

NBER WORKING PAPER SERIES

A MODEL OF BEHAVIORAL MANIPULATION

Daron Acemoglu  
Ali Makhdoumi  
Azarakhsh Malekian  
Asuman Ozdaglar

Working Paper 31872  
<http://www.nber.org/papers/w31872>

NATIONAL BUREAU OF ECONOMIC RESEARCH  
1050 Massachusetts Avenue  
Cambridge, MA 02138  
November 2023

We thank participants at various seminars and conferences for comments and feedback. We are particularly grateful to our discussant Liyan Yang. We also gratefully acknowledge financial support from the Hewlett Foundation, Smith Richardson Foundation, and the NSF. This paper was prepared in part for and presented at the 2023 AI Authors' Conference at the Center for Regulation and Markets (CRM) of the Brookings Institution, and we thank the CRM for financial support as well. The views expressed herein are those of the authors and do not necessarily reflect the views of the National Bureau of Economic Research.

NBER working papers are circulated for discussion and comment purposes. They have not been peer-reviewed or been subject to the review by the NBER Board of Directors that accompanies official NBER publications.

© 2023 by Daron Acemoglu, Ali Makhdoumi, Azarakhsh Malekian, and Asuman Ozdaglar. All rights reserved. Short sections of text, not to exceed two paragraphs, may be quoted without explicit permission provided that full credit, including © notice, is given to the source.

# A Model of Behavioral Manipulation

Daron Acemoglu, Ali Makhdoumi, Azarakhsh Malekian, and Asuman Ozdaglar

NBER Working Paper No. 31872

November 2023

JEL No. D83,D90,D91,L86

## **ABSTRACT**

We build a model of online behavioral manipulation driven by AI advances. A platform dynamically offers one of  $n$  products to a user who slowly learns product quality. User learning depends on a product's "glossiness," which captures attributes that make products appear more attractive than they are. AI tools enable platforms to learn glossiness and engage in behavioral manipulation. We establish that AI benefits consumers when glossiness is short-lived. In contrast, when glossiness is long-lived, users suffer because of behavioral manipulation. Finally, as the number of products increases, the platform can intensify behavioral manipulation by presenting more low-quality, glossy products.

Daron Acemoglu  
Department of Economics, E52-446  
Massachusetts Institute of Technology  
77 Massachusetts Avenue  
Cambridge, MA 02139  
and NBER  
daron@mit.edu

Ali Makhdoumi  
Duke University  
100 Fuqua Drive  
Durham, NC 27708  
ali.makhdoumi@duke.edu

Azarakhsh Malekian  
Rotman School of Management  
University of Toronto  
and LIDS, MIT  
azarakhsh.malekian@rotman.utoronto.ca

Asuman Ozdaglar  
Department of Electrical Engineering  
and Computer Science  
Massachusetts Institute of Technology  
77 Massachusetts Ave, E40-130  
Cambridge, MA 02139  
asuman@mit.edu

# 1 Introduction

The data of billions of individuals are currently being utilized for personalized advertising, product offerings, and pricing, and these practices are set to grow. This exponentially increasing amount of data in the hands of tech companies such as Google, Facebook, and others are raising a host of concerns, including those related to privacy. Many economists, AI experts, and researchers have a fairly optimistic take on this: more data will mean more informative advertising, better product targeting, and finer service differentiation, benefiting users (see, e.g., [Varian \[2018\]](#)).

[Zuboff \[2019\]](#) in *The Age of Surveillance Capitalism* takes a diametrically opposed position. She argues that the economic model of AI “claims human experience as free raw material for hidden commercial practices of extraction, prediction, and sales?” (p. 9). She maintains that the logic of this new system, dominated by Big Tech, turns on the extraction of “behavioral surplus”, created by what she calls “behavioral modification”, meaning the ability of tech companies with vast amounts of data to be able to modify the behavior of users and profit from this. A similar perspective has been discussed in [Hanson and Kysar \[1999\]](#) who note: “Once one accepts that individuals systematically behave in non-rational ways, it follows from an economic perspective that others will exploit those tendencies for gain” (p. 635).

Yet these arguments are insufficient since the modification of behavior may be for the user’s good (directing her towards better products). In this paper, we are interested in the conditions under which behavioral modification harms consumers. We develop a behavioral model to theoretically explore the “bad” type of modification, which we refer to as *behavioral manipulation*. Some examples include retailers forecasting whether a woman is pregnant and presenting them hidden ads for baby products; various companies estimating “prime vulnerability moments” and send ads for beauty products; marketing strategies targeted at more “vulnerable populations” such as the elderly or children; websites favoring products (including credit cards and subscription programs) with delayed costs and apparent short-term benefits; or streaming platforms algorithmically estimating more addictive videos for certain consumer types to maximize engagement. This type of behavioral manipulation would not be possible if consumers were fully rational and understood that with vast amounts of data, platforms may acquire and manipulate relevant information that they do not know about their preferences. However, when the full implications of the superior information of platforms, due to advances in AI and big data, are not properly understood by users, there could be room for much more pernicious types of behavioral modification, which we explore in this paper.

## 1.1 Our Model and Results

We suppose consumers are sufficiently used to and can thus act rationally in the status quo environment. But they are not “hyper-rational” enough to understand how the environment changes after technology platforms can collect vast amounts of data both on them and other consumers with similar characteristics, thus acquiring greater *behavioral predictive power* (meaning prediction of future behavior conditional on current inducements).

We consider a platform that offers one of  $n$  products to a user together with an associated price at any point in time. Each product  $i$  has a quality that is either 1 (high) or 0 (low). When the user consumes a product, she receives a signal about its quality. There is an extraneous factor that impacts the signal: some products may initially be in a state that masks bad news (e.g., people tend to be favorably disposed toward products with some properties such as better appearance, salient attractive characteristics, or hidden and delayed costs). We denote this state for product  $i$  at time 0 by  $\alpha_{i,0}$ , and refer to it as the “glossiness” of the item. We model the dynamics of  $\alpha_{i,t}$  with a continuous-time Markov chain that starts from either 0 or 1. The bad news about the quality of a product will be masked if and only if the initial glossiness state is  $\alpha = 1$  — meaning that the product is glossy. This state is temporary, indicating that a low-quality product may initially appear good, but eventually, the user will understand this. In particular, the continuous-time Markov chain transitions from an initial glossiness state of 1 to an absorbing glossiness state of  $\alpha = 0$  at the rate  $\rho$ . Finally, if a low-quality product is not glossy, the user receives bad news about it at the rate  $\gamma$  (here, bad news means the product is of low quality).

We consider two informational environments: (i) A *pre-AI environment* (or status quo) in which the platform has no informational advantage, and the glossiness state is unknown to both the platform and the user. By observing the signals — specifically, whether bad news has arrived or not — the platform and the user update their belief about both the product’s underlying quality and the product’s state. (ii) A *post-AI environment* in which the platform can estimate the product’s glossiness for this user from the behavior of others with similar characteristics from whom it has collected data.<sup>1</sup> The key (and only) behavioral bias in our model is that users do not understand that the platform knows and can exploit its knowledge of the glossiness state in this new environment. This implies that the user does not learn from the offering and the price decision of the platform.

After characterizing the learning dynamics and the equilibrium strategy of the platform and the user, we turn to our main results. First, for a large enough transition rate  $\rho$  — which means that the initially glossy products remain glossy for a short period — the ex-ante expected user utility, welfare, and platform profits in the post-AI environment are larger than in the pre-AI environment. This confirms the common wisdom that data collected by online platforms can enable better product offerings and higher quality matches for the users. To further understand the intuition for this result, we highlight two opposing forces that impact the user utility in the post-AI environment. There is a *manipulation effect*, resulting from the platform offering a low-quality glossy product ahead of non-glossy products. Counteracting this, there is a *helpfulness effect* that means the platform guides the user towards higher-quality products. With a sufficiently large transition rate  $\rho$ , the helpfulness effect dominates. This is because, for low-quality products, the “manipulation” state in which bad news is masked is short, making the manipulation effect both less profitable for the platform and less consequential for the user’s utility.

Second, in contrast with the modal view in the literature, when the transition rate  $\rho$  is small — the

---

<sup>1</sup>Specifically, the platform can harness the purchase, webpage visit and review data of users with similar characteristics to estimate which products have strong short-term appeal to the relevant demographic group.

In this context, see [Ekmekci \[2011\]](#) and [Che and Hörner \[2018\]](#) for past analysis of the information content of online recommendation systems. In [Acemoglu et al. \[2022a\]](#), even if users are fully Bayesian, they will learn less from such reviews than the platform because they see a summary statistic of past scores.

initially glossy products remain glossy for a long period — the ex-ante expected user utility and welfare in the post-AI environment will be smaller than in the pre-AI environment. While this happens, platform profits always increase because platforms benefit from behavioral manipulation. Intuitively, in this case, with platform superior information, the manipulation effect dominates, as glossy items remain so for a while, enabling the platform to extract significant surplus from the user.<sup>2</sup>

Third, for a sufficiently small transition rate  $\rho$ , as the number of products increases, the ex-ante expected user utility and welfare in the post-AI environment decreases, while the platform continues to benefit. Intuitively, a greater number of products provides more opportunities for the platform to manipulate user behavior by finding glossy items that are profitable. This result implies that as advances in AI simultaneously expand the platform’s ability to collect more information about users and widen their offerings, there may be a “double whammy” from the viewpoint of the consumers.

## 1.2 Related Literature

Our paper relates to the literature on the legal and economic aspects of big data owned by digital platforms. For example, [Pasquale \[2015\]](#) and [Zuboff \[2019\]](#) argue that big data gives digital platforms enormous power for actions that can potentially harm users. Similarly, [Calo \[2014\]](#), [Lin \[2017\]](#), and [Zarsky \[2019\]](#) argue that digital platforms can manipulate consumer behavior and suggest regulatory measures (see also [Tirole \[2020\]](#) for a discussion of regulations and policy implications).

More closely related are a few works investigating whether information harms consumers, either because of surveillance ([Tirole \[2021\]](#)) or because of price discrimination ([Taylor \[2004\]](#), [Acquisti and Varian \[2005\]](#), [Hidir and Vellodi \[2019\]](#), [Bergemann et al. \[2015\]](#), and [Bonatti and Cisternas \[2020\]](#)). See [Acquisti et al. \[2016\]](#), [Bergemann and Bonatti \[2019b\]](#), and [Agrawal et al. \[2018\]](#) for excellent surveys of this literature. We depart from this literature by focusing on behavioral manipulation — distortion of product choice and experimentation — rather than price discrimination.<sup>3</sup> Within this literature, our paper is most closely related to independent work by [Liu et al. \[2023\]](#), which examines the impact of data sharing on consumer welfare and algorithmic inequality. This paper argues that data sharing benefits consumers by providing access to desired products and services, but also enables firms to identify and exploit temptation and self-control problems. This differs but is complementary to our emphasis on platforms identifying sources of mis-learning due to glossiness and other extraneous factors.

Our paper also relates to the literature on information markets such as [Bergemann and Bonatti \[2015\]](#), [Goldfarb and Tucker \[2011\]](#), and [Montes et al. \[2019\]](#) who focus on the use of personal data for improved allocation of online resources and [Admati and Pfleiderer \[1986\]](#), [Anton and Yao \[2002\]](#), [Horner and Skrzypacz \[2016\]](#), [Bergemann et al. \[2018\]](#), and [Begenau et al. \[2018\]](#) who investigate how information

<sup>2</sup>This result also implies that our model derives “endogenous privacy costs” (due to manipulation) when  $\rho$  is small, while there are no such privacy costs when  $\rho$  is large.

<sup>3</sup>Our paper also relates to the works that study privacy concerns in data markets such as [Choi et al. \[2019\]](#), [Bergemann and Bonatti \[2019a\]](#), [Acemoglu et al. \[2022b\]](#), and [Gradwohl \[2017\]](#), who study data externalities; [Fainmesser et al. \[2023\]](#) and [Jullien et al. \[2020\]](#), who consider the negative effects of leaking user’s (private) personalized; and [Ichihashi \[2020b\]](#) and [Ichihashi \[2020a\]](#), who explore the role of information intermediaries and dynamic data collection by platforms. More recent work by [Köbis et al. \[2021\]](#) has empirically investigated broader social implications of extensive data collection and use of machine learning techniques online.

can be monetized either by dynamic sales or optimal mechanisms.

Our work more directly builds on the literature on experimentation and multi-armed bandits. Bandit models, first analyzed in Robbins [1952], have been used for several applications, including pricing (Rothschild [1974]); product search (Bolton and Harris [1999]); research and development (Keller and Rady [2010]); and learning and wage setting (Felli and Harris [1996]). See also Berry and Fristedt [1985] and Krylov [2008] for book-length treatments of bandit problems. Even more closely related to our analysis is the exponential bandit framework of Keller et al. [2005]. The mathematical arguments in our paper are different from those in these past works, because we have to keep track of the beliefs of the platform and the user, which follow different laws of motion.

Finally, our work is inspired by a growing body of evidence that platforms are acquiring a large amount of information about users and influencing their behavior by various strategies (Susser and Grimaldi [2021]). A few other contributions include Acquisti et al. [2015] and Martin and Nissenbaum [2017] who empirically study individuals' ability to control their personal information and Calvo et al. [2020] who study how platform algorithms can enable price manipulation.

The rest of the paper proceeds as follows. Section 2 presents our model. Section 3 provides several results useful for the rest of the analysis. Sections 4 and 5 characterize the equilibrium in the pre-AI and post-AI environments, respectively. Section 6 establishes that manipulation occurs — it is possible for the user to be worse off in the post-AI. Section 7 concludes, while the online Appendix presents the omitted proofs.

## 2 Model

We consider a platform that hosts  $n > 1$  products denoted by  $\mathcal{N} = \{1, \dots, n\}$ . Each product  $i \in \mathcal{N}$  has a quality  $\theta_i \in \{0, 1\}$  and an initial glossiness state  $\alpha_{i,0} \in \{0, 1\}$ . The quantity  $\theta_i$  is fixed and we refer to it as the *quality* of product  $i$  ( $\theta_i = 1$  means a high-quality product). On the other hand,  $\alpha_{i,t}$  changes over time and we refer to  $\alpha_{i,t}$  as the *glossiness state*, or simply *glossiness*, of product  $i$  at time  $t$ .

The platform interacts with a user over an infinite time horizon in continuous time, and at each time  $t \in \mathbb{R}_+$  offers one of the products and an associated price to the user. We let  $x_{i,t} \in \{0, 1\}$  denote the platform's decision about offering product  $i$  at time  $t$  and  $p_{i,t}$  denote the offered price. If the user purchases the product, she receives a signal about its quality and updates her belief (about both  $\theta_i$  and  $\alpha_{i,t}$ ). In particular, if the user purchases product  $i$  at time  $t$ , she observes a signal, described as follows:

1. If  $\theta_i = 1$ , then there is no "bad" news about product quality.
2. If  $\theta_i = 0$  and  $\alpha_{i,t} = 0$ , then bad news (fully revealing signal about product quality) arrives at the rate  $\gamma > 0$ .
3. If  $\theta_i = 0$  and  $\alpha_{i,t} = 1$ , then there is no bad news about product quality.

In summary, a high-quality product never generates bad news (the first case), but a low-quality product may (the second case). In the third case, where  $\alpha_{i,t} = 1$ , the low-quality product is glossy and

does not generate bad news. Glossiness is what makes products appear better than they are (at least in the short run) and enables behavioral manipulation. In what follows, we let  $S_{i,t} = \text{B}$  to designate the arrival of bad news for product  $i$  at time  $t$  and  $S_{i,t} = \text{NB}$  to designate no bad news for product  $i$  at time  $t$ . When purchasing product  $i$ , the user receives no signal about other products  $j \in \mathcal{N} \setminus \{i\}$ .

Glossiness  $\alpha_{i,t}$  follows a homogeneous continuous-time Markov Chain (CTMC) with two states  $\{0, 1\}$  with *transition rate matrix*

$$\begin{bmatrix} 1 & 0 \\ \rho & 1 - \rho \end{bmatrix}.$$

This means that a low-quality product appears high-quality for a while, but this type of glossiness is temporary and come to an end at the rate  $\rho$ .

## 2.1 Information Environment

The platform and the user share a common prior about the initial glossiness state and the quality of each product. In particular, they both understand that a low-quality product can initially be glossy ( $\alpha_{i,0} = 1$ ) or non-glossy ( $\alpha_{i,0} = 0$ ), while high-quality products are always non-glossy. Specifically, for any  $i \in \mathcal{N}$  we have the following as common knowledge:

$$\mathbb{P}[\alpha_{i,0} = 0 \mid \theta_i = 1] = 1, \quad \mathbb{P}[\alpha_{i,0} = 1 \mid \theta_i = 0] = \lambda, \quad \text{and} \quad \mu_{i,0} = \mathbb{P}[\theta_i = 1] \quad (1)$$

for some initial belief  $\mu_{i,0} \in [0, 1]$  and  $\lambda \in (0, 1)$ .

At any time  $t$ , the information available to the user and the platform is whether bad news for product  $i$  has arrived for all  $i \in \mathcal{N}$ . We let

$$\text{NBN}_{i,t} = \{S_{i,\tau} = \text{NB for all } \tau \in [0, t)\}$$

denote the event that there has been no bad news for product  $i$  at times  $\tau \in [0, t)$ . We also use

$$I_{i,t} = \begin{cases} 1 & \text{if } S_{i,\tau} = \text{NB for all } \tau \in [0, t) \\ 0 & \text{otherwise} \end{cases}$$

be the indicator that product  $i$ 's quality at time  $t$  is not 0 from the user's perspective.

The information sets of the platform and the user only differ in the initial signals that the platform observes about the glossiness of different products. In particular, we consider two settings:

1. *Pre-AI environment* in which the platform does not have an information advantage and observes the same signals as the user.
2. *Post-AI environment* in which the platform has an information advantage and observes the initial glossiness state  $\alpha_{i,0}$  for all  $i \in \mathcal{N}$ .

The justification for additional information possessed by the platform in the post-AI environment

was already discussed in the Introduction, but briefly, we consider this as a simple representation of the fact that the platform has a wealth of data about the behavior and preferences of other users with similar characteristics, and thus can estimate which are the products that will (temporarily) appear to the user in question to be more attractive than their quality warrants. For simplicity, these data are not directly informative about the true quality of the product.

A few points about the model are worth mentioning. First, for simplicity, we assume that in the post-AI environment, the platform can perfectly estimate the initial glossiness state  $\alpha_{i,0}$ . Second, the platform's information advantage is about the initial glossiness state and not the product's true quality, which could be idiosyncratic across users. Our main results continue to hold even if we relax these two assumptions. Third, the platform's only information advantage is about the initial glossiness state and not the glossiness trajectory (which is likely to be idiosyncratic across users within the same demographic group). This assumption limits the platform's informational advantage as well. We establish that even with this limited form of informational advantage, the platform can have a very profitable behavioral manipulation strategy. Finally, as we explain below, the users will remain unaware of the platform's ability to estimate the initial glossiness state in the post-AI environment. This could be interpreted as a behavioral bias, but our favorite interpretation is that most users do not understand the informational capabilities of digital platforms in today's economy.

## 2.2 Utilities

If the platform offers item  $i$  at time  $t$  at price  $p_{i,t}$ , the user's instantaneous utility from its purchase is

$$\theta_i - p_{i,t},$$

and the platform's instantaneous utility, if the user purchases, is  $p_{i,t}$ . We let  $b_{i,t} \in \{0, 1\}$  denote the purchase decision of the user if product  $i$  is offered to her. We also assume that both the platform and the user discount the future at rate  $r$  for some  $r \in \mathbb{R}_+$  and update their beliefs using Bayes's rule. The platform's expected profit (utility) at time 0 is given by

$$\mathbb{E}_{P,0} \left[ \int_0^\infty r e^{-rt} \sum_{i=1}^n x_{i,t} p_{i,t} b_{i,t} dt \right],$$

and the user's expected utility at time 0 is

$$\mathbb{E}_{U,0} \left[ \int_0^\infty r e^{-rt} \sum_{i=1}^n x_{i,t} b_{i,t} (\theta_i - p_{i,t}) dt \right],$$

where  $x_{i,t}$  and  $p_{i,t}$  are adapted to the natural filtration of the platform at time  $t$ , denoted by  $\mathcal{F}_{P,t}$  and  $b_{i,t}$  is adapted to the natural filtration of the user's information set at time  $t$ , denoted by  $\mathcal{F}_{U,t}$ .



## 2.3 Equilibrium Concept

In the pre-AI environment, when the platform does not have an information advantage, neither the user nor the platform knows or receives signals about the initial glossiness  $\alpha_{i,0}$  for  $i \in \mathcal{N}$ . Therefore, the platform's decision at time  $t$ ,  $\{(x_{i,t}, p_{i,t})\}_{i \in \mathcal{N}}$  does not contain any extra information about the quality of the product or the glossiness state. In this environment, both the platform and the user are fully Bayesian, and we adopt *Markov Perfect Equilibrium* (MPE) as the solution concept. This means that the platform's and the user's decisions at each time  $t$  are dependent on the game's history through their respective beliefs about the quality of each product and their strategies are conditioned on the payoff-relevant states in this game, which are summarized by these beliefs.

In the post-AI era, the platform possesses superior knowledge about the initial glossiness state  $\alpha_{i,0}$ . Although, the platform does not receive additional information about this variable after the initial date. Nevertheless, the knowledge of  $\alpha_{i,0}$  enables it to estimate  $\alpha_{i,t}$  more accurately than the user for any  $t$ . In this environment, we again use the notion of MPE, except that in this case each player evaluates payoffs according to their beliefs, which differ because they have access to different information. The platform's decision at time  $t$ ,  $\{(x_{i,t}, p_{i,t})\}_{i \in \mathcal{N}}$  contains additional information about the quality of the product and the glossiness state, and the only behavioral aspect of our model is that the user does not recognize this informational superiority and does not update her beliefs on the basis of the platform's offers. This behavioral assumption is in line with the equilibrium concepts proposed in [Eyster and Rabin \[2005\]](#), [Esponda \[2008\]](#), and [Esponda and Pouzo \[2016\]](#). Our equilibrium concept can thus be viewed as the equivalent of the Berk-Nash equilibrium notion of [Esponda and Pouzo \[2016\]](#).<sup>4</sup>

## 3 Belief Dynamics and Equilibrium Characterization

We next determine the belief dynamics in the pre-AI and post-AI environments and then provide a preliminary characterization of equilibrium.

### 3.1 Belief Dynamics

**Pre-AI environment:** At any given time  $t$ , if bad news has been received for a product  $i \in \mathcal{N}$ , both the user and the platform believe with certainty that the product is of low quality. Otherwise, their Bayesian belief will be given by the posterior probability:

$$\mu_{i,t} = \mathbb{P}[\theta_i = 1 \mid \text{NBN}_{i,t}].$$

In forming this expectation, both the user and the platform recognize that the reason why no bad news may have been received so far is that the product is glossy ( $\alpha_{i,t} = 1$ ). For this reason, they also have to

---

<sup>4</sup>(Non-behavioral) Bayesian Nash equilibrium would involve the user drawing inferences from the product offering and pricing decisions of the platform. It is straightforward to show that one (simple) equilibrium in this case would be that the user interprets any deviation from the offering of the highest-belief product as a signal of manipulation, and thus the platform would be unable to use any of its post-AI information.

update the probability that the product is glossy, which is

$$\lambda_{i,t} = \mathbb{P}[\alpha_{i,t} = 1 \mid \theta_i = 0, \text{NBN}_{i,t}].$$

We next characterize the dynamics of the user's belief in the pre-AI environment, which are identical to the platform's belief.<sup>5</sup>

**Lemma 1.** *For any  $\{x_{i,t}, p_{i,t}, b_{i,t}\}_{i \in \mathcal{N}, t}$ , in the pre-AI environment, the beliefs of the user and the platform evolve as follows:*

$$\begin{aligned} d\mu_{i,t} &= x_{i,t} b_{i,t} \mu_{i,t} (1 - \mu_{i,t}) (1 - \lambda_{i,t}) \gamma dt && \text{with initialization } \mu_{i,0} \\ d\lambda_{i,t} &= x_{i,t} b_{i,t} \lambda_{i,t} ((1 - \lambda_{i,t}) \gamma - \rho) dt && \text{with initialization } \lambda_{i,0} = \lambda. \end{aligned}$$

Moreover, the probability of receiving bad news for product  $i$  at time  $t$ , given that bad news has not arrived during  $[0, t)$ , is:

$$\mathbb{P}[I_{i,t+dt} = 0 \mid I_{i,t} = 1] = x_{i,t} b_{i,t} (1 - \mu_{i,t}) (1 - \lambda_{i,t}) \gamma dt.$$

The evolution of these beliefs highlights that the user is Bayesian and knows about the (existence of) glossiness state  $\alpha_{i,t}$ . In particular, the updating equation shows that when  $\lambda_{i,t}$  is smaller,  $\mu_{i,t}$  moves faster because it is recognized that no bad news is more likely to correspond to a high-quality product in this case (whereas with higher  $\lambda_{i,t}$  it may be the product's glossiness that hides bad news).

**Post-AI environment:** Again, at time  $t$ , if bad news about product  $i \in \mathcal{N}$  has arrived, then the belief of the platform and the user is that the product is of low quality with probability 1. Otherwise, the user's posterior belief is

$$\mu_{i,t} = \mathbb{P}[\theta_i = 1 \mid \text{NBN}_{i,t}].$$

In the post-AI environment, the main difference is that the platform knows the initial glossiness state  $\alpha_{i,0}$  of all products and therefore has a different belief about product qualities. We let

$$\mu_{i,t}^{(P)} = \mathbb{P}[\theta_i = 1 \mid \text{NBN}_{i,t}, \alpha_{i,0}]$$

denote the platform's belief (hence, the superscript " $P$ ") about product  $i$ 's quality. Notice that if the initial glossiness state is  $\alpha_{i,0} = 1$ , then the platform knows the product is of low quality with probability 1, and hence  $\mu_{i,t}^{(P)} = 0$ . If the initial glossiness state is  $\alpha_{i,0} = 0$ , then the platform's belief evolves as we characterize in the next lemma. We also denote by

$$\lambda_{i,t}^{(P)} = \mathbb{P}[\alpha_{i,t} = 1 \mid \text{NBN}_{i,t}, \alpha_{i,0}]$$

the platform's belief about the product being in state  $\alpha_{i,t} = 1$  given the initial glossiness state and that no bad news has arrived at time  $t$ . Notice that if the initial glossiness state is  $\alpha_{i,0} = 0$ , then the state

---

<sup>5</sup>Throughout, we follow the standard practice of dropping terms of order  $o(dt)$ .

remains at  $\alpha_{i,t} = 0$  at time  $t$ , and therefore  $\lambda_{i,t}^{(P)} = 0$ . If the initial glossiness state is  $\alpha_{i,0} = 1$ ,  $\lambda_{i,t}^{(P)}$  evolves as we show next.

**Lemma 2.** *For any  $\{x_{i,t}, p_{i,t}, b_{i,t}\}_{i \in \mathcal{N}, t \in \mathbb{R}_+}$ , in the post-AI environment, the user's beliefs are the same as in Lemma 1, and the platform's beliefs evolve as follows:*

- If  $\alpha_{i,0} = 1$ , then

$$\mu_{i,t}^{(P)} = 0 \text{ and } d\lambda_{i,t}^{(P)} = x_{i,t} b_{i,t} \lambda_{i,t}^{(P)} \left( (1 - \lambda_{i,t}^{(P)}) \gamma - \rho \right) dt \text{ with initialization } \lambda_{i,0}^{(P)} = 1.$$

- If  $\alpha_{i,0} = 0$ , then

$$d\mu_{i,t}^{(P)} = x_{i,t} b_{i,t} \mu_{i,t}^{(P)} (1 - \mu_{i,t}^{(P)}) \gamma dt \text{ with initialization } \mu_{i,0}^{(P)} = \frac{\mu_{i,0}}{\mu_{i,0} + (1 - \mu_{i,0})(1 - \lambda)} \text{ and } \lambda_{i,t}^{(P)} = 0.$$

Moreover, from the platform's perspective, the probability of receiving bad news for product  $i$  at time  $t$ , given that bad news has not arrived during  $[0, t)$ , is:

$$\begin{aligned} \mathbb{P}[I_{i,t+dt} = 0 \mid I_{i,t} = 1, \alpha_{i,0} = 0] &= x_{i,t} b_{i,t} \left( 1 - \mu_{i,t}^{(P)} \right) \gamma dt \quad \text{for } \alpha_{i,0} = 0 \\ \mathbb{P}[I_{i,t+dt} = 0 \mid I_{i,t} = 1, \alpha_{i,0} = 1] &= x_{i,t} b_{i,t} \left( 1 - \lambda_{i,t}^{(P)} \right) \gamma dt \quad \text{for } \alpha_{i,0} = 1. \end{aligned}$$

### 3.2 Preliminary Equilibrium Characterization

We first prove a simple lemma that characterizes the platform's equilibrium pricing strategy and the user's equilibrium purchasing strategy.

**Lemma 3.** *The time- $t$  utility of the user from purchasing product  $i$  for which no bad news has arrived ( $I_{i,t} = 1$ ) is*

$$\mathbb{E}[\theta_i \mid \text{NBN}_{i,t}] - p_{i,t} = \mu_{i,t} - p_{i,t}.$$

Thus, the time- $t$  equilibrium price of product  $i$  for which no bad news has arrived is

$$p_{i,t} = \mathbb{E}[\theta_i \mid \text{NBN}_{i,t}] = \mu_{i,t}.$$

For a product  $i$  at time  $t$  for which bad news has arrived ( $I_{i,t} = 0$ ), the equilibrium price is zero.

*Proof of Lemma 3:* The platform's equilibrium pricing strategy is to offer price  $p_{i,t} = \mu_{i,t}$  when  $x_{i,t} = 1$  and the user's equilibrium strategy is to purchase that product. This can be seen by using the one-stage deviation principle: if the platform deviates and offers a lower price, then the continuation game does not change, and the platform would make strictly lower profits. If the platform offers a higher price, then not purchasing is a better response for the user (otherwise, she obtains a negative utility from her perspective), leading to a lower platform's payoff. To establish that offering a price equal to the user's belief is the only MPE, note that the user accepts a price if it is weakly smaller than her belief and, therefore, the platform's pricing in equilibrium is always the user's belief. ■

We next provide expressions for platform and user utilities in the pre-AI and post-AI environments.

**Pre-AI environment:** In the pre-AI environment, the platform's equilibrium strategy is a solution to a stochastic dynamic optimization problem whose state is given by beliefs about the product's quality. Since both the platform and the user share the same information in this setting, we let  $\{\mu_i\}_{i \in \mathcal{N}}$  denote their initial belief. With this notation, the platform's offering strategy in equilibrium denoted by  $\{x_{i,t}^{\text{pre-AI}}\}_{i \in \mathcal{N}, t \in \mathbb{R}_+}$  solves

$$\begin{aligned} \Pi^{\text{pre-AI}}(\{\mu_i\}_{i=1}^n) &= \max_{\{x_{i,t}\}_{i \in \mathcal{N}, t \in \mathbb{R}_+}} \mathbb{E}_{P,0} \left[ \int_0^\infty r e^{-rt} \sum_{i=1}^n x_{i,t} \mu_{i,t} dt \right] \\ d\mu_{i,t} &= I_{i,t} x_{i,t} \mu_{i,t} (1 - \mu_{i,t}) (1 - \lambda_{i,t}) \gamma dt \quad \mu_{i,0} = \mu_i \\ d\lambda_{i,t} &= I_{i,t} x_{i,t} \lambda_{i,t} ((1 - \lambda_{i,t}) \gamma - \rho) dt \quad \lambda_{i,0} = \lambda \\ \mu_{i,t} &= 0 \quad \text{if } I_{i,t} = 0 \\ \mathbb{P}[I_{i,t+dt} = 0 \mid I_{i,t} = 1] &= x_{i,t} (1 - \mu_{i,t}) (1 - \lambda_{i,t}) \gamma dt \quad I_{i,0} = 1, \end{aligned} \quad (2)$$

and the platform's pricing strategy in equilibrium is  $p_{i,t}^{\text{pre-AI}} = \mu_{i,t}$  (see Lemma 3). Here, the subscript  $P, 0$  denotes that the expectation is with respect to the platform's information at time 0, and  $\Pi^{\text{pre-AI}}(\{\mu_i\}_{i=1}^n)$  denotes the platform's payoff starting from user's belief  $\{\mu_i\}_{i=1}^n$ . Notice that the user's purchasing decision  $b_{i,t}$  does not appear in the above expression. This is because, as we proved in Lemma 3, the platform always offers one of the products with a price equal to the user's belief, and the user's equilibrium strategy is always to purchase that product.

We evaluate the user's utility according to the expectations of "true events". Given the platform's equilibrium strategy, the expectation of the user's utility is

$$U^{\text{pre-AI}}(\{\mu_i\}_{i=1}^n) = \mathbb{E}_0 \left[ \int_0^\infty r e^{-rt} \sum_{i=1}^n x_{i,t}^{\text{pre-AI}} (\theta_i - \mu_{i,t}) dt \right]. \quad (3)$$

The expected utilitarian welfare is the sum of the user's utility given in (3) and the platform's payoff given in (2):

$$W^{\text{pre-AI}}(\{\mu_i\}_{i=1}^n) = \mathbb{E}_0 \left[ \int_0^\infty r e^{-rt} \sum_{i=1}^n x_{i,t}^{\text{pre-AI}} \theta_i dt \right].$$

**Post-AI environment:** In contrast with the pre-AI setting, here, the platform knows the initial glossiness state of the product and therefore has a different belief from the user. Starting from any user's belief  $\{\mu_i\}_{i \in \mathcal{N}}$  and the initial glossiness state  $\{\alpha_i\}_{i \in \mathcal{N}}$ , the platform's equilibrium offering strategy  $\{x_{i,t}^{\text{post-AI}}\}_{i \in \mathcal{N}, t \in \mathbb{R}_+}$  solves

$$\begin{aligned} \Pi^{\text{post-AI}}(\{\mu_i, \alpha_i\}_{i=1}^n) &= \max_{\{x_{i,t}\}_{i \in \mathcal{N}, t \in \mathbb{R}_+}} \mathbb{E}_{P,0} \left[ \int_0^\infty r e^{-rt} \sum_{i=1}^n x_{i,t} \mu_{i,t} dt \right] \\ d\mu_{i,t} &= I_{i,t} x_{i,t} \mu_{i,t} (1 - \mu_{i,t}) (1 - \lambda_{i,t}) \gamma dt \quad \mu_{i,0} = \mu_i \end{aligned} \quad (4)$$

$$\begin{aligned}
d\lambda_{i,t} &= I_{i,t}x_{i,t}\lambda_{i,t}((1-\lambda_{i,t})\gamma - \rho) dt \quad \lambda_{i,0} = \lambda \\
d\mu_{i,t}^{(P)} &= I_{i,t}x_{i,t}\mu_{i,t}^{(P)}(1-\mu_{i,t}^{(P)})\gamma dt \quad \text{and} \quad \lambda_{i,t}^{(P)} = 0 \text{ if } \alpha_i = 0 \\
\mu_{i,t}^{(P)} &= 0 \quad \text{and} \quad d\lambda_{i,t}^{(P)} = I_{i,t}x_{i,t}\lambda_{i,t}^{(P)}((1-\lambda_{i,t}^{(P)})\gamma - \rho) dt \text{ if } \alpha_i = 1 \\
\mu_{i,t} &= \mu_{i,t}^{(P)} = 0 \quad \text{if } I_{i,t} = 0 \\
\mathbb{P}[I_{i,t+dt} = 0 \mid I_{i,t} = 1, \alpha_i = 0] &= x_{i,t}b_{i,t} \left(1 - \mu_{i,t}^{(P)}\right) \gamma dt \\
\mathbb{P}[I_{i,t+dt} = 0 \mid I_{i,t} = 1, \alpha_i = 1] &= x_{i,t}b_{i,t} \left(1 - \lambda_{i,t}^{(P)}\right) \gamma dt \text{ with } I_{i,0} = 1,
\end{aligned}$$

and the platform's pricing strategy in equilibrium is  $p_{i,t}^{\text{post-AI}} = \mu_{i,t}$  (see Lemma 3). Notice that  $\mu_{i,0}^{(P)}$  and  $\lambda_{i,0}^{(P)}$ , depending on the initial glossiness state is characterized in Lemma 2. The expectation of the user's utility is then

$$U^{\text{post-AI}}(\{\mu_i, \alpha_i\}_{i=1}^n) = \mathbb{E}_0 \left[ \int_0^\infty r e^{-rt} \sum_{i=1}^n x_{i,t}^{\text{post-AI}} (\theta_i - \mu_{i,t}) dt \right]. \quad (5)$$

Similar to the pre-AI environment, the expected utilitarian welfare is the sum of user utility in (5) and platform profits in (4). With these expressions, we next provide the equilibrium characterizations.

## 4 Equilibrium Characterization in the Pre-AI Environment

Our first theorem characterizes the platform's equilibrium decision in the pre-AI environment.

**Theorem 1.** *Without loss of generality, let  $\mu_{1,0} \geq \dots \geq \mu_{n,0}$  be the initial beliefs for the  $n$  products. Then, in the pre-AI environment, the platform's equilibrium decision is to offer product 1 until bad news occurs for this product (i.e.,  $x_{1,t}^{\text{pre-AI}} = 1$  for all  $t \in [0, \tau_1]$  where  $\tau_1 \in \mathbb{R}_+ \cup \{\infty\}$  is the stochastic time at which bad news for product 1 occurs), and then to offer product 2 and so on.*

*Proof of Theorem 1:* Here, we provide the main steps of the proof, relegating the proof of lemmas to the Appendix. The state of each product is a pair  $(\mu_i, \lambda_i)$ , where  $\mu_i$  is the user and platform belief about the product's quality and  $\lambda_i$  is the probability of having a glossiness  $\alpha = 1$ , given no bad news has occurred. For any product  $i \in \mathcal{N}$ , we introduce the Gittins index:

$$M(\mu_i, \lambda_i) = \sup_{\tau \geq 0} \frac{\mathbb{E}_0 \left[ \int_0^\tau e^{-rt} \mu_{i,t} dt \right]}{\mathbb{E}_0 \left[ \int_0^\tau e^{-rt} dt \right]}$$

with the same evolution as the ones in (2) and where the supremum is over all stopping times. Using the Gittins index theorem (see, e.g., [Gittins et al., 2011, Chapter 2]), which asserts that the platform always chooses the product with the highest Gittins index. We make use of two key properties of the Gittins index in our context: (i) products with higher beliefs have initially higher Gittins indices (Lemma A4), and (ii) in the absence of bad news, the Gittins index for a product increases (Lemma A5). Therefore, the platform begins by offering the product with the highest belief. As long as there is no bad news for this product, its Gittins index continues to grow, maintaining its status as the product with the maximum

Gittins index. Thus, the platform continues to offer this product until bad news occurs, at which point it switches to the product with the second highest belief, and so on. ■

## 5 Equilibrium Characterization in the Post-AI Environment

Here, we determine the platform's equilibrium decision in the post-AI environment.

**Theorem 2.** *Without loss of any generality, let  $\mu_{i_1,0} \geq \dots \geq \mu_{i_{n_0},0}$  be the initial beliefs for products with initial glossiness state  $\alpha_{i,0} = 0$  and  $\mu_{j_1,0} \geq \dots \geq \mu_{j_{n_1},0}$  be the initial beliefs for the products with initial glossiness state  $\alpha_{i,0} = 1$ . In the post-AI environment, the platform's equilibrium strategy involves first offering either  $i_1$  or  $j_1$ . If the platform first offers  $i_1$  [respectively,  $j_1$ ], the platform continues to offer this product until bad news occurs. The platform then offers either  $i_2$  or  $j_1$  [respectively, offers either  $i_1$  or  $j_2$ ] and so on.*

*Proof of Theorem 2:* Again, we provide the main steps of the proof, relegating the detail to the appendix. In the post-AI environment, the state of each product is given by the tuple  $(\mu_i, \lambda_i, \mu_i^{(P)}, \lambda_i^{(P)})$ , where  $\mu_i$  is the user belief about product's quality,  $\lambda_i$  is the probability of having a glossiness state  $\alpha = 1$  given no bad news has occurred and  $(\mu_i^{(P)}, \lambda_i^{(P)})$  are the same quantities but from the platform's perspective. Similar to the analysis of pre-AI, for any product  $i \in \mathcal{N}$ , we introduce the Gittins index as

$$M(\mu_i, \lambda_i, \mu_i^{(P)}, \lambda_i^{(P)}) = \sup_{\tau \geq 0} \frac{\mathbb{E}_0 \left[ \int_0^\tau e^{-rt} \mu_{i,t} dt \right]}{\mathbb{E}_0 \left[ \int_0^\tau e^{-rt} dt \right]} \quad (6)$$

with the same evolution as the ones in (4) and where the supremum is over all stopping times. Using a similar argument to that of Theorem 1, the platform starts by offering the product with the highest user belief either among those with the initial glossiness state  $\alpha = 1$  or among those with the initial glossiness state of 0. If bad news does not occur for this product, the Gittins index of this arm weakly increases while the other Gittins indices remain unchanged. Therefore, it is optimal to keep offering this product until bad news occurs. ■

Theorem 2 characterizes the structure of the platform's equilibrium decision, but it does not pin down the comparison between a product with  $\alpha_{i,0} = 0$  and another product with  $\alpha_{i',0} = 1$ , which we provide next.

In the post-AI environment, the platform no longer offers the product with the highest belief and may opt for a glossy product ( $\alpha_{i,0} = 1$ ) with a lower belief, rather than a higher belief non-glossy product ( $\alpha_{i,0} = 0$ ). We next formalize this observation.

**Proposition 1 (Manipulation Effect).** *Consider two products  $i, i' \in \mathcal{N}$  in the post-AI environment with  $\alpha_{i',0} = 1$  and  $\alpha_{i,0} = 0$ . There exists a function  $\mu_{i,0} \mapsto f(\mu_{i,0}; \lambda, \gamma, r)$  parameterized by  $\lambda, \gamma, r$  that is everywhere below  $\mu_{i,0}$ , and  $\underline{\rho}$ , such that for  $\rho \leq \underline{\rho}$  when*

$$\mu_{i',0} > f(\mu_{i,0}; \lambda, \gamma, r),$$

*the platform prefers to offer product  $i'$  rather than offering product  $i$ .*

We call this phenomenon the *manipulation effect* because the platform is offering a low-quality product with a lower user belief. To understand this effect and why it requires  $\rho$  to be small, observe that there are two forces that determine the platform's choice: (i) when  $\alpha_{i',0} = 1$ , the platform knows that the quality is low, and therefore, once  $\alpha_{i',t}$  becomes 0, the user will receive bad news at the rate  $\gamma$ . This force incentivizes the platform to not offer product  $i'$ ; (ii) when  $\alpha_{i,0} = 1$ , the platform knows that the user's belief increases for a while because bad news will not arrive. This force incentivizes the platform to offer product  $i'$ . For sufficiently small  $\rho$ , the second force dominates, because the glossy state is very persistent and thus bad news for product  $i'$  will not arrive for a long time. In this case, the platform prefers to offer product  $i'$  rather than product  $i$ .

We next show that the opposite phenomenon to the one presented in Proposition 1, which we call the *helpfulness effect*, arises when  $\rho$  is sufficiently large.

**Proposition 2 (Helpfulness Effect).** *Consider two products  $i, i' \in \mathcal{N}$  in the post-AI environment with  $\alpha_{i',0} = 1$  and  $\alpha_{i,0} = 0$ . There exists a function  $\mu_{i',0} \mapsto g(\mu_{i',0}; \lambda, \gamma, r)$  parametrized by  $\lambda, \gamma, r$  that is everywhere below  $\mu_{i',0}$ , and  $\bar{\rho}$ , such that for  $\rho \geq \bar{\rho}$  when*

$$\mu_{i,0} > g(\mu_{i',0}; \lambda, \gamma, r),$$

*the platform prefers to offer product  $i$  relative to offering product  $i'$ .*

In this configuration, the platform prefers not to offer a low-quality with a moderately higher user belief than an alternative product with unknown quality with a lower user belief. This is helpful to the user as it avoids the low-quality product. The platform recognizes the low-quality product simply from the fact that it has an initial glossiness state  $\alpha_{i,0} = 1$ , and deploys this information for the user's benefit. To see the intuition, recall that when  $\rho$  is large, the platform understands that the glossiness of the product will not last, and bad news will start arriving for product  $i'$  (which has initial glossiness state  $\alpha_{i',0} = 1$ ). Therefore, even if  $\mu_{i',0} > \mu_{i,0}$ , so long as the belief  $\mu_{i,0}$  is not too small (ensured by the condition  $\mu_{i,0} > g(\mu_{i',0}; \lambda, \gamma, r)$ ), the platform prefers offering product  $i$  instead of product  $i'$ . We conclude this section by noting that both functions characterized in Propositions 1 and 2, can be written explicitly in terms of the primitives  $(\mu, \lambda, \gamma, r)$ .

## 6 Who Benefits from the Platform's Superior Information?

In this section, we discuss the welfare implications of the platform's superior information and then explore the implications of expanding the set of products available to the platform.

### 6.1 When Users Benefit from the Platform's Superior Information

Our next theorem establishes that for sufficiently large  $\rho$ , the platform's information advantage in the post-AI environment always benefits the user and the platform.



**Theorem 3.** Suppose the initial beliefs  $\{\mu_i\}_{i=1}^n$  are i.i.d. and uniform over  $[0, 1]$ . For any  $r, \gamma, \lambda$ , there exists  $\rho_h$  such that for  $\rho \geq \rho_h$  we have

$$\begin{aligned}\mathbb{E} [\Pi^{\text{post-AI}}(\{\mu_i, \alpha_i\}_{i=1}^n)] &> \mathbb{E} [\Pi^{\text{pre-AI}}(\{\mu_i\}_{i=1}^n)] \\ \mathbb{E} [U^{\text{post-AI}}(\{\mu_i, \alpha_i\}_{i=1}^n)] &> \mathbb{E} [U^{\text{pre-AI}}(\{\mu_i\}_{i=1}^n)] \\ \mathbb{E} [W^{\text{post-AI}}(\{\mu_i, \alpha_i\}_{i=1}^n)] &> \mathbb{E} [W^{\text{pre-AI}}(\{\mu_i\}_{i=1}^n)] ,\end{aligned}$$

where the expectations are over  $\mu_i \sim \text{Unif}[0, 1]$ ,  $\theta_i \sim \text{Bern}(\mu_i)$ , and  $\alpha_i$  drawn according to (1) for all  $i \in \mathcal{N}$ .

The informational advantage of the platform always increases its own profits, as it enables the platform to modify the user's behavior, and also benefits the user when  $\rho$  is large, because in this case the helpfulness effect dominates the manipulation effect. This theorem confirms what might be viewed as the conventional wisdom in the literature: more data enables better allocation of products and, therefore, benefits users.

## 6.2 When Behavioral Manipulation Harms Users

Here, we show that in the post-AI environment, the user's utility decreases because of behavioral manipulation — since the manipulation effect dominates the helpfulness effect.

**Theorem 4.** Suppose the initial beliefs  $\{\mu_i\}_{i=1}^n$  are i.i.d. and uniform over  $[0, 1]$ . For any  $r, \gamma, \lambda$ , there exists  $\rho_l$  such that for  $\rho \leq \rho_l$  we have

$$\begin{aligned}\mathbb{E} [\Pi^{\text{post-AI}}(\{\mu_i, \alpha_i\}_{i=1}^n)] &> \mathbb{E} [\Pi^{\text{pre-AI}}(\{\mu_i\}_{i=1}^n)] \\ \mathbb{E} [U^{\text{post-AI}}(\{\mu_i, \alpha_i\}_{i=1}^n)] &< \mathbb{E} [U^{\text{pre-AI}}(\{\mu_i\}_{i=1}^n)] \\ \mathbb{E} [W^{\text{post-AI}}(\{\mu_i, \alpha_i\}_{i=1}^n)] &< \mathbb{E} [W^{\text{pre-AI}}(\{\mu_i\}_{i=1}^n)] .\end{aligned}$$

In this case with low  $\rho$ , the platform's informational advantage enables it to engage in behavioral manipulation: the user is pushed towards products that have initial glossiness state  $\alpha_{i,0} = 1$ . Because these products do not generate bad news in the short run, the user's belief will become more positive for a while, and this will enable the platform to charge higher prices and increase its profits. Because glossy products are low quality, such behavioral manipulation is bad for user utility and utilitarian welfare. As we saw, when  $\rho$  is large, the platform expects the glossiness of the production to wear off quickly, and thus it is not worthwhile to push the user towards glossy products. But from a welfare point of view, it is more costly to have users consume glossy products when  $\rho$  is small because they will not discover for quite a while that the product is actually not high-quality. It is this feature of behavioral manipulation that reduces user utility and utilitarian welfare.

Finally, we observe that the uniform assumption for the initial beliefs in Theorems 3 and 4 is without loss of generality: the comparisons hold with weak inequalities for any fixed initial beliefs, and the strict inequalities hold whenever the platform's equilibrium strategies in the two environments are different.



### 6.3 Big Data Double Whammy: More Products Negatively Impact User Welfare

The availability of big data provides platforms with valuable insights into predictable patterns of user behavior, which can be leveraged for behavioral manipulation, as we have established thus far. Moreover, the same advances in AI also enable digital platforms to expand the range of products and services they offer. Next, we demonstrate that this combination of greater choice and more platform information may be particularly pernicious — as the number of products increases, the potential for behavioral manipulation increases as well. This result highlights that multiple aspects of the new capabilities of digital platforms closely interact in affecting user welfare.

**Theorem 5.** *Suppose the initial beliefs  $\{\mu_i\}_{i=1}^{n+1}$  are i.i.d. and uniform over  $[\frac{1}{2} - \Delta, \frac{1}{2} + \Delta]$ . For any  $r, \gamma, \lambda$ , there exist  $\bar{\Delta}$  and  $\tilde{\rho}_l$  such that for  $\Delta \leq \bar{\Delta}$ ,  $\rho \leq \tilde{\rho}_l$ , and  $n \geq \lceil 1 - \frac{\log \lambda}{\log(2-\lambda)/(1-\lambda)} \rceil$  in the post-AI environment we have:*

$$\begin{aligned}\mathbb{E} [\Pi^{\text{post-AI}}(\{\mu_i, \alpha_i\}_{i=1}^{n+1})] &> \mathbb{E} [\Pi^{\text{post-AI}}(\{\mu_i, \alpha_i\}_{i=1}^n)] \\ \mathbb{E} [U^{\text{post-AI}}(\{\mu_i, \alpha_i\}_{i=1}^{n+1})] &< \mathbb{E} [U^{\text{post-AI}}(\{\mu_i, \alpha_i\}_{i=1}^n)] \\ \mathbb{E} [W^{\text{post-AI}}(\{\mu_i, \alpha_i\}_{i=1}^{n+1})] &< \mathbb{E} [W^{\text{post-AI}}(\{\mu_i, \alpha_i\}_{i=1}^n)].\end{aligned}$$

The assumption that  $\rho \leq \tilde{\rho}_l$  is adopted for the same reasons as before — this is the range where the platform’s superior information can be costly to the user. The assumptions That  $\Delta \leq \bar{\Delta}$  and  $n$  is large are added are adopted to ensure that the initial beliefs of the user are not very informative about product quality and that the platform has enough products to engage in behavioral manipulation. This combination then leads to a configuration where the presence of more products creates greater opportunities for the platform to select initially glossy items (with initial beliefs that are sufficiently close to the beliefs of non-glossy products). An alternative way of expressing this intuition is that, while in standard choice theory, greater choices are good for the consumer, in the presence of behavioral manipulation, they may be bad precisely because they enable the platform to engage in more intense behavioral manipulation.

## 7 Conclusion

This paper has argued that the vast amounts of data collected by online platforms may also enable behavioral manipulation, harming the users. We view our paper as a first step in the systematic analysis of product choice and consumer experimentation in online markets in which platforms have extensive and growing information about user preferences (and biases). There are several interesting questions for future research, which we list briefly:

- Can behavioral manipulation persist in the long-run even as users become used to the big data environment?
- Can AI tools be used for correcting, rather than exploiting, behavioral biases of users?

- How does dynamic learning by platforms (rather than estimating product characteristics perfectly, as we have assumed) affect the scope for behavioral manipulation?
- How does the presence of users with different levels of sophistication affect platform behavior?
- How does competition between platforms affect behavioral manipulation?
- Last but not least, what types of regulations can be useful for lessening the negative effects of behavioral manipulation?

## References

- D. Acemoglu, A. Makhdoumi, A. Malekian, and A. Ozdaglar. Learning from reviews: The selection effect and the speed of learning. *Econometrica*, 90(6):2857–2899, 2022a.
- D. Acemoglu, A. Makhdoumi, A. Malekian, and A. Ozdaglar. Too much data: Prices and inefficiencies in data markets. *American Economic Journal: Microeconomics*, 14(4):218–256, 2022b.
- A. Acquisti and H. R. Varian. Conditioning prices on purchase history. *Marketing Science*, 24(3):367–381, 2005.
- A. Acquisti, L. Brandimarte, and G. Loewenstein. Privacy and human behavior in the age of information. *Science*, 347(6221):509–514, 2015.
- A. Acquisti, C. Taylor, and L. Wagman. The economics of privacy. *Journal of Economic Literature*, 54(2):442–92, 2016.
- A. R. Admati and P. Pfleiderer. A monopolistic market for information. *Journal of Economic Theory*, 39(2):400–438, 1986.
- A. Agrawal, J. Gans, and A. Goldfarb. *Prediction machines: the simple economics of artificial intelligence*. Harvard Business Press, 2018.
- N. Alon and J. H. Spencer. *The probabilistic method*. John Wiley & Sons, 2016.
- J. J. Anton and D. A. Yao. The sale of ideas: Strategic disclosure, property rights, and contracting. *Review of Economic Studies*, 69(3):513–531, 2002.
- J. Begenau, M. Farboodi, and L. Veldkamp. Big data in finance and the growth of large firms. *Journal of Monetary Economics*, 97:71–87, 2018.
- D. Bergemann and A. Bonatti. Selling cookies. *American Economic Journal: Microeconomics*, 7(3):259–94, 2015.
- D. Bergemann and A. Bonatti. Markets for information: An introduction. *Annual Review of Economics*, 11:85–107, 2019a.
- D. Bergemann and A. Bonatti. Markets for information: An introduction. *Annual Review of Economics*, 11, 2019b.
- D. Bergemann, B. Brooks, and S. Morris. The limits of price discrimination. *American Economic Review*, 105(3):921–57, 2015.
- D. Bergemann, A. Bonatti, and A. Smolin. The design and price of information. *American Economic Review*, 108(1):1–48, 2018.
- D. A. Berry and B. Fristedt. Bandit problems: sequential allocation of experiments (monographs on statistics and applied probability). London: Chapman and Hall, 5(71-87):7–7, 1985.
- P. Bolton and C. Harris. Strategic experimentation. *Econometrica*, 67(2):349–374, 1999.
- A. Bonatti and G. Cisternas. Consumer scores and price discrimination. *The Review of Economic Studies*, 87(2):750–791, 2020.

- R. Calo. Digital market manipulation. *Geo. Wash. L. Rev.*, 82:995, 2014.
- R. A. Calvo, D. Peters, K. Vold, and R. M. Ryan. Supporting human autonomy in ai systems: A framework for ethical enquiry. *Ethics of digital well-being: A multidisciplinary approach*, pages 31–54, 2020.
- Y.-K. Che and J. Hörner. Recommender systems as mechanisms for social learning. *The Quarterly Journal of Economics*, 133(2):871–925, 2018.
- J. P. Choi, D.-S. Jeon, and B.-C. Kim. Privacy and personal data collection with information externalities. *Journal of Public Economics*, 173:113–124, 2019.
- M. Ekmekci. Sustainable reputations with rating systems. *Journal of Economic Theory*, 146(2):479–503, 2011.
- I. Esponda. Behavioral equilibrium in economies with adverse selection. *American Economic Review*, 98(4):1269–91, 2008.
- I. Esponda and D. Pouzo. Berk–nash equilibrium: A framework for modeling agents with misspecified models. *Econometrica*, 84(3):1093–1130, 2016.
- E. Eyster and M. Rabin. Cursed equilibrium. *Econometrica*, 73(5):1623–1672, 2005.
- I. P. Fainmesser, A. Galeotti, and R. Momot. Digital privacy. *Management Science*, 69(6):3157–3173, 2023.
- L. Felli and C. Harris. Learning, wage dynamics, and firm-specific human capital. *Journal of Political Economy*, 104(4):838–868, 1996.
- J. Gittins, K. Glazebrook, and R. Weber. *Multi-armed bandit allocation indices*. John Wiley & Sons, 2011.
- A. Goldfarb and C. Tucker. Online display advertising: Targeting and obtrusiveness. *Marketing Science*, 30(3):389–404, 2011.
- R. Gradwohl. Information sharing and privacy in networks. In *Proceedings of the 2017 ACM Conference on Economics and Computation*, pages 349–350, 2017.
- J. D. Hanson and D. A. Kysar. Taking behavioralism seriously: The problem of market manipulation. *NYUL Rev.*, 74:630, 1999.
- S. Hidir and N. Vellodi. Personalization, discrimination and information disclosure. Technical report, working paper, 2019.
- J. Horner and A. Skrzypacz. Selling information. *Journal of Political Economy*, 124(6):1515–1562, 2016.
- S. Ichihashi. Dynamic privacy choices. In *Proceedings of the 21st ACM Conference on Economics and Computation*, pages 539–540, 2020a.
- S. Ichihashi. Online privacy and information disclosure by consumers. *American Economic Review*, 110(2):569–95, 2020b.
- B. Jullien, Y. Lefouili, and M. H. Riordan. Privacy protection, security, and consumer retention. *Security, and Consumer Retention (June 1, 2020)*, 2020.
- G. Keller and S. Rady. Strategic experimentation with poisson bandits. *Theoretical Economics*, 5(2):275–311, 2010.

- G. Keller, S. Rady, and M. Cripps. Strategic experimentation with exponential bandits. *Econometrica*, 73(1):39–68, 2005.
- N. Köbis, J.-F. Bonnefon, and I. Rahwan. Bad machines corrupt good morals. *Nature Human Behaviour*, 5(6):679–685, 2021.
- N. V. Krylov. *Controlled diffusion processes*, volume 14. Springer Science & Business Media, 2008.
- T. C. Lin. The new market manipulation. *Emory LJ*, 66:1253, 2017.
- J. Z. Liu, M. Sockin, and W. Xiong. Data privacy and algorithmic inequality. *Working Paper*, 2023.
- K. Martin and H. Nissenbaum. Measuring privacy: an empirical test using context to expose confounding variables. *Science and Technology Law Review*, 18(1), Jan. 2017. doi: 10.7916/stlr.v18i1.4015. URL <https://journals.library.columbia.edu/index.php/stlr/article/view/4015>.
- R. Montes, W. Sand-Zantman, and T. Valletti. The value of personal information in online markets with endogenous privacy. *Management Science*, 65(3):1342–1362, 2019.
- F. Pasquale. *The Black Box Society: The Secret Algorithms That Control Money and Information*. Harvard University Press, 2015.
- H. Robbins. Some aspects of the sequential design of experiments. *Bulletin of the American Mathematical Society*, 58(5):527–535, 1952.
- M. Rothschild. A two-armed bandit theory of market pricing. *Journal of Economic Theory*, 9(2):185–202, 1974.
- D. Susser and V. Grimaldi. Measuring automated influence: Between empirical evidence and ethical values. In *Proceedings of the 2021 AAAI/ACM Conference on AI, Ethics, and Society*, pages 242–253, 2021.
- C. R. Taylor. Consumer privacy and the market for customer information. *RAND Journal of Economics*, pages 631–650, 2004.
- J. Tirole. Competition and the industrial challenge for the digital age. 2020.
- J. Tirole. Digital dystopia. *American Economic Review*, 111(6):2007–2048, 2021.
- H. Varian. Artificial intelligence, economics, and industrial organization. In *The economics of artificial intelligence: an agenda*, pages 399–419. University of Chicago Press, 2018.
- T. Z. Zarsky. Privacy and manipulation in the digital age. *Theoretical Inquiries in Law*, 20(1):157–188, 2019.
- S. Zuboff. *The age of surveillance capitalism: The fight for a human future at the new frontier of power: Barack Obama’s books of 2019*. Profile books, 2019.

## A1 Online Appendix

This Appendix includes the omitted proofs from the text.

### Proof of Lemma 1

If  $x_{i,t}b_{i,t} = 0$ , the user belief about  $\theta_i$  does not change. Next, we consider  $x_{i,t}b_{i,t} = 1$  and, for notational convenience, let us suppress subscript  $i$ . Then, by using Bayes' rule, the probability of  $\theta_i = 1$  is

$$\begin{aligned}
\mu_{t+dt} &= \frac{\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 1]\mathbb{P}[\theta = 1]}{\mathbb{P}[\text{NBN}_{t+dt}]} \\
&= \frac{\mathbb{P}[\text{NBN}_t \mid \theta = 1]\mathbb{P}[\text{NBN}_{t,t+dt} \mid \text{NBN}_t, \theta = 1]\mathbb{P}[\theta = 1]}{\mathbb{P}[\text{NBN}_t]\mathbb{P}[\text{NBN}_{t+dt} \mid \text{NBN}_t]} \\
&= \mu_t \frac{\mathbb{P}[\text{NBN}_{t+dt} \mid \text{NBN}_t, \theta = 1]}{\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 1, \text{NBN}_t]\mathbb{P}[\theta = 1 \mid \text{NBN}_t] + \mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t]\mathbb{P}[\theta = 0 \mid \text{NBN}_t]} \\
&\stackrel{(a)}{=} \frac{\mu_t}{\mu_t + (1 - \mu_t)\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t]} \\
&= \frac{\mu_t}{\mu_t + (1 - \mu_t)(\lambda_t + (1 - \lambda_t)(1 - \gamma dt))} \\
&= \frac{\mu_t}{1 - (1 - \mu_t)(1 - \lambda_t)\gamma dt} \\
&\stackrel{(b)}{=} \mu_t (1 + (1 - \mu_t)(1 - \lambda_t)\gamma dt)
\end{aligned} \tag{A1}$$

where (a) follows from

$$\begin{aligned}
\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t] &= \mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t, \alpha_t = 1]\mathbb{P}[\alpha_t = 1 \mid \theta = 0, \text{NBN}_t] \\
&\quad + \mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t, \alpha_t = 0]\mathbb{P}[\alpha_t = 0 \mid \theta = 0, \text{NBN}_t]
\end{aligned}$$

and (b) follows by using Taylor expansion of  $1/(1 - x)$  around 0 and dropping the terms of the order  $(dt)^2$ . From (A1), we then have

$$d\mu_t = \mu_{t+dt} - \mu_t = \mu_t(1 - \mu_t)(1 - \lambda_t)\gamma dt.$$

Again, by using Bayes' rule,

$$\begin{aligned}
\lambda_{t+dt} &= \mathbb{P}[\alpha_{t+dt} = 1 \mid \theta = 0, \text{NBN}_{t+dt}] \\
&= \frac{\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \alpha_{t+dt} = 1]\mathbb{P}[\alpha_{t+dt} = 1 \mid \theta = 0]}{\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0]} \\
&= \mathbb{P}[\alpha_t = 1 \mid \theta = 0]\mathbb{P}[\alpha_{t+dt} = 1 \mid \alpha_t = 1, \theta = 0] \\
&\quad \times \frac{\mathbb{P}[\text{NBN}_t \mid \theta = 0, \alpha_{t+dt} = 1]\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t, \alpha_{t+dt} = 1]}{\mathbb{P}[\text{NBN}_t \mid \theta = 0]\mathbb{P}[\text{NBN}_{t+dt} \mid \text{NBN}_t, \theta = 0]} \\
&= \frac{\lambda_t \mathbb{P}[\alpha_{t+dt} = 1 \mid \alpha_t = 1, \theta = 0]\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t, \alpha_{t+dt} = 1]}{\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t]}
\end{aligned}$$

$$\begin{aligned}
&\stackrel{(a)}{=} \frac{\lambda_t(1 - \rho dt)}{\lambda_t + (1 - \lambda_t)(1 - \gamma dt)} \\
&= \frac{\lambda_t(1 - \rho dt)}{1 - (1 - \lambda_t)\gamma dt} \\
&\stackrel{(b)}{=} \lambda_t - \lambda_t \rho dt + \lambda_t(1 - \lambda_t)\gamma dt
\end{aligned} \tag{A2}$$

where (a) follows from

$$\begin{aligned}
\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t] &= \mathbb{P}[\text{NBN}_{t+dt} \mid \alpha_t = 1, \theta = 0, \text{NBN}_t] \mathbb{P}[\alpha_t = 1 \mid \theta = 0, \text{NBN}_t] \\
&\quad + \mathbb{P}[\text{NBN}_{t+dt} \mid \alpha_t = 0, \theta = 0, \text{NBN}_t] \mathbb{P}[\alpha_t = 0 \mid \theta = 0, \text{NBN}_t]
\end{aligned}$$

and (b) follows by dropping the terms of the order  $(dt)^2$ . From (A2), we now have

$$d\lambda_t = \lambda_t((1 - \lambda_t)\gamma - \rho)dt.$$

Finally, given  $x_{i,t}b_{i,t} = 1$ , we have

$$\mathbb{P}[I_{i,t+dt} = 0 \mid I_{i,t} = 1] = (1 - \mu_{i,t})(1 - \lambda_{i,t})\gamma dt.$$

This completes the proof. ■

## Proof of Lemma 2

Similar to the proof of Lemma 1, if  $x_{i,t}b_{i,t} = 0$ , the user belief about  $\theta_i$  does not change, and when  $x_{i,t}b_{i,t} = 1$ , we let

$$\mu_{i,t}^{(P)} = \mathbb{P}[\theta = 1 \mid \alpha_{i,0} = 0, \text{NBN}_t]$$

be the probability of a product being high quality if the initial glossiness state is zero and no bad news has arrived by time  $t$ , and

$$\lambda_{i,t}^{(P)} = \mathbb{P}[\alpha_{i,t} = 1 \mid \alpha_{i,0} = 1, \text{NBN}_t]$$

be the probability of the glossiness state being 1 if the initial glossiness state is 1 and no bad news has arrived by time  $t$ . Again, to make the notation easier, in this proof, we drop the subscript  $i$ . Notice that conditioning on  $\alpha_0 = 1$ , the platform knows the product is of low quality and therefore  $\mu_t^{(P)} = 0$  for all  $t$ . Also, conditioning on  $\alpha_0 = 0$ , the glossiness state remains at zero and therefore  $\lambda_t^{(P)} = 0$ . We next evaluate  $\mu_t^{(P)}$  conditioning on  $\alpha_0 = 0$  and  $\lambda_t^{(P)}$  conditioning on  $\alpha_0 = 1$ .

By using Bayes' rule, for the probability of  $\theta_i = 1$  when  $\alpha_0 = 0$  we have

$$\begin{aligned}
\mu_{t+dt}^{(P)} &= \frac{\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 1, \alpha_0 = 0] \mathbb{P}[\theta = 1 \mid \alpha_0 = 0]}{\mathbb{P}[\text{NBN}_{t+dt} \mid \alpha_0 = 0]} \\
&= \frac{\mathbb{P}[\text{NBN}_t \mid \theta = 1, \alpha_0 = 0] \mathbb{P}[\text{NBN}_{t+dt} \mid \text{NBN}_t, \theta = 1, \alpha_0 = 0] \mathbb{P}[\theta = 1 \mid \alpha_0 = 0]}{\mathbb{P}[\text{NBN}_t \mid \alpha_0 = 0] \mathbb{P}[\text{NBN}_{t+dt} \mid \text{NBN}_t, \alpha_0 = 0]}
\end{aligned}$$

$$\begin{aligned}
&= \frac{\mu_t^{(P)} \mathbb{P}[\text{NBN}_{t+dt} \mid \text{NBN}_t, \theta = 1, \alpha_0 = 0]}{\mathbb{P}[\text{NBN}_{t+dt} \mid \text{NBN}_t, \alpha_0 = 0]} \\
&\stackrel{(a)}{=} \frac{\mu_t^{(P)}}{\mu_t^{(P)} + (1 - \mu_t^{(P)})(1 - \gamma dt)} \\
&= \frac{\mu_t^{(P)}}{1 - (1 - \mu_t^{(P)})\gamma dt} \\
&\stackrel{(b)}{=} \mu_t^{(P)} \left( 1 + (1 - \mu_t^{(P)})\gamma dt \right)
\end{aligned} \tag{A3}$$

where (a) follows from

$$\begin{aligned}
\mathbb{P}[\text{NBN}_{t+dt} \mid \text{NBN}_t, \alpha_0 = 0] &= \mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 1, \text{NBN}_t, \alpha_0 = 0] \mathbb{P}[\theta = 1 \mid \text{NBN}_t, \alpha_0 = 0] \\
&\quad + \mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t, \alpha_0 = 0] \mathbb{P}[\theta = 0 \mid \text{NBN}_t, \alpha_0 = 0] \\
&= \mu_t^{(P)} + \left( 1 - \mu_t^{(P)} \right) (1 - \gamma dt)
\end{aligned}$$

and (b) follows by dropping the terms of the order  $(dt)^2$ . By using (A3),

$$d\mu_t^{(P)} = \mu_{t+dt}^{(P)} - \mu_t^{(P)} = \mu_t^{(P)}(1 - \mu_t^{(P)})\gamma dt.$$

By using Bayes' rule, when  $\alpha_0 = 1$  we obtain

$$\begin{aligned}
\lambda_{t+dt}^{(P)} &= \mathbb{P}[\alpha_{t+dt} = 1 \mid \theta = 0, \text{NBN}_{t+dt}, \alpha_0 = 1] \\
&= \frac{\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \alpha_{t+dt} = 1, \alpha_0 = 1] \mathbb{P}[\alpha_{t+dt} = 1 \mid \theta = 0, \alpha_0 = 1]}{\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \alpha_0 = 1]} \\
&\stackrel{(a)}{=} \frac{\mathbb{P}[\alpha_t = 1 \mid \theta = 0, \alpha_0 = 1] \mathbb{P}[\alpha_{t+dt} = 1 \mid \alpha_t = 1, \theta = 0, \alpha_0 = 1] \mathbb{P}[\text{NBN}_t \mid \theta = 0, \alpha_{t+dt} = 1, \alpha_0 = 1]}{\mathbb{P}[\text{NBN}_t \mid \theta = 0, \alpha_0 = 1] \mathbb{P}[\text{NBN}_{t+dt} \mid \text{NBN}_t, \theta = 0, \alpha_0 = 1]} \\
&= \frac{\lambda_t^{(P)} \mathbb{P}[\alpha_{t+dt} = 1 \mid \alpha_t = 1, \theta = 0, \alpha_0 = 1]}{\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t, \alpha_0 = 1]} \\
&\stackrel{(b)}{=} \frac{\lambda_t^{(P)}(1 - \rho dt)}{\lambda_t^{(P)} + (1 - \lambda_t^{(P)})(1 - \gamma dt)} \\
&= \frac{\lambda_t^{(P)}(1 - \rho dt)}{1 - (1 - \lambda_t^{(P)})\gamma dt} \\
&\stackrel{(c)}{=} \lambda_t^{(P)} - \lambda_t^{(P)}\rho dt + \lambda_t^{(P)}(1 - \lambda_t)\gamma dt,
\end{aligned} \tag{A4}$$

where (a) follows from  $\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t, \alpha_{t+dt} = 1, \alpha_0 = 1] = 1$ , (b) follows from

$$\begin{aligned}
&\mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t, \alpha_0 = 1] \\
&= \mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t, \alpha_0 = 1, \alpha_t = 1] \mathbb{P}[\alpha_t = 1 \mid \theta = 0, \text{NBN}_t, \alpha_0 = 1] \\
&\quad + \mathbb{P}[\text{NBN}_{t+dt} \mid \theta = 0, \text{NBN}_t, \alpha_0 = 1, \alpha_t = 0] \mathbb{P}[\alpha_t = 0 \mid \theta = 0, \text{NBN}_t, \alpha_0 = 1],
\end{aligned}$$



and (c) follows by dropping the terms of the order  $(dt)^2$ . By using (A4), we obtain

$$d\lambda_t^{(P)} = \lambda_t^{(P)}((1 - \lambda_t^{(P)})\gamma - \rho)dt.$$

Finally, we have

$$\begin{aligned} \mathbb{P}[\text{bad news at } t \mid \alpha_0 = 0, \text{NBN}_t] &= \mathbb{P}[\text{bad news at } t \mid \alpha_0 = 0, \theta = 0, \text{NBN}_t] \mathbb{P}[\theta = 0 \mid \alpha_0 = 0, \text{NBN}_t] \\ &\quad + \mathbb{P}[\text{bad news at } t \mid \alpha_0 = 0, \theta = 1, \text{NBN}_t] \mathbb{P}[\theta = 1 \mid \alpha_0 = 0, \text{NBN}_t] \\ &= \gamma dt \left(1 - \mu_t^{(P)}\right) + 0 \end{aligned}$$

and

$$\begin{aligned} \mathbb{P}[\text{bad news at } t \mid \alpha_0 = 1, \text{NBN}_t] &= \mathbb{P}[\text{bad news at } t \mid \alpha_0 = 1, \text{NBN}_t, \theta = 0] \\ &= \mathbb{P}[\text{bad news at } t \mid \alpha_0 = 1, \theta = 0, \text{NBN}_t, \alpha_t = 1] \mathbb{P}[\alpha_t = 1 \mid \alpha_0 = 1, \theta = 0, \text{NBN}_t] \\ &\quad + \mathbb{P}[\text{bad news at } t \mid \alpha_0 = 1, \theta = 0, \text{NBN}_t, \alpha_t = 0] \mathbb{P}[\alpha_t = 0 \mid \alpha_0 = 1, \theta = 0, \text{NBN}_t] \\ &= 0 + \gamma dt \left(1 - \lambda_t^{(P)}\right). \end{aligned}$$

The initializations follow by using Bayes' rule, completing the proof. ■

## Details of the Proof of Theorem 1

Before presenting the details, let us state a lemma that we use.

**Lemma A1.** *In the pre-AI environment, for any product  $i \in \mathcal{N}$  that has been offered for  $[0, t)$  with  $\text{NBN}_{i,t}$  (no bad news), we have*

$$\lambda_{i,t} = \frac{(\gamma - \rho)\lambda}{\lambda\gamma + e^{(\rho-\gamma)t}((1-\lambda)\gamma - \rho)},$$

and

$$\mu_{i,t} = \frac{\mu_{i,0}(\gamma - \rho)}{\mu_{i,0}(\gamma - \rho) + (1 - \mu_{i,0})(e^{-\rho t}\lambda\gamma + e^{-\gamma t}((1-\lambda)\gamma - \rho))}.$$

Moreover,  $\mu_{i,t}$  is increasing in  $\mu_{i,0}$  and  $t$  and converges to 1 as  $t \rightarrow \infty$ . Additionally:

- If  $\rho > \gamma$ , then  $\lambda_{i,t}$  is monotonically decreasing and limits to zero as  $t \rightarrow \infty$ .
- If  $\rho \leq \gamma$ , then:
  - If  $(1-\lambda)\gamma - \rho \geq 0$ ,  $\lambda_{i,t}$  is monotonically increasing and converges to  $\frac{\gamma-\rho}{\gamma}$  as  $t \rightarrow \infty$ .
  - If  $(1-\lambda)\gamma - \rho < 0$ ,  $\lambda_{i,t}$  is monotonically decreasing and converges to  $\frac{\gamma-\rho}{\gamma}$  as  $t \rightarrow \infty$ .

*Proof of Lemma A1:* The differential equation of the evolution of  $\lambda_{i,t}$ , as we established in Lemma 1, is given by

$$\frac{d}{dt}\lambda_{i,t} = \lambda_{i,t}((1 - \lambda_{i,t})\gamma - \rho)$$

whose solution is

$$\frac{\gamma - \rho}{\gamma + ce^{(\rho - \gamma)t}}$$

for some constant  $c$ . Using the initial condition  $\lambda_{i,0} = \lambda$ , we then have

$$\lambda_{i,t} = \frac{(\gamma - \rho)\lambda}{\lambda\gamma + e^{(\rho - \gamma)t}((1 - \lambda)\gamma - \rho)}.$$

The differential equation of the evolution of  $\mu_{i,t}$ , as in Lemma 1, is given by

$$\frac{d}{dt}\mu_{i,t} = \mu_{i,t} (1 - \mu_{i,t}) (1 - \lambda_{i,t}) \gamma$$

whose solution is

$$\mu_{i,t} = \frac{e^{(\gamma + \rho)t}}{e^{(\gamma + \rho)t} + ce^{\gamma t}\lambda\gamma + ce^{\rho t}((1 - \lambda)\gamma - \rho)}$$

for some constant  $c$ . Using the initial condition  $\mu_{i,0}$ , we obtain

$$\mu_{i,t} = \frac{e^{(\gamma + \rho)t}\mu_{i,0}(\gamma - \rho)}{e^{(\gamma + \rho)t}\mu_{i,0}(\gamma - \rho) + (1 - \mu_{i,0})(e^{\gamma t}\lambda\gamma + e^{\rho t}((1 - \lambda)\gamma - \rho))}.$$

The rest of the proof follows by evaluating the monotonicity of  $\lambda_{i,t}$  and  $\mu_{i,t}$ . In particular, to see that  $\mu_{i,t}$  is increasing in the initial belief  $\mu_{i,0}$ , let us take the derivative of  $\mu_{i,t}$  with respect to  $\mu_{i,0}$  which is

$$\frac{(\gamma - \rho)(e^{-\rho t}\lambda\gamma + e^{-\gamma t}((1 - \lambda)\gamma - \rho))}{(\mu_{i,0}(\gamma - \rho) + (1 - \mu_{i,0})(e^{-\rho t}\lambda\gamma + e^{-\gamma t}((1 - \lambda)\gamma - \rho)))^2}.$$

This expression is nonnegative. To see this, notice that the denominator is always positive, and as we show next, the numerator is also nonnegative. Let us consider the following (exhaustive) three possibilities:

1.  $\gamma(1 - \lambda) - \rho \geq 0$ . In this case, we have  $\gamma - \rho \geq 0$  and therefore both terms in the numerator are nonnegative.
2.  $\gamma(1 - \lambda) - \rho \leq 0$  and  $\gamma - \rho \geq 0$ . In this case,  $e^{-\rho t} \geq e^{-\gamma t}$  and we have

$$\begin{aligned} (\gamma - \rho)(e^{-\rho t}\lambda\gamma + e^{-\gamma t}((1 - \lambda)\gamma - \rho)) &\geq (\gamma - \rho)(e^{-\gamma t}\lambda\gamma + e^{-\gamma t}((1 - \lambda)\gamma - \rho)) \\ &= (\gamma - \rho)e^{-\gamma t}(\lambda\gamma + ((1 - \lambda)\gamma - \rho)) \\ &= (\gamma - \rho)^2 e^{-\gamma t} \geq 0. \end{aligned}$$

3.  $\gamma - \rho \leq 0$ . In this case,  $(\gamma - \rho)((1 - \lambda)\gamma - \rho) \geq 0$  and  $e^{-\gamma t} \geq e^{-\rho t}$  and we have

$$\begin{aligned} (\gamma - \rho)(e^{-\rho t}\lambda\gamma + e^{-\gamma t}((1 - \lambda)\gamma - \rho)) &\geq (\gamma - \rho)e^{-\rho t}\lambda\gamma + (\gamma - \rho)((1 - \lambda)\gamma - \rho)e^{-\rho t} \\ &= (\gamma - \rho)e^{-\rho t}(\lambda\gamma + ((1 - \lambda)\gamma - \rho)) \\ &= (\gamma - \rho)^2 e^{-\rho t} \geq 0. \end{aligned}$$

This completes the proof of the lemma. ■

Lemma A1 characterizes the evolution of user beliefs about the quality of a product, conditioning on no bad news realization until the time in question. This lemma enables us to compare the belief trajectory of different products based on the user's initial belief and their history of experimentation. Building on this property, we now proceed with the proof of the theorem. In what follows, for analogy with the Gittens index literature, we refer to product  $i$  as arm  $i$ , and we refer to offering product  $i$  as pulling arm  $i$ . We next state the properties discussed in the main text and prove them.

**Lemma A2.** *In the pre-AI environment, the platform's equilibrium decision at any time  $t$  is to offer  $i^*$  where*

$$i^* \in \arg \max_{i \in \mathcal{N}} M(\mu_i, \lambda_i).$$

This lemma is an immediate implication of the optimality of the Gittens index established in, e.g., [Gittens et al., 2011, Chapter 2]. We make use of the following three lemmas as well.

**Lemma A3.** *Consider product  $i$  with the state  $(\mu_i, \lambda_i)$ . The optimal stopping time in the definition of Gittens index is to stop pulling arm  $i$  until bad news occurs for it. That is*

$$\tau_i \triangleq \inf \{t \in \mathbb{R}_+ : I_t = 0\}$$

*Proof of Lemma A3:* This lemma follows because  $\mu_{i,t}$  is increasing in  $t$  (Lemma A1). ■

**Lemma A4.** *For two products  $i$  and  $i'$  with  $\lambda_i = \lambda_{i'}$ , we have*

$$M(\mu_i, \lambda_i) \geq M(\mu_{i'}, \lambda_{i'}) \text{ for } \mu_i \geq \mu_{i'}.$$

*Proof of Lemma A4:* First note that by using Lemma A1, we know that  $\mu_{i,t}$  is increasing in  $\mu_{i,0}$ . Therefore,  $\mu_{i,t} \geq \mu_{i',t}$  for all  $t$ . Now consider the first time for which bad news occurs for both arms  $i$  and  $i'$  and let us denote them by  $\tau_i$  and  $\tau_{i'}$ . We claim that  $\tau_i$  first-order stochastically dominates  $\tau_{i'}$ , as

$$\begin{aligned} \mathbb{P}[I_{i,t+dt} = 0 \mid I_{i,t} = 1] &= (1 - \mu_{i,t}) (1 - \lambda_{i,t}) \gamma dt \\ &\stackrel{(a)}{=} (1 - \mu_{i,t}) (1 - \lambda_{i',t}) \gamma dt \\ &\leq (1 - \mu_{i',t}) (1 - \lambda_{i',t}) \gamma dt = \mathbb{P}[I_{i',t+dt} = 0 \mid I_{i',t} = 1], \end{aligned}$$

where (a) follows from the fact that, as established in Lemma A1, given  $\lambda_{i,0} = \lambda_{i',0} = \lambda$ , we have  $\lambda_{i,t} = \lambda_{i',t}$ .

Using Lemma A3 and given the above two properties, we have that

$$M(\mu_i, \lambda_i) = \frac{\mathbb{E}_{t \sim \tau_i}[e^{-rt} \mu_{i,t}]}{\mathbb{E}_{t \sim \tau_i}[e^{-rt}]} \geq \frac{\mathbb{E}_{t \sim \tau_{i'}}[e^{-rt} \mu_{i',t}]}{\mathbb{E}_{t \sim \tau_{i'}}[e^{-rt}]} = M(\mu_{i'}, \lambda_{i'}).$$

This completes the proof. ■

**Lemma A5.** Consider arm  $i$  with initial state  $(\mu_i, \lambda_i)$ . If bad news does not occur, the Gittin's index of this arm increases.

*Proof of Lemma A5:* Letting  $\tau_i$  be the first time for which bad news occurs and using Lemma A3, we can write

$$\begin{aligned} M(\mu_i, \lambda_i) &= \frac{\mathbb{E}_{t \sim \tau_i}[e^{-rt} \mu_{i,t}]}{\mathbb{E}_{t \sim \tau_i}[e^{-rt}]} = \frac{\mathbb{P}[I_{i,0} = 0] \mu_{i,0} + \mathbb{P}[I_{i,0} = 1] \mathbb{E}_{t \sim \tau_i}[e^{-rt} \mu_{i,t} \mid \text{NBN}_{i,0}]}{\mathbb{P}[I_{i,0} = 0] + \mathbb{P}[I_{i,0} = 1] \mathbb{E}_{t \sim \tau_i}[e^{-rt} \mid \text{NBN}_{i,0}]} \\ &\stackrel{(a)}{\leq} \frac{\mathbb{E}_{t \sim \tau_i}[e^{-rt} \mu_{i,t} \mid \text{NBN}_{i,0}]}{\mathbb{E}_{t \sim \tau_i}[e^{-rt} \mid \text{NBN}_{i,0}]}, \end{aligned}$$

where (a) follows from the fact that  $\mu_{i,t}$  is increasing in  $t$ , as established in Lemma A1. ■

We now continue with the detail of the proof of Theorem 1. Without loss of generality, let us suppose  $\mu_{n,0} \geq \dots \geq \mu_{1,0}$ . Using Lemma A4, and given  $\lambda_{i,0} = \lambda$  for all  $i \in \mathcal{N}$ , we have  $M(\mu_n, \lambda_n) \geq M(\mu_i, \lambda_i)$ . Therefore, the platform starts by offering product  $n$ . If bad news does not occur for this product, using Lemma A5, the Gittins index of arm  $n$  weakly increases while the other Gittins indices remain unchanged. Therefore, using Lemma A2, the platform finds it optimal to keep pulling arm  $n$  until bad news occurs. Whenever bad news occurs for product  $n$  (if at all), the platform's problem becomes offering the best arm among  $n - 1$  products, and by induction, again, the platform starts offering product  $n - 1$  and so on. ■

## Details of the Proof of Theorem 2

Let us first state a lemma that we use.

**Lemma A6.** In the post-AI environment, for any product  $i \in \mathcal{N}$  that has been offered for  $[0, t)$  with  $\text{NBN}_{i,t}$  (no bad news), the dynamics of  $\mu_{i,t}$  and  $\lambda_{i,t}$  are the same as Lemma A1, and additionally

- If  $\alpha_{i,0} = 0$ , then

$$\mu_{i,t}^{(P)} = \frac{\mu_{i,0}}{\mu_{i,0} + e^{-\gamma t}(1 - \lambda)(1 - \mu_{i,0})}, \quad \lambda_{i,t}^{(P)} = 0, \quad \mathbb{P}[\text{NBN}_{i,t} \mid \alpha_{i,0} = 0] = \frac{\mu_{i,0} + e^{-\gamma t}(1 - \lambda)(1 - \mu_{i,0})}{1 - \lambda(1 - \mu_{i,0})},$$

and

$$\mathbb{P}[I_{i,t+dt} = 0, I_{i,t} = 1 \mid \alpha_{i,0} = 0] = \frac{e^{-\gamma t}(1 - \lambda)(1 - \mu_{i,0})}{1 - \lambda(1 - \mu_{i,0})} \gamma dt.$$

- If  $\alpha_{i,0} = 1$ , then

$$\mu_{i,t}^{(P)} = 0, \quad \lambda_{i,t}^{(P)} = \frac{\gamma - \rho}{\gamma - e^{(\rho - \gamma)t} \rho}, \quad \mathbb{P}[\text{NBN}_{i,t} \mid \alpha_{i,0} = 1] = \frac{e^{-\rho t} \gamma - e^{-\gamma t} \rho}{\gamma - \rho},$$

and

$$\mathbb{P}[I_{i,t+dt} = 0, I_{i,t} = 1 \mid \alpha_{i,0} = 1] = \frac{(e^{-\rho t} - e^{-\gamma t}) \rho}{\gamma - \rho} \gamma dt.$$

*Proof of Lemma A6:* The differential equation of the evolution of  $\lambda_{i,t}^{(P)}$  when  $\alpha_{i,0} = 1$ , as we established in Lemma 2, is given by

$$\frac{d}{dt}\lambda_{i,t}^{(P)} = \lambda_{i,t}^{(P)} \left( (1 - \lambda_{i,t}^{(P)})\gamma - \rho \right)$$

whose solution is

$$\frac{e^{\gamma t + \rho c}(\gamma - \rho)}{e^{\rho t + \gamma c} - e^{\gamma t + \rho c}\gamma}$$

for some constant  $c$ . Using the initial condition  $\lambda_{i,0}^{(P)} = 1$ , we obtain

$$\lambda_{i,t}^{(P)} = \frac{\gamma - \rho}{\gamma - e^{(\rho - \gamma)t}\rho}$$

The differential equation of the evolution of  $\mu_{i,t}^{(P)}$  when  $\alpha_{i,0} = 0$ , as shown in Lemma 2, is given by

$$\frac{d}{dt}\mu_{i,t}^{(P)} = \mu_{i,t}^{(P)} \left( 1 - \mu_{i,t}^{(P)} \right) \gamma$$

whose solution is

$$\mu_{i,t}^{(P)} = \frac{1}{1 + ce^{-\gamma t}}$$

for some constant  $c$ . Using the initial condition  $\mu_{i,0}^{(P)} = \frac{\mu_{i,0}}{\mu_{i,0} + (1 - \mu_{i,0})(1 - \lambda)}$ , we obtain

$$\mu_{i,t}^{(P)} = \frac{\mu_{i,0}}{\mu_{i,0} + e^{-\gamma t}(1 - \lambda)(1 - \mu_{i,0})}.$$

For  $\alpha_{i,0} = 1$ , we have

$$\mathbb{P}[\text{NBN}_t \mid \alpha_{i,0} = 1] = e^{-\int_0^t \left( 1 - \frac{\gamma - \rho}{\gamma - e^{(\rho - \gamma)s}\rho} \right) \gamma ds} = \frac{e^{-\rho t}\gamma - e^{-\gamma t}\rho}{\gamma - \rho}.$$

Also, the probability of bad news arriving for the first time in  $(t, t + dt]$  is

$$\begin{aligned} \mathbb{P}[I_{t+dt} = 0, I_t = 1 \mid \alpha_{i,0} = 1] &= \mathbb{P}[\text{NBN}_t \mid \alpha_{i,0} = 1] \left( 1 - \frac{\gamma - \rho}{\gamma - e^{(\rho - \gamma)t}\rho} \right) \gamma dt \\ &= \frac{(e^{-\rho t} - e^{-\gamma t})\rho}{\gamma - \rho} \gamma dt. \end{aligned}$$

For  $\alpha_{i,0} = 0$ , we have

$$\mathbb{P}[\text{NBN}_t \mid \alpha_{i,0} = 0] = e^{-\int_0^t \left( 1 - \frac{\mu_{i,0}}{\mu_{i,0} + e^{-\gamma s}(1 - \lambda)(1 - \mu_{i,0})} \right) \gamma ds} = \frac{\mu_{i,0} + e^{-\gamma t}(1 - \lambda)(1 - \mu_{i,0})}{1 - \lambda(1 - \mu_{i,0})}.$$

Also, the probability of bad news arriving for the first time in  $(t, t + dt]$  is

$$\mathbb{P}[I_{t+dt} = 0, I_t = 1 \mid \alpha_{i,0} = 0] = \mathbb{P}[\text{NBN}_t \mid \alpha_{i,0} = 0] \left( 1 - \frac{\mu_{i,0}}{\mu_{i,0} + e^{-\gamma t}(1 - \lambda)(1 - \mu_{i,0})} \right) \gamma dt$$

$$= \frac{e^{-\gamma t}(1-\lambda)(1-\mu_{i,0})}{1-\lambda(1-\mu_{i,0})}\gamma dt.$$

This completes the proof of the lemma. ■

Lemma A6 characterizes the evolution of user and platform beliefs, conditioning on no bad news realization in the post-AI environment. This lemma enables us to compare the belief trajectories of different products based on the user's initial beliefs about them and the initial glossiness state  $\alpha_{i,0}$ . Similar to the proof of Theorem 1, we have the following:

**Lemma A7.** *In the post-AI environment, the platform's equilibrium decision at any time  $t$  is to offer  $i^*$  where*

$$i^* \in \arg \max_{i \in \mathcal{N}} M(\mu_i, \lambda_i, \mu_i^{(P)}, \lambda_i^{(P)}).$$

Using a similar argument to that of Lemma A4, and given  $\lambda_{i,0} = \lambda$  for all  $i \in \mathcal{N}$ , the platform starts by offering the product with the highest user belief either among those with the initial glossiness state  $\alpha = 1$  or among those with the initial glossiness state of 0. If bad news does not occur for this product, again using a similar argument to Lemma A5, establishes that the Gittins index of this arm weakly increases while the other Gittins indices remain unchanged. Therefore, the platform finds it optimal to keep offering this product until bad news occurs. Whenever bad news occurs for this product (if at all), the platform's problem becomes offering the best arm among the  $n - 1$  remaining products, and the proof follows by induction on the number of products. ■

### Proof of Proposition 1

By using Lemma A7, it suffices to consider the platform's problem with two products with initial beliefs  $\mu_{i',0} = \mu_1, \alpha_{i',0} = 1$  and  $\mu_{i,0} = \mu_0, \alpha_{i,0} = 0$  and characterize the platform's equilibrium decision. Also, notice that by using Lemma A5, the platform's equilibrium strategy is to either offer product  $i$  until bad news occurs and then switch to the product  $i'$  or to offer product  $i'$  until bad news occurs and then switch to the product  $i$ . We next evaluate the platform's utility with these two strategies and compare them. We find it useful to write the platform's payoff with these two strategies as a function of  $\rho$ .

Using Lemma A6, the platform's payoff when it offers product  $i$  first and when bad news occurs, switches to  $i'$  is

$$\begin{aligned} U_{0,1}(\rho, \mu_0, \mu_1) &\triangleq \frac{\mu_0}{\mu_0 + (1-\lambda)(1-\mu_0)} \int_0^\infty r e^{-rt} \mu_{i,t} dt \\ &+ \frac{(1-\lambda)(1-\mu_0)}{\mu_0 + (1-\lambda)(1-\mu_0)} \int_0^\infty r e^{-rt} \mu_{i,t} \mathbb{P}[\text{NBN}_{i,t} \mid \theta_i = 0, \alpha_{i,0} = 0] dt \\ &+ \frac{(1-\lambda)(1-\mu_0)}{\mu_0 + (1-\lambda)(1-\mu_0)} \left( \int_0^\infty e^{-rt} \mathbb{P}[I_{i,t+dt} = 0, I_{i,t} = 1 \mid \theta_i = 0, \alpha_{i,0} = 0] \right. \\ &\quad \times \left. \left( \int_0^\infty r e^{-rt} \mu_{i',t} \mathbb{P}[\text{NBN}_{i',t} \mid \alpha_{i',0} = 1] \right) \right) \\ &= \int_0^\infty r e^{-rt} \mu_{i,t} \frac{\mu_0 + (1-\mu_0)(1-\lambda)e^{-\gamma t}}{\mu_0 + (1-\lambda)(1-\mu_0)} dt \end{aligned}$$

$$+ \frac{(1-\lambda)(1-\mu_0)}{\mu_0 + (1-\lambda)(1-\mu_0)} \left( \int_0^\infty e^{-rt} \gamma e^{-\gamma t} dt \right) \left( \int_0^\infty r e^{-rt} \mu_{i',t} \frac{\gamma e^{-\rho t} - \rho e^{-\gamma t}}{\gamma - \rho} dt \right),$$

where, using Lemma A1, we have

$$\mu_{i,t} = \frac{\mu_0(\gamma - \rho)}{\mu_0(\gamma - \rho) + (1 - \mu_0)(e^{-\rho t} \lambda \gamma + e^{-\gamma t}((1 - \lambda)\gamma - \rho))} \text{ and}$$

$$\mu_{i',t} = \frac{\mu_1(\gamma - \rho)}{\mu_1(\gamma - \rho) + (1 - \mu_1)(e^{-\rho t} \lambda \gamma + e^{-\gamma t}((1 - \lambda)\gamma - \rho))}.$$

Again, using Lemma A6, the platform's payoff by offering product  $i'$  first and when bad news occurs then switching to product  $i$  is given by

$$\begin{aligned} U_{1,0}(\rho, \mu_0, \mu_1) &\triangleq \int_0^\infty r e^{-rt} \mu_{i',t} \mathbb{P}[\text{NBN}_{i',t} \mid \alpha_{i',0} = 1] dt \\ &+ \left( \int_0^\infty e^{-rt} \mathbb{P}[I_{i',t+dt} = 0, I_{i,t} = 1 \mid \alpha_{i',0} = 1] \right) \left( \int_0^\infty r e^{-rt} \mu_{i,t} \mathbb{P}[\text{NBN}_{i,t} \mid \alpha_{i,0} = 0] dt \right) \\ &= \int_0^\infty r e^{-rt} \mu_{i',t} \frac{\gamma e^{-\rho t} - \rho e^{-\gamma t}}{\gamma - \rho} dt \\ &+ \left( \int_0^\infty e^{-rt} \gamma \rho \frac{e^{-\rho t} - e^{-\gamma t}}{\gamma - \rho} dt \right) \left( \int_0^\infty r e^{-rt} \mu_{i,t} \frac{\mu_0 + e^{-\gamma t}(1 - \lambda)(1 - \mu_0)}{1 - \lambda(1 - \mu_0)} dt \right). \end{aligned}$$

Notice that both functions are (uniformly) continuous in  $\rho$ , and therefore it suffices to establish the statement in the limit of  $\rho \rightarrow 0$ . We can write

$$\begin{aligned} &\lim_{\rho \rightarrow 0} U_{1,0}(\rho, \mu_0, \mu_1) - U_{0,1}(\rho, \mu_0, \mu_1) \\ &= \int_0^\infty r e^{-rt} \frac{\mu_1 \gamma}{\mu_1 \gamma + (1 - \mu_1)(\lambda \gamma + e^{-\gamma t}(1 - \lambda)\gamma)} dt \\ &- \int_0^\infty r e^{-rt} \frac{\mu_0 \gamma}{\mu_0 \gamma + (1 - \mu_0)(\lambda \gamma + e^{-\gamma t}(1 - \lambda)\gamma)} \frac{\mu_0 + e^{-\gamma t}(1 - \lambda)(1 - \mu_0)}{1 - \lambda(1 - \mu_0)} dt \\ &- \left( \int_0^\infty e^{-rt} \frac{e^{-\gamma t}(1 - \lambda)(1 - \mu_0)}{1 - \lambda(1 - \mu_0)} \gamma dt \right) \left( \int_0^\infty r e^{-rt} \frac{\mu_1 \gamma}{\mu_1 \gamma + (1 - \mu_1)(\lambda \gamma + e^{-\gamma t}(1 - \lambda)\gamma)} dt \right) \\ &= \left( \int_0^\infty r e^{-rt} \frac{\mu_1 \gamma}{\mu_1 \gamma + (1 - \mu_1)(\lambda \gamma + e^{-\gamma t}(1 - \lambda)\gamma)} dt \right) \left( 1 - \frac{(1 - \lambda)(1 - \mu_0)\gamma}{(1 - \lambda(1 - \mu_0))(r + \gamma)} \right) \\ &- \int_0^\infty r e^{-rt} \frac{\mu_0 \gamma}{\mu_0 \gamma + (1 - \mu_0)(\lambda \gamma + e^{-\gamma t}(1 - \lambda)\gamma)} \frac{\mu_0 + e^{-\gamma t}(1 - \lambda)(1 - \mu_0)}{1 - \lambda(1 - \mu_0)} dt \\ &= \left( \int_0^\infty r e^{-rt} \frac{\mu_1 \gamma}{\mu_1 \gamma + (1 - \mu_1)(\lambda \gamma + e^{-\gamma t}(1 - \lambda)\gamma)} dt \right) \left( \int_0^\infty r e^{-rt} \frac{\mu_0 + e^{-\gamma t}(1 - \lambda)(1 - \mu_0)}{1 - \lambda(1 - \mu_0)} dt \right) \\ &- \int_0^\infty r e^{-rt} \frac{\mu_0 \gamma}{\mu_0 \gamma + (1 - \mu_0)(\lambda \gamma + e^{-\gamma t}(1 - \lambda)\gamma)} \frac{\mu_0 + e^{-\gamma t}(1 - \lambda)(1 - \mu_0)}{1 - \lambda(1 - \mu_0)} dt. \tag{A5} \end{aligned}$$

We next prove that the above expression is strictly positive for  $\mu_0 = \mu_1 = \mu$ . This follows from FKG-Harris inequality (see, e.g., [Alon and Spencer, 2016, Chapter 6.2]) that states if  $a : \mathbb{R}_+ \rightarrow \mathbb{R}_+$  is an

increasing and  $b : \mathbb{R}_+ \rightarrow \mathbb{R}_+$  is a decreasing function, then we have

$$\mathbb{E}[a(X)b(X)] \leq \mathbb{E}[a(X)]\mathbb{E}[b(X)]$$

for random variable  $X$  distributed over  $\mathbb{R}_+$ . Moreover, the above inequality is strict if both functions are strictly monotone and  $X$  has a strictly positive mass everywhere. In particular, the positivity of the right-hand side of (A5) follows by invoking the above inequality for

$$a(t) = \frac{\mu\gamma}{\mu\gamma + (1-\mu)(\lambda\gamma + e^{-\gamma t}(1-\lambda)\gamma)}, b(t) = \frac{\mu + e^{-\gamma t}(1-\lambda)(1-\mu)}{1-\lambda(1-\mu)},$$

and  $t$  distributed as  $\mathbb{P}[t = s] = re^{-rs}$ . Therefore, there exists a function  $f(\mu, \lambda, \gamma, r)$  that is everywhere below  $\mu$ , such that for  $\mu_1 \geq f(\mu_0, \lambda, \gamma, r)$ , the right-hand side of inequality (A5) is positive. ■

## Proof of Proposition 2

Similar to the proof of Proposition 1, it suffices to consider the platform's problem with two products with initial beliefs  $\mu_{i',0} = \mu_1, \alpha_{i',0} = 1$  and  $\mu_{i,0} = \mu_0, \alpha_{i,0} = 0$  and characterize the platform's equilibrium decision. Again, using a similar notation to that of Proposition 1, we let  $U_{1,0}(\rho, \mu_0, \mu_1)$  denote the platform's payoff from offering product  $i'$  first and  $U_{0,1}(\rho, \mu_0, \mu_1)$  denote the platform's payoff from offering product  $i$  first. Again, notice that both functions are (uniformly) continuous in  $\rho$  (over  $\mathbb{R}_+ \cup \{\infty\}$ ), and therefore it suffices to establish the statement in the limit of  $\rho \rightarrow \infty$ . We can write

$$\begin{aligned} & \lim_{\rho \rightarrow \infty} U_{0,1}(\rho, \mu_0, \mu_1) - U_{1,0}(\rho, \mu_0, \mu_1) \\ &= \int_0^\infty re^{-rt} \frac{\mu_0}{\mu_0 + e^{-\gamma t}(1-\mu_0)} \frac{\mu_0 + e^{-\gamma t}(1-\lambda)(1-\mu_0)}{1-\lambda(1-\mu_0)} dt \\ &+ \left( \int_0^\infty e^{-rt} \frac{e^{-\gamma t}(1-\lambda)(1-\mu_0)}{1-\lambda(1-\mu_0)} \gamma dt \right) \left( \int_0^\infty re^{-rt} \frac{\mu_1}{\mu_1 + e^{-\gamma t}(1-\mu_1)} e^{-\gamma t} dt \right) \\ &- \int_0^\infty re^{-rt} \frac{\mu_1}{\mu_1 + e^{-\gamma t}(1-\mu_1)} e^{-\gamma t} dt \\ &- \left( \int_0^\infty e^{-rt} \gamma e^{-\gamma t} dt \right) \left( \int_0^\infty re^{-rt} \frac{\mu_0}{\mu_0 + e^{-\gamma t}(1-\mu_0)} \frac{\mu_0 + e^{-\gamma t}(1-\lambda)(1-\mu_0)}{1-\lambda(1-\mu_0)} dt \right) \\ &= \left( \int_0^\infty re^{-rt} \frac{\mu_1}{\mu_1 + e^{-\gamma t}(1-\mu_1)} e^{-\gamma t} dt \right) \left( \frac{\gamma(1-\lambda)(1-\mu_0)}{(1-\lambda(1-\mu_0))(\gamma+r)} - 1 \right) \\ &+ \int_0^\infty re^{-rt} \frac{\mu_0}{\mu_0 + e^{-\gamma t}(1-\mu_0)} \frac{\mu_0 + e^{-\gamma t}(1-\lambda)(1-\mu_0)}{1-\lambda(1-\mu_0)} dt \left( 1 - \frac{\gamma}{r+\gamma} \right). \end{aligned} \tag{A6}$$

For  $\mu_1 = \mu_0 = \mu$ , the right-hand side of (A6) can be written as

$$\left( \frac{r}{r+\gamma} \right)^2 (\mathbb{E}[a(t)b(t)] - \mathbb{E}[a(t)]\mathbb{E}[b(t)])$$



for  $a(t) = \frac{\mu}{\mu + e^{-\gamma t}(1-\mu)}$  and  $b(t) = \frac{\mu e^{\gamma t} + (1-\lambda)(1-\mu)}{1-\lambda(1-\mu)}$  and  $t$  being distributed as  $\mathbb{P}[t = s] = (\gamma + r)e^{-(\gamma+r)s}$ . Again using FKG-Harris inequality for (strictly) increasing functions  $a(\cdot)$  and  $b(\cdot)$ , we conclude that the right-hand side of (A6) is strictly positive. Therefore, there exists a function  $g(\mu, \lambda, \gamma, r)$  that is everywhere below  $\mu$ , such that for  $\mu_0 \geq g(\mu_1, \lambda, \gamma, r)$ , the right-hand side of (A6) is positive. ■

### Proof of Theorem 3

As we established in Theorems 1 and 2, the equilibrium in both pre-AI and post-AI settings start with the offering of one of the products until bad news occurs, and then following bad news, the platform offers another product and so on. Also, as in Propositions 1 and 2, both the expected utility obtained from offering any product and the probability of bad news occurring for a product are continuous functions of  $\rho$ . Therefore, it suffices to prove the theorem statement in the limit of  $\rho \rightarrow \infty$  with strict inequality. This establishes the existence of large enough  $\rho_h$  such that the inequalities hold for  $\rho \geq \rho_h$ .

First, note that the platform's information advantage in the post-AI implies that its expected payoff is higher than in the pre-AI environment. We next prove that the expected user utility in the post-AI is higher than in the pre-AI environment. Since utilitarian welfare is the summation of the user and the platform utilities, we conclude that the expected welfare in the post-AI is higher than in the pre-AI.

Let us consider a realization of the products' beliefs and, without loss of generality, assume

$$\mu_{1,0} \leq \dots \leq \mu_{n,0}.$$

As we established in Theorem 1, in the pre-AI environment, the platform first offers the product with the highest belief, that is, product  $n$ , and if bad news happens, offer the product with the second highest product, and so on. Using Proposition 2, in the post-AI environment, the platform may start offering product  $i$  where  $i < n$ ,  $\alpha_{i,0} = 0$ , and  $\alpha_{n,0} = 1$ . In the post-AI environment, the platform equilibrium strategy may also differ in the order of other products. However, we can obtain the platform strategy in the post-AI from its strategy in the pre-AI by performing a series of *helpful swaps*, defined next.

Suppose  $\mu_{1,0} \leq \dots \leq \mu_{n,0}$  and  $i$  is the highest index for which  $\alpha_{i,0} = 0$ . Consider the strategy of offering products in decreasing order of their beliefs. A *helpful swap* offers product  $i$  first and then offers the rest of the products in decreasing order of their beliefs.

The proof of theorem follows from the following lemma that establishes for large enough  $\rho$ , a helpful swap increases the user expected utility.

**Lemma A8.** *The expected user utility increases for large enough  $\rho$  after performing a helpful swap.*

*Proof of Lemma A8:* It suffices to prove the lemma in the limit of  $\rho \rightarrow \infty$  with strict inequality. This establishes the existence of large enough  $\rho$  for which the statement holds. For any  $j \in \mathcal{N}$ , we let  $\tau_j$  be the stochastic time at which bad news occurs for product  $j$  given  $\theta_j = 0$ . Notice that in the limit of  $\rho \rightarrow \infty$ , for both a product with  $\alpha_{j,0} = 1$  and a low-quality product with  $\alpha_{j,0} = 0$ , the platform knows that bad news arrive with rate  $\gamma$ . We let

$$A_j \triangleq \mathbb{E} \left[ - \int_0^{\tau_j} r e^{-rt} \mu_{j,t} dt \right]$$

for  $j = n, n-1, i$ . Using this notation, and that for  $j = n, n-1, \dots, i+1$ ,  $\delta \triangleq \mathbb{E}_{t \sim \tau_j}[e^{-rt}] < 1$ , the expected utility before the swap is

$$U_1 \triangleq \sum_{j=n}^{i+1} A_j \delta^{n-j} + \delta^{n-i} \left( \mu_{i,0}^{(P)} \int_0^\infty r e^{-rt} (1 - \mu_{i,t}) dt + (1 - \mu_{i,0}^{(P)}) A_i \right) + \delta^{n-i+1} (1 - \mu_{i,0}^{(P)}) (\text{Expected utility from products } i-1, \dots, 0), \quad (\text{A7})$$

where  $\mu_{i,0}^{(P)} = \frac{\mu_{i,0}}{\mu_{i,0} + (1-\lambda)(1-\mu_{i,0})}$  when  $\alpha_{i,0} = 0$ . The expected utility after the swap is

$$\begin{aligned} U_2 &\triangleq \mu_{i,0}^{(P)} \int_0^\infty r e^{-rt} (1 - \mu_{i,t}) dt + (1 - \mu_{i,0}^{(P)}) A_i + \delta (1 - \mu_{i,0}^{(P)}) \sum_{j=n}^{i+1} A_j \delta^{n-j} \\ &\quad + \delta^{n-i+1} (1 - \mu_{i,0}^{(P)}) (\text{Expected utility from products } i-1, \dots, 0) \\ &\stackrel{(a)}{>} \delta^{n-i} \mu_{i,0}^{(P)} \int_0^\infty r e^{-rt} (1 - \mu_{i,t}) dt + (1 - \mu_{i,0}^{(P)}) A_i + \delta (1 - \mu_{i,0}^{(P)}) \sum_{j=n}^{i+1} A_j \delta^{n-j} \\ &\quad + \delta^{n-i+1} (1 - \mu_{i,0}^{(P)}) (\text{Expected utility from products } i-1, \dots, 0) \\ &= U_1 + A_i \left( (1 - \mu_{i,0}^{(P)}) - \delta^{n-i} (1 - \mu_{i,0}^{(P)}) \right) + \sum_{j=n}^{i+1} A_j \left( \delta^{n+1-j} (1 - \mu_{i,0}^{(P)}) - \delta^{n-j} \right) \\ &\stackrel{(b)}{\geq} U_1 + A_i \left( (1 - \mu_{i,0}^{(P)}) - \delta^{n-i} (1 - \mu_{i,0}^{(P)}) \right) + \sum_{j=n}^{i+1} A_i \left( \delta^{n+1-j} (1 - \mu_{i,0}^{(P)}) - \delta^{n-j} \right) \\ &= U_1 - \mu_{i,0}^{(P)} A_i \left( \sum_{j=0}^{n-i+1} \delta^j \right) \stackrel{(c)}{>} U_1 \end{aligned} \quad (\text{A8})$$

where (a) follow from  $\mu_{i,0}^{(P)} \int_0^\infty r e^{-rt} (1 - \mu_{i,t}) dt > 0$  and  $1 > \delta^{n-i}$ , (b) follows from  $A_j < A_i$  and  $\delta^{n+1-j} (1 - \mu_{i,0}^{(P)}) - \delta^{n-j} < 0$  for  $j = n, n-1, i+1$ , and (c) follows from  $A_i < 0$ , completing the proof of lemma. ■

We now continue with the proof of the theorem. If  $\alpha_{n,0} = 0$ , then in both pre-AI and post-AI the platform starts by offering product  $n$ , and the proof follows by induction on the number of products. If  $\alpha_{n,0} = 1$ , then letting  $i$  be the largest index for which  $\alpha_{i,0} = 0$ , with positive probability we have that  $\mu_{i,0} > g(\mu_{n,0}, \lambda, \gamma, r)$  and therefore Proposition 2 implies that the platform starts by offering product  $i$  in the post-AI. The above Lemma establishes the expected user utility increases by this helpful swap. The proof completes by noting that the platform strategy in the post-AI can be obtained from its strategy in the pre-AI by performing a series of helpful swaps. ■

#### Proof of Theorem 4

Again, as we established in Theorems 1 and 2, the equilibrium in both pre-AI and post-AI settings start offering one of the products until bad news occurs for it and then offers another product and so on. Also, as we showed in the proof of Propositions 1 and 2, both the expected utility obtained from

offering any product and the probability of bad news occurring for a product are continuous functions of  $\rho$ . Therefore, it suffices to prove the theorem statement in the limit of  $\rho \rightarrow 0$  with strict inequality. This establishes the existence of small enough  $\rho_l$  such that the inequalities hold for  $\rho \leq \rho_l$ . First, note that the platform's information advantage in the post-AI implies its expected payoff is higher than in the pre-AI environment. We next prove that the expected utilitarian welfare in the post-AI is smaller than in the pre-AI environment. Since utilitarian welfare is the summation of the user and the platform utilities, the expected user utility in the post-AI is smaller than in the pre-AI.

Let us consider a realization of the products' beliefs and, without loss of generality, assume

$$\mu_{1,0} \leq \dots \leq \mu_{n,0}.$$

Using the characterizations of Theorems 1 and 2, we have the following cases:

1.  $\alpha_{n,0} = 1$ : In this case, in both pre-AI and post-AI environments, the platform starts offering product  $n$ . Since  $\rho = 0$ , bad news does not arrive for this product, and the expected welfare in pre-AI and post-AI environments are the same.
2.  $\alpha_{n,0} = 0$ , there are two cases:
  - 2.1. For all products  $j \in \mathcal{N} \setminus \{n\}$  for which  $\alpha_{j,0} = 1$ , we have  $\mu_{j,0} \leq f(\mu_{n,0}, \lambda, \gamma, r)$ . In this case, in both the pre-AI and post-AI environment, the platform starts offering product  $n$ . Therefore, by induction on the number of products, the expected welfare in the post-AI is smaller than in the pre-AI environment.
  - 2.2. There exists product  $j \in \mathcal{N} \setminus \{n\}$  for which  $\alpha_{j,0} = 1$  and  $\mu_{j,0} > f(\mu_{n,0}, \lambda, \gamma, r)$ . We let  $i$  be the largest index among such products. In the post-AI environment, the platform starts offering product  $i$ , and given that bad news does not occur for this product ( $\rho = 0$ ) and  $\theta_i = 0$ , the expected welfare becomes 0. In the pre-AI environment, the platform starts by offering product  $n$ , and therefore, with probability  $\mu_{n,0}$  we have  $\theta_n = 1$ , ensuring an expected welfare of at least  $\mu_{n,0}$ .

Finally, notice that case 2.2 happens with a positive probability, proving that the expected welfare in the post-AI environment is strictly lower than in the pre-AI environment. ■

## Proof of Theorem 5

Similar to the argument of Theorem 4, it suffices to prove the theorem statement for the expected welfare in the limit of  $\rho \rightarrow 0$  with strict inequality. This establishes the existence of small enough  $\tilde{\rho}_l$  such that the inequalities hold for  $\rho \leq \tilde{\rho}_l$ . We let  $\Delta_1$  be small enough so that

$$f\left(\frac{1}{2} + \Delta_1; \lambda, \gamma, r\right) \leq \frac{1}{2} - \Delta_1. \quad (\text{A9})$$

Such  $\Delta_1$  exists because  $f(\mu; \lambda, \gamma, r)$  is continuous in  $\mu$  and is everywhere below  $\mu$  and in particular,  $f(\frac{1}{2}; \lambda, \gamma, r) \leq \frac{1}{2}$ . The continuity of  $f(\cdot)$  follows from the characterization of the function  $f(\cdot)$  derived in Proposition 1 (and in particular, (A5)) and implicit function theorem. Let us compare the expected welfare conditioned on the belief realization for the first  $n$  products. Without loss of generality, we assume  $\mu_{1,0} \leq \mu_{2,0} \leq \dots \mu_{n,0}$ . Let us assume  $\Delta \leq \Delta_1$ . We have the following cases:

1. There exists  $i \in \mathcal{N}$  for which  $\alpha_{i,0} = 1$ : In this case, given  $\Delta \leq \Delta_1$  and (A9), the platform's equilibrium strategy with and without the  $n+1$ -th product is to offer one of the products with the initial glossiness state of 1. Therefore, the expected welfare with and without the extra product is zero.
2.  $\alpha_{i,0} = 0$  for all  $i \in \mathcal{N}$  and  $\alpha_{n+1,0} = 1$ : In this case, given  $\Delta \leq \Delta_1$  and (A9), the platform's equilibrium strategy with the  $n+1$ -th product is to offer that product, and the expected welfare becomes 0. The platform's equilibrium strategy with  $n$  products, however, is to offer them in the decreasing order of their belief whose expected welfare is recursively defined as

$$W(\mu_{n,0}, \dots, \mu_{1,0}) = \mu_{n,0}^{(P)} + (1 - \mu_{n,0}^{(P)})\delta W(\mu_{n-1,0}, \dots, \mu_{1,0}), \quad (\text{A10})$$

where  $\delta \triangleq \mathbb{E}_{t \sim \tau_j} [e^{-rt}] < 1$ ,  $\tau_j$  is the arrival time of bad news for a low-quality product in the initial glossiness state of 0, and  $\mu_{i,0}^{(P)} = \frac{\mu_{i,0}}{\mu_{i,0} + (1-\lambda)(1-\mu_{i,0})}$  when  $\alpha_{i,0} = 0$ .

3.  $\alpha_{i,0} = 0$  for all  $i \in \{1, \dots, n\}$  and  $\alpha_{n+1,0} = 0$ : In this case, again, the expected welfare with  $n$  products is  $W(\mu_{n,0}, \dots, \mu_{1,0})$  as defined in (A10). The expected welfare with  $n+1$  products is

$$W(\mu_{n,0}, \dots, \mu_{i+1,0}, \mu_{n+1,0}, \mu_{i,0}, \dots, \mu_{1,0}) \quad (\text{A11})$$

where  $i$  is such that  $\mu_{i+1,0} \geq \mu_{n+1,0} \geq \mu_{i,0}$ .

With the extra product, the expected welfare in case 2 increases while the expected welfare in case 3 decreases. We next prove that, for small enough  $\Delta_2$ , when  $\mu_i \in [1/2 - \Delta_2, 1/2 + \Delta_2]$ , in expectation, the expected welfare decreases. In case 2, the difference between expected welfare with and without the extra product is upper bounded by

$$-\mu_l \frac{1 - ((1 - \mu_l)\delta)^n}{1 - (1 - \mu_l)\delta}, \text{ where } \mu_l = \frac{\frac{1}{2} - \Delta_2}{\frac{1}{2} - \Delta_2 + (\frac{1}{2} + \Delta_2)(1 - \lambda)},$$

noting that  $\mu_l$  is the smallest belief conditional on the initial glossiness state being 0 obtained by using Bayes' rule. In case 3, the difference between expected welfare with and without the extra product is upper bounded by

$$\mu_h \frac{1 - ((1 - \mu_h)\delta)^{n+1}}{1 - (1 - \mu_h)\delta} - \mu_l \frac{1 - ((1 - \mu_l)\delta)^n}{1 - (1 - \mu_l)\delta}, \text{ where } \mu_h = \frac{\frac{1}{2} + \Delta_2}{\frac{1}{2} + \Delta_2 + (\frac{1}{2} - \Delta_2)(1 - \lambda)},$$

noting that  $\mu_h$  is the largest belief conditional on the initial glossiness state being 0 obtained by using Bayes' rule. Therefore, the expected welfare with  $n + 1$  products minus the expected welfare with  $n$  products is upper bounded by

$$-(1 - \mu_h)\lambda\mu_l \frac{1 - ((1 - \mu_l)\delta)^n}{1 - (1 - \mu_l)\delta} + (1 - (1 - \mu_h)\lambda) \left( \mu_h \frac{1 - ((1 - \mu_h)\delta)^{n+1}}{1 - (1 - \mu_h)\delta} - \mu_l \frac{1 - ((1 - \mu_l)\delta)^n}{1 - (1 - \mu_l)\delta} \right).$$

We claim that the above expression is negative for small enough  $\Delta_2$ . Notice this expression is continuous in  $\Delta_2$  and for  $\Delta_2 = 0$  it becomes

$$\begin{aligned} & \frac{-(2 - \lambda)(1 - \lambda)\lambda + \left(\delta \frac{1 - \lambda}{2 - \lambda}\right)^n ((2 - \lambda)^2 - \delta(1 - \lambda)(2 - (2 - \lambda)\lambda))}{(2 - \delta(1 - \lambda) - \lambda)(2 - \lambda)^2} \\ & \stackrel{(a)}{<} \frac{-(2 - \lambda)(1 - \lambda)\lambda + \left(\frac{1 - \lambda}{2 - \lambda}\right)^n (2 - \lambda)^2}{(2 - \delta(1 - \lambda) - \lambda)(2 - \lambda)^2} \stackrel{(b)}{\leq} 0 \end{aligned}$$

where (a) follows from the fact that the denominator is positive and that  $(2 - \lambda)^2 - \delta(1 - \lambda)(2 - (2 - \lambda)\lambda) \geq (2 - \lambda)^2 - (1 - \lambda)(2 - (2 - \lambda)\lambda) = 2(1 - \lambda) > 0$ ,  $\delta < 1$ , and  $\delta(1 - \lambda)(2 - (2 - \lambda)\lambda) > 0$  and (b) follows from  $n \geq \frac{\log \frac{2 - \lambda}{(1 - \lambda)\lambda}}{\log \frac{2 - \lambda}{1 - \lambda}} = 1 - \frac{\log \lambda}{\log \frac{2 - \lambda}{1 - \lambda}}$ . This proves that for  $\Delta = \min\{\Delta_1, \Delta_2\}$ , the expected welfare (and therefore the expected user utility) decreases as the number of products increases. ■