

Manipulative consumers

MICHAEL RICHTER, Baruch College, CUNY, USA and Royal Holloway, University of London, UK
NIKITA ROKETSKIY, University College London, UK

We study optimal monopoly pricing with evasive consumers. The monopolist uses consumer data to price its products and data-conscious consumers manipulate their data at a cost. We derive the monopolist's gains from using data and characterize the optimal investigation strategy. When there is a large number of observable consumer attributes, we find that the value of data is all-or-nothing: either the data identifies a sizeable number of high-type consumers allowing the monopolist to extract the full surplus from them or it is entirely worthless. We establish a precise condition which delineates when data is valuable, and when it is not.

1 INTRODUCTION

Sellers' use of consumer data for price discrimination is as old as the hills, but, recently, the amount and variety of available data has rapidly multiplied. This would spell trouble for consumers, if not for the fact that many of them are aware of these practices and have some control over their own data. Moreover, with a few exceptions, consumers are not liable for falsifying or manipulating the data that is harvested by sellers.

Thus, one of the key privacy-related questions is this: How are the spoils of trade divided between data-hungry sellers and data-conscious consumers? We answer this question using the canonical monopoly framework in which the seller has access to a full price-discrimination toolkit.

In our model, the seller's optimal pricing is a combination of the two instruments that are commonly used in practice: consumer profiling and menus. The seller employs all information that he has about individual consumers—consumer data—to organize them into specific categories. We refer to these categories as market segments and buyers are treated differently across segments and uniformly within. Because consumer data does not always perfectly explain variation in preferences, the seller screens the remaining intrasegment heterogeneity using a menu of vertically differentiated goods.

When the monopolist employs consumer data to segment the market, he solves a problem akin to a regression of market demand on consumer attributes. The explained part of variation in this regression represents the value of data for the monopolist. As such, the monopolist is interested in maximizing the fraction of customer demand variation that is explained by the consumer data. However, the consumers' interests are opposed. When some parts of their data become overly informative of their demand, the consumers muddle them. This limits the sellers' gains from segment-specific pricing, but it is costly for the consumers to manipulate their data, hence some explanatory power can remain.

Because the data is endogenous to pricing, the derived value of the data is not a function of prior beliefs, but rather a function of the structure of the data and the cost of manipulation. We show that the seller gains from using the data, but this gain can be vanishing when the data becomes very rich—i.e., when it contains a large number of (conditionally) independent

† Richter: michael.richter@baruch.cuny.edu, Roketskiy: n.roketskiy@ucl.ac.uk

variables.¹ Whether this is the case or not, i.e. whether data is valuable, hinges upon the underlying market's economic parameters. We establish a single condition which specifies when data has value, and when it doesn't (Proposition 6).

This condition derives from a key insight. In the limiting case, the usefulness of data is all-or-nothing. If the monopolist can use the data in order to perfectly screen some high-type consumers, then he would do so, fully extracting their surplus. If the monopolist is unable to perfectly separate a non-trivial mass of such consumers, then the usefulness of data to him is vanishing as the data grows large. Put differently, the data cannot be profitably used to partially refine the monopolist's estimate of a consumer's willingness to pay. Rather, either the monopolist can use data to fully establish some consumers' type for the purpose of full profit extraction, or he cannot.

However, even when the data's value vanishes in the limit, the loss of consumer surplus from the monopolist's use of data can remain substantial.

The dimensionality of consumer data takes center stage in this result. Consumers manipulate their data to get a better deal from the seller. The more high-dimensional the data is, the more combinations of attribute values will correspond to desirable deals. Thus, the consumers can rely less on the number of manipulated attributes and more on selecting which attributes to manipulate—i.e., less on the costly and more on the costless aspect of data manipulation.

The hide and seek nature of monopolist-consumer relations means that the seller would benefit from limiting consumers' understanding of how data affects prices. One way to achieve this is to use a random, and therefore unpredictable, subset of variables that are contained in the original data. We show that such a tactic would resolve the issue of vanishing value when the data is rich, but implementing it would require commitment power that sellers may not possess.

2 RELATED LITERATURE

The general problem of inference from muddled data is studied by Frankel and Kartik (2019, 2022) (also, see earlier related work by Kartik, 2009). Ball (2022) studies scoring of strategic agents in which the data for scoring is provided by an intermediary. The latter may serve as a valuable commitment provider to the agency that develops the scoring rule. Ball (2022) shows that randomization by the data intermediary inhibits data manipulation because the target of manipulation becomes more difficult to identify. A distantly related result in the context of information elicitation with verification is obtained by Carroll and Egorov (2019).

Milli, Miller, Dragan and Hardt (2019) study welfare costs associated with threshold binary classifiers when subjects can manipulate classifiers inputs strategically at a cost. They show that a seller who attempts to redesign the classifier to counter strategic manipulations imposes a disproportionately high cost on the subjects. Similar problems are studied by Cunningham and Moreno De Barreda (2022), Perez-Richet and Skreta (2022) and Hardt, Megiddo, Papadimitriou and Wootters (2016).

Incentives to misreport data are not unique to classification problems. Eliaz and Spiegelger (2022) study incentives under regression estimators and Caner and Eliaz (2021) investigate the same question with an additional issue of variable selection and regularization.

¹In our context, we abstract away from the problems associated with finite samples. We assume that seller can sample data from the continuum of consumers and the term "rich" refers to the number of variables, not observations.

Deneckere and Severinov (2022) and Severinov and Tam (2018) study costly misreporting in classic asymmetric information frameworks of signaling and screening. Liang and Madsen (forthcoming) study the use of observables in provision of effort when subjects' productivity is private but correlated with these observables. Dana, Larsen and Moshary (2023), Tan (2023) and Perez-Richet and Skreta (2023) adopt a mechanism design approach in which the agents must report their type (e.g., the consumers must report their individual demand) and they incur a cost when their reports are not truthful. Our interpretation of manipulable consumer data is more literal than theirs. We assume that the seller has to use statistical tools to infer the true demand from the data.

Hu, Immorlica and Vaughan (2019) study strategic manipulation of data by multiple agents. They highlight an externality one subject imposes on others by manipulating their own record. In our framework, similar externality considerations are present.

The use of data by sellers is tightly related to consumer privacy (see Acquisti, Taylor and Wagman (2016) for review of the literature on consumer privacy). Bonatti and Cisternas (2020) investigate how the seller can condition current prices on the consumers' past choices via an aggregator that assigns a score to each consumer. Bhaskar and Roketskiy (2021) consider unrestricted use of past purchases for equilibrium optimal pricing. Conitzer, Taylor and Wagman (2012) study how consumers' control over their past records affects their welfare.

Bonatti, Huang and Villas-Boas (2023) draw a connection between the value of privacy and the concavity of the indirect utility as a function of market beliefs. Their approach is complementary to ours. In their papers, privacy results in pooling of consumption for various types of consumers. In our case, the data is valuable to the seller if it allows him to differentiate his offering. A natural way to measure this value is via the dispersion: an expectation of the quadratic (convex) function of the price.

Acknowledging that some degree of privacy is desirable, Eilat, Eliaz and Mu (2021) suggest a Bayesian measure of privacy protection and use this measure to derive optimal privacy-preserving pricing. They show that an optimal privacy-preserving menu always contains a finite number of alternatives.

We study a monopoly pricing problem where the seller uses statistical inference given increasingly fine data (which consumers can manipulate). Relatedly, Frick, Iijima and Ishii (2024) consider a monopoly pricing problem where the seller can sell a fixed set of goods to a single consumer. The monopolist uses data to estimate how much the consumer is willing to pay for each of these goods.

In their setting, due to the reliability of the consumer data, the first-best is achievable in the limit whereas in our setting, this is often not the case, and in fact, the value of data can even vanish. Given that the first-best is achievable in the limit in their setting, the question becomes how to achieve it, and they find that selling the grand bundle has as good convergence properties as the second-best, while selling goods individually performs worse. Put differently, sufficiently high quality data does away with the pricing complications associated with multi-product monopoly.

We study how the seller can use data to segment the market and tailor prices to the demand in each segment. Our definition of market segment is closely related to the one in Yang (2022). In our setting, the consumer data is available to the seller at a negligible cost. When sellers use consumer data for both pricing and product design, Ichihashi (2020) and Hidir and Vellodi (2021) show that it is possible to find a segmentation that relies on consumers volunteering the private information on their valuations. Ali, Lewis and Vasserman (2022)

study how consumers' control of personal data affects consumer surplus, industry profits and overall welfare (also see Ali, Lewis and Vasserman, 2023).

For a monopoly pricing problem, Bergemann, Brooks and Morris (2015) introduce and solve an information design problem. They characterize all consumer-producer surplus pairs that can be achieved with different information structures. Elliott, Galeotti, Koh and Li (2021) expands upon this approach by considering the case of competing sellers who sell different goods. In contrast, in our setting, the information structure is given, and furthermore, consumers can manipulate it, and we investigate how much use a profit-maximizing seller can make of this given manipulated data?

In general, in the literature on information design, the data that the seller possesses can be simply summarized as a posterior belief. In contrast, in our setting, data is physical, and while it can be manipulated, the costs of such manipulation depends upon the actual data record, and not simply on its informational content.

We study the value that the seller attaches to consumer data. A related question is how to sell the data to the monopolist and what is the resulting price. Taylor (2004), Bergemann and Bonatti (2015), Bergemann, Bonatti and Smolin (2018) and Segura-Rodriguez (2021) answer this question in a variety of settings.

3 THE MODEL

The market consists of a single seller and a mass of consumers indexed by $i \in C = [0, 1]$. A variety of goods can be produced and sold by the seller. Each type of good is characterized by a quality parameter $q \in \mathbb{R}_+$. There are two types of consumers who vary by their taste for quality: those with low marginal willingness to pay for quality, $t_\ell \in \mathbb{R}_+$, and those with a high willingness to pay, $t_h \in \mathbb{R}_+$. For brevity, we will refer to these as low- and high-type consumers, respectively. The difference between these is denoted by $\delta = t_h - t_\ell > 0$ and a consumer i 's marginal willingness to pay for quality is denoted by $\tau(i)$.

Each consumer demands at most one good. We assume that consumers have linear utility $2t(i)q$ from a good of quality q , and production costs are quadratic, q^2 . Thus, if consumer i buys a good of quality q , the following social surplus is realized:²

$$s(i, q) = 2\tau(i)q - q^2$$

A transaction price p determines how the surplus is split between the seller and the buyer:

$$u(i, q, p) = s(i, q) - p$$

$$\pi(i, q, p) = p,$$

where u is the buyer's payoff and π is the seller's profit from this transaction (p is the seller's absolute markup over production cost). The consumers' outside option is valued at zero.

The seller can use consumer data to price the goods. We assume that the data is freely available, but the consumers can, at a cost, secretly manipulate their own records before the seller harvests the data.³ We take a very broad view of what constitutes consumer data: here, it is everything that a seller can observe about individual consumers, ranging from browsing history to computer make/model, operating system, geographic location, occupation

²This is just a convenient normalization of utility being $t(i)q$ and costs being $\frac{1}{2}q^2$ and all results hold irrespective of this normalization.

³The consumers cannot deny the seller access to their data, but they can tamper with their records in an attempt to mislead the seller.

as reported on credit histories, and medical histories. To model this we assume the following. Each consumer i is endowed with K *ex ante* private attributes represented by the vector $\omega(i) \in \mathcal{A} = \{0, 1\}^K$. The vector of *ex post* public attributes—i.e., the attributes that can be observed by the seller—is *chosen* by consumer i and is denoted by $\alpha(i) \in \mathcal{A} = \{0, 1\}^K$. The cost of choosing this vector of attributes is

$$\frac{\|\alpha(i) - \omega(i)\|_1}{K}c,$$

where $c > 0$. This cost is simply c/K multiplied by the *share* of attributes the consumer manipulates. In particular, a consumer can always choose to not manipulate the data, so $\omega(i) = \alpha(i)$, at zero cost.

Consumer data generically contains information about demand. To quantify this, consider a collection of attribute vectors $S \subseteq \mathcal{A}$ and let λ be the Lebesgue measure on C . We denote the mass of consumers with quality valuation t_ℓ and a public vector of attributes from S by $m(S)$ and the mass of consumers with quality valuation t_h and a public vector of attributes from S by $n(S)$ as follows:⁴

$$\begin{aligned} m(S) &= \lambda(\{i \in C : \tau(i) = t_\ell, \alpha(i) \in S\}), \text{ and} \\ n(S) &= \lambda(\{i \in C : \tau(i) = t_h, \alpha(i) \in S\}). \end{aligned}$$

For any set of attribute vectors S , all the seller needs to know about demand to set prices optimally for this set is the proportion of consumers with different willingness to pay for quality. This *hazard ratio* is

$$h(S) = \frac{n(S)}{m(S)}.$$

By $\bar{m}$ and $\bar{n}$ we denote the total masses of consumers with valuations t_ℓ and t_h respectively, i.e., $\bar{m} = m(\mathcal{A})$ and $\bar{n} = n(\mathcal{A})$. The aggregate hazard ratio for the entire market is $\bar{h} = \frac{\bar{m}}{\bar{n}}$.

Finally, we assume that the consumers are anonymous and can be distinguished by the seller only through their *ex post* public attributes α . Thus, the offer made to a consumer i depends only on $\alpha(i)$ and cannot depend upon i explicitly.

3.1 Optimal menu pricing

Suppose the seller offers the same profit-maximizing screening menu to a group of consumers with attribute vectors in $S \subseteq \mathcal{A}$. Let $\pi(i)$ be the profit that results from consumer i purchasing her favorite item from this menu. By $\rho(S)$ we denote the total profit over all consumers i such that $\alpha(i) \in S$, normalized by the number of the consumers with low valuation:

$$\rho(S) = \frac{1}{m(S)} \int_{\alpha(i) \in S} \pi(i) di,$$

As shown by Mussa and Rosen (1978), the profit-maximizing menu consists of two items: A premium item (p_h, q_h) which is designed for consumers with high valuation for quality, and a basic item (p_ℓ, q_ℓ) designed for those with low valuation. As is standard, the participation constraint binds for the lowest valuation and the incentive constraint binds down. That is, the basic item is priced so that low-type consumers get 0 utility from it, $p_\ell = 2t_\ell q_\ell$, and the premium item is designed so that high-type consumers are indifferent between both items,

⁴We assume that functions τ , α and ω are measurable.

$2t_h q_h - p_h = 2t_h q_\ell - p_\ell$. Thus, given these two constraints, the seller solves the following profit-maximization problem:

$$\rho(S) = \max_{q_h \geq q_\ell \geq 0} \{2t_\ell q_\ell - q_\ell^2 - 2h(S)q_\ell \delta + h(S)(2t_h q_h - q_h^2)\}.$$

The solution is

$$\begin{aligned} p_\ell &= t_\ell \max\{0, t_\ell - h(S)\delta\} \\ p_h &= t_h^2 - \delta \max\{0, t_\ell - h(S)\delta\} \\ q_h &= t_h \\ q_\ell &= \max\{0, t_\ell - h(S)\delta\}. \end{aligned}$$

When shopping from this menu, consumers with the low willingness to pay receive no surplus. A high-willingness-to-pay consumer i with $\alpha(i) \in S$ receives surplus

$$u(i) = \max\{0, 2\delta(t_\ell - h(S)\delta)\}.$$

The monopolist's profit is

$$\rho(S) = h(S)t_h^2 + (\max\{0, t_\ell - h(S)\delta\})^2.$$

The maximization appears in the above calculation because in each segment, the monopolist may wish to serve both types of consumers (so the second term equals $[t_\ell - h(S)\delta]^2$) or forgo the low consumers and only serve the high types (so the second term equals 0). The seller forgoes the low types when $t_\ell - h(S)\delta < 0$, i.e. whenever the hazard ratio is above the threshold $H = \frac{t_\ell}{\delta}$.

We now consider the case (Assumption 1) where the economy is relatively balanced, neither with an extremely high or extremely low number of high-types. While this assumption is only on the exogenous aggregate proportion of high types, we will show that it actually binds the endogenous post-manipulation proportion of high types for every segment. Thus, importantly it guarantees that the monopolist serves all consumers in all market segments.

ASSUMPTION 1.

$$\bar{h} \in \left(\frac{c}{\delta^2}, H - \frac{c}{\delta^2} \right).$$

From now on (and until Section 4.3), we will maintain this assumption. Under this assumption, we will characterize the maximal possible value of data, and show that, surprisingly, the value of data vanishes as the number of attributes becomes large.

In Section 4.3, we relax this assumption, and find another surprise, as the number of attributes grows large—the value of data is “all or nothing”, either it vanishes, or it allows for a full separation of some high-types, and may even emulate the first-degree price discrimination outcome, i.e. the first-best.

3.2 Group pricing with data

Eliciting consumers' valuations via a screening menu is costly. Relative to the first-best, the seller reduces the price of the premium item and the quality of the basic item. But, data can help. Specifically, the seller can use consumer data to estimate consumers' valuations and, then rely on menu pricing only to elicit residual uncertainty that is not explained by the available data. Thus, the monopolist can profitably use data by utilizing both second-degree (menu pricing) and third-degree (different menus to different data profiles) price discrimination.

To get a better understanding of the third-degree price discrimination component in the seller's pricing decision, let us consider a notion of market segmentation. Formally, let $\mathcal{S} = \{S_1, S_2, \dots, S_N\}$ be a partition of set $\mathcal{A}$. Every element of this partition corresponds to a market segment that can be identified by the seller using the consumer data.

The seller's total profit under this market segmentation $\mathcal{S}$ is

$$\pi_{\mathcal{S}} = \sum_{S \in \mathcal{S}} m(S) \rho(S) = \pi^* + \delta^2 \sum_{S \in \mathcal{S}} m(S) \left(h(S) - \bar{h} \right)^2,$$

where $\pi^* = \bar{m} \rho(\mathcal{A})$ is the profit of the monopolist without any market segmentation, or equivalently, with the coarsest segmentation possible (i.e., when every consumer gets the same menu). The second term of this expression represents the profit gain the seller could get by segmenting the market. Perhaps not surprisingly, this term is proportional to the weighted variance of the hazard ratio across different market segments. The whole purpose of segmenting the market is that the seller is able to fine tune his offers to the demand in each segment. The variance—a measure of spread—quantifies how different these offers are across the segments. In practice, the seller would prefer to use the finest market segmentation he can given the available consumer data.

3.3 Data manipulation

Naturally, due to third-degree price discrimination, the consumers expect different prices for the same premium quality good across market segments. These differences motivate consumers to perform arbitrage—consumers in high-price segments will manipulate their attributes to “travel” to a segment with a lower price. Note that only high-type consumers find such travel appealing—consumers with low willingness to pay receive their reservation utility in all market segments, and therefore, have no incentive to manipulate their attributes.

So, we focus now on high-type consumers. Recall that under Assumption 1 the price for the high-quality product in segment S is increasing in the hazard ratio in this segment:

$$p_h(S) = h(S) \delta^2 + t_h^2 - t_l \delta.$$

Consider two segments, S_1 and S_2 such that $h(S_1) > h(S_2)$. A consumer's gain from traveling from segment S_1 to S_2 is proportional to the difference between the hazard ratios in these segments $h(S_1) - h(S_2)$. The cost on the other hand is proportional to the share of attributes that need to be changed to travel between the segments. Thus, in equilibrium, the following *no-arbitrage condition* must hold:

$$h(S_1) - h(S_2) \leq \frac{c}{\delta^2 K} \min_{\substack{b \in S_2, \\ a \in S_1: n(a) > 0}} \|a - b\|_1. \quad (1)$$

To understand, recall that the hazard ratio is endogenous in the setup with manipulable data. If the condition does not hold, then there is a positive mass of high valuation consumers with attribute vector $a \in S_1$ who could manipulate it to $b \in S_2$ at a cost that is lower than the gain from switching market segments. When they do so, they reduce the difference in the hazard ratios across the two segments which has distributional consequences for welfare due to the pricing externality imposed on the other high-type consumers in those segments. In particular, if consumers travel from S_1 to S_2 , then the remaining high-type consumers in S_1 benefit due to a price decrease in this segment, the consumers who travel benefit (including

the cost of travel), and high-type consumers in S_2 are harmed due to a price increase in this segment. Overall, for consumers, the gains outweigh the losses, and they benefit on average.

To see this, consider an arbitrage opportunity between S_1 and S_2 :

$$h(S_1) - h(S_2) > \frac{c}{\delta^2 K} \min_{\substack{\mathbf{b} \in S_2, \\ \mathbf{a} \in S_1: n(\mathbf{a}) > 0}} \|\mathbf{a} - \mathbf{b}\|_1 = C.$$

This condition can be rewritten as

$$\frac{n(S_1)}{m(S_1)} - \frac{n(S_2)}{m(S_2)} > C.$$

The total consumer surplus across the two segments after a small mass ϵ of consumers travels from segment S_1 to S_2 is

$$\begin{aligned} \mathfrak{S}(\epsilon) &= \delta \left(t_\ell - \delta \frac{n(S_1) - \epsilon}{m(S_1)} \right) (n(S_1) - \epsilon) + \delta \left(t_\ell - \delta \frac{n(S_2) + \epsilon}{m(S_2)} \right) (n(S_2) + \epsilon) - \epsilon \delta^2 C \\ &= \mathfrak{S}(0) + \delta^2 \epsilon \left(\frac{2n(S_1) - \epsilon}{m(S_1)} + \frac{-2n(S_2) - \epsilon}{m(S_2)} \right) - \epsilon \delta^2 C \\ &= \mathfrak{S}(0) + \epsilon \delta^2 (2h(S_1) - 2h(S_2) - C) - \delta^2 \epsilon^2 \left(\frac{1}{m(S_1)} + \frac{1}{m(S_2)} \right) > \mathfrak{S}(0). \end{aligned}$$

This quantifies the externality due to high-type consumers' travel. They will travel when $h(S_1) - h(S_2) > C$ while aggregate consumer surplus improves as long as $h(S_1) - h(S_2) > C/2$.

4 VALUE OF CONSUMER DATA

By segmenting the market according to $\mathcal{S}$, the seller increases his profit by

$$\delta^2 \sum_{S \in \mathcal{S}} m(S) \left(h(S) - \bar{h} \right)^2.$$

We use this gain to define the value of data for the seller. For analytical purposes it is convenient to represent it in terms of hazard ratios within market segments. However, one could find it more intuitive to represent this gain in terms of dispersion of prices across market segments:

$$\delta^2 \sum_{S \in \mathcal{S}} m(S) \left(h(S) - \bar{h} \right)^2 = \frac{1}{\delta^2} \sum_{S \in \mathcal{S}} m(S) \left(p_h(S) - \frac{1}{\bar{m}} \sum_{\tilde{S} \in \mathcal{S}} m(\tilde{S}) p_h(\tilde{S}) \right)^2,$$

where $p_h(S)$ is the price of the premium item in market segment S and $m(S)$ is the volume of basic items sold in the same segment.

A market segmentation is the seller's choice. Intuitively, the seller should prefer to use *all* available data. To formalize this idea, consider the finest segmentation possible—i.e., the segmentation that associates each attribute vector with a separate market segment:

$$\mathcal{S}^* = \{ \{i \in C \mid \alpha(i) = \mathbf{a}\} \mid \mathbf{a} \in \mathcal{A} \}.$$

PROPOSITION 1. *For any market segmentation $\mathcal{S}$,*⁵

$$\sum_{S \in \mathcal{S}} m(S) \left(h(S) - \bar{h} \right)^2 \leq \sum_{S \in \mathcal{S}^*} m(S) \left(h(S) - \bar{h} \right)^2.$$

⁵One could extend the definition of market segmentation by allowing the seller to assign consumers to segments randomly. The proposition would still hold. We omit this generalization to make the exposition simple.

PROOF. The term $(h(S) - \bar{h})^2$ is convex in the hazard ratios and $\{h(S), m(S)\}_{S \in \mathcal{S}^*}$ is a mean-preserving spread of $\{h(S), m(S)\}_{S \in \mathcal{S}}$ because $\mathcal{S}^*$ is the finest possible segmentation. $\square$

This proposition shows that the seller would always use segmentation $\mathcal{S}^*$ unless he can credibly promise to the consumers to restrict the use of data for price discrimination. The latter would require some commitment mechanism such as outsourcing data collection to an independent intermediary. We consider this possibility in Section 4.2.

Note that under segmentation $\mathcal{S}^*$, if a consumer manipulates any of her attributes, she necessarily changes her market segment as well. Thus, the set of constraints (1) can be simplified to:

$$\forall \mathbf{a}, \mathbf{b} \in \mathcal{A} : |h(\mathbf{a}) - h(\mathbf{b})| \leq \frac{c}{\delta^2} \frac{\|\mathbf{a} - \mathbf{b}\|_1}{K}.$$

The seller's gain from using data depends on the correlation between the consumer attributes and preferences, which is represented by the variance of hazard ratios across different segments of the market. Thus, the gain can be small or even zero if the attributes are independent of the consumer valuations.

On the other hand, no matter how informative the initial attributes are, the gain from using them for price discrimination is bounded by consumers' data manipulation. The higher the informativeness of the attributes, the higher the incentives of the consumers to manipulate them. The very question that we study in this paper is how this tug of war between the seller and the manipulative consumers limits the value of consumer data for the seller.

In pursuit of the maximal potential of consumer data for price discrimination, we define its value as the maximal gain the seller can obtain when consumers can manipulate their attributes. That is, we consider the value of data under the most favorable conditions for the seller.

DEFINITION 1. *The value of consumer data D_K is the largest gain the seller can obtain by segmenting the market according to $\mathcal{S}^*$ across all possible correlations between the attributes and valuations:*

$$D_K = \delta^2 \max_{h: \mathcal{A} \rightarrow \mathbb{R}_+} \sum_{\mathbf{a} \in \mathcal{A}} m(\mathbf{a}) \left(h(\mathbf{a}) - \bar{h} \right)^2 \quad (2)$$

$$s.t. \quad \sum_{\mathbf{a} \in \mathcal{A}} m(\mathbf{a}) \left(h(\mathbf{a}) - \bar{h} \right) = 0, \quad (3)$$

$$\forall \mathbf{a}, \mathbf{b} \in \mathcal{A} : |h(\mathbf{a}) - h(\mathbf{b})| \leq \frac{c}{\delta^2} \frac{\|\mathbf{a} - \mathbf{b}\|_1}{K}. \quad (4)$$

The above conditions are as follows:

- (3) Bayesian plausibility requires that after manipulation, there are the same number of high types as before the manipulation.
- (4) No arbitrage stipulates that no consumer wishes to further manipulate their attributes.

The program that defines D_K highlights the main intuition about the interaction of the data-hungry seller and the data-cautious consumers. On the one hand, it is in the seller's interest to increase the share of preference variation explained by the observed attributes (see objective (2)). On the other hand, if the explanatory power of the attributes increases beyond a certain point, then the consumers erode it via attribute manipulation (see constraint (4)).

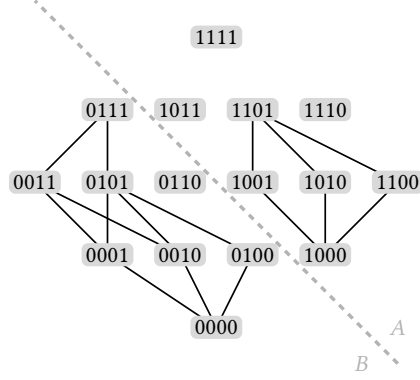


Fig. 1. An example of a constraints graph.

To characterize the program's solution h^* , we need to identify pairs of market segments which are involved in consumer arbitrage. In mathematical terms, we look for program constraints that bind at h^* .⁶ It is convenient to think about the binding constraints as edges of a graph. In particular, for a given h , we define a constraints graph $G(h)$ in the following way.

DEFINITION 2. A graph $G(h)$ with the set of nodes $\mathcal{A}$ is called a constraints graph for h if for every pair of $\mathbf{a}, \mathbf{b} \in \mathcal{A}$ the following statements are equivalent:

- (1) $\mathbf{a}$ and $\mathbf{b}$ are connected,
- (2) $|h(\mathbf{a}) - h(\mathbf{b})| = \frac{c}{\delta^2} \frac{\|\mathbf{a} - \mathbf{b}\|_1}{K}$.

In this graph, every node is a market segment and edges represent arbitrage opportunities, or equivalently, instances of consumers manipulating their attributes. The graph representation of binding constraints allows us to find a simple necessary condition for optimality (2):

PROPOSITION 2. For every solution h^* of the problem (2) the corresponding constraints graph $G(h^*)$ is connected.

PROOF. By contradiction, let h be a solution to the problem (2) for which $G(h)$ is not connected. We can partition the nodes of this graph into two sets A and B in such a way that there are no edges between the two sets, and $h(A) \geq h(B)$.

The objective can be rewritten as:

$$\sum_{\mathbf{a} \in \mathcal{A}} m(\mathbf{a}) (h(\mathbf{a}) - \bar{h})^2 = \overbrace{\sum_{\mathbf{a} \in A} m(\mathbf{a}) (h(\mathbf{a}) - h(A))^2}^{\text{variance within } A} + \overbrace{\sum_{\mathbf{a} \in B} m(\mathbf{a}) (h(\mathbf{a}) - h(B))^2}^{\text{variance within } B} + \overbrace{\sum_{Z \in \{A, B\}} m(Z) (h(Z) - \bar{h})^2}^{\text{variance between } A \text{ and } B}.$$

Because there are no edges in $G(h)$ between A and B , for any $\mathbf{a} \in A$ and $\mathbf{b} \in B$ it holds that $|h(\mathbf{a}) - h(\mathbf{b})| < \frac{c}{\delta^2} \frac{\|\mathbf{a} - \mathbf{b}\|_1}{K}$. Let

$$h_\epsilon(\mathbf{a}) = \begin{cases} h(\mathbf{a}) + \frac{\epsilon}{m(A)}, & \text{if } \mathbf{a} \in A \\ h(\mathbf{a}) - \frac{\epsilon}{m(B)}, & \text{if } \mathbf{a} \in B. \end{cases}$$

⁶Note that the objective function is convex, therefore the solution will be on the boundary of the convex permissible set.

Note that for small enough $\epsilon > 0$, h_ϵ satisfies all the constraints of the problem (2) and increases the value of the objective compared to h because it increases the variance between A and B . Thus, h cannot be a solution to (2). $\square$

This has the following implication in terms of hazard ratios across different market segments:

COROLLARY 1. *If h^* is a solution to the program (2), then the difference in hazard ratios between any two segments is a multiple of $c/(\delta^2 K)$:*

$$\forall \mathbf{a}, \mathbf{b} \in \mathcal{A}, \exists j \in \{0, 1, 2, \dots, K\} : |h^*(\mathbf{a}) - h^*(\mathbf{b})| = \frac{c}{\delta^2} \frac{j}{K}.$$

This corollary sheds light on why dimensionality of the data is important: even though there are 2^K possible realizations of attribute vectors, when data is most informative there are at most K truly distinct market segments. Put differently, many attribute vectors offer identical predictions of the within segment demand.

4.1 More data?

We are now prepared to investigate how the value of data depends on its richness. The parameter that describes the richness of the data is the number of attributes K . However, it is possible to increase the number of attributes and add little or no new information, by making the attributes correlated with each other. A stark example of this is duplicate attributes. To rule this possibility out, we assume that the attributes are conditionally independent. Note that in our setting, the attributes for consumers who have high willingness to pay are endogenous, therefore we only impose the independence condition on the consumers with low willingness to pay.

ASSUMPTION 2. *There exist marginal probabilities $\mu_j : \{0, 1\} \rightarrow \mathbb{R}_+$, $j = 1, \dots, K$, such that for any $\mathbf{a} \in \mathcal{A}$:*

$$m(\mathbf{a}) = \bar{m} \prod_{j=1}^K \mu_j(\mathbf{a}_j).$$

Under Assumption 2, we can use the number of attributes K as a measure of the richness of consumer data because every attribute carries new information and, therefore, improves the estimate of consumer demand. Without loss, we can assume that for every attribute j , it is weakly more likely that a low type has a 0 for that attribute than a 1, that is, $1 - \mu_j(0) = \mu_j(1) \leq \frac{1}{2}$. This is because otherwise the entire problem could be equivalently relabelled by exchanging the values 0 to 1 for attribute j .

This assumption is strong and can be relaxed. In particular, we can allow for some attributes not only to be correlated, but to be identical to each other, as long as the proportion of these attributes does not grow too large when we increase the number of attributes. At the same time, independence allows for a very tractable closed-form characterization of the value of data.

It is important to clarify that we consider the richness of the data in terms of available variables and not observations. In our setup, the seller perfectly understands the data-generating process for any collection of variables. In particular, he understands that different sets of predictors have different predictive power (both exogenously and because of consumers manipulating them).

Because of independence, the solution for program (2) for a large K can be constructed using solutions for the analogous problem for a smaller K . We can use induction on the number of attributes to characterize the solution.

PROPOSITION 3. *If Assumptions 1 and 2 hold, the value of data is*

$$D_K = \frac{1}{K} \bar{m} \left(\frac{c}{\delta} \right)^2 \frac{\sum_{j=1}^K \mu_j(0) \mu_j(1)}{K} \quad (5)$$

PROOF. Recall that objective for the problem (2) for $K + 1$ attributes can be rewritten as

$$\begin{aligned} \sum_{\mathbf{a} \in \mathcal{A}} m(\mathbf{a}) \left(h(\mathbf{a}) - \bar{h} \right)^2 &= \sum_{\mathbf{a} \in A} m(\mathbf{a}) \left(h(\mathbf{a}) - h(A) \right)^2 \\ &\quad + \sum_{\mathbf{a} \in B} m(\mathbf{a}) \left(h(\mathbf{a}) - h(B) \right)^2 \\ &\quad + \sum_{Z \in \{A, B\}} m(Z) \left(h(Z) - \bar{h} \right)^2, \end{aligned} \quad (6)$$

where $A = \{\mathbf{a} \mid \mathbf{a} = (\mathbf{b}, 0), \mathbf{b} \in \mathcal{A}_K\}$ and $B = \{\mathbf{a} \mid \mathbf{a} = (\mathbf{b}, 1), \mathbf{b} \in \mathcal{A}_K\}$. Observe that $\mathcal{A}_{K+1} = A \cup B$ is all attribute vectors with $K + 1$ dimensions, where A is all of those whose last coordinate is 0, and B is all of those whose last coordinate is 1. Note that the first two components of this sum are the values of the objective for the problem (2) for K attributes. We can maximize the third component of the sum and check if the result violates any constraints of the original problem for $K + 1$ attributes. In particular, we solve for the contribution of the $(K + 1)^{\text{th}}$ attribute towards the overall objective:

$$\begin{aligned} V_{K+1} &= \delta^2 \max_{\{h(A), h(B)\}} \sum_{Z \in \{A, B\}} m(Z) \left(h(Z) - \bar{h} \right)^2 \\ &\quad s.t. \quad \sum_{Z \in \{A, B\}} m(Z) \left(h(Z) - \bar{h} \right) = 0 \\ &\quad |h(A) - h(B)| \leq \frac{c}{\delta^2} \frac{1}{K + 1}, \end{aligned}$$

The solution to this problem is

$$\begin{aligned} h(A) - \bar{h} &= \frac{\mu_{K+1}(1)}{(K + 1)} \frac{c}{\delta^2}, \\ h(B) - \bar{h} &= -\frac{\mu_{K+1}(0)}{(K + 1)} \frac{c}{\delta^2}. \end{aligned}$$

Therefore, the maximal contribution of the $(K + 1)^{\text{th}}$ attribute towards the overall objective is:

$$V_{K+1} = \bar{m} \mu_{K+1}(0) \mu_{K+1}(1) \left(\frac{c}{\delta} \frac{1}{K + 1} \right)^2$$

Let $D_K(c)$ be the value of data for K attributes and cost of manipulation c . Combining (6) with the expression for V_{K+1} we get:

$$D_{K+1}(c) = D_K \left(\frac{K}{K+1} c \right) + \bar{m} \mu_{K+1}(0) \mu_{K+1}(1) \left(\frac{c}{\delta} \frac{1}{K+1} \right)^2$$

The solution to this equation is

$$D_K(c) = \frac{1}{K} \bar{m} \left(\frac{c}{\delta} \right)^2 \frac{\sum_{j=1}^K \mu_j(0) \mu_j(1)}{K}$$

To establish if this value is feasible, we construct the solution h_{K+1}^* from a solution to problem with K attributes—i.e., h_K^* . We show that h_{K+1}^* also satisfies all relevant constraints.

Using h_K^* let us define

$$\tilde{h}_{K+1}(\mathbf{a}) = \bar{h} - \frac{K}{K+1} \left(\bar{h} - h_K^*(\mathbf{a}) \right).$$

Note that for any $\mathbf{a}, \mathbf{b} \in \mathcal{A}_K$ and $z \in \{0, 1\}$ the following constraint is satisfied because h^* satisfied all constraints in program (2):

$$\left| \tilde{h}_{K+1}(\mathbf{a}) - \tilde{h}_{K+1}(\mathbf{b}) \right| \leq \frac{c}{\delta^2} \frac{\|(\mathbf{a}, z) - (\mathbf{b}, z)\|_1}{K+1}.$$

The value D_{K+1} is achieved at

$$h_{K+1}(\mathbf{a}) = \begin{cases} \tilde{h}_{K+1}(\mathbf{b}) - \frac{\mu_{K+1}(1)}{K+1} \frac{c}{\delta^2}, & \text{if } \mathbf{a} = (\mathbf{b}, 1) \\ \tilde{h}_{K+1}(\mathbf{b}) + \frac{\mu_{K+1}(0)}{K+1} \frac{c}{\delta^2}, & \text{if } \mathbf{a} = (\mathbf{b}, 0). \end{cases}$$

Note that the proposed solution h_{K+1} is obtained from h_K by applying the same transformation. It has the following features:

(1) for any $\mathbf{a}, \mathbf{b} \in \mathcal{A}_{K+1}$ this transformation ensures that the constraint

$$|h_{K+1}(\mathbf{a}) - h_{K+1}(\mathbf{b})| \leq \frac{c}{\delta^2} \frac{\|\mathbf{a} - \mathbf{b}\|_1}{K+1}$$

is satisfied. In particular, if

(a) $a_{K+1} = b_{K+1}$, this constraint is implied by the corresponding constraint for h_K ,

(b) $a_{K+1} \neq b_{K+1}$, this constraint is satisfied because $\frac{\mu_{K+1}(1)}{K+1} \frac{c}{\delta^2} + \frac{\mu_{K+1}(0)}{K+1} \frac{c}{\delta^2} = \frac{c}{\delta^2} \frac{1}{K+1}$.

(2) it maximizes the objective because it maximizes all three components of the sum in (6). □

Note that Assumption 1 ensures that the constructed solution h is

- (i) positive; and
- (ii) below H .

In Section 4.3, we relax this assumption. It turns out that (ii) is crucial for the result in Propositions 3 and 4, but (i) is not. In particular, note that if we relax the requirement that $h(\cdot) \geq 0$ and allow for negative hazard ratios for some segments, the resulting value of the program will be the valid upper bound. Later, in Proposition 4, we can use this observation and ignore the zero bound on the solution h .

There are two features present in expression (5) worthy of attention. First, the value of data is proportional to the average variance of the attributes for consumers with low valuation:

$$\frac{1}{K} \sum_{j=1}^K \mu_j(0)\mu_j(1).$$

There is a simple intuition for this. While consumers with high valuations manipulate their data, consumers with low valuation stay passive. Market segments with low concentrations of low valuation consumers are precisely the segments that consumers with high valuations are trying to avoid. Thus, the seller benefits from an as equal distribution of low valuation consumers across segments as possible. This attribute diversity among low valuation consumers is achieved by increasing the variance of their attributes.

Across attributes, a higher spread in the probabilities $\mu_j(1)$ leads to a lower value of data because the variance of the binomial distribution is a concave function of the underlying probability. The following corollary formalizes this observation.

COROLLARY 2. *Consider two distinct sequences, $\{\mu_j(1)\}_{j=1}^K \neq \{v_j(1)\}_{j=1}^K$, that induce the same value of data. Their convex combination $\{\alpha\mu_j(1) + (1 - \alpha)v_j(1)\}_{j=1}^K$ has a strictly higher value of data.*

This has strong implications: the value of data is lower when one attribute is very evenly distributed and one is very unevenly distributed (e.g., $\mu_1(1) = 0.1$, $\mu_2(1) = 0.5$) than when both attributes are intermediately distributed (e.g., $\mu_1(1) = \mu_2(1) = 0.3$). This happens because in the first case, the first attribute is disproportionately susceptible to manipulation (in terms of the total number of high-type consumers who manipulate their value from 1 to 0).

Second, the value of data depends on the number of attributes in a stark way. Increasing the number of attributes has two conflicting effects: the ex-ante predictive power of personal data increases; while at the same time, the possibilities for manipulation grow which might make the data less reliable in aggregate. The following proposition proves that, surprisingly, the second effect dominates.

PROPOSITION 4. *(Rich data is worthless) When the number of consumer attributes becomes arbitrarily large, the value of data becomes arbitrarily small:*

$$\lim_{K \rightarrow \infty} D_K = 0.$$

PROOF. This is immediate as the numerator of the value of data grows, at most, at the speed of K (it can grow slower depending upon the underlying μ sequence), while the denominator of the value always grows at the speed of K^2 . Formally, since $\mu_j(1)\mu_j(0) \leq \frac{1}{4}$ for any j , the following simple bound holds:

$$D_K = \frac{1}{K} \bar{m} \left(\frac{c}{\delta} \right)^2 \frac{\sum_{j=1}^K \mu_j(0)\mu_j(1)}{K} \leq \frac{1}{K} \frac{\bar{m}}{4} \left(\frac{c}{\delta} \right)^2.$$

Thus, $\lim_{K \rightarrow \infty} D_K = 0$ uniformly fast, for any μ sequence. $\square$

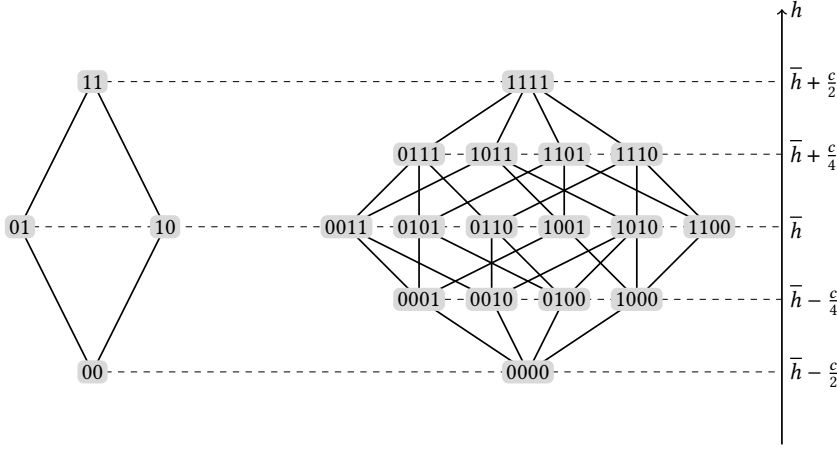


Fig. 2. The optimal market segmentation for $K = 2$ (on the left) and $K = 4$ (on the right) under symmetry.

To gain more intuition about this result, consider the symmetric case in which every segment contains the same mass of consumers with low willingness to pay. That is, $\mu_j(0) = \mu_j(1) = \frac{1}{2}$ for every $j = 1, \dots, K$, and $m(\mathbf{a}) = \frac{\bar{m}}{2^K}$. When the number of attributes K increases, the share of segments with the hazard ratio close to $\bar{h}$ grows whereas the share of segments with either large or small hazard ratios vanishes. This occurs because the mesh of binding no-arbitrage constraints becomes tighter (see Figure 2). Because the data is used to predict deviations in demand from the average, the total value of this prediction becomes smaller and smaller as K grows large.

To summarize, the larger the number of attributes, the more combinations of these attributes correspond to approximately “average” market segments, i.e., those market segments in which the composition of consumers is similar to that of the entire market. This is reminiscent of the weak law of large numbers.

What are the welfare implications of richer consumer data? First, let us set aside the direct welfare costs associated with data manipulation and concentrate on the creation and division of surplus *after* the consumers changed their attributes. When the seller segments the market using consumer data, total welfare is reduced by:

$$\Delta W = -\delta^2 \sum_{\mathbf{a} \in \mathcal{A}} m(\mathbf{a}) \left(h_K(\mathbf{a}) - \bar{h} \right)^2 = -\frac{1}{K} \bar{m} \left(\frac{c}{\delta} \right)^2 \frac{\sum_{j=1}^K \mu_j(0) \mu_j(1)}{K}.$$

Therefore, consumer surplus is also reduced:

$$\Delta \mathfrak{S}_{min} = -\frac{2}{K} \bar{m} \left(\frac{c}{\delta} \right)^2 \frac{\sum_{j=1}^K \mu_j(0) \mu_j(1)}{K}.$$

Just like the gain in the monopoly profit, these values become arbitrarily small when the number of attributes becomes arbitrarily large, i.e., $\Delta W, \Delta \mathfrak{S}_{min} \rightarrow 0$.

The previous expression excludes the cost of data manipulations—it accounts only for quality distortions and higher prices. The aggregate cost of manipulation depends on how informative the *ex-ante* data is (i.e., the private data represented by $\omega(i)$). For example, one extreme scenario is that the *ex-ante* attributes are such that no consumer wants to change them. In that case, the change in consumer surplus is $\Delta\mathfrak{S}_{min}$. To capture the range of welfare losses, consider the upper bound on the total cost of manipulation:

$$\sum_{\mathbf{a} \in \mathcal{A}} m(\mathbf{a}) h_K(\mathbf{a}) \frac{\|\mathbf{a} - \mathbf{a}_m\|_1}{K} c,$$

where $\mathbf{a}_m = \arg \max_{\mathbf{a} \in \mathcal{A}} h_K(\mathbf{a})$. To understand how this upper bound depends on the richness of the data, again consider the symmetric case in which $\mu_i(0) = \mu_i(1) = 1/2$ for all $i = 1, \dots, K$. In this case,

$$h_K(\mathbf{a}) = \left(\frac{K}{2} - \sum_{i=1}^K a_i \right) \frac{c}{K\delta^2} + \bar{h},$$

and $\mathbf{a}_m = (0, \dots, 0)$. For maximal losses we get⁷

$$\sum_{\mathbf{a} \in \mathcal{A}} m(\mathbf{a}) h_K(\mathbf{a}) \frac{\|\mathbf{a} - \mathbf{a}_m\|_1}{K} c = \frac{\bar{m}c^2}{K\delta^2} \left(\sum_{k=1}^K \binom{K}{k} \frac{k}{K} \left(\frac{K}{2} - k + \frac{K\delta^2}{c} \bar{h} \right) \right) = \frac{\bar{m}c}{2} - \frac{1}{K} \frac{\bar{m}c^2}{2\delta^2}$$

If we combine this cost with the expression for consumer surplus we obtained earlier, we get the maximal reduction in consumer surplus:

$$\Delta\mathfrak{S}_{max} = -\frac{\bar{m}c}{2}.$$

Thus, the presence of data reduces the consumer surplus, and in contrast to the effect on the seller's profit, the reduction in consumer surplus may not vanish when the number of consumer attributes becomes large.

4.2 Less data!

The value of data vanishes when consumer attributes become numerous because of the increase in opportunities for manipulation. There are ways for the seller to limit these opportunities. One natural possibility is to commit to using less data: the seller can promise the consumers to use only a small fraction of the variables for pricing. However, this improves a seller's ability to make personalized offers only if the seller is secretive about which attributes are used for pricing. In this section we consider an extreme example of such a policy. The seller commits to using only one attribute without disclosing to the consumers which one.

The result that we obtain in this section echoes similar observations made in other settings about limiting the use of data: Frankel and Kartik (2022) and Ball (2022) show that such a commitment, often implemented by introducing a data intermediary, improves endogenous precision of the data.

In our setting, the seller randomly chooses the attribute for the purpose of segmenting the market. If a particular attribute is not chosen in equilibrium, it becomes very informative, due to the consumers not manipulating it, and, therefore, using it would be a profitable deviation. Recall Assumption 1 from before:

⁷For binomial sums, see Boros and Moll (2004).

$$\bar{h} \in \left(\frac{c}{\delta^2}, \frac{t_\ell}{\delta} - \frac{c}{\delta^2} \right).$$

We now replace that assumption with the following:

ASSUMPTION 1*.

$$\mu_j(0) \in (\psi, 1 - \psi) \text{ where } 0 < \psi < 1$$

and

$$\bar{h} \in \left(\frac{c}{2\delta^2\sqrt{\psi(1-\psi)}}, \frac{t_\ell}{\delta} - \frac{c}{2\delta^2\sqrt{\psi(1-\psi)}} \right).$$

This assumption serves the same purpose as Assumption 1 did previously. It ensures that in each segment, both types of consumers are served.

By γ_j denote the probability of the seller using attribute $j = 1, \dots, K$ for market segmentation. When γ_j is sufficiently large, the no-arbitrage condition (4) becomes

$$|h(a_j = 0) - h(a_j = 1)| \leq \frac{c}{\delta^2} \frac{1}{\gamma_j K}.$$

Assumption 1* guarantees that γ_j is sufficiently large. Then, the largest possible gain from market segmentation based on attribute j is:

$$\bar{m} \left(\frac{c}{\delta} \right)^2 \frac{\mu_j(0)\mu_j(1)}{\gamma_j^2 K^2}.$$

Because the firm chooses the attribute randomly, at the optimum, the gain must be the same for any two attributes. We can find the probabilities γ_j from this condition:

$$\gamma_j = \frac{\sqrt{\mu_j(0)\mu_j(1)}}{\sum_{k=1}^K \sqrt{\mu_k(0)\mu_k(1)}}.$$

The likelihood of the seller using the attribute j for pricing is increasing in the variance of this attribute among low-type consumers, i.e. $\mu_j(0)\mu_j(1)$. As we pointed out in the previous section, this variance measures how (un)attractive the particular attribute is for the purposes of manipulation. If the variance is low, then the returns to manipulation are high.

PROPOSITION 5. *If the seller commits to use only a single attribute for market segmentation without disclosing which attribute exactly, the value of consumer data is*

$$D_K^r = \bar{m} \left(\frac{c}{\delta} \frac{\sum_{j=1}^K \sqrt{\mu_j(0)\mu_j(1)}}{K} \right)^2 \geq \bar{m} \frac{c^2 \psi(1-\psi)}{\delta^2}.$$

If the seller adopts this data policy, then the consumers' expected return to data manipulation becomes smaller. The reason is simple: when a consumer changes the value of attribute i , with probability $1 - \gamma_i$ she does not gain anything because the seller does not use this attribute for pricing. This implies that when the seller does use the attribute for pricing, it contains a great amount of information. This is true for every attribute, and therefore, the value of data is larger compared to the case when the seller uses all available attributes simultaneously.

Note that D_K^r is larger than D_K by a factor of $\sum_{j=1}^K \sqrt{\mu_j(0)\mu_j(1)}$ which is bounded below by $K\sqrt{\psi(1-\psi)}$. This implies that D_K^r does not vanish when K becomes large and in this case, using less data results in higher overall value.

4.3 More data, Less assumptions

In this section, we investigate what happens to the value of data when Assumption 1 does not hold. As it turns out, the value of consumer data responds asymmetrically. If the proportion of high-type consumers is low, there is still no value in the limit for consumer data.

On the other hand, we show that if the proportion of high-type consumers is high, i.e. $\bar{h}$ is above a certain threshold, then the value of data does not vanish as the number of attributes becomes arbitrarily large. To understand why, suppose that there is a segment with a proportion of high-type consumers so large that the monopolist sells only to them and forgoes the low types altogether. The utility of the high-type consumers in this segment is zero and adding more high-types to this segment does not decrease their utility any further. This creates an opportunity: holding all other hazard ratios fixed, the hazard ratio can be increased arbitrarily for this segment—and thus increase the predictive power of the data—without violating the no-arbitrage constraints. Thus, the main result of this section (Proposition 6) complements Proposition 4 by pointing out that data has value in the limit only as far as there is a combination of attributes that precisely identifies high-type consumers.

Recall that Assumption 1 guaranteed that the seller wishes to offer products to both high-type and low-type consumers. As mentioned above, without this assumption, it is possible that on some segments the hazard ratio is high enough that a seller wishes to completely forgo the low-type consumers and focus solely on extracting from high-type consumers. This occurs on any segment where the hazard ratio is above the threshold $H = \frac{t_\ell}{\delta}$. To account for this possibility, we extend the definition of the value of consumer data:

DEFINITION 3. *The value of consumer data D_K is a solution to the following program:*

$$D_K = \delta^2 \max_{h: \mathcal{A} \rightarrow \mathbb{R}_+} \sum_{\mathbf{a} \in \mathcal{A}} m(\mathbf{a}) \left[\left(h(\mathbf{a}) - \bar{h} \right)^2 + (\max\{0, \bar{h}\delta - t_\ell\})^2 - (\max\{0, h(\mathbf{a})\delta - t_\ell\})^2 \right] \quad (7)$$

$$\text{s.t.} \quad \sum_{\mathbf{a} \in \mathcal{A}} m(\mathbf{a}) \left(h(\mathbf{a}) - \bar{h} \right) = 0, \quad (8)$$

$$\forall \mathbf{a}, \mathbf{b} \in \mathcal{A} : |\min\{H, h(\mathbf{a})\} - \min\{H, h(\mathbf{b})\}| \leq \frac{c}{\delta^2} \frac{\|\mathbf{a} - \mathbf{b}\|_1}{K} \quad (9)$$

$$\forall \mathbf{a} \in \mathcal{A} : h(\mathbf{a}) \geq 0 \quad (10)$$

Before proceeding to our main result, Proposition 6, we will present a sequence of lemmas which will be helpful. The first lemma characterizes the solution to program (7) when $h(\mathbf{a}) \geq H$ for at least one segment $\mathbf{a}$. In particular, for such a solution, if at least one segment has a high hazard ratio ($> H$), then all other hazard ratios are set below this threshold, and as low as possible while respecting the no-arbitrage constraints. It also has an important implication that in any segmentation which induces maximal value of the data, there will be at most one segment $\mathbf{a}$ for which $h(\mathbf{a}) \geq H$.

LEMMA 1. *Let h be a solution to program (7). If for some segment $\mathbf{a}$ it holds that $h(\mathbf{a}) \geq H$, then for any other segment $\mathbf{b}$,*

$$h(\mathbf{b}) = \max \left\{ 0, H - \frac{c\|\mathbf{a} - \mathbf{b}\|_1}{K\delta^2} \right\}.$$

PROOF. We proceed by induction on $\|\mathbf{a} - \mathbf{b}\|_1$.

Base Case: It must be that for every $\mathbf{b} \neq \mathbf{a}$, $h(\mathbf{b}) \leq \max\{0, H - \frac{c}{K\delta^2}\}$. If not, then h can be reduced to $\max\{0, H - \frac{c}{K\delta^2}\}$ at every segment for which $h(\mathbf{b}) > \max\{0, H - \frac{c}{K\delta^2}\}$ and $h(\mathbf{a})$ is simultaneously increased. In words, high types are shifted from segments which violate this condition to the segment $\mathbf{a}$ where they can be fully extracted. This shift causes some high types to be fully extracted when they were not before, thus increasing the value of program (7).

Importantly, notice that no-arbitrage condition is satisfied by this shift because:

- (i) No consumer wishes to deviate to a segment whose hazard ratio is reduced since their utility from doing so is at most the cost to manipulate a single attribute.
- (ii) No consumer wishes to deviate from a segment whose hazard ratio is reduced because those consumers are better off than they were before.
- (iii) No consumer wishes to deviate between unaltered segments because the no-arbitrage condition already held between these segments.

Induction Step: Let i be a positive integer. Suppose that for every $\mathbf{b}$ such that $\|\mathbf{b} - \mathbf{a}\|_1 \leq i$, it holds that $h(\mathbf{b}) = \max\{0, H - \frac{\|\mathbf{b} - \mathbf{a}\|_1 c}{K\delta^2}\}$. Then h can be reduced for every segment which satisfies $h(\mathbf{b}) > \max\{0, H - \frac{(i+1)c}{K\delta^2}\}$ and $\|\mathbf{b} - \mathbf{a}\|_1 \geq i+1$ to satisfy $h(\mathbf{b}) = \max\{0, H - \frac{(i+1)c}{K\delta^2}\}$ and $h(\mathbf{a})$ is simultaneously increased. As above, this is a mean-preserving spread of the hazard ratios h , increasing the value of program (7). This establishes that $h(\mathbf{b}) = \max\{0, H - \frac{\|\mathbf{b} - \mathbf{a}\|_1 c}{K\delta^2}\}$ for all $\mathbf{b}$ such that $\|\mathbf{b} - \mathbf{a}\|_1 \leq i+1$. Similar to the base case, this shift does not disturb the no-arbitrage condition. In words, just as in the base case, high types are removed from segments as much as possible and placed into the segment $\mathbf{a}$ where they are fully extracted. $\square$

We refer to the segment $\mathbf{a}$ for which $h(\mathbf{a}) \geq H$ as the “apex”. The next lemma shows that, without loss of generality, in an optimum, the apex can always be taken to be $\mathbf{a} = (1, \dots, 1)$.

LEMMA 2. *Let h be a solution to program (7). If for some segment $\mathbf{a}$, $h(\mathbf{a}) \geq H$, then there is an optimal program for which $h(1, \dots, 1) \geq H$.*

PROOF. Let h be a solution to program (7) with apex $\mathbf{a}$.

For any $j = 1, \dots, K$, if $\mu_j(1) = \frac{1}{2}$, then without loss of generality, $a_j = 1$. If not, then h can be re-defined by $h(b_{-j}, 0) = h(b_{-j}, 1)$ and $h(b_{-j}, 1) = h(b_{-j}, 0)$ for all b . This does not change the mean nor the no-arbitrage conditions because the underlying segments $(b_{-j}, 0)$ and $(b_{-j}, 1)$ have an equal number of low types, because $\mu_j(1) = \mu_j(0) = 1/2$.

Next, we show that for any $j = 1, \dots, K$ if $\mu_j(1) < \frac{1}{2}$ then it must be the case that $a_j = 1$.

Suppose not. Let j be an attribute for which $a_j = 0$ and $\mu_j(1) < \frac{1}{2}$. We construct a hazard ratio $\tilde{h}$ with a higher value to program (7) by swapping high types as follows.

$$\begin{aligned}\tilde{h}(\mathbf{b}_{-j}, 0) &= h(\mathbf{b}_{-j}, 1) \text{ and} \\ \tilde{h}(\mathbf{b}_{-j}, 1) &= h(\mathbf{b}_{-j}, 0)\end{aligned}$$

The total quantity of high types decreases, i.e.

$$m(\mathbf{b}_{-j}, 1)h(\mathbf{b}_{-j}, 1) + m(\mathbf{b}_{-j}, 0)h(\mathbf{b}_{-j}, 0) > m(\mathbf{b}_{-j}, 1)h(\mathbf{b}_{-j}, 0) + m(\mathbf{b}_{-j}, 0)h(\mathbf{b}_{-j}, 1)$$

because $h(\mathbf{b}_{-j}, 0) \geq h(\mathbf{b}_{-j}, 1)$ (since $\|\mathbf{a} - (\mathbf{b}_{-j}, 0)\|_1 + 1 = \|\mathbf{a} - (\mathbf{b}_{-j}, 1)\|_1$) and

$$m(\mathbf{b}_{-j}, 0) = \frac{\mu_j(0)}{\mu_j(1)} m(\mathbf{b}_{-j}, 1) > m(\mathbf{b}_{-j}, 1).$$

So, to allocate all high types in $\tilde{h}$, the remainder are placed at $(\mathbf{a}_{-j}, 1)$. Thus, regarding $\mathbf{a}$, we have that

$$\begin{aligned}\tilde{h}(\mathbf{a}) &= h(\mathbf{a}_{-j}, 1) \text{ and} \\ \tilde{h}(\mathbf{a}_{-j}, 1) &> h(\mathbf{a})\end{aligned}$$

where the last inequality is because all excess high types from all swaps are placed at $(\mathbf{a}_{-j}, 1)$ and $m(\mathbf{a}_{-j}, 1) < m(\mathbf{a})$. This is a mean-preserving spread, and therefore increases the value of program (8). Finally, outside of $\mathbf{a}$ and $(\mathbf{a}_{-j}, 1)$, the no-arbitrage condition is maintained by construction. $\square$

The above lemma can be understood as follows. The segment $(1, \dots, 1)$ is the “lightest” in the sense that it has the fewest number of low types (i.e. it minimizes $\prod \mu_j(a_j)$). Since the advantage of a high hazard rate is that it permits full extraction of high types, it is sensible to place them where they most stick out due to the relative paucity of low types.

The following notation is useful in describing how the solution looks like when there is a large number of attributes; it denotes the number of low type consumers who have i attributes equal to 0:

$$w(i) = \sum_{\mathbf{a}: \|\mathbf{a}\|_1 = i} \prod_{j=1}^K \mu_j(a_j).$$

In the solution given in Proposition 3, all of these consumers will be in segments with the same hazard ratio (i.e., one can think of them as being in the same segment de facto).

The following equation establishes a threshold which tests if it is possible for a segment to have hazard ratio above H :

$$h' = \sum_{i=0}^K w(i) \max \left\{ 0, H - \frac{ic}{\delta^2 K} \right\}. \quad (11)$$

LEMMA 3. *There is a solution to program (7) with $h(\mathbf{a}) \geq H$ for some $\mathbf{a}$ only if $\bar{h} \geq h'$.*

PROOF. By Lemma 2, if $h(\mathbf{a}) \geq H$, then there is a solution with $\mathbf{a} = (1, \dots, 1)$. By Lemma 1, h is established for all other segments, and $h(1, \dots, 1) \geq H$ only if $\bar{h} \geq h'$. $\square$

We will now employ the above lemmas to analyze the value of data as the number of attributes grows large. Before proceeding, we assume that the average of the μ_i is well-defined. This covers several cases, including that the μ_i parameters are constant, periodic, or have a well-defined limit.

ASSUMPTION 3. *The sequence $\mu_1(0), \mu_2(0), \dots$ has a well-defined average. That is, there is a $\mu(0) \in [0, 1]$ such that:*

$$\lim_{K \rightarrow \infty} \frac{1}{K} \sum_{i=1}^K \mu_i(0) = \mu(0).$$

There are two implications of the above assumption:

- (i) Define $\mu(1) = 1 - \mu(0)$. It holds that $\lim_{K \rightarrow \infty} \frac{1}{K} \sum_{i=1}^K \mu_i(1) = \mu(1)$.
- (ii) On average, a low-type consumer will have about $\mu(1)$ fraction of its many attributes equal to one.

The importance of this second item is that it allows us to define a threshold h^∞ which serves as an analogue for h' when the number of attributes becomes sufficiently large. Furthermore, Lemma 3 is strengthened for this analogue—when K is large, the optimal program has a segment with hazard ratio above H if and only if $\bar{h} \geq h^\infty$.

LEMMA 4. *The following equality holds:*

$$\max \left\{ 0, H - \frac{\mu(1)c}{\delta^2} \right\} = \lim_{K \rightarrow \infty} \sum_{i=0}^K w(i) \max \left\{ 0, H - \frac{ic}{\delta^2 K} \right\}.$$

Define h^∞ to be the above quantity.

PROOF. Let $\mathcal{I}(\epsilon, K) = \{i \in \mathbb{N} : \frac{i}{K} \in (\mu(1) - 2\epsilon, \mu(1) + 2\epsilon)\}$. Then, for any $\epsilon > 0$:

$$\lim_{K \rightarrow \infty} \sum_{i \in \mathcal{I}(\epsilon, K)} w(i) = 1.$$

To see why, let S_K be a random variable which takes value i with probability $w(i)$. Chebyshev's inequality states that

$$\Pr \left(\left| \frac{S_K - \sum_{i=1}^K \mu_i(1)}{K} \right| \geq \epsilon \right) \leq \frac{1}{(\epsilon K)^2} \text{Var}(S_K) = \frac{1}{(\epsilon K)^2} \sum_{i=1}^K \mu_i(0) \mu_i(1) \leq \frac{1}{(\epsilon K)^2} \frac{K}{4} = \frac{1}{4\epsilon^2 K}.$$

By Assumption 3, for large enough K , $\left| \frac{\sum_{i=1}^K \mu_i(1)}{K} - \mu(1) \right| \leq \epsilon$, and therefore

$$\Pr \left(\left| \frac{S_K}{K} - \mu(1) \right| \geq 2\epsilon \right) \leq \Pr \left(\left| \frac{S_K - \sum_{i=1}^K \mu_i(1)}{K} \right| \geq \epsilon \right) \leq \frac{1}{4\epsilon^2 K} \rightarrow 0. \quad \square$$

As mentioned above, h^∞ will serve as an analogue for h' when the number of attributes becomes sufficiently large. Lemma 3 showed that h' determines when a solution can have a segment for which the monopolist targets only high types, i.e. $h(a) \geq H$. Likewise, h^∞ determines when a solution will have such a segment in the limit. The next proposition use this fact to show that, unlike in Section 4.1, when the number of attributes grows, it is possible that information has value in the limit. Specifically:

Information has value in the limit *if and only if* the proportion of high-types exceeds this threshold.

The proposition also calculates the gain in profit from first-degree price discrimination. Interestingly, and as seen in Figure 3, the consumer data obtains this value when $h^\infty = 0$.

PROPOSITION 6. *Given Assumptions 2 and 3, as the number of attributes grows large, i.e., when $K \rightarrow \infty$, the value of consumer data converges to:*

$$D_K \rightarrow \bar{m} \begin{cases} \delta^2(H - h^\infty)^2 & = \min \left\{ t_\ell^2, \left(\frac{\mu(1)c}{\delta} \right)^2 \right\} & \text{if } \bar{h} > H; \\ \delta^2[(H - h^\infty)^2 - (H - \bar{h})^2] = \min \left\{ t_\ell^2, \left(\frac{\mu(1)c}{\delta} \right)^2 \right\} - (t_\ell - \delta\bar{h})^2 & \text{if } \bar{h} \in (h^\infty, H]; \\ 0 & \text{if } \bar{h} \leq h^\infty. \end{cases}$$

which is bounded above by the gain in profit from second- to first-degree price discrimination:

$$F = \bar{m} \begin{cases} \delta^2 H^2 & = t_\ell^2 & \text{if } \bar{h} > H; \\ \delta^2 [H^2 - (H - \bar{h})^2] & = t_\ell^2 - (t_\ell - \delta \bar{h})^2 & \text{if } \bar{h} \leq H. \end{cases}$$

The gap in profits between first-degree price discrimination and optimal pricing with consumer data converges to

$$\bar{m} \begin{cases} \delta^2 [H^2 - (H - h^\infty)^2] & \text{if } \bar{h} > h^\infty; \\ \delta^2 [H^2 - (H - \bar{h})^2] & \text{if } \bar{h} \leq h^\infty. \end{cases}$$

In particular, data can do as well as first-degree price discrimination if and only if $h^\infty = 0$ or equivalently $\mu(1)c \geq \delta z$.

PROOF. First, note that h^∞ is a weighted average of terms between H and 0, so the above formulation is well-defined.

If $\bar{h} < h^\infty$, then no segment $\mathbf{a}$ can satisfy $h(\mathbf{a}) > H$, and Proposition 4 shows that $D_K \rightarrow 0$.

If $\bar{h} > h^\infty$, then for evaluating the value of data in the limit, by the above observation we can focus on solutions for which $h(\mathbf{a}) > H$ for some segment $\mathbf{a}$. By Lemmas 1 and 2, we know that there exists a solution which has $h(1, \dots, 1) > H$ and $h(\cdot) < H$ otherwise. Lemma 1 shows that this solution is the same as the solution we get in Proposition 4 for a setting with the average hazard ratio equal to h^∞ except for one segment, namely $(1, \dots, 1)$. Compared to the solution from Proposition 4, this segment contains an additional mass $\bar{m}[\bar{h} - h^\infty]$ of high-type consumers. Note that in this segment, these high-type consumers are sold $q = t_h$ and fully extracted, thus, their contribution to the monopolists profit is $\bar{m}[\bar{h} - h^\infty]t_h^2$.

By the same Proposition 4, the value of data converges to zero when the average hazard ratio is h^∞ . Therefore, the value of data converges to

$$\bar{m}[\bar{h} - h^\infty]t_h^2 + \bar{m} \left(h^\infty t_h^2 + (t_\ell - \delta h^\infty)^2 - \bar{h}t_h^2 - (\max\{0, t_\ell - \delta \bar{h}\})^2 \right),$$

where the second term is the difference between no-information profits for the markets with average hazard ratios $\bar{h}$ and h^∞ . We rearrange terms to get the following expression for the value of data: $\bar{m}[(t_\ell - \delta h^\infty)^2 - (\max\{0, t_\ell - \delta \bar{h}\})^2]$.

To get the gain in profit between first- and second-degree price discrimination, note that the maximum total surplus is $\bar{m}[t_\ell^2 + \bar{h}t_h^2]$. $\square$

To understand the value of data in the proposition above, the following intuition may be helpful. Information is valuable to the seller insofar that he can use it to separate the high-type consumers from the low-type ones so as to save paying information rent to the former. When there is a large overall proportion of high-type consumers, this is done by placing as many of them as possible in the apex segment. The apex is special because it is essentially comprised of only high types (the number of low types converges exponentially quickly to 0 on it) and hence high-types are fully extracted on it. On all other segments, high-type consumers are not fully separated and not fully extracted. Therefore, efficiency losses occur due to information rents which are due to this imperfect separation, and the magnitude of these losses depend upon how many high-type consumers are imperfectly separated.

In summary, at the optimum, there is one apex segment which is essentially populated with high-type consumers and a collection of segments that on average have a proportion

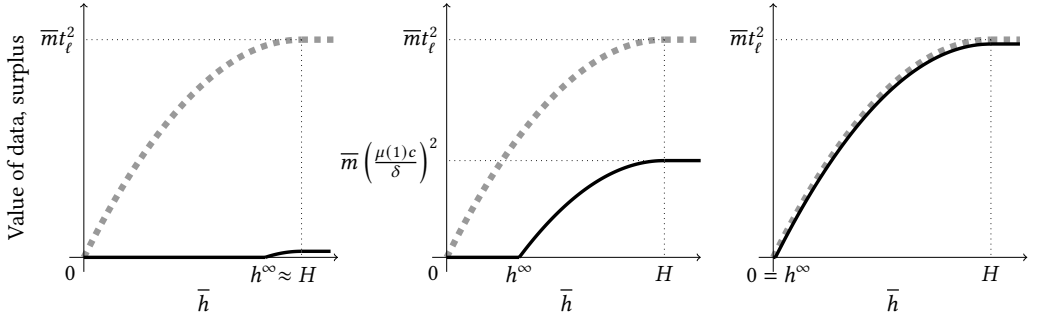


Fig. 3. The value of data for $K \rightarrow \infty$ and near-zero cost c (left), intermediate cost $c \in (0, \frac{\delta t_l}{\mu(1)})$ (center) and high cost $c \geq \frac{\delta t_l}{\mu(1)}$ (right); The black solid line is the value of data (D_∞) and the grey dashed line represents the difference in the monopolist's profit between first- and second-degree price discrimination (F).

of the high types equal to h^∞ . Because the number of high type consumers in the apex is $\bar{m}(\bar{h} - h^\infty)$, the value of data is increasing in $\bar{h}$ as shown on Figure 3.

One can think of h^∞ as an environment-specific, indirect measure data impurity (due to manipulation). Indeed, h^∞ is decreasing in the cost of manipulation c (higher manipulation costs leads to purer data). When h^∞ is low, there is a bigger distance between the apex and the rest of the segments, and therefore a clearer separation between the different types of consumers. In the extreme, when the cost of manipulation is above a certain threshold, and therefore the data is very reliable, the outcome is first-best for the seller, i.e. first-price discrimination. The monopolist is able to perfectly identify consumers' willingness to pay using data, offer the personalized efficient level of quality and charge a price equal to the entire surplus (see Figure 3, right pane).⁸ In particular, the seller can obtain his first-best outcome if $h^\infty = 0 \Leftrightarrow t_l \leq \frac{\mu(1)c}{\delta}$.

Note that the expression for the value of data in the Proposition 6 incorporates our earlier result (see Proposition 4) and illustrates its scope more precisely than it is done by Assumption 1. If there is an insufficient proportion of high types, i.e., if $\bar{h} < h^\infty$, then no apex segment exists and information is worthless. This occurs because no non-negligible proportion of high-type consumers can be fully separated from the low-type ones. The threshold h^∞ is always weakly below H and the distance between H and h^∞ depends on the cost of manipulation, c , the average proportion of attributes with value 1 among the low-type consumers, $\mu(1)$, and the difference in willingness to pay between the high and the low types of the consumers, δ .

Let us examine the first term in the expression for the value of data more closely:

$$\bar{m} \min \left\{ t_l^2, \left(\frac{\mu(1)c}{\delta^2} \right)^2 \right\}.$$

⁸On the other hand, naturally, when data is extremely unreliable ($c \approx 0$), the value of the data is zero (see Figure 3, left pane).

This is the extra profit that the seller is able to secure by selling to the low-type consumers because of the data.⁹ When the cost of data manipulation is large and this expression is equal to the maximal total surplus that can be generated by selling to the low-type consumers: $\overline{m}t_l^2$.

A standard narrative in the literature on optimal pricing is that private information possessed by consumers with high-willingness to pay is an obstacle to efficiency and surplus extraction. One could intuit that the main role of consumer data is to reduce this information rent. In light of this (tempting, but incorrect) intuition, our expression for the value of data appears paradoxical.

Our discovery is that, when the number of attributes is large, data cannot be used to partially reduce information rent. The role of data is to help the seller serve low-type consumers without taking into account the offer that he makes to the large portion of the high-type consumers. Put differently, it is “all-or-nothing” property: the data is useful only if it separates the consumers sufficiently for the seller to create an apex—a sizable homogeneous segment populated by the high-type consumers. This is done to ensure that these consumers do not get in his way when he designs the products and sets prices for the low-type consumers in all other segments.

ACKNOWLEDGMENTS

We thank audiences at 2023 IIOC, 2023 CEPR-JIE Applied IO Conference, 2023 SAET Conference, 2023 EARIE Annual conference, 2023 EEA-ESEM Congress, Christmas Micro Theory Conference in Bonn, 2024 ESSET Gerzensee and seminar participants at the University of Cambridge for their helpful comments.

REFERENCES

- Acquisti, A., Taylor, C., Wagman, L. 2016. The economics of privacy. *Journal of economic Literature*. 54 (2), 442–492.
- Ali, S. N., Lewis, G., Vasserman, S. 2022. Voluntary Disclosure and Personalized Pricing. *The Review of Economic Studies*. 90 (2), 538–571.
- Ali, S. N., Lewis, G., Vasserman, S. 2023. Consumer control and privacy policies. *AEA Papers and Proceedings*. 113, 204–09.
- Ball, I. 2022. Scoring strategic agents.
- Bergemann, D., Bonatti, A. 2015. Selling cookies. *American Economic Journal: Microeconomics*. 7 (3), 259–294.
- Bergemann, D., Bonatti, A., Smolin, A. 2018. The design and price of information. *American economic review*. 108 (1), 1–48.
- Bergemann, D., Brooks, B., Morris, S. 2015. The limits of price discrimination. *American Economic Review*. 105 (3).
- Bhaskar, V., Roketskiy, N. 2021. Consumer privacy and serial monopoly. *The RAND Journal of Economics*. 52 (4), 917–944.
- Bonatti, A., Cisternas, G. 2020. Consumer scores and price discrimination. *The Review of Economic Studies*. 87 (2), 750–791.
- Bonatti, A., Huang, Y., Villas-Boas, J. M. 2023. A theory of the effects of privacy.
- Boros, G., Moll, V. 2004. *Irresistible integrals: symbolics, analysis and experiments in the evaluation of integrals*. Cambridge University Press.
- Caner, M., Eliaz, K. 2021. Non-manipulable machine learning: The incentive compatibility of lasso. *arXiv preprint arXiv:2101.01144*.
- Carroll, G., Egorov, G. 2019. Strategic communication with minimal verification. *Econometrica*. 87 (6), 1867–1892.
- Conitzer, V., Taylor, C. R., Wagman, L. 2012. Hide and seek: Costly consumer privacy in a market with repeat purchases. *Marketing Science*. 31 (2), 277–292.

⁹The second term—i.e., $\overline{m}\delta^2(H - \bar{h})$ —represents the difference in the no-data profit that the monopolist is able to secure by catering to both high- and low-type consumers, rather than just the high-type ones.

- Cunningham, T., Moreno De Barreda, I. 2022. Effective signal-jamming.
- Dana, J., Larsen, B., Moshary, S. 2023. Fake ids and arbitrage: Price discrimination with misreporting costs.
- Deneckere, R., Severinov, S. 2022. Signalling, screening and costly misrepresentation. *Canadian Journal of Economics*. 55 (3), 1334–1370.
- Eilat, R., Eliaz, K., Mu, X. 2021. Bayesian privacy. *Theoretical Economics*. 16 (4), 1557–1603.
- Eliaz, K., Spiegler, R. 2022. On incentive-compatible estimators. *Games and Economic Behavior*. 132, 204–220.
- Elliott, M., Galeotti, A., Koh, A., Li, W. 2021. Market segmentation through information.
- Frankel, A., Kartik, N. 2019. Muddled information. *Journal of Political Economy*. 127 (4), 1739–1776.
- Frankel, A., Kartik, N. 2022. Improving information from manipulable data. *Journal of the European Economic Association*. 20 (1), 79–115.
- Frick, M., Iijima, R., Ishii, Y. 2024. Multidimensional screening with rich consumer data.
- Hardt, M., Megiddo, N., Papadimitriou, C., Wootters, M. 2016. Strategic classification. In *Proceedings of the 2016 ACM conference on innovations in theoretical computer science*, 111–122.
- Hidir, S., Vellodi, N. 2021. Privacy, personalization, and price discrimination. *Journal of the European Economic Association*. 19 (2), 1342–1363.
- Hu, L., Immorlica, N., Vaughan, J. W. 2019. The disparate effects of strategic manipulation. In *Proceedings of the Conference on Fairness, Accountability, and Transparency*, 259–268.
- Ichihashi, S. 2020. Online privacy and information disclosure by consumers. *American Economic Review*. 110 (2), 569–595.
- Kartik, N. 2009. Strategic communication with lying costs. *The Review of Economic Studies*. 76 (4), 1359–1395.
- Liang, A., Madsen, E. forthcoming. Data and incentives. *Theoretical Economics*.
- Milli, S., Miller, J., Dragan, A. D., Hardt, M. 2019. The social cost of strategic classification. In *Proceedings of the Conference on Fairness, Accountability, and Transparency*, 230–239.
- Mussa, M., Rosen, S. 1978. Monopoly and product quality. *Journal of Economic theory*. 18 (2), 301–317.
- Perez-Richet, E., Skreta, V. 2022. Test design under falsification. *Econometrica*. 90 (3), 1109–1142.
- Perez-Richet, E., Skreta, V. 2023. Fraud-proof non-market allocation mechanisms.
- Segura-Rodriguez, C. 2021. Selling data.
- Severinov, S., Tam, Y. C. 2018. Screening under fixed cost of misrepresentation.
- Tan, T. Y. 2023. Price discrimination with manipulable observables.
- Taylor, C. R. 2004. Consumer privacy and the market for customer information. *RAND Journal of Economics*. 631–650.
- Yang, K. H. 2022. Selling consumer data for profit: Optimal market-segmentation design and its consequences. *American Economic Review*. 112 (4), 1364–1393.