mathcal{E}^*(\mu^*(h^t))$, consistent with the buyer's payoff from h^t onwards. That is, from amongst all the mechanisms that satisfy the participation and incentive compatibility constraints given the buyer's equilibrium payoffs, $\mathbf{M}_t^*$ delivers the seller the highest payoff at h^t .

As in [Equation 6](#), [Equation 7](#) shows that the seller's mechanism at h^t can be thought of as a distribution over posteriors, $\tau^{\mathbf{M}_t^*}$, and a probability of trade, $q^{\mathbf{M}_t^*}$. $(\tau^{\mathbf{M}_t^*}, q^{\mathbf{M}_t^*})$ are chosen so as to maximize a version of the virtual surplus, evaluated at the seller's prior belief at h^t , $\mu^*(h^t)$. As in the discussion of [Equation 6](#), the virtual value $\hat{v}_L$ is evaluated at $\mu^*(h^t)$ because $\mu^*(h^t)$ is the probability the seller assigns to the buyer's valuation being v_H at h^t .

Moreover, by expanding the above expressions and using the Bayes' consistency conditions, we can show that if $h^{t+1} = (h^t, \mathbf{M}_t^*, 1, \mu', (0, x^{\mathbf{M}_t^*}(\mu')))$ is a history on

¹⁷[Equation 7](#) is obtained by replacing the corresponding versions of equations 3-5 at history h^t .

the path of $(\Gamma^*, (\pi_v^*, r_v^*)_{v \in V})$, then:

$$\begin{aligned}
U_S^*(h^{t+1}) + \left(\frac{\mu'}{1-\mu'} - \frac{\mu^*(h^t)}{1-\mu^*(h^t)} \right) (1-\mu') U_{H|L}^*(h^{t+1}) = \\
\sum_{\tilde{\mu} \in \Delta(V)} \tau^{\mathbf{M}_{t+1}^*}(\mu', \tilde{\mu}) \left[q^{\mathbf{M}_{t+1}^*}(\tilde{\mu}) (\tilde{\mu} v_H + (1-\tilde{\mu}) \hat{v}_L(\mu^*(h^t))) + (1-q^{\mathbf{M}_{t+1}^*}(\tilde{\mu})) \times \right. \\
\left. \delta \left(U_S^*(h^{t+2}) + \left(\frac{\tilde{\mu}}{1-\tilde{\mu}} - \frac{\mu^*(h^t)}{1-\mu^*(h^t)} \right) (1-\tilde{\mu}) U_{H|L}^*(h^{t+2}) \right) \right] \quad (8)
\end{aligned}$$

where we use that in history h^{t+1} the seller assigns probability μ' to the buyer's valuation being v_H .

The above expression shows how the term $\left(\frac{\mu'}{1-\mu'} - \frac{\mu^*(h^t)}{1-\mu^*(h^t)} \right) (1-\mu') U_{H|L}^*(h^{t+1})$ adjusts the seller's continuation values. If at history h^t , the seller's mechanism does not sell the good to the buyer with some probability, he understands that the buyer receives continuation rents $U_{H|L}^*(h^{t+1})$, which have to be reflected in the payments of $\mathbf{M}_t^*$. Because the seller assigns probability $\mu^*(h^t)$ to the buyer's valuation being v_H , he pays those rents with probability $\mu^*(h^t)$. Hence, he adjusts the continuation values to reflect his perceived probability of paying these future rents, so that from his perspective, the virtual value $\hat{v}_L(\cdot)$ is computed at $\mu^*(h^t)$ at all periods after t .

By contrast, [Equation 7](#) evaluated at period $t+1$ shows that when the seller's prior is μ' , he chooses his mechanism taking into account that he leaves rents with probability μ' to the buyer. Whenever $\mu' \neq \mu^*(h^t)$, the seller in period $t+1$ does not internalize the cost he may impose on his "predecessor" because he evaluates leaving rents to the high-valuation buyer differently than how the seller in period t does.

One last implication of incentive efficiency is that, along the equilibrium path, the seller's beliefs together with the buyer's rents pin down the strategy profile:

Corollary 1. Let $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$ be the PBE assessment that delivers payoff $u_S^*(\mu_0)$. If there are two histories, h^τ, h^t , on the path of play such that $\mu^*(h^t) = \mu^*(h^\tau)$ and $U_{H|L}^*(h^t) = U_{H|L}^*(h^\tau)$, then $\Gamma^*(h^\tau)$ and $\Gamma^*(h^t)$ are payoff equivalent for the seller.

[Corollary 1](#) implies that along the equilibrium path, we can think of $U_S^*(h^t)$ and $U_{H|L}^*(h^t)$ as functions of the seller's belief $\mu^*(h^t)$. In a slight abuse of notation, we denote them by $U_S^*(\mu^*(h^t))$ and $U_{H|L}^*(\mu^*(h^t))$.

Corollary 1, Equation 7, and Equation 8 are the basis for the recursive formulation we use to characterize the seller's highest equilibrium payoff, $u_S^*(\mu_0)$. Indeed, let h^t be a history such that the seller assigns probability μ^* to v_H . Define

$$R^*(\mu', \mu^*) = U_S^*(\mu') + \left(\frac{\mu'}{1-\mu'} - \frac{\mu^*}{1-\mu^*}\right)(1-\mu')U_{H|L}^*(\mu'),$$

and let

$$R^*(\mu^*, \mu^*) = \sum_{\mu' \in \Delta(V)} \tau^{\mathbf{M}_i^*}(\mu^*, \mu') (q^{\mathbf{M}_i^*}(\mu')(\mu'v_H + (1-\mu')\hat{v}_L(\mu^*)) + (1-q^{\mathbf{M}_i^*}(\mu'))\delta R^*(\mu', \mu^*)). \quad (9)$$

Consider now the following problem:

Problem 1. Find a policy $(\tau, q) : \Delta(V) \mapsto \Delta(\Delta(V)) \times [0, 1]^{\Delta(V)}$ and a value function $R^{(\tau, q)}$ such that¹⁸

1. For all $\mu_0 \in \Delta(V)$, $\int \mu' \tau(\mu_0, d\mu') = \mu_0$,
2. For all $\mu_0, \mu' \in \Delta(V)$ such that $\mu_0 < 1$,

$$R^{(\tau, q)}(\mu', \mu_0) = \int \left[q(\mu', \tilde{\mu})(\tilde{\mu}v_H + (1-\tilde{\mu})\hat{v}_L(\mu_0)) + \delta(1-q(\mu', \tilde{\mu}))R^{(\tau, q)}(\tilde{\mu}, \mu_0) \right] \tau(\mu', d\tilde{\mu}),$$

3. For all $0 \leq \mu_0 < 1$,

$$R^{(\tau, q)}(\mu_0, \mu_0) \geq \int \left[q'(\mu_0, \mu')(\mu'v_H + (1-\mu')\hat{v}_L(\mu_0)) + \delta(1-q'(\mu_0, \mu'))R^{(\tau, q)}(\mu', \mu_0) \right] \tau'(\mu_0, d\mu')$$

for all $\tau'(\mu_0, \cdot) \in \Delta\Delta(V)$ such that $\int \mu' \tau'(\mu_0, d\mu') = \mu_0$ and all $q'(\mu_0, \cdot) : \Delta(V) \mapsto [0, 1]$. For $\mu_0 = 1$, $q(1, 1) = 1$.

A solution to Problem 1 specifies (i) for each prior, μ_0 , a Bayes' plausible distribution over posteriors (part 1) and a probability of trade for each induced posterior, and (ii) continuation values that are consistent with the policy (part 2). Moreover, $\langle (\tau, q), R^{(\tau, q)} \rangle$ is a solution if for each prior belief the seller may hold, he does not have a one-shot deviation from the policy, (τ, q) , given the continuation values (part 3).

¹⁸In what follows, we abuse notation and denote $q(\mu_0)(\mu') \in [0, 1]$ by $q(\mu_0, \mu')$.

Standard dynamic programming arguments imply that if we find a solution, $\langle (\tau^*, q^*), R^{(\tau^*, q^*)} \rangle$, to [Problem 1](#), we will have found a solution to the system defined by [Equation 7](#) for every h^t on the path of play. Because this system of equations satisfies the necessary conditions in [Proposition 1](#), $R^{(\tau^*, q^*)}(\mu_0, \mu_0)$ is an upper bound on $u_S^*(\mu_0)$. If we then show that there is a PBE that attains $R^{(\tau^*, q^*)}(\mu_0, \mu_0)$, we have characterized the seller's best PBE payoff using the solution to the recursive program.

3.2 Intrapersonal game

[Problem 1](#) is not a decision problem, but a *game*.¹⁹ The players in this game represent the seller holding different beliefs about the buyer's valuation being v_H . Each player chooses a distribution over posteriors that averages out to his prior and, for each posterior that is induced, a probability of trade. In this way, [Problem 1](#) is a formal representation of an idea already present in the literature on the sale of a durable good: when the seller cannot commit not to lower the prices in the future, the seller today competes against his future selves.²⁰

Contrast the analysis so far with the solution if we did not require that the seller's strategy be sequentially rational. An analogue of parts 3 and 4 of [Proposition 1](#) would also hold in that case. Moreover, we would use them to replace

¹⁹While at first glance [Problem 1](#) looks like a dynamic programming formulation of the seller's problem in [Equation 6](#), it is not. After all we cannot apply Bellman's principle of optimality since the strategy that is optimal for the seller at $t = 0$ is not necessarily the one that continues to be optimal as time goes by and the seller accumulates information about the buyer's type. Instead [Problem 1](#) takes the perspective that the seller in each period is a distinct player who chooses his action so as to maximize his payoff, based on the anticipation of what others will do. [Problem 1](#) corresponds then to a sequential game with infinitely many players (one for each belief the seller may have), who each act only once. This approach is common in the literature on time inconsistency that followed the seminal work of [Strotz \(1955\)](#); see [Vieille and Weibull \(2009\)](#) and [Björk and Murgoci \(2014\)](#) for excellent discussions of the differences with (time-consistent) dynamic programming.

²⁰[McAfee and Wiseman \(2008\)](#) put forward this intuition in the introduction to their paper. In a sense, it already appears in [Coase \(1972\)](#) where Coase points out

[...] In these circumstances, why should the landowner continue to hold MQ off the market? The original landowner could obviously improve his position by selling more land since he could by these means acquire more money. It is true that this would reduce the value of the land OM owned by those who had previously bought land from him loss would fall on them, not on him.

the transfers out of the seller's problem to obtain a program in which the seller maximizes the virtual surplus. However, under commitment, the maximization of the virtual surplus is now a *decision problem*: in the event that he does not sell the good, the seller does not need to consider that his mechanism may not be optimal given the information he has now learned.

Notwithstanding the aforementioned difference, [Problem 1](#) shares a useful similarity with the usual approach when the designer has commitment: in both cases, the buyer's behavior has been reduced to a system of equations and is no longer a player.²¹ Moreover, under the assumption of transferable utility, these equations allow us to express the seller's problem only in terms of choosing the allocation and the communication device. This result is a consequence of the application of the revelation principle in [Doval and Skreta \(2018\)](#), and thus, we expect that a similar simplification is feasible in settings other than the one we study here.

Problems like [Problem 1](#) have a long tradition in economics. Indeed, it is an example of an intrapersonal equilibrium (see [Pollak \(1968\)](#); [Peleg and Yaari \(1973\)](#); [Harris and Laibson \(2001\)](#); [Bernheim et al. \(2015\)](#)).²² For future reference, we record this observation in [Definition 3](#) below.

Definition 3. A policy $(\tau, q) : \Delta(V) \mapsto \Delta(\Delta(V)) \times [0, 1]^{\Delta(V)}$ and a value function $R^{(\tau, q)}$ constitute an *intrapersonal equilibrium* if they solve [Problem 1](#).

The main result of this section shows that an intrapersonal equilibrium exists in the game that we study. This is stated formally in [Theorem 2](#) below (note the sequence in the statement is the one described after the statement of [Theorem 1](#)).

Theorem 2. *There exists a unique intrapersonal equilibrium, $\langle (\tau^*, q^*), R^{(\tau^*, q^*)} \rangle$. Indeed, there exists a sequence $0 = \bar{\mu}_0 < \bar{\mu}_1 = v_L/v_H < \dots < \bar{\mu}_n < \dots$ such that if $\mu_0 \in [\bar{\mu}_i, \bar{\mu}_{i+1})$,*

1. *and $i = 0$, then $\tau^*(\mu_0, \mu_0) = 1, q^*(\mu_0, \mu_0) = 1$, while*

²¹Contrast this to the approach in [Freixas et al. \(1985\)](#); [Laffont and Tirole \(1987, 1988\)](#); [Bester and Strausz \(2001\)](#); [Gerardi and Maestri \(2018\)](#) among others, where the designer can offer menus of contracts and observes the agent's choice out of the menu. In these papers, the offered menu together with the agent's choice out of the menu determine the designer's beliefs about the agent's type. This makes it difficult to separate the design of the allocation from the "design" of the information that is revealed by the choice out of the menu.

²²Indeed, the system defined by [Equation 7](#) and [Equation 8](#) has analogues in the work of [Harris and Laibson \(2001\)](#) and [Bernheim et al. \(2015\)](#) on hyperbolic discounting.

2. if $i \geq 1$, then $\tau^*(\mu_0, 1) = 1 - \tau^*(\mu_0, \bar{\mu}_{i-1}) = (\mu_0 - \bar{\mu}_{i-1}) / (1 - \bar{\mu}_{i-1})$ and $q^*(\mu_0, 1) = 1 = 1 - q^*(\mu_0, \bar{\mu}_{i-1})$.

Theorem 2 says that when the seller's prior is low (i.e., $\mu_0 < \bar{\mu}_1$), he sells with probability 1 and transmits no information, whereas when the seller's prior is high (i.e., $\bar{\mu}_1 \leq \mu_0$), he induces two posteriors: one in which trade happens ($\mu' = 1$) and one in which trade is delayed ($\mu' = \bar{\mu}_{i-1}$).

Theorem 2 is the basis for the equilibrium constructed in **Theorem 1**. Indeed, the communication device and the probability of trade used by the seller to achieve his maximum PBE payoff are those from the intrapersonal equilibrium.²³ **Theorem 2**, however, says nothing about the transfers or the buyer's behavior in the PBE assessment. After all, in the intrapersonal game, the seller is the only player.

The proof of **Theorem 2** is constructive. The rest of this section sketches out the main steps and provides intuition for **Theorem 2**. The reader interested in understanding how we move from the intrapersonal equilibrium to a PBE assessment that delivers payoff $u_S^*(\mu_0)$ may skip straight to **Section 3.3**, where we tackle the last step of our construction.

Necessary conditions: To show that an intrapersonal equilibrium exists and is unique, we proceed as follows. **Proposition 2** and **Proposition 3** below describe necessary conditions that an intrapersonal equilibrium must satisfy. We use these properties to build a candidate policy, (τ^*, q^*) , and its continuation values, $R^{(\tau^*, q^*)}$, and show they constitute an equilibrium.

Proposition 2. Let $\langle (\tau^*, q^*), R^{(\tau^*, q^*)} \rangle$ denote an intrapersonal equilibrium. Then, the following hold:

1. If $\mu_0 < \bar{\mu}_1$, then $q^*(\mu_0, \mu') = 1$ for all $\mu' \in \Delta(V)$,
2. If $\mu_0 > \bar{\mu}_1$, then the seller places positive probability on two beliefs, $\{\mu_D(\mu_0), 1\}$, and sets $q^*(\mu_0, \mu_D(\mu_0)) = 0$ and $q^*(\mu_0, 1) = 1$.

An incomplete intuition for **Proposition 2** is as follows. To understand part 1, recall that the solution when the seller has commitment and his prior is below $\bar{\mu}_1$ is to sell to both types of the buyer with probability 1. This is still a solution under limited commitment and is clearly the best that the seller can do. To understand

²³More precisely, the communication device used by the seller is derived from the distribution over posteriors in the intrapersonal equilibrium.

part 2, note that in the intrapersonal equilibrium, the seller's problem is as follows:

$$\begin{aligned} R^{(\tau^*, q^*)}(\mu_0, \mu_0) &= \max_{\tau, q} \int_0^1 \left(q(\mu')(\mu' v_H + (1 - \mu') \hat{v}_L(\mu_0)) + (1 - q(\mu')) \delta R^{(\tau^*, q^*)}(\mu', \mu_0) \right) \tau(d\mu') \\ &= \max_{\tau} \int_0^1 \max\{\mu' v_H + (1 - \mu') \hat{v}_L(\mu_0), \delta R^{(\tau^*, q^*)}(\mu', \mu_0)\} \tau(d\mu'). \end{aligned} \quad (10)$$

Indeed, given the continuation values, the seller can obtain whatever is best between selling today (with value $\mu' v_H + (1 - \mu') \hat{v}_L(\mu_0)$) and delaying trade (with value $\delta R^{(\tau^*, q^*)}(\mu', \mu_0)$), by choosing $q(\cdot)$ appropriately.

Equation 10 shows the seller's problem is like an *information design* problem.²⁴ Starting from an equilibrium $\langle (\tau^*, q^*), R^{(\tau^*, q^*)} \rangle$, it follows that the seller always has a best response in which he uses at most two posteriors. Conditional on inducing two posteriors, it follows that the seller sets $q = 0$ for one and $q = 1$ for the other (Proposition C.1 in Section C.1). That the seller induces a posterior of 1 when he sets $q = 1$ is intuitive: the largest payoff he can get when he sells is v_H . By setting $q = 0$ for at most one posterior, the seller maximizes the probability with which he induces a posterior of 1.²⁵

The reason this intuition is incomplete is that we have only argued that there is a best response in which the seller uses at most two posteriors, not that *every* best response involves two posteriors. While the seller with belief μ_0 is indifferent between all his best responses, this does not mean that a seller with belief μ'_0 is indifferent between all the solutions to Equation 10. Thus, starting from an intrapersonal equilibrium $\langle (\tau^*, q^*), R^{(\tau^*, q^*)} \rangle$, we cannot just replace the policy $(\tau^*(\mu_0, \cdot), q^*(\mu_0, \cdot))$ for the one that solves the maximization in Equation 10 and that induces at most two posteriors. This may affect the best response condition of a seller with belief $\mu'_0 \neq \mu_0$ that either (i) assigns positive probability to μ_0 in the candidate equilibrium and sees his payoff changed as a result of the change in the policy for μ_0 , or (ii) does not assign positive probability to μ_0 , but may now prefer to do so.

Key to our result is then to show that the seller uses at most two posteriors for any prior μ_0 in any solution to Problem 1 (Proposition III.2 in Appendix III

²⁴However, we cannot just invoke concavification-style arguments because the continuation values, $R^{(\tau^*, q^*)}$, may fail to be upper semi-continuous (see the discussion at the end of this section).

²⁵The linearity of the seller's payoff in the allocation and in the posterior beliefs in the event of trade matters for this argument. The same linearity obtains when the buyer's valuations are drawn from a continuum.

in Doval and Skreta (2019)). **Proposition C.1** implies that whenever the seller induces a posterior at which the allocation is $q = 0$, his beliefs drop, i.e., the induced posterior is below the prior. That is, whenever the seller does not trade, he becomes more pessimistic about the buyer's valuation being high, which increases his temptation to lower the price to v_L .

Policies (τ, q) where the seller uses more than two posteriors could help the seller slow down the rate at which his beliefs drop. Consider, for instance, a seller with belief μ_0 who induces three posteriors $\mu'_1 < \mu'_2 < \mu_0 < 1$. Any seller with prior $\mu'_0 > \mu_0$ who induces a posterior equal to μ_0 prefers the "three-way" split over the policy in which the seller with belief μ_0 just induces posteriors μ'_1 and 1, since the three-way split implies a slower decay in beliefs. This kind of slow and probabilistic "rationing" exploits analogous forces as those in the construction in Ausubel and Deneckere (1989) and is behind the optimality of rationing in Denicolo and Garella (1999). Ruling out that these more complex policies can be part of an intrapersonal equilibrium is probably where the assumption of binary valuations matters the most. The proof of **Proposition III.2** exploits the properties of intrapersonal equilibria where for each prior, μ_0 , the seller induces at most two posteriors listed in **Proposition 3** below, to which we turn next.

The result in **Proposition 2** implies we can reduce the construction of the intrapersonal equilibrium to the construction of the belief $\mu_D(\mu_0)$ at which the seller sets $q = 0$ when $\mu_0 > \bar{\mu}_1$. Suppose $\langle \mu_{D^*}, R^{\mu_{D^*}} \rangle$ is an intrapersonal equilibrium, where $R^{\mu_{D^*}}$ are the continuation values implied by the policy that sets $q^*(\mu_0, \mu_{D^*}(\mu_0)) = 0$. The martingale property of beliefs implies

$$\tau^*(\mu_0, \mu_{D^*}(\mu_0)) = \frac{1 - \mu_0}{1 - \mu_{D^*}(\mu_0)},$$

so that part 2 of **Problem 1** implies

$$R^{\mu_{D^*}}(\mu_0, \mu_0) = \frac{\mu_0 - \mu_{D^*}(\mu_0)}{1 - \mu_{D^*}(\mu_0)} v_H + \frac{1 - \mu_0}{1 - \mu_{D^*}(\mu_0)} R^{\mu_{D^*}}(\mu_{D^*}(\mu_0), \mu_0). \quad (11)$$

Proposition 3 characterizes the properties of μ_{D^*} , under the assumption that such an equilibrium exists:

Proposition 3. *Suppose $\langle \mu_{D^*}, R^{\mu_{D^*}} \rangle$ is an intrapersonal equilibrium. Then, the following hold:*

1. *If $\bar{\mu}_1 < \mu_0$ is such that $\mu_{D^*}(\mu_0) < \bar{\mu}_1$, then $\mu_{D^*}(\mu_0) = 0$.*

2. If $\bar{\mu}_1 < \mu_0$, then there exists $N_{\mu_0} < \infty$ such that the seller's belief drops below $\bar{\mu}_1$ after N_{μ_0} periods, i.e., $\mu_{D^*}^{(N_{\mu_0})}(\mu_0) \leq \bar{\mu}_1$.
3. If $\mu_0 < \mu'_0$, then $\mu_{D^*}(\mu_0) \leq \mu_{D^*}(\mu'_0)$.
4. If $N_{\mu_0} = N_{\mu'_0}$, then $\mu_{D^*}(\mu_0) = \mu_{D^*}(\mu'_0)$.

The proof is in Appendix C. In what follows, we provide intuition for the result, which in turn also provides intuition for the structure of the intrapersonal equilibrium in Theorem 2.

Part 1 is immediate. Recall from Proposition 2 that a seller with prior $\mu'_0 < \bar{\mu}_1$ trades immediately. Thus, if a seller with prior $\bar{\mu}_1 < \mu_0$ sets $\mu_{D^*}(\mu_0) = \tilde{\mu} \in (0, \bar{\mu}_1)$, he knows that he trades with v_L in the next period. That is, $R^{\mu_{D^*}}(\tilde{\mu}, \mu_0) = \tilde{\mu}v_H + (1 - \tilde{\mu})\hat{v}_L(\mu_0)$. By setting $\mu_{D^*}(\mu_0) = 0$, the seller with prior μ_0 maximizes the probability of trading with the high-valuation buyer today.

Part 2 says that trade with the buyer when her valuation is v_H happens in finitely many periods. It does not say, however, that the same is true when the buyer's valuation is v_L , unless $0 < \bar{\mu}_1$. We now explain why the seller taking infinitely many periods for his belief to update below $\bar{\mu}_1$, conditional on the event of not allocating the good, cannot be part of an intrapersonal equilibrium. Starting from any prior, $\bar{\mu}_1 < \mu_0$, the repeated application of the function μ_{D^*} induces a decreasing sequence of beliefs, $\mu_m = \mu_{D^*}^{(m)}(\mu_0)$, all of them above $\bar{\mu}_1$. But if μ_m remains above $\bar{\mu}_1$ for every finite m , the probability $\tau^*(\mu_m, 1)$ becomes small. Furthermore, a seller with a prior μ_m always has the possibility of placing weight μ_m on 1 and the remaining weight on 0. We show that for m large enough, this deviation is profitable. Thus, whereas the seller with prior belief, μ_0 , would benefit from his successors never splitting their beliefs below $\bar{\mu}_1$, taking infinitely many periods to update below $\bar{\mu}_1$ cannot be part of an intrapersonal equilibrium.

Part 2 echoes the results in Fudenberg et al. (1985), Gul et al. (1986), and Ausubel and Deneckere (1989) that in the “gap” case, trade happens with both types of the buyer in finitely many periods.²⁶ We do not phrase it in these terms because when $v_L = 0$, which corresponds to the no-gap case in our model, part 2 holds; however, in the unique intrapersonal equilibrium, the seller with prior $\mu_0 = 0$ does not allocate the good to the buyer whose valuation is v_L . Because once we reach a posterior of 0 beliefs do not change, trade never happens with the low-valuation

²⁶The formal arguments are reminiscent to those in De Fraja and Muthoo (2000) and Condorelli et al. (2016).

buyer.

Parts 3-4 say the policy is monotone in that if $\mu_0 < \mu'_0$, the seller with belief μ'_0 chooses a higher posterior at which to delay trade than the seller with belief μ_0 . This is intuitive. When $0 < \bar{\mu}_1$, the rents for the high-valuation buyer are determined by how long it takes to trade with v_L : the longer it takes, the lower the rents. The seller's prior is also the probability with which he pays rents to the buyer. Thus, the higher the prior, the higher the incentive for the seller to delay trade with v_L so as to make the rents for v_H smaller.²⁷

Part 4 is important in the construction of the intrapersonal equilibrium when $0 < \bar{\mu}_1$. We can classify the priors, μ_0 , according to how many periods it takes for μ_0 to drop below $\bar{\mu}_1$. That is, define

$$D_n = \{\mu_0 \in \Delta(V) : \mu_{D^*}^{(n)}(\mu_0) = 0\}, \quad (12)$$

where we appeal to Proposition 3 to say that the “final” belief at which the seller updates is 0 and we define $D_0 = \{0\}$. Part 4 says that if $\mu_0, \mu'_0 \in D_n$, they delay trade at the same posterior. Otherwise, by setting $\mu_{D^*}(\cdot)$ to be $\min\{\mu_{D^*}(\mu_0), \mu_{D^*}(\mu'_0)\}$, the seller (weakly) increases the probability of immediately trading with v_H without affecting how long it takes to trade with v_L (both $\mu_{D^*}(\mu_0), \mu_{D^*}(\mu'_0)$ correspond to beliefs for which the seller trades in $n - 1$ periods with v_L .)

Part 4 suggests the “smallest” element of D_n ²⁸ is of relevance: this prior, denoted by $\bar{\mu}_n$ in what follows, is the “largest” prior in D_{n-1} and the “smallest” in D_n . The equilibrium policy, (τ^*, q^*) , in Theorem 2 selects the policy of $\bar{\mu}_n$ so that it is an element of D_n . We break the indifference between n and $n - 1$ periods until updating to 0 in favor of n .

The sequence $\{\bar{\mu}_n\}_{n \geq 0}$: We now construct the cutoffs $\bar{\mu}_n$. We set $\bar{\mu}_0 = 0$ and recall that $\bar{\mu}_1$ denotes the ratio v_L/v_H . Note that it satisfies that

$$\bar{\mu}_1 v_H + (1 - \bar{\mu}_1) \hat{v}_L(\bar{\mu}_1) = \bar{\mu}_1 v_H + (1 - \bar{\mu}_1) \delta \underbrace{\hat{v}_L(\bar{\mu}_1)}_{R^{\mu_{D^*}}(0, \bar{\mu}_1)}, \quad (13)$$

because $\hat{v}_L(\bar{\mu}_1) = 0$. Equation 13 shows that, when his prior is $\bar{\mu}_1$, the seller is indifferent between trading today with both types of the buyer, or trading today

²⁷Note that when $\bar{\mu}_1 = 0$, the seller sets $\mu_D(\mu_0) = 0$ for all $0 < \mu_0$: in this case, the seller never pays rents to the buyer, so that for each prior, the seller sells to the high-valuation buyer with the maximum possible probability, μ_0 .

²⁸Technically, at this point, an infimum rather than a minimum.

with the buyer whose type is v_H and waiting until tomorrow to trade with the buyer whose type is v_L . Set $\mu_{D^*}(\bar{\mu}_1) = 0 = \bar{\mu}_0$ and $\tau(\bar{\mu}_1, 0) = 1 - \bar{\mu}_1$. That is, we break the indifference of the seller with belief $\bar{\mu}_1$ between trading immediately with v_L and delaying trade with v_L for one period in favor of delaying trade.

For $n \geq 1$, define inductively $\bar{\mu}_{n+1}$ to be the prior of the seller such that

$$\frac{\bar{\mu}_{n+1} - \bar{\mu}_n}{1 - \bar{\mu}_n} v_H + \frac{1 - \bar{\mu}_{n+1}}{1 - \bar{\mu}_n} \delta R^{\mu_{D^*}}(\bar{\mu}_n, \bar{\mu}_{n+1}) = \frac{\bar{\mu}_{n+1} - \bar{\mu}_{n-1}}{1 - \bar{\mu}_{n-1}} v_H + \frac{1 - \bar{\mu}_{n+1}}{1 - \bar{\mu}_{n-1}} \delta R^{\mu_{D^*}}(\bar{\mu}_{n-1}, \bar{\mu}_{n+1}). \quad (14)$$

That is, when the seller's prior is $\bar{\mu}_{n+1}$, he is indifferent between taking $n + 1$ periods to trade with v_L or taking n periods to trade with v_L . Note that Equation 14 only depends on the policy for $\bar{\mu}_m, m \leq n$, which has already been defined.

Lemma C.3 in the Appendix shows $\{\bar{\mu}_n\}_{n \geq 0}$ is an increasing sequence. These cutoffs are precisely those in the statement of **Theorem 2**. Indeed, we set $D_n = [\bar{\mu}_n, \bar{\mu}_{n+1})$ for $n \geq 0$ and $\mu_{D^*}(\mu_0) = \bar{\mu}_{n-1}$ for $\mu_0 \in D_n$. The policy (τ^*, q^*) satisfies all the properties of **Proposition 3**. We verify in **Section C.3** that it constitutes an intrapersonal equilibrium. Uniqueness follows from the results in Appendix C and Appendix III in Doval and Skreta (2019), and in particular, **Proposition III.1**, where we show that although seller incarnations with beliefs $\{\bar{\mu}_n\}_{n \geq 1}$ have multiple best responses, the one specified above is the only one that can be part of an intrapersonal equilibrium.

Although the intrapersonal game has a unique equilibrium, the game between the seller and the buyer does not (see **Appendix V** in Doval and Skreta (2019)). This finding is a reflection of a deeper observation: we derived the intrapersonal game from necessary conditions for revenue maximization, whereas not all equilibria in the game between the seller and the buyer give the seller his best equilibrium payoff.

Tie-breaking and upper semicontinuity: Before proceeding with the construction of the PBE assessment, $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$, we make a final observation about the properties of the intrapersonal equilibrium. **Figure 1** below illustrates that when $\mu_0 \in [\bar{\mu}_n, \bar{\mu}_{n+1})$, the continuation values $R^{(\tau^*, q^*)}(\cdot, \mu_0)$ fail to be upper semicontinuous at beliefs $\{\bar{\mu}_k : k \geq n + 1\}$.²⁹

²⁹**Lemma IV.2** in Doval and Skreta (2019) shows formally that upper semicontinuity fails at these beliefs.

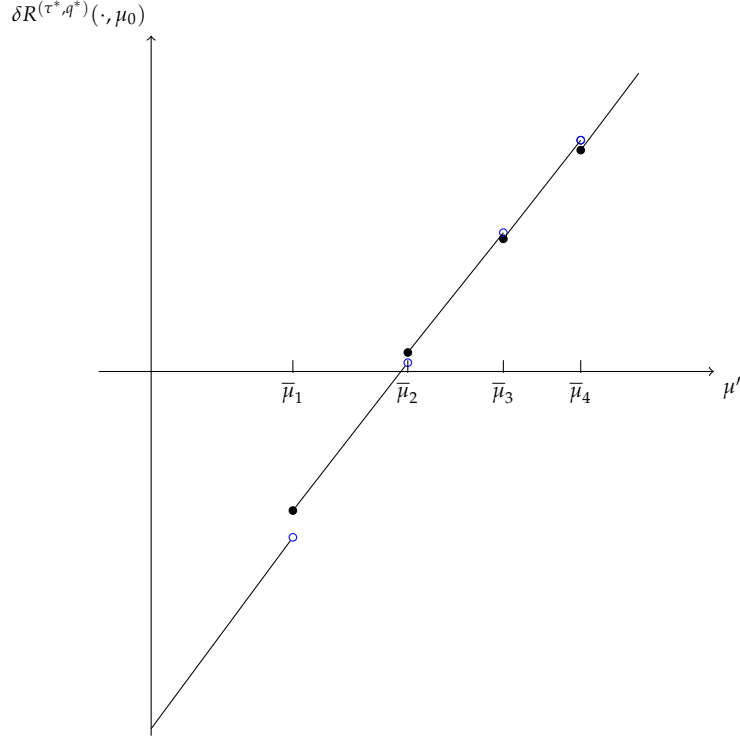


Figure 1: Illustration of the failure of upper-semicontinuity for $\mu_0 \in [\bar{\mu}_2, \bar{\mu}_3)$: blue empty circles indicate the discontinuity

Conceptually, the failure of upper semicontinuity illustrates that sellers with different prior beliefs *disagree* about how ties should be broken. If $\mu_0 \in [\bar{\mu}_n, \bar{\mu}_{n+1})$, the seller with belief μ_0 finds that sellers with beliefs $\bar{\mu}_k, k \geq n+1$ take too long to trade with the buyer when her valuation is v_L . However, inducing posteriors greater than $\bar{\mu}_{n+1}$ when $\mu_0 < \bar{\mu}_{n+1}$ is never optimal: any weight on such posteriors can be split between $\bar{\mu}_{n-1}$ and 1, which are the optimal choices for the seller when his prior $\mu_0 \in [\bar{\mu}_n, \bar{\mu}_{n+1})$. This observation, together with the property that seller incarnations with priors lower than μ_0 break ties in favor of the seller with prior μ_0 , guarantees an intrapersonal equilibrium exists.

Note that conflicts in tie-breaking are usually an issue for equilibrium existence in the literature on intrapersonal games. A similar issue arises in sequential voting games (see [Duggan \(2006\)](#) and the references therein). Two properties of our game seem important in managing to sidestep these issues: (i) the agreement between a seller with belief μ_0 with how sellers with lower priors break ties, and (ii) the

seller with belief $\mu_0 \in [\bar{\mu}_n, \bar{\mu}_{n+1})$ can always choose not to put a seller with belief $\mu' \geq \bar{\mu}_{n+1}$ on the move.

3.3 From the intrapersonal equilibrium to the PBE assessment

Section 3.3 describes how we use the intrapersonal equilibrium policy, (τ^*, q^*) , to construct an assessment, $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$, that gives the seller his highest equilibrium payoff in $G^\infty(\mu_0)$.

Theorem 2 delivers an allocation rule and a distribution over posteriors for each prior belief the seller may hold. We now use the seller's policy in the intrapersonal equilibrium to construct a mapping, $\gamma^* : \Delta(V) \mapsto \mathcal{M}_C$, that assigns a mechanism to each belief the seller may have about the buyer (omitted details of this construction are in Section D.1). Given the seller's prior, μ_0 , we can use the distribution over posteriors, $\tau^*(\mu_0, \cdot)$, to construct a communication device $\beta_{\mu_0}^*$, which under the assumption of truthtelling, implements the same distribution over posteriors as τ^* . Indeed, if $\mu_0 \in D_n$,

$$\beta_{\mu_0}^*(\bar{\mu}_{n-1}|v_L) = 1 \text{ and } \beta_{\mu_0}^*(\bar{\mu}_{n-1}|v_H) = \frac{\bar{\mu}_{n-1}\tau^*(\mu_0, \bar{\mu}_{n-1})}{\mu_0}.$$

Moreover, we can use the allocation rule $q^*(\mu_0, \cdot)$ to construct the probability of trade:

$$q_{\mu_0}^*(\mu') = q^*(\mu_0, \mu') \text{ for } \mu' \in \{\bar{\mu}_{n-1}, 1\}.$$

Finally, using Equations 3 and 4, we can use the allocation rule, $q^*(\mu_0, \cdot)$, and the communication device, $\beta_{\mu_0}^*$, to construct the transfers. Indeed, if $\mu_0 \in D_n$,

$$x_{\mu_0}^*(\bar{\mu}_{n-1}) = 0 \text{ and } x_{\mu_0}^*(1) = v_L + (1 - \delta^n)\Delta v. \quad (15)$$

The mapping γ^* then maps each prior, μ_0 , to $\langle (V, \beta_{\mu_0}^*, \Delta(V)), (q_{\mu_0}^*(\cdot), x_{\mu_0}^*(\cdot)) \rangle$. Moreover, we construct the buyer's rents under this mechanism as a function of the seller's prior. Indeed, if $\mu_0 \in D_n$, the buyer's continuation payoff when her valuation is v_H is given by:

$$u_H^*(\mu_0) = \delta^n \Delta v. \quad (16)$$

To construct a PBE assessment, we need to specify the buyer's and the seller's strategy after every history. We do so in Sections E.1 and E.2 in Appendix E. In

particular, [Section E.2](#) shows how to jointly construct a strategy for the seller, Γ^* , and a system of beliefs, μ^* , such that after every history, the seller's choice of mechanism is determined by $\Gamma^*(h^t) = \gamma^*(\mu^*(h^t))$. Moreover, $\mu^*(h^t)$ is consistent with the buyer's strategy constructed in [Section E.1](#).

The PBE assessment we construct has the following property at all histories h^t . If at history h^t , the seller has beliefs $\mu^*(h^t)$, the seller's payoff is $R^{(\tau^*, q^*)}(\mu^*(h^t), \mu^*(h^t))$ and the buyer's payoff is 0 if her valuation is v_L and is $u_H^*(\mu^*(h^t))$ if her valuation is v_H (recall equation (16)). [Section E.3](#) verifies that given these continuation payoffs, at history h^{t-1} , neither the buyer nor the seller have a one-shot deviation from the equilibrium strategy profile. In the language of [Abreu et al. \(1990\)](#), the strategy at h^{t-1} , together with the continuation payoffs at h^t , decompose the payoffs at h^{t-1} . However, without knowing whether $(R^{(\tau^*, q^*)}(\mu^*(h^t), \mu^*(h^t)), 0, u_H^*(\mu^*(h^t)))$ is itself an equilibrium payoff in $G^\infty(\mu^*(h^t))$, this is not enough to conclude $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$ is a PBE assessment.

In [Doval and Skreta \(2019\)](#), we lay out the definitions and statements needed to show that self-generation techniques apply to our setting, so that the above steps verify $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$ is a PBE assessment. Ours is a dynamic incomplete information game with persistent types, like the one studied in [Athey and Bagwell \(2008\)](#).³⁰ The authors show that under a restriction on the solution concept, which they term Public PBE, equilibrium payoffs can be characterized using techniques like in [Abreu et al. \(1990\)](#).³¹ More precisely, a Public PBE is a PBE assessment where each player's strategy depends on the public history of the game and their current private information.

Although in general the set of Public PBE payoffs is a strict subset of the PBE payoffs, it is not in the game we study.³² As we discussed in [Section 3.1](#), it follows from [Doval and Skreta \(2018\)](#) that it is without loss of generality to focus on PBE assessments where the buyer's strategy depends only on her valuation and the public history. Thus, self-generation techniques can be applied to our game to characterize not only the seller's best equilibrium payoff, but all equilibrium payoffs in our game.

³⁰To be sure, in [Athey and Bagwell \(2008\)](#), types are persistent, but not fully persistent as in our game.

³¹[Cole and Kocherlakota \(2001\)](#) introduce similar techniques, but in dynamic games with *full support*. That is, the information structure in the game is such that no player can infer from his signal realizations that another player has deviated.

³²Notwithstanding this difficulty, [Athey and Bagwell \(2008\)](#) show that, under some conditions, the most collusive outcome can be achieved by using Public PBE strategy profiles.

We close [Section 3](#) by illustrating some qualitative features of the model by means of pictures and discussing the case $v_L = 0$:

Illustrations: Although the recursive nature of the equations defining the cut-offs $\{\bar{\mu}_n\}_{n \geq 0}$ make performing comparative statics somewhat difficult, in what follows, we reproduce the price path and the seller's payoff as a function of the seller's prior and the discount factor for a specific parameterization of the model. In particular, the figures are plotted using $v_L = 1$ and $v_H = 5$ (which implies $\bar{\mu}_1 = 0.2$).

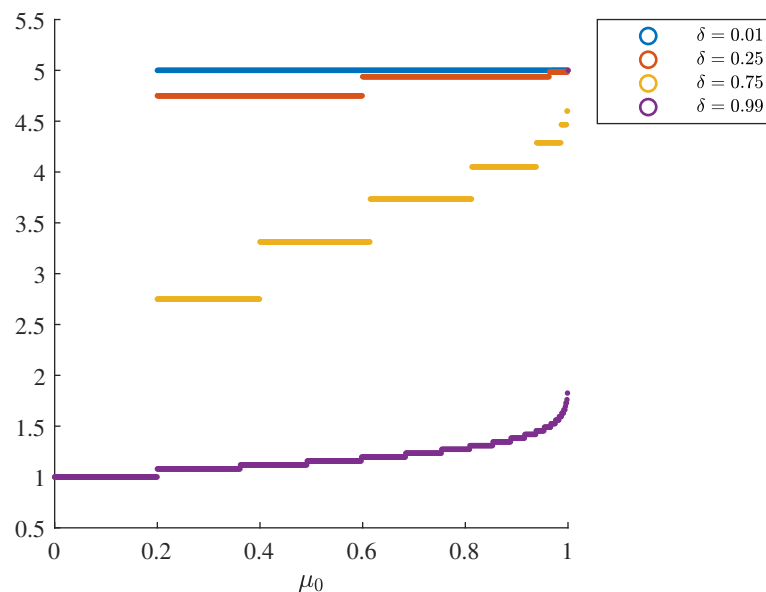


Figure 2: Price dynamics as a function of seller's prior – different discount factors.

[Figure 2](#) plots the price dynamics as a function of the discount factor.³³ Note that for low values of the discount factor ($\delta = 0.01$ in blue), the solution is almost like the commitment solution. In a sense, the seller's impatience affords him commitment power, so that he can keep the prices close to v_H as long as his prior is above $\bar{\mu}_1$. At the other extreme, when the seller's discount factor is high ($\delta = 0.99$ in purple), the logic behind the Coase conjecture applies and the seller quickly low-

³³More precisely, [Figure 2](#) plots the seller's initial price offer as a function of his prior. Given the structure of the PBE assessment that we characterize, the equilibrium price dynamics can be read from the figure.

ers the prices. As a consequence, the seller's revenue is lower for higher discount factors. We illustrate this in [Figure 3](#):

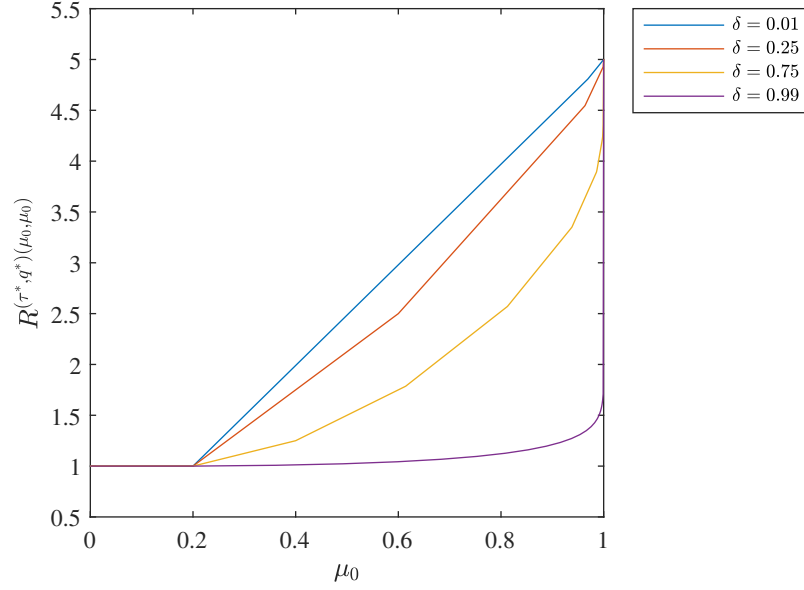


Figure 3: Seller's payoff as a function of seller's prior – different discount factors.

[Figure 3](#) may give the impression that the seller's revenue is monotonically decreasing in the discount factor for a given prior. It is not, however, as [Figure 4](#) shows:

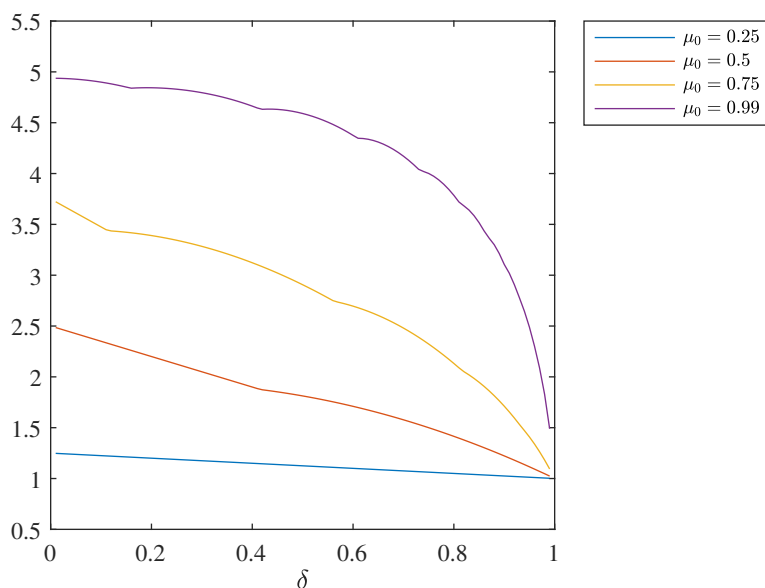


Figure 4: Seller's payoff as a function of the discount factor – different priors.

The case $v_L = 0$ The case of $v_L = 0$ shares features with the *no gap* case in [Ausubel and Deneckere \(1989\)](#). The 'no gap' case refers to the possibility that the seller's value is above the lowest possible buyer's value (in our case this means $v_L \leq 0$). As we argue next, when $v_L = 0$, the seller can sustain the full commitment profit even with limited commitment.

In particular, at the seller optimal equilibrium revenue is $\mu_0 v_H$, regardless of the discount factor (and not just in the limit as in [Ausubel and Deneckere \(1989\)](#)). From the perspective of the intrapersonal equilibrium, regardless of his belief, μ_0 , as long as $\mu_0 > 0$, the seller prefers that the seller with belief 0 breaks ties in favor of no trade. In that case, each seller with belief $\mu_0 > 0$ can trade with maximal probability with v_H because the "threat" of no further trade is credible.

The reader may then be tempted to draw these conclusions more broadly. However, that the seller can achieve the commitment profits when $v_L = 0$ for all discount factors is an artifact of the binary case. With finitely many types, the lowest of which is below 0, the seller would eventually trade with all positive ones in a price-posting game. Thus, with finitely many types, the seller can always avoid dropping the price below the smallest, yet positive, valuation.

4 Conclusions and further directions

This paper makes two contributions to the literature on mechanism design with limited commitment. First, we characterize the optimal mechanism for the sale of a durable good when the seller has limited commitment and interacts with the buyer for infinitely many periods and show that posted prices are optimal.

Second, we provide a *recipe* for solving problems of mechanism design with limited commitment with transferable utility, which applies to settings other than the ones we study here. Indeed, the steps followed in the proof of our main result closely follow the standard steps in classical mechanism design. To wit, the prototypical procedure to analyze problems of mechanism design with commitment and transferable utility follows these steps: (i) Apply the revelation principle to simplify the space of mechanisms and the features of the agent’s strategy profile (participation and truthtelling), (ii) appeal to quasilinearity and single-crossing to derive an envelope representation of payoffs that replaces the transfers out of the designer’s payoff, (iii) solve the designer’s *decision problem*, and (iv) verify that the solution satisfies all constraints that (may) have been ignored.

Note that these steps are mainly the ones we followed in [Section 3](#). Indeed, thanks to the result in our previous work, [Doval and Skreta \(2018\)](#), we are able to reduce the buyer’s behavior to a series of participation and truthtelling constraints, just as we do in the case of full commitment. The main difference is step (iii): with limited commitment, the designer’s problem is an *intrapersonal game*. This is a reflection of a deeper observation: optimal mechanisms under full commitment often fail to be sequentially rational, whereas the best response condition of the intrapersonal equilibrium captures the restrictions imposed by sequential rationality on the designer’s behavior.

The analysis so far has focused on the case of binary valuations, which is consistent with recent papers in the literature that studies infinite-horizon principal-agent problems with limited commitment, like [Strulovici \(2017\)](#) and [Gerardi and Maestri \(2018\)](#). The case of binary valuations allows us to bring to the forefront the conceptual innovations that arise in an infinite horizon game between two long-run players, where one of them, the seller, has a rich action space. Nevertheless, the extension to the case in which buyer valuations are drawn from a continuum is of interest, and we plan to address it in future work.

While the optimality of posted prices maybe specific to our setting, some of the forces we illustrate in the paper are more broadly applicable. Indeed, the richer set

of mechanisms we endow the seller with allow him to design how much he learns about the buyer's type and also how to use this information (subject, of course, to the buyer's incentives). The seller's ability to design his beliefs, and therefore, the demand he faces is relevant in other settings, as is the case, for instance, when the seller interacts repeatedly with the consumer across different transactions. While our setting is not rich enough to capture some of the nuances introduced by repeated purchases, we believe that future work can leverage the framework we have developed to study these issues.

References

- ABREU, D., D. PEARCE, AND E. STACCHETTI (1990): "Toward a theory of discounted repeated games with imperfect monitoring," *Econometrica*, 58, 1041–1063.
- ACHARYA, A. AND J. ORTNER (2017): "Progressive learning," *Econometrica*, 85, 1965–1990.
- ATHEY, S. AND K. BAGWELL (2008): "Collusion with persistent cost shocks," *Econometrica*, 76, 493–540.
- AUSUBEL, L. M. AND R. J. DENECKERE (1989): "Reputation in bargaining and durable goods monopoly," *Econometrica*, 57, 511–531.
- BECCUTI, J. AND M. MÖLLER (2018): "Dynamic adverse selection with a patient seller," *Journal of Economic Theory*, 173, 95–117.
- BERNHEIM, B. D., D. RAY, AND Ş. YELTEKIN (2015): "Poverty and self-control," *Econometrica*, 83, 1877–1911.
- BESTER, H. AND R. STRAUZ (2001): "Contracting with imperfect commitment and the revelation principle: the single agent case," *Econometrica*, 69, 1077–1098.
- (2007): "Contracting with imperfect commitment and noisy communication," *Journal of Economic Theory*, 136, 236–259.
- BJÖRK, T. AND A. MURGOCI (2014): "A theory of Markovian time-inconsistent stochastic control in discrete time," *Finance and Stochastics*, 18, 545–592.
- BOARD, S. AND M. PYCIA (2014): "Outside options and the failure of the Coase conjecture," *American Economic Review*, 104, 656–71.

- BULOW, J. I. (1982): "Durable-goods monopolists," *Journal of Political Economy*, 90, 314–332.
- BURGUET, R. AND J. SAKOVICS (1996): "Reserve prices without commitment," *Games and Economic Behavior*, 15, 149–164.
- CAILLAUD, B. AND C. MEZZETTI (2004): "Equilibrium reserve prices in sequential ascending auctions," *Journal of Economic Theory*, 117, 78–95.
- COASE, R. H. (1972): "Durability and monopoly," *The Journal of Law and Economics*, 15, 143–149.
- COLE, H. L. AND N. KOCHERLAKOTA (2001): "Dynamic games with hidden actions and hidden states," *Journal of Economic Theory*, 98, 114–126.
- CONDORELLI, D., A. GALEOTTI, AND L. RENOU (2016): "Bilateral trading in networks," *The Review of Economic Studies*, 84, 82–105.
- CONLISK, J., E. GERSTNER, AND J. SOBEL (1984): "Cyclic pricing by a durable goods monopolist," *The Quarterly Journal of Economics*, 99, 489–505.
- DE FRAJA, G. AND A. MUTHOO (2000): "Equilibrium partner switching in a bargaining model with asymmetric information," *International Economic Review*, 41, 849–869.
- DEB, R. AND M. SAID (2015): "Dynamic screening with limited commitment," *Journal of Economic Theory*, 159, 891–928.
- DENICOLO, V. AND P. G. GARELLA (1999): "Rationing in a durable goods monopoly," *The RAND Journal of Economics*, 44–55.
- DILMÉ, F. AND F. LI (Forthcoming): "Revenue Management without Commitment: Dynamic Pricing and Periodic Flash Sales," *The Review of Economic Studies*.
- DOVAL, L. AND V. SKRETA (2018): "Mechanism Design with Limited Commitment," *arXiv preprint arXiv:1811.03579*.
- (2019): "Supplement to "Optimal mechanism for the sale of a durable good"," [Click here](#).
- DUGGAN, J. (2006): "Endogenous voting agendas," *Social Choice and Welfare*, 27, 495–530.

- FIOCO, R. AND R. STRAUZ (2015): "Consumer standards as a strategic device to mitigate ratchet effects in dynamic regulation," *Journal of Economics & Management Strategy*, 24, 550–569.
- FREIXAS, X., R. GUESNERIE, AND J. TIROLE (1985): "Planning under incomplete information and the ratchet effect," *The Review of Economic Studies*, 52, 173–191.
- FUDENBERG, D., D. LEVINE, AND J. TIROLE (1985): "Infinite horizon models of bargaining with one-sided uncertainty," in *Game Theoretic Models of Bargaining*, Cambridge University Press, vol. 73, 79.
- FUDENBERG, D. AND E. MASKIN (1986): "The Folk Theorem in Repeated Games with Discounting or with Incomplete Information," *Econometrica*, 54, 533–554.
- FUDENBERG, D. AND J. TIROLE (1991): *Game theory*, MIT Press.
- GARRETT, D. F. (2016): "Intertemporal price discrimination: Dynamic arrivals and changing values," *American Economic Review*, 106, 3275–99.
- GERARDI, D. AND L. MAESTRI (2018): "Dynamic Contracting with Limited Commitment and the Ratchet Effect," Tech. rep.
- GUL, F., H. SONNENSCHN, AND R. WILSON (1986): "Foundations of dynamic monopoly and the coase conjecture," *Journal of Economic Theory*, 39, 155 – 190.
- HARRIS, C. AND D. LAIBSON (2001): "Dynamic choices of hyperbolic consumers," *Econometrica*, 69, 935–957.
- HART, O. D. AND J. TIROLE (1988): "Contract renegotiation and Coasian dynamics," *The Review of Economic Studies*, 55, 509–540.
- HOLMSTRÖM, B. AND R. B. MYERSON (1983): "Efficient and durable decision rules with incomplete information," *Econometrica*, 1799–1819.
- LAFFONT, J.-J. AND J. TIROLE (1987): "Comparative statics of the optimal dynamic incentive contract," *European Economic Review*, 31, 901–926.
- (1988): "The dynamics of incentive contracts," *Econometrica*, 1153–1175.
- (1990): "Adverse selection and renegotiation in procurement," *The Review of Economic Studies*, 57, 597–625.

- LIU, Q., K. MIERENDORFF, X. SHI, AND W. ZHONG (2019): "Auctions with limited commitment," *American Economic Review*, 109, 876–910.
- MCAFEE, R. P. AND D. VINCENT (1997): "Sequentially optimal auctions," *Games and Economic Behavior*, 18, 246–276.
- MCAFEE, R. P. AND T. WISEMAN (2008): "Capacity choice counters the Coase conjecture," *The Review of Economic Studies*, 75, 317–331.
- MYERSON, R. B. (1982): "Optimal coordination mechanisms in generalized principal–agent problems," *Journal of Mathematical Economics*, 10, 67–81.
- ORTNER, J. (2017): "Durable goods monopoly with stochastic costs," *Theoretical Economics*, 12, 817–861.
- PELEG, B. AND M. E. YAARI (1973): "On the existence of a consistent course of action when tastes are changing," *The Review of Economic Studies*, 40, 391–401.
- PESKI, M. (2019): "Alternating-offer bargaining with incomplete information and mechanisms," .
- POLLAK, R. A. (1968): "Consistent planning," *The Review of Economic Studies*, 35, 201–208.
- SIMON, L. K. AND W. R. ZAME (1990): "Discontinuous games and endogenous sharing rules," *Econometrica*, 861–872.
- SKRETA, V. (2006): "Sequentially optimal mechanisms," *The Review of Economic Studies*, 73, 1085–1111.
- (2015): "Optimal auction design under non-commitment," *Journal of Economic Theory*, 159, 854–890.
- SOBEL, J. (1991): "Durable goods monopoly with entry of new consumers," *Econometrica*, 59, 1455–1485.
- SOBEL, J. AND I. TAKAHASHI (1983): "A multistage model of bargaining," *The Review of Economic Studies*, 50, 411–426.
- STOKEY, N. L. (1979): "Intertemporal price discrimination," *The Quarterly Journal of Economics*, 355–371.

- (1981): “Rational expectations and durable goods pricing,” *The Bell Journal of Economics*, 112–128.
- STROTZ, R. H. (1955): “Myopia and inconsistency in dynamic utility maximization,” *The Review of Economic Studies*, 23, 165–180.
- STRULOVICI, B. (2017): “Contract negotiation and the Coase conjecture: A strategic foundation for renegotiation-proof contracts,” *Econometrica*, 85, 585–616.
- VIEILLE, N. AND J. W. WEIBULL (2009): “Multiple solutions under quasi-exponential discounting,” *Economic Theory*, 39, 513.

A Omitted formal statements and equations

A.1 Bayes’ rule where possible

In this section, we define formally what we mean by Bayes’ rule where possible in [Section 2](#). We do so using $G^\infty(\mu_0)$, but it should be clear that the same applies to $G_{\mathcal{M}}^\infty(\mu_0)$. Note that in the game under consideration, the public history h^t represents an information set for the seller with nodes $h_B^t \in H_B^t(h^t)$. In what follows, we use the notation y, y' to denote nodes in the game.

Fix a strategy profile $(\Gamma, (\pi_v, r_v)_{v \in V})$, and two nodes y and y' such that y precedes y' . We can use the strategy profile to define a probability, $P^{(\Gamma, (\pi_v, r_v)_{v \in V})}(y'|y)$, of reaching node y' conditional on being at node y . Extend this probability to all nodes by making it 0 for nodes y' that do not succeed y .

Say that information set h^t precedes information set h^{t+1} , or that h^t, h^{t+1} are consecutive information sets if there exists a mechanism, $\mathbf{M}$, such that either of the following hold:

- (a) there is a posterior, μ' , such that $\sum_{v \in V} \beta^{\mathbf{M}}(\mu'|v) > 0$ ³⁴ and $h^{t+1} = (h^t, \mathbf{M}, 1, \mu', (0, x^{\mathbf{M}}(\mu')), \omega_{t+1})$, or
- (b) $h^{t+1} = (h^t, \mathbf{M}, 0, \emptyset, (0, 0), \omega_{t+1})$.

Fix an assessment, $\langle \Gamma, (\pi_v, r_v)_{v \in V}, \mu \rangle$, and two consecutive information sets h^t, h^{t+1} . Say that h^{t+1} is reached with positive probability from h^t under $\langle \Gamma, (\pi_v, r_v)_{v \in V}, \mu \rangle$,

³⁴In [Doval and Skreta \(2018\)](#), we argue that it is without loss of generality to prune from the tree all histories that correspond to posteriors that cannot be generated with positive probability by the mechanism.

if

$$P^{\langle \Gamma, (\pi_v, r_v)_{v \in V}, \mu \rangle} (h^{t+1} | h^t) \equiv \sum_{y \in h^t, y' \in h^{t+1}} \mu^*(y | h^t) P^{\langle \Gamma, (\pi_v, r_v)_{v \in V} \rangle} (y' | y) > 0$$

Say that h^{t+1} can be reached from h^t through a deviation by the seller if there exists Γ' such that $P^{\langle \Gamma', (\pi_v, r_v)_{v \in V}, \mu \rangle} (h^{t+1} | h^t) > 0$.

Definition 4. An assessment $\langle \Gamma, (\pi_v, r_v)_{v \in V}, \mu \rangle$ satisfies Bayes' rule where possible if for all $t \geq 0$ and for all consecutive h^t, h^{t+1} , $\mu(y' | h^{t+1})$ is obtained via Bayes' rule from $\mu(\cdot | h^t)$ and $\langle \Gamma, (\pi_v, r_v)_{v \in V} \rangle$ if either

1. $P^{\langle \Gamma, (\pi_v, r_v)_{v \in V}, \mu \rangle} (h^{t+1} | h^t) > 0$, or
2. h^{t+1} can be reached from h^t through a deviation by the seller.

A.2 Participation, truthtelling, and Bayes' plausibility constraints

Lemma 1 implies that for any payoff $(u_S, u_H, u_L) \in \mathcal{E}^*(\mu_0)$, we can find a PBE assessment, $\langle \Gamma, (\pi_v, r_v)_{v \in V}, \mu \rangle$, such that the following constraints are satisfied. For each period t and each public history h^t , the buyer must find it optimal to participate in the mechanism chosen by the seller:

$$\begin{aligned} & \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}_t^*}(\mu' | v) (v q^{\mathbf{M}_t^*}(\mu') - x^{\mathbf{M}_t^*}(\mu') + (1 - q^{\mathbf{M}_t^*}(\mu')) \delta U_v(h^t, \mathbf{M}_t^*, 1, \mu', (0, x^{\mathbf{M}_t^*}))) \\ & \geq \delta U_v((h^t, \mathbf{M}_t^*, \emptyset, (0, 0))), \end{aligned} \quad (\text{PC}_{v, h^t})$$

where we use $U_v(h^{t+1})$ to denote the buyer's continuation payoffs at history h^{t+1} when her valuation is v under strategy profile $(\Gamma^*, (\pi_v^*, r_v^*)_{v \in V})$. Second, the buyer must find it optimal to report her type truthfully. That is, for all $v \in V$ and $v' \neq v$, we have that

$$\begin{aligned} & \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}_t^*}(\mu' | v) (v q^{\mathbf{M}_t^*}(\mu') - x^{\mathbf{M}_t^*}(\mu') + (1 - q^{\mathbf{M}_t^*}(\mu')) \delta U_v(h^t, \mathbf{M}_t^*, 1, \mu', (0, x^{\mathbf{M}_t^*}))) \\ & \geq \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}_t^*}(\mu' | v') (v q^{\mathbf{M}_t^*}(\mu') - x^{\mathbf{M}_t^*}(\mu') + (1 - q^{\mathbf{M}_t^*}(\mu')) \delta U_v(h^t, \mathbf{M}_t^*, 1, \mu', (0, x^{\mathbf{M}_t^*}))). \end{aligned} \quad (\text{IC}_{v, h^t})$$

Finally, if the mechanism outputs posterior μ' , it must be that

$$\mu' \left[\sum_{v \in V} \beta^{\mathbf{M}_t^*}(\mu' | v) \mu(h^t)(v) \right] = \mu(h^t) \beta^{\mathbf{M}_t^*}(\mu' | v_H). \quad (\text{BC}_{\mu(h^t)})$$

That is, from the communication device, we can infer a distribution over posteriors $\tau^{\mathbf{M}_t^*}(\mu(h^t), \cdot) \in \Delta(\Delta(V))$ such that $\tau^{\mathbf{M}_t^*}(\mu(h^t), \mu') = \sum_{v \in V} \mu(h^t)(v) \beta^{\mathbf{M}_t^*}(\mu'|v)$.

B Proof of Proposition 1

Part 1: The proof of this part is immediate, and hence we omit it.

Part 2: For this proof, we rely on several facts and results in Doval and Skreta (2019). First, by Proposition I.1, we know that if $\langle \Gamma, (\pi_v, r_v)_{v \in V}, \mu \rangle$ is a PBE assessment of $G^\infty(\mu_0)$, then the continuation payoffs at history h^t are equilibrium payoffs of $G^\infty(\mu(h^t))$. Second, because of the public randomization device, it is without loss of generality to assume that the seller does not randomize his choice of mechanism. Let $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$ denote the PBE assessment that gives the seller payoff $u_S^*(\mu_0)$. Toward a contradiction, suppose that $h^t = (h^{t-1}, \omega, \mathbf{M}_t, 1, \mu', (0, x^{\mathbf{M}_t}))$ is a history on the path of play such that $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle|_{h^t}$ is not incentive-efficient. Thus, $U_S^*(h^t) < \max\{u_S : (u_S, U_H^*(h^t), U_L^*(h^t)) \in \mathcal{E}^*(\mu^*(h^t))\}$. Then, there is a strategy profile in $G^\infty(\mu^*(h^t))$ that gives the seller a higher payoff and the buyer the same payoff at h^t . Consider then modifying $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$ at h^t so that the latter strategy profile is used. Clearly, this increases the seller's payoff without changing the buyer's payoff. Moreover, at history $(h^{t-1}, \omega, \mathbf{M}_t)$, the buyer still has an incentive to participate and truthfully report her type, since we have not modified her continuation payoffs in the mechanism. Finally, note that at (h^{t-1}, ω) , the seller does not have an incentive to deviate: we have not changed the buyer's strategy at $(h^{t-1}, \omega, \mathbf{M}_t')$, for $\mathbf{M}_t' \neq \mathbf{M}_t$. Therefore, since the seller had no incentive to deviate in the original assessment, he also has no incentive to deviate at the new one. The new assessment is thus a PBE assessment in $G^\infty(\mu_0)$ and the seller obtains a higher payoff than $u_S^*(\mu_0)$, a contradiction.³⁵

Part 3: We first show that in the first period, v_L has to be indifferent between participating or not in the mechanism. Let $U_L^*(h^0, \mathbf{M}_0^*)$ denote the buyer's equilibrium payoff at the initial history in the PBE when the seller offers mechanism $\mathbf{M}_0^*$ according to $\Gamma^*(h^0)$. If $U_L^*(h^0, \mathbf{M}_0^*) > 0$, consider the following modification to the seller's strategy profile. At the initial history, instead of offer mechanism $\mathbf{M}_0^*$, the seller offers mechanism $\mathbf{M}_0'$ that coincides with mechanism $\mathbf{M}_0^* \equiv \Gamma^*(\emptyset)$, except that all transfers $x^{\mathbf{M}_0'}(\mu')$ such that $\sum_{v \in V} \beta^{\mathbf{M}_0'}(\mu'|v) > 0$, they are raised by $U_L^*(h^0, \mathbf{M}_0^*)$. Modify the buyer's strategy at h^0 so that $\pi_v^*(h^0, \mathbf{M}_0') =$

³⁵Technically, we just showed that there cannot be a positive measure of realizations of the correlating device for which $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$ is not incentive-efficient.

$\pi_v^*(h^0, \mathbf{M}_0^*)$ and $r_v^*(h^0, \mathbf{M}'_0) = r_v^*(h^0, \mathbf{M}_0^*)$; leave the buyer's strategy everywhere else unchanged. Finally, modify the seller's beliefs so that when the buyer rejects $\mathbf{M}'_0$, he assigns probability 1 to the buyer's valuation being v_H . It is standard to check that the seller's and buyer's sequential rationality constraints hold³⁶ and the seller's payoff increases by $U_L^*(h^0) > 0$.

To see that part 3 holds after every history on the path of play, let h^t denote the shortest-length public history on the path of play such that a mechanism $\mathbf{M}_t^*$ is offered by the seller at h^t such that $U_L^*(h^t, \mathbf{M}_t^*) > 0$. Following similar steps as above, we can modify the mechanism at h^t , so that the buyer's payoff starting from h^t is 0 when her valuation is v_L and the seller's payoff at h^t goes up by $U_L^*(h^t, \mathbf{M}_t^*)$; adjusting the strategies of the buyer so that it is a best response to participate and truthfully report her type in the modified mechanism. Letting $h^t = (h^{t-1}, \mathbf{M}_{t-1}^*, \mu', 0, x^{\mathbf{M}_{t-1}^*}(\mu'))$, modify $\mathbf{M}_{t-1}^*$ so that transfers are lowered by $\delta(1 - q^{\mathbf{M}_{t-1}^*}(\mu'))U_L^*(h^t, \cdot)$. Clearly, the change is revenue neutral for the seller, it satisfies that $U_L^*(h^t, \cdot) = U_L^*(h^{t-1}, \cdot) = 0$, and it is standard to check that the buyer's incentives and payoffs remain unchanged.

Part 4: Let $\langle \Gamma^*, (\pi_v^*, r_v^*)_{v \in V}, \mu^* \rangle$ denote the PBE assessment that delivers the highest payoff for the seller. Clearly, if the mechanism used in the initial history satisfies that

$$|\text{supp } \mu_0(v)\beta^{\mathbf{M}_0^*}(\cdot|v)| = 1,$$

then the buyer is indifferent between reporting v_L and v_H at h^0 .

Consider then the case in which $|\text{supp } \mu_0(v)\beta^{\mathbf{M}_0^*}(\cdot|v)| > 1$. Then, Bayes' consistency implies that two beliefs, μ'_H, μ'_L , exist such that $\beta^{\mathbf{M}_0^*}(\mu'_H|v_H) > \beta^{\mathbf{M}_0^*}(\mu'_H|v_L)$ and $\beta^{\mathbf{M}_0^*}(\mu'_L|v_H) < \beta^{\mathbf{M}_0^*}(\mu'_L|v_L)$. Let $\epsilon_H = U_H^*(h^0, \mathbf{M}_0^*) - U_{H|L}^*(h^0, \mathbf{M}_0^*)$, where $U_H^*(h^0, \mathbf{M}_0^*)$ denotes the equilibrium payoff of the high type at history h^0 and $U_{H|L}^*$ denotes the buyer's payoff when her type is v_H , she reports v_L into the mechanism $\mathbf{M}_0^*$, and then plays according to her equilibrium strategy. By assumption, $\epsilon_H > 0$.

³⁶Because the original assessment is a PBE, sequential rationality of the seller's strategy profile implies that he does not have an incentive to deviate to mechanisms different from $\mathbf{M}_0^*$ and $\mathbf{M}'_0$. Clearly, given the buyer's strategy profile, offering $\mathbf{M}'_0$ dominates offering $\mathbf{M}_0^*$. As for the buyer, note that conditional on participating in $\mathbf{M}'_0$, reporting her valuation truthfully is a best response because it was a best response in the original PBE. Because in the new assessment, participating in $\mathbf{M}'_0$ yields a non-negative payoff for the buyer, whereas rejection guarantees a payoff of 0, participating is a best response for the buyer.

Then, define $\Delta_{\mu'_H}, \Delta_{\mu'_L}$ to be such that

$$\begin{aligned}\beta^{\mathbf{M}_0^*}(\mu'_H|v_L)\Delta_{\mu'_H} &= \beta^{\mathbf{M}_0^*}(\mu'_L|v_L)\Delta_{\mu'_L}, \\ \epsilon_H &= \beta^{\mathbf{M}_0^*}(\mu'_H|v_H)\Delta_{\mu'_H} - \beta^{\mathbf{M}_0^*}(\mu'_L|v_H)\Delta_{\mu'_L}.\end{aligned}$$

Modify the transfers of mechanism $\mathbf{M}_0^*$ so that $x^{\mathbf{M}_0'}(\mu'_L) = x^{\mathbf{M}_0^*}(\mu'_L) - \Delta_{\mu'_L}$, $x^{\mathbf{M}_0'}(\mu'_H) = x^{\mathbf{M}_0^*}(\mu'_H) + \Delta_{\mu'_H}$. Modify the buyer's strategy at h^0 so that $\pi_v^*(h^0, \mathbf{M}_0') = \pi_v^*(h^0, \mathbf{M}_0^*)$ and $r_v^*(h^0, \mathbf{M}_0', 1) = r_v^*(h^0, \mathbf{M}_0^*, 1)$; leave the buyer's strategy everywhere else unchanged. Finally, modify the seller's beliefs so that when the buyer rejects $\mathbf{M}_0'$, he assigns probability 1 to the buyer's valuation being v_H . The seller's revenue goes up by $\mu_0\epsilon_H$ and the buyer's best response is still to participate in the mechanism and truthfully report her type.³⁷

To see that along the path of play, v_H must be indifferent between truthfully reporting and reporting v_L , let h^t denote the shortest-length history on the path of play such that the seller offers a mechanism $\mathbf{M}_t^*$ such that the buyer strictly prefers to report v_H than to report v_L when her type is v_H .³⁸ As before, at least two posteriors must be generated in the mechanism. Let $\epsilon_H(h^t) > 0$ denote the difference in payoffs between reporting v_H and reporting v_L . Following the construction in the previous paragraph, modify the strategy profile at h^t so that the seller, instead of offering $\mathbf{M}_t^*$, offers the modified version $\mathbf{M}_t'$. Mechanism $\mathbf{M}_t'$ is as $\mathbf{M}_t^*$, except that the transfers have been modified as we did for the case $t = 0$, so that the buyer is now indifferent between reporting v_H and v_L , when her valuation is v_H . Modify the buyer's strategy at h^t so that $\pi_v^*(h^t, \mathbf{M}_t') = \pi_v^*(h^t, \mathbf{M}_t^*)$ and $r_v^*(h^t, \mathbf{M}_t', 1) = r_v^*(h^t, \mathbf{M}_t^*, 1)$; for now, leave the buyer's strategy everywhere else unchanged. Finally, modify the seller's beliefs so that when the buyer rejects $\mathbf{M}_t'$, he assigns probability 1 to the buyer's valuation being v_H . To keep the incentives at the histories preceding h^t the same, we proceed as follows. Let $h^t = (h^{t-1}, \mathbf{M}_{t-1}^*, 1, \mu', (0, x^{\mathbf{M}_{t-1}^*}(\mu')))$. Modify the strategy at h^{t-1} so that instead of offering $\mathbf{M}_{t-1}^*$, the seller offers $\mathbf{M}_{t-1}'$. This mechanism coincides with $\mathbf{M}_{t-1}^*$ except that $x^{\mathbf{M}_{t-1}^*}(\mu')$ is lowered by $\Delta_{\mu'} = \delta(1 - q^{\mathbf{M}_{t-1}^*}(\mu'))\epsilon_H(h^t)$. Note that this is

³⁷To see that participating with probability one and truthtelling are best responses for the buyer, note that when the buyer announces v_L , her expected payment (averaging out over posteriors) is unchanged, whereas when she announces v_H her expected payment goes up by ϵ_H . This means that when her valuation is v_H the buyer is now indifferent between reporting v_L and v_H , whereas the buyer now has a strict incentive to report v_L when her valuation is v_L .

³⁸Because without loss of generality we can focus on assessments where the buyer is truthful, if she is not indifferent, she must strictly prefer to tell the truth.

payoff neutral for the seller: his increase in payoff at $(h^t, \mathbf{M}'_t)$ is exactly offset by his payoff decrease at $h^{t-1}, \Delta_{\mu'}$. However, by lowering $x^{\mathbf{M}'_{t-1}}(\mu')$, we may have relaxed the participation constraint of v_L if $1 - \mu^*(h^{t-1}) > 0$, in which case, we increase all transfers by $\epsilon_L = \beta^{\mathbf{M}'_t}(\mu'|v_L)\Delta_{\mu'} > 0$. Again, modify the buyer's strategy so that her best response at $(h^{t-1}, \mathbf{M}'_{t-1})$ is the same as at $(h^{t-1}, \mathbf{M}^*_{t-1})$ and modify the seller's beliefs so that when the buyer rejects $\mathbf{M}_{t-1}$, he assigns probability 1 to the buyer's valuation being v_H . The best response conditions from the original assessment imply the new assessment is also a PBE assessment. In this new assessment, the seller's payoff (weakly) increases, which contradicts incentive efficiency if $1 - \mu^*(h^{t-1}) > 0$.

C Proofs of Section 3.2

Instead of proving Proposition 2, Proposition 3, and Theorem 2 *chronologically*, we proceed as follows:

1. Proposition C.1 in Section C.1 shows that in any intrapersonal equilibrium
 - (a) If $\mu_0 < \bar{\mu}_1$, then $q(\mu_0, \cdot) = 1$ (Proposition 2, Part 1)
 - (b) If $\mu_0 > \bar{\mu}_1$, then
 - i. If the seller induces a posterior at which he sets $q(\mu_0, \mu') = 1$, then $\mu' = 1$
 - ii. The seller never induces a posterior in $(\mu_0, 1)$.
2. Proposition C.2 shows that if there is an intrapersonal equilibrium such that
 - (a) If $\mu_0 < \bar{\mu}_1$, then $q(\mu_0, \cdot) = 1$,
 - (b) For all $\mu_0 \geq \bar{\mu}_1$, there exists $\mu_{D^*}(\mu_0) < \mu_0$ such that $\tau^*(\mu_0, \mu_{D^*}(\mu_0)) = 1 - \tau^*(\mu_0, 1)$ and $q^*(\mu_0, 1) = 1 = 1 - q^*(\mu_0, \mu_{D^*}(\mu_0))$.
 then it satisfies the properties stated in Proposition 3.
3. Theorem C.1 shows that the policy described in Theorem 2 is an intrapersonal equilibrium.

To complete the proofs of the results in Section 3.2 it only remains to show that the intrapersonal equilibrium is unique. This is done in Appendix III in Doval and Skreta (2019). Section III.1 shows that there is a unique equilibrium in which the

seller uses at most two posteriors. [Section III.2](#) shows that there is no equilibrium in which the seller uses more than two posteriors.

C.1 Proof of [Proposition C.1](#)

Proposition C.1. *Let $\langle (\tau^*, q^*), R^{(\tau^*, q^*)} \rangle$ be an intrapersonal equilibrium. Then,*

1. *For all $\mu_0 < \bar{\mu}_1$, $q^*(\mu_0, \cdot) = 1$,*
2. *For all $\mu_0 > \bar{\mu}_1$, if $q^*(\mu_0, \mu') = 1$ and the seller induces μ' with positive probability, then $\mu' = 1$.*
3. *For all $\mu_0 > \bar{\mu}_1$, $\int_{\mu_0}^1 \tau^*(\mu_0, d\mu') = \tau^*(\mu_0, 1)$.*

The proof of [Proposition C.1](#) uses the following lemma:

Lemma C.1 (Properties of the continuation values, $R^{(\tau, q)}$). *Let $R^{(\tau, q)}$ denote the continuation values induced by a policy (τ, q) . Then,*

1. *If $\mu_0 < \bar{\mu}_1$, then $R^{(\tau, q)}(\mu', \mu_0) \leq \mu' v_H + (1 - \mu') \hat{v}_L(\mu_0)$,*
2. *If $\mu_0 > \bar{\mu}_1$, then $R^{(\tau, q)}(\mu', \mu_0) \leq \mu' v_H$.*

Proof. Note the contraction mapping theorem implies $R^{(\tau, q)}$ is well defined for any policy, (τ, q) . Fix a prior μ_0 , and consider the operator T_{μ_0} that takes bounded function $w : \Delta(V) \mapsto \mathbb{R}$ into

$$T_{\mu_0}(w)(\mu') = \int (q(\mu', \tilde{\mu})(\tilde{\mu} v_H + (1 - \tilde{\mu}) \hat{v}_L(\mu_0)) + (1 - q(\mu', \tilde{\mu})) \delta w(\tilde{\mu})) \tau(\mu', d\tilde{\mu})$$

To show part 1 holds, let w be a bounded function such that $w(\tilde{\mu}) \leq \tilde{\mu} v_H + (1 - \tilde{\mu}) \hat{v}_L(\mu_0)$. Then,

$$\begin{aligned} T_{\mu_0}(w)(\mu') &\leq \int (q(\mu', \tilde{\mu})(\tilde{\mu} v_H + (1 - \tilde{\mu}) \hat{v}_L(\mu_0)) + (1 - q(\mu', \tilde{\mu})) \delta(\tilde{\mu} v_H + (1 - \tilde{\mu}) \hat{v}_L(\mu_0))) \tau(\mu', d\tilde{\mu}) \\ &\leq \mu' v_H + (1 - \mu') \hat{v}_L(\mu_0), \end{aligned}$$

where the inequalities follow from (i) the bound on w , (ii) $\hat{v}_L(\mu_0) > 0$, and (iii) Bayes' plausibility. This implies that for all $\mu' \in \Delta(V)$, $R^{(\tau, q)}(\mu', \mu_0) \leq \mu' v_H + (1 - \mu') \hat{v}_L(\mu_0)$.

Similarly, to show part 2 holds, let w be a bounded function such that $w(\tilde{\mu}) \leq \tilde{\mu}v_H$. Then,

$$\begin{aligned} T_{\mu_0}(w(\mu')) &\leq \int (q(\mu', \tilde{\mu})(\tilde{\mu}v_H + (1 - \tilde{\mu})\hat{v}_L(\mu_0)) + (1 - q(\mu', \tilde{\mu}))\delta\tilde{\mu}v_H) \tau(\mu', d\tilde{\mu}) \\ &\leq \mu'v_H, \end{aligned}$$

where the inequalities follow from (i) the bound on w , (ii) $\hat{v}_L(\mu_0) < 0$, and (iii) Bayes' plausibility. This implies that for all $\mu' \in \Delta(V)$, $R^{(\tau, q)}(\mu', \mu_0) \leq \mu'v_H$. $\square$

Proof of Proposition C.1. To prove part 1, note Lemma C.1 implies that

$$R^{(\tau^*, q^*)}(\mu_0, \mu_0) \leq \mu_0v_H + (1 - \mu_0)\hat{v}_L(\mu_0),$$

and this bound is achieved by setting $q^*(\mu_0, \cdot) = 1$.

To prove part 2, note that whenever $\mu_0 > \bar{\mu}_1$, for any posterior μ' it follows that

$$\mu'v_H + (1 - \mu')\hat{v}_L(\mu_0) < \mu'v_H + (1 - \mu')\delta\hat{v}_L(\mu_0), \quad (\text{C.1})$$

because $\hat{v}_L(\mu_0) < 0$ and $\delta < 1$. Thus, instead of setting $q^*(\mu_0, \mu') = 1$, which yields the payoff on the left-hand side of Equation C.1, the seller can split μ' between 0 and 1, setting $q^*(\mu_0, 1) = 1 = 1 - q^*(\mu_0, 0)$, which yields the payoff on the right-hand side of Equation C.1.

To prove part 3, let $A = (\mu_0, 1)$ and suppose that $\int_A \tau^*(\mu_0, d\mu') > 0$. Then,

$$\int_A \delta R^{(\tau^*, q^*)}(\mu', \mu_0) \tau^*(\mu_0, d\mu') \geq \int_A R^{(\tau^*, q^*)}(\mu', \mu_0) \tau^*(\mu_0, d\mu'),$$

otherwise, the seller with prior μ_0 would be better off by generating $\mu' \in A$ and imitating at that point the policy that the seller with belief μ' uses (rather than generating μ' and $q^*(\mu_0, \mu') = 0$, which has payoff $R^{(\tau^*, q^*)}(\mu', \mu_0)$).

Since A has positive measure under $\tau^*(\mu_0, \cdot)$, it follows that the set of posteriors $\mu' > \mu_0$ such that $R^{(\tau^*, q^*)}(\mu', \mu_0) \leq 0$ has positive measure. For any such posterior μ' , we have:

$$\mu'v_H + (1 - \mu')\hat{v}_L(\mu_0) \leq \delta R^{(\tau^*, q^*)}(\mu', \mu_0) \leq 0,$$

where the first inequality follows because it has to be that the seller prefers to induce μ' and set $q^*(\mu_0, \mu') = 0$ over $q^*(\mu_0, \mu') = 1$. Since $\mu_0 < \mu'$, we then have

$$v_L \equiv \mu_0v_H + (1 - \mu_0)\hat{v}_L(\mu_0) < \mu'v_H + (1 - \mu')\hat{v}_L(\mu_0) \leq 0,$$

a contradiction. Thus, it cannot be that $\tau^*(\mu_0, A) > 0$, which completes the proof. $\square$

C.2 Proof of Proposition C.2

Proposition C.2. Suppose $\langle \mu_{D^*}, R^{\mu_{D^*}} \rangle$ is an intrapersonal equilibrium with the properties that:

1. For all $\mu_0 < \bar{\mu}_1$, $q^*(\mu_0, \mu') = 1$ for all $\mu' \in \Delta(V)$.
2. For all $\mu_0 \geq \bar{\mu}_1$, there exists $\mu_{D^*}(\mu_0) < \mu_0$ such that $\tau^*(\mu_0, \mu_{D^*}(\mu_0)) = 1 - \tau^*(\mu_0, 1)$ and $q^*(\mu_0, 1) = 1 = 1 - q^*(\mu_0, \mu_{D^*}(\mu_0))$.

Then, it satisfies parts 1-4 of Proposition 3.

Part 1: Note that if $\mu_{D^*} \equiv \mu_{D^*}(\mu_0) < \bar{\mu}_1$, then $R^{\mu_{D^*}}(\mu_{D^*}, \mu_0) = (1 - \mu_{D^*})\hat{v}_L(\mu_0) + \mu_{D^*}v_H$. Consider the following alternative policy for μ_0 , where $\mu_{D'}(\mu_0) = 0$. This is achieved by $\tau'(\mu_0, 1) = \tau^*(\mu_0, 1) + \mu_{D^*}$ and $\tau'(\mu_0, 0) = 1 - \tau^*(\mu_0, 1) - \mu_{D^*}$, where $\tau^*(\mu_0, \cdot)$ are the weights on μ_{D^*} and 1 in the original best response. The difference in payoffs is given by:

$$\begin{aligned} & \mu_{D^*}v_H + (1 - \tau^*(\mu_0, 1))\delta[\hat{v}_L(\mu) - \mu_{D^*}v_H - (1 - \mu_{D^*})\hat{v}_L(\mu_0)] - \mu_{D^*}\delta\hat{v}_L(\mu_0) = \\ & = \mu_{D^*}v_H + (1 - \tau^*(\mu_0, 1))\delta\mu_{D^*}[\hat{v}_L(\mu_0) - v_H] - \mu_{D^*}\delta\hat{v}_L(\mu) \\ & = \mu_{D^*}(v_H - \delta\hat{v}_L(\mu_0)) - (1 - \tau^*(\mu_0, 1))\mu_{D^*}\delta(v_H - \hat{v}_L(\mu_0)) \\ & = \mu_{D^*}v_H(1 - \tau^*(\mu_0, 1)\delta) - \mu_{D^*}\delta\hat{v}_L(\mu_0)\tau^*(\mu_0, 1) > 0, \end{aligned}$$

where the last inequality follows from noting $\hat{v}_L(\mu_0) < 0$. This contradicts that $\mu_{D^*}(\mu_0)$ is a best response. Hence, part 1 holds.

Part 2: Let $\bar{\mu}_1 \leq \mu_0$ denote the seller's prior. Toward a contradiction, assume that for all finite n , $\mu_n = \mu_{D^*}^{(n)}(\mu_0) > \bar{\mu}_1$. Fix $\epsilon > 0$. Note that there exists M_ϵ such that for all $M_\epsilon \leq m$, we have

$$\mu_m - \mu_{m+1} < \epsilon(1 - \mu_{m+1}), \quad (\text{C.2})$$

so that as m goes to infinity, the seller puts smaller weight on posterior 1. To see that Equation C.2 holds, assume to the contrary that for all m , we have

$$\epsilon(1 - \mu_{m+1}) \leq \mu_m - \mu_{m+1}.$$

Adding up from $m = 0$ to $m = N - 1$, we obtain

$$\epsilon \sum_{i=1}^N (1 - \mu_i) \leq \mu_0 - \mu_N. \quad (\text{C.3})$$

Now, the right-hand side of Equation C.3 is bounded above by $\mu_0 - \bar{\mu}_1$ because $\bar{\mu}_1 < \mu_N$. The left-hand side of Equation C.3 is bounded below by $N\epsilon(1 - \mu_0)$ because $\mu_0 \geq \mu_i$ for all $i \geq 0$. We then obtain the following chain of inequalities:

$$N\epsilon(1 - \mu_0) < \epsilon \sum_{i=1}^N (1 - \mu_i) \leq \mu_0 - \mu_N < \mu_0 - \bar{\mu}_1,$$

and this cannot possibly hold for all N . Thus, Equation C.2 holds.

Now fix $\epsilon > 0$ such that

$$\frac{\epsilon}{1 - \delta(1 - \epsilon)} < \bar{\mu}_1(1 - \delta). \quad (\text{C.4})$$

Note the right-hand side of Equation C.4 is bounded above by $\mu_m(1 - \delta)$ for all $m \geq 0$. Take $M_\epsilon \leq m$ and note

$$R^{\mu_{D^*}}(\mu_m, \mu_m) < \epsilon v_H + (1 - \epsilon)\delta R^{\mu_{D^*}}(\mu_{m+1}, \mu_m). \quad (\text{C.5})$$

To see that Equation C.5 holds, note

$$\begin{aligned} R^{\mu_{D^*}}(\mu_m, \mu_m) &= \frac{\mu_m - \mu_{m+1}}{1 - \mu_{m+1}} v_H + \frac{1 - \mu_m}{1 - \mu_{m+1}} \delta R^{\mu_{D^*}}(\mu_{m+1}, \mu_m) \\ &= \frac{\mu_m - \mu_{m+1}}{1 - \mu_{m+1}} (v_H - \delta R^{\mu_{D^*}}(\mu_{m+1}, \mu_m)) + \delta R^{\mu_{D^*}}(\mu_{m+1}, \mu_m) \\ &< \epsilon v_H + (1 - \epsilon)\delta R^{\mu_{D^*}}(\mu_{m+1}, \mu_m), \end{aligned}$$

where the inequality follows from $v_H > \delta R^{\mu_{D^*}}(\mu_{m+1}, \mu_m)$ and $M_\epsilon \leq m$. Because Equation C.5 holds for all $M_\epsilon \leq m$, applying it recursively, we obtain

$$R^{\mu_{D^*}}(\mu_m, \mu_m) < \frac{v_H \epsilon}{1 - \delta(1 - \epsilon)} + \lim_{N \rightarrow \infty} (\delta(1 - \epsilon))^N R^{\mu_{D^*}}(\mu_{m+N}, \mu_m) = \frac{v_H \epsilon}{1 - \delta(1 - \epsilon)},$$

because $R^{\mu_{D^*}}(\cdot, \mu_m)$ is bounded and $\delta(1 - \epsilon) < 1$. Now, the following policy is always feasible for μ_m : place probability μ_m on 1 and $1 - \mu_m$ on 0. This yields

$$\mu_m(1 - \delta)v_H + v_L \delta,$$

when $\bar{\mu}_1 > 0$ and

$$\mu_m v_H,$$

when $\bar{\mu}_1 = 0$. Suppose $0 < \bar{\mu}_1$ (the other case is similar). Because splitting between 0 and 1 is not optimal, we must have

$$\mu_m(1 - \delta)v_H + v_L\delta \leq R^{\mu_{D^*}}(\mu_m, \mu_m) < \frac{\epsilon}{1 - \delta(1 - \epsilon)}v_H < v_H(1 - \delta)\bar{\mu}_1,$$

where the last inequality follows from [Equation C.4](#). This is a contradiction to $\langle \mu_{D^*}, R^{\mu_{D^*}} \rangle$ being an intrapersonal equilibrium. Thus, starting from μ_0 , a finite N exists such that $\mu_{D^*}^{(N)}(\mu_0) \leq \bar{\mu}_1$.

[Lemma C.2](#) below is used to show that if an intrapersonal equilibrium like the one described in [Proposition C.2](#) exists, then it satisfies part 3 of [Proposition 3](#), and it is the first step in the induction in the proof that it also satisfies part 4:

Lemma C.2. *Let $\bar{\mu}_1 < \mu_0 < \mu_0''$ be such that $\mu_{D^*}(\mu_0) = \mu_{D^*}(\mu_0'') = 0$. Then, for all $\mu_0' \in (\mu_0, \mu_0'')$, $\mu_{D^*}(\mu_0') = 0$.*

Proof. Suppose $\mu_0 < \mu_0' < \mu_0''$ are such that $\mu_{D^*}(\mu_0) = \mu_{D^*}(\mu_0'') = 0$, $\mu_D \equiv \mu_{D^*}(\mu_0') \geq \bar{\mu}_1$.³⁹ Then, the following must be true:

$$\frac{\mu_0' - \mu_D}{1 - \mu_D}v_H + \frac{1 - \mu_0'}{1 - \mu_D}\delta R^{\mu_{D^*}}(\mu_D, \mu_0') \geq \mu v_H + (1 - \mu_0')\delta \hat{v}_L(\mu_0').$$

Rewriting the above expression, we obtain:

$$g(\mu_0') \equiv \left[\frac{\mu_0' - \mu_D}{1 - \mu_D} - \mu_0' \right] v_H + \frac{1 - \mu_0'}{1 - \mu_D} \delta R^{\mu_{D^*}}(\mu_D, \mu_0') + (1 - \mu_0')\delta(-\hat{v}_L(\mu_0')) \geq 0. \quad (\text{C.6})$$

In [Section IV.1](#) in [Doval and Skreta \(2019\)](#), we show the continuation values $R^{\mu_{D^*}}$ can be written as:

$$R^{\mu_{D^*}}(\mu_D, \mu_0') = v_H \underbrace{\left(\frac{\mu_D - \mu_1}{1 - \mu_1} + (1 - \mu_D) \sum_{i=1}^{N_D-1} \frac{\delta^i}{1 - \mu_i} \frac{\mu_i - \mu_{i+1}}{1 - \mu_{i+1}} \right)}_{\alpha(\mu_D)} + \underbrace{(1 - \mu_D)\delta^{N_D}}_{\gamma(\mu_D)} \hat{v}_L(\mu_0'),$$

where $\mu_1 > \mu_2 > \dots > \mu_{N_D-1} \geq \bar{\mu}_1 > 0 = \mu_{N_D}$ are the posteriors generated by the policy at which trade is delayed. Note $R^{\mu_{D^*}}(\mu_D, \mu_0')$ is differentiable in μ_0' and

$$\frac{\partial}{\partial \mu} R^{\mu_{D^*}}(\mu_D, \mu)|_{\mu=\mu_0'} = -\gamma(\mu_D) \frac{\Delta v}{(1 - \mu_0')^2}.$$

³⁹By [Proposition C.2](#), part 1, if $\mu_D(\mu_0') \neq 0$, then $\mu_D(\mu_0') \geq \bar{\mu}_1$.

Therefore, g is differentiable with respect to μ'_0 , and we obtain the following expression for $g'(\mu'_0)$:

$$\begin{aligned}
g'(\mu'_0) &= \frac{\mu_D}{1 - \mu_D} v_H - \frac{\delta R^{\mu_D^*}(\mu_D, \mu'_0)}{1 - \mu_D} + \delta \hat{v}_L(\mu'_0) + \delta(1 - \mu'_0) \left[1 - \frac{\gamma(\mu_D)}{1 - \mu_D} \right] \frac{\Delta v}{(1 - \mu'_0)^2} \\
&= v_H \frac{\mu_D}{1 - \mu_D} - \frac{\delta}{1 - \mu_D} (\alpha(\mu_D) v_H + \gamma(\mu_D) \hat{v}_L(\mu'_0)) + \delta \hat{v}_L(\mu'_0) + \delta \frac{\Delta v}{1 - \mu'_0} \left[1 - \frac{\gamma(\mu_D)}{1 - \mu_D} \right] \\
&= v_H \left[\frac{\mu_D - \delta \alpha(\mu_D)}{1 - \mu_D} \right] + \delta \left(\hat{v}_L(\mu'_0) + \frac{\Delta v}{1 - \mu'_0} \right) \left[1 - \frac{\gamma(\mu_D)}{1 - \mu_D} \right] \\
&= v_H \left[\frac{\mu_D - \delta \alpha(\mu_D)}{1 - \mu_D} \right] + \delta v_H \left[1 - \frac{\gamma(\mu_D)}{1 - \mu_D} \right].
\end{aligned}$$

Now, note $\alpha(\mu_D) \leq \mu_D$. Intuitively, the largest probability that μ_D can trade with v_H is μ_D . To see this formally, we use that⁴⁰

$$\begin{aligned}
\alpha(\mu_D) &= \frac{\mu_D - \mu_1}{1 - \mu_1} + (1 - \mu_D) \sum_{i=1}^{N_D-1} \frac{\delta^i}{1 - \mu_i} \frac{\mu_i - \mu_{i+1}}{1 - \mu_{i+1}} \\
&\leq \frac{\mu_D - \mu^1}{1 - \mu^1} + (1 - \mu_D) \frac{\mu_1}{1 - \mu_1} \sum_{i=1}^{N_D-1} \delta_i,
\end{aligned}$$

where the inequality follows from noting that the expression inside the sum is increasing in μ_i and decreasing in μ_{i+1} , and that $\mu_D - \mu_1 + (1 - \mu_D) \mu_1 \sum_{i=1}^{N_D-1} \delta_i \leq \mu_D(1 - \mu_1)$.

It then follows that $g'(\mu'_0) \geq 0$, which contradicts that μ''_0 finds it optimal to set $\mu_{D^*}(\mu'') = 0$. $\square$

Part 3: We use the following property of a policy where $\mu_{D^*}(\cdot)$ is weakly increasing. If we let N_{μ_0} denote the smallest value of n such that $\mu_{D^*}^{(n)}(\mu_0) \leq \bar{\mu}_1$, then N_{μ_0} is weakly increasing in μ_0 .⁴¹ To see this, let $\mu_0 < \mu'_0$, and inductively define $\mu_{i+1} = \mu_{D^*}(\mu_i)$ and similarly, $\mu'_{i+1} = \mu_{D^*}(\mu'_i)$. Clearly, we have that $\mu_0 > \mu_1 > \dots > \mu_N > \dots$ and similarly, $\mu'_0 > \mu'_1 > \dots > \mu'_N > \dots$. Moreover, monotonicity implies $\mu_1 = \mu_{D^*}(\mu_0) \leq \mu'_1 = \mu_{D^*}(\mu'_0)$ and inductively $\mu_N < \mu'_N$. Lemma C.2 implies that if $\mu_{D^*}(\mu'_N) = 0$, then letting n be the smallest m such that $\mu_m \leq \mu'_N$, we have $\mu_{D^*}(\mu_m) = 0$. Hence, $N_{\mu_0} \leq N_{\mu'_0}$.

⁴⁰See Section IV.1 in Doval and Skreta (2019).

⁴¹Proposition C.2 implies $\mu_{D^*}(\mu_{D^*}^{(N_{\mu_0}-1)}) = 0$.

Let μ_0 denote the smallest prior above $\bar{\mu}_1$ such that there exists $\mu'_0 > \mu_0$ with $\mu_{D^*}(\mu_0) > \mu_{D^*}(\mu'_0)$.⁴²

Now consider a policy for μ_0 that splits the weight $\tau(\mu_0, \mu_{D^*}(\mu_0))$ between $\mu_{D^*}(\mu'_0)$ and 1. The payoff from this policy is:

$$v_H \left[\frac{\mu_0 - \mu_{D^*}(\mu_0)}{1 - \mu_{D^*}(\mu_0)} + \frac{1 - \mu_0}{1 - \mu_{D^*}(\mu_0)} \frac{\mu_{D^*}(\mu_0) - \mu_{D^*}(\mu'_0)}{1 - \mu_{D^*}(\mu'_0)} \right] + \delta \frac{1 - \mu_0}{1 - \mu_{D^*}(\mu'_0)} R^{\mu_{D^*}}(\mu_{D^*}(\mu'_0), \mu_0).$$

Consider the difference between the payoff of the above policy and that of the policy that the seller employs when his prior is μ_0 :

$$\begin{aligned} & \frac{1 - \mu_0}{1 - \mu_{D^*}(\mu_0)} \frac{\mu_{D^*}(\mu_0) - \mu_{D^*}(\mu'_0)}{1 - \mu_{D^*}(\mu'_0)} v_H + \delta \frac{1 - \mu_0}{1 - \mu_{D^*}(\mu'_0)} R^{\mu_{D^*}}(\mu_{D^*}(\mu'_0), \mu_0) \\ & - \delta \frac{1 - \mu_0}{1 - \mu_{D^*}(\mu_0)} R^{\mu_{D^*}}(\mu_{D^*}(\mu_0), \mu_0). \end{aligned} \quad (C.7)$$

Optimality of μ'_0 's policy implies

$$\begin{aligned} & \frac{\mu'_0 - \mu_{D^*}(\mu'_0)}{1 - \mu_{D^*}(\mu'_0)} v_H + \delta \frac{1 - \mu'_0}{1 - \mu_{D^*}(\mu'_0)} R^{\mu_{D^*}}(\mu_{D^*}(\mu'_0), \mu'_0) \\ & \geq \frac{\mu'_0 - \mu_{D^*}(\mu_0)}{1 - \mu_{D^*}(\mu_0)} v_H + \delta \frac{1 - \mu'_0}{1 - \mu_{D^*}(\mu_0)} R^{\mu_{D^*}}(\mu_{D^*}(\mu_0), \mu'_0), \end{aligned}$$

or, equivalently,

$$\delta \left[\frac{R^{\mu_{D^*}}(\mu_{D^*}(\mu'_0), \mu'_0)}{1 - \mu_{D^*}(\mu'_0)} - \frac{R^{\mu_{D^*}}(\mu_{D^*}(\mu_0), \mu_0)}{1 - \mu_{D^*}(\mu_0)} \right] \geq - \frac{(\mu_{D^*}(\mu_0) - \mu_{D^*}(\mu'_0))}{(1 - \mu_{D^*}(\mu_0))(1 - \mu_{D^*}(\mu'_0))} v_H.$$

To show μ_0 can improve by using the new policy, we only need to show

$$\begin{aligned} & \frac{R^{\mu_{D^*}}(\mu_{D^*}(\mu'_0), \mu_0)}{1 - \mu_{D^*}(\mu'_0)} - \frac{R^{\mu_{D^*}}(\mu_{D^*}(\mu_0), \mu_0)}{1 - \mu_{D^*}(\mu_0)} \geq \frac{R^{\mu_{D^*}}(\mu_{D^*}(\mu'_0), \mu'_0)}{1 - \mu_{D^*}(\mu'_0)} - \frac{R^{\mu_{D^*}}(\mu_{D^*}(\mu_0), \mu'_0)}{1 - \mu_{D^*}(\mu_0)} \\ & \Leftrightarrow \\ & (\delta^{N_{\mu_{D^*}(\mu'_0)}} - \delta^{N_{\mu_{D^*}(\mu_0)}})(-\hat{v}_L(\mu_0)) \leq (\delta^{N_{\mu_{D^*}(\mu'_0)}} - \delta^{N_{\mu_{D^*}(\mu_0)}})(-\hat{v}_L(\mu'_0)), \end{aligned}$$

where recall that $N_{\cdot}$ denotes the first time at which updating leads to a posterior below $\bar{\mu}_1$ for a given prior. Note that a sufficient condition for the above

⁴²Note $\mu_0 > \bar{\mu}_1$ because we know $\mu_{D^*}(\bar{\mu}_1) = 0$.

inequality to hold is that $N_{\mu_{D^*}(\mu'_0)} < N_{\mu_{D^*}(\mu_0)}$. That this is indeed the case follows from the observation before the proof because the policy below μ_0 (and hence for $\mu_{D^*}(\mu_0), \mu_{D^*}(\mu'_0)$) satisfies monotonicity.

Because $N_{\mu_{D^*}(\mu'_0)} < N_{\mu_{D^*}(\mu_0)}$ implies the expression in [Equation C.7](#) is non-negative, we see that without loss of generality, we can have $\mu_{D^*}(\mu_0) \leq \mu_{D^*}(\mu'_0)$ and strictly so when $N_{\mu_{D^*}(\mu_0)} \neq N_{\mu_{D^*}(\mu'_0)}$.

Part 4: The proof is by induction on $n \geq 1$. Define $D_0 = \{0\}$, and for $n \geq 1$,

$$D_n = \{\mu_0 : \mu_{D^*}^{(n)} = 0\},$$

to be the set of priors, μ_0 , such that the seller updates his prior to 0 in n periods when his prior is μ_0 . Let $\mathbf{P}(n)$ denote the following inductive statement:

P(n): If $\mu_0, \mu'_0 \in D_n$, then $\mu_{D^*}(\mu_0) = \mu_{D^*}(\mu'_0) \in D_{n-1}$.

[Lemma C.2](#) implies the inductive statement is true for $n = 1$, i.e., $\mathbf{P}(1) = 1$.

Suppose $\mathbf{P}(n) = 1$, and we show that $\mathbf{P}(n+1) = 1$. Let $\mu_0, \mu'_0 \in D_{n+1}$ and assume without loss of generality that $\mu_0 < \mu'_0$. Toward a contradiction, suppose $\mu_{D^*}(\mu_0) \neq \mu_{D^*}(\mu'_0)$; by part 3, it follows that $\mu_{D^*}(\mu_0) < \mu_{D^*}(\mu'_0)$. The payoff for the seller when his prior is μ'_0 from following his policy is

$$\frac{\mu'_0 - \mu_{D^*}(\mu'_0)}{1 - \mu_{D^*}(\mu'_0)} v_H + \delta \frac{1 - \mu'_0}{1 - \mu_{D^*}(\mu'_0)} \left[\frac{\mu_{D^*}(\mu'_0) - \bar{\mu}_{n-1}}{1 - \bar{\mu}_{n-1}} v_H + \delta \frac{1 - \mu_{D^*}(\mu'_0)}{1 - \bar{\mu}_{n-1}} R^{\mu_{D^*}}(\bar{\mu}_{n-1}, \mu'_0) \right],$$

where $\bar{\mu}_{n-1} = \mu_{D^*}(\mu_{D^*}(\mu'_0))$. By the inductive hypothesis, $\mu_{D^*}(\mu_{D^*}(\mu'_0)) = \mu_{D^*}(\mu_{D^*}(\mu_0)) = \bar{\mu}_{n-1}$.

Instead, if the seller follows μ_0 's policy when his posterior is μ'_0 , his payoff would be

$$\frac{\mu'_0 - \mu_{D^*}(\mu_0)}{1 - \mu_{D^*}(\mu_0)} v_H + \delta \frac{1 - \mu'_0}{1 - \mu_{D^*}(\mu_0)} \left[\frac{\mu_{D^*}(\mu_0) - \bar{\mu}_{n-1}}{1 - \bar{\mu}_{n-1}} v_H + \delta \frac{1 - \mu_{D^*}(\mu_0)}{1 - \bar{\mu}_{n-1}} R^{\mu_{D^*}}(\bar{\mu}_{n-1}, \mu'_0) \right],$$

where we use the inductive statement that $\mu_{D^*}(\mu_0), \mu_{D^*}(\mu'_0) \in D_{n-1}$ and hence $\mu_{D^*}(\mu_{D^*}(\mu_0)) = \mu_{D^*}(\mu_{D^*}(\mu'_0)) = \bar{\mu}_{n-1}$.

Taking differences, we have:

$$\begin{aligned} & v_H \left[\frac{\mu'_0 - \mu_{D^*}(\mu'_0)}{1 - \mu_{D^*}(\mu'_0)} + \delta \frac{1 - \mu'_0}{1 - \mu_{D^*}(\mu'_0)} \frac{\mu_{D^*}(\mu'_0) - \bar{\mu}_{n-1}}{1 - \bar{\mu}_{n-1}} - \frac{\mu'_0 - \mu_{D^*}(\mu_0)}{1 - \mu_{D^*}(\mu_0)} - \delta \frac{1 - \mu'_0}{1 - \mu_{D^*}(\mu_0)} \frac{\mu_{D^*}(\mu_0) - \bar{\mu}_{n-1}}{1 - \bar{\mu}_{n-1}} \right] \\ &= \frac{v_H(1 - \mu'_0)}{(1 - \mu_{D^*}(\mu_0))(1 - \mu_{D^*}(\mu'_0))} ((\delta - 1)(\mu_{D^*}(\mu'_0) - \mu_{D^*}(\mu_0))) < 0, \end{aligned}$$

which contradicts the optimality of μ'_0 's policy. Thus, $\mu_{D^*}(\mu_0) = \mu_{D^*}(\mu'_0)$ (by part 3 we cannot have $\mu_{D^*}(\mu'_0) < \mu_{D^*}(\mu_0)$). Hence, $\mathbf{P}(\mathbf{n} + \mathbf{1}) = 1$.

C.3 Proof of Theorem C.1

Theorem C.1. *The policy described in Theorem 2 defines an intrapersonal equilibrium.*

Proof of Theorem C.1. As in the proof of Proposition C.2, define

$$D_n = \{\mu_0 : \mu_{D^*}^{(n)}(\mu_0) = 0\},$$

to be the set of seller priors such that trade with v_L happens in n periods.

The proof of Theorem C.1 is structured as follows. First, Lemma C.3 shows the cutoffs $\{\bar{\mu}_n\}_{n \geq 0}$ form a strictly increasing sequence. Second, we construct the payoffs from (τ^*, q^*) . Finally, we show $\langle (\tau^*, q^*), R^{(\tau^*, q^*)} \rangle$ is an intrapersonal equilibrium.

Lemma C.3. *The sequence $\{\bar{\mu}_n\}_{n \geq 0}$ is strictly increasing in n for $n \geq 1$.*

Proof. Recall that we are defining $\bar{\mu}_0 = 0$ and $\bar{\mu}_1 = v_L/v_H$. Now let $n \geq 1$ and let $\mu_n \in D_n$ and $\mu_{n-1} \in D_{n-1}$. Fix a prior, $\mu_0 \geq \mu_n$. We claim that the difference

$$\Delta_n(\mu_0; \mu_n, \mu_{n-1}) = \left[\frac{\mu_0 - \mu_n}{1 - \mu_n} v_H + \frac{1 - \mu_0}{1 - \mu_n} \delta R^{\mu_{D^*}}(\mu_n, \mu_0) \right] - \left[\frac{\mu_0 - \mu_{n-1}}{1 - \mu_{n-1}} v_H + \frac{1 - \mu_0}{1 - \mu_{n-1}} \delta R^{\mu_{D^*}}(\mu_{n-1}, \mu_0) \right],$$

is increasing in μ_0 for any monotone simple policy μ_{D^*} .

The proof follows similar steps to those in the proof of Lemma C.2. The arguments in that proof imply Δ_n is differentiable in μ_0 . Then,

$$\begin{aligned} \frac{\partial}{\partial \mu_0} \Delta_n(\mu_0; \mu_n, \mu_{n-1}) &= \\ &= \frac{v_H(\mu_n - \mu_{n-1})}{(1 - \mu_n)(1 - \mu_{n-1})} - \delta \left[\frac{R^{\mu_{D^*}}(\mu_n, \mu_0)}{1 - \mu_n} - \frac{R^{\mu_{D^*}}(\mu_{n-1}, \mu_0)}{1 - \mu_{n-1}} \right] + \delta(1 - \mu_0)[\delta^{n-1} - \delta^n] \frac{\Delta v}{(1 - \mu_0)^2} \\ &= \frac{v_H(\mu_n - \mu_{n-1})}{(1 - \mu_n)(1 - \mu_{n-1})} - \delta \left[\frac{v_H \alpha(\mu_n)}{1 - \mu_n} - \frac{v_H \alpha(\mu_{n-1})}{1 - \mu_{n-1}} \right] + \delta \hat{v}_L(\mu_0)(\delta^{n-1} - \delta^n) + \delta[\delta^{n-1} - \delta^n] \frac{\Delta v}{(1 - \mu_0)} \\ &= \frac{v_H(\mu_n - \mu_{n-1})}{(1 - \mu_n)(1 - \mu_{n-1})} - \delta v_H \left[\frac{\mu_n - \mu_{n-1}}{(1 - \mu_n)(1 - \mu_{n-1})} \right] + \delta(\delta^{n-1} - \delta^n) v_H > 0, \end{aligned}$$

where the last equality follows from using the definition of $\alpha(\cdot)$.

Now recall [Equation 14](#), which we reproduce below for easy reference:

$$\frac{\bar{\mu}_{n+1} - \bar{\mu}_n}{1 - \bar{\mu}_n} v_H + \frac{1 - \bar{\mu}_{n+1}}{1 - \bar{\mu}_n} \delta R^{(\tau^*, q^*)}(\bar{\mu}_n, \bar{\mu}_{n+1}) = \frac{\bar{\mu}_{n+1} - \bar{\mu}_{n-1}}{1 - \bar{\mu}_{n-1}} v_H + \frac{1 - \bar{\mu}_{n+1}}{1 - \bar{\mu}_{n-1}} \delta R^{(\tau^*, q^*)}(\bar{\mu}_{n-1}, \bar{\mu}_{n+1}).$$

Note that the cutoff $\bar{\mu}_{n+1}$ is defined by $\Delta_n(\bar{\mu}_{n+1}; \bar{\mu}_n, \bar{\mu}_{n-1}) = 0$.

Note that $\bar{\mu}_{n+1} \neq \bar{\mu}_n$ if $n \geq 1$. If $\bar{\mu}_n = \bar{\mu}_{n+1}$, then

$$0 = \Delta_n(\bar{\mu}_{n+1}; \bar{\mu}_n, \bar{\mu}_{n-1}) = \delta R^{(\tau^*, q^*)}(\bar{\mu}_n, \bar{\mu}_n) - R^{(\tau^*, q^*)}(\bar{\mu}_n, \bar{\mu}_n) \leq 0,$$

a contradiction, unless $\delta = 1$. Because $\Delta_n(\cdot; \bar{\mu}_n, \bar{\mu}_{n-1})$ is increasing, $\bar{\mu}_{n+1} > \bar{\mu}_n$. This completes the proof. $\square$

An implication of [Lemma C.3](#) is that $\bar{\mu}_{n+1} = \inf\{\mu_0 : \mu_0 \in D_{n+1}\} = \sup\{\mu_0 : \mu_0 \in D_n\}$. The equilibrium we construct makes it so that $\bar{\mu}_{n+1} = \min\{\mu_0 : \mu_0 \in D_{n+1}\}$.

Given (τ^*, q^*) as in the statement of [Theorem 2](#), we construct the continuation payoffs, $R^{(\tau^*, q^*)}$, for the seller for each of his priors, μ_0 , and posteriors he may induce, μ' . If $\mu' < \bar{\mu}_1$,

$$R^{(\tau^*, q^*)}(\mu', \mu_0) = \mu' v_H + (1 - \mu') \hat{v}_L(\mu_0), \quad (\text{C.8})$$

whereas if $\mu' \in [\bar{\mu}_n, \bar{\mu}_{n+1})$ for $n \geq 1$, we have

$$R^{(\tau^*, q^*)}(\mu', \mu_0) = \frac{\mu' - \bar{\mu}_{n-1}}{1 - \bar{\mu}_{n-1}} v_H + \delta \frac{1 - \mu'}{1 - \bar{\mu}_{n-1}} R^{(\tau^*, q^*)}(\bar{\mu}_{n-1}, \mu_0) \quad (\text{C.9})$$

$$R^{(\tau^*, q^*)}(\bar{\mu}_n, \mu_0) = v_H \left[\frac{\bar{\mu}_n - \bar{\mu}_{n-1}}{1 - \bar{\mu}_{n-1}} + (1 - \bar{\mu}_n) \sum_{i=1}^{n-1} \frac{\delta^{n-i}}{1 - \bar{\mu}_i} \frac{\bar{\mu}_i - \bar{\mu}_{i-1}}{1 - \bar{\mu}_{i-1}} \right] + \delta^n (1 - \bar{\mu}_n) \hat{v}_L(\mu_0)$$

We now verify that for each $\mu_0 \geq \bar{\mu}_1$, $\langle \tau^*(\mu_0, \cdot) q^*(\mu_0, \cdot) \rangle$, is optimal given $R^{(\tau^*, q^*)}(\cdot, \mu_0)$. Recall that given $R^{(\tau^*, q^*)}(\cdot, \mu_0)$, the seller solves:

$$\max_{\tau, q} \int_0^1 [q(\mu_0, \mu') (\mu' v_H + (1 - \mu') \hat{v}_L(\mu_0)) + \delta (1 - q(\mu_0, \mu')) R^{(\tau^*, q^*)}(\mu', \mu_0)] d\tau(\mu_0, \mu') \quad (\text{C.10})$$

subject to the constraints that $q(\mu_0, \mu') \in [0, 1]$ and $\int_0^1 \mu' d\tau(\mu_0, \mu') = \mu_0$.

Consider first the case in which $\mu_0 = \bar{\mu}_1$. Note that $\mu'v_H + (1 - \mu')\hat{v}_L(\bar{\mu}_1) = \mu'v_H \geq \delta R^{(\tau^*, q^*)}(\mu', \bar{\mu}_1) = \delta \alpha(\mu')v_H$ strictly so when $\mu' > 0$. Thus, $q^*(\bar{\mu}_1, \mu') = 1$, except possibly at $\mu = 0$. Moreover, the payoff of setting $q^*(\bar{\mu}_1, 0) = 0$ and $q^*(\bar{\mu}_1, 1) = 1$ (and splitting beliefs in this way) equals $\bar{\mu}_1 v_H$ and is exactly the same as the payoff of setting $q^*(\bar{\mu}_1, \cdot) = 1$ everywhere. Thus, $\tau^*(\bar{\mu}_1, 0) = 1 - \bar{\mu}_1 = 1 - \tau^*(\bar{\mu}_1, 1)$ is optimal.

Consider next $\mu_0 \in D_n$ for $n \geq 1$, if $n = 1$, then $\mu_0 > \bar{\mu}_1$. Part 3 of [Proposition C.1](#) implies that the seller never places positive probability on $(\mu_0, 1)$. Moreover, [Proposition C.2](#) implies that $\tau(\mu_0, \cdot)$ has finite support since it cannot assign positive probability to two posteriors $\mu'_1, \mu'_2 \in D_m$ for some $m \leq n$. It also implies that if $\mu' \in D_m$ and $\tau(\mu_0, \mu') > 0$, then $\mu' = \bar{\mu}_m$. Finally, it must be that $\tau(\mu_0, \bar{\mu}_n) = 0$ in a best response. By definition, if $\mu_0 \in D_n$, then $\Delta_n(\mu_0; \bar{\mu}_n, \bar{\mu}_{n-1}) < 0$.

We now show that

$$\frac{\mu_0 - \bar{\mu}_{n-1}}{1 - \bar{\mu}_{n-1}}v_H + \frac{1 - \mu_0}{1 - \bar{\mu}_{n-1}}\delta R^{(\tau^*, q^*)}(\bar{\mu}_{n-1}, \mu_0) \geq \tau(\mu_0, 1)v_H + \sum_{m=0}^{n-1} \tau(\mu_0, \bar{\mu}_m)\delta R^{(\tau^*, q^*)}(\bar{\mu}_m, \mu_0) \quad (\text{C.11})$$

for any $\tau(\mu_0, \cdot) \in \Delta(V)$ such that

$$\begin{aligned} \text{supp } \tau(\mu_0, \cdot) &\subset \{\bar{\mu}_0, \dots, \bar{\mu}_{n-1}, 1\} \\ \tau(\mu_0, 1) + \sum_{m=0}^{n-1} \tau(\mu_0, \bar{\mu}_m)\bar{\mu}_m &= \mu_0. \end{aligned}$$

Using the properties of $\tau(\mu_0, \cdot)$, one can write the RHS of [Equation C.11](#) as follows:

$$\begin{aligned} v_H &\left[\frac{\mu_0 - \bar{\mu}_{n-1} + \sum_{m=0}^{n-2} \tau(\mu_0, \bar{\mu}_m)(\bar{\mu}_{n-1} - \bar{\mu}_m)}{1 - \bar{\mu}_{n-1}} \right] + \delta \frac{1 - \mu_0}{1 - \bar{\mu}_{n-1}} R^{(\tau^*, q^*)}(\bar{\mu}_{n-1}, \mu_0) \\ &+ \sum_{m=0}^{n-2} \tau(\mu_0, \bar{\mu}_m)(1 - \bar{\mu}_m) \left[\frac{\delta R^{(\tau^*, q^*)}(\bar{\mu}_m, \mu_0)}{1 - \bar{\mu}_m} - \frac{\delta R^{(\tau^*, q^*)}(\bar{\mu}_{n-1}, \mu_0)}{1 - \bar{\mu}_{n-1}} \right] \end{aligned}$$

Thus, to show that [Equation C.11](#) holds, we need to show that

$$\sum_{m=0}^{n-2} \tau(\mu_0, \bar{\mu}_m)(1 - \bar{\mu}_m) \left[\frac{\delta R^{(\tau^*, q^*)}(\bar{\mu}_{n-1}, \mu_0)}{1 - \bar{\mu}_{n-1}} - \frac{\delta R^{(\tau^*, q^*)}(\bar{\mu}_m, \mu_0)}{1 - \bar{\mu}_m} - \frac{v_H(\bar{\mu}_{n-1} - \bar{\mu}_m)}{(1 - \bar{\mu}_{n-1})(1 - \bar{\mu}_m)} \right] \geq 0 \quad (\text{C.12})$$

That Equation C.12 holds follows from the following observation. For all $m \leq n - 1$:

$$\begin{aligned}
& \frac{\bar{\mu}_m - \bar{\mu}_{m-1}}{1 - \bar{\mu}_m} v_H + \frac{1 - \bar{\mu}_m}{1 - \bar{\mu}_{m-1}} \delta R^{(\tau^*, q^*)}(\bar{\mu}_{m-1}, \mu_0) = \\
& = \frac{\bar{\mu}_m - \bar{\mu}_{m-1}}{1 - \bar{\mu}_{m-1}} v_H + \frac{1 - \bar{\mu}_m}{1 - \bar{\mu}_{m-1}} \delta R^{(\tau^*, q^*)}(\bar{\mu}_{m-1}, \bar{\mu}_m) + \delta^m (\hat{v}_L(\mu_0) - \hat{v}_L(m)) \\
& = \frac{\bar{\mu}_m - \bar{\mu}_{m-2}}{1 - \bar{\mu}_{m-2}} v_H + \frac{1 - \bar{\mu}_m}{1 - \bar{\mu}_{m-2}} \delta R^{(\tau^*, q^*)}(\bar{\mu}_{m-2}, \bar{\mu}_m) + \delta^m (\hat{v}_L(\mu_0) - \hat{v}_L(m)) \\
& = \frac{\bar{\mu}_m - \bar{\mu}_{m-2}}{1 - \bar{\mu}_{m-2}} v_H + \frac{1 - \bar{\mu}_m}{1 - \bar{\mu}_{m-2}} \delta R^{(\tau^*, q^*)}(\bar{\mu}_{m-2}, \mu_0) + \delta^{m-1} (\hat{v}_L(m) - \hat{v}_L(\mu_0)) + \delta^m (\hat{v}_L(\mu_0) - \hat{v}_L(m)) \\
& = \frac{\bar{\mu}_m - \bar{\mu}_{m-2}}{1 - \bar{\mu}_{m-2}} v_H + \frac{1 - \bar{\mu}_m}{1 - \bar{\mu}_{m-2}} \delta R^{(\tau^*, q^*)}(\bar{\mu}_{m-2}, \mu_0) + \delta^{m-1} (\hat{v}_L(m) - \hat{v}_L(\mu_0)) (1 - \delta),
\end{aligned}$$

where the third line uses the indifference condition that defines $\bar{\mu}_m$. This is if, and only if, for all $m \leq n - 1$

$$\frac{\delta R^{(\tau^*, q^*)}(\bar{\mu}_{m-1}, \mu_0)}{1 - \bar{\mu}_{m-1}} - \frac{\delta R^{(\tau^*, q^*)}(\bar{\mu}_{m-2}, \mu_0)}{1 - \bar{\mu}_{m-2}} \geq v_H \frac{(\bar{\mu}_{m-1} - \bar{\mu}_{m-2})}{(1 - \bar{\mu}_{m-2})(1 - \bar{\mu}_{m-1})}. \quad (\text{C.13})$$

Note that the inequality is strict whenever $m < n - 1$. Successive application of Equation C.13 implies that Equation C.12 holds. Indeed, this can be verified by writing the term inside the summation in Equation C.12 as:

$$\frac{\delta R^{(\tau^*, q^*)}(\bar{\mu}_{n-1}, \mu_0)}{1 - \bar{\mu}_{n-1}} - \frac{\delta R^{(\tau^*, q^*)}(\bar{\mu}_m, \mu_0)}{1 - \bar{\mu}_m} = \sum_{l=m}^{n-2} \left(\frac{\delta R^{(\tau^*, q^*)}(\bar{\mu}_{l+1}, \mu_0)}{1 - \bar{\mu}_{l+1}} - \frac{\delta R^{(\tau^*, q^*)}(\bar{\mu}_l, \mu_0)}{1 - \bar{\mu}_l} \right),$$

noting that Equation C.13 implies that the RHS of the above expression is bounded below by

$$\sum_{l=m}^{n-2} \frac{v_H (\bar{\mu}_{l+1} - \bar{\mu}_l)}{(1 - \bar{\mu}_{l+1})(1 - \bar{\mu}_l)} = v_H \frac{\bar{\mu}_{n-1} - \bar{\mu}_m}{(1 - \bar{\mu}_m)(1 - \bar{\mu}_{n-1})}.$$

This completes the proof. $\square$

D Collected objects from intrapersonal equilibrium

In this section, we construct the following objects using the ingredients from the intrapersonal equilibrium:

1. The mapping $\gamma^* : \Delta(V) \mapsto \mathcal{M}_C$ that associates to each belief the seller may hold, μ_0 , a mechanism $\gamma^*(\mu_0) \in \mathcal{M}_C$,
2. The mapping $u_H^* : \Delta(V) \mapsto \mathbb{R}$ that associates to each belief the seller may hold, μ_0 , the buyer's payoff when her valuation is v_H and the seller employs the mechanism derived from the intrapersonal equilibrium,
3. The correspondence $\mathcal{U}_H^* : \Delta(V) \rightrightarrows \mathbb{R}$ that associates to each belief the seller may hold, μ_0 , the set of feasible buyer's payoffs when her valuation is v_H , the seller employs the mechanism derived from the intrapersonal equilibrium, but does not necessarily break ties as in the intrapersonal equilibrium. That is, $\mathcal{U}_H^*(\mu_0) = u_H^*(\mu_0)$ whenever $\mu_0 \neq \bar{\mu}_i, i = 0, 1, \dots$.

D.1 Mechanism and buyer's payoffs in the intrapersonal equilibrium

In this section, we construct the transfers and the continuation rents for the buyer, when her type is v_H , implied by the intrapersonal equilibrium. (The communication device and the probability of trade are the ones specified in [Section 3.3](#).) We proceed “backwards,” starting from $\mu_0 < \bar{\mu}_1$ and then $\mu_0 \in D_n$ for $n \geq 1$.

Let $\mu_0 < \bar{\mu}_1$. Using the binding participation constraint for the buyer when her type is v_L , we have

$$x_{\mu_0}^*(\mu_0) = v_L.$$

Note that the utility of the buyer of type v_H when the seller has belief $\mu_0 < \bar{\mu}_1$ is $u_H^*(\mu_0) = \Delta v$.

Fix $n \geq 1$. Consider now $\mu_0 \in [\bar{\mu}_n, \bar{\mu}_{n+1})$. Using the binding participation constraint for the buyer when her valuation is v_L , we have

$$x_{\mu_0}^*(\bar{\mu}_{n-1}) = 0.$$

To determine $x_{\mu_0}^*(1)$, we use the incentive compatibility constraint for the high type:

$$\begin{aligned} \beta_{\mu_0}^*(1|v_H)(v_H - x_{\mu_0}^*(1)) + \beta_{\mu_0}^*(\bar{\mu}_{n-1}|v_H)\delta u_H^*(\bar{\mu}_{n-1}) &= \delta u_H^*(\bar{\mu}_{n-1}) \\ x_{\mu_0}^*(1) &= v_H - \delta u_H^*(\bar{\mu}_{n-1}). \end{aligned}$$

We then construct $x_{\mu_0}^*(1)$ recursively. Set $n = 1$. Then, for $\mu \in [\bar{\mu}_1, \bar{\mu}_2)$,

$$\begin{aligned} x_{\mu_0}^*(1) &= v_H - \delta \Delta v = v_L + (1 - \delta) \Delta v \\ u_H^*(\mu_0) &= \delta \Delta v, \end{aligned}$$

and for $n \geq 2$, we obtain that for $\mu_0 \in [\bar{\mu}_n, \bar{\mu}_{n+1})$,

$$\begin{aligned} x_{\mu_0}^*(1) &= v_H - \delta^n \Delta v = v_L + (1 - \delta^n) \Delta v \\ u_H^*(\mu_0) &= \delta^n \Delta v. \end{aligned}$$

Note the construction of the transfers highlights that if the seller posts a price of $x_{\mu_0}^*(1)$ when his belief is μ_0 , the buyer, when her type is v_H , is indifferent between accepting this price and rejecting it; whereas the low type is always better off rejecting this price because it lies above v_L .

Summing up, the intrapersonal equilibrium determines a map $\gamma^* : \Delta(V) \mapsto \mathcal{M}_C$ such that

$$\gamma^*(\mu_0) = \langle (V, \beta_{\mu_0}^*, \Delta(V)), (q_{\mu_0}^*, x_{\mu_0}^*) \rangle. \quad (\text{D.1})$$

The buyer's payoffs when her valuation is v_H are given by:

$$u_H^*(\mu_0) = \delta^i \Delta v, \mu_0 \in D_i, i = 0, \dots$$

Let $\mathcal{U}_H^*$ denote the following correspondence:

$$\mathcal{U}_H^*(\mu_0) = \begin{cases} u_H^*(\mu_0) & \text{if } \mu_0 \neq \bar{\mu}_i, i \geq 1 \\ [\delta^i \Delta v, \delta^{i-1} \Delta v] & \text{if } \mu_0 = \bar{\mu}_i \end{cases}. \quad (\text{D.2})$$

For future use, note $\mathcal{U}_H^*$ is upper-hemicontinuous, convex-valued, and compact-valued.

E Proofs of Section 3.3

Appendix E is organized in three parts. Section E.1 constructs the buyer's strategy profile so as to specify her best responses after every history in the game. Section E.2 performs the same exercise for the seller, using the strategy from the intrapersonal equilibrium (recall Equation D.1). Finally, in Section E.3, we show that if continuation values are specified as in the intrapersonal equilibrium, and

the buyer's strategy is completed as in the first step, the seller has no one-shot deviations from the equilibrium strategy. We also argue the buyer has no one-shot deviations. [Appendix I](#) in [Doval and Skreta \(2019\)](#) lays out the formalisms that justify the use of self-generation arguments in our setting; this, in turn, implies that guaranteeing the absence of one-shot deviations given the continuation values is enough to conclude we have indeed constructed a PBE of $G^\infty(\mu_0)$, which gives the seller the same payoff as in the intrapersonal equilibrium.

E.1 Completing the buyer's strategy

Fix a public history h^t and let μ_0 denote the seller's beliefs at that public history.⁴³ To complete the buyer's strategy, we classify mechanisms, $\mathbf{M}$, in four categories:⁴⁴

0. Mechanisms that given the intrapersonal equilibrium continuation values satisfy participation and truthtelling, that is, $\mathbf{M}$ such that the following hold:

$$\begin{aligned} & \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu' | v_H) [v_H q^{\mathbf{M}}(\mu') - x^{\mathbf{M}}(\mu') + \delta(1 - q^{\mathbf{M}}(\mu') u_H^*(\mu'))] \\ & \geq \max\{0, \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu' | v_L) [v_H q^{\mathbf{M}}(\mu') - x^{\mathbf{M}}(\mu') + \delta(1 - q^{\mathbf{M}}(\mu') u_H^*(\mu'))]\}, \end{aligned}$$

and

$$\sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu' | v_L) [v_L q^{\mathbf{M}}(\mu') - x^{\mathbf{M}}(\mu')] \geq \max\{0, \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu' | v_H) [v_L q^{\mathbf{M}}(\mu') - x^{\mathbf{M}}(\mu')]\}.$$

Denote the set of these mechanisms $\mathcal{M}_C^0$.

1. Mechanisms not in $\mathcal{M}_C^0$ such that the buyer, when her type is v_H , has no reporting strategy, $\rho \in \Delta(V)$, such that $\sum_{v'} \beta^{\mathbf{M}}(\mu' | v') \rho(v') [v_H q^{\mathbf{M}}(\mu') - x^{\mathbf{M}}(\mu')] \geq 0$; let $\mathcal{M}_C^1$ denote the set of these mechanisms.
2. Mechanisms not in $\mathcal{M}_C^0$ such that the buyer, when her type is v_H , has a reporting strategy, $\rho \in \Delta(V)$, such that $\sum_{v'} \beta^{\mathbf{M}}(\mu' | v') \rho(v') [v_H q^{\mathbf{M}}(\mu') - x^{\mathbf{M}}(\mu')] \geq 0$, but not when her type is v_L ; let $\mathcal{M}_C^2$ denote the set of these mechanisms.
3. Mechanisms not in $\mathcal{M}_C^0$ such that both types have such a reporting strategy; let $\mathcal{M}_C^3$ denote the set of these mechanisms.

⁴³This prior will, of course, be an equilibrium object, but we suppress this from the notation to keep things simple.

⁴⁴[Gerardi and Maestri \(2018\)](#) use a similar trick to complete the worker's strategy in their paper.

If $\mathbf{M} \in \mathcal{M}_C^1$, specify that the buyer rejects the mechanism for both types. Hence, under this strategy, the seller does not update his beliefs after observing a rejection. If, however, the buyer accepts, the seller believes $v = v_H$. Note that, in this case, continuation payoffs for the buyer are 0 from then on, regardless of her type. For each type v , let $r_v^*(\mathbf{M}, 1)$ denote an element of⁴⁵

$$\arg \max_{\rho \in \Delta(V)} \sum_{v' \in V} \rho(v') \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu' | v') (v q^{\mathbf{M}}(\mu') - x^{\mathbf{M}}(\mu')).$$

If $\mathbf{M} \in \mathcal{M}_C^2$, specify that the buyer rejects when her type is v_L . Hence, without loss of generality, we can specify that if the seller observes that the buyer accepts the mechanism, the buyer's type is v_H . For type v_H , let $r_{v_H}^*(\mathbf{M}, 1)$ satisfy that if $r_{v_H}^*(\mathbf{M}, 1)(v) > 0$, then

$$\sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu' | v) [v_H q^{\mathbf{M}}(\mu') - x^{\mathbf{M}}(\mu')] \geq \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu' | v') [v_H q^{\mathbf{M}}(\mu') - x^{\mathbf{M}}(\mu')]$$

for all $v' \in V$. Note we are using that the buyer's continuation payoffs are 0 conditional on her accepting the mechanism.

Then, v_H 's payoff from participating in the mechanism $\mathbf{M}$ is given by

$$U_{1,v_H}(\mathbf{M}) = \sum_{v' \in V} r_{v_H}^*(\mathbf{M}, 1)(v') \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu' | v') [v_H q^{\mathbf{M}}(\mu') - x^{\mathbf{M}}(\mu') + \delta(1 - q^{\mathbf{M}}(\mu')) \times 0],$$

whereas the payoff from rejecting is

$$U_{0,v_H}(\pi_{v_H}, f) = \delta f(v_2(\mu_0, \pi_{v_H})).$$

where

$$v_2(\mu_0, \pi_{v_H}) = \frac{\mu_0(1 - \pi_{v_H})}{\mu_0(1 - \pi_{v_H}) + 1 - \mu_0}.$$

is the seller's belief that the buyer is of type v_H when observing a rejection, according to Bayes' rule, and $f(v_2(\mu_0, \pi_{v_H}))$ is a measurable selection from $\mathcal{U}_H^*(v_2(\mu_0, \pi_{v_H}))$. Note the payoff from rejecting is specified under the assumption that in the continuation, the equilibrium path coincides with that of the intrapersonal equilibrium

⁴⁵Even if the buyer does not participate on the equilibrium path, we still need to guarantee the reporting strategy is sequentially rational.

when beliefs are $\nu_2(\mu_0, \pi_{v_H})$. We use, however, a selection from $\mathcal{U}_H^*$ to ensure that if needed, the seller randomizes between the posted prices when indifferent to help make the buyer's continuation problem well-behaved.

Now, $(\pi_{v_H}^*(\mathbf{M}), f_2^*(\mathbf{M}))$ are chosen so that

$$\pi_{v_H}^* \in \arg \max_{p \in [0,1]} (1-p)U_{0,v_H}(\pi_{v_H}^*, f_2^*(\mathbf{M})) + pU_{1,v_H}(\mathbf{M})$$

The main result in [Simon and Zame \(1990\)](#) implies a solution to the above problem exists, given the properties of $\mathcal{U}_H^*$ and the linearity in p of the objective. Let $\pi_{v_H}^*(\mathbf{M})$ denote this fixed point. Now, if $\pi_{v_H}^*(\mathbf{M}) < 1$ and $\nu_2(\mu_0, \pi_{v_H}^*(\mathbf{M})) = \bar{\mu}_i$ for some $i \geq 1$, then there exists a weight $\phi_2(\bar{\mu}_i, \mathbf{M}) \in [0, 1]$ that solves the following:

$$f_2^*(\mathbf{M}, \bar{\mu}_i) = (1 - \phi)\delta^{i-1}\Delta v + \phi\delta^i\Delta v. \quad (\text{E.1})$$

This weight captures the probability with which the seller, when his prior is $\bar{\mu}_i$, mixes between $\gamma^*(\bar{\mu}_i)$ and $\gamma^*(\bar{\mu}_{i-1})$.

For a mechanism in $\mathcal{M}_C^2$, let $r_{v_L}^*(\mathbf{M}, 1)$ denote an element in

$$\arg \max_{\rho \in \Delta(V)} \sum_{v' \in V} \rho(v') \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu'|v') (v_L q^{\mathbf{M}}(\mu') - x^{\mathbf{M}}(\mu')).$$

Finally, if $\mathbf{M} \in \mathcal{M}_C^3$, specify that the buyer accepts for both types. If the seller observes that the buyer rejects the mechanism, he assigns probability 1 to the buyer's valuation being v_H . Thus, upon rejection, continuation payoffs are 0 for both types. Note that mechanisms in $\mathcal{M}_C^3$ satisfy the participation constraint given the continuation values of the intrapersonal equilibrium. Now, let m_L^* satisfy⁴⁶

$$\sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu'|m_L^*) (v_L q^{\mathbf{M}}(\mu') - x^{\mathbf{M}}(\mu')) \geq \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu'|v) (v_L q^{\mathbf{M}}(\mu') - x^{\mathbf{M}}(\mu')),$$

for all $v \in V$. Set $r_{v_L}^*(\mathbf{M}, 1)(m_L^*) = 1$. Let $\{m_H^*\} = V \setminus \{m_L^*\}$ and define

$$\begin{aligned} \Delta(V)^H &= \{\mu' : \beta^{\mathbf{M}}(\mu'|m_H^*) > 0\} \\ \Delta(V)^L &= \{\mu' : \beta^{\mathbf{M}}(\mu'|m_L^*) > 0\}. \end{aligned}$$

⁴⁶We cannot ensure that truthtelling will hold for mechanisms in $\mathcal{M}_C^3$, which is why we need the extra piece of notation.

Let μ_0 denote the seller's belief about the buyer's valuation being v_H . Let $r \in [0, 1]$ denote the weight the buyer assigns to m_L^* when her valuation is v_H :

$$v_3(\mu_0, \mu', r, \delta_{m_L^*}) = \begin{cases} 1 & \text{if } \mu' \in \Delta(V)^H \setminus \Delta(V)^L \\ \frac{\mu_0(r\beta^{\mathbf{M}}(\mu'|m_L^*) + (1-r)\beta^{\mathbf{M}}(\mu'|m_H^*))}{\mu_0(r\beta^{\mathbf{M}}(\mu'|m_L^*) + (1-r)\beta^{\mathbf{M}}(\mu'|m_H^*)) + (1-\mu_0)\beta^{\mathbf{M}}(\mu'|m_L^*)} & \text{if } \mu' \in \Delta(V)^H \cap \Delta(V)^L \\ \frac{\mu_0 r}{\mu_0 r + (1-\mu_0)} & \text{if } \mu' \in \Delta(V)^L \setminus \Delta(V)^H \end{cases},$$

where we assume that if the seller observes an output message that is not consistent with m_L^* , he believes it was generated by the high-valuation buyer. This specification of beliefs does not conflict with Bayes' rule where possible: either m_H^* has positive probability in the optimal reporting strategy of the buyer when her valuation is v_H , in which case, v_3 would be consistent with Bayes' rule, or it does not, in which case, Bayes' rule where possible places no restrictions on $v_3(\mu_0, \mu', \cdot)$ for $\mu' \in \Delta(V)^H \setminus \Delta(V)^L$.

Given v_3 , the buyer when her valuation is v_H obtains a payoff of

$$U_{1,v_H}(r, m, f) = \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu'|m) (v_H q^{\mathbf{M}}(\mu') - x^{\mathbf{M}}(\mu') + \delta(1 - q^{\mathbf{M}}(\mu')) f(v_3(\mu_0, \mu', r, \delta_{m_L^*}))),$$

when she reports $m \in V$, where f is a selection from $\mathcal{U}_H^*$. We want to find $(r_{v_H}^*, f_3^*(\mathbf{M}))$ so that

$$r_{v_H}^* \in \arg \max_{r \in [0,1]} r U_{1,v_H}(r_{v_H}^*, m_L^*, f_3^*(\mathbf{M})) + (1-r) U_{1,v_H}(r_{v_H}^*, m_H^*, f_3^*(\mathbf{M})). \quad (\text{E.2})$$

Again, we appeal to the result in [Simon and Zame \(1990\)](#) to argue that the properties of $\mathcal{U}_H^*$ and the objective function above imply the existence of such an $(r^*, f_3^*(\mathbf{M}))$. Set $r_{v_H}^*(\mathbf{M}, 1)(m_L^*) = r_{v_H}^*$.

As we did before, whenever $v_3(\mu_0, \mu', r_{v_H}^*, \delta_{m_L^*}) = \bar{\mu}_i$ for $i \geq 1$, we can define $\phi_3(\mathbf{M}, \bar{\mu}_i)$ as the weight on $\gamma^*(\bar{\mu}_i)$ implied by $f_3^*(\mathbf{M}, \bar{\mu}_i)$.

Summing up, the buyer's strategy at any history h_B^t is given by

$$\pi_v^*(h_B^t, \mathbf{M}) = \begin{cases} 1 & \text{if } \mathbf{M} \in \mathcal{M}_C^0 \cup \mathcal{M}_C^3 \\ \pi_{v_H}^*(\mathbf{M}) & \text{if } \mathbf{M} \in \mathcal{M}_C^2 \text{ and } v = v_H \\ 0 & \text{otherwise} \end{cases} \quad (\text{E.3})$$

and

$$r_v^*(h_B^t, \mathbf{M}, 1) = \begin{cases} \delta_v & \text{if } \mathbf{M} \in \mathcal{M}_C^0 \\ r_v^*(\mathbf{M}, 1) & \text{otherwise} \end{cases}. \quad (\text{E.4})$$

E.2 Full specification of the PBE assessment

To complete the PBE assessment for $G^\infty(\mu_0)$, we now specify the seller's strategy and the belief system. We introduce two pieces of notation: the first one allows us to keep track of the last payoff-relevant event; the second one allows us to keep track of how the seller's beliefs evolve given the buyer's strategy.

From the beginning of period t of the game until the end of that period, the following things are determined: (i) the seller's choice of mechanism, $\mathbf{M}$; (ii) the buyer's participation decision, $p \in (0, 1)$; and (iii) if the buyer participates of the mechanism, the allocation and the output message. We let $z_\emptyset(\mathbf{M}) = (\mathbf{M}, 0, \emptyset, 0, 0)$ to denote the outcome when mechanism $\mathbf{M}$ is chosen, the buyer does not participate, no output message is produced, and no trade and no transfers occur. We let $z_q(\mathbf{M}, \mu') = (\mathbf{M}, 1, \mu', q, x)$ denote the outcome at the end of period t when mechanism $\mathbf{M}$ is chosen, the buyer participates (1), the output message is μ' , and the allocation is (q, x) . Of course, if $q = 1$, the game ends. Note that any public history h^t can be written as (h^{t-1}, z) for some z as defined above. Given h^t , let $z(h^t)$ denote the corresponding outcome z . Given any prior μ_0 , define

$$T(\mu_0, z) = \begin{cases} \mu' & \text{if } z = z_q(\mathbf{M}, \mu') \text{ and } \mathbf{M} \in \mathcal{M}_C^0 \\ & \text{if } z = z_\emptyset(\mathbf{M}) \text{ and } \mathbf{M} \in \mathcal{M}_C^0 \cup \mathcal{M}_C^3 \\ 1 & \text{or} \\ & z = z_q(\mathbf{M}, \mu') \text{ and } \mathbf{M} \in \mathcal{M}_C^1 \cup \mathcal{M}_C^2 \\ \mu_0 & \text{if } z = z_\emptyset(\mathbf{M}) \text{ and } \mathbf{M} \in \mathcal{M}_C^1 \\ v_2(\mu_0, \pi_{v_H}^*(\mathbf{M})) & \text{if } z = z_q(\mathbf{M}, \mu') \text{ and } \mathbf{M} \in \mathcal{M}_C^2 \\ v_3(\mu_0, \mu', r_{v_H}^*(\mathbf{M}, 1), r_{v_L}^*(\mathbf{M}, 1)) & \text{if } z = z_q(\mathbf{M}, \mu') \text{ and } \mathbf{M} \in \mathcal{M}_C^3 \end{cases} \quad (\text{E.5})$$

Let μ_0 denote the seller's prior in $G^\infty(\mu_0)$. Define $\mu^*(\emptyset) = \mu_0$ and $\Gamma^*(\emptyset) = \gamma^*(\mu_0)$, where γ^* is the strategy in the intrapersonal equilibrium as defined in Equation D.1. For any public history h^t , define

$$\mu^*(h^t) = T(\mu^*(h^{t-1}), z(h^t)). \quad (\text{E.6})$$

If either (i) $z = z_q(\mathbf{M}, \cdot)$ and $\mathbf{M} \in \mathcal{M}_C^0 \cup \mathcal{M}_C^1 \cup \mathcal{M}_C^2$, (ii) $z = z_\emptyset(\mathbf{M})$, $\mu^*(h^t) \notin \{\bar{\mu}_i\}_{i \geq 1}$ and $\mathbf{M} \in \mathcal{M}_C^2$, (iii) $z = z_q(\mathbf{M}, \cdot)$, $\mu^*(h^t) \notin \{\bar{\mu}_i\}_{i \geq 1}$ and $\mathbf{M} \in \mathcal{M}_C^3$, or (iv) $z = z_\emptyset(\mathbf{M})$ and $\mathbf{M} \in \mathcal{M}_C^0 \cup \mathcal{M}_C^1 \cup \mathcal{M}_C^3$, define

$$\Gamma^*(h^t)(\mathbf{M}) = \mathbb{1}[\mathbf{M} = \gamma^*(T(\mu^*(h^{t-1}), z(h^t)))]. \quad (\text{E.7})$$

If $z(h^t) = z_\emptyset(\mathbf{M}, \cdot)$ for $\mathbf{M} \in \mathcal{M}_C^2$ and $\mu^*(h^t) = \bar{\mu}_i, i \geq 1$, define

$$\Gamma^*(h^t)(\mathbf{M}) = \begin{cases} \phi_2(\mathbf{M}, \bar{\mu}_i) & \mathbf{M} = \gamma^*(\bar{\mu}_i) \\ 1 - \phi_2(\mathbf{M}, \bar{\mu}_i) & \mathbf{M} = \gamma^*(\bar{\mu}_{i-1}) \\ 0 & \text{otherwise} \end{cases} . \quad (\text{E.8})$$

Finally, if $z(h^t) = z_q(\mathbf{M}, \cdot)$ for $\mathbf{M} \in \mathcal{M}_C^3$ and $\mu^*(h^t) = \bar{\mu}_i, i \geq 1$, define

$$\Gamma^*(h^t)(\mathbf{M}) = \begin{cases} \phi_3(\mathbf{M}, \bar{\mu}_i) & \mathbf{M} = \gamma^*(\bar{\mu}_i) \\ 1 - \phi_3(\mathbf{M}, \bar{\mu}_i) & \mathbf{M} = \gamma^*(\bar{\mu}_{i-1}) \\ 0 & \text{otherwise} \end{cases} . \quad (\text{E.9})$$

Equations E.3-E.4, together with equations E.6- E.9, define a PBE assessment where the system of beliefs is derived from the strategy profile where possible.

E.3 Seller's and buyer's sequential rationality

We now check that at all histories, neither the seller nor the buyer have a one-shot deviation from the prescribed strategy profile, given the continuation values for the seller and the buyer constructed using the intrapersonal equilibrium. This implies that the payoffs in the intrapersonal equilibrium together with the belief operator defined in Equation E.5 belong in the self-generating set. Appendix I in Doval and Skreta (2019) lays out the formalisms that justify the use of techniques à la Abreu et al. (1990) for the game under consideration.

Let $\mu^*(h^t)$ denote the seller's belief at history h^t that the buyer's type is v_H . We now show the seller cannot achieve a payoff higher than $R^{(\tau^*, q^*)}(\mu^*(h^t), \mu^*(h^t))$. Note that among mechanisms of type 0, the one that corresponds to the intrapersonal equilibrium is, by definition, the best that the seller can do. Hence, to show that there are no deviations, we need to show the seller cannot benefit from offering mechanisms of types 1-3.

Given the buyer's strategy, the payoff from offering a mechanism in $\mathcal{M}_C^1$ is $\delta R^{(\tau^*, q^*)}(\mu^*(h^t), \mu^*(h^t)) \leq R^{(\tau^*, q^*)}(\mu^*(h^t), \mu^*(h^t))$, and hence this deviation is not profitable. Let $\mathbf{M}$ denote a mechanism in $\mathcal{M}_C^2$ and let $\pi_{v_H}^*(h_B^t, \mathbf{M}), r_{v_H}^*(h_B^t, \mathbf{M}, 1)$ denote the buyer's best responses as constructed in Section E.1. Denote by $\nu_2 \equiv \nu_2(\mu^*(h^t), \pi_{v_H}^*(h_B^t, \mathbf{M}))$ the seller's belief that he is facing a buyer with valuation

v_H when he observes non-participation. The seller's payoff is then

$$\begin{aligned} & \mu^*(h^t) \pi_{v_H}^*(h_B^t, \mathbf{M}) \sum_{\mu' \in \Delta(V)} \left(\sum_{v \in V} \beta^{\mathbf{M}}(\mu'|v) r_{v_H}^*(h_B^t, \mathbf{M}, 1)(v) \right) [x^{\mathbf{M}}(\mu') + \delta(1 - q^{\mathbf{M}}(\mu'))v_H] \\ & + (1 - \mu^*(h^t) \pi_{v_H}^*(h_B^t, \mathbf{M})) R^{(\tau^*, q^*)}(v_2, v_2), \end{aligned} \quad (\text{E.10})$$

where the continuation values after rejection are constructed using the policy from the intrapersonal equilibrium once the seller has posterior $v_2(\mu^*(h^t), \pi_{v_H}^*(h_B^t, \mathbf{M}))$.⁴⁷ Now, consider the following alternative mechanism, $\mathbf{M}'$:

$$\begin{aligned} \beta^{\mathbf{M}'}(1|v_H) &= \pi_{v_H}^*(\mathbf{M}) \sum_{v \in V, \mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu'|v) r_{v_H}^*(\mathbf{M})(v) \\ \beta^{\mathbf{M}'}(v_2|v_H) &= (1 - \pi_{v_H}^*(\mathbf{M})) \\ \beta^{\mathbf{M}'}(v_2|v_L) &= 1 \\ q^{\mathbf{M}'}(v_2) &= x^{\mathbf{M}'}(v_2) = 0 \\ x^{\mathbf{M}'}(1) &= \sum_{v \in V} \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu'|v) r_{v_H}^*(\mathbf{M})(v) x^{\mathbf{M}}(\mu') \\ q^{\mathbf{M}'}(1) &= \sum_{v \in V} \sum_{\mu' \in \Delta(V)} \beta^{\mathbf{M}}(\mu'|v) r_{v_H}^*(\mathbf{M})(v) q^{\mathbf{M}}(\mu') \end{aligned}$$

Note that if the buyer participates and truthfully reports her type, $\mathbf{M}'$ gives the seller a payoff equal to the expression in [Equation E.10](#). Moreover, participating and truthfully reporting her type is a best response for the buyer. Then, we can write the payoff to mechanism, $\mathbf{M}'$, as

$$\sum_{\mu' \in \{v_2, 1\}} \tau^{\mathbf{M}'}(\mu^*(h^t), \mu') [x^{\mathbf{M}'}(\mu') + \delta(1 - q^{\mathbf{M}'}(\mu')) R^{(\tau^*, q^*)}(\mu', \mu')],$$

where, as in the main text,

$$\tau^{\mathbf{M}'}(\mu^*(h^t), \mu') = \sum_{v \in V} \mu^*(h^t)(v) \beta^{\mathbf{M}'}(\mu'|v).$$

⁴⁷To keep notation simple, we do not use that the seller in the continuation may randomize in his choice.

Now, truth-telling implies $x^{\mathbf{M}'}(1) \leq v_H q^{\mathbf{M}'}(1) - \delta u_H^*(\mu'(\pi_{v_H}^*(\mathbf{M})))$, so that

$$\begin{aligned} & \sum_{\mu' \in \{v_2, 1\}} \tau^{\mathbf{M}'}(\mu^*(h^t), \mu') [x^{\mathbf{M}'}(\mu') + \delta(1 - q^{\mathbf{M}'}(\mu')) R^{(\tau^*, q^*)}(\mu', \mu')] \\ & \leq \sum_{\mu' \in \{v_2, 1\}} \tau^{\mathbf{M}'}(\mu^*(h^t), \mu') \left\{ q^{\mathbf{M}'}(\mu') (\mu' v_H + (1 - \mu') \hat{v}_L(\mu^*(h^t))) \right. \\ & \quad \left. + \delta(1 - q^{\mathbf{M}'}(\mu')) \left[R^{(\tau^*, q^*)}(\mu', \mu') + (1 - \mu') \left(\frac{\mu'}{1 - \mu'} - \frac{\mu^*(h^t)}{1 - \mu^*(h^t)} \right) u_H^*(\mu') \right] \right\}, \end{aligned}$$

which, by definition, is weakly lower than $R^{(\tau^*, q^*)}(\mu^*(h^t), \mu^*(h^t))$.

Similar steps imply the seller does not have a one-shot deviation to a mechanism of type 3.

Finally, note the construction of π_v^*, r_v^* in [Section E.1](#) (Equations [E.3](#) and [E.4](#)) ensures the buyer is best responding after every history given the continuation values.