

STRATEGIC COMPLEXITY UNDER MANDATED DISCLOSURE

Agata Farina

University of Chicago, Booth School of Business

July 7, 2026

ABSTRACT

Mandated disclosure can be ineffective when firms retain control over how information about their products is presented. This freedom allows firms to strategically disclose information in formats that are costly for consumers to understand. We ask whether stronger market alternatives mitigate or amplify this policy limitation, and how the answer depends on the cost consumers face to process what firms disclose. In an information design setting, we show that better market alternatives and lower information-processing costs need not increase information transparency, and, in some cases, can even increase the use of costly-to-process information. This offers a new rational-consumer explanation for excessive disclosure complexity in competitive markets and a reason why mandated disclosure policies may have heterogeneous effects.

I am extremely grateful to Maher Said for his guidance and invaluable advice on this project. I thank Matthew Mitchell and two anonymous referees for insightful comments and suggestions. I am also grateful to Anna Airoidi, Francesco Fabbri, Jimena Galindo, Sam Kapon, Semih Kara, Mario Leccese, Erik Madsen, Marta Mojoli, Paula Onuchic, Debraj Ray, Lawrence White, and the participants in seminars at New York University for their helpful comments. All errors are my own.

1 Introduction

In many consumer markets, firms are required to disclose information that is relevant to consumers’ decisions. For example, credit card issuers must disclose their pricing terms, digital apps must provide information about their privacy policies and data practices, and prescription-drug advertisers must disclose risk information. However, in practice, firms can retain substantial control over which information is effectively conveyed to consumers by strategically shaping *how* it is presented. While regulations often try to limit such discretion through standardized disclosure formats—such as the Schumer box, requirements that privacy notices be clear and accessible, and rules governing the presentation of risk information in direct-to-consumer drug advertising—it is often infeasible to fully regulate all dimensions of disclosure design. As a result, firms may be able to *de facto* withhold information by making it costly or difficult for consumers to understand.¹

This pattern contributes to the proliferation of obfuscation practices and to the persistence of disclosure formats that are costly to process and can impair information transmission. The limits of disclosure regulation alone in preventing these inefficiencies raise a natural question: can market forces discipline firms’ incentives to obfuscate key information about their products? This paper studies how increased market competition, modeled as a shift in the distribution of consumers’ market alternatives toward stronger outside options, affects firms’ incentives to make information excessively costly to understand.² Throughout the paper, we refer to this strategic choice as the use of information complexity, or as complex disclosure. Importantly, we examine how this effect depends on the information-processing cost consumers face after complex disclosure, and how this cost interacts with our measure of competition to shape firms’ incentives to obfuscate.

Using an information design framework, we show that, even with fully rational consumers, strong market alternatives, which we refer to as outside options, and low information-processing costs, do not necessarily discipline firms toward more transparent disclosure, and in some markets may even have the opposite effect. For some distributions, when consumers’ outside options are concentrated at high values, lowering the cost of understanding complex disclosures

¹The disclosure decisions of mutual funds constitute another example of strategic obfuscation. [DeHaan et al. \(2021\)](#) document how investors’ poor choices are partially due to managers creating unnecessary complexity.

²We model competition in reduced form, through the distribution of outside options, rather than as strategic interaction among competing information providers, as in models of persuasion with multiple senders or competitive disclosure (e.g., [Kamenica and Gentzkow, 2017](#); [Board and Lu, 2018](#)). This isolates how consumers’ available alternatives interact with costly information processing.

leads the firm to use such a strategic tool more often. The same market force that is often expected to discipline firms—the availability of attractive alternatives for consumers—can interact with costly information acquisition in a way that preserves, and in some cases strengthens, firms’ incentives to hide information behind complex disclosure. These results suggest that the distribution of consumers’ outside options is a key driver of firms’ willingness to make some information inaccessible. This generates new insights on how to evaluate the effectiveness of mandated disclosure regulations, and may help explain the observed heterogeneity of such policies.

We formalize these mechanisms in a sender-receiver model of mandated disclosure with costly consumer learning. A firm privately observes the quality of his product—which translates into the consumer’s value, net of price—and commits *ex-ante* to a disclosure policy that, for each realization of product quality, sends either a transparent or a complex message with some probability. We refer to this disclosure strategy as a “*complex disclosure rule*”. Transparent messages fully reveal the firm’s private information, while complex messages are costly to process and pool together all realizations that are not disclosed transparently. In what follows, we equivalently refer to the choice of using a complex message as the disclosure of complex information or the sender’s use of complexity/complex disclosure, and to an increase in the likelihood of such a choice as more complex disclosure. Consumers have heterogeneous outside options, which they privately observe. Each consumer takes the complex disclosure rule as given, observes the message sent by the firm, and chooses whether to buy the product, take the outside option, or, in case of a complex message, whether to pay a fixed information-processing cost to learn the firm’s private information before deciding. As mentioned, we interpret the distribution of outside options as a reduced-form measure of market competition: when outside options are strong, consumers can more easily shift to good market alternatives.

The fixed information-processing cost should be interpreted as a feature of the market environment, rather than as an additional instrument available to the firm. In practice, understanding a complex disclosure may require time, attention, expertise, or even third-party assistance: consumers may need to read a credit-card agreement, interpret a privacy policy, or learn about health risks of a drug before making an informed decision. We capture these frictions with a common fixed processing cost, which cannot be affected by the firm. Variation in this cost should therefore be read as variation across markets, regulations, disclosure technologies, or institutional features, for example, because disclosures are more standardized, easier to compare, or interpretable by third-party tools.

We show that the use of complex disclosure can be optimal for firms even in such an envi-

ronment with fully rational consumers, who make payoff-maximizing choices when presented with hard-to-understand information.³ When consumers are presented with complex information, assessing product attributes becomes costly, and consumers’ willingness to pay this cost depends on the alternatives available to them, i.e., their outside option. Learning is valuable only when it is likely to change the consumer’s decision: after observing a complex message, the consumer compares the value of buying based on the information already available, the value of taking the outside option, and the value of paying the processing cost to learn the exact quality before deciding. If other good products are available, consumers may rely on those “safer” alternatives. Similarly, if no good alternative exists, costly learning may be pointless because the firm’s product is already sufficiently attractive. Hence, consumers with very low or very high outside options may optimally choose not to pay the information-processing cost, while consumers with intermediate outside options may want to make a more informed purchase. The fixed processing cost and the firm’s complex disclosure rule, therefore, jointly affect the value of processing information for different groups of consumers. By changing which quality realizations are pooled behind the complex message, the firm changes not only consumers’ beliefs but also which consumers find it worthwhile to learn. This introduces a new strategic channel: the firm can target the group of consumers who will exert effort to process complex information, changing marginal costs and benefits of such strategic choice.

This novel channel can lead to the counterintuitive result that even without exploiting consumer unsophistication, complex disclosure may become more likely when markets are highly competitive and information-processing costs are low. To show this possibility, we restrict attention to two meaningful distributions of consumers’ outside options: unimodal distributions, representing markets with homogeneous consumers, sharing access to similar alternatives; and U-shaped distributions, representing segmented markets, with consumers having access to either very high or very low outside options. The main mechanism we identify is the following. When information-processing costs are extremely high, complex disclosure is effectively information censorship, and large availability of attractive outside options disciplines the firm’s disclosure choices by shifting consumers toward better alternatives. When processing costs are lower, however, this “*walk-away*” threat is mitigated: consumers with very attractive outside options may still decide to process complex disclosures and buy if the product turns out to be good, decreasing the firm’s cost of using complexity. As a consequence, even if low

³A large literature has focused, instead, on departures from rationality that can support information withholding. In particular, if consumers fail to draw sufficiently skeptical inferences from missing or hard-to-understand information, classic unraveling logic may fail even when information is verifiable (e.g., [Spiegler, 2016](#); [Jin, Luca and Martin, 2021](#)).

information-processing costs can intuitively reduce the benefit of complexity, in some markets they also reduce its cost, possibly leading to a reduction in transparency. As a result, depending on the underlying competitive structure, policies that reduce processing costs or that target markets where learning is already accessible may lead to an increase in equilibrium complexity. In this sense, the paper contributes to the literature—and potentially to policy discussions—by showing that the interaction between the structure of the market and the information-processing cost can overturn the standard disciplining logic of competitive forces. Hence, in practice, heterogeneity in processing ability is not the only factor to account for when designing regulations.⁴ We show that heterogeneity in available outside options is also key, as it links learning incentives to the available market alternatives.

In addition to this main result, we also discuss how the structure of receivers’ outside options affects the type of information the sender is likely to disclose in a complex way. Indeed, when markets tend to be more homogeneous, and the outside option follows a unimodal distribution, we show that information complexity is mainly used to aggregate high-quality realizations and maximize the consumers’ positive perception of a complex message. Instead, when the outside option follows a U-shaped distribution, the segmented nature of the market leads the sender to use complex information to hide low qualities and transparently disclose the high ones to sometimes capture harder-to-reach consumers. This pattern confirms a result already discussed in the literature that information complexity is not only a feature of low-quality products (e.g., [Asriyan, Foarta and Vanasco, 2023](#); [Aghamolla and Smith, 2023](#)). However, we motivate it from an information design perspective, showing how the choice of complexity can be a tool to attract different types of consumers and affect their learning choices.

From a methodological point of view, we study an information design setting in which receivers can acquire, at a cost, more information than what is conveyed by the sender’s message. Using the structure of the receiver’s information-acquisition problem, we derive necessary optimality conditions for the sender’s commitment to a complex disclosure rule.⁵ These

⁴Heterogeneity in information-processing cost alone would be closer to standard costly-search mechanisms; our channel instead links the incentive to learn to the value of walking away.

⁵The sender’s commitment assumption mainly plays a methodological role: it fixes the interpretation of complex messages and allows consumers to compare the costs and benefits of learning. It allows us to abstract from strategic uncertainty about the meaning of complex messages and to focus on consumers’ rational learning decisions. From an application perspective, a firm’s commitment power can be interpreted as a reduced-form representation of markets with recurring disclosure interactions, in which consumers can form rational expectations about firms’ disclosure behavior over time, as in credit-card agreement formats or digital-app privacy terms ([Best and Quigley, 2024](#); [Mathevet, Pearce and Stacchetti, 2025](#)).

conditions identify when the sender has incentives to reveal a quality realization transparently or pool it behind a complex message. Because the general problem admits no general closed-form solution under arbitrary distributions, these conditions are most useful in restricting the form of any candidate rule. This also distinguishes the analysis from standard censorship problems, since changing the use of complexity affects not only consumers' beliefs but also their incentives to pay the information-processing cost. We then apply these necessary conditions to specific parametric environments, triangular and Beta distributions, to analytically derive the structure of the optimal disclosure rule and implement numerical simulations. Our numerical simulations allow us to perform clean comparative statics in complexity and consumers' welfare and highlight the role of strong outside options in triggering increases in complex disclosure. Such an effect on the firm's strategic incentives may prevent a strict improvement in consumers' welfare when the information-processing cost decreases.

Overall, our results show that mandated disclosure policies may not lead to full information transmission when firms retain control over how information is presented. Their effectiveness can depend on the competitive structure of the market and on the cost consumers face to process complex information. Policies that make disclosures easier to understand may improve consumer learning, but they may also backfire by strengthening firms' incentives to obfuscate information. Accounting for the interaction between these forces can provide useful guidance to regulators when designing policies aimed at increasing information transparency.

The rest of the paper is organized as follows. Section 1.1 reviews the related literature; Section 2 presents the model; Section 3 characterizes the optimal information acquisition choices of the consumer; Section 4 studies the firm's choice of complex disclosure strategy and reports the main analytical results; Section 5 discusses our key numerical results; Section 6 discusses the possible implications for regulatory policies; Section 7 concludes.⁶

1.1 Related Literature

This paper is closely related to the price-obfuscation literature, which provides mixed findings on how the competitive structure of the market affects firms' incentives to make price structures more complex to understand.⁷ Carlin (2009) shows how price dispersion in retail financial markets can be the outcome of complexity strategically added to products' prices. As in our paper, increased competition can lead to more complexity. However, this channel relies on

⁶All the proofs can be found either in the Appendix or in the Online Appendix.

⁷For a more general survey of the choice complexity induced by firms' strategic choices, see Spiegel (2016).

complexity affecting investors’ ability to compare products and on the endogenous nature of investors’ lack of sophistication, which relaxes competitive pressure. Similar results are derived by [Chioveanu and Zhou \(2013\)](#), who study price framing and framing complexity in competitive markets.⁸ In a duopoly setting, [Gu and Wenzel \(2014\)](#) and [Chioveanu \(2019\)](#) show, in different environments, that more prominent firms, which face weaker competitive pressure, tend to use more price complexity. Different patterns are derived by [Carlin and Manso \(2011\)](#), who study a dynamic model in which higher competition increases investors’ sophistication and decreases obfuscation.⁹ A common feature of many of these works is the presence of a share of unsophisticated, inattentive, or confused consumers, which can be affected by firms’ complexity choices and exploited to extract additional rent. Our paper, instead, focuses on the learning decisions of fully rational consumers when heterogeneous alternatives are available. Higher complexity cannot change the distribution of consumers’ outside options—our reduced-form notion of competitive pressure—but it can change which consumers find it optimal to learn.

A few papers study the use of complexity in disclosure choices. Both [Jin, Luca and Martin \(2022\)](#) and [de Clippel and Rozen \(2025\)](#) study theoretically and experimentally sender-receiver games in which, as in our paper, disclosure is mandated but can be obfuscated by excessive complexity. [Jin, Luca and Martin \(2022\)](#) reinterpret the standard disclosure game proposed by [Milgrom \(1981\)](#) by replacing information withholding with unnecessary complexity. [de Clippel and Rozen \(2025\)](#) study a related mandated-disclosure environment in which messages can be transparent or obfuscated, and obfuscated messages require costly perception. In their model, strategic inference and costly perception coexist because the sender’s intended communication goal is realized only imperfectly.¹⁰ Other papers study related forms of product or disclosure complexity. [Asriyan, Foarta and Vanasco \(2023\)](#) study the joint determination of product complexity and quality, showing that complexity is not necessarily a signal of low quality. A similar link between complexity and high quality is derived by [Aghamolla and Smith \(2023\)](#), who develop a model to rationalize the documented non-monotonicity between firms’ performance and disclosure complexity. In their model, complexity can either provide more precise information to investors or obfuscate unfavorable news, and investors’ sophisti-

⁸[Chioveanu \(2020\)](#) also studies competition when firms can set price complexity. [Piccione and Spiegler \(2012\)](#) investigate how price formats affect product comparability in a Bertrand setting.

⁹Other papers study price obfuscation as a way to soften market constraints by exploiting myopia, search costs, or limited attention. See, among others, [Gabaix and Laibson \(2006\)](#), [Wilson \(2010\)](#), [Ellison and Wolitzky \(2012\)](#), [Ellison \(2005\)](#), [Ellison and Ellison \(2009\)](#), [Martin \(2017\)](#), and [Janssen and Kasinger \(2024\)](#). Relatedly, [Monte and Linhares \(2023\)](#) and [Xu \(2025\)](#) study information design problems in which the sender can directly affect the cost of information acquisition.

¹⁰[Perez-Richet and Prady \(2011\)](#) show how certification systems generate excessive product complexity.

cation determines which effect dominates.

From a methodological point of view, this paper relates to the broad information design literature. The sender’s ability to commit to a complexity rule follows a broad stream of persuasion papers, with pioneer works [Kamenica and Gentzkow \(2011\)](#), and [Rayo and Segal \(2010\)](#).¹¹ However, in our paper, the feasible information structures are substantially restricted, receivers are heterogeneous in their payoffs, and can acquire additional information at a cost. The private heterogeneity in receivers’ preferences through the outside option relates to [Rayo and Segal \(2010\)](#) and [Kolotilin \(2018\)](#), who extend information design problems by considering heterogeneous receiver types. [Kolotilin, Mylovanov and Zapechelnyuk \(2022\)](#) derive necessary and sufficient conditions for the optimality of upper-censorship rules, where favorable information is pooled together by the sender. The high-cost limit of our model reproduces the patterns in [Ginzburg \(2019\)](#), who studies optimal censorship when information can either be revealed or hidden from heterogeneous receivers. This pattern arises in our environment when the information-processing cost becomes large enough, often leading to a similar structure for the complex disclosure rule compared to the case in which costs are, instead, low. However, as mentioned in Section 1, learning opportunities introduce different channels in the use of complexity that we explore in this paper. Especially close on the role of outside options is [Lyu \(2023\)](#), who studies information design for search goods with privately known outside options. The main difference is that in this model search is necessary for purchase, while in ours consumers can buy without learning and only decide whether to pay a fixed cost to collect more payoff-relevant information.¹²

This work also contributes to the stream of information design papers in which the receiver has the option to acquire costly information. Most related to us, [Bizzotto, Rüdiger and Vigier \(2020\)](#) study an information design game in which homogeneous receivers can acquire additional information of a given precision at a cost. [Matyskova and Montes \(2023\)](#) also allow receivers to acquire more information than that designed by the sender, but at a flexible rate and cost, as in the rational-inattention literature. They show that the sender provides information that prevents the receiver from gathering independent information in equilibrium. However, they also show how such a no-learning result would not necessarily survive under a fixed information acquisition cost. Recent work by [Dall’Ara \(2024\)](#) studies persuasion of a privately informed receiver with costly attention, but such attention affects the probability of observing the realized message rather than the decision to interpret it.¹³

¹¹See [Kamenica \(2019\)](#) for a survey of the Bayesian persuasion literature.

¹²See also [Kolotilin et al. \(2017\)](#) and [Candogan and Strack \(2023\)](#).

¹³Other papers combine information design and rational-inattention problems. For instance, [Bloedel and Segal](#)

Our sender needs to decide whether to reveal the information transparently or to make it costly to process. For consumers who do not pay the information cost, this is equivalent to a standard disclosure decision when the information is verifiable. This feature relates the paper to the broad voluntary disclosure literature, starting with [Grossman \(1981\)](#), [Milgrom \(1981\)](#), [Jovanovic \(1982\)](#), and [Dye \(1985\)](#). However, in this paper, the standard unraveling result does not apply because the firm commits ex-ante to a disclosure rule rather than deciding whether to disclose each realized quality after having observed it. The introduction of such an assumption connects our work to [Onuchic \(2025\)](#), who studies a disclosure problem in which advisors with private preferences can commit to disclosure rules about the profitability of their products, and to [Fr chet te, Lizzeri and Perego \(2022\)](#), who investigate the interaction of commitment power and information verifiability in communication.

2 The Model

There are two players, a firm (sender, he) and a consumer (receiver, she). The sender aims to sell a product of quality, net of price, x .¹⁴ In what follows, we refer to x as the quality or as the state. The sender's quality is distributed according to a distribution $\mathcal{F}(\cdot)$, with support $[\underline{x}, 1]$, for $\underline{x} \geq 0$ and continuous and positive density $f(\cdot)$. Such a distribution is common knowledge for both players.

While the quality is private information of the sender, he is subject to a mandatory disclosure policy forcing him to verifiably reveal it to the receiver. However, the sender retains control over the format of the disclosure. In particular, for each realization x , the sender has access to the set of messages $M(x) = \{x, m^c\}$. The sender can either disclose the state in an immediately readable format, denoted by x , or disclose a verifiable but costly-to-process message, denoted by m^c . Before any information processing, all states that are made complex are observationally equivalent to the receiver and are therefore denoted by the common message m^c . After paying an information-processing cost k , however, the receiver would perfectly learn the underlying state x . We refer to the decision to report x as transparent disclosure and to the decision to report m^c as complex disclosure. Given these primitives, we can describe the strategy of the sender through a *complex disclosure rule* $c : [\underline{x}, 1] \rightarrow [0, 1]$, which assigns to each x the

(2021) and [Wei \(2021\)](#) study settings in which a sender tries to persuade a rationally inattentive receiver, who pays an attention cost to process the received message.

¹⁴For simplicity, we do not model the effect of information on prices. However, most of our analysis would generalize to the presence of a fixed exogenous price. [Appendix E](#) discusses this extension.

likelihood of disclosing the complex message m^c (see Figure 1 for some examples).¹⁵ Hence, $c(x) = 0$ implies that the state x is transparently revealed, while $c(x) = 1$ represents a decision to fully pool such a state behind the complex message. The sender commits to a complex disclosure rule before observing the realization of x , and the choice of the rule is known to the receiver. This implies that when the message m^c realizes, the receiver would be unable to distinguish between all x such that $c(x) > 0$, but she would be able to use $c(\cdot)$ to compute $\mathbb{E}[x|m^c] = \hat{x} = \frac{\int_{\underline{x}}^1 xc(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx}$.¹⁶

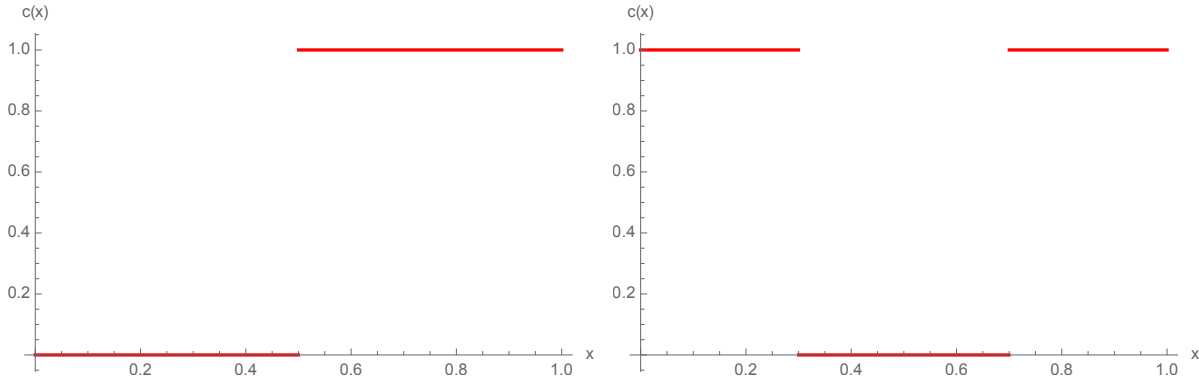


Figure 1: Examples of complex disclosure rules

In the left panel, every $x > 0.5$ is pooled in the complex message m^c with probability one. In the right panel, every $x < 0.3$ and every $x > 0.7$ is pooled in the complex message m^c with probability one.

The receiver has to decide whether to buy the product or not. In addition to the product, the receiver has an outside option y , which is her private information, and to which we will refer to as her type. Such outside options follow a distribution $G(\cdot)$, with support $[0, 1]$ —or $(0, 1)$ in case the distribution is not defined at the extremes—and density $g(\cdot)$. We assume that the density is positive and continuous over the domain, with the possible exception of the endpoints.¹⁷ $G(\cdot)$ is common knowledge for both players, and can be interpreted as the distribution of outside options of a heterogeneous population of receivers, which are then randomly matched to the sender. Before making any purchase decision, the receiver learns her outside option y and the message. Given the complex disclosure rule c and the observed message, she has to decide

¹⁵Throughout the paper, we restrict attention to measurable functions $c : [\underline{x}, 1] \rightarrow [0, 1]$.

¹⁶Such expectation is well-defined only when the chosen disclosure rule involves a positive complexity mass.

¹⁷This allows us to accommodate distributions like triangular distributions, for which $g(0) = g(1) = 0$ or U-shaped beta distributions, for which the density is not defined at the endpoints. In this last case, the domain should be defined as $(0, 1)$. However, given the absence of mass points in the distribution of x , everything in the analysis follows.

whether to buy the product or keep the outside option. If the realized message is transparent, the receiver can directly compare the state with the outside option. If, instead, the message is complex, she can decide to acquire more information before making a purchase decision. To gather more information, the receiver needs to pay a fixed information-processing cost $k > 0$. Once the cost is paid, the player perfectly learns the state.

It is important to highlight that, depending on the value of k , the model can accommodate full information transmission and pure censorship. Indeed, when $k \rightarrow 0$, the extra information the receiver can acquire is free, and the sender cannot effectively hide any state. When $k \rightarrow \infty$, learning becomes prohibitively costly for the receiver, and any state pooled in the complex message is censored, i.e., cannot be learned by the receiver (Ginzburg, 2019).

If the product is sold, the sender gets 1 and the receiver gets x . If, instead, no transaction takes place, the sender gets 0, and the receiver gets y . In case the receiver decides to pay the information-processing cost, k is deducted from the realized payoff.

Timing and Equilibrium Concept Given the primitives discussed above, the timing of the game is as follows:

1. Before the state x realizes, the sender commits to a complex disclosure rule $c(\cdot)$;
2. The state x realizes and the message $m \in \{x, m^c\}$ is sent;
3. The receiver learns her outside option y and the message m ;
4. If $m = m^c$, the receiver decides whether to pay the information-processing cost k ;
5. Given the available information, the receiver decides whether to buy or not;
6. Payoffs realize.

We study Perfect Bayesian equilibria of the game. Since the sender commits to a complex disclosure rule before observing the realization of x , the equilibrium analysis can be carried out by backward induction. We first characterize the receiver's optimal behavior after each message, given the realization of the outside option y and the beliefs induced by the disclosure rule c . Then, given the receiver's behavior, we study the sender's optimal complex disclosure rule.

3 Receiver's Problem

The main receiver's decision is whether to buy the good or keep the outside option y . To make this decision, the receiver relies on the information conveyed by the firm through the complex disclosure rule $c(\cdot)$ and the realized message m . If the message is transparent, i.e., $m = x$, the receiver directly learns the state and buys whenever $y \leq x$. If, instead, the message is complex, i.e., $m = m^c$, the receiver does not immediately observe x , but can compute $\hat{x} = \frac{\int_{\underline{x}}^1 xc(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx}$. This information captures the expected value the receiver can extract from the interaction with the sender, and such value, jointly with the outside option y , is also key to deciding whether to process the complex information at a cost k .

To characterize such a choice, we can first define the value function of a receiver with outside option y conditional on having received the complex message m^c , and assuming that no information is acquired: $V^{no}(y) = \max\{\hat{x}, y\}$. Equivalently,

$$V^{no}(y) = \begin{cases} \hat{x} & \text{if } y \leq \hat{x} \\ y & \text{if } y > \hat{x} \end{cases}$$

The function $V^{no}(\cdot)$ clearly highlights how, without any additional information, the expected value of purchasing the product is the right benchmark to compare to the value provided by the outside option. If, instead, the receiver decides to process the information contained in the complex message m^c , her value function becomes $V^i(y) = \mathbb{E}[\max\{x, y\} | m^c] - k$. Equivalently:

$$V^i(y) = \begin{cases} \hat{x} - k & \text{if } y \leq \underline{x} \\ y \frac{\int_{\underline{x}}^y c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} + \frac{\int_y^1 xc(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} - k & \text{if } y > \underline{x} \end{cases}$$

After paying the information-processing cost k , the receiver learns the state x , and so can perfectly assess the quality of the product. This implies that the expected value of an information acquisition choice needs to be computed, accounting for the cases in which learning would lead to a purchase decision ($x \geq y$) and for the cases in which learning would lead to consumption of the outside option ($x < y$). Notice that both $V^{no}(y)$ and $V^i(y)$ are weakly increasing in y , as better outside options lead to better expected outcomes, no matter the information acquisition decision. The direct comparison of $V^{no}(y)$ and $V^i(y)$ determines whether it is optimal for a receiver of type y to process information after observing the complex message $m = m^c$ (see Figure 2 for two examples).

At this point, it is straightforward to define the value function conditional on observing a complex message m^c , which, to simplify the notation, we denote as $V(\cdot)$. For each $y \in$

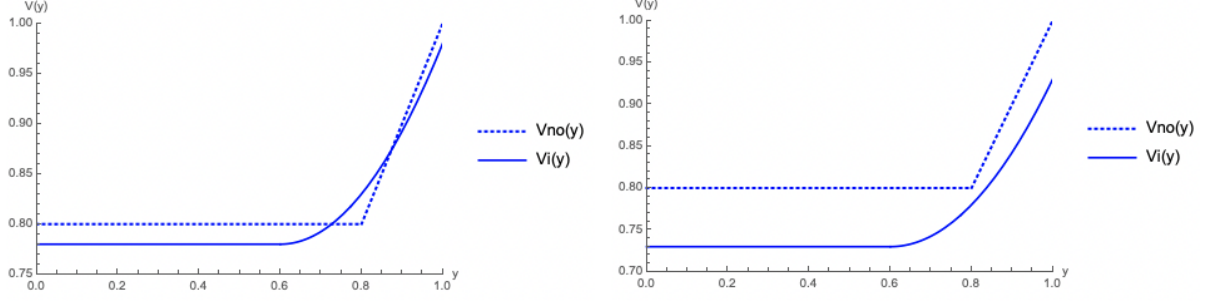


Figure 2: Value functions for the receivers conditional on receiving $m = m^c$

In both panels, $F \sim \mathcal{U}[0, 1]$ and $c(x) = 0$ for all the $x \in [0, 0.6)$ and $c(x) = 1$ for all the $x \in [0.6, 1]$. In the left panel $k = 0.01$ and in the right panel $k = 0.07$.

$[0, 1]$, $V(y) = \max\{V^{no}(y), V^i(y)\}$. We can use this value function to characterize the information-processing decisions of each receiver of type y given a complex disclosure rule c and an information-processing cost k , and no matter the shape of the distributions $F(\cdot)$ and $G(\cdot)$. It is important to highlight that such a decision matters only if the receiver observes a complex message, so the whole analysis takes as given that $m = m^c$.

Proposition 1. *Fix any complex disclosure rule c and information-processing cost $k > 0$. There exist $\underline{x} \leq \underline{y} < \hat{x} < \bar{y} \in [0, 1]$ such that*

$$V(y) = \begin{cases} V^{no}(y) & \text{if } y < \underline{y} \text{ or } y > \bar{y} \\ V^i(y) & \text{if } \underline{y} \leq y \leq \bar{y} \end{cases} \quad (1)$$

if and only if $V^i(\hat{x}) > V^{no}(\hat{x})$. Furthermore, $V(y) = V^{no}(y)$ for all $y \in [0, 1]$ if and only if $V^i(\hat{x}) \leq V^{no}(\hat{x})$. At the boundary where $V^i(\hat{x}) = V^{no}(\hat{x})$, the learning region disappears, and the two thresholds coincide, such that $\underline{y} = \bar{y} = \hat{x}$.

Proposition 1 achieves two important goals: (1) it provides us with a simple necessary and sufficient condition to assess whether information acquisition can occur under the complex disclosure rule c ; (2) it identifies the pattern of such learning, which only involves receivers with intermediate outside options (Figure 3 provides a graphical representation).

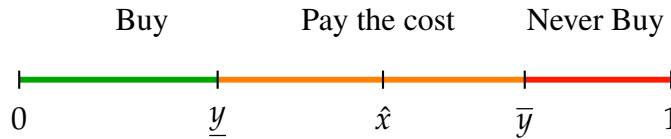


Figure 3: Consumers' learning behavior after observing a complex message m^c

Regarding the first point, the condition $V^i(\hat{x}) > V^{no}(\hat{x})$ can be conveniently expressed as

$$k < \frac{\int_{\underline{x}}^{\hat{x}} (\hat{x} - x)c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} = \frac{\int_{\hat{x}}^1 (x - \hat{x})c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx}.$$

This suggests that some information acquisition can be sustained in equilibrium only when a receiver of type $y = \hat{x}$ finds learning appealing. The economic intuition behind this condition is straightforward. The receiver type who faces the most challenging purchase decision in case of a complex message m^c is exactly the type with $y = \hat{x}$. Indeed, when the receiver's outside option coincides with the average quality behind a complex message, the receiver is initially indifferent between the two purchase decisions, and any additional information can help break the tie. This makes information very valuable, possibly exceeding the information-processing cost k . This type of reasoning would extend to receivers with outside options close to $\hat{x}$, for which information can substantially affect the decision. When, instead, the value of the outside option gets further away from $\hat{x}$, extra information is unlikely to change the purchase decision of the receiver, reducing the value of costly learning. As a consequence, verifying the willingness of type $\hat{x}$ to acquire information is enough to understand whether learning can be rationally sustained, given a disclosure choice of the sender and the information cost k .

The pattern of the receivers' learning decision directly follows from this discussion. Indeed, in case a complex message m^c triggers some positive learning, Proposition 1 identifies two threshold types $\underline{y} < \bar{y}$ which are sufficient to fully describe the information acquisition pattern. Specifically, the receiver types who will find it optimal to acquire information are the ones close to $\hat{x}$, i.e., the intermediate y for which a more precise signal is likely to affect the purchase choice and, so, be valuable. Instead, the receiver types below $\underline{y}$ and the ones above $\bar{y}$ are too confident about their choice, making the extra information too costly compared to its expected value. For this reason, such extreme types do not engage in information acquisition and, instead, generate two key groups from the sender's perspective: (1) a group that certainly buys the product ($y < \underline{y}$); (2) a group that can never be reached, as always consumes the alternative ($y > \bar{y}$).

The above discussion takes as given the choice of a complex disclosure rule c . However, c is strategically chosen by the sender, making $\underline{y}$ and $\bar{y}$ endogenous objects of the model. Using the structure of $V(y)$, we can derive two equations that pin down $\underline{y}$ and $\bar{y}$ under the assumption that a complex disclosure message triggers learning for a mass of receiver types:

$$\underline{y} \frac{\int_{\underline{x}}^{\underline{y}} c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} + \frac{\int_{\underline{y}}^1 xc(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} - k = \hat{x} \quad (2)$$

$$\underline{y} \frac{\int_{\underline{x}}^{\underline{y}} c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} + \frac{\int_{\underline{y}}^1 xc(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} - k = \underline{y} \quad (3)$$

In both equation (2) and equation (3), the left-hand side constitutes the expected value for a receiver of type y from processing the information contained in the complex message, which, as we saw before, grows with y and involves a constant cost k . The right-hand side, instead, represents the expected value in the absence of costly information acquisition. We can detect a key difference in the two equations. In equation (2), which describes the problem for the low-outside-options receiver group, the no-learning value is $\hat{x}$, as the best uninformed decision would be to buy the product. Instead, in equation (3), which represents the problem of high-outside-option receivers, the best no-learning decision would be keeping the outside option. This implies that $\underline{y}$ is defined as the highest receiver's outside option for which information acquisition is suboptimal before a purchase. Instead, $\bar{y}$ is the lowest value of the outside option for which a receiver relies on the available alternative without investigating the value hidden behind the complex disclosure. Anyone in the middle finds it optimal to engage in some learning, as the gap between the outside option and the expected value of the product is not large enough to make an uninformed decision. The described pattern can be visualized in the first panel of Figure 2, where the comparison between the cost and the benefits of the information is captured by the crossing of the two value functions conditional on the complex message m^c .

Finally, while the two thresholds are endogenously determined by the sender's choice of c , they also depend on the exogenous information acquisition cost k . Intuitively, for any given complex disclosure rule, a rise in the information-processing cost k triggers an increase in $\underline{y}$ and a decrease in $\bar{y}$, reducing the mass of receivers who process complex information. This is formalized in the following proposition.

Proposition 2. *Fix any complex disclosure c . For any value of k for which the learning region is non-empty, the threshold $\underline{y}$, as defined in Equation (2), is increasing in k , while the threshold $\bar{y}$, as defined in Equation (3), is decreasing in k .*

The discussion in this section shows that, for each complex disclosure rule c , and each information-processing cost k , the receiver's learning decisions can be described by two thresholds $\underline{y}$ and $\bar{y}$. This allows us to define the probability of sale from the sender's perspective for each value of x and for each possible disclosure decision. In particular, if $m = x$, the probability of a sale will simply be $G(x)$, i.e., the mass of receiver types who have worse outside

options than the transparently disclosed x . If, instead, $m = m^c$, and receivers learn, the probability of a sale will be $G(\underline{y})$ for $x \leq \underline{y}$, $G(x)$ for $x \in (\underline{y}, \bar{y})$ and $G(\bar{y})$ for $x \geq \bar{y}$. These cases map into the learning decisions of the receivers. After a complex message, all types $y \leq \underline{y}$ would buy the product, and if the state x is worse than $\underline{y}$, that would be the only group from which the sender can secure a sale. Similarly, any type above $\bar{y}$ would never interact with the sender, making the maximal sale mass $G(\bar{y})$ even when the quality is above such threshold. Finally, for intermediate qualities, the sender would secure the group who always buy, plus the share of types who would assess that the quality of the product is high enough through learning. In case, instead, the message is complex but no learning occurs, the probability of a sale will simply be $G(\hat{x})$, i.e., the share of receivers with outside option below the average complexity value. We use this structure to study the sender's problem in Section 4.

4 Sender's Complex Disclosure Rule

Given the payoff structure, the sender's goal is to maximize the probability of making a sale, taking as given the distribution of the outside options G and accounting for the effect of the choice of a complex disclosure rule c on the information acquisition strategy of the receiver. To make this dependence explicit in the general problem of the sender, we will summarize the pair of thresholds in the receiver's learning strategy as $(\underline{y}(c), \bar{y}(c))$. For every possible disclosure rule c , and corresponding $\hat{x}(c) = \mathbb{E}[x|m^c]$, Proposition 1 provides us with a clear condition for positive information acquisition, i.e., that at least the most intermediate type $\hat{x}(c)$ finds it profitable to acquire information. Recall that we can express such a condition as

$$\frac{\int_{\underline{x}}^{\hat{x}(c)} (\hat{x}(c) - x) c(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} > k. \quad (4)$$

When solving his optimization problem, the sender has to keep in mind that every time condition (4) holds, a non-empty set of customers $(\underline{y}(c), \bar{y}(c))$ is willing to acquire information, affecting the structure of his objective function. In particular, the payoff the sender can achieve through a complex message depends on the learning choices triggered for the receiver. Without any learning, the sender can effectively reach, through complex disclosure, any type below $\hat{x}(c)$, while when learning becomes part of the receiver's strategy, the sender can only achieve the payoff discussed at the end of Section 3. This implies that the exact problem the sender solves depends on condition (4), which is, itself, endogenously determined by the disclosure rule c .

This prevents us from deriving optimality conditions that only depend on the parameters of the model, but we need to consider these multiple specifications for the sender's objective function. The approach we follow is to use the optimal information acquisition choices of the receiver to set up the appropriate sender's problem and derive necessary conditions which, depending on the shape of G , allow us to numerically identify the optimal complex disclosure rule for the sender. To achieve this goal, we split the problem of the sender into three subproblems, listed below, depending on the induced information-processing decision for the receiver. It is important to highlight that the degenerate rule in which no complexity is used is considered as a separate candidate, which guarantees the sender a payoff of $\int_{\underline{x}}^1 G(x)f(x)dx$. For the problems below to be well-defined, we need some positive complex disclosure mass.

1. The “*Low-Cost Problem*”, in which the cost k and the complex disclosure rule c are such that a positive mass of receivers decide to acquire information, i.e., where condition (4) is satisfied with strict inequality. Formally, this optimization problem can be expressed as

$$\begin{aligned} \max_{c: [\underline{x}, 1] \rightarrow [0, 1]} & \int_{\underline{x}}^1 G(x)f(x)dx + \int_{\underline{x}}^{\underline{y}(c)} \left(G(\underline{y}(c)) - G(x) \right) c(x)f(x)dx \\ & - \int_{\underline{y}(c)}^1 (G(x) - G(\bar{y}(c))) c(x)f(x)dx \end{aligned} \quad (5)$$

$$\text{s.t. } \frac{\int_{\underline{x}}^{\hat{x}(c)} (\hat{x}(c) - x)c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} > k \quad (6)$$

$$\int_{\underline{x}}^{\underline{y}(c)} (\underline{y}(c) - x)c(x)f(x)dx - k \int_{\underline{x}}^1 c(x)f(x)dx = 0 \quad (7)$$

$$\int_{\underline{y}(c)}^1 (x - \bar{y}(c))c(x)f(x)dx - k \int_{\underline{x}}^1 c(x)f(x)dx = 0 \quad (8)$$

where the objective function (5) represents the probability of making a sale under the disclosure rule c , equation (6) is the (slack) learning constraint for the receiver and equations (7) and (8) rearrange the definition of $\underline{y}$ and $\bar{y}$, given the previously defined equations (2) and (3), under the assumption that $\int_{\underline{x}}^1 c(x)f(x)dx > 0$, which is necessary for a positive mass of receivers to acquire information. We study this problem in detail in Section 4.1.

2. The “*High-Cost Problem*”, in which the cost k and the complex disclosure rule c are such that receivers find it strictly optimal not to acquire any additional information;¹⁸

$$\max_{c: [\underline{x}, 1] \rightarrow [0, 1]} \int_{\underline{x}}^1 G(x)f(x)dx + \int_{\underline{x}}^1 (G(\hat{x}(c)) - G(x))c(x)f(x)dx \quad (9)$$

¹⁸This subproblem is equivalent to [Ginzburg \(2019\)](#), where there is no information processing, and the sender can fully censor some of the information.

$$\text{s.t. } \frac{\int_{\underline{x}}^{\hat{x}(c)} (\hat{x}(c) - x) c(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} < k \quad (10)$$

$$\int_{\underline{x}}^1 (\hat{x}(c) - x) c(x) f(x) dx = 0 \quad (11)$$

where the objective function (9) represents the probability of making a sale under the disclosure rule c , equation (10) is the (slack) no-learning constraint of the receiver (i.e. the condition for which $\hat{x}$ strictly prefers not to learn) and equation (11) captures the definition of $\hat{x}$ under the assumption that $\int_{\underline{x}}^1 c(x) f(x) dx > 0$, which is necessary for having a well-defined $\mathbb{E}[x|m^c]$. We study this problem in detail in Section 4.2.

3. The “*Binding Problem*”, in which the complex disclosure rule c is chosen to make the most intermediate type $\hat{x}$ indifferent between acquiring and not acquiring additional information at cost k . In this case the sender’s choice of c is made to disincentivize any learning for the receiver. The structure of this problem is equivalent to the High-Cost Problem, with the only difference that equation (10) needs to be satisfied with equality. We discuss this problem in Section 4.3.

4.1 Sender’s Low-Cost Problem

We can start from the sender’s Low-Cost Problem (LCP), which involves positive learning and differentiates our setting from a standard censorship one. Before looking at a possible approach to solve this problem, we can more carefully discuss the economic interpretation behind the sender’s objective function. Indeed, the objective (5) contains three objects that have a clear economic meaning. The first addend, $\int_{\underline{x}}^1 G(x) f(x) dx$, corresponds to the full transparency payoff, which only depends on the outside option distribution G and on the quality distribution $\mathcal{F}$. The second addend, $\int_{\underline{x}}^{\underline{y}(c)} (G(\underline{y}(c)) - G(x)) c(x) f(x) dx$, captures the gain in sale probability, compared to the full transparency scenario, from the use of information complexity. This gain only realizes among the types below $\underline{y}$, who end up buying without acquiring additional information and can mistakenly end up with a lower quality value than their outside option. Finally, the third addend, $\int_{\bar{y}(c)}^1 (G(x) - G(\bar{y}(c))) c(x) f(x) dx$, measures the loss in sale probability compared to the full transparency benchmark due to the group of consumers who stay uninformed and keep their outside option. Compared to the case with full information, receivers with type above $\bar{y}$ may avoid products that would guarantee a higher value than the outside option, leading to a lower sender’s payoff even when the state is high enough. From this interpretation, it is easy to see how the optimal complex disclosure rule

needs to be designed to balance these possible trade-offs: maximize the gains from low-outside-option consumers, while controlling the cost of losing high-outside-option ones.

Given the endogeneity of the two information-processing thresholds, deriving a solution to the LCP problem is not an easy task. Indeed, any change in c would generate a deviation in $\underline{y}(c)$ and $\bar{y}(c)$, and this makes it hard to derive straightforward conditions to characterize the solution.¹⁹ To mitigate this challenge, we can notice that the structure of the problem indirectly allows the sender to choose $\underline{y}$ and $\bar{y}$ through the disclosure rule c , which is converted into the strategic response of the receiver by equations (7) and (8). We can exploit this structure and write an equivalent simplified version, which we refer to as LCP2, in which the sender chooses three variables: $\underline{y}$, $\bar{y}$, and c . Formally, the new problem is

$$\begin{aligned} \max_{c: [\underline{x}, 1] \rightarrow [0, 1], \underline{y}, \bar{y}} \quad & \int_{\underline{x}}^1 G(x) f(x) dx + \int_{\underline{x}}^{\underline{y}} (G(\underline{y}) - G(x)) c(x) f(x) dx \\ & - \int_{\bar{y}}^1 (G(x) - G(\bar{y})) c(x) f(x) dx \end{aligned} \quad (12)$$

$$\text{s.t. } \bar{y} - \underline{y} > 0 \quad (13)$$

$$\int_{\underline{x}}^{\underline{y}} (\underline{y} - x) c(x) f(x) dx - k \int_{\underline{x}}^1 c(x) f(x) dx = 0 \quad (14)$$

$$\int_{\bar{y}}^1 (x - \bar{y}) c(x) f(x) dx - k \int_{\underline{x}}^1 c(x) f(x) dx = 0 \quad (15)$$

$$\int_{\underline{x}}^1 c(x) f(x) dx > 0 \quad (16)$$

In the LCP2, $\underline{y}$ and $\bar{y}$ do not explicitly depend on c anymore, but they are chosen by the sender to satisfy the two defining constraints (14) and (15). In addition, together with condition (16), that requires strictly positive complexity, Proposition 1 allows us to replace the learning condition (6) with the simpler inequality (13): information is processed every time a positive mass of receivers decides to learn, which translates into a strict ranking between the learning thresholds, $\underline{y} < \bar{y}$. From this discussion, the equivalence of LCP and LCP2 is straightforward (Online Appendix D.3 provides a formal argument for this equivalence).

Given the absence of any meaningful difference between the LCP and the LCP2, we can derive optimality conditions through the simplified problem, using the Lagrange multiplier approach. In doing so, we assume that conditions (13) and (16) are satisfied with strict inequality. The reason for this choice is that the focus of our theoretical analysis is not to derive full optimality conditions for the LCP, which would be unfeasible given the lack of concavity of

¹⁹Online Appendix C discusses a calculus of variation approach to the problem.

the objective function, but to derive necessary conditions that can be used to identify the structure of the optimal disclosure rule when G and $\mathcal{F}$ take specific forms (see Section 5). In this section, the goal is to derive such conditions for the case in which some learning occurs.

Proposition 3. *Let $(c, \underline{y}, \bar{y})$ be a solution to the LCP2 with positive learning. Then, it needs to satisfy the following necessary conditions:*

1. *For all $x \in [\underline{x}, \underline{y}]$, $c(x) = 1$ if $\int_{\underline{x}}^{\underline{y}} (g(z) - g(\underline{y})) dz + k(g(\underline{y}) - g(\bar{y})) > 0$ and $c(x) = 0$ if $\int_{\underline{x}}^{\underline{y}} (g(z) - g(\underline{y})) dz + k(g(\underline{y}) - g(\bar{y})) < 0$;*
2. *For all $x \in (\underline{y}, \bar{y})$, $c(x) = 1$ if $k(g(\underline{y}) - g(\bar{y})) > 0$ and $c(x) = 0$ if $k(g(\underline{y}) - g(\bar{y})) < 0$;*
3. *For all $x \in [\bar{y}, 1]$, $c(x) = 1$ if $-\int_{\bar{y}}^x (g(z) - g(\bar{y})) dz + k(g(\underline{y}) - g(\bar{y})) > 0$ and $c(x) = 0$ if $-\int_{\bar{y}}^x (g(z) - g(\bar{y})) dz + k(g(\underline{y}) - g(\bar{y})) < 0$;*
4. *$\underline{y}$ is such that $\int_{\underline{x}}^{\underline{y}} (\underline{y} - x)c(x)f(x)dx - k \int_{\underline{x}}^1 c(x)f(x)dx = 0$;*
5. *$\bar{y}$ is such that $\int_{\bar{y}}^1 (x - \bar{y})c(x)f(x)dx - k \int_{\underline{x}}^1 c(x)f(x)dx = 0$.*

Every time the relevant conditions in (1)-(3) are satisfied with equality, $c(x)$ can take any value in $[0, 1]$.

Proposition 3 gives us the necessary conditions that any LCP2 solution—and so any LCP solution—with positive learning must satisfy.²⁰ As mentioned before, given the non-concavity of the objective function, these conditions are not sufficient. In particular, they may identify learning thresholds that do not correspond to globally optimal ones. However, the derived conditions can give us important insights into the structure of the optimal disclosure rule and guide the use of numerical optimizations to identify the optimal disclosure strategy.

It is useful to discuss the economic interpretation of these necessary conditions. Conditions (1)-(3) govern the specific choice of the complex disclosure rule c . Instead, conditions (4) and (5) discipline the choice of $\underline{y}$ and $\bar{y}$ to make it consistent with the receiver's optimization. We

²⁰The deterministic structure of the optimal rule reflects the linearity of the sender's problem in c : for fixed learning thresholds, the objective and constraints are linear in c , so optimal rules can be found at the extreme points of the feasible set. This connects our conditions to extreme-point arguments used in mechanism and information design (Börgers, 2015; Kleiner, Moldovanu and Strack, 2021). Because the learning thresholds are themselves endogenous to c , we work directly with the necessary conditions rather than through this approach.

can start by analyzing in more detail condition (2) for the intermediate values of x , meaning, the ones in between the induced learning thresholds. The condition shows that marginally increasing the likelihood of complex disclosure for intermediate states only affects the sender's payoff through changes in the learning thresholds. Indeed, from conditions (4) and (5), we can see that marginally increasing the probability of complex disclosure for some $x \in (\underline{y}, \bar{y})$ leads to an increase in $\underline{y}$ and a decrease in $\bar{y}$.²¹ This translates into more uninformed sales but also into more uninformed walk-away decisions following a complex message. This implies that the shape of G , and in particular, the magnitudes of the marginal changes $g(\underline{y})$ and $g(\bar{y})$ determine whether the gain below $\underline{y}$ outweighs the cost above $\bar{y}$, making this trade-off the main driver of the choice of $c(x)$ for the intermediate values of quality x .

Instead, increasing complexity for $x \leq \underline{y}$ or $x \geq \bar{y}$ has ambiguous effects on the learning thresholds, and also directly affects the likelihood of a sale in the two extreme groups. Regarding the information acquisition part, the value of acquiring information for any type y may decrease or increase, depending on which states have been affected by the change in complexity. This effect is summarized by two terms. The first one is the same as before, $k \left(g(\underline{y}) - g(\bar{y}) \right)$, and reflects the effect of the shift in the total complexity mass on the value of acquiring information, which, alone, would lead to an increase in $\underline{y}$ and a decrease in $\bar{y}$. The second one, instead, is $-g(\underline{y})(\underline{y} - x)$ for $x < \underline{y}$ and $g(\bar{y})(x - \bar{y})$ for $x > \bar{y}$, and captures the fact that, for extreme quality realizations, a marginal increase in complexity also changes the value of becoming informed relative to the value of the relevant alternative. Indeed, such an increase for quality realizations below $\underline{y}$ makes learning more attractive for the marginal threshold type, since acquiring information would allow such a type to avoid an ex-post suboptimal purchase. In terms of the sender's payoff, this translates into a negative effect, because it reduces the mass of consumers who buy without acquiring extra information. Instead, such an increase above $\bar{y}$ makes learning more attractive for the marginal upper type, because adding a high-quality realization to the complex pool raises the value of learning, as information may reveal that the quality realization is high enough to lead to a purchase. In terms of the sender's payoff, this has a positive effect as it reduces the mass of high-outside-option consumers who walk away after a complex message. The final impact is given by the combination of these different channels.

Finally, in both cases, we have an additional term which reflects the direct impact of marginal changes in complexity on the sender's probability of making a sale: $G(\underline{y}) - G(x)$ for $x < \underline{y}$

²¹This can be easily seen through the conditions (14) and (15). An increase in complex disclosure between $\underline{y}$ and $\bar{y}$ would increase $k \int_{\underline{x}}^1 c(x)f(x)dx$ in both equations, while the other addends would be unaffected by the change in c . This would translate into an increase in $\underline{y}$ and a decrease in $\bar{y}$.

and $-(G(x) - G(\bar{y}))$ for $x > \bar{y}$. The intuition for these terms is very simple. Increasing complexity for low qualities increases sales since consumers below $\underline{y}$ see the complex message more often, and buy without acquiring information also realizations $x < \underline{y}$ that were first disclosed transparently and only purchased by the types below x . Instead, increasing complexity for high realizations of x tends to reduce sales among types above $\bar{y}$, who would otherwise buy after observing that the state is good. Conditions (1) and (3) simply combine these simultaneous effects to identify whether the gains of complexity exceed the costs.

The solution procedure discussed so far applies only to candidate solutions satisfying constraint (13), i.e., $\underline{y} < \bar{y}$. However, the global solution to the sender's problem need not satisfy this condition. In particular, the sender may optimally choose a disclosure rule that induces no information acquisition, or the cost k may be too high for any learning to be feasible. This motivates the need to discuss procedures that also account for no-learning and boundary cases. We complete this analysis in Section 4.2 and Section 4.3.

4.2 Sender's High-Cost Problem

As we discussed in Section 4.1, for the Low-Cost Problem to be well defined, we need the positive learning condition (4) to be satisfied with strict inequality. When, instead, acquiring the information becomes excessively costly for the receiver, we can end up in a situation in which making information complex is equivalent to a censorship decision. This changes the structure of our optimization to the sender's High-Cost Problem (HCP) discussed above. Whether receivers can optimally ignore complex information depends, again, on the cost of the information k and on the complex disclosure rule c . Since this problem can be relevant in numerically searching for a global solution, we formally study it and derive, as for the LCP, necessary conditions that can help identify the optimal complex disclosure rule under specific assumptions over the distribution of outside options G .

First, as we did for the LCP, it is useful to discuss the economic interpretation behind the different ingredients of the optimization. The objective function (9) presents two addends. As in the LCP, the first one, $\int_{\underline{x}}^1 G(x)f(x)dx$, represents the benchmark from full information transparency. Instead, the second addend, $\int_{\underline{x}}^1 (G(\hat{x}(c)) - G(x))c(x)f(x)dx$, measures the gain in payoff resulting from the use of the complex message m^c compared to full transparency. Specifically, when the information is complex, all the types with outside option below $\hat{x}$ find it optimal to purchase the good, and all the other ones just walk away. In terms of constraints, equation (10) has the role to guarantee that no receiver learning happens in equilibrium, while

equation (11) defines the average complexity value $\hat{x}$.

This problem presents the same challenges discussed for the LCP. Similarly to what we did in Section 4.1, we can set up an alternative version of the problem, the HCP2, in which the sender directly chooses both $\hat{x}$ and c , under the constraints that $\hat{x}$ follows the right definition.

$$\max_{c: [\underline{x}, 1] \rightarrow [0, 1], \hat{x}} \int_{\underline{x}}^1 G(x) f(x) dx + \int_{\underline{x}}^1 (G(\hat{x}) - G(x)) c(x) f(x) dx \quad (17)$$

$$\text{s.t. } k - \frac{\int_{\underline{x}}^{\hat{x}} (\hat{x} - x) c(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} > 0 \quad (18)$$

$$\int_{\underline{x}}^1 (\hat{x} - x) c(x) f(x) dx = 0 \quad (19)$$

The equivalence between these two problems is straightforward.²² As in Section 4.1, we can assume that constraint (18) is satisfied with strict inequality and focus on deriving necessary conditions, through the Lagrange multiplier approach, within the problem in which learning is strictly suboptimal. Again, these necessary conditions will be used in Section 5 to numerically identify the optimal complex disclosure rule under some parametric assumptions.

Proposition 4. *Let $(c, \hat{x})$ be a solution to the HCP2. Then it needs to satisfy the following necessary conditions:*

1. *For every $x \in [\underline{x}, 1]$, $c(x) = 1$ if $\int_{\underline{x}}^{\hat{x}} (g(z) - g(\hat{x})) dz > 0$ and $c(x) = 0$ if $\int_{\underline{x}}^{\hat{x}} (g(z) - g(\hat{x})) dz < 0$;*
2. $\int_{\underline{x}}^1 (\hat{x} - x) c(x) f(x) dx = 0$.

Every time the relevant condition in (1) is satisfied with equality, $c(x)$ can take any value in $[0, 1]$.

Before moving on to the binding case, it is important to discuss the economic interpretation of the optimality condition (1), which captures some of the mechanisms already highlighted in Section 4.1. When marginally increasing the likelihood of complex disclosure, the sender affects two aspects of the problem, which are directly compared by the first derivative in condition (1): (i) the threshold $\hat{x}$, which modifies the set of types who are willing to make a purchase after the complex message m^c ; (ii) the likelihood of such a message. We can decompose condition (1) as $-g(\hat{x})(\hat{x} - x) + (G(\hat{x}) - G(x))$ for $x \leq \hat{x}$ and as $g(\hat{x})(x - \hat{x}) - (G(x) - G(\hat{x}))$

²²The equivalence of the constraints is immediate, and the objective function remains the same. Moreover, all the arguments from Online Appendix D.3 can be generalized to this case.

for $x > \hat{x}$. From this decomposition, it is immediate to see how the first addend captures the effect the change in c has on $\hat{x}$, modifying the purchase decision of the marginal receiver type. Such a change negatively impacts the sender payoff when complexity increases among low states and positively impacts it when it increases among high states. Instead, as for the LCP case, the second captures the direct effect on sales, which either increase or decrease, depending on whether complexity increases among low states, triggering more uninformed purchases from types below $\hat{x}$, or among high states, leading to more frequent uninformed walk-aways among receivers with outside option above $\hat{x}$. In both cases, the gain/loss in sale is with respect to $G(x)$, i.e., the sale probability the state x would generate if disclosed transparently.

4.3 Binding Problem

In our analysis of the LCP and HCP, we have assumed that the receiver's learning constraint never binds, and all receiver types are either strictly better off or strictly worse off in acquiring additional information when the observed message is the complex one. This excludes a meaningful case: the one in which the learning constraint is satisfied with equality, i.e.

$$\frac{\int_{\underline{x}}^{\hat{x}(c)} (\hat{x}(c) - x) c(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} = k.$$

This case has a special interpretation: given k and c , the type $\hat{x}$ is indifferent between acquiring and not acquiring information, generating a bridge between the HCP and the LCP with $\underline{y} = \bar{y} = \hat{x}$. This case can be interpreted, from an economic perspective, as the version of the problem in which the sender uses his discretion to choose a complex disclosure rule c that aims to marginally discourage learning.

Different from the LCP and the HCP, the Lagrange multiplier approach in the binding case does not deliver a simple characterization independent of the multipliers. For this reason, the necessary conditions we can derive are more complex and less easily interpretable. However, they are still an important piece of the numerical analysis we perform in Section 5. For completeness, Lemma 1 reports such necessary conditions.

Lemma 1. *Let $(c, \hat{x})$ be a solution to the sender's Binding Problem. Then, it needs to satisfy the following necessary conditions, where $\mu \geq 0$ is the multiplier on the no-learning constraint (18), and $\lambda \in \mathbb{R}$ is the multiplier on the constraint defining $\hat{x}$, captured by equation (19):*

1. *For every $x \in [\underline{x}, \hat{x}]$, $c(x) = 1$ if $\int_{\underline{x}}^{\hat{x}} [g(z) - (\mu + \lambda)] dz + \mu k > 0$ and $c(x) = 0$ if $\int_{\underline{x}}^{\hat{x}} [g(z) - (\mu + \lambda)] dz + \mu k < 0$;*

2. For every $x \in [\hat{x}, 1]$, $c(x) = 1$ if $-\int_{\hat{x}}^x (g(z) - \lambda) dz + \mu k > 0$ and $c(x) = 0$ if $-\int_{\hat{x}}^x (g(z) - \lambda) dz + \mu k < 0$;
3. $(g(\hat{x}) - \lambda) \int_{\underline{x}}^1 c(x) f(x) - \mu \int_{\underline{x}}^{\hat{x}} c(x) f(x) = 0$;
4. $\int_{\underline{x}}^1 (\hat{x} - x) c(x) f(x) dx = 0$.

Every time the relevant condition in (1)-(2) is satisfied with equality, $c(x)$ can take any value in $[0, 1]$.

It is easy to see how, simply setting $\mu = 0$, which must be true when the learning constraint is slack, we can recover the same conditions discussed in Proposition 4. However, in this binding version of the learning constraints, the decision not to acquire information is somewhat harder to induce, and the sender needs to make the most intermediate type indifferent between paying or not the information acquisition cost. Clearly, this can have an effect on the structure of the complex disclosure rule, which is captured by the additional multiplier in the problem.

4.4 The Solution to the General Problem

The previous sections provide a decomposition of the sender's general problem. However, all parts need to be jointly considered when deriving the global solution. Just to summarize the analysis, for any non-degenerate disclosure rule c , we define the receiver's net gain of learning after observing the complex message as

$$\Delta(c, k) = \frac{\int_{\underline{x}}^{\hat{x}(c)} (\hat{x}(c) - x) c(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} - k.$$

Then every candidate solution to the general problem falls into one of three cases. If $\Delta(c, k) > 0$, learning is strictly optimal for some receiver types, and the candidate solution belongs to the Low-Cost Problem. If $\Delta(c, k) < 0$, learning is strictly suboptimal, and the candidate solution belongs to the High-Cost Problem. Finally, if $\Delta(c, k) = 0$, the receiver's learning constraint binds, and the candidate solution belongs to the binding problem. Therefore, although a closed-form solution cannot be obtained in general, the analysis of Sections 4.1, 4.2, and 4.3 provides a useful roadmap for solving the general problem. Under additional assumptions on the outside-option distribution G , one can derive more structure for the disclosure rules from the necessary conditions associated with each sub-problem, and then look for a numerical optimizer under the appropriate learning constraint. Finally, one can compare the sender's payoff across the feasible candidate solutions. Section 5 implements this approach for triangular

outside-option distributions and for standard two-parameter Beta distributions. More generally, the theoretical restrictions apply to bounded-support outside option distributions whose densities are monotone, unimodal, or U-shaped.

5 Applications: Parametric Outside Option Distributions

We now impose additional structure on the key distribution in our model and study how the receiver’s equilibrium behavior, characterized in Section 3, together with the necessary conditions for the sender’s optimality derived in Section 4, can be used to derive the structure of the optimal complex disclosure rule. We then examine how this rule changes as competition intensifies, that is, as the distribution of outside options shifts towards higher values.²³ This allows us to connect the general solution characterization to the main question of the paper: whether stronger market alternatives discipline the sender’s use of complex disclosure or, instead, can amplify it. To this end, we first assume that quality is uniformly distributed on $[0, 1]$, that is, $x \sim \mathcal{U}[0, 1]$. Although this assumption does not affect the main logic of our analysis, and importantly, the structure of the optimal complex disclosure rule, it allows us to derive closed-form solutions for the equilibrium thresholds $\underline{y}$ and $\bar{y}$. This makes the characterization of the solution more tractable and yields conditions with a clear economic interpretation.²⁴

Second, we focus on parametric outside-option distributions with bounded support and economically meaningful shape restrictions. We consider two classes of distributions. The first is the triangular distribution, which provides a tractable unimodal case and allows us to study effectively how changes in the mode affect the sender’s disclosure incentives. The second is the Beta distribution, which provides a flexible family of bounded-support distributions and allows us to capture both unimodal and U-shaped densities. This distinction is useful because the main economic forces depend primarily on the shape of the outside-option distribution. Unimodal densities describe relatively homogeneous markets, in which receivers’ outside options are generally concentrated on some specific value. One possible example is the market for social network apps, where most consumers in a geographic area have access to the same type of products. U-shaped densities, instead, describe segmented markets, in which receivers’ outside options tend to be more polarized and concentrated near the extreme points of the support. We can think about these markets as ones in which different groups of receivers have access to very different alternatives, like for instance financial and credit-card markets, creating more

²³In our simulations, this corresponds to a shift in a first-order stochastic dominance sense.

²⁴Appendix B.3 suggests that our main comparative statics do not depend on this distributional assumption on x .

polarized outside values. The triangular distribution gives a simple way to study the unimodal case, while the Beta family allows us to extend the analysis to both types of market structures.

Using these distributions, we first show that different shapes of the outside option can lead to very different strategic incentives for the sender, and, as a consequence, to different formats for the optimal complex disclosure rule. Second, we numerically explore the relationship between the competitive structure of the market, modeled by the outside-option distribution G , the accessibility of complex information, modeled by the cost k , and the strategic use of complexity. We show that better outside options and lower information-processing costs do not necessarily lead to more transparent disclosure; instead, the outcome depends on receivers' learning choices.²⁵ We also show that this can translate into no improvement in receivers' welfare and, in some cases, even into welfare losses, although the magnitudes we find are relatively small.

5.1 Homogeneous Markets: Unimodal G

We start by assuming that the outside-option distribution G has a unimodal density g with mode $\gamma \in [0, 1]$, i.e., we focus on cases in which g is strictly increasing on $[0, \gamma)$ and strictly decreasing on $(\gamma, 1]$. For this class of distributions, the analysis in [Ginzburg \(2019\)](#) suggests that the optimal complex disclosure rule in the HCP has an upper-censorship structure: there exists a threshold $a^{cens} \in [0, 1]$ such that all states $x < a^{cens}$ are transparently disclosed, while all states $x \geq a^{cens}$ are pooled under a complex message. The exact position of the censorship threshold depends on the distribution, but it is optimally located weakly below the mode γ . This has a straightforward implication. When g is monotonically decreasing, $\gamma = 0$, and the sender optimally withholds all information. Conversely, when g is monotonically increasing, $\gamma = 1$, the conditions in [Proposition 4](#) rule out any complex disclosure rule with a positive complexity mass. This implies that the sender transparently discloses all states. These results are well established in the literature ([Kolotilin et al., 2017](#); [Kolotilin, 2018](#); [Ginzburg, 2019](#)), and the economic intuition behind them connects directly to our market interpretation. When $\gamma = 0$, most receiver types are concentrated near low values of the outside-option distribution. As a result, the outside alternatives tend to be unattractive, even when transparent information is unavailable. This makes complex disclosure optimal for the sender. When, instead, $\gamma = 1$, most receiver types are concentrated near high values of the outside option, making alternative products very attractive. In this case, complex disclosure would discourage a large mass of consumers from making a purchase, making transparency preferable for the sender. When

²⁵Appendix [F](#) theoretically documents this possibility for triangular outside options.



Figure 4: Upper-censorship disclosure rule

$\gamma \in (0,1)$, these two forces interact, and a^{cens} is chosen to balance the gains and losses: making good states complex is the most effective way to extract rent from the receiver, as it guarantees sales among low types without incurring too many losses among high ones.

The pattern described above does not consider the possibility of information acquisition on the receiver's side, leaving open the question of how learning interacts with the sender's disclosure incentives. The analysis in the previous sections allows us to characterize the structure of the solution to the sender's LCP.

Theorem 1. *Assume that G has a unimodal density g with mode γ . Any optimal disclosure rule in the sender's LCP follows an upper-censorship structure, that is, there exists $a \in [0, \gamma]$ such that*

$$c(x) = \begin{cases} 0 & \text{if } x < a \\ 1 & \text{if } x \geq a \end{cases}$$

Hence, the sender's maximization collapses to a one-dimensional one.

Theorem 1 provides two main results. First, even when learning is possible for the receiver, the sender maximizes the sale probability by revealing low-value states and pooling together high-value states under a complex message (Figure 4 provides a graphical representation of this structure). Second, the theorem characterizes the position of the optimal threshold, which, as in the no-learning case, is located weakly below the mode γ .²⁶

We can use the insights of Theorem 1 to discuss the extreme cases of monotone densities when learning is allowed for the receiver. It is easy to see from the conditions in Proposition 3 that when $\gamma = 1$, the sender transparently reveals all states even in the LCP. This is not surprising, since the incentives for full transparency were already present in the absence of information processing. What is more surprising, instead, is that when $\gamma = 0$ and the density g is monotonically decreasing, the sender's optimal complex disclosure strategy is still to withhold all information. However, unlike the HCP, when the cost k is low enough, receivers acquire

²⁶The restriction $a \leq \gamma$ follows the same logic as in the censorship benchmark: above the mode, the marginal gain from improving the value of the complex message is too small relative to the loss from disclosing the marginal state transparently.

some information. The robustness of this no-disclosure result stems from the receiver's learning pattern. Although learning reduces the share of achievable sales when a complex message is used, as some receivers pay the information cost, most receiver types are concentrated near zero, where the same logic as in the censorship case applies. For these receivers, the outside option is sufficiently unattractive that even for low information-processing cost, they would not learn, treating the complex states as censored. For this reason, fully withholding information through complexity remains the optimal sender's strategy.

At this point, we can turn to the intermediate cases with $\gamma \in (0, 1)$, and discuss why the upper-censorship structure survives when receivers can learn. Even if receivers are now able to process complex information, the behavior of the groups below $\underline{y}$ and above $\bar{y}$ preserves the main trade-off the sender faces in the pure-censorship case. The lower-outside-option receivers buy after a complex message without learning, while the higher-outside-option receivers walk away without learning. The intermediate types are the new group: conditional on learning, they behave as if the information were transparent even if complex disclosure is used, and cannot be influenced by the disclosure choices. Therefore, learning changes the position of the optimal threshold by changing the composition of these groups, but it does not overturn the basic threshold logic. The sender still balances the gain from using complexity to attract lower-outside-option receivers against the cost of losing higher-outside-option ones.

Before moving to our simulations to compare the optimal thresholds in the LCP and HCP and discuss the intuition behind such comparison, it is useful to discuss how the receiver's learning behavior changes with the sender's complexity threshold. Denote by a the threshold characterizing the sender's upper-censorship disclosure rule. Given this structure and the assumption that $x \sim \mathcal{U}[0, 1]$, we can derive closed-form expressions for the two learning thresholds $\underline{y}$ and $\bar{y}$ that characterize the receiver's optimal strategy. In particular, by solving equations (2) and (3), we obtain

$$\underline{y} = a + \sqrt{2k(1-a)}, \quad \bar{y} = 1 - \sqrt{2k(1-a)}$$

with $\underline{y} < \bar{y}$ any time $a < 1 - 8k$. These formulas show that, for a given k , an increase in the complex-disclosure threshold a reduces the likelihood of a complex message but also shifts both $\underline{y}$ and $\bar{y}$ upward. The intuition is straightforward. A higher a increases the expected value of the complex message, making uninformed purchases more attractive and reducing the set of receiver types who would simply walk away without acquiring information. However, a change in a affects not only the purchase decision, but also the learning region. As a increases, the interval $[\underline{y}, \bar{y}]$ shrinks. Thus, the set of receiver types who acquire information shifts toward higher values of the outside option and contracts in length. This can positively affect the

sender's payoff, since some intermediate receivers would now buy directly after observing the complex message, and some high- y receivers would now find it worthwhile to acquire information before deciding whether to purchase rather than walk away. This pattern is key in the sender's consideration and shows how changes in the complex disclosure rule can effectively change the composition of receivers who are willing to learn, and introduce new mechanisms in the sender's disclosure choices.

Finally, Proposition 2 shows that, fixing the complex disclosure rule, a reduction in the information-processing cost k increases the set of receivers who acquire information. This can be verified directly from the formulas above: when k decreases, $\underline{y}$ decreases and $\bar{y}$ increases at the same absolute rate. Hence, unlike an increase in a , which reduces complexity, lower information-processing costs expand the learning region on both sides, reducing uninformed purchases, but also reducing uninformed walk-aways.

To conclude the discussion of unimodal outside-option distributions, we need to describe the structure of the solution when the receiver's learning constraint binds. While this case is less directly interpretable than the LCP and HCP cases, it must be considered in the numerical optimization, because it represents the boundary between learning and no learning. The following lemma summarizes the resulting threshold structure.

Lemma 2. *Assume that G has a unimodal density g with mode γ . The optimal disclosure rule in the sender's Binding Problem can be described with five (possibly degenerate) thresholds $0 \leq a_1 \leq a_2 \leq a_3 \leq a_4 \leq a_5 \leq 1$, such that*

$$c(x) = \begin{cases} 0 & \text{if } x < a_1 \text{ or } a_2 \leq x \leq a_3 \text{ or } a_4 \leq x \leq a_5 \\ 1 & \text{if } a_1 \leq x \leq a_2 \text{ or } a_3 \leq x \leq a_4 \text{ or } x > a_5 \end{cases}$$

Because some of these thresholds may coincide, the binding class nests simpler threshold rules as degenerate cases, including the previously described upper-censorship one. The numerical comparison below treats it as a candidate family, accounting for the possible degenerate cases, and evaluates it as a possible solution to the sender's problem.

5.1.1 Unimodal Outside Option: Numerical Results

Given the theoretical results derived in Section 5.1, we can show the following important result.

Result 1. *An increase in k has an ambiguous effect on the optimal threshold a . In particular, in markets with more favorable outside options, low information-processing costs can increase*

the likelihood of complex disclosure, and so decrease the mass of quality realizations that are transparently disclosed, compared to the limit censorship case with $k \rightarrow \infty$.

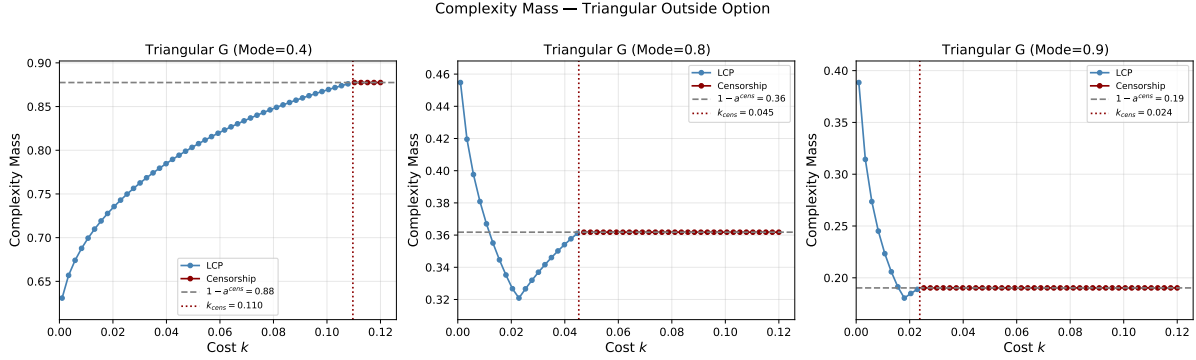


Figure 5: Probability of complex disclosure ($1 - a$) as a function of k , when G is a triangular distribution and $x \sim \mathcal{U}[0, 1]$

Figure 5 reports the predicted probability of complex disclosure, labeled as complexity mass, as k varies, for triangular outside option distributions with increasing modes, and setting the quality distribution to $\mathcal{U}[0, 1]$. Given such parametric assumptions, the represented complexity probability is simply $1 - a$, i.e., the x mass above the optimal upper-censorship threshold. The results in Figure 5 are obtained by numerically solving the sender's problem. The numerical simulations account for the possible structure of the problem (LCP, HCP and the Binding Problem) and select the payoff-maximizing solution (see Appendix B for more details). In all simulations studied in the paper, the binding candidate is never identified as the optimal one; the payoff-maximizing solution strictly falls into either the LCP or HCP, confirming that the optimal rule takes a simple upper-censorship form. In the figure, the solution to the HCP is labeled as censorship, to make a direct parallel with the situation in which learning is not feasible, for which we report the optimal upper-censorship threshold a^{cens} and the value k^{cens} at which learning becomes unfeasible.

Two facts need to be highlighted. First, in all the simulations, as k becomes large enough, the problem collapses to the standard censorship, and the optimal threshold is the same as in the case in which receivers cannot learn. Such a switching point depends on the mode of the distribution and, unsurprisingly, tends to be lower as γ grows, and the distribution shifts in a first-order stochastically dominant way: when the outside options become more competitive, lower information-processing costs are needed to justify the effort to acquire more information.

Second, when the outside option is concentrated towards lower values, e.g., $\gamma = 0.4$, lower information-processing costs lead to a decrease in complex disclosure. Instead, this is not the case when the outside option is subject to a first-order stochastic dominance shift, e.g., $\gamma = 0.8$

or $\gamma = 0.9$, where we obtain higher complexity mass for small k , compared to the case of information censorship.²⁷ After having established such a possibility, we are left to discuss why such non-monotonicity occurs, and what features of the receiver's outside-option distribution interact with the information-processing cost in a way that triggers more information complexity when learning becomes easier. The non-monotonicity in Figure 5 comes from the way costly learning changes the marginal cost and benefit of complex disclosure.

To better understand the mechanism, fix any upper censorship threshold a and let $\hat{x}(a) = \frac{1+a}{2}$ denote the expected quality value given that the complex message is used. In the censorship problem, receivers do not learn. Hence, upon observing the complex message, all receivers with outside options below $\hat{x}(a)$ buy the product, while all receivers with outside options above $\hat{x}(a)$ just decide to keep their outside option. This implies that the sender's payoff in this HCP problem can be written as $\Pi_{HCP} = \int_0^a G(x)dx + (1-a)G(\hat{x}(a))$. Therefore, lowering a , and so increasing the mass of complex disclosure, has two effects:

$$-\frac{d\Pi_{HCP}}{da} = \underbrace{G(\hat{x}) - G(a)}_{\text{Direct Benefit on Sales}} - \underbrace{\frac{1-a}{2}g(\hat{x})}_{\text{Direct Cost}}.$$

First, it increases the likelihood that lower types buy the product with probability one, as they see the complex message m^c more often. This improves the sender's payoff relative to $G(a)$, the sales that the marginal state would have generated if transparently disclosed. Second, lowering a makes the quality behind the complex message less attractive to high-outside-option receivers. This increases the share of receivers who end up keeping the outside option and walking away from the purchase, negatively impacting the sender payoff.²⁸ This second force is especially key when the outside-option distribution places substantial mass on receivers with attractive outside options, who are willing to buy only if the complex message remains sufficiently attractive. In the no-learning benchmark, these receivers cannot investigate the hidden state realizations, so a decrease in the value of the complex message translates directly into walk-away decisions.

When k decreases and learning becomes possible, this whole dynamic changes, as we need to include the learning decisions. This makes the payoff equal to $\Pi_{LCP} = \int_0^a G(x)dx + G(\underline{y}(a))(\underline{y}(a) - a) + \int_{\underline{y}(a)}^{\bar{y}(a)} G(x)dx + G(\bar{y}(a))(1 - \bar{y}(a))$. Starting from a threshold a , and lowering it, there are three distinct effects that need to be considered, which are derived

²⁷Appendix B.2 reports simulations for a unimodal Beta outside option. We find that the non-monotonicity of the complexity in k is still possible, even if at smaller magnitudes.

²⁸These forces coincide with the ones discussed to interpret the conditions of Proposition 4. However, the threshold structure makes the intuition easier to convey.

accounting for the definitions of $\underline{y}$ and $\bar{y}$:

$$-\frac{\partial \Pi_{LCP}}{\partial a} = \underbrace{G(\underline{y}) - G(a)}_{\text{Direct Benefit on Sales}} - \underbrace{g(\underline{y}) \sqrt{2k(1-a)}}_{\text{Direct Cost}} + \underbrace{k(g(\underline{y}) - g(\bar{y}))}_{\text{Benefit/Cost From Info Acquisition}}.$$

First, we record the same type of positive effect discussed before on the probability of making a sale after a complex message. The key difference is that now the relevant threshold is $\underline{y}$ and not $\hat{x}$, making it clear how the possibility of learning directly reduces the gain of complexity from the sender. This is the most intuitive mechanism: when receivers acquire information, complex disclosure no longer affects their purchase decision, since they act as under transparent disclosure, no matter the sender's choice. As in the HCP, lowering the upper censorship threshold also makes the complex message less attractive. The second term captures the reduction in uninformed purchases: adding the lower marginal state to the complex pool reduces the expected quality behind the complex message, and so the threshold $\underline{y}$ below which receivers buy without processing the information. Intuitively, this is the same effect at play in the no-learning case, with the only difference that the marginal uninformed buyer is now the type $\underline{y}$ rather than $\hat{x}$.

All the remaining effects of the change in complex disclosure are captured by the information-acquisition channel.²⁹ Indeed, the marginal payoff change above presents a completely new term compared to the no-learning benchmark, which makes the effect of learning explicit in the sender's problem. Depending on the outside option distribution, the cost of more complexity can be mitigated by the information acquisition channel that is absent in the no-learning case. As already mentioned at the end of Section 5.1, better learning opportunities (e.g., lower k) can reduce the threat of losing higher-value customers after a complex message. While easier information acquisition does not counterbalance the loss of low receiver types after an increase in complexity—which decreases its value through a lower threshold a —it can partially offset the walk-away threat at the top of the outside option distribution. This can convert unconditional losses of higher value receivers into conditional ones, depending on the outcome of the learning. The term $k(g(\underline{y}) - g(\bar{y}))$ captures this learning-margin trade-off: it measures how

²⁹The distinction between the last two terms follows directly from the closed-form thresholds derived in Section 5.1, $\underline{y} = a + \sqrt{2k(1-a)}$ and $\bar{y} = 1 - \sqrt{2k(1-a)}$. A marginal decrease in a moves $\underline{y}$ through two components. First, a one-to-one change with a : the added marginal state reduces the expected quality behind the complex message, directly reducing the share of receivers who always buy. This generates the direct cost term $g(\underline{y})(\underline{y} - a) = g(\underline{y}) \sqrt{2k(1-a)}$, the analogue of the cost $(\hat{x} - a)g(\hat{x}) = \frac{1-a}{2}g(\hat{x})$ in the censorship benchmark. Second, through the change in $\sqrt{2k(1-a)}$: the complex message becomes more likely and, since the cost k is paid only when such a message realizes, the receivers' expected information-processing cost rises, making learning less attractive at both margins. Since $\bar{y}$ depends on a only through this second component, these effects are collected in the term $k(g(\underline{y}) - g(\bar{y}))$. The balance between these changes determines the final effect on the sender's payoff.

the endogenous movements of the lower and upper learning thresholds affect the sender's pay-off, and therefore whether information acquisition mitigates or reinforces the marginal cost of complexity. In the unimodal case, this term is strictly positive at the optimum, and the channel it captures is especially key in markets in which outside options are concentrated at high values.³⁰ Intuitively, in such markets, the receivers who discipline the use of complexity by unconditionally keeping the outside option are exactly the ones that learning can convert into potential buyers: when the information-processing cost is low, only a small mass of receivers keeps relying on the available alternatives without processing the complex message, possibly reducing the sender's losses from more complex disclosure.³¹

Such a mechanism explains why the realized comparative statics are affected by the shape of G . When outside options are concentrated at low or intermediate values, the main effect of receiver learning is to reduce uninformed purchase after a complex message: some receivers who would have bought the product now learn and reject low realizations. In this case, learning mostly reduces the gain from complex disclosure and pushes the optimal threshold upward (as shown in Figure 5 for $\gamma = 0.4$). By contrast, when outside options are concentrated at high values, the no-learning censorship threshold is mainly disciplined by the large mass of receivers who may walk away after a complex message. If, instead, information-processing costs are low, a part of these receivers start acquiring information rather than simply rejecting the purchase. This decreases the marginal cost of withholding worse information, because a portion of the high types still buy when the realized quality is high enough. Such a reduction in cost can dominate when the no-learning walk-away cost was large enough to discipline complexity, but low information-processing costs induce enough high-outside-option receivers to learn rather than reject the purchase immediately. In these cases, in response to such learning opportunities, the sender may lower a and use complex disclosure more often, as documented in our numerical simulations.

³⁰The comparison with the censorship benchmark clarifies the role of this channel in high-mode markets. Without learning, the mass of receivers who walk away after a complex message is $1 - G(\hat{x})$, which is large when outside options are concentrated at high values. With learning, the only receivers who unconditionally walk away are the ones above $\bar{y} = 1 - \sqrt{2k(1-a)}$, a mass that vanishes as k decreases.

³¹Appendix Figure 9 reports the equilibrium disclosure and learning thresholds for the triangular simulations, illustrating how the learning region changes with k and how the disclosure threshold approaches the censorship benchmark as learning disappears.

5.2 Segmented Markets: U-Shaped Outside Options

The unimodal benchmark above describes markets in which consumer outside options are concentrated around a single typical value. A natural alternative case is a segmented market in which outside options are concentrated near both extremes of the distribution: some consumers have very low outside options and would buy almost anything instead, others have very high outside options and would only buy if the product is of very high quality. We capture this case by letting G have a U-shaped density: g is strictly decreasing on $[0, \gamma)$ and strictly increasing on $(\gamma, 1]$, with $\gamma \in (0, 1)$. As the main application, we can consider the Beta with support $(0, 1)$ and both shape parameters below one. We can perform the same analysis as for the unimodal distribution and discuss the structure of the optimal complex disclosure rule in the different formulations of the sender's problem. We show that the different shape of the outside option completely changes the incentives of the sender, leading to a lower-censorship rule in both the HCP and LCP (Figure 6 provides a graphical representation).³²

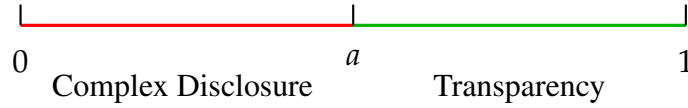


Figure 6: Lower-censorship disclosure rule

Theorem 2. *Assume that G has a U-shaped density g , which reaches the minimum point at γ . In the sender's HCP and LCP, any optimal disclosure rule follows a lower-censorship structure, that is, there exists $a \in [\gamma, 1]$ such that*

$$c(x) = \begin{cases} 1 & \text{if } x < a \\ 0 & \text{if } x \geq a \end{cases}$$

Hence, the sender's maximization collapses to a one-dimensional one.

Theorem 2 shows that, differently from the unimodal case, when outside options tend to be polarized at the extreme values, complexity becomes a feature of low-quality products. It is easy to see how this changes the learning pattern, generating different solutions for equations (2) and (3), which take the form of

$$\underline{y}(a, k) = \sqrt{2ka}, \quad \bar{y}(a, k) = a - \sqrt{2ka}.$$

³²Ginzburg (2019) provides the optimal structure of the complex disclosure rule also with bimodal distributions, which can be extended to this U-shaped case. We report the analysis of the HCP to specifically discuss the case with a U-shaped distribution.

with $\underline{y} < \bar{y}$ any time $a > 8k$. There are a few things worth highlighting. First, the main intuition behind the different structure of the disclosure rule is straightforward. When the distribution is U-shaped, a large mass of receivers is located towards the lowest and the highest points of the domain. Those are also the groups that, when observing a complex message m^c , either decide to buy the product without extra information or decide to walk away and keep their outside option. The first group can be obtained even with a low expected value of complex disclosure. Instead, the second group is extremely hard to reach, no matter the value of $\hat{x}$. Changing the direction of the censorship choice guarantees the sales from the low part of the distribution, but allows for the possibility that higher-value receivers are reached when good realizations of x occur. A similar type of reasoning applies to the HCP, where the structure of the complex disclosure rule is the same.

To conclude the discussion of U-shaped outside-option distributions, we need to describe the structure of the solution when the receiver's learning constraint binds. The following lemma summarizes the resulting threshold structure.

Lemma 3. *Assume that G has a U-shaped density g . The optimal disclosure rule in the sender's Binding Problem can be described with five (possibly degenerate) thresholds $0 \leq a_1 \leq a_2 \leq a_3 \leq a_4 \leq a_5 \leq 1$, such that*

$$c(x) = \begin{cases} 1 & \text{if } x < a_1 \text{ or } a_2 \leq x \leq a_3 \text{ or } a_4 \leq x \leq a_5 \\ 0 & \text{if } a_1 \leq x \leq a_2 \text{ or } a_3 \leq x \leq a_4 \text{ or } x > a_5 \end{cases}$$

5.2.1 U-Shaped Outside Option: Numerical Results

We repeat the same numerical analysis as in the unimodal case for U-shaped outside option distributions. The results are reported in Figure 7, where the distributions $Beta(0.7, 0.3)$ and $Beta(0.5, 0.5)$ first-order stochastically dominate the distribution $Beta(0.3, 0.7)$. In this case, the optimal threshold structure is reversed relative to the unimodal case. Complex disclosure is concentrated on the lower side of the quality distribution, so the mass of complex disclosure is increasing in the threshold rather than decreasing in it. Therefore, in reading the figure, more complexity corresponds to a higher lower-censorship threshold.

Despite this reversal in the threshold structure, the economic mechanism behind the result is consistent with the one described in the unimodal case. More specifically, expanding the complex disclosure region adds marginally higher states to the previous ones, and it increases the average value of x upon a complex message m^c , $\hat{x}$. This has two effects. First, a positive one, which captures the increase in value behind the complex message and the consequent

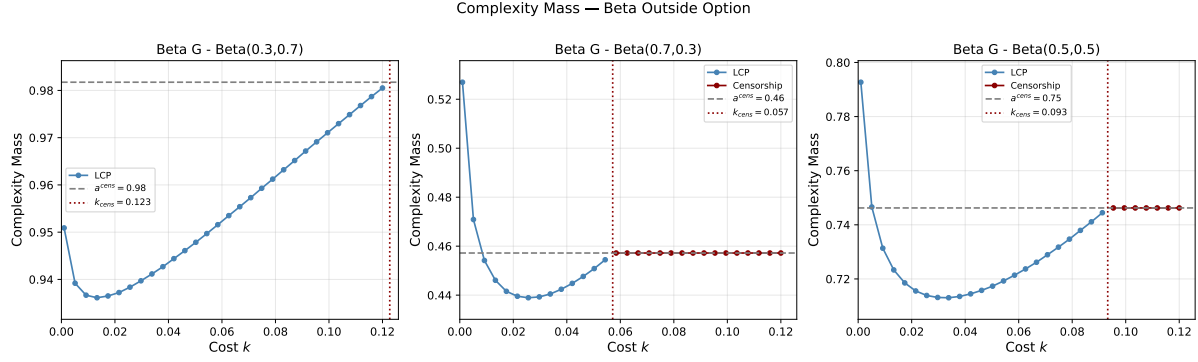


Figure 7: Probability of complex disclosure a as a function of k , when G is a U-shaped Beta distribution and $x \sim \mathcal{U}[0, 1]$

marginal growth in the mass of receivers who would make the purchase when the information is complex. Second, a negative one, which captures the direct loss in sales due to the fact that a decrease in transparency triggers a higher probability of no purchase among the higher-outside-option types, who walk away after a complex message. Similarly to before, when learning becomes possible, the negative effect of more complexity can be partially mitigated, since for low enough k , a positive mass of receiver types would acquire information and buy whenever the realized quality is high enough.³³ As a result, consistent with the previously stated Result 1 from Section 5.1.1, low information-processing costs can increase the sender's incentive to use complex disclosure even in the U-shaped case. This occurs for the Beta distributions that place more mass on high outside options (the $Beta(0.7, 0.3)$ and $Beta(0.5, 0.5)$ cases in Figure 7), while for the $Beta(0.3, 0.7)$ case, in which most receivers have weak alternatives, lower costs reduce complexity, mirroring the role of the mode in the unimodal case.

6 Regulatory Implications

In this section, we discuss what our results imply for the relationship between market structure and strategic complexity, for receivers' welfare, and for the design of information policies.

³³Formally, with the reversed rule the complex pool is $[0, a]$, the HCP and LCP payoff would look like $\Pi_{HCP} = aG(\hat{x}(a)) + \int_a^1 G(x)dx$, with $\hat{x}(a) = \frac{a}{2}$, and $\Pi_{LCP} = G(\underline{y}(a))\underline{y}(a) + \int_{\underline{y}(a)}^{\bar{y}(a)} G(x)dx + G(\bar{y}(a))(a - \bar{y}(a)) + \int_a^1 G(x)dx$. This implies that $\frac{d\Pi_{HCP}}{da} = \underbrace{-(G(a) - G(\hat{x}))}_{\text{Direct Cost on Sales}} + \underbrace{\frac{a}{2}g(\hat{x})}_{\text{Direct Benefit}}$ and $\frac{\partial \Pi_{LCP}}{\partial a} = \underbrace{-(G(a) - G(\bar{y}))}_{\text{Direct Cost on Sales}} + \underbrace{g(\bar{y})\sqrt{2ka}}_{\text{Direct Benefit}} + \underbrace{k(g(\underline{y}) - g(\bar{y}))}_{\text{Benefit/Cost From Info Acquisition}}$. Note that the direct cost and benefit now depend on the upper learning threshold $\bar{y}$, since the marginal state enters the complex pool from the top.

How the market structure affects the choice of strategic complexity. Given the analysis in Section 5, we can highlight a few policy relevant facts. First, we showed that the sender tends to transparently disclose negative realizations when the outside option is unimodal, i.e., when receivers are homogeneous in their product evaluation, and, instead, to transparently disclose positive realizations when the outside option is U-shaped, i.e., when the market tends to be more segmented. This implies that, depending on the competitive structure of the market, we can expect different uses of information complexity, aiming to satisfy different strategic objectives of the senders. In practice, this means that complexity is not only a way to hide unfavorable information. It can also be used to affect the skepticism of receivers, associating a positive meaning to complex messages. This possibility can affect the evaluation of such a strategic tool depending on the characteristics of the market.

Second, our numerical simulations show that, for some specifications of outside-option distributions, decreasing the information-processing cost k can have non-monotonic effects on the probability of complex disclosure. In our simulations, this occurs when outside-option distributions place substantial mass on higher values of the domain. The reason is that, when consumers are more willing to process complex information, the sender may find it profitable to use complexity more often, even when market alternatives are good, as it no longer induces as many sales losses from the unconditional walk-away of higher types. This implies that policies aiming to make the information more accessible to consumers may have the unintended effect of increasing the amount of complex information firms produce.

Effect of strategic complexity on receivers' welfare. The ambiguous effect of the cost k on the amount of complex disclosure also translates into possible ambiguous effects on the receiver's welfare. We can focus attention on the unimodal distribution case and define the receiver's loss due to complexity as the difference between her average transparency payoff—i.e., the expected payoff accounting for the distribution of possible types when each state is revealed transparently—and the average payoff under the equilibrium complex disclosure rule. To better interpret this measure, we can then normalize it using the gap between the receiver's average transparency payoff and the receiver's payoff when no information about x is revealed, which is equivalent to the payoff when every $x \in [0, 1]$ is censored. This loss measures how far receivers are from the benchmark, in which all product qualities are transparently revealed. Hence, we derive the following expression for the welfare loss function

$$L(a, k) = \begin{cases} \frac{1}{2} \int_a^{\underline{y}} (y - a)^2 g(y) dy + \frac{1}{2} \int_{\bar{y}}^1 (1 - y)^2 g(y) dy + (1 - a)k \left(G(\bar{y}) - G(\underline{y}) \right) & \text{if } \underline{y} < \bar{y} \\ \frac{1}{2} \int_a^{\hat{x}} (y - a)^2 g(y) dy + \frac{1}{2} \int_{\hat{x}}^1 (1 - y)^2 g(y) dy & \text{otherwise} \end{cases}$$

while the normalization parameter is equal to $\frac{1}{2} \int_0^\mu y^2 g(y) dy + \frac{1}{2} \int_\mu^1 (1-y)^2 g(y) dy$, with μ representing the unconditional average quality $\mathbb{E}[x]$. This term is useful because it captures the maximal growth in the expected receiver's payoff, aggregating over the whole outside option distribution, that can be achieved through information disclosure. In particular, in both cases in the function, the first term represents the loss from suboptimal uninformed purchases, while the second one captures the loss from uninformed walk-away. When learning occurs, there is a third term that captures the aggregate information-processing costs paid by receivers who decide to learn. The lower the value of $L(a, k)$, the higher the surplus for the receivers: $L(a, k) \geq 0$, and the highest possible surplus is reached when $L(a, k) = 0$, in the full information case.

Result 2. For any k , $L(a, k)$ is decreasing in a .

This result is very intuitive: the more information is transparently conveyed to the receivers, the higher the final surplus, since the payoff gets closer to the full information. We do not have the same clear comparative statics when we assess a change in k , which, in equilibrium, also triggers a possible non-monotone change in the optimal threshold a .

Result 3. An increase in k has an ambiguous effect on $L(a, k)$.

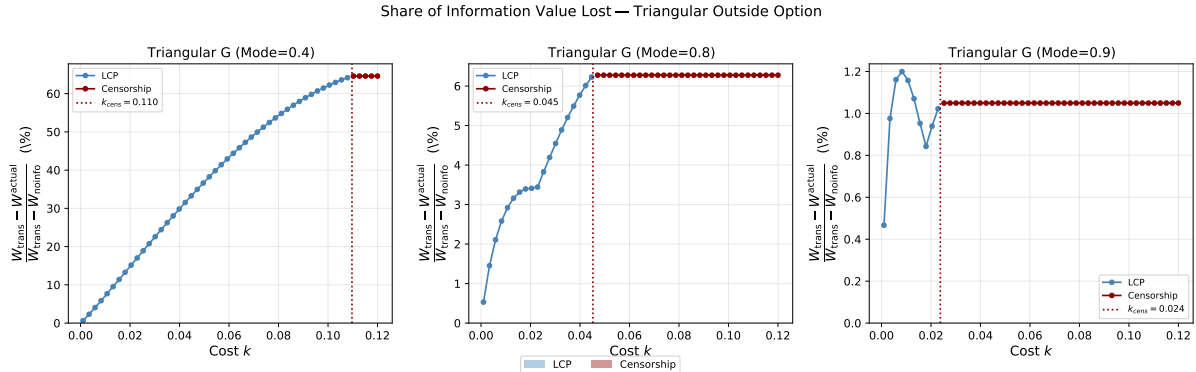


Figure 8: Receiver's normalized loss for different values of k , when G is a triangular distribution with mode γ

This third result suggests that an increase in the information-processing cost might not reduce the welfare of the receivers. Intuitively, a growth in k increases the cost paid by the types who decide to learn, but it also shrinks the mass of the receivers acquiring the information $[y, \bar{y}]$. In addition to this immediate effect, our simulations show that in some cases, a growth in k might lead the sender to disclose information more transparently, making information more readily available. This might have an overall positive effect on the surplus of the receivers, even if the cost to pay to process information is higher.

Figure 8 uses our simulations to illustrate this point in the triangular outside option setting.

The first panel, with $\gamma = 0.4$, shows how the receiver’s welfare is monotonically decreasing in k : higher information-processing costs always negatively impact the receiver’s welfare, inducing a growth in the loss. When the mode grows to $\gamma = 0.8$, a higher k still induces larger losses, but the curve tends to be much flatter around $k = 0.02$. From Figure 5 it is clear how that area coincides with one in which increases in k also decrease complexity, mitigating the negative impact of higher information-processing costs. Finally, when $\gamma = 0.9$, we find evidence that welfare losses may be strictly non-monotone in k , and can exceed the no-learning benchmark loss. Again, such a pattern coincides with the non-monotone complexity areas in Figure 5.

It is important to highlight a few caveats of this analysis. First, in our welfare simulations, the magnitude of the non-monotonicity is small.³⁴ For this reason, it is better to frame the result in terms of a lack of welfare improvements from reducing information-processing costs, rather than in terms of large harms to receivers. Second, the same type of non-monotonicity does not appear in the Beta distributions we study, either in the unimodal or in the U-shaped case (see Online Appendix H). In these cases, an increase in k always reduces the receiver’s welfare. Such a result is likely the consequence of weaker changes in the complex disclosure rule after changes in k (see Figures 7 and 10).

Overall, our analysis shows that, depending on the structure of the receiver’s outside option, and so of the market, changes in k can have unintended consequences in terms of information transparency and, even if in smaller magnitudes, on the receiver’s welfare. This suggests the need to carefully consider these mechanisms before designing regulatory interventions to improve information transmission.

7 Conclusion

This paper studies possible limitations of mandated disclosure that can arise when firms retain control over how product information is disclosed to consumers. This allows firms to strategically present information in formats that are costly for consumers to understand. We investigate whether a more competitive structure of the market, modeled as a more favorable distribution of outside alternatives for the receivers, disciplines or instead strengthens firms’ incentives to use complex disclosure, and how this answer depends on the consumers’ information-processing costs. Consumers are assumed to be fully rational and can decide whether to pay a cost to understand the information that is made complex by the firm. In such an environment, disclosure format affects not only consumers’ beliefs about product quality but also their in-

³⁴This is especially clear in Figure H.1, which reports the non-normalized welfare losses in Online Appendix H.

centives to learn at a cost. This creates a strategic channel absent from standard disclosure and censorship models.

We show that more favorable market alternatives—captured by outside-option distributions that place substantial mass on high values—and lower information-processing costs may fail to discipline the sender’s choices and even lead to less transparent information disclosure. The main reason for this result is that, when information is accessible, the strategic use of complexity may be less effective but also less costly, as consumers with good outside options may decide to gather information about the sender’s product rather than just keep the outside option when presented with complex information. As a result, making disclosures easier to understand can improve learning conditional on a complex message, but also makes complex messages more attractive for the firm to use. The same force that improves consumer information acquisition can therefore decrease the equilibrium amount of transparent disclosures.

We also show that the structure of consumers’ outside options determines how complexity is used. In homogeneous markets, where outside options are concentrated around a single value, the sender finds it optimal to disclose worse information more transparently. Instead, in segmented markets, where a large share of consumers have access to either very weak or very strong alternatives, complexity is used to hide lower-quality states while high-quality states are disclosed transparently. Complexity is therefore not simply a symptom of low quality, but a tool the sender can use to design information depending on the structure of the market.

These results can have implications for disclosure regulation. In particular, when firms retain discretion over disclosure format, mandated disclosure policies may be ineffective, and interventions aiming to help consumers understand complex disclosures may increase information acquisition but also make the information disclosed less transparent. The effectiveness of such policies depends on the competitive structure of the market. This suggests that empirical evaluations should account for heterogeneity across markets both in terms of competitive structure and in terms of information accessibility before designing regulatory interventions.

References

- Aghamolla, Cyrus, and Kevin Smith.** 2023. “Strategic complexity in disclosure.” *Journal of Accounting and Economics*, 76(2-3): 101635.
- Asriyan, Vladimir, Dana Foarta, and Victoria Vanasco.** 2023. “The Good, the bad, and the complex: product design with imperfect information.” *American Economic Journal*:

- Microeconomics*, 15(2): 187–226.
- Best, James, and Daniel Quigley.** 2024. “Persuasion for the long run.” *Journal of Political Economy*, 132(5): 1740–1791.
- Bizzotto, Jacopo, Jesper Rüdiger, and Adrien Vigier.** 2020. “Testing, disclosure and approval.” *Journal of Economic Theory*, 187: 105002.
- Bloedel, Alexander W, and Ilya Segal.** 2021. “Persuading a rationally inattentive agent.” *Working Paper*.
- Board, Simon, and Jay Lu.** 2018. “Competitive information disclosure in search markets.” *Journal of Political Economy*, 126(5): 1965–2010.
- Börgers, Tilman.** 2015. *An Introduction to the Theory of Mechanism Design*. Oxford University Press.
- Candogan, Ozan, and Philipp Strack.** 2023. “Optimal disclosure of information to privately informed agents.” *Theoretical Economics*, 18(3): 1225–1269.
- Carlin, Bruce I.** 2009. “Strategic price complexity in retail financial markets.” *Journal of Financial Economics*, 91(3): 278–287.
- Carlin, Bruce Ian, and Gustavo Manso.** 2011. “Obfuscation, learning, and the evolution of investor sophistication.” *The Review of Financial Studies*, 24(3): 754–785.
- Chioveanu, Ioana.** 2019. “Prominence, complexity, and pricing.” *International Journal of Industrial Organization*, 63: 551–582.
- Chioveanu, Ioana.** 2020. “A more general model of price complexity.” *International Journal of Industrial Organization*, 69: 102563.
- Chioveanu, Ioana, and Jidong Zhou.** 2013. “Price competition with consumer confusion.” *Management Science*, 59(11): 2450–2469.
- Dall’Ara, Pietro.** 2024. “Persuading an inattentive and privately informed receiver.” *arXiv preprint arXiv:2408.01250*.
- de Clippel, Geoffroy, and Kareen Rozen.** 2025. “Communication, perception, and strategic obfuscation.” *Working Paper*.

- DeHaan, Ed, Yang Song, Chloe Xie, and Christina Zhu.** 2021. “Obfuscation in mutual funds.” *Journal of Accounting and Economics*, 72(2-3): 101429.
- Dye, Ronald A.** 1985. “Disclosure of nonproprietary information.” *Journal of accounting research*, 123–145.
- Ellison, Glenn.** 2005. “A model of add-on pricing.” *The Quarterly Journal of Economics*, 120(2): 585–637.
- Ellison, Glenn, and Alexander Wolitzky.** 2012. “A search cost model of obfuscation.” *The RAND Journal of Economics*, 43(3): 417–441.
- Ellison, Glenn, and Sara Fisher Ellison.** 2009. “Search, obfuscation, and price elasticities on the internet.” *Econometrica*, 77(2): 427–452.
- Fréchette, Guillaume R, Alessandro Lizzeri, and Jacopo Perego.** 2022. “Rules and commitment in communication: An experimental analysis.” *Econometrica*, 90(5): 2283–2318.
- Gabaix, Xavier, and David Laibson.** 2006. “Shrouded attributes, consumer myopia, and information suppression in competitive markets.” *The Quarterly Journal of Economics*, 121(2): 505–540.
- Ginzburg, Boris.** 2019. “Optimal information censorship.” *Journal of Economic Behavior & Organization*, 163: 377–385.
- Grossman, Sanford J.** 1981. “The informational role of warranties and private disclosure about product quality.” *The Journal of Law and Economics*, 24(3): 461–483.
- Gu, Yiquan, and Tobias Wenzel.** 2014. “Strategic obfuscation and consumer protection policy.” *The Journal of Industrial Economics*, 62(4): 632–660.
- Janssen, Aljoscha, and Johannes Kasinger.** 2024. “Obfuscation and rational inattention.” *The Journal of Industrial Economics*, 72(1): 390–428.
- Jin, Ginger Zhe, Michael Luca, and Daniel Martin.** 2021. “Is no news (perceived as) bad news? An experimental investigation of information disclosure.” *American Economic Journal: Microeconomics*, 13(2): 141–173.
- Jin, Ginger Zhe, Michael Luca, and Daniel Martin.** 2022. “Complex disclosure.” *Management Science*, 68(5): 3236–3261.

- Jovanovic, Boyan.** 1982. “Truthful disclosure of information.” *The Bell Journal of Economics*, 36–44.
- Kamenica, Emir.** 2019. “Bayesian persuasion and information design.” *Annual Review of Economics*, 11: 249–272.
- Kamenica, Emir, and Matthew Gentzkow.** 2011. “Bayesian persuasion.” *American Economic Review*, 101(6): 2590–2615.
- Kamenica, Emir, and Matthew Gentzkow.** 2017. “Competition in persuasion.” *Review of economic studies*, 84(1): 300–322.
- Kleiner, Andreas, Benny Moldovanu, and Philipp Strack.** 2021. “Extreme Points and Majorization: Economic Applications.” *Econometrica*, 89(4): 1557–1593.
- Kolotilin, Anton.** 2018. “Optimal information disclosure: A linear programming approach.” *Theoretical Economics*, 13(2): 607–635.
- Kolotilin, Anton, Tymofiy Mylovanov, and Andriy Zapechelnjuk.** 2022. “Censorship as optimal persuasion.” *Theoretical Economics*, 17(2): 561–585.
- Kolotilin, Anton, Tymofiy Mylovanov, Andriy Zapechelnjuk, and Ming Li.** 2017. “Persuasion of a privately informed receiver.” *Econometrica*, 85(6): 1949–1964.
- Lyu, Chen.** 2023. “Information design for selling search goods and the effect of competition.” *Journal of Economic Theory*, 213: 105722.
- Martin, Daniel.** 2017. “Strategic pricing with rational inattention to quality.” *Games and Economic Behavior*, 104: 131–145.
- Mathevet, Laurent, David Pearce, and Ennio Stacchetti.** 2025. “Reputation and information design.” *Working Paper*.
- Matyskova, Ludmila, and Alfonso Montes.** 2023. “Bayesian persuasion with costly information acquisition.” *Journal of Economic Theory*, 211: 105678.
- Milgrom, Paul R.** 1981. “Good news and bad news: Representation theorems and applications.” *The Bell Journal of Economics*, 380–391.
- Monte, Daniel, and Luis Henrique Linhares.** 2023. “Strategic obfuscation: information costs as a tool to influence consumers.” *Working Paper*.

- Onuchic, Paula.** 2025. “Advisors with hidden motives.” *Games and Economic Behavior*.
- Perez-Richet, Eduardo, and Delphine Prady.** 2011. “Complicating to persuade.” *Working Paper*.
- Piccione, Michele, and Ran Spiegler.** 2012. “Price competition under limited comparability.” *The Quarterly Journal of Economics*, 127(1): 97–135.
- Rayo, Luis, and Ilya Segal.** 2010. “Optimal information disclosure.” *Journal of Political Economy*, 118(5): 949–987.
- Spiegler, Ran.** 2016. “Choice complexity and market competition.” *Annual Review of Economics*, 8: 1–25.
- Wei, Dong.** 2021. “Persuasion under costly learning.” *Journal of Mathematical Economics*, 94: 102451.
- Wilson, Chris M.** 2010. “Ordered search and equilibrium obfuscation.” *International Journal of Industrial Organization*, 28(5): 496–506.
- Xu, Shuo.** 2025. “Persuasion through information cost design.” *Economic Modelling*, 107348.

Appendix

A Proofs

A.1 Proof of Proposition 1

Proposition 1 is a direct consequence of Lemma D.4 and Lemma D.5, which, for reasons of space, are proved in Online Appendix D.

If Part. We want to prove that if $V^i(\hat{x}) > V^{no}(\hat{x})$, then there exist $\underline{y}, \bar{y} \in [\underline{x}, 1]$ with $\underline{y} < \bar{y}$ such that

$$V(y) = \begin{cases} V^{no}(y) & \text{if } y < \underline{y} \text{ or } y > \bar{y} \\ V^i(y) & \text{if } \underline{y} \leq y \leq \bar{y} \end{cases}$$

Lemma D.4 guarantees that there are some $y \in [\underline{x}, 1]$ such that $V(y) = V^i(y)$. Lemma D.5 instead guarantees that the two functions cross exactly twice, once in $[\underline{x}, \hat{x})$ and the other in $(\hat{x}, 1]$. Let's denote these intersection points respectively as $\underline{y}$ and $\bar{y}$. Since we know that $V^i(0) < V^{no}(0)$ and $V^i(1) < V^{no}(1)$, it must be

- $V^i(y) \leq V^{no}(y)$ for $y \in [0, \underline{y})$ and $y \in (\bar{y}, 1]$
- $V^i(y) \geq V^{no}(y)$ for $y \in [\underline{y}, \bar{y}]$

From these relations, we get the definition in Equation (1). We are left to prove that when $V^i(\hat{x}) = V^{no}(\hat{x})$, $\underline{y} = \bar{y} = \hat{x}$ and $V(y) = V^{no}(y)$. Lemma D.4 and Lemma D.5 guarantee that the two value functions only cross once, at $\hat{x}$, and that $V^i(y) \leq V^{no}(y)$ for all $y \in [0, 1]$. This guarantees that $\underline{y} = \bar{y} = \hat{x}$ and that $V(y) = V^{no}(y)$.

Finally, the no-learning characterization follows from the same arguments: when $V^i(\hat{x}) < V^{no}(\hat{x})$, Lemma D.4 guarantees that there is no y for which $V(y) = V^i(y)$, so that $V(y) = V^{no}(y)$ for every $y \in [0, 1]$, while the converse is immediate from the definition of V .

Only If Part. We want to show that if Equation (1) holds for a pair $\underline{y} < \bar{y}$, then $V^i(\hat{x}) > V^{no}(\hat{x})$. The only if part of Lemma D.4 guarantees that $V^i(\hat{x}) \geq V^{no}(\hat{x})$. However, the fact that the two value functions intersect twice and the proof of Lemma D.5 guarantee directly that $V^i(\hat{x}) > V^{no}(\hat{x})$. We want to show that $\underline{y} = \bar{y} = \hat{x}$ and $V(y) = V^{no}(y)$ for every $y \in [0, 1]$ implies that $V^i(\hat{x}) = V^{no}(\hat{x})$. This trivially follows from the definition of Equation (1).

We are left to prove that the condition $V^i(\hat{x}) \geq V^{no}(\hat{x})$ can be expressed as

$$k \leq \frac{\int_{\underline{x}}^{\hat{x}} (\hat{x} - x)c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} = \frac{\int_{\hat{x}}^1 (x - \hat{x})c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx}.$$

This comes directly from the definitions of the two value functions. Indeed, $V^i(\hat{x}) \geq V^{no}(\hat{x})$ implies $\hat{x} \frac{\int_{\underline{x}}^{\hat{x}} c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} + \frac{\int_{\hat{x}}^1 xc(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} - \hat{x} \geq k$. Notice that $\hat{x} = \hat{x} \frac{\int_{\underline{x}}^{\hat{x}} c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx}$. Just plugging the relation inside the expression above and rearranging some terms, we get that $k \leq \frac{\int_{\hat{x}}^1 (x - \hat{x})c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx}$.

If instead, rather than plugging $\hat{x} = \hat{x} \frac{\int_{\underline{x}}^{\hat{x}} c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx}$, we exploit the definition of $\hat{x}$ and we plug $\hat{x} = \frac{\int_{\underline{x}}^1 xc(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx}$, rearranging our initial expression we get $k \leq \frac{\int_{\underline{x}}^{\hat{x}} (\hat{x} - x)c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx}$. Since we are just using different formulations for $\hat{x}$ to manipulate the same expression,

$$\frac{\int_{\underline{x}}^{\hat{x}} (\hat{x} - x)c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} = \frac{\int_{\hat{x}}^1 (x - \hat{x})c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx}.$$

This concludes the proof of Proposition 1.

A.2 Proof of Proposition 2

Take as given the complex disclosure rule $c : [\underline{x}, 1] \rightarrow [0, 1]$. Let's start from Equation (2), which defines $\underline{y}$. We can rewrite the expression as

$$\frac{\int_{\underline{x}}^{\underline{y}} (\underline{y} - x)c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} = k.$$

We first show that the left-hand side is weakly increasing in $\underline{y}$, and strictly increasing at the relevant positive level. Take $\underline{x} < y_1 < y_2 < \hat{x}$. Then

$$\begin{aligned} & \int_{\underline{x}}^{y_2} (y_2 - x)c(x)f(x)dx - \int_{\underline{x}}^{y_1} (y_1 - x)c(x)f(x)dx \\ &= (y_2 - y_1) \int_{\underline{x}}^{y_1} c(x)f(x)dx + \int_{y_1}^{y_2} (y_2 - x)c(x)f(x)dx \geq 0. \end{aligned}$$

Hence the left-hand side of Equation (2) is weakly increasing in $\underline{y}$. Moreover, if y_1 solves Equation (2) for some $k > 0$, $\int_{\underline{x}}^{y_1} (y_1 - x)c(x)f(x)dx = k \int_{\underline{x}}^1 c(x)f(x)dx > 0$. Since the integrand is nonnegative, this implies that $\int_{\underline{x}}^{y_1} c(x)f(x)dx > 0$. Therefore, for any $y_2 > y_1$,

$$\int_{\underline{x}}^{y_2} (y_2 - x)c(x)f(x)dx - \int_{\underline{x}}^{y_1} (y_1 - x)c(x)f(x)dx$$

$$= (y_2 - y_1) \int_{\underline{x}}^{y_1} c(x)f(x)dx + \int_{y_1}^{y_2} (y_2 - x)c(x)f(x)dx > 0.$$

Hence, the left-hand side of Equation (2) is strictly increasing in $\underline{y}$. It follows that an increase in k must determine an increase in the threshold $\underline{y}$ to balance the two sides of Equation (2).

Similarly, we can prove that $\bar{y}$ is decreasing in k . Let's consider Equation (3), which can be rewritten as $\frac{\int_{\bar{y}}^1 (x - \bar{y})c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} = k$. We first show that the left-hand side is weakly decreasing in $\bar{y}$, and strictly decreasing at the relevant positive level. Take $y_1 < y_2$. Then

$$\begin{aligned} & \int_{y_1}^1 (x - y_1)c(x)f(x)dx - \int_{y_2}^1 (x - y_2)c(x)f(x)dx \\ &= \int_{y_1}^{y_2} (x - y_1)c(x)f(x)dx + (y_2 - y_1) \int_{y_2}^1 c(x)f(x)dx \geq 0. \end{aligned}$$

Hence, the left-hand side of Equation (3) is weakly decreasing in $\bar{y}$. Moreover, if y_2 solves Equation (3) for some $k > 0$, then

$$\int_{y_2}^1 (x - y_2)c(x)f(x)dx = k \int_{\underline{x}}^1 c(x)f(x)dx > 0.$$

Since the integrand is nonnegative, this implies that $\int_{y_2}^1 c(x)f(x)dx > 0$. Therefore, for any $y_1 < y_2$,

$$\begin{aligned} & \int_{y_1}^1 (x - y_1)c(x)f(x)dx - \int_{y_2}^1 (x - y_2)c(x)f(x)dx \\ &= \int_{y_1}^{y_2} (x - y_1)c(x)f(x)dx + (y_2 - y_1) \int_{y_2}^1 c(x)f(x)dx > 0. \end{aligned}$$

Hence, the left-hand side of Equation (3) is strictly decreasing in $\bar{y}$. It follows that an increase in k must determine a decrease in the threshold $\bar{y}$ to balance the two sides of Equation (3).

A.3 Proof of Proposition 3

In what follows, we ignore the constraints $\underline{y} \leq \bar{y}$ and $\int_{\underline{x}}^1 c(x)f(x)dx > 0$. Indeed, the goal of the Proposition is to derive necessary conditions for the structure of an optimal complex disclosure rule under the assumption that some information is processed in equilibrium. For this reason, we assume that constraint (13) is satisfied with strict inequality.

We are ready to set up our Lagrangian. We can denote by $\lambda_1, \lambda_2 \in \mathbb{R}$ the Lagrange multipliers on constraints (14) and (15). We can write the Lagrangian as:

$$\mathcal{L}(c, \underline{y}, \bar{y}) = \int_{\underline{x}}^{\underline{y}} (G(\underline{y}) - G(x)) c(x)f(x)dx - \int_{\bar{y}}^1 (G(x) - G(\bar{y})) c(x)f(x)dx$$

$$\begin{aligned}
& + \lambda_1 \left(- \int_{\underline{x}}^{\underline{y}} (\underline{y} - x) c(x) f(x) dx + k \int_{\underline{x}}^1 c(x) f(x) dx \right) \\
& + \lambda_2 \left(\int_{\underline{y}}^1 (x - \underline{y}) c(x) f(x) dx - k \int_{\underline{x}}^1 c(x) f(x) dx \right)
\end{aligned}$$

where we have omitted all the terms not depending on c , $\underline{y}$ and $\bar{y}$ and so not relevant for the optimization. This problem consists of finding a function c and two values $\underline{y}$ and $\bar{y}$ which satisfy the constraints with the equality and maximize the objective function. Given the mathematical structure of the problem, we can interpret it as an optimal control problem, in which the control variable is the function c , and the state variable is constant and independent of the control. We can construct an ancillary function, the control Hamiltonian $L(c(x), \underline{y}, \bar{y})$, which allows us to solve a continuum of static problems with choice variable $c(x)$.³⁵ We can define the control Hamiltonian as

$$\begin{aligned}
L(c(x), \underline{y}, \bar{y}) = & \mathbb{I} \left(x \in [\underline{x}, \underline{y}] \right) \left(G(\underline{y}) - G(x) - \lambda_1 \underline{y} + \lambda_1 x \right) c(x) f(x) \\
& - \mathbb{I} \left(x \in [\bar{y}, 1] \right) \left(G(x) - G(\bar{y}) - \lambda_2 x + \lambda_2 \bar{y} \right) c(x) f(x) \\
& + (\lambda_1 - \lambda_2) k c(x) f(x)
\end{aligned} \tag{20}$$

Taking the First Order Conditions with respect to $c(x)$ for every $x \in [\underline{x}, 1]$, given $\underline{y}$ and $\bar{y}$, we get the following conditions:

- For all $x \in [\underline{x}, \underline{y}]$, $\frac{\partial L(c(x), \underline{y}, \bar{y})}{\partial c(x)} = f(x) \left[\int_{\underline{x}}^{\underline{y}} (g(z) - \lambda_1) dz + k(\lambda_1 - \lambda_2) \right];$
- For all $x \in [\underline{y}, \bar{y}]$, $\frac{\partial L(c(x), \underline{y}, \bar{y})}{\partial c(x)} = k f(x) (\lambda_1 - \lambda_2)$
- For all $x \in [\bar{y}, 1]$, $\frac{\partial L(c(x), \underline{y}, \bar{y})}{\partial c(x)} = f(x) \left[- \int_{\bar{y}}^x (g(z) - \lambda_2) dz + k(\lambda_1 - \lambda_2) \right]$

However, our problem also consists of choosing $\underline{y}$ and $\bar{y}$ and it is clear how the conditions above depend on these two thresholds. We can get first derivatives with respect to $\underline{y}$ and $\bar{y}$ directly from our Lagrangian:

- $\frac{\partial \mathcal{L}(c, \underline{y}, \bar{y})}{\partial \underline{y}} = \left(g(\underline{y}) - \lambda_1 \right) \int_{\underline{x}}^{\underline{y}} c(x) f(x) dx$
- $\frac{\partial \mathcal{L}(c, \underline{y}, \bar{y})}{\partial \bar{y}} = (g(\bar{y}) - \lambda_2) \int_{\bar{y}}^1 c(x) f(x) dx$

At this point, there are a few things to notice.

³⁵Performing this continuum of optimizations provides an optimal control policy function, which describes the behavior of c for every x .

1. Given the linearity of the problem in $c(x)$, the structure of the complex disclosure rule is deterministic every time the First Order Conditions are different from zero—and whether $c(x) \in \{0, 1\}$ depends on their sign. If, for some x , the FOC is equal to zero, then $c(x) \in [0, 1]$.
2. The first derivatives with respect to $\underline{y}$ and $\bar{y}$ need to be equal to zero in the optimum, since those solutions are internal (we are studying the game in which learning happens). Because $k > 0$ and $\int_{\underline{x}}^1 c(x)f(x)dx > 0$, constraints (14) and (15) imply that there is positive complex-message mass below $\underline{y}$ and above $\bar{y}$. Hence,

$$\int_{\underline{x}}^{\underline{y}} c(x)f(x)dx > 0 \quad \text{and} \quad \int_{\bar{y}}^1 c(x)f(x)dx > 0.$$

This allows us to pin down the value of the two multipliers: $\lambda_1 = g(\underline{y})$ and $\lambda_2 = g(\bar{y})$.

3. Given the structure of our problem, we already know that the constraints will be binding:

- $\int_{\underline{x}}^{\underline{y}} (\underline{y} - x)c(x)f(x)dx - k \int_{\underline{x}}^1 c(x)f(x)dx = 0$
- $\int_{\bar{y}}^1 (x - \bar{y})c(x)f(x)dx - k \int_{\underline{x}}^1 c(x)f(x)dx = 0$

Combining observations 1-3 with the first-order conditions, and noticing that the value of $f(x)$ does not affect the sign of the FOCs for $c(x)$, we can directly write the necessary conditions (1)-(5). This concludes the proof of Proposition 3.

A.4 Proof of Proposition 4

Let's assume that the no-learning constraint (18) is satisfied and ignore it in deriving the necessary condition. In addition, let's assume a non-degenerate complex disclosure rule, with $\int_{\underline{x}}^1 c(x)f(x)dx > 0$. Denoting by λ the multiplier on constraint (19), we can write down the Lagrangian of this problem as

$$\begin{aligned} \mathcal{L}(c, \hat{x}) = & - \int_{\underline{x}}^1 G(x)c(x)f(x)dx + G(\hat{x}) \int_{\underline{x}}^1 c(x)f(x)dx \\ & + \lambda \left[\int_{\underline{x}}^1 xc(x)f(x)dx - \hat{x} \int_{\underline{x}}^1 c(x)f(x)dx \right] \end{aligned}$$

where we omit all the terms that are not relevant for the maximization.

At this point, we can notice that we can again apply the control Hamiltonian approach to optimize with respect to c . This allows us to take derivatives of the control Hamiltonian with

respect to $c(x)$. Defining

$$L(c(x), \hat{x}) = -G(x)c(x)f(x) + G(\hat{x})c(x)f(x) + \lambda(x - \hat{x})c(x)f(x)$$

we can derive the following derivatives:

- for every $x \in [\underline{x}, 1]$, $\frac{\partial L(c(x), \hat{x})}{\partial c(x)} = f(x) \int_x^{\hat{x}} (g(z) - \lambda) dz$;
- $\frac{\partial \mathcal{L}(c, \hat{x})}{\partial \hat{x}} = g(\hat{x}) \int_{\underline{x}}^1 c(x)f(x)dx - \lambda \int_{\underline{x}}^1 c(x)f(x)$

Through the same analysis as in Section 4.1—and noticing that $\hat{x} \in (0, 1)$ for every non-degenerate disclosure rule c —we can conclude that the sign of the first derivatives with respect to $c(x)$ will identify the possible corner solutions for the disclosure rule for every $x \in [\underline{x}, 1]$. In addition, the FOC with respect to $\hat{x}$ needs to be satisfied with the equality and pins down the value of the multiplier $\lambda = g(\hat{x})$. Finally, we already know that constraint (19) has to hold with equality, since $\hat{x}$ needs to satisfy its definition. Combining these observations with the FOCs, we derive the conditions reported in the statement of Proposition 4.

A.5 Proof of Lemma 1

Denoting by $\mu \geq 0$ the multiplier on the no-learning constraint (18) and by $\lambda \in \mathbb{R}$ the multiplier on constraint (19) we can write down the Lagrangian of this problem as

$$\begin{aligned} \mathcal{L}(c, \hat{x}) = & - \int_{\underline{x}}^1 G(x)c(x)f(x)dx + G(\hat{x}) \int_{\underline{x}}^1 c(x)f(x)dx \\ & + \lambda \left[\int_{\underline{x}}^1 xc(x)f(x)dx - \hat{x} \int_{\underline{x}}^1 c(x)f(x)dx \right] \\ & + \mu \left[k \int_{\underline{x}}^1 c(x)f(x)dx - \int_{\underline{x}}^{\hat{x}} (\hat{x} - x)c(x)f(x)dx \right] \end{aligned}$$

where we omit all the terms that are not relevant for the maximization.

At this point, we can notice that we can again apply the control Hamiltonian approach to optimize with respect to c . This allows us to take derivatives of the control Hamiltonian with respect to $c(x)$. Defining

$$\begin{aligned} L(c(x), \hat{x}) = & -G(x)c(x)f(x) + G(\hat{x})c(x)f(x) + \lambda(x - \hat{x})c(x)f(x) \\ & + \mu [k \cdot c(x)f(x) - (\hat{x} - x)c(x)f(x)\mathbb{I}[x \leq \hat{x}]] \end{aligned}$$

we can derive the following derivatives:

- for every $x \in [\underline{x}, \hat{x}]$, $\frac{\partial L(c(x), \hat{x})}{\partial c(x)} = f(x) \int_x^{\hat{x}} [g(z) - (\lambda + \mu)] dz + f(x) \mu k$;
- for every $x \in [\hat{x}, 1]$, $\frac{\partial L(c(x), \hat{x})}{\partial c(x)} = -f(x) \int_{\hat{x}}^x (g(z) - \lambda) dz + f(x) \mu k$;
- $\frac{\partial \mathcal{L}(c, \hat{x})}{\partial \hat{x}} = g(\hat{x}) \int_{\underline{x}}^1 c(x) f(x) dx - \lambda \int_{\underline{x}}^1 c(x) f(x) dx - \mu \int_{\underline{x}}^{\hat{x}} c(x) f(x) dx$

Through the same analysis as in Proposition 4, we can conclude that the sign of the first derivatives with respect to $c(x)$ will identify the possible corner solutions for the disclosure rule for every $x \in [\underline{x}, 1]$. Also, we already know that constraints (18) and (19) have to be satisfied with equality. Combining these observations with the FOCs, we derive the conditions reported in the statement of Lemma 1. It is important to highlight that the conditions above allow us to derive a few more facts that will be useful in the following proofs. First, since $g(\hat{x}) > 0$ —we are assuming continuous and non-zero densities—it cannot be that $\lambda = \mu = 0$. Second, if $\mu = 0$, the problem collapses to the HCP problem. This implies that, in order for the binding constraint to matter, it must be that $\mu > 0$. Finally, the derivative of the Lagrangian with respect to $\hat{x}$ allows us to conclude that $g(\hat{x}) < \lambda + \mu$, otherwise the expression could not be equal to zero in the optimum. Finally from the same expression, we can also see that $g(\hat{x}) - \lambda > 0$, otherwise, again, we could not have that the expression equals zero in the optimum.

A.6 Proof of Theorem 1

In order to prove Theorem 1 we need to examine the optimality conditions derived in the general problem. We will just focus on the conditions derived for the LCP. Indeed, this result has already been proved in the context of the HCP, when no receiver is willing to learn (see Ginzburg (2019)).

The proof will be based on three steps:

1. First we can prove that for $x \in (\underline{y}, \bar{y})$, in optimum it must be $\frac{\partial L(x, c)}{\partial c(x)} > 0$;
2. Given step 1, we can show that when $x \in [0, \underline{y}]$, $\frac{\partial L(x, c)}{\partial c(x)}$ is positive in a neighborhood of $\underline{y}$ and crosses zero at most in one point;
3. Given step 1, we can show that when $x \in [\bar{y}, 1]$, $\frac{\partial L(x, c)}{\partial c(x)}$ is always positive.

Step 1: in optimum it must be that $k(g(\underline{y}) - g(\bar{y})) > 0$.

Assume this is not true and $k(g(\underline{y}) - g(\bar{y})) \leq 0$. This means that $g(\underline{y}) \leq g(\bar{y})$. Since $\underline{y} < \bar{y}$ and g is unimodal, this implies that $\underline{y} < \gamma$.

At this point, we need to look at the FOC for $x \leq \underline{y}$, $\int_{\underline{x}}^{\underline{y}} (g(z) - g(\underline{y})) dz + k(g(\underline{y}) - g(\bar{y}))$. If $\underline{y} < \gamma$, then $\int_{\underline{x}}^{\underline{y}} (g(z) - g(\underline{y})) dz < 0$, since the relevant x 's are located on the increasing part of the density function. However, this would mean that both addends in the FOC are non-positive and that there is no $x \leq \underline{y}$ such that $c(x) > 0$. This leads to a contradiction, because under the condition that some complexity is used, there must be some $c(x) > 0$ for $x \leq \underline{y}$ to satisfy constraint (7) in the sender's problem, which defines $\underline{y}$. This means that in optimum it must be that $k(g(\underline{y}) - g(\bar{y})) > 0$. Given the structure of the density, this also implies that $\bar{y} > \gamma$, otherwise it could never be that $g(\underline{y}) > g(\bar{y})$. This step suggests that for all $x \in (\underline{y}, \bar{y})$, $\frac{\partial L(x, c)}{\partial c(x)} > 0$, and so $c(x) = 1$ for all the intermediate qualities.

Step 2: At this point, we need to focus on the first derivative of the Lagrangian for $x \in [0, \underline{y}]$:

$$\frac{\partial L(x, c)}{\partial c(x)} = \int_{\underline{x}}^{\underline{y}} (g(z) - g(\underline{y})) dz + k(g(\underline{y}) - g(\bar{y})).$$

Notice that the second term is strictly positive, as proved in Step 1, and independent of x . The first term can be either positive or negative. However, when $x \rightarrow \underline{y}$, $\int_{\underline{x}}^{\underline{y}} (g(z) - g(\underline{y})) dz \rightarrow 0$. By continuity and by the fact that $k(g(\underline{y}) - g(\bar{y})) > 0$, there must be a neighborhood of $\underline{y}$ in which $\frac{\partial L(x, c)}{\partial c(x)} > 0$.

At this point, there are three cases to consider:

1. Suppose first that $\underline{y} \leq \gamma$. Then g is increasing on $[0, \underline{y}]$, so $g(x) - g(\underline{y}) < 0$ for every $x \in [0, \underline{y}]$. Hence $\frac{\partial}{\partial x} \left[\int_{\underline{x}}^{\underline{y}} (g(z) - g(\underline{y})) dz \right] = g(\underline{y}) - g(x) > 0$. This implies that the first derivative of the Lagrangian is increasing in x . Therefore, either it is positive for every $x \in [0, \underline{y}]$, or it crosses zero at most once.
2. Suppose now that $\underline{y} > \gamma$ and that $g(x) - g(\underline{y}) \geq 0$ for every $x \in [0, \gamma]$. Since g is decreasing on $[\gamma, \underline{y}]$, we also have $g(x) - g(\underline{y}) \geq 0$ for every $x \in [\gamma, \underline{y}]$. Therefore, $\int_{\underline{x}}^{\underline{y}} (g(z) - g(\underline{y})) dz \geq 0$ for every $x \in [0, \underline{y}]$. Since the second term is strictly positive, it follows that $\frac{\partial L(x, c)}{\partial c(x)} > 0$ for every $x \in [0, \underline{y}]$.
3. Finally, suppose that $\underline{y} > \gamma$ and that $g(x) - g(\underline{y}) < 0$ for some $x \in [0, \gamma]$. By unimodality, the function $g(x) - g(\underline{y})$ crosses zero once before γ , remains non-negative after that point, and equals zero again at $\underline{y}$. Therefore, for all values of x above this first crossing point, the integral term is non-negative, and the first derivative of the Lagrangian is strictly positive. For values of x below this first crossing point, instead, we have $g(x) - g(\underline{y}) < 0$, and therefore $\frac{\partial}{\partial x} \left[\int_{\underline{x}}^{\underline{y}} (g(z) - g(\underline{y})) dz \right] = g(\underline{y}) - g(x) > 0$.

Hence, the first derivative of the Lagrangian is strictly increasing in this region and can cross zero at most once.

We have proved that there exists a threshold $a \in [0, \underline{y}]$ such that

$$c(x) = \begin{cases} 0 & \text{if } x < a \\ 1 & \text{if } a \leq x \leq \underline{y} \end{cases}$$

Step 3: we need to focus on the first derivative for $x \in [\bar{y}, 1]$:

$$\frac{\partial L(x, c)}{\partial c(x)} = - \int_{\bar{y}}^x (g(z) - g(\bar{y})) dz + k(g(\underline{y}) - g(\bar{y}))$$

From Step 1, it is clear that $\bar{y} > \gamma$. This implies that the segment $[\bar{y}, x]$ is located on the decreasing part of the function and so $g(z) - g(\bar{y}) < 0$ for every $z \in [\bar{y}, 1] \implies - \int_{\bar{y}}^x (g(z) - g(\bar{y})) > 0$ for every $x \in [\bar{y}, 1]$. Given that $k(g(\underline{y}) - g(\bar{y})) > 0$, we can conclude that the first derivative for $x \geq \bar{y}$ is always positive. This allows us to conclude that $c(x) = 1$ for every $x \in [\bar{y}, 1]$. Summing up all the steps, we can conclude that there exists a threshold $a \in [0, 1]$ such that

$$c(x) = \begin{cases} 0 & \text{if } x < a \\ 1 & \text{if } x \geq a \end{cases}$$

We are left to show that $a \leq \gamma$. This follows directly from Step 2. If the first derivative is positive for every $x \in [0, \underline{y}]$, then $a = 0$, and therefore $a \leq \gamma$. If instead the first derivative crosses zero, there are two cases. When $\underline{y} \leq \gamma$, the threshold satisfies $a \leq \underline{y} \leq \gamma$. When $\underline{y} > \gamma$, Step 2 shows that the first derivative is positive on $[\gamma, \underline{y}]$, so any crossing must occur below γ . Hence $a < \gamma$. Therefore, in all cases, $a \leq \gamma$. This concludes the proof of the second part of the Theorem. From the steps discussed above, it is straightforward to verify that when, instead, g is monotonically decreasing or monotonically increasing, the extreme cases discussed in Section 5.1 are the solution to the sender's LCP.

It is important to highlight that the argument above has a straightforward implication on the structure of the maximization problem for the sender. Indeed, given the formulas for $\underline{y}$ and $\bar{y}$ derived in Section 5.1 and the simplified structure for the complex disclosure rule derived above, we can rewrite the sender's LCP as one in which he simply chooses the optimal value of a :

$$\max_a \int_{\underline{x}}^1 G(x) dx + \int_a^{a + \sqrt{2k(1-a)}} \left(G\left(a + \sqrt{2k(1-a)}\right) - G(x) \right) dx$$

$$- \int_{1-\sqrt{2k(1-a)}}^1 \left(G(x) - G\left(1 - \sqrt{2k(1-a)}\right) \right) dx \quad (21)$$

$$\text{s.t.} \quad a < 1 - 8k \quad (22)$$

The objective function takes into account both the structure of the optimal disclosure rule and the strategy of the receivers through the threshold types. Constraint (22) ensures that positive learning occurs.³⁶

A.7 Proof of Theorem 2

In order to prove the Theorem, we need to refer to the necessary conditions derived for the HCP and the LCP, adapted to our parametric choice of $x \sim \mathcal{U}[0, 1]$ and to the case in which g is U-shaped. Let γ denote the lowest point of g , so that g is strictly decreasing on $[0, \gamma]$ and strictly increasing on $[\gamma, 1]$. We can start from the HCP, assuming that the no-learning constraint does not bind. From the necessary conditions, we can rewrite the first derivatives as

- for every $x \in [0, \hat{x}]$, $\frac{\partial \mathcal{L}^{HCP}(c, \hat{x})}{\partial c(x)} = \int_x^{\hat{x}} [g(z) - g(\hat{x})] dz$;
- for every $x \in [\hat{x}, 1]$, $\frac{\partial \mathcal{L}^{HCP}(c, \hat{x})}{\partial c(x)} = - \int_{\hat{x}}^x [g(z) - g(\hat{x})] dz$.

Since optimality requires $c(x) = 1$ when the first derivative is positive and $c(x) = 0$ when it is negative, we only need to study its sign.

First, we show that we cannot have $\hat{x} > \gamma$. Suppose, by contradiction, that $\hat{x} > \gamma$. Then, for $x \in [\gamma, \hat{x}]$, since g is increasing on $[\gamma, 1]$, we would have $g(z) \leq g(\hat{x})$ for all $z \in [x, \hat{x}]$. This would imply that $\int_x^{\hat{x}} [g(z) - g(\hat{x})] dz \leq 0$ for $x \in [\gamma, \hat{x}]$. Similarly, for $x \in [\hat{x}, 1]$, because g is increasing on $[\gamma, 1]$, we have $g(z) \geq g(\hat{x})$ for all $z \in [\hat{x}, x]$. This would imply that $-\int_{\hat{x}}^x [g(z) - g(\hat{x})] dz \leq 0$. Therefore, the first derivative with respect to $c(x)$ would be negative for all $x \in [\gamma, 1]$. This would lead to $c(x) = 0$ for $x \in [\gamma, 1]$. However this would contradict $\hat{x} > \gamma$. Hence, $\hat{x} \leq \gamma$. In addition, we can argue that such inequality must be strict. If $\hat{x} = \gamma$, then g is strictly increasing above $\hat{x}$, so the first derivative would be strictly negative for every $x > \gamma$, implying $c(x) = 0$ on $(\gamma, 1]$. The complex pool would then lie in $[0, \gamma]$ and, since it has positive mass, its expected value would be strictly below γ , contradicting $\hat{x} = \gamma$. Hence, $\hat{x} < \gamma$.

³⁶A similar expression of the problem can be derived for the case in which learning does not occur, keeping in mind that $\hat{x} = \frac{1+a}{2}$ and that the no-learning constraint would require $a \geq 1 - 8k$. [Ginzburg \(2019\)](#) studies the properties of the optimal threshold rule when the problem of the sender is a censorship one.

The fact we just discussed allows us to say that for $x \in [0, \hat{x}]$, g would be decreasing, and so $g(z) \geq g(\hat{x})$ for all $z \in [x, \hat{x}]$. This would imply that $\int_x^{\hat{x}} [g(z) - g(\hat{x})] dz \geq 0$ for all $x \in [0, \hat{x}]$. Similarly, for $x \in [\hat{x}, \gamma]$, since g is decreasing, we would have $g(z) \leq g(\hat{x})$ for all $z \in [\hat{x}, x]$. This would imply that $-\int_{\hat{x}}^x [g(z) - g(\hat{x})] dz \geq 0$. Therefore, for any $\hat{x}$ the first derivative would be positive for any $x \leq \gamma$, and so $c(x) = 1$.

Let's now consider $x \in [\gamma, 1]$. For such values, the derivative of $-\int_{\hat{x}}^x [g(z) - g(\hat{x})] dz$ with respect to x is $g(\hat{x}) - g(x)$. Since g is decreasing before γ and increasing after, and since $\hat{x} < \gamma$, $g(\hat{x}) - g(x)$ is either always positive, or first positive and then negative. This implies that, by continuity, as $x \rightarrow \hat{x}$, $-\int_{\hat{x}}^x [g(z) - g(\hat{x})] dz$ is first positive, then it grows and, possibly, it decreases to become negative. However, given the structure of the density, the change in sign can happen only once. Hence there exists a threshold $a^{HCP} \in [0, 1]$ such that $c(x) = 1$ for $x \in [0, a^{HCP}]$, and $c(x) = 0$ for $x \in [a^{HCP}, 1]$. This allows us to conclude that the optimal complex disclosure rule in the HCP is a lower-censorship rule. Moreover, the argument above guarantees that $a^{HCP} > \gamma$.

We can now study the optimal complex disclosure rule in the LCP. Again, from the necessary conditions derived in Proposition 3, we know that $c(x)$ is determined by the sign of the first derivatives:

- for every $x \in [0, \underline{y}]$, $\frac{\partial \mathcal{L}^{LCP}(c, \underline{y}, \bar{y})}{\partial c(x)} = \int_x^{\underline{y}} [g(z) - g(\underline{y})] dz + k [g(\underline{y}) - g(\bar{y})]$;
- for every $x \in [\underline{y}, \bar{y}]$, $\frac{\partial \mathcal{L}^{LCP}(c, \underline{y}, \bar{y})}{\partial c(x)} = k [g(\underline{y}) - g(\bar{y})]$;
- for every $x \in [\bar{y}, 1]$, $\frac{\partial \mathcal{L}^{LCP}(c, \underline{y}, \bar{y})}{\partial c(x)} = -\int_{\bar{y}}^x [g(z) - g(\bar{y})] dz + k [g(\underline{y}) - g(\bar{y})]$.

We can prove the statement with similar steps as in Theorem 1.

Step 1: In optimum it must be that $k [g(\underline{y}) - g(\bar{y})] > 0$.

We can show that $g(\underline{y}) > g(\bar{y})$, so that $k [g(\underline{y}) - g(\bar{y})] > 0$ and $c(x) = 1$ for all $x \in [\underline{y}, \bar{y}]$. Assume by contradiction that this is not the case and that $g(\underline{y}) \leq g(\bar{y})$. Given that $\underline{y} < \bar{y}$ and that g is U-shaped, this implies that $\bar{y} > \gamma$. However, this also necessarily implies that $g(\bar{y}) < g(z)$ for all $z \in [\bar{y}, 1]$, making the FOC negative for all $x \in [\bar{y}, 1]$. However, this would imply that there is no complexity used above $\bar{y}$, contradicting the definition of $\bar{y}$ in equation (8), which requires some positive complexity above $\bar{y}$. Hence, it must be that $k [g(\underline{y}) - g(\bar{y})] > 0$ and $c(x) = 1$. This conclusion also implies that $\underline{y} \leq \gamma$. Indeed, if $\underline{y} > \gamma$, then both $\underline{y}$ and $\bar{y}$ would lie on the increasing part of g , and since $\underline{y} < \bar{y}$, we would have $g(\underline{y}) < g(\bar{y})$, a contradiction.

Step 2: Given Step 1, we can directly show that $c(x) = 1$ for all $x \in [0, \underline{y}]$. Indeed, since $\underline{y} \leq \gamma$, it must be that $g(z) > g(\underline{y})$ for all $x \in [0, \underline{y}]$. This, together with $k [g(\underline{y}) - g(\bar{y})] > 0$, implies that the FOC for $x \in [0, \underline{y}]$ is always positive, and, as a consequence, $c(x) = 1$.

Step 3: We can now show that above $\bar{y}$, there is at most one point where the relevant derivative changes sign. Taking the derivative of $-\int_{\bar{y}}^x [g(z) - g(\bar{y})] dz$ with respect to x we get $g(\bar{y}) - g(x)$. We need to consider two possible cases. If $\bar{y} \geq \gamma$, then g is increasing on $[\bar{y}, 1]$. Hence $g(x) \geq g(\bar{y})$ for all $x \in [\bar{y}, 1]$, and the addend is negative and decreasing in x . Since by continuity it is positive at $x = \bar{y}$, it is either always positive or first positive and then negative. This gives at most one point in which the derivative changes sign. If instead $\bar{y} < \gamma$, then g is decreasing on $[\bar{y}, \gamma]$ and increasing on $[\gamma, 1]$. On $[\bar{y}, \gamma]$, we have $g(x) \leq g(\bar{y})$, so the integral would be positive and increasing in x . This implies that the first derivative remains positive on $[\bar{y}, \gamma]$. On $[\gamma, 1]$, the density g is increasing, following the same pattern as before. This implies that the first derivative can change sign at most once. Again, this gives at most one threshold where the complex disclosure rule goes from complexity to transparency.

Overall, this implies that there exists a threshold a^{LCP} such that $c(x) = 1$ for $x \in [0, a^{LCP}]$, and $c(x) = 0$ for $x \in [a^{LCP}, 1]$. Thus the LCP is also a lower-censorship rule. Moreover, since the first derivative is positive for all $x \leq \gamma$, the threshold satisfies $a^{LCP} \geq \gamma$. Therefore, when g is U-shaped, both the HCP and the LCP induce lower-censorship rules, with a threshold to the right of the minimal point of the density.

B Other Simulation Results

To solve the sender's problem numerically, we exploit the theoretical characterizations derived in Sections 4 and 5. For a fine grid of information-processing costs k , we evaluate the relevant candidate solutions across the learning and no-learning cases. For the LCP, we impose the single-threshold structure identified by the theory: upper censorship for unimodal outside-option distributions and lower censorship for U-shaped outside-option distributions. For the no-learning side, we evaluate the corresponding HCP benchmark whenever feasible and also search over the multi-threshold structure associated with the Binding Problem. For each candidate family, we first perform a grid search over feasible threshold values and then locally refine the best candidates numerically. Finally, we select the threshold configuration that delivers the highest sender payoff.

B.1 Equilibrium Learning Threshold: Triangular Outside Option Distributions

Figure 9 reports the equilibrium upper-censorship threshold and the induced receiver learning thresholds for the triangular simulations discussed in Section 5.1.1. The figure is useful for visualizing the mechanism behind Result 1. When learning occurs, the interval between the two learning thresholds identifies the receiver types who process the complex message. When k is low, there is a lot of learning, and few receivers decide to make the choice of walking away without information. This leads to more complexity when the outside option distribution places large mass on higher values. As k increases, this learning region shrinks and eventually disappears, at which point the equilibrium coincides with the censorship benchmark.

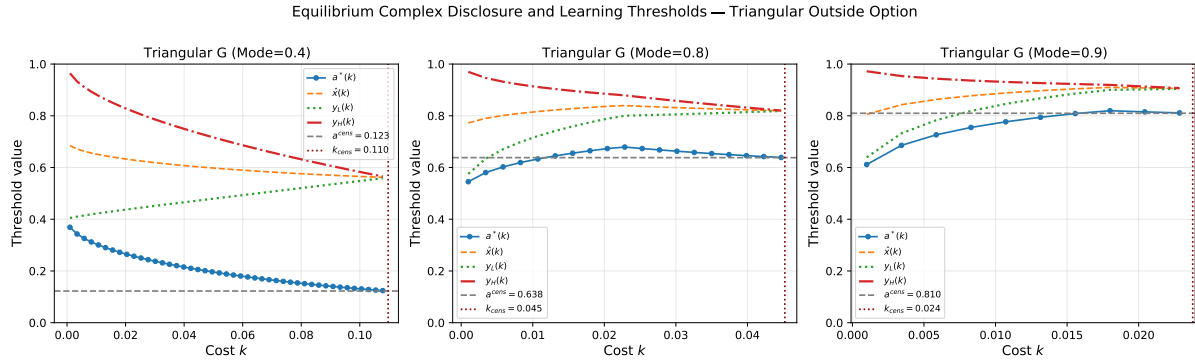


Figure 9: Equilibrium upper-censorship a^* and learning thresholds ($\underline{y} = y_L$, $\bar{y} = y_H$) for triangular outside-option distributions

B.2 Unimodal Beta Outside Option: Simulations

Figure 10 shows the possibility of Result 1 for unimodal Beta distributions, even if with smaller magnitudes. To do so, we compare a Beta(6,2) to a Beta(2,2) distribution, as the former first-order stochastically dominates the latter, providing us with more favorable market alternatives.

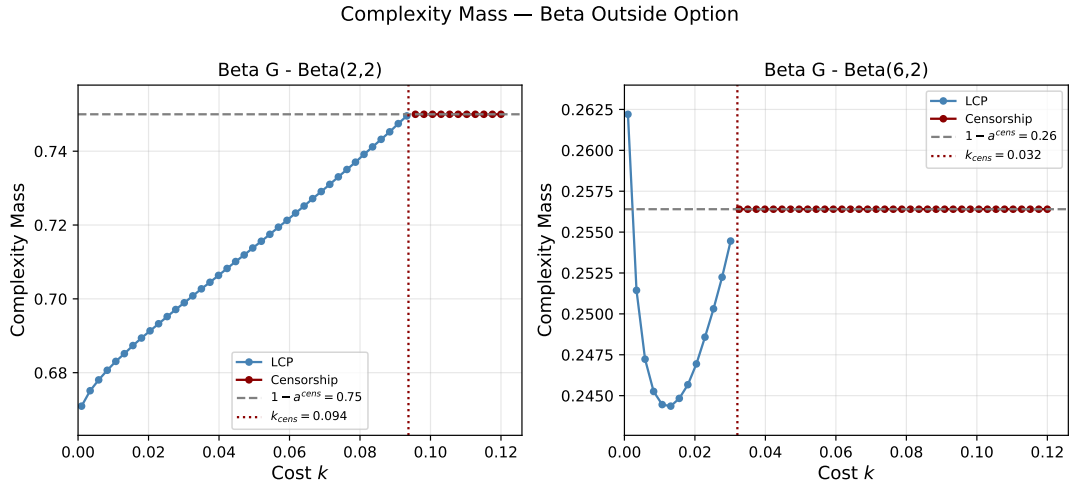


Figure 10: Probability of complex disclosure ($1 - a$) as a function of k , when G is a unimodal Beta distribution and $x \sim \mathcal{U}[0, 1]$

B.3 Non-Uniform Quality Distribution: Simulations

Figures 11 and 12 provide evidence of Result 1 even when we relax the assumption made in Section 5 that the quality x is uniformly distributed. Specifically, these figures plot the likelihood of complex disclosure in the game where learning is possible and in the no-learning benchmark, this time accounting for the alternative distribution of x —which makes the complexity mass different from $1 - a$.

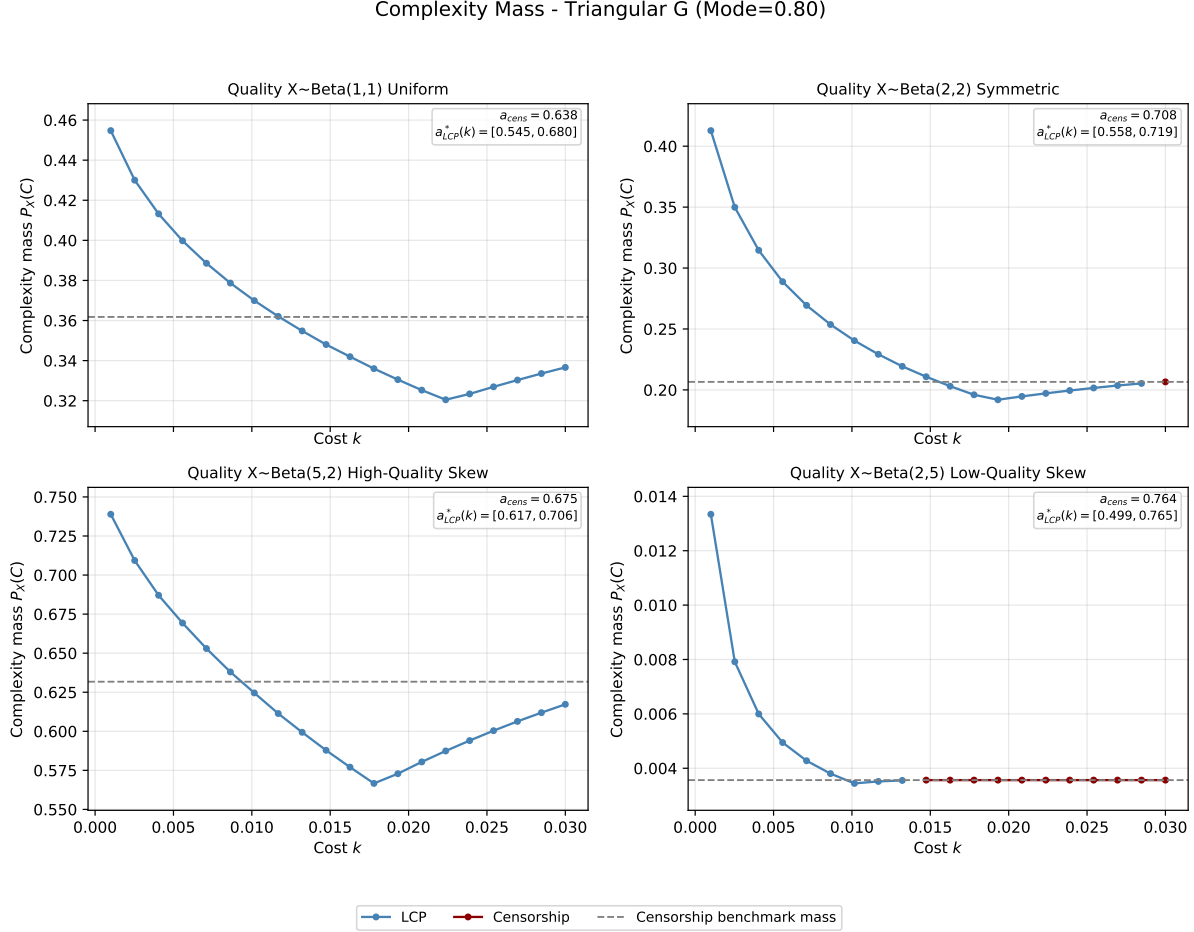


Figure 11: Probability of complex disclosure as a function of k , when G is a triangular distribution with mode 0.8 and $\mathcal{F}$ is a Beta distribution

Complexity Mass - Triangular G (Mode=0.90)

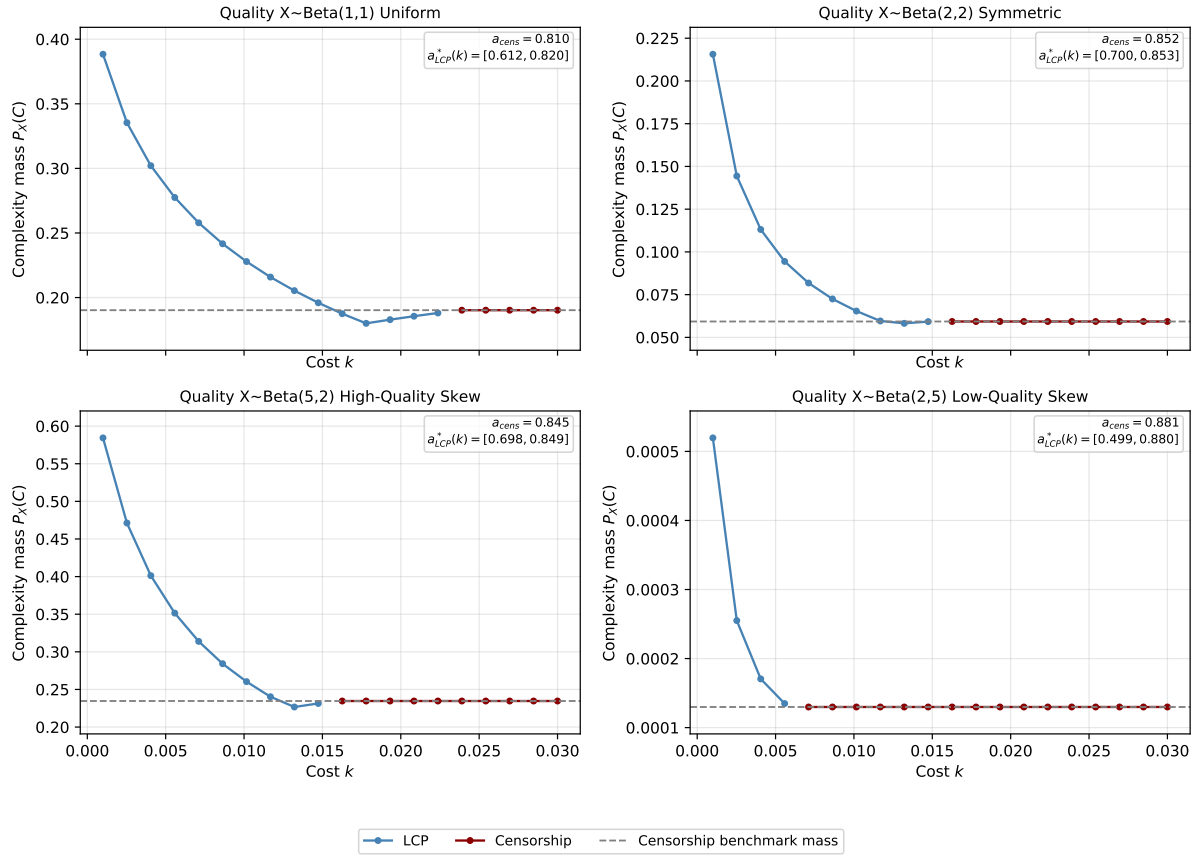


Figure 12: Probability of complex disclosure as a function of k , when G is a triangular distribution with mode 0.9 and $\mathcal{F}$ is a Beta distribution

Online Appendix for

STRATEGIC COMPLEXITY UNDER MANDATED DISCLOSURE

Agata Farina

University of Chicago

Contents

1	Introduction	1
1.1	Related Literature	5
2	The Model	8
3	Receiver's Problem	11
4	Sender's Complex Disclosure Rule	15
4.1	Sender's Low-Cost Problem	17
4.2	Sender's High-Cost Problem	21
4.3	Binding Problem	23
4.4	The Solution to the General Problem	24
5	Applications: Parametric Outside Option Distributions	25
5.1	Homogeneous Markets: Unimodal G	26
5.1.1	Unimodal Outside Option: Numerical Results	29
5.2	Segmented Markets: U-Shaped Outside Options	34
5.2.1	U-Shaped Outside Option: Numerical Results	35
6	Regulatory Implications	36
7	Conclusion	39

A	Proofs	45
A.1	Proof of Proposition 1	45
A.2	Proof of Proposition 2	46
A.3	Proof of Proposition 3	47
A.4	Proof of Proposition 4	49
A.5	Proof of Lemma 1	50
A.6	Proof of Theorem 1	51
A.7	Proof of Theorem 2	54
B	Other Simulation Results	56
B.1	Equilibrium Learning Threshold: Triangular Outside Option Distributions . . .	57
B.2	Unimodal Beta Outside Option: Simulations	57
B.3	Non-Uniform Quality Distribution: Simulations	59
C	Calculus of Variation Approach	1
D	Additional Proofs and Theoretical Results	2
D.1	Properties of Receiver's Value Function	2
D.2	Intermediate Results for Proposition 1	4
D.3	Equivalence of LCP and LCP2	7
D.4	Proof of Lemma 2	10
D.5	Proof of Lemma 3	12
D.6	Proof of Result 2	13
E	Extension: Exogenous Price	14
F	Formal Discussion of Simulations with Triangular G	18
G	Sender Payoff: Simulations	26
H	Welfare Losses due to Complexity: Simulations	28

C Calculus of Variation Approach

We can study the LCP of the sender through a calculus of variations approach. We can compute the Fréchet derivative of the objective function (5) at c , and study the sign of such a derivative for any possible perturbation h . In this way, we can understand which directions of perturbation are beneficial for the objective function and which ones are not, identifying the optimal structure of the function c . However, there is a caveat to this approach: any change in c affects also $\underline{y}(c)$ and $\bar{y}(c)$ (as clear from constraints (7) and (8) in the main problem). Without knowing the structure of c (which is an endogenous object), we cannot derive an explicit relation between the two thresholds and the disclosure rule to be used in Equation (5).

To overcome such an issue, it is enough to notice that when choosing c it might be enough for the sender to know how a change in the disclosure rule modifies the thresholds. To determine such variation, we can use the two constraints (7) and (8), apply the implicit function theorem, and derive the effect of any perturbation of c on the pair $(\underline{y}(c), \bar{y}(c))$.

Applying the implicit function theorem to Equation (7) we get that:

$$\frac{\partial \underline{y}(c)}{\partial c} = \frac{k \int_{\underline{x}}^1 h(x) f(x) dx - \int_{\underline{x}}^{\underline{y}(c)} (\underline{y}(c) - x) h(x) f(x) dx}{\int_{\underline{x}}^{\underline{y}(c)} c(x) f(x) dx}$$

where h represents the relevant perturbation of the function c .

Similarly, we can apply the implicit function theorem to Equation (8) and we get that

$$\frac{\partial \bar{y}(c)}{\partial c} = \frac{\int_{\bar{y}(c)}^1 (x - \bar{y}(c)) h(x) f(x) dx - k \int_{\underline{x}}^1 h(x) f(x) dx}{\int_{\bar{y}(c)}^1 c(x) f(x) dx}$$

where again h represents the relevant perturbation of the function c .

At this point, it is possible to derive the first-order conditions of the objective function $\mathcal{L}(x, c)$ (given by Equation (5)) with respect to c , again perturbing the function c by a given function h and measuring the effect on the thresholds through the expressions above. Following this procedure, we get the FOC

$$\begin{aligned} \frac{\partial \mathcal{L}(x, c)}{\partial c} = & \int_{\underline{x}}^{\underline{y}(c)} \left(G(\underline{y}(c)) - G(x) \right) h(x) f(x) dx - \int_{\bar{y}(c)}^1 \left(G(x) - G(\bar{y}(c)) \right) h(x) f(x) dx \\ & + g(\underline{y}(c)) \left(k \int_{\underline{x}}^1 h(x) f(x) dx - \int_{\underline{x}}^{\underline{y}(c)} (\underline{y}(c) - x) h(x) f(x) dx \right) \end{aligned}$$

$$+ g(\bar{y}(c)) \left(\int_{\bar{y}(c)}^1 (x - \bar{y}(c)) h(x) f(x) dx - k \int_{\underline{x}}^1 h(x) f(x) dx \right)$$

Notice that the derivative above is computed for a generic h . This implies that we can study the effect on the objective function for any feasible perturbation from the starting point c .

The optimal disclosure rule c will be the one that cannot be improved by any feasible perturbation h . It is important to notice that the FOC's are still dependent on c through the two thresholds $\underline{y}(c)$ and $\bar{y}(c)$. This means that, in principle, we could have multiple solutions for this problem, provided by different disclosure rules that imply different strategies for the receivers.

D Additional Proofs and Theoretical Results

D.1 Properties of Receiver's Value Function

To prove the main results in Section 3, we first need to show additional features of the receiver's value function.

Lemma D.1. $V(y) = \max\{V^{no}(y), V^i(y)\}$ is weakly increasing in y .

Proof. Since the max is a monotonic operator, it is enough to prove that both $V^{no}(y)$ and $V^i(y)$ are weakly increasing in y . $V^{no}(y)$ is trivially weakly increasing in y . We are left to prove that $V^i(y)$ is weakly increasing in y as well. Consider any pair $y_1 < y_2$. There are three cases to consider:

1. $y_1 < y_2 \leq \underline{x}$: $V^i(y_1) = V^i(y_2) = \hat{x} - k \implies V^i(y_1) = V^i(y_2)$;
2. $y_1 \leq \underline{x} < y_2$: $V^i(y_1) = \hat{x} - k$ while $V^i(y_2) = y_2 \frac{\int_{\underline{x}}^{y_2} c(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} + \frac{\int_{y_2}^1 xc(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} - k$.

By the definition of $\hat{x}$, it trivially follows that $V^i(y_1) \leq V^i(y_2)$;

3. $\underline{x} < y_1 < y_2$: $V^i(y_1) = y_1 \frac{\int_{\underline{x}}^{y_1} c(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} + \frac{\int_{y_1}^1 xc(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} - k$ and $V^i(y_2) = y_2 \frac{\int_{\underline{x}}^{y_2} c(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} + \frac{\int_{y_2}^1 xc(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} - k$. It is easy to see that

$$\begin{aligned} y_1 \frac{\int_{\underline{x}}^{y_1} c(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} + \frac{\int_{y_1}^1 xc(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} &\leq y_2 \frac{\int_{\underline{x}}^{y_1} c(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} + \frac{\int_{y_1}^{y_2} xc(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} + \frac{\int_{y_2}^1 xc(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} \\ &\leq y_2 \frac{\int_{\underline{x}}^{y_1} c(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} + y_2 \frac{\int_{y_1}^{y_2} c(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} + \frac{\int_{y_2}^1 xc(x) f(x) dx}{\int_{\underline{x}}^1 c(x) f(x) dx} \end{aligned}$$

$$= y_2 \frac{\int_{\underline{x}}^{y_2} c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} + \frac{\int_{y_2}^1 xc(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx}$$

This implies that $V^i(y_1) \leq V^i(y_2)$.

Since $y_1 < y_2$ were chosen arbitrarily, we can conclude that $V^i(y)$ is weakly increasing over its domain. In particular, it is constant for all $y \leq \underline{x}$ and weakly increasing for all $y > \underline{x}$. This concludes the proof that $V(y)$ is weakly increasing in y .

□

Lemma D.2. $V(y) = \max\{V^{no}(y), V^i(y)\}$ is continuous at every $y \in [0, 1]$.

Proof. It is enough to prove that $V^i(y)$ is continuous at every $y \in [0, 1]$ (since it is immediate that $V^{no}(y)$ is continuous at every point of its domain). For every $y < \underline{x}$, this is trivially true.

We can easily see from the definition

$$V^i(y) = \begin{cases} \hat{x} - k & \text{if } y \leq \underline{x} \\ y \frac{\int_{\underline{x}}^y c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} + \frac{\int_y^1 xc(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} - k & \text{if } y > \underline{x} \end{cases}$$

that the function is also continuous at $\underline{x}$. We need to prove that continuity applies for every $y' > \underline{x}$. We want to show that $\forall \varepsilon > 0, \exists \delta > 0$ such that $|y - y'| < \delta \implies |V^i(y) - V^i(y')| < \varepsilon$.

Take $y' \in [0, 1]$ and consider wlog $y > y' > \underline{x}$. $|y - y'| < \delta \implies y < y' + \delta$. Denote the denominator, which does not depend on y , by C . We have that

$$\begin{aligned} & C^{-1} \left| y \int_{\underline{x}}^y c(x)f(x)dx + \int_y^1 xc(x)f(x)dx - y' \int_{\underline{x}}^{y'} c(x)f(x)dx - \int_{y'}^1 xc(x)f(x)dx \right| \\ &= C^{-1} \left| (y - y') \int_{\underline{x}}^{y'} c(x)f(x)dx + y \int_{y'}^y c(x)f(x)dx - \int_{y'}^y xc(x)f(x)dx \right| \\ &< C^{-1} \left| \delta \int_{\underline{x}}^{y'} c(x)f(x)dx + \int_{y'}^{y'+\delta} (y' + \delta - x)c(x)f(x)dx \right| \\ &\leq C^{-1} \left| \delta \int_{\underline{x}}^{y'} f(x)dx + (y' + \delta - y') \int_{y'}^{y'+\delta} f(x)dx \right| \\ &\leq C^{-1} |\delta + \delta| = 2C^{-1}\delta \leq \varepsilon \end{aligned}$$

Hence, for every $\varepsilon > 0$, choose $\delta = \min\left\{\frac{C\varepsilon}{2}, y' - \underline{x}\right\} > 0$. The argument above covers the case $y > y'$. The same argument applies for $y < y'$, and the choice of δ ensures that $y > \underline{x}$ whenever $|y - y'| < \delta$. Therefore, $|V^i(y) - V^i(y')| < \varepsilon$ whenever $|y - y'| < \delta$. Since $y' > \underline{x}$ was chosen arbitrarily, we can conclude that $V^i(y)$ is continuous at every $y > \underline{x}$.

This proves that V^i is continuous at every $y \in [0, 1]$. Since V is the maximum between two functions continuous on their domain, we can conclude that it is continuous as well.

□

Lemma D.3. *There exists no interval $[y_1, y_2]$, with $y_1 < y_2$, over which $V(y)$ is strictly concave in y .*

Proof. We know that the maximum of two convex functions is convex as well. Hence, it is enough to prove that $V^{no}(y)$ and $V^i(y)$ are both convex for $y \in [0, 1]$. Since $V^{no}(y) = \max\{\hat{x}, y\}$ V^{no} is the maximum of two affine functions. Hence V^{no} is convex.

It is left to prove that $V^i(y)$ is convex as well. Notice first that, for $y \leq \underline{x}$, $V^i(y)$ is constant equal to $\hat{x} - k$. For $y > \underline{x}$, differentiating V^i with respect to y gives

$$\frac{dV^i(y)}{dy} = \frac{\int_{\underline{x}}^y c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx}.$$

This derivative is weakly increasing in y . Indeed, for any $\underline{x} < y_1 < y_2$, we have

$$\frac{\int_{\underline{x}}^{y_2} c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} - \frac{\int_{\underline{x}}^{y_1} c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} = \frac{\int_{y_1}^{y_2} c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} \geq 0.$$

Since the first derivative is weakly increasing in y , then it must be that V^i is convex on $(\underline{x}, 1]$.

Finally, since V^i is constant on $[0, \underline{x}]$, its slope on that interval is equal to zero. This implies that V^i is convex on $[0, 1]$. Since V^{no} and V^i are convex, their maximum $V(y) = \max\{V^{no}(y), V^i(y)\}$ is convex. Therefore there exists no interval $[y_1, y_2]$, with $y_1 < y_2$, over which $V(y)$ is strictly concave.

□

D.2 Intermediate Results for Proposition 1

To prove Proposition 1 we need the following lemmas.

Lemma D.4. *There exist some $y \in [0, 1]$ such that $V(y) = V^i(y)$ if and only if $V^i(\hat{x}) \geq V^{no}(\hat{x})$.*

Proof. If Part. If $V^i(\hat{x}) = V^{no}(\hat{x})$, there is nothing to prove, since the statement is trivially true for $y = \hat{x}$. Assume that $V^i(\hat{x}) > V^{no}(\hat{x})$. We want to show that there must exist some $y \in [0, 1]$ such that $V(y) = V^i(y)$. Notice that $V^i(y)$ is weakly increasing and continuous

for every $y \in [0, 1]$ and $V^{no}(y) = \hat{x}$ for every $y \in [0, \hat{x}]$. Since $V^i(0) < V^{no}(0)$ and $V^i(\hat{x}) > V^{no}(\hat{x})$, this means that the two functions cross at least once: there exists $\underline{y} < \hat{x}$ such that $V^i(\underline{y}) = V^{no}(\underline{y}) = \hat{x}$ and $V^i(y) > V^{no}(y)$ for $y \in (\underline{y}, \hat{x}] \implies V(y) = V^i(y)$ for all $y \in [\underline{y}, \hat{x}]$. This proves that there exist some $y \in [0, 1]$ such that $V(y) = V^i(y)$.

Only If Part. We want to show that if there exist some $y \in [0, 1]$ such that $V(y) = V^i(y)$, it must be that $V^i(\hat{x}) \geq V^{no}(\hat{x})$. Assume by contradiction that $V^i(\hat{x}) < V^{no}(\hat{x})$ but there exists some y such that $V(y) = V^i(y)$. Since $V^i(\cdot)$ is weakly increasing, $V^i(0) < V^{no}(0) = \hat{x}$ and $V^i(\hat{x}) < V^{no}(\hat{x}) = \hat{x}$, no $y \leq \hat{x}$ can satisfy $V(y) = V^i(y)$.

Now consider any $y > \hat{x}$. Since $V^{no}(y) = y$, we have

$$V^i(y) - V^{no}(y) = \frac{\int_y^1 (x - y)c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} - k.$$

Moreover it is easy to see that,

$$\begin{aligned} & \int_{\hat{x}}^1 (x - \hat{x})c(x)f(x)dx - \int_y^1 (x - y)c(x)f(x)dx = \\ &= \int_{\hat{x}}^y (x - \hat{x})c(x)f(x)dx + \int_y^1 (x - \hat{x})c(x)f(x)dx - \int_y^1 (x - y)c(x)f(x)dx \\ &= \int_{\hat{x}}^y (x - \hat{x})c(x)f(x)dx + (y - \hat{x}) \int_y^1 c(x)f(x)dx \geq 0. \end{aligned}$$

Recalling that $V^i(\hat{x}) - V^{no}(\hat{x}) = \frac{\int_{\hat{x}}^1 (x - \hat{x})c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} - k$, this implies that $V^i(y) - V^{no}(y) \leq V^i(\hat{x}) - V^{no}(\hat{x}) < 0$. Thus $V^i(y) < V^{no}(y)$ for every $y > \hat{x}$ as well. Hence, there exists no $y \in [0, 1]$ such that $V(y) = V^i(y)$, contradicting the assumption and proving the desired result.

□

Lemma D.5. *If V^i and V^{no} cross, we can have two scenarios:*

- they cross twice, once in $[\underline{x}, \hat{x})$ and once in $(\hat{x}, 1]$;
- they cross once, at $\hat{x}$.

Proof. From Lemma D.4 we know that for V^i and V^{no} to intersect, it must be that $V^i(\hat{x}) \geq V^{no}(\hat{x})$. We need to prove that only the cases listed in the statement are possible.

Let's first prove that they can cross at most once in $[\underline{x}, \hat{x}]$. For $y \leq \hat{x}$, $V^{no}(y) = \hat{x}$. Hence, given the definition of $\hat{x}$, any crossing in this region must satisfy

$$\frac{\int_{\underline{x}}^y (y-x)c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} = k.$$

Take $\underline{x} < y_1 < y_2 \leq \hat{x}$. Then

$$\begin{aligned} & \int_{\underline{x}}^{y_2} (y_2-x)c(x)f(x)dx - \int_{\underline{x}}^{y_1} (y_1-x)c(x)f(x)dx \\ &= (y_2-y_1) \int_{\underline{x}}^{y_1} c(x)f(x)dx + \int_{y_1}^{y_2} (y_2-x)c(x)f(x)dx \geq 0. \end{aligned}$$

Moreover, if both y_1 and y_2 were points in which the two value functions are crossing, the two sides would both be equal to $k \int_{\underline{x}}^1 c(x)f(x)dx$. Since $k > 0$ and by assumption some complexity is used by the sender, this common value is strictly positive. However, if y_1 is a crossing point, it must be that $c(x)f(x) > 0$ for a positive mass of $x < y_1$. This implies that $\int_{\underline{x}}^{y_1} c(x)f(x)dx > 0$ and so that $\int_{\underline{x}}^{y_2} (y_2-x)c(x)f(x)dx > \int_{\underline{x}}^{y_1} (y_1-x)c(x)f(x)dx$, making it impossible for both functions to equate $k \int_{\underline{x}}^1 c(x)f(x)dx$. Therefore V^i and V^{no} cross at most once in $[\underline{x}, \hat{x}]$.

Similarly, we can show that they can cross at most once in $[\hat{x}, 1]$. For $y \geq \hat{x}$, $V^{no}(y) = y$. Hence, any crossing in this region must satisfy

$$\frac{\int_y^1 (x-y)c(x)f(x)dx}{\int_{\underline{x}}^1 c(x)f(x)dx} = k.$$

Take $\hat{x} \leq y_1 < y_2 \leq 1$. Then

$$\begin{aligned} & \int_{y_1}^1 (x-y_1)c(x)f(x)dx - \int_{y_2}^1 (x-y_2)c(x)f(x)dx \\ &= \int_{y_1}^{y_2} (x-y_1)c(x)f(x)dx + (y_2-y_1) \int_{y_2}^1 c(x)f(x)dx \geq 0. \end{aligned}$$

As before, if both y_1 and y_2 were points in which the two value functions are crossing, the two sides would both be equal to $k \int_{\underline{x}}^1 c(x)f(x)dx$. Since $k > 0$ and by assumption some complexity is used by the sender, this common value is strictly positive. However, if y_2 is a crossing point, it must be that $c(x)f(x) > 0$ for a positive mass of $x > y_2$. This implies that $\int_{y_2}^1 c(x)f(x)dx > 0$ and so that $\int_{y_1}^1 (x-y_1)c(x)f(x)dx > \int_{y_2}^1 (x-y_2)c(x)f(x)dx$, making it impossible for both functions to equate $k \int_{\underline{x}}^1 c(x)f(x)dx$. Therefore V^i and V^{no} cross at most once in $[\hat{x}, 1]$.

- Assume that $V^i(\hat{x}) > V^{no}(\hat{x})$. Since $V^i(0) < V^{no}(0)$, $V^i(\hat{x}) > V^{no}(\hat{x})$, and both functions are continuous, it must be that V^i and V^{no} intersect strictly below $\hat{x}$. We also know that $V^i(1) < V^{no}(1)$, so this means that they have to intersect again. We have already proved that the two curves can only intersect once below $\hat{x}$ and at most once above $\hat{x}$. Therefore, when $V^i(\hat{x}) > V^{no}(\hat{x})$, V^i and V^{no} intersect twice, once below and once above $\hat{x}$.
- Assume that $V^i(\hat{x}) = V^{no}(\hat{x}) = \hat{x}$. Then the two functions cross at $\hat{x}$. Since we have shown that they can cross at most once in $[\underline{x}, \hat{x}]$ and at most once in $[\hat{x}, 1]$, there can be no other crossing. Therefore, in this case, the two functions cross only once, at $\hat{x}$.

This concludes the proof that if V^i and V^{no} intersect, they either intersect once at $\hat{x}$ or twice (once below and once above $\hat{x}$).

□

D.3 Equivalence of LCP and LCP2

Lemma D.6. *The LCP and the LCP2 are equivalent. Hence, they provide us with the same set of solutions.*

Proof. To prove the equivalence between the two problems we can start proving that the two feasible sets coincide. Remember that in the LCP we are just optimizing with respect to the disclosure rule c , while in the LCP2 we are optimizing with respect to both c and the pair of thresholds $(\underline{y}, \bar{y})$. Hence, what we need to show is that

- every c which is feasible in the LCP determines a pair $(\underline{y}(c), \bar{y}(c))$ such that the triple $(c, \underline{y}(c), \bar{y}(c))$ is feasible in the LCP2;
- every triple $(c, \underline{y}, \bar{y})$ which is feasible in the LCP2 is such that the disclosure rule c is feasible in the LCP and $(\underline{y}, \bar{y})$ are exactly the thresholds induced by c .

Once we have proved that the two feasible sets coincide, the equivalence of the problems is straightforward. Assume that the disclosure rule c is feasible in the LCP. This means that such rule satisfies constraints (6), (7) and (8). The latter pin down two thresholds $\underline{y}(c)$ and $\bar{y}(c)$.² We need to show that such c and the implied $\underline{y}(c)$ and $\bar{y}(c)$ are feasible also in the LCP2, i.e. they

²Proposition 1 guarantees that for every disclosure rule c feasible in the LCP there exist only two thresholds $\underline{y}(c)$ and $\bar{y}(c)$.

also satisfy constraints (13), (14) and (15). The argument is straightforward for constraints (14) and (15). Indeed, in the LCP the relevant $\underline{y}(c)$ and $\bar{y}(c)$ are chosen to make (7) and (8) satisfied with equality. Since constraints (14) and (15) describe the same relation between c and $(\underline{y}, \bar{y})$ as (7) and (8), we can conclude that they are satisfied if evaluated in the triple $(c, \underline{y}(c), \bar{y}(c))$. We are left to prove that the triple under analysis is able to satisfy constraint (13). Notice that constraints (7) and (8) imply that

$$\int_{\underline{x}}^{\underline{y}(c)} (\underline{y}(c) - x)c(x)f(x)dx = \int_{\bar{y}(c)}^1 (x - \bar{y}(c))c(x)f(x)dx = k \int_{\underline{x}}^1 c(x)f(x)dx$$

In addition, constraint (6) and the definition of $\hat{x}$ imply that

$$\int_{\underline{x}}^{\hat{x}} (\hat{x} - x)c(x)f(x)dx = \int_{\hat{x}}^1 (x - \hat{x})c(x)f(x)dx > k \int_{\underline{x}}^1 c(x)f(x)dx$$

Given the fact that $\int_{\underline{x}}^{\underline{y}(c)} (\underline{y}(c) - x)c(x)f(x)dx$ strictly increases with $\underline{y}(c)$ and $\int_{\bar{y}(c)}^1 (x - \bar{y}(c))c(x)f(x)dx$ strictly decreases with $\bar{y}(c)$, to make both relations above satisfied we need $\underline{y}(c) < \hat{x}$ and $\bar{y}(c) > \hat{x}$. By transitivity, it must be $\underline{y}(c) < \bar{y}(c)$. This implies that the pair $(\underline{y}(c), \bar{y}(c))$ satisfies constraint (13) in the LCP2. This argument allows us to conclude that any feasible disclosure rule c in the LCP, and the implied thresholds $\underline{y}(c)$ and $\bar{y}(c)$, are a feasible triple for the LCP2.

Consider a triple $(c, \underline{y}, \bar{y})$ which is feasible in the LCP2. We need to show that the rule c is feasible in the LCP and that the pair $(\underline{y}, \bar{y})$ coincides with the thresholds determined by c through constraints (7) and (8). For the triple $(c, \underline{y}, \bar{y})$ to be feasible, it must satisfy constraints (13), (14), (15) and (16). We can immediately argue that constraints (7) and (8) are satisfied as well. Indeed, they describe the same relation between c and $(\underline{y}, \bar{y})$ as constraints (14) and (15). This also directly implies that the pair $(\underline{y}, \bar{y})$ is the one that corresponds to the assumed disclosure rule c in the LCP. Indeed, as argued in the previous step of this proof, $\underline{y}$ and $\bar{y}$ are respectively determined through constraints (7) and (8), and they are unique. We are left to prove that (13), (14) and (15) imply constraint (6). Notice that from constraints (14) and (15) we know that

$$\int_{\underline{x}}^{\underline{y}} (\underline{y} - x)c(x)f(x)dx = \int_{\bar{y}}^1 (x - \bar{y})c(x)f(x)dx = k \int_{\underline{x}}^1 c(x)f(x)dx$$

This fact together with $\underline{y} < \bar{y}$ guarantees that $\underline{y} < \hat{x} < \bar{y}$. Indeed, remember that by the definition of $\hat{x}$

$$\int_{\underline{x}}^{\hat{x}} (\hat{x} - x)c(x)f(x)dx = \int_{\hat{x}}^1 (x - \hat{x})c(x)f(x)dx$$

and that $\int_{\underline{x}}^{\underline{y}} (\underline{y} - x)c(x)f(x)dx$ is strictly increasing in $\underline{y}$ and $\int_{\underline{y}}^1 (x - \underline{y})c(x)f(x)dx$ is strictly decreasing in $\underline{y}$. Assume now by contradiction that $\underline{y} < \bar{y} < \hat{x}$. It follows that

$$\int_{\underline{x}}^{\hat{x}} (\hat{x} - x)c(x)f(x)dx > \int_{\underline{x}}^{\underline{y}} (\underline{y} - x)c(x)f(x)dx = \int_{\underline{y}}^1 (x - \underline{y})c(x)f(x)dx > \int_{\hat{x}}^1 (x - \hat{x})c(x)f(x)dx$$

but this contradicts $\int_{\underline{x}}^{\hat{x}} (\hat{x} - x)c(x)f(x)dx = \int_{\hat{x}}^1 (x - \hat{x})c(x)f(x)dx$. Similarly, assume that $\hat{x} < \underline{y} < \bar{y}$. It follows that

$$\int_{\underline{x}}^{\hat{x}} (\hat{x} - x)c(x)f(x)dx < \int_{\underline{x}}^{\underline{y}} (\underline{y} - x)c(x)f(x)dx = \int_{\underline{y}}^1 (x - \underline{y})c(x)f(x)dx < \int_{\hat{x}}^1 (x - \hat{x})c(x)f(x)dx$$

but again this contradicts $\int_{\underline{x}}^{\hat{x}} (\hat{x} - x)c(x)f(x)dx = \int_{\hat{x}}^1 (x - \hat{x})c(x)f(x)dx$. Hence, it must be $\underline{y} < \hat{x} < \bar{y}$. This implies that

$$\int_{\underline{x}}^{\hat{x}} (\hat{x} - x)c(x)f(x)dx > \int_{\underline{x}}^{\underline{y}} (\underline{y} - x)c(x)f(x)dx = k \int_{\underline{x}}^1 c(x)f(x)dx$$

and so the constraints (13), (14) and (15) imply constraint (6). This argument allows us to conclude that any feasible triple $(c, \underline{y}, \bar{y})$ in the LCP2, will be feasible in the LCP, i.e. will provide us with a feasible c which will induce $\underline{y}(c) = \underline{y}$ and $\bar{y}(c) = \bar{y}$.

The two problems (LCP and LCP2) have the same set of feasible alternatives. In addition, they are characterized by the same objective function. The only difference is the variables over which we are optimizing. We can then argue that the two problems will provide us with the same solution.

Assume that c^* is a solution to the LCP and that $(\underline{y}(c^*), \bar{y}(c^*))$ is the pair of receiver's thresholds arising from such disclosure rule. Define $\mathcal{L}(c^*, \underline{y}(c^*), \bar{y}(c^*))$ as the value of the objective function given the chosen disclosure rule. Since c^* is a solution to the problem it must be that $\mathcal{L}(c^*, \underline{y}(c^*), \bar{y}(c^*)) \geq \mathcal{L}(c, \underline{y}(c), \bar{y}(c))$ for all the other feasible c . We want to show that $(c^*, \underline{y}(c^*), \bar{y}(c^*))$ constitutes a solution to the LCP2. Assume by contradiction that this is not the case. This means that there exists another alternative $(c, \underline{y}, \bar{y})$ which is feasible and dominates $(c^*, \underline{y}(c^*), \bar{y}(c^*))$. There are two possibilities:

- The new triple has the same disclosure rule but (at least one) different receiver's thresholds, i.e., it can be described as $(c^*, \underline{y}, \bar{y})$. For the argument made above, this would imply a contradiction in feasibility, since (at least one of) constraints (14) and (15) would be violated. Hence, we get a contradiction;
- The new triple has a different disclosure rule c , but the same receiver's thresholds, i.e., it can be described by $(c, \underline{y}(c^*), \bar{y}(c^*))$. By the argument made above, if this solution is

feasible, it was available also under LCP and the fact that was dominated by our solution implies that $F(c^*, \underline{y}(c^*), \bar{y}(c^*)) > F(c, \underline{y}(c), \bar{y}(c))$. Again, we get a contradiction. The same argument applies to a new triple with different disclosure rules and receiver's thresholds.

Hence, if c^* solves the LCP, then $(c^*, \underline{y}(c^*), \bar{y}(c^*))$ solves the LCP2. Assume now that the triple $(c^*, \underline{y}^*, \bar{y}^*)$ is a solution to the LCP2. This implies that $\mathcal{L}(c^*, \underline{y}^*, \bar{y}^*) \geq \mathcal{L}(c, \underline{y}, \bar{y})$ for every other feasible triple $(c, \underline{y}, \bar{y})$. We want to show that c^* is a solution to the LCP as well. Assume by contradiction this is not the case, i.e., there exists another disclosure rule c that dominates c^* in the LCP. This disclosure rule will induce the two thresholds $(\underline{y}(c), \bar{y}(c))$. By the equivalence of the feasible sets proved above, if c is feasible in the LCP, $(c, \underline{y}(c), \bar{y}(c))$ must be feasible in the LCP2. However, the fact that $(c^*, \underline{y}^*, \bar{y}^*)$ is a solution in the LCP2 implies that $\mathcal{L}(c^*, \underline{y}^*, \bar{y}^*) > \mathcal{L}(c, \underline{y}(c), \bar{y}(c))$. Since the two problems share the same objective function, we get a contradiction again. Hence, if $(c^*, \underline{y}^*, \bar{y}^*)$ solves the LCP2, then c^* solves the LCP.

We can conclude that the two problems are equivalent and provide us with the same set of solutions.

□

D.4 Proof of Lemma 2

In order to prove the lemma, we use the necessary conditions derived in Lemma 1, adapted to the case $x \sim \mathcal{U}[0, 1]$ and to a unimodal density g with mode γ . Denote the first derivative with respect to $c(x)$ for every x as:

$$M_-(x) = \int_x^{\hat{x}} [g(z) - (\lambda + \mu)] dz + \mu k \quad \text{for } x \leq \hat{x},$$

and

$$M_+(x) = - \int_{\hat{x}}^x [g(z) - \lambda] dz + \mu k \quad \text{for } x \geq \hat{x}.$$

At points where such a derivative is strictly positive, optimality requires $c(x) = 1$; at points where it is strictly negative, optimality requires $c(x) = 0$. Points at which it is equal to zero are tie points and do not affect the threshold characterization.

First note that $M_-(\hat{x}) = M_+(\hat{x}) = \mu k > 0$. Thus, by continuity, $c(x) = 1$ in a neighborhood of $\hat{x}$. From the proof of Lemma 1 we also know that $\lambda < g(\hat{x})$. Therefore, $\lambda < g(\hat{x}) < \lambda + \mu$. We now characterize the possible sign changes of M_- and M_+ . To do so, we compute

the derivatives of the two expressions with respect to x , which are $M'_-(x) = \lambda + \mu - g(x)$ and $M'_+(x) = \lambda - g(x)$.

Consider first $x \leq \hat{x}$. If $\hat{x} < \gamma$, then $g(x) \leq g(\hat{x}) < \lambda + \mu$ for all $x \in [0, \hat{x}]$. Therefore $M'_-(x) > 0$ on $[0, \hat{x}]$, and M_- can change sign at most once. Since $M_-(\hat{x}) > 0$, the sign pattern is either always positive or negative and then positive.

If instead $\hat{x} \geq \gamma$, since $g(\hat{x}) < \lambda + \mu$, the set of points in $[0, \hat{x}]$ at which $g(x) > \lambda + \mu$ is either empty or an interval around the mode. Therefore $M'_-(x)$ is first positive, then possibly negative, and then positive again when x approaches $\hat{x}$. Hence, M_- can change sign at most three times, again possibly starting negative. Thus, allowing for degenerate thresholds, the part of the disclosure rule below $\hat{x}$ can be written as follows: there exists $0 \leq a_1 \leq a_2 \leq a_3 \leq \hat{x}$ such that $c(x) = 0$ for $x \in [0, a_1] \cup [a_2, a_3]$, and $c(x) = 1$ for $x \in [a_1, a_2] \cup [a_3, \hat{x}]$.

Now consider $x \geq \hat{x}$. If $\lambda \leq 0$, then the derivative $M'_+(x) = \lambda - g(x) \leq 0$ for all $x \in [\hat{x}, 1]$. Thus, M_+ is weakly decreasing and starts positive, so it can change sign at most once. Therefore, there exists at most one threshold $a_4 \in [\hat{x}, 1]$ such that $c(x) = 1$ for $x \in [\hat{x}, a_4]$, and $c(x) = 0$ for $x \in [a_4, 1]$. Combining this with the characterization below $\hat{x}$, the disclosure rule has at most four non-degenerate thresholds: $0 \leq a_1 \leq a_2 \leq a_3 \leq a_4 \leq 1$, with $c(x) = 0$ for $x \in [0, a_1] \cup [a_2, a_3] \cup [a_4, 1]$, and $c(x) = 1$ for $x \in [a_1, a_2] \cup [a_3, a_4]$.

If instead $\lambda > 0$, recall from the proof of Lemma 1 that $\lambda < g(\hat{x})$. This immediately implies that $M'_+(\hat{x}) = \lambda - g(\hat{x}) < 0$, so M_+ is strictly decreasing in a neighborhood of $\hat{x}$. To characterize the behavior of M_+ on the whole interval $[\hat{x}, 1]$, we need to identify the set of points at which $g(x) > \lambda$, since $M'_+(x) < 0$ exactly on this set. We claim that this set is an interval of the form $[\hat{x}, z]$, for some $z \in (\hat{x}, 1]$. There are two cases to consider. If $\hat{x} \geq \gamma$, then g is strictly decreasing on $[\hat{x}, 1]$ and $g(\hat{x}) > \lambda$: hence g can fall below λ at most once, at some point z , and remains below it afterwards. If instead $\hat{x} < \gamma$, then g is strictly increasing on $[\hat{x}, \gamma]$, where it stays above $g(\hat{x}) > \lambda$, and strictly decreasing on $[\gamma, 1]$, where it can fall below λ at most once. In both cases, the claim holds, with $z = 1$ whenever g remains above λ on the entire interval. As a consequence, M_+ is strictly decreasing on $[\hat{x}, z]$ and strictly increasing on $(z, 1]$. Since $M_+(\hat{x}) = \mu k > 0$, this shape allows for at most two sign changes: M_+ starts positive, can become negative at most once while decreasing, and can return positive at most once while increasing. Denoting by $a_4 \leq a_5$ the two (possibly degenerate) points at which the sign changes—with $a_5 = 1$ when the second change does not occur—we can conclude that $c(x) = 1$ for $x \in [\hat{x}, a_4] \cup [a_5, 1]$ and $c(x) = 0$ for $x \in [a_4, a_5]$.

Combining this with the characterization below $\hat{x}$, the disclosure rule has at most five (possibly degenerate) thresholds: $0 \leq a_1 \leq a_2 \leq a_3 \leq a_4 \leq a_5 \leq 1$, with $c(x) = 0$ for

$x \in [0, a_1] \cup [a_2, a_3] \cup [a_4, a_5]$, and $c(x) = 1$ for $x \in [a_1, a_2] \cup [a_3, a_4] \cup [a_5, 1]$.

D.5 Proof of Lemma 3

In order to prove the lemma, we use the necessary conditions derived in Lemma 1, adapted to the case $x \sim \mathcal{U}[0, 1]$ and to a U-shaped density g with lowest point γ . Thus, g is decreasing on $[0, \gamma]$ and increasing on $[\gamma, 1]$. Define the first derivative with respect to $c(x)$ for every x as

$$M_-(x) = \int_x^{\hat{x}} [g(z) - (\lambda + \mu)] dz + \mu k \quad \text{for } x \leq \hat{x},$$

and

$$M_+(x) = - \int_{\hat{x}}^x [g(z) - \lambda] dz + \mu k \quad \text{for } x \geq \hat{x}.$$

At points where such a derivative is strictly positive, optimality requires $c(x) = 1$; at points where it is strictly negative, optimality requires $c(x) = 0$. Points at which it is equal to zero are tie points and do not affect the threshold characterization.

As in the previous lemma, first note that $M_-(\hat{x}) = M_+(\hat{x}) = \mu k > 0$. Thus, by continuity, $c(x) = 1$ in a neighborhood of $\hat{x}$. From the proof of Lemma 1 we also know that $\lambda < g(\hat{x})$. Therefore, $\lambda < g(\hat{x}) < \lambda + \mu$.

We now characterize the possible sign changes of M_- and M_+ . As before, their derivatives with respect to x are $M'_-(x) = \lambda + \mu - g(x)$ and $M'_+(x) = \lambda - g(x)$.

Consider first $x \leq \hat{x}$. If $\hat{x} < \gamma$, then g is decreasing on $[0, \hat{x}]$. Since $g(\hat{x}) < \lambda + \mu$, the set of points in $[0, \hat{x}]$ at which $g(x) > \lambda + \mu$ is either empty or an interval starting at 0. Therefore $M_-(x)$ can first decrease and then increase again. Since $M_-(x)$ can start positive and since $M_-(\hat{x}) > 0$, the possible sign pattern is either always positive, negative and then positive, or positive, then negative, and then positive.

If instead $\hat{x} \geq \gamma$, then g is decreasing on $[0, \gamma]$ and increasing on $[\gamma, \hat{x}]$. Since $g(\hat{x}) < \lambda + \mu$, we have $g(x) < g(\hat{x}) < \lambda + \mu$ for every $x \in [\gamma, \hat{x}]$. As before, the set of points in $[0, \hat{x}]$ at which $g(x) > \lambda + \mu$ is again either empty or an interval starting at 0. Therefore $M'_-(x)$ can be first negative and then positive. As before, M_- can start positive and then needs to be positive again at $\hat{x}$, so it can change sign at most twice.

Thus, allowing for degenerate thresholds, the part of the disclosure rule below $\hat{x}$ can be written as follows: there exists $0 \leq a_1 \leq a_2 \leq \hat{x}$ such that $c(x) = 1$ for $x \in [0, a_1] \cup [a_2, \hat{x}]$, and $c(x) = 0$ for $x \in [a_1, a_2]$.

Now consider $x \geq \hat{x}$. If $\lambda \leq 0$, then $M'_+(x) = \lambda - g(x) \leq 0$ for all $x \in [\hat{x}, 1]$. Thus M_+ is weakly decreasing and can change sign at most once. Since $M_+(\hat{x}) > 0$, there can

exist a threshold $a_3 \in [\hat{x}, 1]$ such that $c(x) = 1$ for $x \in [\hat{x}, a_3]$, and $c(x) = 0$ for $x \in [a_3, 1]$. Combining this with the characterization below $\hat{x}$, the disclosure rule has at most three non-degenerate thresholds: $0 \leq a_1 \leq a_2 \leq a_3 \leq 1$, with $c(x) = 1$ for $x \in [0, a_1] \cup [a_2, a_3]$, and $c(x) = 0$ for $x \in [a_1, a_2] \cup [a_3, 1]$.

If instead $\lambda > 0$, recall from the proof of Lemma 1 that $\lambda < g(\hat{x})$. This immediately implies that $M'_+(\hat{x}) = \lambda - g(\hat{x}) < 0$, so M_+ is strictly decreasing in a neighborhood of $\hat{x}$. As before, to characterize the behavior of M_+ on $[\hat{x}, 1]$, we need to identify the set of points at which $g(x) > \lambda$, since $M'_+(x) < 0$ exactly on this set. There are two cases to consider. If $\hat{x} \geq \gamma$, then g is strictly increasing on $[\hat{x}, 1]$ and $g(\hat{x}) > \lambda$: hence $g(x) > \lambda$ for every $x \in [\hat{x}, 1]$, and M_+ is strictly decreasing on the whole interval. Since $M_+(\hat{x}) = \mu k > 0$, M_+ can change sign at most once, which gives at most one threshold above $\hat{x}$. If instead $\hat{x} < \gamma$, then g is strictly decreasing on $[\hat{x}, \gamma]$ and strictly increasing on $[\gamma, 1]$. Since $g(\hat{x}) > \lambda$, the set of points at which $g(x) > \lambda$ is the union of an interval $[\hat{x}, z_1)$ and an interval $(z_2, 1]$, with $\hat{x} < z_1 \leq z_2 \leq 1$. In other words, g can fall below λ at most once while decreasing, at some point z_1 , and can rise above λ at most once while increasing, at some point z_2 . We set $z_1 = z_2 = 1$ whenever g remains above λ on the entire interval, and $z_2 = 1$ whenever g falls below λ but does not rise above it again. As a consequence, M_+ can be strictly decreasing on $[\hat{x}, z_1)$, strictly increasing on (z_1, z_2) , and strictly decreasing again on $(z_2, 1]$. Since $M_+(\hat{x}) = \mu k > 0$, this shape allows for at most three sign changes: M_+ starts positive, can become negative at most once while decreasing, return positive at most once while increasing, and become negative again at most once in the final decreasing region. Denoting by $a_3 \leq a_4 \leq a_5$ the (possibly degenerate) points at which the sign changes, we can conclude that $c(x) = 1$ for $x \in [\hat{x}, a_3] \cup [a_4, a_5]$ and $c(x) = 0$ for $x \in [a_3, a_4] \cup [a_5, 1]$.

Combining this with the characterization below $\hat{x}$, the disclosure rule has at most five non-degenerate thresholds: $0 \leq a_1 \leq a_2 \leq a_3 \leq a_4 \leq a_5 \leq 1$, with $c(x) = 1$ for $x \in [0, a_1] \cup [a_2, a_3] \cup [a_4, a_5]$, and $c(x) = 0$ for $x \in [a_1, a_2] \cup [a_3, a_4] \cup [a_5, 1]$.

D.6 Proof of Result 2

Recall that the receiver surplus loss when positive learning takes place is

$$L(a, k) = \frac{1}{2} \int_a^{\bar{y}} (y - a)^2 g(y) dy + \frac{1}{2} \int_{\bar{y}}^1 (1 - y)^2 g(y) dy + (1 - a)k \left(G(\bar{y}) - G(\underline{y}) \right),$$

where $\underline{y}(a) = a + \sqrt{2k(1-a)}$ and $\bar{y}(a) = 1 - \sqrt{2k(1-a)}$. Differentiating $L(a, k)$ with respect to a , we get:

$$\begin{aligned} \frac{\partial L(a, k)}{\partial a} &= \frac{1}{2}(\underline{y} - a)^2 g(\underline{y}) \underline{y}'(a) - \int_a^{\underline{y}} (y - a) g(y) dy - \frac{1}{2}(1 - \bar{y})^2 g(\bar{y}) \bar{y}'(a) \\ &\quad - k \left(G(\bar{y}) - G(\underline{y}) \right) + (1 - a)k \left[g(\bar{y}) \bar{y}'(a) - g(\underline{y}) \underline{y}'(a) \right] \end{aligned}$$

Since $(\underline{y} - a)^2 = (1 - \bar{y})^2 = 2k(1 - a)$, we can see that

$$\frac{1}{2}(\underline{y} - a)^2 g(\underline{y}) \underline{y}'(a) - \frac{1}{2}(1 - \bar{y})^2 g(\bar{y}) \bar{y}'(a) = k(1 - a) \left[g(\underline{y}) \underline{y}'(a) - g(\bar{y}) \bar{y}'(a) \right].$$

Therefore, the derivative of the surplus is

$$\frac{\partial L(a, k)}{\partial a} = - \int_a^{\underline{y}} (y - a) g(y) dy - k \left(G(\bar{y}) - G(\underline{y}) \right).$$

Both terms are weakly negative and, since $G(\bar{y}) > G(\underline{y})$ as we are considering a setting with positive learning, the derivative is strictly negative: $\frac{\partial L(a, k)}{\partial a} < 0$. Thus, the receiver surplus loss is decreasing in a .

Finally, in the case with no learning, the expected quality is $\hat{x} = \frac{1+a}{2}$ and the loss is

$$L(a, k) = \frac{1}{2} \int_a^{\hat{x}} (y - a)^2 g(y) dy + \frac{1}{2} \int_{\hat{x}}^1 (1 - y)^2 g(y) dy.$$

Differentiating with respect to a , we get

$$\frac{\partial L(a, k)}{\partial a} = \frac{1}{4}(\hat{x} - a)^2 g(\hat{x}) - \int_a^{\hat{x}} (y - a) g(y) dy - \frac{1}{4}(1 - \hat{x})^2 g(\hat{x})$$

Notice that $(\hat{x} - a)^2 = (1 - \hat{x})^2 = (\frac{1-a}{2})^2$. Therefore, we are left with

$$\frac{\partial L(a, k)}{\partial a} = - \int_a^{\hat{x}} (y - a) g(y) dy$$

which is strictly negative. This proves Result 2.

E Extension: Exogenous Price

We now discuss how the analysis changes when the product is sold at a fixed price p . Since the price is fixed, the sender's expected payoff is p times the probability of sale. Thus the fixed price changes neither the sender's maximization objective nor the receiver's decision problem, except through the support of the payoff-relevant net quality $x = q - p$, where q is the gross

value of the product and p the fixed price. A receiver with outside option y will buy the product whenever $x = q - p \geq y$. Thus, introducing a fixed price in the model is equivalent to considering the net product value, as mentioned in Section 2. The only substantive difference in the analysis is that the support of the net product value x now depends on the price p and on the distribution q , and it may now not overlap with the support of the outside option y , making the analysis less straightforward than in the paper.

Without loss of generality, we can keep the support of y as $[0, 1]$, and, instead, consider $x \sim \mathcal{F}$ with support $[x_l, x_h]$ to be the distribution of the net value of the product accounting for the fixed price p , with $x_l \in \mathbb{R}$ and $x_h \in \mathbb{R}$. In doing so, we generalize the domain of the quality distribution discussed in the main body of the paper. At this point, we can notice that there are a few cases to study depending on the position of $[x_l, x_h]$ relative to the outside option domain $[0, 1]$. We ignore the cases in which $x_h < 0$ or $x_l > 1$ as they would provide us with degenerate cases in which sales are either impossible or happen with certainty.

Case 1: $x_l \geq 0$ and $x_h = 1$. This is the case equivalent to the one discussed in the main body of the paper, where $x \in [\underline{x}, 1]$. Economically, this case represents one in which the product of the firm can be at least as good as the outside option, sharing the same upper bound of the distribution, which could be interpreted as the maximal point of the consumer's value.

Case 2: $x_l \geq 0$ and $x_h < 1$. This case would represent a situation in which the purchased product cannot deliver less than the worst outside option, but in which there are outside options that are strictly better than the higher value the sender's product can deliver. Again, this analysis would coincide with the one in the main body of the paper, with the caveat that all the receiver types with outside option $[x_h, 1]$ would never interact with the sender. Hence, we can define a new outside option distribution which removes those types from the pool and simply solve the problem we have defined in Section 2.

Case 3: $x_l < 0$ or $x_h > 1$ or both. This case is the one we need to discuss more extensively, since it can allow for the possibility in which the net value of the sender's product is strictly below any value of the receiver's outside option or strictly above. Hence, this case represents the economic scenario in which the sender's product can have extremely good or extremely bad realizations, which can affect the structure of the complex disclosure rule and the willingness to transparently disclose these extreme realizations.

The receiver's problem has the same form as in the baseline model, after replacing the original distribution of product values with the distribution of net values. In what follows, we focus on the interior case in which the relevant receiver learning thresholds lie in $[0, 1]$. What changes is, instead, the sender's problem, because the probability of a sale under transparent

disclosure—i.e., $c(x) = 0$ —is now constant outside $[0, 1]$. Indeed, $G(x) = 0$ for every $x \leq 0$ and $G(x) = 1$ for every $x \geq 1$, as no outside option lies below 0 or above 1, making the product either unattractive or attractive for all receivers.

The structure of the LCP and the HCP stays unchanged, but the objective functions described in (5) and (9) need to accommodate the extended domain of x on $[x_l, x_h]$, recalling that $G(x) = 0$ for $x \in [x_l, 0]$ and $G(x) = 1$ for $x \in [1, x_h]$:

$$\begin{aligned} \text{LCP: } \max_c \int_{x_l}^{x_h} G(x)f(x)dx + \int_{x_l}^{\underline{y}} (G(\underline{y}) - G(x)) c(x)f(x)dx - \int_{\underline{y}}^{x_h} (G(x) - G(\bar{y})) c(x)f(x)dx \\ \text{HCP: } \max_c \int_{x_l}^{x_h} G(x)f(x)dx + \int_{x_l}^{x_h} (G(\hat{x}) - G(x)) c(x)f(x)dx, \end{aligned}$$

subject to the same constraints as in Section 4. Since the form of both objective functions and constraint are the same, the necessary conditions of Proposition 3 and Proposition 4 still guide the definition of the optimal complex disclosure rule. The only key difference is that the integrals defining these conditions can involve the value of g on regions where it is not defined. This implies that they must be adapted to the cases in which $x < 0$ or $x > 1$.

In the LCP, for $x \leq \underline{y}$ the relevant condition is given by the sign of

$$\int_x^{\underline{y}} (g(z) - g(\underline{y})) dz + k(g(\underline{y}) - g(\bar{y})) = G(\underline{y}) - G(x) - g(\underline{y})(\underline{y} - x) + k(g(\underline{y}) - g(\bar{y})).$$

When $x < 0$, $G(x) = 0$ and this reduces to $G(\underline{y}) - g(\underline{y})(\underline{y} - x) + k(g(\underline{y}) - g(\bar{y}))$. Symmetrically, for $x \geq \bar{y}$ the relevant condition is the sign of $-\int_{\bar{y}}^x (g(z) - g(\bar{y})) dz + k(g(\underline{y}) - g(\bar{y})) = -(G(x) - G(\bar{y})) + g(\bar{y})(x - \bar{y}) + k(g(\underline{y}) - g(\bar{y}))$. When $x > 1$, $G(x) = 1$, and the expression becomes

$$-(1 - G(\bar{y})) + g(\bar{y})(x - \bar{y}) + k(g(\underline{y}) - g(\bar{y})).$$

The same adjustments apply to the HCP conditions of Proposition 4: for $x < 0$, the relevant expression is $G(\hat{x}) - g(\hat{x})(\hat{x} - x)$ and for $x > 1$, $-(1 - G(\hat{x})) + g(\hat{x})(x - \hat{x})$.

The economic interpretation behind these new conditions is the same as before, with the additional caveat that the cost/benefit of changing the disclosure rule stays fixed in the extreme intervals. In particular, for $x < 0$, increasing the complexity of a state improves the payoff of the sender by $G(\underline{y})$ in the LCP and $G(\hat{x})$ in the HCP, but then induces a cost of decreasing the value of complex disclosure—i.e., $g(\underline{y})(\underline{y} - x)$ or $g(\hat{x})(\hat{x} - x)$ —which varies with x and tends to be larger the worse is the state pooled under m^c . The symmetric interpretation applies for the states above 1, where the cost is the loss of receivers who would buy for sure if the extremely good states were revealed, $1 - G(\bar{y})$, while the benefit is the improvement in the

expected quality upon a complex message—i.e., $g(\bar{y}) (x - \bar{y})$ or $g(\hat{x}) (x - \hat{x})$. The effect of learning in the LCP problem has the same interpretation as discussed in the main body.

At this point, as we did in Section 5.1 and Section 5.2, we can use the necessary condition to derive the structure of the optimal complex disclosure rule when the outside options are either distributed as a unimodal distribution or as a U-shaped one.

Unimodal Outside Option. Just applying the same steps from Theorem 1 to the adapted first order conditions, it is easy to see that for both the LCP and the HCP, the structure of the disclosure rule is unchanged, featuring transparency at the bottom and complexity at the top through a new upper-censorship threshold $a \in [x_l, \gamma]$.³ This result suggests that even when the quality values can be very extreme compared to the outside option, the sender finds it optimal to reveal bad news and pool good news, making complex disclosure a feature of high-value products. When G is unimodal, receivers tend to be concentrated towards intermediate values. This implies that increasing the value behind m^c through the non-disclosure of $x > 1$ can have similar effects to disclosing extremely high x . In the same way, pushing down such an average value through the non-disclosure of $x < 0$ would have a similar negative impact. The intuition is the same one provided in the main body of the paper.

U-shaped Outside Option. In this parametric case, instead, the more general definition of the domain of x can change the structure of the optimal complex disclosure rule. In the main body of the paper, we find that the structure of c is reversed compared to the unimodal case: the sender transparently discloses high values of x and, instead, tends to hide lower values of the state. However, using the necessary conditions in Propositions 3 and 4, it is easy to see that the sender's incentives to use complex disclosure can change more than once. Depending on the position of x_h , we have multiple structures of the complex disclosure rule that can arise. To keep the discussion simple, we discuss one possibility, which can be a common structure for the LCP and the HCP problem. When $x_l < 0$ and $x_h > 1$, the complex disclosure rule could take a form that can be described by three thresholds $a_1 \leq a_2 \leq a_3$:

$$c(x) = \begin{cases} 0 & \text{if } x < a_1 \\ 1 & \text{if } a_1 \leq x < a_2 \\ 0 & \text{if } a_2 \leq x < a_3 \\ 1 & \text{if } x \geq a_3 \end{cases}$$

with $a_1 < 0$, $a_2 \in (\gamma, 1)$ and $a_3 > 1$. Intuitively, the sender finds it optimal to transparently

³The steps of the proof would closely follow the ones in Theorem 1, together with the continuity of the first derivative at $x = 0$ and $x = 1$. For this reason, we omit the proof of this claim.

disclose extremely low and high realizations— $x < a_1$ and $a_2 \leq x < a_3$ —, and, instead, to hide under complex information low and the extremely high ones— $a_1 \leq x < a_2$ and $x \geq a_3$. The economic intuition for this result is straightforward. When G is U-shaped, receivers are concentrated at the two extremes of the outside option distribution. This implies that the main goal of the sender is to try to attract the low-value and high-value types. Complex disclosure is an effective way to attract the former, but it only works if the value of m^c is large enough. Adding extremely low realizations can be detrimental, and the reduction in value does not justify a small increase in sale probability. However, this negative effect can be mitigated by adding extremely high x in the states revealed with complexity. While these states would lead to a sale probability of 1 when transparently disclosed, they also positively affect the value of complexity, also reducing the share of receivers who walk away without buying. Clearly, this is only worth it for very high x , which has a huge impact on the expectation behind complexity.

The fixed-price extension therefore preserves the main learning channel in the interior of the support. It can enrich the treatment of extreme net-value realizations, especially in segmented markets, but it does not appear to overturn the key mechanism studied in the main text.

F Formal Discussion of Simulations with Triangular G

This Appendix formally proves the results discussed through the simulations for the triangular distribution outside options and uniform quality. In particular, we show that if the mode of the distribution is high enough and k is small, the LCP involves a lower threshold, and so more complexity mass, than the problem in which $k \rightarrow \infty$ and so information is *de facto* censored (Ginzburg, 2019). Formally, we prove the following Proposition:

Proposition F.1. *Let G_γ be the triangular distribution on $[0, 1]$ with mode γ and let $x \sim \mathcal{U}[0, 1]$. Let $a_L^*(k, \gamma)$ denote any optimal threshold of the LCP, over the feasible set $[0, 1 - 8k]$, and let $a^{cens}(\gamma) = \max \left\{ 0, \frac{3\gamma - 2(1-\gamma)\sqrt{\gamma}}{4-\gamma} \right\}$ denote the optimal censorship threshold when $k \rightarrow \infty$. If $\gamma > \frac{17+\sqrt{33}}{32}$, then there exists $\kappa(\gamma) > 0$ such that, for every $k \in (0, \kappa(\gamma))$, $a_L^*(k, \gamma) < a^{cens}(\gamma)$. In addition, when k is low enough, the LCP's payoff is larger than the HCP's payoff, and so, in optimum, the sender chooses more complexity than in a censorship problem.*

Proof. The proof follows three steps. First we prove that for high-enough γ and low-enough k , $a_{LCP}^*(k; \gamma) < a^{cens}(\gamma)$. Second, we show that the sender does not want to prevent learning

for low k , and so the relevant problem to compare with the censorship one is indeed the LCP problem. Last, we collect the necessary restrictions on k used in the previous two steps.

Fix a triangular distribution on $[0, 1]$ with mode $\gamma \in (0, 1)$, CDF G_γ and density g_γ :

$$G_\gamma(x) = \begin{cases} \frac{x^2}{\gamma}, & 0 \leq x \leq \gamma, \\ 1 - \frac{(1-x)^2}{1-\gamma}, & \gamma \leq x \leq 1, \end{cases}$$

$$g_\gamma(x) = \begin{cases} \frac{2x}{\gamma}, & 0 \leq x \leq \gamma, \\ \frac{2(1-x)}{1-\gamma}, & \gamma \leq x \leq 1. \end{cases}$$

For an LCP threshold rule with cutoff $a \leq 1 - 8k$, let $\delta(a, k) = \sqrt{2k(1-a)}$. Then the receiver's learning thresholds can be written as

$$\underline{y} = a + \delta(a, k), \quad \bar{y} = 1 - \delta(a, k).$$

Step 1: Let $\Delta(a, k)$ denote the sender's payoff gain in the LCP with cutoff a relative to full transparency—i.e. $c(x) = 0$ for all $x \in [0, 1]$. We first show that this gain has a simple expression. To simplify notation, set $\delta = \delta(a, k)$. It is easy to see that, given the learning decisions of the receivers, the upper-censorship rule only affects the payoff of the sender, relative to full transparency, for $x \in [a, a + \delta]$ and $x \in [1 - \delta, 1]$. Indeed, below a , the quality values are disclosed, and the receiver's choices are the same. In case of complexity, receivers with $y \in [\underline{y}, \bar{y}]$ learn, and make the same decision as before, while the other groups of consumers do not learn, triggering an increase in sale below $\underline{y} = a + \delta$ and a decrease in sale above $\bar{y} = 1 - \delta$. This implies that

$$\Delta(a, k) = \int_a^{a+\delta} (G_\gamma(a + \delta) - G_\gamma(x)) dx - \int_{1-\delta}^1 (G_\gamma(x) - G_\gamma(1 - \delta)) dx.$$

The first addend can be rewritten as

$$\int_a^{a+\delta} [G_\gamma(a + \delta) - G_\gamma(x)] dx = \int_0^\delta t g_\gamma(a + t) dt.$$

Indeed, notice that

$$G_\gamma(a + \delta) - G_\gamma(x) = \int_x^{a+\delta} g_\gamma(z) dz$$

This makes the whole integral equal to

$$\int_a^{a+\delta} \int_x^{a+\delta} g_\gamma(z) dz dx.$$

Given that g is continuous, bounded, and non-negative, we can switch the order of integration. To do that, notice that $a \leq x \leq z \leq a + \delta$. Therefore, if we first integrate over x and then over z , it must be that $x \in [a, z]$ and that $z \in [a, a + \delta]$. Hence, we get

$$\int_a^{a+\delta} \int_a^z g_\gamma(z) dx dz.$$

At this point, we can notice that $g_\gamma(z)$ does not depend on x and so the integral can be simplified to

$$\int_a^{a+\delta} g_\gamma(z)(z - a) dz.$$

At this point, we can do one more variable substitution by defining $z - a = t$. This directly implies that $z = a + t$ and that $dz = dt$. Also, when $z = a$, $t = 0$ and when $z = a + \delta$, $t = \delta$, giving us

$$\int_0^\delta t g_\gamma(a + t) dt.$$

Following the same steps and defining $t = 1 - z$, we can show that

$$\int_{1-\delta}^1 (G_\gamma(x) - G_\gamma(1 - \delta)) dx = \int_0^\delta t g_\gamma(1 - t) dt.$$

This implies that

$$\Delta(a, k) = \int_0^{\delta(a, k)} t [g_\gamma(a + t) - g_\gamma(1 - t)] dt.$$

We now use this expression to derive a useful approximation of the LCP payoff gain. Since $t \in [0, \delta]$, for $\delta = \sqrt{2k(1 - a)}$, when k is small, also t is small and since g is continuous the terms $g_\gamma(a + t)$ and $g_\gamma(1 - t)$ are close to $g_\gamma(a)$ and $g_\gamma(1)$, respectively. We can add and subtract $g_\gamma(a)$ and $g_\gamma(1)$ inside the integral, to have:

$$\Delta(a, k) = \int_0^\delta t [g_\gamma(a) - g_\gamma(1)] dt + \int_0^\delta t [g_\gamma(a + t) - g_\gamma(a) - (g_\gamma(1 - t) - g_\gamma(1))] dt.$$

Since g is a triangular distribution, $g_\gamma(1) = 0$, and this allows us to simplify the first addend above as

$$\int_0^\delta t [g_\gamma(a) - g_\gamma(1)] dt = g_\gamma(a) \frac{\delta^2}{2}.$$

Using $\delta^2 = 2k(1 - a)$, this becomes

$$g_\gamma(a) \frac{\delta^2}{2} = k(1 - a) g_\gamma(a).$$

We can now work on the second addend, which in our approximation will constitute the error term. It is easy to see that the triangular density is Lipschitz-continuous on $[0, 1]$. This means

that there is a constant $C_\gamma = \max\{2/\gamma, 2/(1-\gamma)\}$ such that, for every $t \in [0, \delta]$, $|g_\gamma(a+t) - g_\gamma(a)| \leq C_\gamma t$ and $|g_\gamma(1-t) - g_\gamma(1)| \leq C_\gamma t$. Therefore,

$$\begin{aligned} & \left| \int_0^\delta t [g_\gamma(a+t) - g_\gamma(a) - (g_\gamma(1-t) - g_\gamma(1))] dt \right| \leq \\ & \int_0^\delta t (|g_\gamma(a+t) - g_\gamma(a)| + |g_\gamma(1-t) - g_\gamma(1)|) dt \leq \\ & \int_0^\delta t (2C_\gamma t) dt = 2C_\gamma \int_0^\delta t^2 dt = \frac{2C_\gamma}{3} \delta^3 \leq \frac{2C_\gamma}{3} (2k)^{3/2} = \frac{2^{5/2}}{3} k^{3/2} C_\gamma. \end{aligned}$$

Since we are operating in a case in which k is low, this implies that we can interpret this term as an error of order $k^{3/2}$, and label it as $r_\gamma(a, k)$. Hence, we can rewrite the LCP payoff gain as

$$\Delta(a, k) = k(1-a)g_\gamma(a) + r_\gamma(a, k).$$

where $|r_\gamma(a, k)| \leq \frac{2^{5/2}}{3} k^{3/2} C_\gamma$. We can now focus on $(1-a)g_\gamma(a)$ and notice that, for the triangular density it is equal to,

$$(1-a)g_\gamma(a) = \begin{cases} \frac{2a(1-a)}{\gamma} & \text{if } a \leq \gamma, \\ \frac{2(1-a)^2}{1-\gamma} & \text{if } a \geq \gamma. \end{cases}$$

At this point, focus on cases with $\gamma > 1/2$. We can see that $\frac{2a(1-a)}{\gamma}$ is maximized at $a = 1/2$, while $\frac{2(1-a)^2}{1-\gamma}$ is decreasing in a , so the largest value is with $a = \gamma$. Moreover, $\frac{1}{2\gamma} > 2(1-\gamma)$ whenever $\gamma > 1/2$. Therefore, $(1-a)g_\gamma(a)$ is uniquely maximized at $a = 1/2$.

We now need to consider the censorship problem. Let $a^{cens}(\gamma)$ denote the maximizer of the objective function

$$\int_0^a G_\gamma(x) dx + (1-a)G_\gamma\left(\frac{1+a}{2}\right).$$

Notice that in the censorship problem, learning is not possible, so we do not need to worry about the existence of other constraints. For $\gamma > 1/2$, the maximizer lies in the region $a \in [2\gamma - 1, \gamma]$. We can show that when $\gamma > 1/2$,

$$a^{cens}(\gamma) = \max \left\{ 0, \frac{3\gamma - 2(1-\gamma)\sqrt{\gamma}}{4-\gamma} \right\}.$$

Moreover, it is easy to see that $a^{cens}(\gamma) > \frac{1}{2}$, if and only if $\gamma > \frac{17+\sqrt{33}}{32}$.⁴

⁴To derive this value of γ , it is enough to notice that $a^{cens}(\gamma) > 1/2$ and defining $\sqrt{\gamma} = t$ implies the inequality $4t^3 + 7t^2 - 4t - 4 = (t+2)(4t^2 - t - 2) > 0$. At this point, we can simply solve $4t^2 - t - 2 > 0$ and redefine $\gamma = t^2$.

At this point, we can fix $\gamma > \frac{17+\sqrt{33}}{32}$. Given this, we have that $a^{cens}(\gamma) > 1/2$. Since $(1-a)g_\gamma(a)$ is strictly decreasing to the right of $1/2$, we have $1/2g_\gamma(1/2) > (1-a)g_\gamma(a)$ for all $a \geq a^{cens}(\gamma)$. Therefore, for every $a \geq a^{cens}(\gamma)$ with $a \leq 1-8k$,

$$\Delta(1/2, k) - \Delta(a, k) = k[1/2g_\gamma(1/2) - (1-a)g_\gamma(a)] + (r_\gamma(1/2, k) - r_\gamma(a, k))$$

and

$$\frac{1}{2}g_\gamma(1/2) - (1-a)g_\gamma(a) \geq \frac{1}{2}g_\gamma(1/2) - (1-a^{cens}(\gamma))g_\gamma(a^{cens}(\gamma)) > 0.$$

Using the bound on the remaining terms,

$$|r_\gamma(1/2, k) - r_\gamma(a, k)| \leq |r_\gamma(1/2, k)| + |r_\gamma(a, k)| \leq 2\frac{2^{5/2}}{3}C_\gamma k^{3/2}$$

We can conclude that

$$\Delta(1/2, k) - \Delta(a, k) \geq k \left[\frac{1}{2}g_\gamma(1/2) - (1-a)g_\gamma(a) \right] - 2\frac{2^{5/2}}{3}C_\gamma k^{3/2}.$$

The first term is strictly positive and of order k , for all $a \geq a^{cens}(\gamma)$, while the second term is of order $k^{3/2}$. Therefore, for all sufficiently small k , $\Delta(1/2, k) > \Delta(a, k)$ for every $a \geq a^{cens}(\gamma)$ such that $a \leq 1-8k$. This implies that the LCP payoff at $a = 1/2$ is higher than the LCP payoff at any $a \geq a^{cens}$. Since for sufficiently small k , $a = 1/2$ is feasible, then we can conclude that, when $\gamma > \frac{17+\sqrt{33}}{32}$ no LCP maximizer can lie weakly above $a^{cens}(\gamma)$. Thus, for sufficiently small k , $a_L^*(k, \gamma) < a^{cens}(\gamma)$, i.e., more complexity is optimal in the LCP problem. This proves the first part of Proposition F.1.

Step 2: We now show that, for small k , the LCP payoff is strictly larger than the best payoff among rules that deter learning, so that $a_L^*(k, \gamma)$ is the optimal threshold of the general sender's problem. Let $B_\gamma = \int_0^1 G_\gamma(x) dx$ denote the full-transparency payoff. Consider any rule $c : [0, 1] \rightarrow [0, 1]$ that leads to no learning. Let $D = \int_0^1 c(x) dx$ be the mass of the complex pool. If $D = 0$, the rule coincides with full transparency, so the payoff gain relative to B_γ is zero. We therefore focus on $D > 0$. Define $A(c) = \int_0^{\hat{x}} (\hat{x} - x)c(x) dx$. Since $\hat{x}$ is the average quality given the complex message m^c , then $\int_0^{\hat{x}} (\hat{x} - x)c(x) dx = \int_{\hat{x}}^1 (x - \hat{x})c(x) dx$ and so

$$\int_0^{\hat{x}} (\hat{x} - x)c(x) dx + \int_{\hat{x}}^1 (x - \hat{x})c(x) dx = \int_0^1 |x - \hat{x}|c(x) dx = 2A(c).$$

By Proposition 1, if the rule c leads to no-learning, then it must be $A(c) \leq kD$.

We first show that this feasibility condition forces D to be small. Since $c(x) \leq 1$ and $x \sim \mathcal{U}[0, 1]$, the interval $[\hat{x} - t, \hat{x} + t]$ can contain at most $2t$ of complex mass. Hence, for

every $t \leq D/2$, the mass outside this interval is at least $D - 2t$. At this point, we can notice that for every x , we can write $|x - \hat{x}| = \int_0^{|x - \hat{x}|} dt = \int_0^\infty \mathbb{I}\{|x - \hat{x}| > t\} dt$. Therefore,

$$\int_0^1 |x - \hat{x}| c(x) dx = \int_0^1 \left[\int_0^\infty \mathbb{I}\{|x - \hat{x}| > t\} dt \right] c(x) dx.$$

Since the integrand is bounded and non-negative, we can switch the order of integration and obtain

$$\int_0^1 |x - \hat{x}| c(x) dx = \int_0^\infty \left[\int_0^1 \mathbb{I}\{|x - \hat{x}| > t\} c(x) dx \right] dt.$$

Equivalently,

$$\int_0^1 |x - \hat{x}| c(x) dx = \int_0^\infty \left[\int_{|x - \hat{x}| > t} c(x) dx \right] dt.$$

Since $\int_0^1 |x - \hat{x}| c(x) dx = 2A(c)$, this leaves us with

$$2A(c) = \int_0^\infty \left[\int_{|x - \hat{x}| > t} c(x) dx \right] dt.$$

We can now use the previous reasoning to bound $\int_{|x - \hat{x}| > t} c(x) dx$ from below. For any t , the set of types with distance at most t from $\hat{x}$ is contained in the interval $[\hat{x} - t, \hat{x} + t]$. As mentioned before, this interval has length at most $2t$. Since $c(x) \leq 1$, the complex-message mass inside this interval is at most $2t$. Since the total complex-message mass is D , it follows that, for every $t \leq D/2$,

$$\int_{|x - \hat{x}| > t} c(x) dx \geq D - 2t.$$

Therefore, since $2A(c) = \int_0^\infty \left[\int_{|x - \hat{x}| > t} c(x) dx \right] dt$ and since the interior integral is always non-negative:

$$2A(c) \geq \int_0^{D/2} (D - 2t) dt = \frac{D^2}{4}.$$

Thus, $A(c) \geq \frac{D^2}{8}$. Combining this inequality with $A(c) \leq kD$ gives $D \leq 8k$.

We can now bound the payoff gain from any complex disclosure rule c that prevents receiver learning. Since $g_\gamma(x) \leq 2$ for all x , the distribution function G_γ is Lipschitz continuous with $C_\gamma = 2$. Indeed, consider any $x, x' \in [0, 1]$ and suppose, without loss of generality, that $x < x'$. Then

$$G_\gamma(x') - G_\gamma(x) = \int_x^{x'} g_\gamma(z) dz$$

with, given the triangular density, $g_\gamma(z) \leq 2$. Therefore,

$$G_\gamma(x') - G_\gamma(x) = \int_x^{x'} g_\gamma(z) dz \leq \int_x^{x'} 2 dz = 2(x' - x).$$

The same argument applies when $x' < x$, so $|G_\gamma(x') - G_\gamma(x)| \leq 2|x' - x|$ and so $C_\gamma = 2$.

At this point, we can write the gap between the sender's no-learning payoff under the rule c and the transparency payoff, as we did in the LCP problem, to get

$$\Delta_H(c) = \int_0^1 [G_\gamma(\hat{x}) - G_\gamma(x)] c(x) dx.$$

Indeed, it is easy to see how the two payoffs only differ for the qualities that are disclosed with complexity—for which $c(x) > 0$ —for which the achievable payoff is $G_\gamma(\hat{x})$. Given the argument above, we can also claim that

$$\Delta_H(c) = \int_0^1 [G_\gamma(\hat{x}) - G_\gamma(x)] c(x) dx \leq 2 \int_0^1 |x - \hat{x}| c(x) dx.$$

Using the result that $\int_0^1 |x - \hat{x}| c(x) dx = 2A(c)$, we obtain $\Delta_H(c) \leq 4A(c)$. Since $A(c) \leq kD$ and $D \leq 8k$, then $A(c) \leq 8k^2$, and so $\Delta_H(c) \leq 32k^2$. Given that this argument was made in general for any disclosure rule that prevents learning, it must also apply for the optimal one c_H^* , and so $\Delta_H(c_H^*) \leq 32k^2$.

We can now compare this upper bound for the sender's payoff without learning to a feasible LCP upper-censorship rule. Consider the LCP complex disclosure rule with cutoff $a = 1/2$ that we worked with before. Recall that $\delta(a, k) = \sqrt{2k(1-a)}$, and so $\delta(1/2, k) = \sqrt{k}$. Using the expression for the LCP gain derived above, we have

$$\Delta(1/2, k) = \int_0^{\sqrt{k}} t [g_\gamma(1/2 + t) - g_\gamma(1 - t)] dt.$$

For sufficiently small k , since $\gamma > 1/2$, we can take k small enough so that $\sqrt{k} < \gamma - \frac{1}{2}$ and $\sqrt{k} < 1 - \gamma$. This would imply that, for every $t \in [0, \sqrt{k}]$, $\frac{1}{2} + t \leq \frac{1}{2} + \sqrt{k} < \gamma$, so $1/2 + t$ lies on the left part of the triangular density, giving us $g_\gamma(1/2 + t) = \frac{2(1/2+t)}{\gamma} = \frac{1+2t}{\gamma}$. Similarly, for every $t \in [0, \sqrt{k}]$, $1 - t \geq 1 - \sqrt{k} > \gamma$, so $1 - t$ lies on the right part of the triangular density, giving $g_\gamma(1 - t) = \frac{2(1-(1-t))}{1-\gamma} = \frac{2t}{1-\gamma}$. Substituting these expressions into the LCP payoff gain gives us

$$\Delta(1/2, k) = \int_0^{\sqrt{k}} t \left[\frac{1+2t}{\gamma} - \frac{2t}{1-\gamma} \right] dt.$$

And directly solving the integral would give us

$$\Delta(1/2, k) = \frac{k}{2\gamma} + \frac{2}{3} \left[\frac{1}{\gamma} - \frac{1}{1-\gamma} \right] k^{3/2}.$$

Since $\gamma > 1/2$, the coefficient multiplying $k^{3/2}$ is negative, since $1/\gamma < 1/(1-\gamma)$. However, consistently to what we showed before, this term is of order $k^{3/2}$, and so smaller in

order than k . Therefore, we can choose a sufficiently small k so that $\left| \frac{2}{3} \left[\frac{1}{\gamma} - \frac{1}{1-\gamma} \right] k^{3/2} \right| \leq \frac{k}{4\gamma}$, allowing us to say

$$\Delta(1/2, k) \geq \frac{k}{2\gamma} - \frac{k}{4\gamma} = \frac{k}{4\gamma}.$$

Since $\Delta_H(c_H^*) \leq 32k^2$, and since, when k is sufficiently small, $k/(4\gamma) > 32k^2$, it follows that $\Delta(1/2, k) > \Delta_H(c_H^*)$. Thus, when the information-processing cost is low enough, the payoff from the LCP rule with upper-censorship cutoff $a = 1/2$ is strictly larger than the payoff of any rule that prevents learning—that is, than the best candidate of both the HCP and the Binding Problem. Since the optimal LCP payoff is at least as high as the one at $a = 1/2$, the sender strictly prefers a rule that does not deter learning. This implies that for small enough k , the relevant sender's problem is the LCP one, which, as we showed before, induces an upper-censorship threshold below the one that would be induced in the censorship game in which learning is not feasible. This concludes the proof of the second step.

Step 3: We can now collect the small- k restrictions used above. Fix $\gamma > \frac{17+\sqrt{33}}{32}$. The previous steps impose finitely many upper bounds on k . First, k must be small enough for $a = 1/2$ to be feasible in the LCP problem—i.e. $k \leq 1/16$; second, it must be small enough for the order- k term to dominate the order- $k^{3/2}$ term in Step 1; third, it must be small enough that $1/2 + t$ and $1 - t$ remain on two opposite sides of the triangular density for all $t \in [0, \sqrt{k}]$; fourth, it must be small enough that $\left| \frac{2}{3} \left[\frac{1}{\gamma} - \frac{1}{1-\gamma} \right] k^{3/2} \right| \leq \frac{k}{4\gamma}$; and finally, it must be small enough that $\frac{k}{4\gamma} > 32k^2$. Each of these restrictions gives a strictly positive upper bound on k for fixed value of $\gamma > \frac{17+\sqrt{33}}{32}$. Therefore, taking the minimum of these bounds, we can define $\kappa(\gamma) > 0$ such that all the inequalities used in Step 1 and Step 2 hold for every $k \in (0, \kappa(\gamma))$.

Combining Step 1 and Step 2, for every $k \in (0, \kappa(\gamma))$, the sender's optimal complex disclosure rule is derived from the LCP problem, and its upper-censorship threshold satisfies

$$a_L^*(k, \gamma) < a^{cens}(\gamma).$$

This concludes the proof of Proposition F.1. □

There are a few aspects of the analysis that are key to highlight. First, the proof only shows that the possibility of an increased complexity when consumers are allowed to learn is a theoretical possibility when the outside option is large enough, and the information-processing cost is low enough. The bounds derived in the proof of Proposition F.1 are very restrictive, imposing much stricter conditions on k than what is needed to derive the result in practice (see Section 5.1.1). Second, the proof only considers a specific parametrization of the outside option and the quality distribution. This is clearly a limitation of the analysis, which does not allow

us to formally derive conditions of such distributions that lead to this comparative statics in information complexity.

G Sender Payoff: Simulations

This section reports the results of our main simulations in terms of sender payoff. In reading these figures, it is important to highlight that the labels LCP and HCP refer to candidate payoff, rather than necessarily to the equilibrium subproblem selected at each value of k . The LCP line reports the highest payoff the sender can obtain among the disclosure rules that induce positive information acquisition. The HCP line reports the highest payoff among the disclosure rules that induce no positive learning. For low values of k , this no-learning payoff corresponds to the one in the Binding Problem, in which the sender chooses the complex disclosure rule c to discourage information acquisition by making the receiver's learning constraint bind. When k becomes large enough, the unconstrained no-learning solution becomes feasible, and the HCP line coincides with the censorship benchmark discussed in the main text, corresponding to [Ginzburg \(2019\)](#).

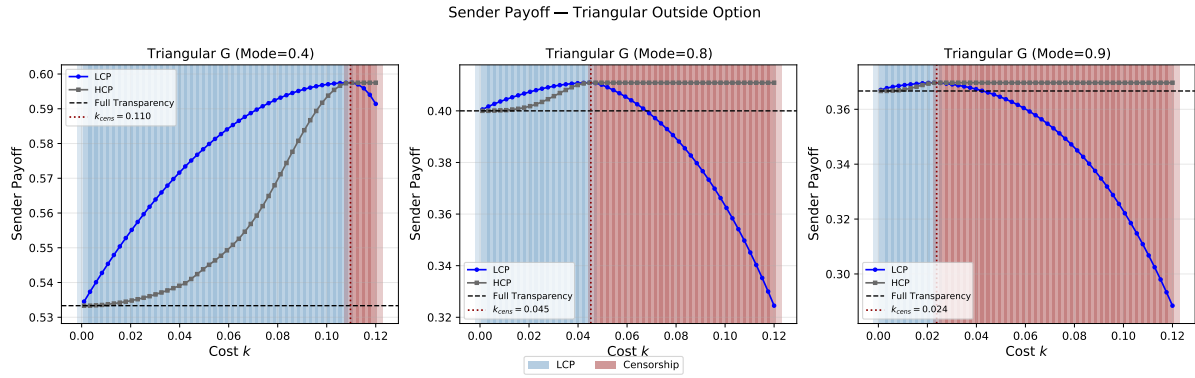


Figure G.1: Sender's payoff as a function of k , when G is a triangular distribution and $x \sim \mathcal{U}[0, 1]$

Sender Payoff — Beta Outside Option

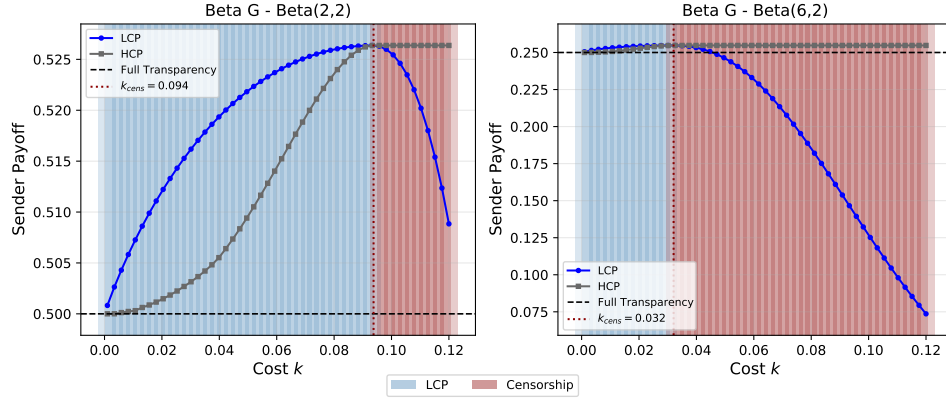


Figure G.2: Sender's payoff as a function of k , when G is a unimodal Beta distribution and $x \sim \mathcal{U}[0,1]$

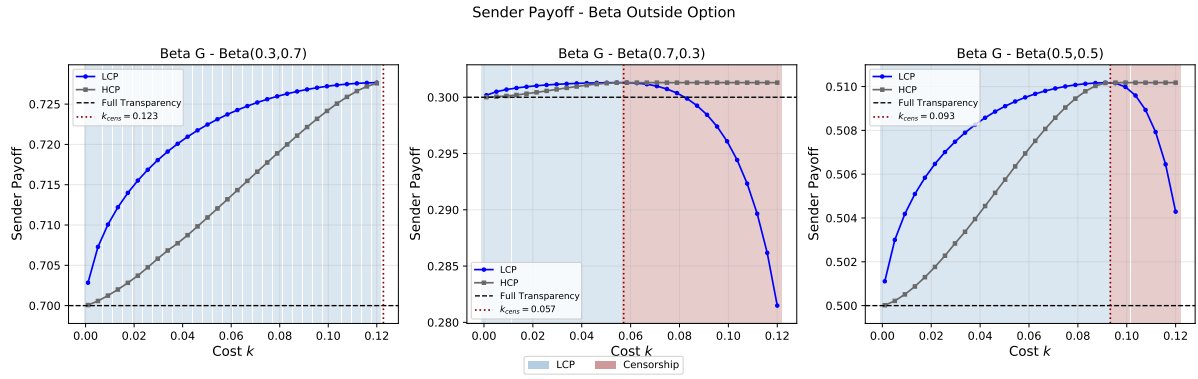


Figure G.3: Sender's payoff as a function of k , when G is a U-shaped Beta distribution and $x \sim \mathcal{U}[0,1]$

H Welfare Losses due to Complexity: Simulations

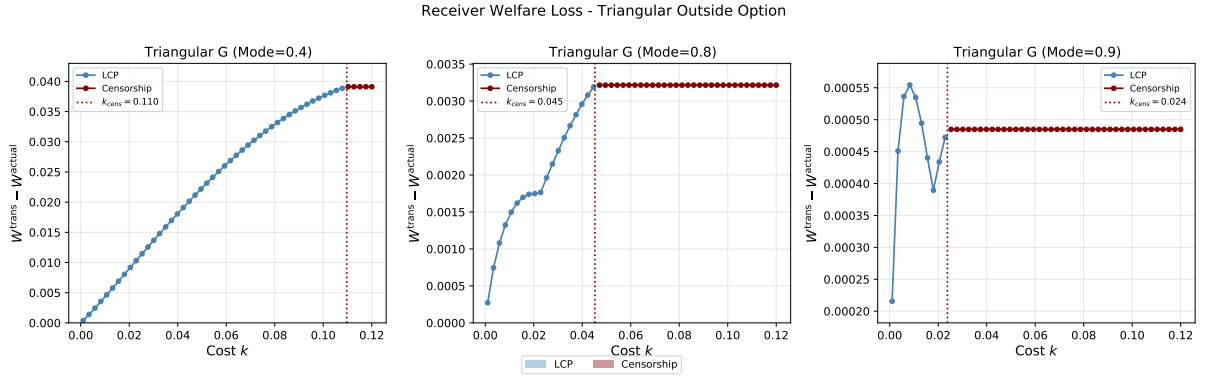


Figure H.1: Receiver's loss as a function of k , when G is a triangular distribution and $x \sim \mathcal{U}[0, 1]$

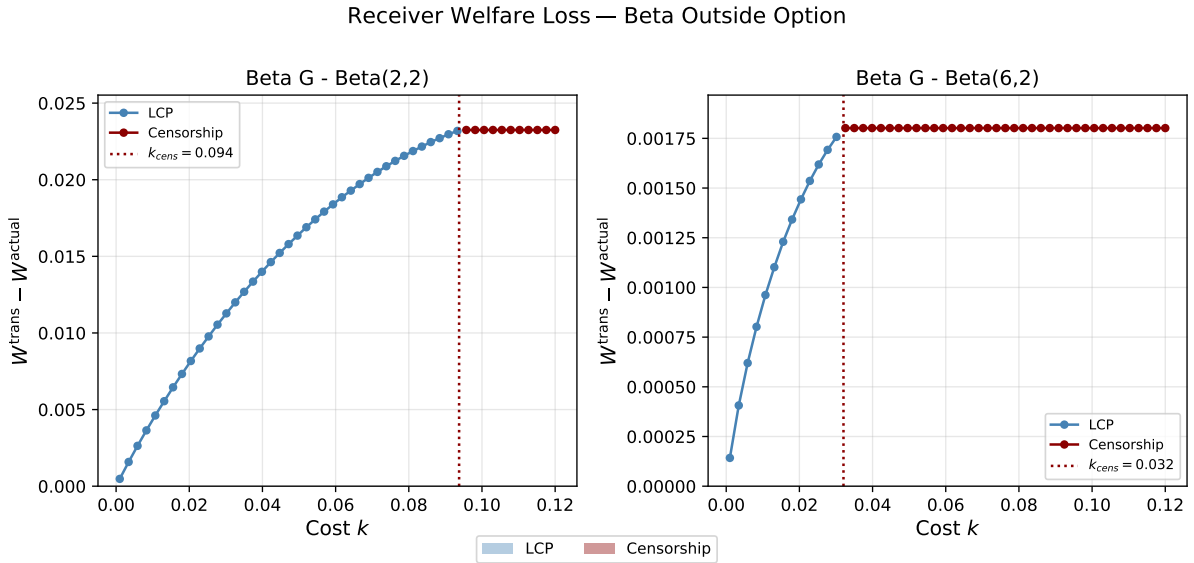


Figure H.2: Receiver's loss as a function of k , when G is a unimodal Beta distribution and $x \sim \mathcal{U}[0, 1]$

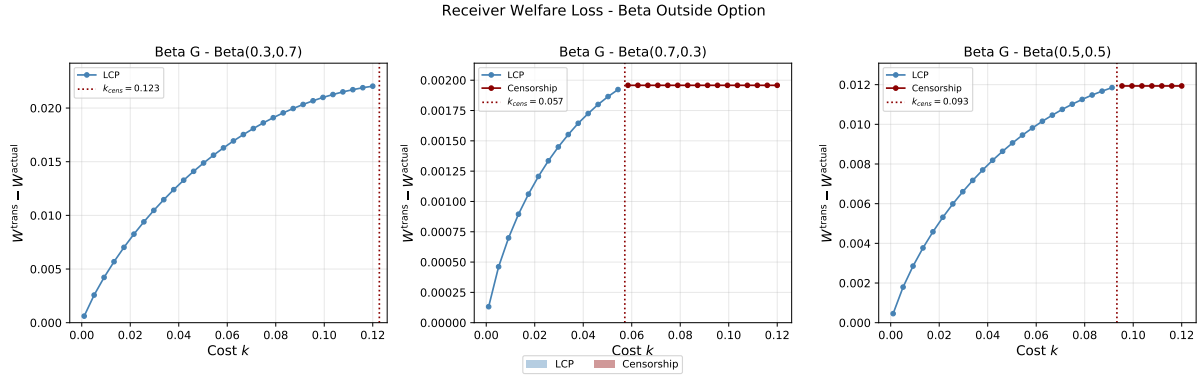


Figure H.3: Receiver's loss as a function of k , when G is a U-shaped Beta distribution and $x \sim \mathcal{U}[0,1]$