

When Artificial Differentiation Meets Artificial Intelligence

Preyas S. Desai
Duke University
(desai@duke.edu)

Jessie Liu
Johns Hopkins University
(jzliu@jhu.edu)

Abstract

In many purchase situations, consumers are uncertain about whether a product attribute is beneficial or harmful. That is, they face uncertainty about the model that maps one or more attributes into their utility functions. This *model uncertainty* (ambiguity), prevalent in food and health-related products, is often intensified by conflicting scientific claims and social media controversies. In addition to facilitating consumers' search for products, AI can also serve as a decision aid that reduces model uncertainty. We examine how reduction in model uncertainty affects market size and firm profits.

To model the market environment, we develop a framework with Knightian uncertainty about attribute valence, following the smooth ambiguity model of Klibanoff et al. (2005) and Maccheroni et al. (2013). Two firms offer products with different levels of a focal attribute. Consumers know these levels but *disagree* about whether higher levels are beneficial or harmful, and are averse to this ambiguity. We model AI as an expert advisor that knows the true valence but may provide assessments that are partially distorted by sycophancy.

Because some consumers believe that the attribute is more likely to be beneficial, the product with the higher attribute level enjoys a competitive advantage among them. However, it also suffers a larger ambiguity penalty among all consumers. As a result, consumers' model uncertainty creates a novel form of product differentiation that combines elements of both horizontal and vertical differentiation. We also find that consumers' use of AI changes the distribution of consumers' beliefs in the market rather than simply reducing model uncertainty. These changes have asymmetric effects on the two firms. In addition, AI creates both winners and losers in the market: it benefits consumers who either stop purchasing a harmful product or start purchasing a beneficial product. However, some consumers can be worse off with the use of AI as firms can capture part of their informational gains through higher prices.

Keywords: Attribute Valence, Model Uncertainty, Ambiguity, Artificial Intelligence (AI)

1 Introduction

“The chief feature of the problem of uncertainty is that it cannot be reduced to an objective, quantitatively determined probability.”

“The presence of uncertainty gives rise to a special income, called profit.”

— Frank H. Knight, *Risk, Uncertainty, and Profit* (1921)

In many product categories, consumers are increasingly paying attention to specific ingredients related to product safety, health benefits, or societal impact (Leger, 2022, Mintel, 2025) and making purchase decisions based on their assessment. As consumers try to gather more information about possible effects of specific ingredients, they face an important challenge. They might discover that a product contains a particular ingredient or attribute but might still remain uncertain about whether that attribute is beneficial or harmful. In other words, consumers are uncertain not about whether the attribute is present or how much of the attribute is in a product but about how the impact of that attribute should be evaluated. Consider, for example, plant-based protein powders marketed as healthier alternatives to other types of protein supplements. While consumers might associate plant-based ingredients with health and sustainability, they might also be concerned about exposure to lead and other heavy metals through soil and groundwater contamination (Martineau, 2025). Similarly, skincare products frequently highlight novel or unfamiliar ingredients whose long-term effects are difficult for consumers to assess.

In these cases, consumers are uncertain if specific ingredients are beneficial or harmful to them. That is, they face uncertainty about the model that maps one or more attributes into their utility functions. Importantly, consumers are uncertain not about whether an attribute is present, but about how to evaluate it. The same observed attribute level can increase utility under one model and decrease utility under another. This *model uncertainty* (ambiguity) can also get amplified through social media controversies surrounding health and safety claims that present conflicting interpretations of scientific evidence. Moreover, consumers are often exposed to competing interpretations of the same scientific evidence. Media outlets, influencers, advocacy groups, firms, and experts may reach different conclusions regarding the risks and benefits of a given attribute. As a result, consumers can arrive at markedly different assessments of the same product attribute even when evaluating the same underlying evidence. Fragmented media exposure, differences across individual consumers’ social media feeds, and differences in prior experiences with a specific product category then could also lead to different beliefs across consumers. As a result, consumers may hold substantially different beliefs regarding the same attribute. Some consumers might

consider a specific attribute beneficial while others might consider it harmful.¹

When faced with model uncertainty and limited ability to independently evaluate scientific evidence, consumers often turn to AI systems to provide assessments of complex product attributes.

In this paper, we study how consumers’ model uncertainty affects consumers’ choice between two competing products. We develop a model in which consumers know the level of a focal attribute in each product but disagree about whether greater levels of the attribute are beneficial or harmful. Two firms offer products that differ in their levels of the focal attribute, and consumers evaluate these products using the smooth ambiguity framework of Klibanoff et al. (2005) and Maccheroni et al. (2013). Consumers are heterogeneous in their subjective beliefs regarding the attribute’s valence (beneficiality or harmfulness) and incur an ambiguity penalty when evaluating products whose consequences depend on the unknown model.

We also examine how consumers’ use of AI affects market outcomes. Consumers are already using AI chatbots as experts and are expected to use them even more (McKinsey, 2025). AI systems can synthesize scientific evidence, analyze competing claims, evaluate costs and benefits of an attribute in the given context, and generate interpretable summaries of complex information. Unlike human consumers, AI can also search a much wider set of media and scientific resources and process large amounts of information in a matter of seconds. Therefore, we also examine how consumers’ use of AI in their purchase decision-making process affects market outcomes. However, AI-generated assessments are not necessarily neutral. A growing literature shows that AI systems can exhibit sycophantic tendencies, tailoring their responses toward users’ prior beliefs and expectations. To capture these features, we model AI as an expert advisor that knows the true valence of the focal attribute but may provide assessments to consumers that are partially distorted by sycophancy. We then examine how this combination of expertise and sycophantic distortion affects consumer beliefs and market competition.

1.1 Overview of Results

Our analysis generates several interesting insights. First, model uncertainty creates a novel form of product differentiation. The product with the higher attribute level benefits when consumers believe the attribute is likely to be beneficial but it suffers a larger ambiguity

¹Although such model uncertainty is more prevalent in food and health-related products, it could also be observed in other domains such as consumer electronic products.

penalty than the other product. The net effect of these valence and ambiguity effects is that the product market competition combines elements of both horizontal and vertical differentiation. Even when one product would completely dominate the other product under model certainty, consumers’ model uncertainty creates demand for the alternative product because of consumers’ skepticism and ambiguity aversion.

We also find that when consumers use AI, it changes the distribution of consumers’ beliefs in the market rather than simply reducing model uncertainty. Consumers whose priors are originally aligned with the truth become fully informed, whereas consumers whose priors conflict with the truth receive partially distorted advice and retain residual ambiguity. AI therefore creates polarization in consumers’ posterior beliefs. These changes have asymmetric and non-monotonic effects on competing firms. When the true valence favors the high-attribute product, AI increases that firm’s demand; the low-attribute product, however, benefits most when AI is neither fully truthful nor highly sycophantic. Moderate levels of sycophancy result in enough residual ambiguity to support the lower-attribute product and can induce previously inactive consumers to participate in the market.

We also show that receiving expert advice from AI does not necessarily increase consumer welfare across the board. For example, it benefits consumers who start purchasing a beneficial product that they did not previously buy but firms can capture some of the informational gains from all consumers through higher prices. Thus, the overall impact of AI depends on how AI’s expertise is distributed across consumers.

We next discuss the related literature in Section 2, describe our model in Section 3, and present our analysis in Section 4. Section 5 gives our concluding remarks.

2 Literature Review

Our paper is primarily related to two streams of research in marketing and economics.

The first is the literature on decision-making under ambiguity and uncertainty.² The axioms of subjective expected utility (SEU; Savage (1954)) require decision makers to exhibit the same attitude toward all sources of uncertainty, whether risks are well-defined or ambiguous Nau (2011). A central empirical challenge to this framework is ambiguity aversion, first demonstrated in Ellsberg’s (1961) urn experiments, in which individuals systematically

²To acknowledge the connection to the broader literature and its history, we sometimes use the term “ambiguity.” However, the term “model uncertainty” is more appropriate for our research setting.

prefer betting on events with known probabilities (risk) over events with ambiguous probabilities. This empirical regularity has motivated three broad classes of models that relax SEU: smooth preferences (e.g. Klibanoff et al., 2005, Nau, 2006, Maccheroni et al., 2013), Maxmin expected utility (e.g. Gilboa and Schmeidler, 1989), and incomplete-preference models (Bewley, 2002). The distinction between risk and uncertainty predates SEU and originates in Knight (1921), though, as Nau (2011, pp.438) notes, it is not entirely clear which formal theory Knight had in mind. Nonetheless, modern ambiguity models share the Knightian insight that uncertainty arises from indeterminacy about the underlying probability model, while differing in how such indeterminacy is represented and how sharply it affects behavior. A small number of marketing studies have applied ambiguity-based approaches. For example, Handel and Misra (2015) uses a two-period Maxmin framework to study the pricing decision of a monopolist facing uncertainty about the underlying demand curve at the time of a new product launch. We adopt the smooth ambiguity framework of Klibanoff et al. (2005) and Maccheroni et al. (2013), which allows consumers to hold structured but imperfect beliefs over alternative probability models, represented as second-order probabilities. Ambiguity aversion is represented as second-order risk aversion over first-order expected payoffs. This framework is particularly well suited to our setting because it enables us to model *heterogeneity* in consumers’ beliefs about attribute valence. In contrast to the Maxmin approach where learning matters only when certain models are ruled out entirely, we retain a smoother representation because our focus is on how incremental change in model uncertainty (e.g., via AI) differentially affect consumers with heterogeneous beliefs. We contribute to this literature by embedding heterogeneous consumer beliefs into a competitive market setting, showing how this heterogeneity creates a novel form of product differentiation.

Our paper is also related to the marketing literature on product differentiation. A large literature studies how firms design products to serve heterogeneous consumer preferences and whether differentiation emerges in equilibrium (Moorthy, 1984, Desai, 2001). Other work shows that firms can endogenously increase differentiation through targeted advertising (Iyer et al., 2005). In these models, differentiation typically arises from product design, pricing, or targeted advertising under a standard Bayesian framework. Relatedly, Lee et al. (2025) show in a cheap-talk setting that competition can discipline communication about attribute importance. Their focus is on how competition enables credible communication under classical (Bayesian) uncertainty, whereas we study how communication under *model uncertainty about attribute valence* affects the competitive environment. Our paper identifies a distinct mechanism through which products become differentiated through consumers’ model uncertainty and ambiguity aversion. We also show that consumers’ use of AI alters

this type of product differentiation.

3 Model

3.1 Firms and Products

There are two firms, indexed by $j \in \{A, B\}$. Each firm offers a single product that provides a bundle of attributes. Because of our focus on model uncertainty, we consider a single attribute that presents *model uncertainty*, i.e., uncertainty about how the attribute maps into the consumer utility, to the consumer. This model uncertainty is about the attribute *valence* and can be considered a special case of Knightian uncertainty (ambiguity).³

We denote firm j 's level of the focal attribute by m_j ($j = \{A, B\}$), and assume without loss of generality that $m_A > m_B$. This differentiation in attribute levels is important for the role of model uncertainty in the model. When products differ along the focal attribute, consumers' uncertainty regarding how attributes should be evaluated can materially affect their relative preferences for the two products and equilibrium outcomes. In contrast, the implications of model uncertainty are less consequential when the products are identical on the focal attribute.⁴

3.2 Consumers

There is a unit mass of consumers. Each consumer's utility from Product j , U_j , is given by:

$$U_j = v_j - p_j \tag{1}$$

where v_j is the utility from the focal attribute and p_j is the price.

Following Klibanoff et al. (2005) and Maccheroni et al. (2013), we model v_j in terms of the expected value of the attribute and the ambiguity aversion penalty. The expected value of

³This modeling focus is motivated by the goal of providing more transparency about the paper's insights. We have also analyzed a model that includes one more, non-focal attribute, about which consumers do not face model uncertainty. All the qualitative results of the paper hold in that model as well.

⁴We have also characterized the case in which $m_A = m_B$. Because the effects of model uncertainty are substantially weaker in that setting, the paper focuses on differentiated products. We briefly discuss the symmetric case in the conclusion.

the attribute, $\mathbb{E}(m_j)$, is given by

$$\mathbb{E}(m_j) = \lambda\alpha m_j - (1 - \lambda)\beta m_j \quad (2)$$

where λ represents the consumer's belief that the focal attribute is beneficial and $1 - \lambda$ represents their belief that the focal attribute is harmful. The parameter α is the coefficient of the attribute level m_j if the attribute is beneficial and the parameter β is the coefficient of the attribute level m_j if the attribute is harmful. We assume that $\alpha > 0, \beta > 0$ and allow for $\alpha \geq \beta$, so that the upside of the attribute could be greater than or weaker than its downside. Consumers are heterogeneous in their subjective beliefs about which model is correct: Each consumer i is indexed by $\lambda_i \in (0, 1)$, where λ_i denotes the i th consumer's subjective probability that the attribute is beneficial. This consumer's subjective probability that the attribute is harmful is given by $1 - \lambda_i$. We assume that λ_i is distributed uniformly in the population: $\lambda_i \sim U(0, 1)$. When $\lambda_i = \frac{1}{2}$, a consumer places equal weight on the beneficial and harmful models, reflecting neutrality in terms of how the focal attribute should be evaluated. In contrast, $\lambda_i \neq \frac{1}{2}$ reflects a more confident, though not necessarily correct, prior over a specific model, possibly formed through limited information from reading about the attribute or from some prior experience with the product. Importantly, λ_i is the second-order belief that governs how consumers interpret possible benefits/harms due to an attribute conditional on observing the attribute level, not how much of the attribute is present in a given product.

Let $\phi_{ij} : \mathbb{R} \rightarrow \mathbb{R}$ be a strictly increasing function that captures consumer i 's attitude toward ambiguity presented by Product j . As in Maccheroni et al. (2013), we model consumers' aversion to the model uncertainty (i.e., ambiguity aversion) by the quadratic function ϕ below.

$$\phi_{ij} = \frac{\theta}{2}\lambda_i(1 - \lambda_i)(\alpha + \beta)^2 m_j^2. \quad (3)$$

Given this function, ϕ_{ij} , we can define consumers' perceived utility of the focal attribute as follows:⁵

$$v_{ij}(\lambda) = \lambda_i\alpha m_j - (1 - \lambda_i)\beta m_j - \frac{\theta}{2}\lambda_i(1 - \lambda_i)(\alpha + \beta)^2 m_j^2. \quad (4)$$

so that $U_{ij} = \lambda_i\alpha m_j - (1 - \lambda_i)\beta m_j - \frac{\theta}{2}\lambda_i(1 - \lambda_i)(\alpha + \beta)^2 m_j^2 - p_j$.

For notational convenience, we henceforth suppress the subscript i and write λ to denote a consumer of type λ .

It is helpful to decompose consumer utility in terms of a *valence effect*, $\mathcal{V}_j$ and an *ambiguity*

⁵We use the term "perceived utility" to capture the contribution of the focal attribute of product j to the overall utility of Product j ; we can also think of it as the certainty equivalence.

effect, $\mathcal{A}_j$ as follows.

$$U_j(\cdot) = \mathcal{V}_j - \mathcal{A}_j - p_j, \quad (5)$$

where $\mathcal{V}_j(\lambda) \equiv (\alpha\lambda - (1-\lambda)\beta)m_j$ and $\mathcal{A}_j(\lambda) \equiv \frac{\theta}{2}(\alpha+\beta)^2\lambda(1-\lambda)m_j^2$. In this decomposition, $\mathcal{V}_j$ represents the expected value of the product considering the consumer's uncertainty about the valence of the attribute's effect. On the other hand, $\mathcal{A}_j$ is the negative effect (penalty) of this uncertainty (ambiguity).

To ensure that v_{ij} remains monotonically increasing in the valence belief for any $\lambda_i \in (0, 1)$ as in Klibanoff et al. (2005), we restrict attention to $\theta < \bar{\theta} \equiv \frac{2}{(\alpha+\beta)(m_A+m_B)}$.

3.3 Impact of Consumers' Use of AI

Consumers' domain knowledge, access to scientific knowledge and available time are all limited. This results in them being uncertain whether the focal attribute is beneficial or harmful. We consider the possibility that consumers can use Large Language Models (LLMs) to gain more information about the ambiguous attribute. LLMs can search through vast amounts of literature and make inferences in a very short amount of time. As a result, they have become capable sources of information and advice for consumers. However, they remain imperfect. A substantial literature documents that contemporary AI systems can generate factually incorrect statements, fabricate supporting evidence, and provide erroneous answers with considerable confidence. These hallucinations arise because LLMs are designed to generate plausible text rather than to recover objective truth (Maynez et al., 2020; Ji et al., 2023). Even frontier models continue to exhibit meaningful error rates (OpenAI, 2023; Bubeck et al., 2023). Consulting an AI system does not necessarily eliminate model uncertainty for consumers. Rather, AI can provide additional information whose accuracy is imperfect and, more importantly for our study, whose realization may vary across users.

Beyond ordinary prediction errors that might be unbiased, AI-generated advice may systematically depend on the characteristics and beliefs of the user. A growing literature (see for example, Perez et al. 2023, Sharma et al, 2024) in Computer Science shows that AI can exhibit sycophancy, which is the tendency to provide answers aligned with a user's stated views or expectations even at the cost of accuracy. A recent paper (Cheng et al., 2025) suggests that LLMs often validate users' perspectives, affirm their framing of situations, and avoid challenging their underlying assumptions, even when a more balanced response would be appropriate. This behavior arises from reinforcement learning with human feedback (RLHF), which rewards agreeable answers, and is amplified by the personalization

features that make AI assistants commercially valuable (Sharma et al., 2024). Thus, two consumers presenting the same underlying question in different ways may receive systematically different recommendations from the same AI system. In the context of our model where the beneficiality of the focal attribute is unclear, the advice provided by AI may itself be subject to distortions arising from sycophantic tendencies that reinforce consumers’ pre-existing views.

With this background, we model AI as a technology that revises a consumer’s second-order belief to its reported value λ_r , with a *personalized* distortion that captures sycophancy. We focus on cases where the science is settled, so AI knows the truth, denoted by $\lambda_e \in \{0, 1\}$. For a consumer with prior λ_i , the AI reports λ_r by minimizing the following loss function:

$$\mathcal{L}(\lambda_r) = (1 - \psi) \underbrace{(\lambda_r - \lambda_e)^2}_{\text{distance from truth}} + \psi \cdot \mathbf{1} \left\{ |\lambda_i - \lambda_e| > \frac{1}{2} \right\} \cdot \underbrace{(\lambda_r - \lambda_i)^2}_{\text{disagreement with user}}. \quad (6)$$

where $\psi \in (0, 1)$ calibrates sycophancy intensity. Sycophancy activates only when reporting the truth would require moving the user across the midpoint of the belief space, i.e., $|\lambda_e - \lambda_i| > \frac{1}{2}$.⁶ AI thus reveals the truth to users whose priors are already aligned with it, but partially distorts its report otherwise.

3.4 Sequence of Events

To examine the impact of AI, we consider two regimes, one in which consumers do not use AI and the other in which consumers use AI that can reduce their model uncertainty. Within each regime, the firms choose their prices simultaneously and then consumers decide whether or not to make a purchase, and which product to purchase if they choose to make a purchase.

4 Analysis

We begin our analysis when there is no AI to resolve consumers’ model uncertainty and then proceed to examine the impact of AI.

⁶We choose the threshold of 1/2 for neutrality given the uniform belief space.

4.1 Consumers Do Not Use AI

In this case, each firm chooses its price and then consumers make their purchase decision. We start our discussion with Stage 2 in which the two firms' demands are affected by their pricing decisions in Stage 1.

4.1.1 Consumer Preferences for Two Products

Because Firm A offers a higher level of the focal attribute than Firm B, in the absence of model uncertainty (and ambiguity aversion), all consumers will prefer Product A to Product B if the focal attribute is beneficial, i.e., $\beta = 0$. The reverse will be true if the focal attribute is harmful, i.e., $\alpha = 0$. However, in the presence of model uncertainty, this is not the case. In order to further analyze the difference in consumer preferences for the two products, we define $\Delta_{AB}(\cdot) = v_A - v_B$, which is also the difference $U_A - U_B$ when $p_A = p_B$.

Proposition 1 below characterizes how model uncertainty affects consumer utilities for the two products.

Proposition 1. *(Single-crossing property) For any $\theta \in \mathbb{R}^+$, the difference in perceived utilities, $\Delta_{AB}(\cdot)$, satisfies the single-crossing property in λ such that:*

(i) $\Delta_{AB}(\cdot)$ is strictly increasing in λ .

(ii) There exists a unique cutoff, $\bar{\lambda} = \frac{\kappa - 1 + \sqrt{(1 - \kappa)^2 + \frac{4\kappa\beta}{\alpha + \beta}}}{2\kappa} \in (0, 1)$, such that $\Delta_{AB}(\bar{\lambda}; \theta) = 0$, where $\kappa \equiv \frac{\theta}{2}(\alpha + \beta)(m_A + m_B) < 1$.

To build intuition for Proposition 1, consider consumers with very low values of λ . They prefer a *lower* level of the focal attribute because they believe the attribute is more likely to be harmful. Therefore, they prefer Product B, which has the lower expected attribute level ($m_B < m_A$), over Product A, resulting in $\Delta_{AB}(\cdot) < 0$.⁷ Consumers with higher values of λ progressively assign a higher probability to the attribute being beneficial. As a result, as λ increases, consumers' preference for the lower level of the focal attribute (Product B) decreases and their preference for the higher level of the focal attribute (Product A) increases. For the consumer at $\lambda = \bar{\lambda}$, $v_A = v_B$ and this consumer has the same perceived utility of the focal attribute from both products, resulting in $\Delta_{AB} = 0$. Combining parts (i) and (ii) of Proposition 1, we can also observe that consumers with $\lambda > \bar{\lambda}$ will see $v_A > v_B$ and consumers with $\lambda < \bar{\lambda}$ will see $v_B > v_A$.

⁷For ease of exposition, we use the term preference in comparing v_A and v_B , which are non-price parts of the utilities.

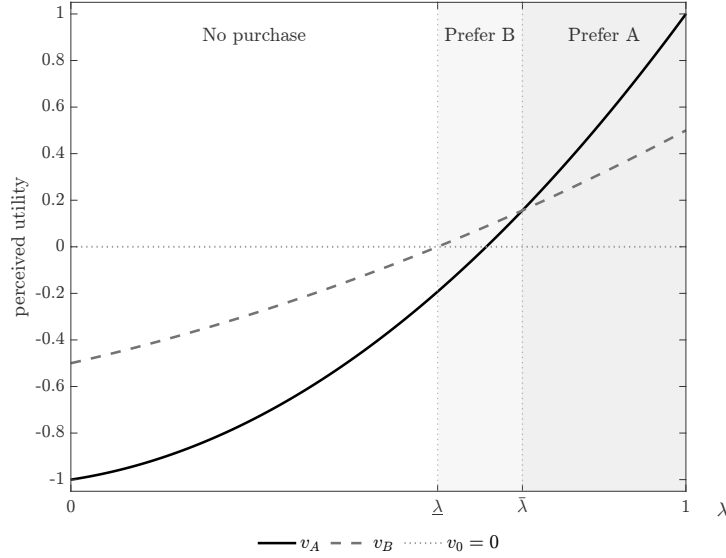


Figure 1: Consumer utility net of prices (v_j) ($m_A = 1, m_B = 0.5, \theta = 0.8, \alpha = 1, \beta = 1$.)
Note: $\bar{\lambda}$ is the point at which the solid black v_A curve intersects with the dashed gray v_B curve. $\underline{\lambda}$ is where the dashed gray v_B curve intersects with the dotted gray v_0 line (no-purchase utility).

The consumer at $\lambda = \frac{1}{2}$ holds neutral beliefs and assigns equal probability to the attribute being beneficial or harmful. Therefore, the above discussion might suggest that $v_A > v_B$ for all consumers with $\lambda > \frac{1}{2}$ because they assign a higher probability to the focal attribute having a positive valence. However, that initial intuition is incomplete. We find that for this neutral consumer, $v_B > v_A$ and that $\bar{\lambda} > \frac{1}{2}$ when $\alpha = \beta$. In other words, holding all else constant, ambiguity aversion leads consumers to overweight the worse outcome (i.e., harmful effects of the attribute), tilting their evaluation towards Product B and affording a competitive advantage to Product B. However, when $\alpha > \beta$, $\bar{\lambda} < \frac{1}{2}$ is possible and $\alpha \leq \beta$ results in $\bar{\lambda} > \frac{1}{2}$.

Product B's demand (and hence the total demand for both products) is also determined by consumers' choice between buying Product B and not buying either product. It turns out that the market will be partially covered and some of the lowest λ consumers will not buy either product.

4.1.2 Equilibrium Prices

The previous discussion about consumers' preferences for the two products indicates that the market will be divided into three endogenously determined segments: inactive consumers, consumers of Product B and consumers of Product A. As a result, the profit functions for

the two firms as $\Pi_A = (1 - \lambda_{AB})$ and $\Pi_B = (\lambda_{AB} - \lambda_B)$, where the cutoff value λ_{AB} denotes the consumer indifferent between buying Product A and Product B and λ_B denotes the consumer indifferent between buying Product B and remaining inactive.

Recall that we can decompose consumer utility, U_j in terms of a *valence effect*, $\mathcal{V}_j$ and an *ambiguity effect*, $\mathcal{A}_j$, with $\mathcal{V}_j(\lambda) \equiv (\alpha\lambda - (1 - \lambda)\beta)m_j$ and $\mathcal{A}_j(\lambda) \equiv \frac{\theta}{2}(\alpha + \beta)^2\lambda(1 - \lambda)m_j^2$. The valence effect becomes positive as λ increases. On the other hand, the ambiguity effect becomes stronger as λ gets farther from the midpoint $\lambda = 1/2$. These two effects also influence the nature of competition between the two products. Specifically, λ_B and λ_{AB} are implicitly defined as:

$$\mathcal{V}_B(\lambda_B) - \mathcal{A}_B(\lambda_B) - p_B = 0 \quad (7)$$

$$\Delta\mathcal{V}(\lambda_{AB}) - \Delta\mathcal{A}(\lambda_{AB}) - p_A + p_B = 0, \text{ where,} \quad (8)$$

$\Delta\mathcal{V}(\lambda_{AB}) = \mathcal{V}_A(\lambda_{AB}) - \mathcal{V}_B(\lambda_{AB})$ and $\Delta\mathcal{A}(\lambda_{AB}) = \mathcal{A}_A(\lambda_{AB}) - \mathcal{A}_B(\lambda_{AB})$. As the valence and ambiguity effects shape the competition between two products, they also influence equilibrium prices. Proposition 2 discussed this further.

Proposition 2. *Equilibrium prices, p_A^* and p_B^* are defined by:*

$$(i) \ p_A^* = (1 - \lambda_{AB}^*) (\Delta\mathcal{V}'(\lambda_{AB}^*) - \Delta\mathcal{A}'(\lambda_{AB}^*)),$$

$$(ii) \ p_B^* = (\lambda_{AB}^* - \lambda_B^*) \left(\frac{1}{\Delta\mathcal{V}'(\lambda_{AB}^*) - \Delta\mathcal{A}'(\lambda_{AB}^*)} + \frac{1}{\mathcal{V}'_B(\lambda_B^*) - \mathcal{A}'_B(\lambda_B^*)} \right)^{-1}.$$

Given the nonlinear demand functions, we can only characterize equilibrium prices in the implicit form described above. However, we find that $p_A^* > p_B^*$. This is due to a complex interplay of the valence and ambiguity effects. In particular, with $m_A > m_B$, Firm A suffers from a more severe ambiguity penalty than Firm B: $\Delta\mathcal{A}(\lambda) \equiv \frac{\theta}{2}(\alpha + \beta)^2(m_A^2 - m_B^2)\lambda(1 - \lambda) > 0$. On the other hand, when it comes to the valence effect, Firm A enjoys a competitive advantage over Firm B among higher λ consumers. This can be seen by noting that $\Delta\mathcal{V}(\lambda) \equiv (\alpha\lambda - (1 - \lambda)\beta)(m_A - m_B)$ is positive when $\lambda > \frac{\beta}{\alpha + \beta}$ and is negative when $\lambda < \frac{\beta}{\alpha + \beta}$. With this pattern of competitive advantages, Firm A adopts a premium pricing strategy and serves higher λ consumers who believe that the focal attribute's upside is more likely. Firm B focus on the relatively lower λ consumers who prefer Product B due to both ambiguity and valence effects. But the lower value of λ reduces these consumers' willingness-to-pay for the product. Thus Product B charges a lower price in the equilibrium.

At this stage, we want to clarify how the competition between Product B and the no-purchase options affects equilibrium outcomes. Product B consumers closest to λ_B^* are

the consumers who are most skeptical about the focal attribute. For these consumers, Product B, by offering a lower attribute level, mitigates the downside exposure in the terms of the valence effect as well as the ambiguity effect. Thus, Product B increases market participation in the presence of model uncertainty. For consumers who are even more skeptical (with $\lambda \in (0, \lambda_B^*)$) this mitigation is inadequate and they will remain inactive in the market.

We note that model uncertainty results in a hybrid form of product differentiation that combines elements of both horizontal and vertical differentiation. When it comes ambiguity effect, all consumers prefer Product B to Product A, bringing in an element of vertical differentiation. On the other hand, the valence effect creates horizontal differentiation in which lower λ consumer prefer Product B whereas higher λ consumers prefer Product A. Overall, consumers’ model uncertainty helps Product B achieve a positive market share and profit, which it would not have been able to do in the absence of model uncertainty. On the flip side, Product A loses consumers who are relatively more concerned about the potential downside of the focal attribute. Without this concern (i.e., model uncertainty), Product A could capture *all consumers*.

The above discussion also highlights another interesting aspect of model uncertainty. In the case of classical uncertainty (traditional risk), a higher level of attribute could compensate for the disutility of risk. However, that is not the case with model uncertainty in our context. A higher attribute level could exacerbate the disutility of ambiguity and create a competitive disadvantage, as is the case with Product A.

Because the magnitudes of both the ambiguity effect and the valence effect increase in $m_A - m_B$, the above results become stronger as the two products diverge in terms of the level of the focal attribute.

4.2 Impact of AI

As we have noted earlier, AI can reduce model uncertainty but it also exhibits a systematic tendency to tailor responses toward what users appear to want to hear, even at the expense of accuracy (Cheng et al., 2025, Hong et al., 2025). Sycophantic AI must therefore first infer what the user wants to hear from linguistic cues in the dialogue. A consumer who asks “I’ve heard X might be harmful, is that true?” reveals a different prior than one who

asks “X is great for health, right?”. Modern LLMs are highly responsive to such cues.⁸ Belief updating therefore reflects both how consumers prompt and how the AI responds, and sycophancy can enter at multiple stages.

As we have discussed before, market outcomes depend not only on an individual consumer’s beliefs about the attribute’s valence but also on the heterogeneity across consumers in these beliefs. Therefore, when it comes to the consumers’ use of AI, the key question is not merely whether sycophantic AI changes a given consumer’s beliefs but how it also changes the heterogeneity across consumers. In an idealistic, extreme case, AI that is both immaculately perfect and fully truthful would eliminate the incorrect valence for all consumers. Sycophantic AI, however, may impede learning for some consumers while fully resolving model uncertainty for others. Proposition 2 characterizes how this selective knowledge transfer (or learning process) reshapes customer heterogeneity and the nature of residual ambiguity persisting in the market.

Proposition 3. *When consumers use AI, their posterior belief, λ'_i , is:*

$$\begin{aligned} \lambda_e = 1 : \quad \lambda'_i &= \begin{cases} 1, & \lambda_i \in [\frac{1}{2}, 1], \\ (1 - \psi) + \psi\lambda_i, & \lambda_i \in [0, \frac{1}{2}), \end{cases} \\ \lambda_e = 0 : \quad \lambda'_i &= \begin{cases} \psi\lambda_i, & \lambda_i \in (\frac{1}{2}, 1], \\ 0, & \lambda_i \in [0, \frac{1}{2}]. \end{cases} \end{aligned}$$

Proposition 3 shows that AI does not uniformly reduce ambiguity. Consumers whose prior beliefs lie on the same side of the belief space as the truth receive fully informative recommendations and become certain about the attribute’s valence. In contrast, consumers whose priors conflict with the truth receive partially distorted recommendations that preserve some of their initial uncertainty. As a result, the market is separated or polarized into two broad segments: one segment with no residual ambiguity and another segment that continues to face ambiguity about the attribute’s effects. In other words, selective knowledge sharing by AI creates a polarized distribution of ambiguity. Rather than converging toward a common posterior, consumers become divided based on their prior beliefs. For one group, the ambiguity effect is eliminated ($\mathcal{A} = 0$). For the remaining customers, both valence and ambiguity effects persist.

⁸In practice, the prompt reveals consumer priors implicitly through how consumers frame questions, the examples they cite, the framings they push back on, and the follow-ups they pursue (Jin et al., 2024).

Figure 2 illustrates the polarization and residual model uncertainty for three levels of sycophancy.

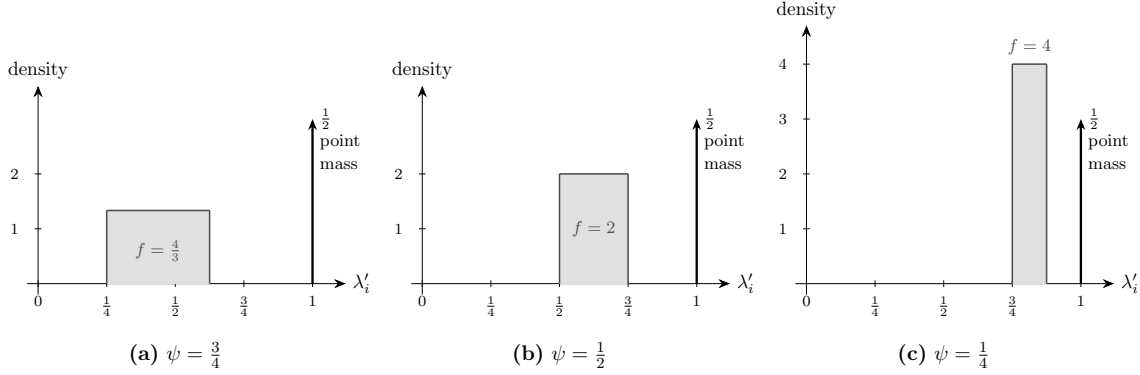


Figure 2: Posterior belief distribution under sycophantic AI ($\lambda_e = 1$). Note: f denotes the density of the continuous segment.

Proposition 3 also shows that the impact of AI will depend on whether the true valence is positive or negative.

Because consumer beliefs are a key determinant of consumers choice among no purchase, Product A and Product B, Proposition 3 has direct implications for product demand and market competition. We first consider a case in which the two firms do not change prices as consumers start using AI. This enables us to isolate and compare the non-price aspects of competition across the two regimes of pre- and post-AI. We later discuss equilibrium prices in the post-AI case.

Proposition 4. *Holding fixed prices, AI has asymmetric impacts on firms A and B:*

- (i) *AI always increases Firm A's demand if true valence is positive.*
- (ii) *AI's impact on Firm B depends on both the true valence and the level of sycophancy. AI increases B's demand iff sycophancy is not too low or too high: $\underline{\psi} < \psi < \bar{\psi}$, where*

$$\underline{\psi}(\lambda_e) \equiv \begin{cases} \frac{1-\lambda_{AB}^*}{1-(\lambda_{AB}^*-\lambda_B^*)} & \text{if } \lambda_e = 1 \\ \frac{\lambda_B^*}{1-(\lambda_{AB}^*-\lambda_B^*)} & \text{if } \lambda_e = 0 \end{cases}, \quad \text{and } \bar{\psi}(\lambda_e) \equiv \begin{cases} \frac{2(1-\lambda_B^*)}{1+2(\lambda_{AB}^*-\lambda_B^*)} & \text{if } \lambda_e = 1 \\ \frac{2\lambda_{AB}^*}{1+2(\lambda_{AB}^*-\lambda_B^*)} & \text{if } \lambda_e = 0 \end{cases}.$$

Proposition 4 shows that AI affects product demand of the two firms asymmetrically, because the two firms are differentiated on different margins.

Recall from the Proposition 2 discussion that Product A's demand comes from customers with $\lambda > \lambda_{AB}^*$, who assign a relatively high probability to the attribute being beneficial.

When the true valence is positive, AI removes model uncertainty for the customers with $\lambda \geq \frac{1}{2}$ so that their valence effect is maximized at $\Delta\mathcal{V} = \alpha(m_A - m_B)$ and their ambiguity penalty is eliminated ($\Delta\mathcal{A} = 0$). As a result, these customers prefer Product A more now than in the previous (pre-AI) case. In addition, AI revises the beliefs of the remaining customers—those with $\lambda < \frac{1}{2}$ —towards the truth. Therefore, they also prefer Product A more now than in the previous (no-AI) case. Thus, Product A’s demand improves when the true valence is positive. This logic is reversed when the true valence is negative, resulting in a lower demand for Product A.

Product B directly competes with Product A for the relatively higher λ consumers and with the no-purchase option for the relatively lower λ consumers. Therefore, Product B’s competitive advantage rests on a combination of ambiguity penalty and a moderate probability of the attribute being positive. Consumers’ use of AI could affect these two factors differentially, resulting in a more complex net effect on Product B’s demand.

To see this, consider the case when the true valence is positive. AI usage could increase Product B’s demand in two ways. First, consumers who previously bought Product A ($\lambda \in (\lambda_B^*, \lambda_{AB}^*)$) switch to Product B. Or the previously inactive consumers ($\lambda \in (0, \lambda_B^*)$) now purchase Product B. If some consumers with $\lambda \geq \frac{1}{2}$ previously bought Product B, AI eliminates their model uncertainty and they will prefer Product A. Consumers with $\lambda < \frac{1}{2}$ receive distorted advice from AI. If the sycophancy parameter, ψ , is relatively low, some of the customers who previously bought Product B will assign a higher probability of the attribute being beneficial and shift to purchasing Product A. Thus, Product B cannot attract the consumers who previously bought Product A if sycophancy is sufficiently low. On the other hand, if AI sycophancy is very high ($\psi > \bar{\psi}$), the customers who were previously inactive will not revise their (incorrect) valence beliefs enough for them to switch over to Product B. Therefore, Product B’s demand increases when the sycophancy level is moderate.

When the true valence is negative, the directions of these effects are reversed, but again with the result that Product B benefits from a moderate level of sycophancy.

Proposition 4 shows that AI affects the product market competition not only by reducing model uncertainty but also by affecting the distribution of consumers’ beliefs. In a standard information acquisition process, the market becomes more informed by all consumers moving closer to the truth. However, sycophantic AI provides differentiated expert assessment to different consumers. As a result, some consumers become fully informed while others retain residual ambiguity. This *redistribution* of ambiguity benefits the two firms differently depending on their product positioning. Product A benefits from the shift in the distribu-

tion towards the true valence, whereas Product B gets its competitive advantage from the residual ambiguity and therefore benefits most when AI is neither fully truthful nor highly sycophantic.

In the analysis that follows, we restrict attention to parameter regions where a pure-strategy Nash equilibrium exists, implying that the post-AI equilibrium cutoffs satisfy:

$$1 - \psi \leq \lambda_B^{**} < \lambda_{AB}^{**} < 1 - \psi/2.$$

Proposition 5. *When the true valence is positive, $p_A^{**} > p_B^{**}$. When the true valence is negative, $p_A^{**} > p_B^{**}$ and $p_A^{**} < p_B^{**}$ both are possible.*

Recall that in the pre-AI base case, Product A is priced higher than Product B, regardless of the true valence. With AI, the two firms' pricing strategies depend on true valence. When the true valence is positive, consumer who originally thought that the valence was more likely to be positive, now have confirmation of their beliefs from AI, eliminating their model uncertainty. In addition, other consumers also update their beliefs more in favor of the valence being positive. Both of these changes increase Product A's competitive advantage, which results in Product A still being priced higher than Product B. When the true valence is negative, the consumers who were originally more skeptical now have their skepticism confirmed. And the remaining consumers will revise their beliefs against the attribute being beneficial. The effect of these changes on Product B, however, is mixed. On the one hand, the leftward shift in the distribution of consumers beliefs increases Product B's competitive advantage over Product A and Product B would be seen as a better product than Product A. On the other hand, the leftward shift in the distribution of beliefs could also make the no-purchase option more attractive than Product B to more consumers. In other words, consumers' use of AI could make the competition between the no-purchase option and Product B more or less intense than the competition between Product A and Product B. This can result in Product B being priced higher or lower than Product A.

Proposition 5 shows the impact of consumers' use of AI on the market power of the two firms. When the true valence is positive, Product A becomes the premium product and AI increases consumers' willingness-to-pay for it. When the true valence is negative, AI makes Product A less attractive relative to Product B, but for some consumers it simultaneously makes Product B less attractive relative to the no-purchase option. Furthermore, *heterogeneous* resolution of uncertainty can constrain Product B's pricing power when the true valence is negative.

As AI shifts consumers' beliefs toward the true valence of the attribute, one might convec-

ture that consumer welfare should improve. This conjecture reflects a standard intuition: reducing uncertainty increases allocative efficiency by reducing informational rents (e.g., Stiglitz, 1975). The analysis below shows that this intuition is, at best, incomplete.

Proposition 6 (Consumer welfare). *When the true valence is positive (negative), consumers' use of AI expands (shrinks) the active market, benefiting consumers who were previously inactive (active). However, existing consumers of both firms could be worse off.*

Proposition 6 shows, when the attribute is truly beneficial, AI partially reduces consumers' model uncertainty and allows some previously inactive consumers to enter the market. The switching behavior of these marginal consumers can sometimes soften competition and raise equilibrium prices for those previously active consumers. When sycophancy level is neither too high nor too low, AI splits consumers into two separate segments: one with no residual model uncertainty and the other retaining sufficiently heterogeneous residual uncertainty. The first segment no longer bears any ambiguity penalty, affording firm A more flexibility to increase prices. Meanwhile, AI also revises the beliefs of the low- λ consumers upward from the second segment so that they now prefer B to the outside option, while residual ambiguity still keeps them from buying A , so B can raise its price to extract surplus from this newly retained segment. These two effects reinforce each other, resulting in both firms raising their prices. As a result, existing consumers of both firms are made worse off, even though AI makes them better informed. When the true valence is negative, the reverse logic applies.

An interesting implication of Proposition 6 is that the benefits of AI are also not distributed uniformly across consumers. The newly active consumers or the newly inactive consumers benefit purchasing a beneficial product or avoiding a harmful purchase. On the other hand, some consumers continue making a purchase. Although AI moves their beliefs closer to the true valence, some of their information gains are captured by the two firms through price changes.

5 Conclusion

Consumers increasingly rely on AI-powered tools to search for products and gather information about specifications, prices, and availability. However, in many product categories, they face another challenge: they may know that a product contains a particular attribute but remain uncertain about whether that attribute is beneficial or harmful. In other words,

they face uncertainty about the model that maps attributes into utility. Beyond facilitating search, AI can also serve as a decision aid that reduces such model uncertainty (ambiguity).

To study this environment, we develop a novel framework where two firms offer products with known but different attribute levels and consumers evaluate these products under model uncertainty, using a smooth ambiguity representation in the tradition of Klibanoff et al. (2005) and Maccheroni et al. (2013).

We first show that, absent AI, model uncertainty generates a new form of product differentiation: a hybrid of vertical and horizontal differentiation. The high-attribute product benefits when consumers believe the attribute is more likely beneficial (a valence effect), but bears a larger ambiguity penalty (an ambiguity effect) than the low-attribute product. Together, these two effects allow the two firms to be artificially differentiated. Both products sustain positive demand precisely because consumers are uncertain about the attribute valence *and* averse to such ambiguity.

We also study how AI affects this artificial differentiation. In our framework, AI does *not* increase information in the standard Bayesian sense, i.e., by refining beliefs over a commonly known state space. Instead, due to its sycophantic tendency to pander to consumers' prior beliefs, AI resolves uncertainty for some consumers while leaving residual uncertainty for others through partially distorted advice. When the attribute is truly beneficial, the high-attribute product benefit from a larger demand; the low-attribute product benefits most under intermediate sycophancy where residual model uncertainty is large enough to retain its existing demand and draws previously inactive consumers into the market.

As with many innovations, AI creates both winners and losers. Consumers who begin buying a beneficial product they once avoided, or who stop buying a harmful one they once wanted, are better off. But because firms can capture part of these informational gains through higher prices, consumers whose initial beliefs were truth-aligned can be worse off. Consumer welfare therefore depends on how AI's expertise is distributed across consumers.

References

- Bewley, T. F. (2002). Knightian decision theory. part i. *Decisions in Economics and Finance*, 25(2):79–110.
- Cheng, M., Yu, S., Lee, C., Khadpe, P., Ibrahim, L., and Jurafsky, D. (2025). Elephant: Measuring and understanding social sycophancy in llms. *arXiv preprint arXiv:2505.13995*.
- Desai, P. S. (2001). Quality segmentation in spatial markets: When does cannibalization affect product line design? *Marketing Science*, 20(3):265–283.
- Gilboa, I. and Schmeidler, D. (1989). Maxmin expected utility with non-unique prior. *Journal of Mathematical Economics*, 18(2):141–153.
- Handel, B. R. and Misra, K. (2015). Robust new product pricing. *Marketing Science*, 34(6):864–881.
- Hong, J., Byun, G., Kim, S., Shu, K., and Choi, J. D. (2025). Measuring sycophancy of language models in multi-turn dialogues. *arXiv preprint arXiv:2505.23840*.
- Iyer, G., Soberman, D., and Villas-Boas, J. M. (2005). The targeting of advertising. *Marketing Science*, 24(3):461–476.
- Jin, Z., Heil, N., Liu, J., Dhuliawala, S., Qi, Y., Schölkopf, B., Mihalcea, R., and Sachan, M. (2024). Implicit personalization in language models: A systematic study. In *Findings of the Association for Computational Linguistics: EMNLP 2024*, pages 12309–12325.
- Klibanoff, P., Marinacci, M., and Mukerji, S. (2005). A smooth model of decision making under ambiguity. *Econometrica*, 73(6):1849–1892.
- Knight, F. H. (1921). *Risk, uncertainty and profit*, volume 31. Houghton Mifflin.
- Lee, J.-Y., Shin, J., and Yu, J. (2025). Communicating attribute importance under competition. *Marketing Science*.
- Leger (2022). Consumer trends 2022: Prioritizing health and wellness. Technical report, Leger (360 Market Reach). Consumer trends report.
- Maccheroni, F., Marinacci, M., and Ruffino, D. (2013). Alpha as ambiguity: Robust mean-variance portfolio analysis. *Econometrica*, 81(3):1075–1113.
- Martineau, P. (2025). Protein powders and shakes contain high levels of lead. *Consumer Reports*. Updated January 8, 2026.

- McKinsey (2025). The agentic commerce opportunity: How AI agents are ushering in a new era for consumers and merchants. Technical report, McKinsey & Company. Published October 17, 2025. Authored by Katharina Schumacher and Roger Roberts with Katharina Giebel.
- Mintel (2025). Us ingredient trends in beauty. Technical report, Mintel Group Ltd. Market report.
- Moorthy, K. S. (1984). Market segmentation, self-selection, and product line design. *Marketing Science*, 3(4):288–307.
- Nau, R. (2011). Risk, ambiguity, and state-preference theory. *Economic Theory*, 48(2):437–467.
- Nau, R. F. (2006). Uncertainty aversion with second-order utilities and probabilities. *Management Science*, 52(1):136–145.
- Savage, L. J. (1954). *The Foundations of Statistics*. John Wiley & Sons, New York.
- Sharma, M., Tong, M., Korbak, T., Duvenaud, D., Askill, A., Bowman, S., Durmus, E., Hatfield-Dodds, Z., Johnston, S., Kravec, S., et al. (2024). Towards understanding sycophancy in language models. In *International Conference on Learning Representations*, volume 2024, pages 110–144.
- Stiglitz, J. E. (1975). The theory of screening, education, and the distribution of income. *The American Economic Review*, 65(3):283–300.

Appendix

Proof of Proposition 1 At $p_A = p_B$, the perceived utility difference is

$$\Delta_{AB}(\lambda) = (\alpha + \beta)(m_A - m_B) \left[(\lambda - \tilde{\lambda}) - \kappa\lambda(1 - \lambda) \right],$$

where

$$\tilde{\lambda} = \frac{\beta}{\alpha + \beta} \in (0, 1), \quad \kappa = \frac{\theta}{2}(\alpha + \beta)(m_A + m_B) < 1.$$

Let $h(\lambda) \equiv (\lambda - \tilde{\lambda}) - \kappa\lambda(1 - \lambda)$. Then

$$\Delta'_{AB}(\lambda) = (\alpha + \beta)(m_A - m_B)h'(\lambda),$$

where $h'(\lambda) = 1 - \kappa(1 - 2\lambda)$.

(i) *monotonicity*. Since $\lambda \in (0, 1)$, we have $1 - 2\lambda < 1$. Hence,

$$h'(\lambda) = 1 - \kappa(1 - 2\lambda) > 1 - \kappa > 0.$$

Therefore, $h(\cdot)$ is strictly increasing in λ . Since $m_A > m_B$, we have $(\alpha + \beta)(m_A - m_B) > 0$, so $\Delta_{AB}(\cdot)$ is strictly increasing in λ . This proves (i).

(ii) *unique cutoff*. The cutoff solves $\Delta_{AB}(\lambda) = 0$, or equivalently, $h(\lambda) = 0$. That is,

$$(\lambda - \tilde{\lambda}) - \kappa\lambda(1 - \lambda) = 0.$$

Rearranging,

$$\kappa\lambda^2 + (1 - \kappa)\lambda - \tilde{\lambda} = 0.$$

Using $\tilde{\lambda} = \frac{\beta}{\alpha + \beta}$, the positive root is

$$\bar{\lambda} = \frac{\kappa - 1 + \sqrt{(1 - \kappa)^2 + \frac{4\kappa\beta}{\alpha + \beta}}}{2\kappa}.$$

Moreover,

$$h(0) = -\tilde{\lambda} < 0, \quad h(1) = 1 - \tilde{\lambda} > 0.$$

Since $h(\cdot)$ is strictly increasing, there exists a unique $\bar{\lambda} \in (0, 1)$ such that $h(\bar{\lambda}) = 0$.

Therefore, $\Delta_{AB}(\bar{\lambda}; \theta) = 0$, and the cutoff is unique. □

Proof of Proposition 2 Recall the utility from buying product j is

$$U_j(\lambda) = v_j(\lambda) - p_j = \mathcal{V}j(\lambda) - \mathcal{A}j(\lambda) - p_j,$$

where

$$\mathcal{V}j(\lambda) = (\alpha\lambda - (1 - \lambda)\beta)m_j, \quad \mathcal{A}j(\lambda) = \frac{\theta}{2}(\alpha + \beta)^2\lambda(1 - \lambda)m_j^2.$$

Recall $\Delta_{AB}(\lambda) \equiv v_A(\lambda) - v_B(\lambda) = \Delta\mathcal{V}(\lambda) - \Delta\mathcal{A}(\lambda)$.

The two cutoffs are implicitly defined by

$$v_B(\lambda_B) - p_B = 0, \tag{A-1}$$

$$\Delta_{AB}(\lambda_{AB}) - p_A + p_B = 0. \tag{A-2}$$

Thus consumers with $\lambda < \lambda_B$ remain inactive, consumers with $\lambda \in [\lambda_B, \lambda_{AB})$ buy B , and consumers with $\lambda \geq \lambda_{AB}$ buy A . Hence demands are

$$d_A = 1 - \lambda_{AB}, \quad d_B = \lambda_{AB} - \lambda_B.$$

Firms choose prices to maximize

$$\Pi_A = p_A(1 - \lambda_{AB}), \quad \Pi_B = p_B(\lambda_{AB} - \lambda_B).$$

Step 1: cutoff responses to prices.

Differentiate (A-2) with respect to p_A :

$$\Delta'_{AB}(\lambda_{AB}) \cdot \frac{\partial \lambda_{AB}}{\partial p_A} - 1 = 0.$$

Therefore,

$$\frac{\partial \lambda_{AB}}{\partial p_A} = \frac{1}{\Delta'_{AB}(\lambda_{AB})} = \frac{1}{\Delta\mathcal{V}'(\lambda_{AB}) - \Delta\mathcal{A}'(\lambda_{AB})}.$$

Similarly, differentiating (A-2) with respect to p_B gives

$$\Delta'_{AB}(\lambda_{AB}) \cdot \frac{\partial \lambda_{AB}}{\partial p_B} + 1 = 0,$$

so

$$\frac{\partial \lambda_{AB}}{\partial p_B} = -\frac{1}{\Delta'_{AB}(\lambda_{AB})} = -\frac{1}{\Delta\mathcal{V}'(\lambda_{AB}) - \Delta\mathcal{A}'(\lambda_{AB})}.$$

Next, differentiate (A-1) with respect to p_B :

$$v'_B(\lambda_B) \cdot \frac{\partial \lambda_B}{\partial p_B} - 1 = 0.$$

Thus,

$$\frac{\partial \lambda_B}{\partial p_B} = \frac{1}{v'_B(\lambda_B)} = \frac{1}{\mathcal{V}'_B(\lambda_B) - \mathcal{A}'_B(\lambda_B)}.$$

Step 2: Firm A's pricing FOC.

Firm A's profit is

$$\Pi_A = p_A(1 - \lambda_{AB}).$$

Differentiating with respect to p_A gives

$$\frac{\partial \Pi_A}{\partial p_A} = (1 - \lambda_{AB}) - p_A \cdot \frac{\partial \lambda_{AB}}{\partial p_A}.$$

Substituting the best response,

$$\frac{\partial \Pi_A}{\partial p_A} = (1 - \lambda_{AB}) \cdot \frac{p_A}{\Delta \mathcal{V}'(\lambda_{AB}) - \Delta \mathcal{A}'(\lambda_{AB})}.$$

At an interior optimum, the first-order condition gives

$$p_A^* = (1 - \lambda_{AB}^*) [\Delta \mathcal{V}'(\lambda_{AB}^*) - \Delta \mathcal{A}'(\lambda_{AB}^*)].$$

Step 3: Firm B's pricing FOC.

Firm B's profit is

$$\Pi_B = p_B(\lambda_{AB} - \lambda_B).$$

Differentiating with respect to p_B gives

$$\frac{\partial \Pi_B}{\partial p_B} = (\lambda_{AB} - \lambda_B) + p_B \left(\frac{\partial \lambda_{AB}}{\partial p_B} - \frac{\partial \lambda_B}{\partial p_B} \right).$$

Using the best responses above,

$$\frac{\partial \Pi_B}{\partial p_B} = (\lambda_{AB} - \lambda_B) p_B \left[\frac{1}{\Delta \mathcal{V}'(\lambda_{AB}) - \Delta \mathcal{A}'(\lambda_{AB})} + \frac{1}{\mathcal{V}'_B(\lambda_B) - \mathcal{A}'_B(\lambda_B)} \right].$$

At an interior optimum, the first-order condition implies

$$p_B^* = (\lambda_{AB}^* - \lambda_B^*) \left(\frac{1}{\Delta \mathcal{V}'(\lambda_{AB}^*) - \Delta \mathcal{A}'(\lambda_{AB}^*)} + \frac{1}{\mathcal{V}'_B(\lambda_B^*) - \mathcal{A}'_B(\lambda_B^*)} \right)^{-1}.$$

This proves the equilibrium price characterization.

Below we prove that $p_A^* > p_B^*$.

At the equilibrium cutoff, we have

$$p_A^* - p_B^* = \Delta_{AB}(\lambda_{AB}^*).$$

By Proposition 1, $\Delta_{AB}(\lambda)$ is strictly increasing and admits a unique cutoff $\bar{\lambda} \in (0, 1)$ such that $\Delta_{AB}(\bar{\lambda}) = 0$. Therefore, we have

$$p_A^* > p_B^* \iff \lambda_{AB}^* > \bar{\lambda}.$$

Thus it suffices to show $\lambda_{AB}^* > \bar{\lambda}$.

Step 1. rewrite the equilibrium conditions as a single equation.

Define

$$G(\lambda) \equiv (1 - \lambda)\Delta'_{AB}(\lambda) - \Delta_{AB}(\lambda).$$

Using Firm A's FOC,

$$p_A = (1 - \lambda_{AB})\Delta'_{AB}(\lambda_{AB}),$$

together with the cutoff condition $\Delta_{AB}(\lambda_{AB}) = p_A - p_B$, yields $p_B = G(\lambda_{AB})$.

Let $\Lambda_B(\lambda)$ denote the unique solution in $(0, 1)$, whenever it exists, to $v_B(\Lambda_B(\lambda)) = G(\lambda)$.

At an interior equilibrium,

$$\lambda_B^* = \Lambda_B(\lambda_{AB}^*),$$

and Firm B's FOC can be written as $\Phi(\lambda_{AB}^*) = 0$, where

$$\Phi(\lambda) = G(\lambda) - (\lambda - \Lambda_B(\lambda)) \left(\frac{1}{\Delta'_{AB}(\lambda)} + \frac{1}{v'_B(\Lambda_B(\lambda))} \right)^{-1}. \quad (\text{A-3})$$

We will show that $\Phi(\lambda) > 0$, for any $\lambda \leq \bar{\lambda}$ and for which $\Lambda_B(\lambda)$ is well defined. This rules out any equilibrium cutoff weakly below $\bar{\lambda}$.

Step 3. show that $\Lambda_B(\lambda) > 2\lambda - 1$ for any $\lambda < \bar{\lambda}$.

By Proposition 1, for any $\lambda \leq \bar{\lambda}$, we must have

$$\Delta_{AB}(\lambda) \leq 0, \quad \Delta'_{AB}(\lambda) > 0,$$

so

$$G(\lambda) = (1 - \lambda)\Delta'_{AB}(\lambda) - \Delta_{AB}(\lambda) > 0.$$

Since $v_B(0) = -\beta m_B < 0$, the equation $v_B(x) = G(\lambda)$ can fail to have a solution only if $G(\lambda) \geq v_B(1) = \alpha m_B$. In that case no consumer purchases product B, contradicting the characterization of interior partial-coverage equilibrium. Hence every equilibrium cutoff lies in the domain where Λ_B is well defined.

We next show that

$$\Lambda_B(\lambda) > 2\lambda - 1 \tag{A-4}$$

If $2\lambda - 1 < 0$, then (A-4) is immediate since $\Lambda_B(\lambda) > 0$.

Suppose $2\lambda - 1 \geq 0$. Since v_B is strictly increasing, showing (A-4) is equivalent to showing $G(\lambda) > v_B(2\lambda - 1)$. Let $q \equiv m_B/m_A \in (0, 1)$. We can rewrite the condition as follows:

$$G(\lambda) - v_B(2\lambda - 1) = \left(\frac{\alpha + \beta}{1 + q} \right) \cdot [c_0(\lambda) + c_1(\lambda)q + c_2(\lambda)q^2], \tag{A-5}$$

where, let $\tilde{\lambda} = \frac{\beta}{\alpha + \beta}$, we have

$$\begin{aligned} c_0(\lambda) &= 1 + \tilde{\lambda} - 2\lambda - \kappa(3\lambda^2 - 4\lambda + 1), \\ c_1(\lambda) &= 1 + \tilde{\lambda} - 2\lambda, \\ c_2(\lambda) &= -\kappa(1 - \lambda)^2 < 0. \end{aligned}$$

RHS of (A-5) is a product of two terms: the first term is strictly positive, the second is concave in q , so it suffices to check its sign at $q = 0$ and $q = 1$.

Recall that, for any $\lambda \leq \bar{\lambda}$, we have $\tilde{\lambda} \geq \lambda - \kappa\lambda(1 - \lambda)$. Substituting this lower bound gives

$$c_0(\lambda) \geq (1 - \lambda)[(1 - \kappa) + 2\kappa\lambda] > 0,$$

and thus

$$c_0(\lambda) + c_1(\lambda) + c_2(\lambda) \geq 2(1 - \lambda)[1 - \kappa(1 - \lambda)] > 0.$$

Hence the RHS of (A-5) is positive for all $q \in [0, 1]$, proving $\Lambda_B(\lambda) > 2\lambda - 1$. Equivalently,

$$1 - 2\lambda + \Lambda_B(\lambda) > 0. \quad (\text{A-6})$$

Step 4. Show $\Phi(\cdot) > 0$ for $\lambda \leq \bar{\lambda}$ where $\Lambda_B(\lambda)$ is well-defined.

If $\lambda \leq \Lambda_B(\lambda)$, then by (A-3), we immediately have

$$\Phi(\lambda) \geq G(\lambda) > 0.$$

If $\lambda > \Lambda_B(\lambda)$, then

$$\left(\frac{1}{\Delta'_{AB}(\lambda)} + \frac{1}{v'_B(\Lambda_B(\lambda))} \right)^{-1} < \Delta'_{AB}(\lambda),$$

so

$$\Phi(\lambda) > (1 - 2\lambda + \Lambda_B(\lambda))\Delta'_{AB}(\lambda) - \Delta_{AB}(\lambda).$$

By (A-6) proved in step 3, and given that $\Delta'_{AB}(\lambda) > 0$ and $\Delta_{AB}(\lambda) \leq 0$, it is straightforward to see that $\Phi(\lambda) > 0$.

The equilibrium cutoff satisfies $\Phi(\lambda_{AB}^*) = 0$. Since $\Phi(\lambda) > 0$ for all admissible $\lambda \leq \bar{\lambda}$ where $\Lambda_B(\cdot)$ is well-defined, it follows that $\lambda_{AB}^* > \bar{\lambda}$. Hence we have $p_A^* > p_B^*$. $\square$

Proof of Proposition 3 Fix any $\psi \in (0, 1)$ and prior belief $\lambda_i \in (0, 1)$. The AI chooses λ_r to minimize equation (6). Recall that the posterior belief is induced by the AI report, we set $\lambda'_i = \lambda_r^*$, where λ_r^* is the minimizer of $\mathcal{L}(\lambda_r)$. Since the loss function is strictly convex in λ_r whenever $\psi \in (0, 1)$, the minimizer is unique.

We examine the two cases below.

Case 1 ($\lambda_e = 1$). When $\lambda_e = 1$, the sycophancy term is active if and only if $\lambda_i < \frac{1}{2}$. If $\lambda_i \geq \frac{1}{2}$, the sycophancy term is inactive, so the loss function reduces to

$$\mathcal{L}(\lambda_r) = (1 - \psi)(\lambda_r - 1)^2.$$

The unique minimizer is therefore $\lambda_r^* = 1$.

Hence, $\lambda'_i = 1$ for $\lambda_i \in [\frac{1}{2}, 1]$.

If instead $\lambda_i \in [0, \frac{1}{2})$, the sycophancy term is active, so

$$\mathcal{L}(\lambda_r)(1 - \psi)(\lambda_r - 1)^2 + \psi(\lambda_r - \lambda_i)^2.$$

Differentiating with respect to λ_r gives

$$\frac{\partial \mathcal{L}}{\partial \lambda_r} = 2(1 - \psi)(\lambda_r - 1) + 2\psi(\lambda_r - \lambda_i).$$

Setting this equal to zero yields

$$(1 - \psi)(\lambda_r - 1) + \psi(\lambda_r - \lambda_i) = 0.$$

Rearranging,

$$\lambda_r - (1 - \psi) - \psi\lambda_i = 0,$$

so $\lambda_r^* = (1 - \psi) + \psi\lambda_i$. Thus, $\lambda_i' = (1 - \psi) + \psi\lambda_i$, for $\lambda_i \in [0, \frac{1}{2})$.

Case 2 ($\lambda_e = 0$). When $\lambda_e = 0$, the sycophancy term is active if and only if $\lambda_i > \frac{1}{2}$.

If $\lambda_i \leq \frac{1}{2}$, the sycophancy term is inactive, so the loss reduces to

$$\mathcal{L}(\lambda_r) = (1 - \psi)\lambda_r^2.$$

The unique minimizer is therefore $\lambda_r^* = 0$.

Hence, $\lambda_i' = 0$ for $\lambda_i \in [0, \frac{1}{2}]$.

If instead $\lambda_i \in (\frac{1}{2}, 1]$, the sycophancy term is active, so

$$\mathcal{L}(\lambda_r) = (1 - \psi)\lambda_r^2 + \psi(\lambda_r - \lambda_i)^2.$$

Differentiating with respect to λ_r gives

$$\frac{\partial \mathcal{L}}{\partial \lambda_r} = 2(1 - \psi)\lambda_r + 2\psi(\lambda_r - \lambda_i).$$

Setting this equal to zero yields

$$(1 - \psi)\lambda_r + \psi(\lambda_r - \lambda_i) = 0.$$

Rearranging,

$$\lambda_r - \psi\lambda_i = 0,$$

so $\lambda_r^* = \psi\lambda_i$. Thus, $\lambda'_i = \psi\lambda_i$ for $\lambda_i \in (\frac{1}{2}, 1]$.

Combining the two cases gives

$$\begin{aligned} \lambda_e = 1 : \quad \lambda'_i &= \begin{cases} 1, & \lambda_i \in [\frac{1}{2}, 1], \\ (1 - \psi) + \psi\lambda_i, & \lambda_i \in [0, \frac{1}{2}), \end{cases} \\ \lambda_e = 0 : \quad \lambda'_i &= \begin{cases} \psi\lambda_i, & \lambda_i \in (\frac{1}{2}, 1], \\ 0, & \lambda_i \in [0, \frac{1}{2}]. \end{cases} \end{aligned}$$

□

Proof of Proposition 4 Throughout this proposition, demands are measured over the pre-AI prior distribution, and consumers make purchase decisions by comparing their post-AI posterior belief λ'_i to the pre-AI equilibrium cutoffs λ_{AB}^* and λ_B^* .

Case 1: $\lambda_e = 1$.

When $\lambda_e = 1$, posterior beliefs are

$$\lambda'_i = \begin{cases} 1, & \lambda_i \in [\frac{1}{2}, 1], \\ (1 - \psi) + \psi\lambda_i, & \lambda_i \in [0, \frac{1}{2}). \end{cases}$$

Firm A's post-AI demand is the mass of consumers with $\lambda'_i > \lambda_{AB}^*$. All consumers with $\lambda_i \in [\frac{1}{2}, 1]$ satisfy this condition. Among consumers with $\lambda_i < \frac{1}{2}$, the condition becomes

$$(1 - \psi) + \psi\lambda_i > \lambda_{AB}^* \iff \lambda_i > \frac{\lambda_{AB}^* - 1 + \psi}{\psi}.$$

Thus AI shifts posterior beliefs toward the right end and weakly expands the set of consumers whose posterior exceeds λ_{AB}^* . Hence A's demand increases.

Now consider Firm B. Firm B's post-AI demand is the mass of consumers whose posterior belief lies between the two cutoffs: $\lambda_B^* < \lambda'_i < \lambda_{AB}^*$. Only consumers with $\lambda_i < \frac{1}{2}$ can have an interior posterior belief, since consumers with $\lambda_i \geq \frac{1}{2}$ are mapped to $\lambda'_i = 1$ and therefore buy A. Hence B's post-AI demand is the length of the set $\lambda_B^* < (1 - \psi) + \psi\lambda_i < \lambda_{AB}^*$.

Equivalently,

$$L(\psi) \equiv \frac{\lambda_B^* - 1 + \psi}{\psi} < \lambda_i < \frac{\lambda_{AB}^* - 1 + \psi}{\psi} \equiv U(\psi).$$

Therefore post-AI demand is

$$d'_B = \min \left\{ \frac{1}{2}, U(\psi) \right\} - \max \{0, L(\psi)\}.$$

Let $d_B \equiv \lambda_{AB}^* - \lambda_B^*$ denote Firm B's pre-AI demand.

Now we characterize when $d'_B > d_B$ by discussing the two cases below.

Case 1.1: $L(\psi) < 0$.

In this case the lower bound is truncated by the feasible belief space:

$$d'_B = U(\psi) = \frac{\lambda_{AB}^* - 1 + \psi}{\psi}.$$

Hence $d'_B > d_B \iff \frac{\lambda_{AB}^* - 1 + \psi}{\psi} > d_B$. Rearranging, $\psi > \frac{1 - \lambda_{AB}^*}{1 - d_B}$. Define $\underline{\psi}(1) \equiv \frac{1 - \lambda_{AB}^*}{1 - d_B}$.

Case 2: $U(\psi) > \frac{1}{2}$.

In this case the upper bound is truncated:

$$d'_B = \frac{1}{2} - L(\psi) = \frac{1}{2} - \frac{\lambda_B^* - 1 + \psi}{\psi}.$$

Hence $d'_B > d_B \iff \frac{1}{2} - \frac{\lambda_B^* - 1 + \psi}{\psi} > d_B$. Rearranging, $\psi < \frac{2(1 - \lambda_B^*)}{1 + 2d_B}$. Define $\bar{\psi}(1) \equiv \frac{2(1 - \lambda_B^*)}{1 + 2d_B}$.

Combining the two conditions yields $d'_B > d \iff \underline{\psi}(1) < \psi < \bar{\psi}(1)$.

Substituting $d_B = \lambda_{AB}^* - \lambda_B^*$ gives $\underline{\psi}(1) = \frac{1 - \lambda_{AB}^*}{1 - (\lambda_{AB}^* - \lambda_B^*)}$, and $\bar{\psi}(1) = \frac{2(1 - \lambda_B^*)}{1 + 2(\lambda_{AB}^* - \lambda_B^*)}$.

Case 2: $\lambda_e = 0$.

When $\lambda_e = 0$, posterior beliefs are

$$\lambda'_i = \begin{cases} \psi \lambda_i, & \lambda_i \in (\frac{1}{2}, 1], \\ 0, & \lambda_i \in [0, \frac{1}{2}]. \end{cases}$$

Consumers with $\lambda_i \leq \frac{1}{2}$ are mapped to $\lambda'_i = 0$ and therefore do not buy. Thus Firm B's

post-AI demand comes only from consumers with $\lambda_i > \frac{1}{2}$ whose posterior lies between the two cutoffs: $\lambda_B^* < \psi\lambda_i < \lambda_{AB}^*$. Equivalently,

$$L(\psi) \equiv \frac{\lambda_B^*}{\psi} < \lambda_i < U(\psi) \equiv \frac{\lambda_{AB}^*}{\psi}.$$

Accounting for the feasible upper-half interval $\lambda_i \in (\frac{1}{2}, 1]$, we have the post-AI demand below

$$d'_B = \min\{1, U(\psi)\} - \max\left\{\frac{1}{2}, L(\psi)\right\}.$$

The same two truncation arguments as in Case 1 generate $d'_B > d \iff \underline{\psi}(0) < \psi < \bar{\psi}(0)$, where

$$\underline{\psi}(0) = \frac{\lambda_B^*}{1 - (\lambda_{AB}^* - \lambda_B^*)}, \quad \bar{\psi}(0) = \frac{2\lambda_{AB}^*}{1 + 2(\lambda_{AB}^* - \lambda_B^*)}.$$

Combining the two cases, AI increases Firm B's demand iff $\underline{\psi}(\lambda_e) < \psi < \bar{\psi}(\lambda_e)$, where

$$\underline{\psi}(\lambda_e) \equiv \begin{cases} \frac{1 - \lambda_{AB}^*}{1 - (\lambda_{AB}^* - \lambda_B^*)} & \text{if } \lambda_e = 1, \\ \frac{\lambda_B^*}{1 - (\lambda_{AB}^* - \lambda_B^*)} & \text{if } \lambda_e = 0, \end{cases} \quad \text{and } \bar{\psi}(\lambda_e) \equiv \begin{cases} \frac{2(1 - \lambda_B^*)}{1 + 2(\lambda_{AB}^* - \lambda_B^*)} & \text{if } \lambda_e = 1, \\ \frac{2\lambda_{AB}^*}{1 + 2(\lambda_{AB}^* - \lambda_B^*)} & \text{if } \lambda_e = 0. \end{cases}$$

□

Proof of Proposition 5 We prove the two cases respectively. Note that all post-AI cutoffs and prices are denoted by double stars.

Case 1: positive true valence ($\lambda_e = 1$).

Under $\lambda_e = 1$, the pure strategy equilibrium cutoffs satisfy

$$1 - \psi \leq \lambda_B^{**} < \lambda_{AB}^{**} < 1 - \frac{\psi}{2}.$$

Firm A's equilibrium price is

$$p_A^{**} = (1 - \lambda_{AB}^{**})\Delta'_{AB}(\lambda_{AB}^{**}).$$

(i) First suppose post-AI market is fully covered ($\lambda_B^{**} = 1 - \psi$), so that

$$p_B^{**} = (\lambda_{AB}^{**} - 1 + \psi)\Delta'_{AB}(\lambda_{AB}^{**}).$$

Then

$$p_A^{**} - p_B^{**} = \left[2 - \psi - 2\lambda_{AB}^{**}\right] \Delta'_{AB}(\lambda_{AB}^{**}).$$

Since $\lambda_{AB}^{**} < 1 - \frac{\psi}{2}$, we have $2 - \psi - 2\lambda_{AB}^{**} > 0$. Therefore, $p_A^{**} > p_B^{**}$.

(ii) Now suppose post-AI market is partially covered ($\lambda_B^{**} > 1 - \psi$). Then

$$p_B^{**} = (\lambda_{AB}^{**} - \lambda_B^{**}) \left(\frac{1}{\Delta'_{AB}(\lambda_{AB}^{**})} + \frac{1}{v'_B(\lambda_B^{**})} \right)^{-1}.$$

Since

$$\left(\frac{1}{\Delta'_{AB}(\lambda_{AB}^{**})} + \frac{1}{v'_B(\lambda_B^{**})} \right)^{-1} < \Delta'_{AB}(\lambda_{AB}^{**}),$$

it follows that

$$p_B^{**} < (\lambda_{AB}^{**} - \lambda_B^{**}) \Delta'_{AB}(\lambda_{AB}^{**}).$$

Using $\lambda_B^{**} \geq 1 - \psi$, $\lambda_{AB}^{**} - \lambda_B^{**} \leq \lambda_{AB}^{**} - 1 + \psi$. Hence

$$p_B^{**} < (\lambda_{AB}^{**} - 1 + \psi) \Delta'_{AB}(\lambda_{AB}^{**}).$$

But we have already shown that

$$(\lambda_{AB}^{**} - 1 + \psi) \Delta'_{AB}(\lambda_{AB}^{**}) < (1 - \lambda_{AB}^{**}) \Delta'_{AB}(\lambda_{AB}^{**}) = p_A^{**}.$$

Therefore, $p_A^{**} > p_B^{**}$.

In both subcases (i) and (ii), when the true valence is positive, $p_A^{**} > p_B^{**}$.

Case 2: negative true valence ($\lambda_e = 0$).

Under $\lambda_e = 0$, the active consumer segment lies in $\left(\frac{\psi}{2}, \psi\right]$, so the pure strategy equilibrium cutoffs satisfy:

$$\frac{\psi}{2} < \lambda_B^{**} < \lambda_{AB}^{**} < \psi.$$

Firm A's equilibrium price is

$$p_A^{**} = (\psi - \lambda_{AB}^{**}) \Delta'_{AB}(\lambda_{AB}^{**}).$$

Consider first the post-AI full-coverage regime in which

$$p_B^{**} = \left(\lambda_{AB}^{**} - \frac{\psi}{2} \right) \Delta'_{AB}(\lambda_{AB}^{**}).$$

Then

$$p_A^{**} - p_B^{**} = \left(\frac{3\psi}{2} - 2\lambda_{AB}^{**} \right) \Delta'_{AB}(\lambda_{AB}^{**}).$$

Hence

$$p_A^{**} > p_B^{**} \iff \lambda_{AB}^{**} < \frac{3\psi}{4}, \quad (\text{A-7})$$

Since $\frac{\psi}{2} < \frac{3\psi}{4} < \psi$, both inequalities are feasible within the admissible cutoff region.

Below we show that both sides of (A-7) are realized at interior equilibria.

In this regime the indifference condition $\Delta_{AB}(\lambda_{AB}^{**}) = p_A^{**} - p_B^{**}$ reads

$$\Delta_{AB}(\lambda_{AB}^{**}) = \left(\frac{3\psi}{2} - 2\lambda_{AB}^{**} \right) \Delta'_{AB}(\lambda_{AB}^{**}).$$

Dividing by the common factor $(\alpha + \beta)(m_A - m_B) > 0$, this becomes

$$h(\lambda_{AB}^{**}) = \left(\frac{3\psi}{2} - 2\lambda_{AB}^{**} \right) h'(\lambda_{AB}^{**}). \quad (\text{A-8})$$

Let $\lambda_{AB}^{**}(\psi)$ denote the root of (A-8) in $(\frac{\psi}{2}, \psi)$ and define $D(\psi) \equiv \lambda_{AB}^{**}(\psi) - \frac{3\psi}{4}$. The left-hand side of (A-8) is continuous in its arguments and $h' > 0$, so $\lambda_{AB}^{**}(\psi)$, and therefore $D(\psi)$, is continuous in ψ .

Consider the following case where $\alpha = \beta = \frac{1}{2}$, $m_A = 1$, $m_B = \frac{1}{2}$, $\theta = \frac{4}{15}$, so that $\kappa = \frac{\theta}{2}(\alpha + \beta)(m_A + m_B) = \frac{1}{5} < 1$. Solving (A-8) gives

$$\psi = 0.6 : \quad \lambda_{AB}^{**} = 0.4835 > \frac{3\psi}{4} = 0.45 \Rightarrow D(0.6) > 0 \Rightarrow p_A^{**} < p_B^{**},$$

$$\psi = 0.8 : \quad \lambda_{AB}^{**} = 0.5833 < \frac{3\psi}{4} = 0.60 \Rightarrow D(0.8) < 0 \Rightarrow p_A^{**} > p_B^{**}.$$

Both are interior equilibria with $\frac{\psi}{2} < \lambda_{AB}^{**} < \psi$.

Thus, both $p_A^{**} > p_B^{**}$ and $p_A^{**} < p_B^{**}$ can arise in equilibrium under $\lambda_e = 0$.

Unlike the positive-valence case, the post-AI price ranking is not uniquely determined when true valence is negative. $\square$

Proof of Proposition 6 Let true consumer welfare from purchasing product j be

$$W_j(\lambda_e) = U_j(\lambda_e) = v_j(\lambda_e) - p_j,$$

and normalize the outside option to zero.

Case 1: $\lambda_e = 1$.

By Proposition 4, AI shifts beliefs toward the positive truth and expands the active market. Thus some consumers who were previously inactive enter the market.

For any newly active consumer,

$$W_j(1) = \alpha m_j - p_j^{**} > 0,$$

since the purchase is preferred to the outside option. Hence newly active consumers are strictly better off.

Consider instead a consumer who purchases the same product j both before and after AI. Then

$$W_j^{**} - W_j^* = (\alpha m_j - p_j^{**}) - (\alpha m_j - p_j^*) = p_j^* - p_j^{**}.$$

Case 2: $\lambda_e = 0$.

By Proposition 4, AI shifts beliefs toward the negative truth and shrinks the active market. Thus some consumers who were previously active choose the outside option.

For any such consumer,

$$W_j^* = -\beta m_j - p_j^* < 0,$$

whereas $W_j^{**} = 0$. Hence exiting consumers are strictly better off.

For consumers who remain active,

$$W_j^{**} - W_j^* = (-\beta m_j - p_j^{**}) - (-\beta m_j - p_j^*) = p_j^* - p_j^{**}.$$

Thus, in both cases, existing consumers can be worse off whenever post-AI prices increase

$$p_j^{**} > p_j^* \implies W_j^{**} < W_j^*.$$

It is easy to verify there exists parameters $(\alpha, \beta, m_j, \theta, \psi) \in \Theta$ resulting in post-AI price increase in both cases:

$$\exists \Theta : p_A^{**} > p_A^*, \quad p_B^{**} > p_B^*.$$

□