
INSIDE THE ENGINE ROOM OF DIGITAL PLATFORMS: REVIEWS, RATINGS, AND RECOMMENDATIONS¹

Paul BELLEFLAMME

Martin PEITZ

Abstract

The rise and success of digital platforms (such as Airbnb, Amazon, Booking, Expedia, Ebay, and Uber) rely, to a large extent, on their ability to address two major issues. First, to effectively facilitate transactions, platforms need to resolve the problem of trust in the implicit or explicit promises made by the counterparties; they post reviews and ratings to pursue this objective. Second, as platforms operate in marketplaces where information is abundant, they may guide their users towards the transactions that these users may have an interest in; recommender systems are meant to play this role. In this article, we elaborate on review, rating, and recommender systems. In particular, we examine how these systems generate network effects on platforms.

Keywords: Platforms, network effects, ratings, recommender systems, digital economics.

JEL: L80.

¹ We thank Juanjo Ganuza, Gerard Llobet, and Markus Reisinger for helpful comments. Martin Peitz gratefully acknowledges financial support from Deutsche Forschungsgemeinschaft (DFG) through CRC TR 224.

I. INTRODUCTION

Platforms can be defined as undertakings whose core mission is to enable and to generate value from interactions between users. Although platforms can operate off-line, Internet and digital technologies greatly contribute to reducing transaction costs, which explains why digital platforms are so prevalent nowadays. Digital platforms typically provide a number of services that generate so-called “platform-specific network effects,” insofar as the attractiveness of a particular platform increases with the volume of interactions that the platform manages. Roughly speaking, the platform becomes more attractive the more it is used, and, as a result, each user cares about the participation of other users.²

The participation of other users may matter for a few reasons. First, their active evaluation of products and services, or the information contained in their actions, provides guidance for a user’s action; second, the information contained in the users’ actions enables the platform to provide better services or add specific offerings, both of which potentially benefit all users. In this article,³ we focus on the former reason and analyze platforms’ deployment of review, rating, and recommender systems. These non-price strategies allow platforms to generate within-group and/or cross-group external effects, that are (as we will argue below) platform-specific: the disclosure, aggregation and interpretation of information provided by the participants steer trade on the platform, thereby affecting the overall attractiveness of participating on the platform.

How are rating and recommender systems instrumental in producing network effects? Consider, for instance, the case of Amazon, which publishes product reviews and average ratings. Arguably, the more consumers that are active on Amazon, the more informative are the reviews and ratings, thus allowing consumers to make a better-informed decision. Amazon also provides recommendations by matching product descriptions with consumers’ interests. Similarly, the more consumers that are active on the platform and the larger the volume of transactions they generate, the better the data that Amazon has about consumer characteristics and, so, the better the matches it can suggest; the quality of recommendations increases thus with the number of consumers, which in many cases will lead to a higher expected net consumer benefit. These mechanisms point to positive within-group external effects.

On two-sided platforms, positive cross-group external effects might arise. For instance, a high-quality seller thinking of participating on Ebay, Amazon

² For a justification of this broad notion of what constitutes a platform (*i.e.*, a managed marketplace featuring network effects), see, for instance, Belleflamme and Peitz (2018b).

³ We use material from Chapters 2 and 5 of Belleflamme and Peitz (2018a).

Marketplace or some other B2C platform cares about the ease with which it can build its reputation. The more buyers active on the platform, the more precise the information about the seller type at a given point in time (assuming truthful consumer ratings). Thus, there is a positive cross-group external effect from buyers to high-quality sellers. Similarly, the more buyers on a platform, the better the matching between buyers and sellers (in terms of horizontal characteristics). This, in particular, reduces the expected number of products returned to the sellers. Thus, thanks to the recommender system, there is a positive cross-group external effect from buyers to sellers. This effect is strengthened by more detailed data on each consumer, as this improves the expected match quality.

Ratings are intended to help consumers make choices based on the quality or value-for-money dimension. Recommendations can also serve this purpose; they also have the potential to address buyer heterogeneity if they are personalized. This does not mean that some degree of personalization is impossible in the context of a rating system. In fact, several platforms offer the option of personalization; by, for instance, showing ratings and reviews only of buyers with certain profiles. Such rating selection can provide better guidance because what is good for one group of buyers is not necessarily good for others. For example, a business traveler may have different needs and preferences than a family on vacation and, thus, may prefer to see only reviews and ratings by fellow business travelers.

In the rest of this article, we analyze the economics behind the ratings, reviews and recommendations that have become mainstream on digital platforms. We start in Section [second](#), with rating and review systems. These systems provide platform users with information about either products or their counterparties to a transaction. Of crucial importance is, of course, the informativeness of these systems, which depends not only on the users' actions but also on the specific design chosen by the platforms. We then turn, in Section [third](#), to recommender systems, which aim to reduce users' search cost by pointing them towards transactions that may better match their tastes. Besides the ability of such systems to generate network effects, we also discuss their effects on the distribution of sales between 'mass-market' and 'niche' products, as well as the incentives that platforms may have to distort their informativeness. We conclude in Section [fourth](#).

II. RATINGS AND REVIEWS

Ratings and reviews are prevalent on digital platforms. Platforms acting as vertically integrated retailers (such as Amazon.com) generally ask buyers to rate products or services and often give buyers the chance to write reviews. In such a

case, we speak of product ratings and *product reviews*. For platforms that host buyers and sellers (such as Amazon Marketplace), users on either side are often asked to rate and comment on the counterparty to the transaction. These we call seller (or buyer) ratings and reviews.

1. Asymmetric information and network effects

Before analyzing the economics of rating and review systems, we consider their significance for digital platforms. Unquestionably, the main function of ratings and reviews is to respond to asymmetric information problems. At the same time, they are also an important source of network effects, which makes them instrumental in platforms' efforts to gain market shares. We describe these two aspects in turn.

1.1. Asymmetric information

Asymmetric information problems are prominent on platforms that facilitate the trade of experience goods, as buyers typically have less information than sellers about the quality of the goods or services offered for sale. In this section, we focus on those asymmetric information problems that arise with experience goods.⁴

A traditional instrument to address asymmetric information problems is the use of *certification and warranties*. When a seller wants to transact with a buyer, third parties may provide certification, and platforms are a natural candidate for such certification services. Certification is an *ex ante* solution to asymmetric information problems, as it may ensure a minimum quality provided on the platform; lower-quality sellers are not admitted or worse-performing sellers are expelled from the platform. Certification can be mandatory or voluntary. For instance, Uber checks the records of its drivers to make sure that they are eligible to drive; such certification is mandatory. Airbnb offers the sellers of accommodation services the option to certify the authenticity of photos of the announced property, thus reducing the risk of unpleasant surprises for the buyer; such certification is voluntary. As for warranties, they may, in principle, be provided by sellers themselves, but platforms are often in a better position to provide them, since they interact more frequently and directly with buyers.

⁴ We argue in [Section three](#) that asymmetric information problems may also apply to search goods. In this case, even if buyers can ascertain quality before purchase, they may lack information prior to investing time and effort to obtain relevant product information. Here, platforms can use ratings and reviews (on top of other instruments) to lower buyers' search costs and to improve the match between buyers and products/sellers.

Asymmetric information problems can also be addressed ex post through *insurance and guarantees*. For instance, Airbnb insures sellers against vandalism by buyers. Another example is Ebay's guarantee to buyers (introduced in 2010) to compensate them if the seller does not deliver as advertised (see Hui *et al.*, 2016).

Rating and review systems complement these classic instruments and tend to become relatively more effective than them, the larger the number of transactions that the platforms facilitate. Indeed, the ability of rating and review systems to tackle information problems faced by buyers (and possibly sellers) increases with the volume, variety, and velocity of the data that platforms can collect about their users and the transactions they conduct.⁵

1.2. Network effects

As just argued, ratings and reviews can be an important source of network effects: the more users that are active on a platform—and, thus, the more ratings and reviews that are available—the better-informed other users are prior to making their purchase decisions. In the following sections, we will clearly identify the various forms that these network effects can take. What we want to stress here is that, although users often have access to ratings and reviews whether or not they purchase on a particular platform, network effects tend to be 'platform-specific' for a number of reasons.

First, some users may not consider purchasing on a platform different from the one on which they obtain information. In this case, even if a featured product is available on multiple platforms, it matters on which platform better information is available. For instance, in the early 2000s, buyers in the U.S. may have accessed ratings and reviews available on books at Amazon and then purchased the book from Barnes & Noble. However, as we discuss below, the positive sales effect of high ratings is more pronounced on the same platform than across platforms. This suggests that a substantial fraction of buyers only took note of reviews and ratings only on the platform on which they terminated their purchase.

Second, when buyers rate sellers on a two-sided platform, a seller may (at least partially) condition its behavior on the distribution channel picked by the user. In this case, the seller's reputation is actually conditional on the transaction on a platform. For example, a hotel may be more accommodating to the wishes and requests of a guest who booked on a particular platform. To give another example, a seller may exert particular effort to speedy delivery of a product ordered through a particular platform.

⁵ The veracity of the data is also crucial, as we discuss in point 4 of Section II.

Third, the identity of a seller may be platform-specific, or it may be costly for the user to identify the same seller across platforms. For instance, it may be difficult to verify that the seller name on Ebay or Amazon Marketplace corresponds to the seller name on some other distribution channel. If this is the case, network effects are, by construction, platform-specific. For all these reasons, we can safely record the following finding.

Finding 1. Because they generate platform-specific network effects, rating and review systems fuel self-reinforcing mechanisms that, other things being equal, make successful platforms even more successful, at the expense of their smaller rivals.

We now turn to an in-depth analysis of rating and review systems on products and services (point 2 of Section II), and on transaction counterparties (point 3 Section II). We then address the fundamental issue of the informativeness of these systems (point 4 Section II).

2. Product rating and review systems

Many online retailers have established rating and review systems (or ‘rating systems’ for short) that allow buyers to rate and comment on particular products. Absent such a rating system, we would not classify an online retailer as a platform, since, given prices, a buyer’s purchase intention would not be affected by other buyers’ purchases. However, the presence of a rating system renders the retailer a platform, as it is a source of network effects, and its design affects the strength of network effects.

Finding 2. Product rating systems have the potential to solve asymmetric information problems. In an e-commerce context in which buyers rate products, as more buyers on a platform make the average product rating more informative, a platform with a product rating system features positive network effects among buyers.

To illustrate this point, we consider a firm that carries products sourced at marginal cost c and sold at price p . Neither the firm nor the buyers know the quality of any product prior to consumption. What is known is that quality q may be either high ($q=H$) or low ($q=L$) with probability $1/2$, and that this probability is drawn independently across products. Buyer valuations for high and low quality (respectively, v_H and v_L) satisfy $v_H > c > v_L$ and $(v_H + v_L)/2 > c$. The first set of inequalities tells us that if information were complete, only high-quality products would be traded (as buyers value the low quality below its marginal cost). The second inequality tells us that when buyers are

uninformed, trade will nevertheless take place, as the average valuation of a product is above the marginal cost.

Suppose that there are k buyers, who arrive in random order at each product. Each buyer is inclined to leave a review (if the firm provides a rating system) with some probability ρ , which is independent of the actual quality of a product. Furthermore, suppose that buyers perfectly observe product quality after purchase and report this quality truthfully if they write a review.

Absent a product rating system, a monopoly firm sets its price equal to the average valuation, $p = (v_H + v_L)/2$, and all buyers make a purchase. With a product rating system and under the assumption of a uniform price, the firm has to set the price such that buyers buy the product even when no review is available. This price is the same as without a rating system, as a buyer who does not observe any review is willing to pay up to the average valuation—i.e., $(v_H + v_L)/2$.

At such a price, a buyer buys the product as long as no review of low quality has been posted (i.e., if either no review is available, or if only positive reviews are available). If the product is of high quality, regardless of the order in which buyers appear, there will be no negative review posted. If the product is of low quality, a buyer in position k encounters with probability $(1-\rho)^{k-1}$ that none of the previous $k-1$ buyers left a review. Thus, the overall probability that a buyer in a market with a total of n_b buyers does not see a negative review is $P_H + P_L$, where $P_H = 1/2$ is the probability that the product is of high quality (and it does not matter then whether or not buyers wrote a review), and $P_L = \sum_{k=0}^{n_b-1} (1-\rho)^k / (2n_b) = [1 - (1-\rho)^{n_b}] / (2\rho n_b)$ is the cumulative probability that none of the previous buyers left a review and the product is of low quality. Importantly, P_L decreases as the number of buyers, n_b , increases (it converges to 0 as n_b turns to infinity). The expected surplus of a buyer is then equal to $U^e = P_H(v_H - p) + P_L(v_L - p)$. As $p = (v_H + v_L)/2 > v_L$, it follows that $U^e = (P_H - P_L)(v_H - v_L)/2$, which is increasing in n_b . Thus, a platform with a product rating system is more informative the larger the number of buyers and, therefore, exhibits positive network effects.⁶

In the above example, the rating system generates positive network effects among buyers; such effects are generally called ‘within-group’ or ‘one-sided’ network effects. Does this imply that retailers with a rating system do not feature two-sidedness? In general, one- or two-sidedness is often a matter of

⁶ In the example, a monopoly firm makes a lower profit with a rating system because it sells at the same price to fewer buyers. However, if buyer participation necessitates an up-front fixed cost for buyers, there is a hold-up problem absent a rating system. In this case, establishing a rating system limits the hold-up problem and, in equilibrium, may lead to higher profits for a firm with a rating system, since the market breaks down absent a rating system. In this case, a monopoly firm has the incentive to establish a rating system.

the concrete circumstances. This is also the case with rating systems, as we now show in the following three examples.

In the first example, we consider a stylized two-period setting in which some users simultaneously make purchase decisions in period 1, and other users simultaneously make purchase decisions in period 2. Suppose that a fraction of the former group posts a rating. Thus, period-2 buyers can make better-informed decisions, as the number of period-1 users increases. This means that due to the ratings system, there are positive cross-group external effects from period-1 users to period-2 users.

In the second example, we consider another stylized setting that features two types of buyers. For the first type, products are experience goods (quality is observed with some noise after purchase) and for the second type, they are credence goods (quality is not observed, even after consumption). Suppose that only users who learn the product rate the product (truthfully) and that those who do not learn the quality do not leave a rating. If users buy different products over time and base their decisions on average ratings, they benefit from a retailer attracting more type-1 buyers, as additional rankings allow for better-informed choices. Thus, there exist positive within-group external effects for type-1 buyers and positive cross-group external effects from type-1 to type-2 buyers. To the extent that type-1 buyers can draw on their own previous experience, informative ratings are less essential than for type-2 buyers, and, thus, the cross-group external effects generated by type-1 buyers are stronger than their within-group external effects.

Turning to the third example, consider now that, depending on the group a buyer belongs to, she leaves reviews with different probabilities; let λ_j denote the review probability in group j . If n_j^i buyers of group j participate on platform i , the expected number of reviews on platform i is $m^i = \lambda_1 n_1^i + \lambda_2 n_2^i$. More reviews make a platform more attractive to buyers. This benefit can be captured by an increasing and concave function $f(m^i)$. In this setting, there are positive within-group external effects for each group of buyers. In addition, there are positive cross-group external effects between the two groups of different strength (if $\lambda_1 \neq \lambda_2$).

As argued above, rating systems help buyers make more-informed choices. With a rating system in place, the empirical prediction is that a more-highly-rated product should see its sales increase compared to a less-highly-rated product. Chevalier and Mayzlin (2006) analyze the effect of book reviews on the sales patterns of the two leading online booksellers in the USA (at that point in time), Amazon and Barnes & Noble.⁷ Both offer buyers the opportunity to post book

⁷ Our exposition is almost identical to that in Belleflamme and Peitz (2015, Chapter 15).

reviews on their site. The central question of the study is whether an additional negative report on Amazon leads to a decline in sales at Amazon relative to the sales at Barnes & Noble. If the answer is 'yes,' this means that book reviews carry relevant information that affect sales. To answer this question, Chevalier and Mayzlin use the 'differences-in-differences' approach—that is, they take differences between the relative sales of a book at the two retailers to control for possible effects of unobserved book characteristics on book sales and reviews. Data were publicly available: they cover a random selection of book titles with certain characteristics in three short periods—two-day periods in May and August 2003 and May 2004.

Chevalier and Mayzlin regress the natural logarithm of the sales rank of book i at retailer j (which serves as a proxy for sales) on a number of variables including fixed effects, prices at Amazon and Barnes & Nobles and the share of positive (5-star) and negative (1-star) reviews. Chevalier and Mayzlin show that an additional positive review for a particular book at one retailer leads to an increase in the sales of this book at that retailer relative to the other. There is also some evidence that an additional negative review is more powerful in decreasing book sales than an additional positive review is in increasing sales (measured by the sales rank). The fact that the length of reviews also matters suggests that buyers not only use summary statistics but actually take a look at the reviews; this also suggests that they take the content of the review explicitly into account (perhaps to evaluate how much to trust a particular review or because there is uncertainty with respect to the fit of the match, which is buyer-specific).

Vana and Lambrecht (2018) use product review data from an UK online retailer. They identify the effect of the content of individual reviews, since the position at which reviews are placed is exogenous in their setting (placement by the date of being posted). When a new review appears, all existing reviews are shifted downward by one position. This shift occurs regardless of the content and rating of any review. As the authors show, the rating of the first displayed reviews have a strong effect of purchase likelihood. In particular, if these reviews come with a high rating (four or five stars out of five) the estimated purchase probability increases significantly.

3. Seller rating systems

So far, we have considered rating systems by a retailer that interacts with consumers. We now turn to rating systems of two-sided platforms: B2C and C2C platforms bring sellers and buyers together. Here, rating systems are a

solution to the general trust problems encountered by buyers. Should they trust the quality claims that sellers make about their products on offer? Should they trust the service promises? Possibly, these trust problems also exist the other way round. In a bilateral relationship, such trust problems can be solved through repeated interaction. When buyers are likely to provide reviews and/or ratings and these are informative, the trust problem can (at least, partially) also be solved in anonymous markets. Here, the rating and review system (or 'reputation system') serves as a substitute for personal experience: an individual buyer can draw on the collective experience of other buyers.

If you have ever booked a room in a hotel and learned upon late arrival that all the rooms were occupied, you may appreciate booking platforms that provide feedback from other buyers on the reliability of the information provided by the hotel. Perhaps more importantly, hotels have to worry about their reputation if they do not treat their guests well. For this reason, reputation systems are an important driver of the success of platforms as enablers to transaction—they may generate trust for at least one of the parties involved and resolve asymmetric information problems.

A rating system may be one-sided or two-sided. For instance, Amazon Marketplace has a one-sided rating system according to which buyers rate sellers. The initial Ebay system was two-sided, and so are the systems of Airbnb and Uber. Here, each transaction partner can rate, and leave a review about, the partner on the other side.

Rating systems can tackle adverse selection and moral hazard problems. For instance, accommodations on Airbnb that suffer from some unexpected problems can be singled out by reviews and ratings. To the extent that these unexpected problems are inherent to the property, this reveals the quality of the accommodation and resolves adverse selection problems. Unexpected problems can also arise if the seller does not exert effort; here, ratings and reviews can help to solve the associated moral hazard problem.

If reviews and ratings are noisy, a platform with few transactions per seller does not provide very reliable information. Given the number of sellers, the more buyers that are active on the platform, the more precise is the information on any seller since the average valuation tends to converge to the true valuation. This suggests that there exist positive network effects on the buyer side—we will discuss and qualify this finding below (as the informativeness of the ratings depends on their truthfulness).

Finding 3. Seller rating systems have the potential to solve asymmetric information problems. In a buyer-seller context in which buyers rate sellers, as more buyers on a platform make the rating system more informative, a platform with a rating system features positive within-group external effects on the buyer side.

For a given number of buyers, the rating system's informativeness tends to increase in the response rate of buyers. Here, the rating system may be designed to encourage buyers to leave a review or rating. Response rates may depend positively on the ease of use of the platform, and on the community feeling that it creates. The platform may also provide non-monetary or monetary incentives to leave reviews. As an example of the former, Tripadvisor awards a number of badges depending on review activity. Regarding the latter, Fradkin, Grewal, and Holtz (2017) ran a field experiment on Airbnb in which they provided monetary incentives for leaving reviews and showed that this can be effective. A seller reputation system may also suffer from low response rates by buyers who are afraid to rate a seller after a bad experience—more on this below when we discuss the informativeness of ratings and reviews.

A number of empirical works have shown that more reputable sellers are more successful—that is, reputation pays. Reputable sellers may be able to ask for a premium and/or they may enjoy higher transaction volumes—in particular, they may also be able to successfully sell products that buyers *a priori* deem to be risky to buy.

Resnick *et al.* (2006) run a controlled field experiment to investigate the price premium of reputation: they sell a number of identical products (collectible postcards); some of them are randomly assigned to an established seller with a good record and some to a seller with little track record. They estimate an 8% price premium for a seller with 2,000 positive and one negative ratings, compared to a seller with ten positive and zero negative ones. Cabral and Hortacsu (2010) collect a large data set of seller histories on Ebay. Unfortunately, they do not observe the number of a seller's past completed transactions and assume that the frequency of a seller's feedback is a good proxy for the frequency of actual transactions.⁸ According to their estimates, a seller's weekly sales growth rate drops from a positive rate of 5% to a negative rate of 8% upon receiving his first negative rating.⁹

⁸ This assumption may seem innocuous. However, as discussed below, different seller types are likely to have different rates by which buyers give reviews and ratings.

⁹ A potential drawback is that they do not include price effects, but they may actually be small. Other early empirical work on auction sites includes McDonald and Slawson (2002), Melnik and Alm (2002), Livingston (2005), and Jin and Kato (2006). For a summary of this and other work, see Bajari and Hortacsu (2004) and Tadelis (2016).

Some platforms started off without a rating system. For instance, the Chinese auction site Eachnet operated initially (1999-2001) without such a system. A certain degree of bilateral trust between seller and buyer was established through communication between the two parties, which eventually led to a physical meeting. Thus, the buyer could inspect the product before paying, and the seller could make sure that the seller made the payment. While this does not resolve all asymmetric information problems *ex ante*, some of the most unpleasant surprises for buyer and seller could be avoided even without a rating and review system. In 2001, Eachnet introduced a rating and review system. Cai *et al.* (2014) empirically investigate how a seller's "reputation" affects outcome, depending on whether a rating and review system is in place. A seller's reputation is approximated by the cumulative success rate of its listings. A seller's listing is successful if it led to at least one transaction. One may expect that if a buyer and a seller successfully complete a transaction, they may be more likely to interact again in the future. This may hold, in particular, for "reputed" sellers (*i.e.*, those with a high cumulative success rate). Indeed Cai *et al.* (2014) find a positive correlation between sellers' cumulative success rate and the fraction of repeat buyers. The important finding here is that this correlation weakens after the introduction of the rating system. This suggests that the rating system makes the asymmetric information problem faced by occasional buyers less severe and, thus, serves as a partial substitute to reputation within a bilateral relationship.

The introduction or redesign of a rating system may have an impact on the sellers' decision of whether to join a platform (and on the scale of its activities). For instance, if the rating system leads to better-informed buyers, low-quality sellers may abstain from participating. It might also affect the behavior of sellers beyond whether (and with what intensity) to participate. For instance, if a misrepresentation of product quality is punished through a negative rating that is easily observable to potential buyers, a seller may be more careful in drafting his announcements. In short, a rating system may affect participation (and, thus, affect the amount of adverse selection) and behavior, given participation (and, thus, the degree to which the moral hazard problem plays out). Klein, Lambertz, and Stahl (2016) investigate the effects of Ebay's redesign of its reputation system in May 2008, when Ebay introduced one-sided feedback that is not subject to retaliation and, thus, can be seen as more accurately reflecting a buyer's experience (below, see more on retaliation). Since, prior to that date, in May 2007, Ebay introduced an anonymous details seller rating (DSR) on top of its rating system, Klein, Lambertz, and Stahl could use this DSR before and after the change to a one-sided rating system as a measure of buyer satisfaction. They found a significant increase in buyer satisfaction with the introduction of the one-sided rating system, but did not observe a significant change in the sellers' exit rate. This can be seen as evidence that, in this instance, the redesign of the rating system was successful in reducing moral

hazard but did not significantly affect the composition of sellers. In the case of Ebay, this seems conceivable, as a low-quality product may find its buyer even if quality is revealed since there may be a market for such low-quality products. The effect of the redesign of the rating system would then encourage truthful announcements by sellers but would not remove their incentive to participate.

Finding 4. In the case of hidden-information problems, sellers are affected differentially by seller rating systems: high-quality sellers enjoy a positive cross-group external effect from more buyers leaving ratings, while low-quality sellers suffer a negative cross-group external effect from more buyers leaving ratings. In the case of hidden-action problems, all sellers may benefit, as buyers understand that the system disciplines sellers.

4. The informativeness of ratings and reviews

Rankings and reviews can be relevant for buyers only if they contain relevant information. Clearly, if they are informative about the (price-adjusted) quality of a product, buyers must, at least to some degree, have a common perception of the (price-adjusted) quality, and buyers must be able and willing to report their experiences with the product.

We identify three sets of reasons why the informativeness of ratings and reviews may be limited due to decisions by buyers and sellers¹⁰²: (i) noisy ratings and reviews; (ii) strategically distorted ratings and reviews; and (iii) asymmetric herding behavior. We discuss these, in turn, before examining how platforms can act to make rating systems more—or less—informative.

4.1 Noise

We describe here four reasons that buyers may leave noisy ratings and reviews: bad understanding, idiosyncratic tastes, uncontrollable shocks, and price variations.

— *Bad understanding*

Buyers may leave noisy ratings and reviews simply because they fail to understand what they are asked. While this is often easily identified after reading a review, buyers who rely on summary statistics may not be able to identify that ratings are based on irrelevant experiences. For instance, this applies to

¹⁰² For other overviews, see Aral (2014) and Tadelis (2016).

product ratings on Amazon. Here, some reviewers do not base their rating on the quality and characteristics of the product they bought, but on such factors as Amazon's delivery service, which can be considered orthogonal to the product sold by Amazon. For example, the 2010 edition of our textbook *Industrial Organization: Markets and Strategies*, received a 5-star rating by one reviewer on Amazon.com with the following review: "It's my first time to buy used books. And it has definitely met my expectation. Well kept just few marks. Like it very much."¹¹ While we are happy that the reviewer gave a 5-star rating, we are not so sure if this actually reflects his or her quality assessment of the book rather than the physical appearance of the used copy.

— *Idiosyncratic tastes*

Ratings may also be noisy for potential buyers because of idiosyncratic tastes. While rating systems are supposed to capture the quality of a product or seller, reviewers may comment on horizontal characteristics or on vertical characteristics for which they have heterogeneous willingness to pay. In other words, ratings that aggregate tastes of other buyers may not strongly correlate with one's own taste. For instance, a reviewer may give a negative product rating because she does not like the color of the product, but other potential buyers may not share this negative feeling.

— *Uncontrollable shocks*

Relatedly, there may be shocks that are not under the seller's control. If a reviewer leaves a negative seller rating because of late delivery, this may not have been under the seller's control if, say, the transport company did not deliver in time. One would expect that such shocks to product and service satisfaction wash out if there is a large number of reviewers. Thus, the informativeness increases with the number of fellow users, a source of the network effects mentioned above.

— *Price variations*

Product and seller reviews are often likely to be based on how satisfied a buyer is when taking into account how much she paid. However, products may be sold at different prices over time and space. Thus, what looks like a rather bad deal at a high price may be a good deal at a low price. Therefore, with price variation (over time and space), the informativeness of ratings suffers.

¹¹ As Tadelis (2016, p. 328) notes, confusion is likely with multiple review targets: "Multiple review targets may create an inference problem that confuses between the seller's quality of executing the sale and the quality of the product."

4.2. Strategic distortions by buyers or sellers

Buyers or sellers may take actions that systematically distort seller or product ratings. Clearly, since sellers benefit from a positive reputation, they may pay others to leave positive reviews and ratings about their offers; they may also pay others to leave negative reviews about the offers of close competitors. First, we examine such ‘fake reviews,’ and then we consider the specific problems that may emerge from ‘two-sided rating systems,’ in which both counterparties to a transaction are invited to rate one another.

— Fake reviews

The unsuspecting reader may think that fake reviews are an issue cooked up by economists who believe in incentive theory. However, there is evidence that fake reviews are widespread and that markets for such fake reviews have been created (see, e.g., Xu, Chen and Winston, 2015).¹²

Generating such fake reviews is costly. Costs and benefits from fake reviews depend on the particular site. As Ott, Cardie, and Hancock (2012) argue in case of hotels, the costs of a fake review are high if a user is required to purchase a product prior to reviewing it. For instance, hotel booking platforms Booking and Expedia require an actual purchase, whereas Tripadvisor (which, as a referral website, does not monitor transactions) allows anyone who claims to have made a booking to post reviews about a hotel. Thus, fake reviews are more costly on Booking and Expedia than on Tripadvisor. The expected benefit depends on the attention that a particular review attracts. Everything else being given, the benefit on a website with many visitors is greater, while on a website with many other reviews the expected benefit, it is smaller. Hence, in an environment in which the ratio of reviews to traffic is the same across websites, it is not clear on which website the expected benefit is the largest. We note that posting a fake review on a website with a quickly growing visitor base and a small stock of reviews is particularly attractive. This suggests that newcomer platforms must think hard about how to design their rating system right from the start.

Providing evidence on the extent of fake reviews is hard, since actual fakes are difficult to spot. Mayzlin, Dover, and Chevalier (2014) exploit different policies by hotel information and booking sites about who can leave feedback: Expedia requires the reviewer to have booked a hotel on its site, while Tripadvisor does not (as it only referred to booking sites). Thus, we would expect to see

¹² Since fake reviews are costly to generate, a more benign view of the use of positive, paid-for reviews and ratings is that they can be seen as a seller’s costly advertising and may be used as a signal of high quality—the seminal paper on advertising as a quality signal is Milgrom and Roberts (1986). For an empirical analysis of such behavior on the platform Taobao, see Li, Tadelis, and Zhou (2016).

more fake reviews on Tripadvisor. Consider a geographic area in which hotels compete for business travelers. It is in the strategic interest of any hotel in this area to improve its ranking relative to that of competing hotels in the same area. A hotel can achieve this by inflating its own rating with fake positive reviews and by deflating the rating of hotels in its vicinity with fake negative reviews.

Mayzlin, Dover, and Chevalier argue that independent hotels are more likely to sponsor fake reviews, as their cost from being detected is less severe than if such a review was sponsored by a hotel belonging to a chain. Thus, the prediction is that hotels in the vicinity of such independent hotels have more negative reviews on Tripadvisor relative to Expedia, and independent hotels have more positive reviews on Tripadvisor relative to Expedia. These predictions are confirmed in their dataset. And fake reviews are not unique to hotels; for instance, Luca and Zervas (2016) analyze fake restaurant reviews on Yelp.

— *Two-sided rating systems*

Problems of systematic misrepresentation and, possibly, underreporting of negative experiences may arise with two-sided rating systems in which both buyer and seller leave feedback. Such two-sided ratings appear to be desirable if both parties have private information and/or choose private actions. In its early days, Ebay used a two-sided system, arguably because sellers would like to know which buyers to trust. In particular, a buyer may place the highest bid but then refuse to make the promised payment. With developments in electronic payments, this risk for the seller could be eliminated. This has removed the main reason to use two-sided ratings on Ebay. Other platforms continue to employ two-sided rating systems. This applies, in particular, to platforms in the sharing economy because here, not only the payment, but also the way a buyer uses a product matters to the seller. For instance, somebody renting out an apartment on Airbnb may worry about whether the renter will create a mess or damage some furniture.

Although two-sided rating systems do not necessarily distort ratings, the past system on Ebay did. The Ebay rating system had the design feature that buyers and sellers had a time window during which they could leave a feedback. When one party left a feedback, it was disclosed to the other party. This opened up the possibility of retaliation for a negative rating. Bolton, Greiner, and Ockenfels (2013) analyze rating behavior on the old Ebay and document that the two ratings in buyer-seller pairs are highly positively correlated. They also document that sellers typically wait for the buyer to leave a rating and respond promptly. This supports the view that sellers use their feedback as an implicit threat to leave a negative rating if they receive a negative one. This makes it more painful for buyers to give negative ratings and, effectively, distorts the

distribution of ratings received by sellers.¹³ Indeed, as Nosko and Tadelis (2015) report, using internal Ebay data, a buyer is three times more likely to complain to Ebay's customer service than to give a negative rating. This suggests a severe underreporting of negative experiences. As mentioned above, Ebay eventually switched to a one-sided rating system.

Airbnb also has a two-sided rating system.¹⁴ Initially, reviews were immediately made public, allowing the possibility of retaliation. Fradkin, Grewal, and Holtz (2017) run field experiments and find that those who do not provide reviews tend to have worse experiences than those who do. They conclude that strategic reviewing behavior has occurred on Airbnb, although the overall bias appears to be small. Also, since buyer and seller may interact socially, they may be less inclined to leave negative reviews.

Airbnb no longer makes reviews public as long as the counterparty still has the option of posting a review and has not yet done so. While one party does not observe the counterparty's review prior to uploading her own review, there remain reasons for strategically underreporting negative experiences (in addition to the social interaction reason given above). Reviews are not anonymized, so somebody who rents out a flat can check the track record of somebody wanting to rent the flat. If that person tends to leave negative reviews, a future landlord may be less inclined to confirm the request. Anticipating this, the potential renter may be less harsh and leave positively biased reviews or no review at all.

A platform has various design options that affect the response rate and the informativeness of review and rating systems. For our purposes, we summarize the insights obtained so far by the following finding.

Finding 5. Rating systems may suffer from a lack of informativeness due to noise and bias introduced through the actions of buyers and sellers. In particular, platform users may game the system. This tends to reduce the strength of network effects.

4.3 Asymmetric herding behavior

A tendency to provide positive feedback, but to refrain from providing negative feedback, does not necessarily arise due to strategic considerations or independent mistakes by reviewers. It may also be the result of asymmetric

¹³ There is, of course, an easy way for the platform to avoid such retaliation possibilities: ratings may be disclosed only after the other party has provided the rating, or the time window to leave ratings has closed.

¹⁴ For descriptive statistics on Airbnb's rating system, see Zervas, Proserpio, and John (2015).

herding behavior. Muchnik, Aral, and Taylor (2013) conduct a randomized field experiment with fake ratings of comments on posted articles on a news website and analyze the dynamics of future feedback. They observe an asymmetric response to a fake positive rating compared to a fake negative rating. They find that a fake positive rating increases the probability of accumulating positive herding by 25%. While a fake negative rating also increases subsequent negative votes, this was neutralized by offsetting positive votes. Thus, there is herding on positive but not on negative ratings—Muchnik, Aral, and Taylor call this a ‘social influence bias.’

These results were obtained in a news setting and not in shopping contexts, but they are suggestive of reviewer behavior also in the latter contexts. This suggests that paid-for fake positive reviews can generate positive herding on B2C and C2C platforms. Thus, the damage done from a positive fake review would not be corrected if the fake report were not removed immediately but at some later time (see Aral, 2014). As pointed out above, there are other reasons that ratings and reviews do not provide accurate information. This may also give rise to long-term effects thanks to herding.

4.4 Design of the rating system

In the analysis above, we identified reasons that rankings and reviews lose informativeness because of the actions taken by the transaction partners. The assumption was that the platform aims to maximize the informativeness, possibly battling against errors and gaming. While more-informative rankings and reviews tend to make the platform more attractive (and are a source of positive network effects), a for-profit platform is ultimately interested in maximizing profit. It may, then, have an incentive to sacrifice informativeness if that increases its revenues. In addition to measures taken by the platform that affect the aggregate rankings of products or sellers, the platform may vary the ordering and display of individual reviews. The findings by Vana and Lambrecht (2018) provide some indications how a different design of the listing of reviews can affect purchase probability.

The literature on certifying intermediaries provides some insights into the design of rating systems by a profit-maximizing platform. In particular, platforms may deliberately design their system so as to avoid the worst offending behavior—that is, it features a minimum quality threshold—but to offer few clues about product quality otherwise. In such a case, rating inflation and presumed design flaws that limit the informativeness of a rating system would actually indicate that a profit-maximizing intermediary with market power sacrifices buyer participation in favor of higher margins. This is the lesson one can draw

from the work on certifying intermediaries by Lizzeri (1999), who shows in an adverse selection environment that a platform discloses only whether a product satisfies a minimum quality threshold.¹⁵ In his setting, a monopoly intermediary charges a fee to sellers for providing its certification service.¹⁶ As a result, the intermediary certifies minimum quality for products that are traded via the intermediary. Translated into the context of rating systems, the platform commits to its rating system and charges sellers for being listed. Thus, Lizzeri's result says that the rating system is designed in such a way that only the worst offenders disappear from the platform.

Finding 6. A profit-maximizing platform may deliberately design its rating system so as to limit its informativeness. As a result, sellers of rather low quality may do better on such a platform than on a platform that maximizes the quality of its rating system, while high-quality sellers do worse.

Bouvard and Levy (2016) further investigate the potential tension between informativeness and rent extraction. In their setting, the platform cannot commit to a certification technology and establishes a reputation for accuracy; for its service, it charges a fixed fee to participating sellers upfront. Applied to ratings systems, this means that the platform can redesign features that reflect the rating system's accuracy; and the fixed fee corresponds to a listing fee charged to sellers, as is observed, for example, on some price search engines.

Sellers have different opportunity costs of providing high quality. While higher accuracy attracts high-quality sellers, it repels low-quality sellers. As a result, the profit of a platform is first increasing and then decreasing in the level of accuracy it provides to sellers seeking certification. Thus, a profit-maximizing platform provides an intermediate level of accuracy. Applied to rating systems, instead of offering certification, a platform may make use of buyer reviews and ratings to (noisily) reveal quality. The design decisions regarding the rating system then affect its accuracy.

Platform competition improves the information available to buyers when sellers have to make a discrete choice between platforms: it enables full disclosure in the Lizzeri's (1999) setting and increases accuracy in Bouvard and Levy's

¹⁵ Similarly, Albano and Lizzeri (2001) analyze a moral hazard problem.

¹⁶ The timing is as follows: first, the intermediary sets its fee and commits to an information disclosure policy. Second, after observing the intermediary's decision, sellers decide whether to pay the fee, offer their products through the intermediary, and submit their product for testing. Third, consumers observe all previous decisions, and the seller makes a take-it-or-leave-it offer.

(2016) setting. By contrast, under seller multihoming, Bouvard and Levy (2016) show that platforms have weaker incentives for accuracy under competition.

III. RECOMMENDATIONS

As we discussed in the previous section, buyers can obtain valuable information from reviews and ratings by other buyers. In this case, the role of the platform is twofold: first, it invites buyers to evaluate various offers that have proved successful or popular with others; second, it organizes the exchange of the information across users (possibly combined with some policing so as to ensure that abuses are contained and mistakes are corrected). Since buyers actively provide and access the information, we may consider ratings and reviews as part of a platform's *information-pull* strategy.

In this section, we examine an alternative strategy of platforms, which consists of making recommendations to specific buyers. Such recommendations, based on popularity and on other sources of information, are an attempt to reduce search costs. Hence, platforms pursue an *information-push* strategy, as they advertise specific products to buyers based on their characteristics and observed behavior. Naturally, information-pull and -push strategies are not mutually exclusive—quite the contrary, as ratings and reviews often serve as inputs for recommendation algorithms. For instance, Amazon makes product suggestions, and buyers then access additional information before making their purchase decision.

In what follows, we first analyze how recommender systems, such as rating systems, generate network effects (Section 1). Next, we examine how recommender systems affect the distribution of sales (Section 2): do they contribute to making popular products even more popular, or do they drive consumers to discover niche products? Finally, we look into platforms' incentives to manipulate recommender systems (Section 3).

1. Product recommender systems and network effects

In this Section, we argue that product recommender systems are the source of positive network effects. This insight is easily established when buyers have homogeneous tastes and make mistakes, and the recommender system is based on the popularity of a product. Suppose that there are two products that can be ranked by their attractiveness. Product *A* is more attractive than product *B*; more specifically, suppose that product *A* gives a net benefit of 1 and product *B* of -1 . Consumers arrive sequentially and can be of two types:

'amateur' or 'expert.' An amateur consumer bases her decision on popularity, while an expert consumer acquires information about product features and makes a purchase based on that information.

To construct a numerical example, suppose that 50% of buyers follow a recommendation if they receive one and otherwise do not buy, while the remaining 50% collect information and, with 80% probability, make the right choice—i.e., with 20% probability, they erroneously choose the inferior product. The recommender system recommends the product that is purchased more. We will show that the last buyer is better off if there are more fellow buyers. Let us start with two buyers. If buyer 2 is an amateur, she makes an expected benefit $0.5(0.8-0.2)=0.3$, as, with 50% probability, buyer 1 was an expert (that is, buyer 1 purchased and, thus, indirectly recommended, the 'good' product with 80% probability and the 'bad' product with 20% probability). If buyer 2 is an expert, she makes an expected benefit of $0.8-0.2=0.6$. Hence, the expected benefit of buyer 2 is 0.45 (i.e., the average of 0.3 and 0.6, as she has equal chances of being either type).

Now consider the case with three buyers. If the third buyer is an expert, her expected benefit continues to be 0.6 (as the recommender system has no influence on her decision). If the third buyer is an amateur, she purchases only if the recommender system points her to the most popular product. For this to happen, the two previous buyers must have purchased one product more than the other. Let us examine when this does and does not happen. Four cases have to be distinguished according to the type of the successive buyers; each case has the same probability of occurrence—25%. The first case is the succession of two amateurs: as neither of them purchased, the recommender system remains silent, and the third buyer does not purchase either, yielding her a benefit of zero. Second, if the first buyer is an amateur (who, therefore, did not purchase) and the second is an expert, then the system recommends the good product with an 80% probability, and the bad product with a 20% probability, yielding the third buyer an expected benefit of $0.8-0.2=0.6$. Third, if the first buyer is an expert and the second an amateur, the configuration is similar to the previous one (as the second buyer follows the recommendation resulting from the first buyer's purchase decision); the expected benefit of the third buyer is again equal to 0.6. Finally, if there is a succession of two experts, both must have made the same choice for the recommender system to be informative (and so for the third buyer to purchase); this is so if they both decide to buy the good product (with 64% probability) or the bad product (with 4% probability); the third buyer's benefit in this case is then equal to $0.64-0.04=0.6$. In sum, if the third buyer is an amateur, her expected benefit is $0.25 \times 0 + 3 \times 0.25 \times 0.6 = 0.45$. Hence, the expected benefit of the third buyer is $0.5 \times 0.6 + 0.5 \times 0.45 = 0.525$.

Comparing the two cases, we observe that the last of three buyers has a larger expected benefit (0.525) than the last of two buyers (0.45). Hence, we have established that the last buyer benefits if more previous buyers are around and that buyers, prior to knowing their position in the sequence, are also better off if more fellow buyers are present. In this example, amateurs benefit from more buyers, as it becomes more likely that an expert has been around previously.

Finding 7. By recommending more-popular products, product recommender systems have the potential to provide purchase-relevant information to amateur buyers. In an e-commerce context, they have the potential to generate network effects, as a buyer is better off the more fellow buyers that are around.

A recommender system may also help to reduce the search cost. Suppose that there are several products, some of which are considered clear failures and a few that can be considered serious options. Absent recommendations based on popularity, a consumer may have to inspect quite a large number of products. With such recommendations, the consumer can restrict her search to the subset of serious options and, thus, reduce her expected search costs.

Finding 8. Product recommender systems have the potential to reduce search costs. In an e-commerce context, they have the potential to generate network effects, as a larger number of buyers provides more reliable information about which products are serious options.

If some consumers are frequent shoppers, while others buy only occasionally, the former make larger contributions to the functioning of the recommender system than the latter. As an illustration, suppose that frequent shoppers buy several products from a large set, whereas occasional buyers buy only one. The shopping behavior of frequent buyers allows the recommendation system to help other frequent shoppers to more easily find other products of interest. Thus, the recommender system generates positive within-group external effects among frequent shoppers.

If the recommender system can access additional information on occasional shoppers (e.g., that they are close to certain frequent shoppers in a friendship network), information gathered on frequent shoppers may also allow for useful recommendations to casual shoppers. In this case, there is a positive cross-group external effect from frequent shoppers to occasional shoppers. By contrast, information on purchase decisions by occasional shoppers is of little or

no help in making better recommendations to other shoppers. More generally, not only the total number of users, but the composition of the recommendation network, matter for the functioning of the recommender system.

Recommender systems can also be important on *two-sided platforms*. Here, the platform can make recommendations to both sides with the aims of reducing search costs and improving expected match quality. These recommendations may be based not only on observables of the two individual users on either side, but also on the behavior of other users on both sides.

Finding 9. Partner recommender systems have the potential to reduce search costs. In a two-group matching context, they have the potential to generate positive cross-group external effects, as more participation by one group generates the chance for the platform to propose matches that are more attractive for members of the other group, and **viceversa**.

We note that while both sides tend to benefit from such cross-group external effects, the benefits may vary depending on the terms of transaction between users on both sides. These terms of transaction for a particular user may also depend on participation levels on the same side. For instance, if buyers for collectibles receive better recommendations, they may drive up the price and, thus, receive a smaller fraction of the generated surplus.

2. Product recommender systems and the long tail

In many internet markets, a limited number of items (often a few hundred) account for the bulk of sales, while the vast majority of items (which constitute the tail of the distribution) sell only very few units. It has been argued that internet markets have a longer tail in the sales distribution than traditional markets.¹⁷ The question we address in this section is how recommender systems affect the distribution of sales: do they reinforce the skewness of the distribution, or do they make the tail longer, or thicker? We first discuss the main effects that recommender systems can have; we then formalize the intuition in a specific model, before reviewing recent empirical work.

— *Heterogeneous tastes and recommendations*

Since buyers often do not have homogeneous tastes, a recommender system reporting the popularity of different products may provide information

¹⁷ For an informal account, see Anderson (2006).

about which types of consumers may like a specific product. In particular, some buyers may be aware that they have a taste for niche products in a certain product category, whereas others may realize their preference for the standard products that cater to the taste of the mass market. Recommender systems may be based on popularity information—that is, information displaying in relative terms how often a product has been purchased. As a fictional example, consider a supermarket selling different types of cheese and providing popularity information. If you are new to the store and know that you like to avoid unpleasant surprises, you may opt for the popular cheese varieties. However, if you know that you like new taste experiences, you may opt for cheese varieties that are bought less frequently. In such a situation, the fact that a product has or has not been sold often provides valuable information to new buyers. A buyer with a niche taste may buy products that sold little in the past, whereas a buyer with a mass-market taste will purchase products that sold a lot in the past.

In practice, buyers may encounter products with mass or niche appeal and, in addition, suffer from not being able to judge product quality *ex ante*. It may then appear to be difficult to disentangle popularity information as a proxy for quality from popularity information as an indication of whether a product is a mass-market product—one that provides a good fit to the taste of many buyers—or a niche market product—one that provides a good fit to the taste of only few buyers.

There are two borderline cases. In the first, all buyers have the same taste and care only about quality. High quality proves to be more “popular” and accounts for a larger volume of sales if some consumers are informed about the product quality and buy only high quality, whereas others are not and, thus, have to randomize over several products of different qualities. Higher quality, then, turns out to be more popular. To resolve the asymmetric information problem, a platform may want to resort to a rating system, as analyzed in the previous section. Thus, the effect of such a rating system is to divert demand from a low-quality product to a high-quality product. In the other borderline case, buyers are uncertain only about whether the product better serves the mass or the niche market, leading to the outcome above.

A different situation arises if buyers observe whether a product is meant to cater to the mass or to the niche market, but they do not observe the product quality. To address the role of popularity information in guiding buyer behavior in such a situation, we present a simple model in which firm behavior is treated as exogenous—in particular, the prices of all products are fixed. As we will show, in such a scenario—in which consumers know in advance whether some product features fit their taste but are not fully informed about a quality dimension of

the product—a recommender system reporting the popularity of a product may also provide valuable information to consumers.

— *A specific model*

The model goes as follows.¹⁸ Suppose that consumers face a choice problem of buying one unit of two products offered by two different sellers; they may buy none, one, or both. Prices are fixed throughout the analysis. With probability $\lambda > 1/2$, a consumer thinks more highly of product 1 than of product 2; consequently, product 1 can be called a mass-market product and product 2 a niche product. Each product can also be of high or low quality with equal probability.

The consumer's utility depends both on the quality of the product and on whether the product matches her taste. A high-quality product that provides the wrong match is assumed to give net utility $v_H=1$ and a low-quality product, $v_L=0$. A product with the right match gives the previous net utilities augmented by t . These utilities are gross of the opportunity cost z that a consumer incurs when visiting a seller (e.g., clicking onto its website). A consumer knows her match value and receives a noisy private signal about quality. The noisy quality signal may come from noisy information in the public domain, such as publicly revealed tests. The ex ante probability of high quality is assumed to be $1/2$. The probability that the signal provides the correct information is ρ , which, for the signal to be informative but noisy, lies between $1/2$ and 1 . Hence, with a positive signal realization, the posterior belief that the product is of high quality is ρ . It follows that if a consumer who prefers product i receives a high-quality signal and buys from seller j , she obtains expected utility $U_{Hg} \equiv \rho + t - z$ if $i=j$ (i.e., if seller j offers the product that matches consumer i 's taste), and $U_{Hb} \equiv \rho - z$ if $i \neq j$. Correspondingly, with a low-quality signal, expected utility is $U_{Lg} \equiv (1 - \rho) + t - z$ if $i=j$ and $U_{Lb} \equiv (1 - \rho) - z$ if $i \neq j$. Table 1 displays the four possible levels of expected utility.

TABLE 1
EXPECTED UTILITY ACCORDING TO SIGNAL AND MATCH

	Good match	Bad match
High-quality signal	$U_{Hg} \equiv \rho + t - z$	$U_{Hb} \equiv \rho - z$
Low-quality signal	$U_{Lg} \equiv (1 - \rho) + t - z$	$U_{Lb} \equiv (1 - \rho) - z$

¹⁸ The model exposition is, in large part, identical to the one in Belleflamme and Peitz (2015, Chapter 15). It is based on Tucker and Zhang (2011).

For a given match, $\rho > 1/2$ implies that the consumer is better off with a high-quality signal: $U_{Hk} > U_{Bk}$ for $k=g, b$. Also, for a given signal, $t > 0$ implies that the consumer prefers to have a good match: $U_{Kg} > U_{Kb}$ for $K=H, L$. What is unclear is how the consumer balances the quality of the match with the quality of the signal. The consumer finds the quality of the match more important if $U_{Lg} > U_{Hb}$, which means that she is better off with a low-quality signal and a good match than with a high-quality signal and a bad match. This is so if $1+t > 2\rho$. Otherwise, the quality of the signal outweighs the quality of the match.

We first consider the product choice of a single buyer—this is the situation encountered by buyers when no recommender system is available. A buyer purchases the product independently of the signal realization and match value if $U_{Lb} > 0$; that is, the opportunity cost of visiting a seller is sufficiently small, $z < z_{Lb} \equiv 1 - \rho$. By contrast, if the opportunity cost is too large, the consumer will never buy. This is the case if $U_{Hg} < 0$, or, equivalently, if $z > z_{Hg} \equiv \rho + t$. Hence, we focus on the intermediate range where $z \in [z_{Lb}, z_{Hg}]$. A product with a good match but a low-quality signal is bought if $U_{Lg} \geq 0$, or, equivalently, if $z \leq z_{Lg} \equiv 1 - \rho + t$. A product with a bad match but a high-quality signal is bought if $U_{Hb} \geq 0$ or $z \leq z_{Hb} \equiv \rho$.

As indicated above, two scenarios are possible. In the first scenario, the buyer sees the quality of the match as more important; the inequality $U_{Lg} > U_{Hb}$ is equivalent to $z_{Lg} > z_{Hb}$, which becomes $1+t > 2\rho$. Thus, for this scenario to apply, consumer tastes must be sufficiently heterogeneous (t large) and signals sufficiently noisy (ρ small). In the second scenario, the quality of the signal matters more; we have $U_{Lg} < U_{Hb}$, or, equivalently, $z_{Lg} < z_{Hb}$. Thus, for this scenario to apply, consumer tastes must be sufficiently homogeneous (t small) and signals sufficiently informative (ρ large). Consumer choice can be fully described depending on whether $z_{Lg} > z_{Hb}$ or the reverse inequality holds.¹⁹

Second, we analyze buyer behavior in the presence of a *recommender system* that provides popularity information. For a recommender system to have any impact, we need at least another consumer who makes her choice after obtaining the information generated by the first consumer's choice. The

¹⁹ For $z_{Lg} > z_{Hb}$, we obtain that a product is bought by a consumer who does not observe a low-quality signal and a bad match if $z \in (z_{Lb}, z_{Hb})$; it is bought by a consumer who observes a good match if $z \in (z_{Hb}, z_{Lg})$; and it is bought by a consumer who observes a good match and a high-quality signal if $z \in (z_{Lg}, z_{Hg})$. For $z_{Lg} < z_{Hb}$, we obtain that a product is bought by a consumer who observes neither a low-quality signal nor a bad match if $z \in (z_{Lb}, z_{Lg})$; it is bought by a consumer who does not observe a low-quality signal if $z \in (z_{Lg}, z_{Hb})$; and it is bought by a consumer who observes a good match and a high-quality signal if $z \in (z_{Hb}, z_{Hg})$. Interestingly, in the first scenario, if $z \in (z_{Hb}, z_{Lg})$, consumer choice is determined purely by the match quality, whereas in the second scenario, if $z \in (z_{Lg}, z_{Hb})$, consumer choice is determined purely by the signal realization.

recommender system here simply reports the choice of the first consumer. The second consumer knows the parameters of the model but neither the signal realization nor the type of the first consumer. We assume that all random variables are i.i.d. across consumers (concerning the quality signal, this is conditional on true quality).

To analyze whether a recommender system favors mass-market products or niche products, we consider two cases: $z \in (z_{Lb}, \min\{z_{Hb}, z_{Lg}\})$ and $z \in (\max\{z_{Hb}, z_{Lg}\}, z_{Hg})$.²⁰ The former case is characterized by a relatively low cost of visiting sellers. Here, a consumer who observes a good match with a particular product always visits the corresponding seller. The consumer visits the seller of the product with a bad match only in case of high-quality information. This implies that click and purchase data still contain some useful information for the second consumer. The second consumer knows whether she has a taste for the niche product or the mass-market product. Hence, if she has a taste for the niche product, she knows that it is unlikely that the first consumer had the same taste. Therefore, it is quite likely that the first consumer's visit or purchase was driven by a positive realization of the quality signal. The opposite reasoning applies to a consumer who has a taste for the mass-market product. Here, click and purchasing data are less informative, thus implying that sellers of niche products benefit more from information on visits or purchases.

In the latter case, in which $z \in (\max\{z_{Hb}, z_{Lg}\}, z_{Hg})$, information on a *lack* of visits or purchases hurts the seller of the mass-market product more. While niche sellers are at a disadvantage matching consumer tastes, this disadvantage becomes an asset when it comes to consumer inferences about product quality. It increases the benefit due to favorable popularity information and reduces the loss due to unfavorable popularity information.²¹

Tucker and Zhang (2011) provide support for this theory in a field experiment. A website that lists wedding service vendors switched from an

²⁰ In addition, there are two intermediary cases—that is, $z \in (z_{Hb}, z_{Lg})$ for $1+t>2p$ and $z \in (z_{Lg}, z_{Hb})$ for $1+t<2p$. In the first case, in which $z \in (z_{Hb}, z_{Lg})$, the first consumer's choice does not reveal anything about her private signal. Hence, the recommender system does not contain any valuable information for the second consumer. In the second case, where $z \in (z_{Lg}, z_{Hb})$, the first consumer's choice is determined solely by the signal realization. The second consumer will then use the information provided by the recommender system to update her beliefs: she updates her quality perception upwards if a particular product has been bought (purchase data) or if the seller has been visited (click data). This implies that a previous visit or purchase increases the chance of subsequent visits and purchases. Here, the recommender system favors the sale of high-quality products.

²¹ An interesting question, which we do not analyze here, is the possibility of rational herding. This is a situation in which consumers ignore their private information and rely fully on the aggregate information provided by the system. This means that learning stops at some point. A seminal paper on rational herding is Banerjee (1992). Tucker and Zhang (2011) also address herding in the present context.

alphabetical listing to a popularity-based ranking in which offers are ranked by the number of clicks the vendor receives. The authors measure vendors when located in towns with a large population as having broad appeal and when located in small towns as having narrow appeal. Tucker and Zhang find strong evidence that narrow-appeal vendors receive more clicks than broad-appeal vendors when ranked similarly in the popularity-based ranking.

Finding 10. Product recommender systems reporting product popularity may affect mass-market and niche products differently. Given a similar ranking, niche products tend to do relatively better with such a recommender system.

A prominent mix of various recommender systems is in place at Amazon.com. Perhaps the most notable example (at least in product categories in which consumers do not search among product substitutes) is that, when listing a particular product, Amazon recommends other products that consumers have purchased together with the displayed product. The economics of such a recommender system are different from a system that merely reports the popularity of products. It allows consumers to discover products that serve similar tastes and, thus, is likely to produce good matches at low search costs. Such a recommender system is based on previous sales and appears to be particularly useful in consumer decision-making for products that enjoy complementary relationships. It implies that products with no or limited sales will receive little attention. This reasoning suggests that recommender systems may work against the long tail, an argument in contrast to the view that people discover better matches on recommender systems. The latter view is based on the observation that consumers with very special tastes more easily find products that provide a good match to their tastes, so that they do not need to resort to very popular products or buy at random.

However, these two views are not necessarily contradictory. While the long-tail story refers to the diversity of aggregate sales, the discovery of better matches refers to diversity at the individual level. It might well be the case that people discover better matches through recommender systems but that they discover products that are already rather popular among the whole population. Hence, sales data in the presence of recommender systems may show more concentration at the aggregate level.²²

²² This point is made in the numerical analyses of Fleder and Hosanagar (2009). However, in their model, the recommendation network essentially provides information about the popularity of a product and does not allow for more fine-tuned recommendations.

— *Empirical work on recommender systems*

While the previous discussion brings interesting insights, empirical analyses will have to show whether recommender systems, indeed, lead to more concentrated sales; or whether the directed search, which is inherent in recommender systems, reduces users' search costs to the extent that they feel more encouraged to search outside of known products that they like, with the effect that diversity also increases at the aggregate level. Indeed, as can be shown formally, if the consumer population is characterized by taste heterogeneity, a recommender system that provides personalized recommendations may lead to a 'thicker' tail in the aggregate, meaning that less-popular products receive a larger share of sales after the introduction of a recommender system.²³ A likely outcome, then, is that more niche products will be put on the market and that product variety in the market will, therefore, increase.

Oestreicher-Singer and Sundararajan (2012a, 2012b) shed some light on this issue.²⁴ They collected a large data set, starting in 2005, of more than 250,000 books from more than 1,400 categories sold on Amazon.com.

They restrict their analysis to categories with more than 100 books, leaving them with more than 200 categories. For all the books, they obtain detailed daily information, including copurchase links—that is, information on titles that other consumers bought together with the product in question (and which Amazon prominently communicates to consumers). These copurchase links exploit possible demand complementarities. Since these links arise from actual purchases and not from statements by consumers, they can be seen as providing reliable information about what other consumers like. By reporting these links, Amazon essentially provides a personalized shelf for each consumer according to what she was looking at last. This allows consumers to perform a directed search based on their starting point. Oestreicher-Singer and Sundararajan (2012b) find that if a copurchase relationship becomes visible, this leads, on average, to a three-fold increase in the influence that complementary products have on each others' demand.

The question, then, is how these copurchase links affect sales. In particular: which products make relative gains in such a recommendation network? Are these the products that already have mass appeal (because they are linked to

²³ See Hervas-Drane (2015) for a formal analysis.

²⁴ Other relevant empirical work has been done by Brynjolfsson, Hu and Simester (2011) and Elberse and Oberholzer-Gee (2007). Brynjolfsson, Hu and Simester (2011) compare online and offline retailing and find that online sales are more dispersed. While compatible with the hypothesis that recommender networks lead to more-dispersed sales, other explanations can be given. Elberse and Oberholzer-Gee (2007), comparing DVD sales in 2005 to those in 2000, find that the tail had got longer in 2005. However, they also find that a few blockbusters enjoy even more sales; this is like a superstar effect. Again, the role of recommender systems is not explicit.

other products) or, rather, niche products? To answer this question, one must measure the strength of the links that point to a particular product. For this, it is important to count the number of links pointing to a product and to know the popularity of the products from which a link originates. Hence, a web page receives a high ranking if the web pages of many other products point to it or if highly ranked pages point to it. This is measured by a weighted page rank based on Google's initial algorithm. Oestreicher-Singer and Sundararajan (2012a) construct the Gini coefficient for each product category as a measure of demand diversity within a category. They regress this measure of demand diversity on the page rank (averaged within a category), together with a number of other variables. In their 30-day sample, they find that categories with a higher page rank are associated with a significantly lower Gini coefficient. This means that in a product category in which, on average, recommendations play an important role, niche products within this category do relatively better in terms of sales, whereas popular products perform relatively worse than in a product category where this is not the case. This is seen as evidence in support of the theory of the long tail.²⁵

The finding that a recommender system favors products in the long tail suggests that such a system may encourage participation on the seller side, as it becomes more attractive for niche players to become active. Since an increase in the number of buyers improves the granularity of the recommender system, a platform with a well-designed recommender system features positive cross-group external effects from buyers to marginal sellers.

Recommender systems may use information that is different from the actual purchases, but may also use hints of purchase intentions. For instance, Amazon can recommend products based on clicking behavior. If many people who looked at one product also took a close look at another product, this may suggest that the two products are closely related (as substitutes or complements) and that potential buyers benefit from cross-recommendations. We note that recommender systems may also have a future in physical retailing, provided that shoppers use a device that can provide personalized recommendations. For instance, in-shop displays may make personalized recommendations based on a shopper's history and the histories of fellow shoppers.

3. Search engine bias and quality degradation

As in the design of review and rating systems, platforms may have incentives that are not aligned with those of buyers. In particular, a profit-

²⁵ To take into account possible unobserved heterogeneity in the data, Oestreicher-Singer and Sundararajan (2012a) also construct a panel data set. The estimation results are confirmed with panel data techniques.

maximizing platform may have an incentive to distort the recommender system or make it less informative. The theoretical literature has uncovered several reasons that platforms operating as search engines may have an incentive to bias their search results. First, a platform may favor search results from which it can extract larger profits. Second, partial integration of the platform with some sellers or content providers may reinforce the previous motivation. Finally, a platform may discourage search so as to reduce competition among sellers. We examine these three motivations, in turn, and comment on empirical results when available.

— *Search bias to favor more-profitable sellers*

A platform may bias the order of recommendations if different offers lead to different commissions or to different purchase probabilities. Regarding the former, such higher margins occur if the platform has a specific partner program for which it charges higher commissions. Regarding the latter, if an offer is available on different distribution channels and some buyers multihome, these multihoming buyers are likely to purchase elsewhere if offers on alternative distribution channels are available at a lower price. Therefore, a profit-maximizing platform would place offers that were cheaper elsewhere in a lower position than if such lower-priced alternatives were not available.²⁶

Given such motivations, it is interesting to ask whether platforms list search results in the best interest of consumers. Hunold, Kesler, and Laitenberger (2017) empirically investigate this issue in the context of hotel booking sites. Booking and Expedia use a default to place their recommendations—Expedia calls this list “Recommended” and Booking “Top Picks.” These platforms do not provide clear information on how they construct the lists; this is in contrast to other listings that a user can obtain and that are based on price or reviewer ratings. Thus, platforms maintain discretion over how they order the available offers in the list. The authors use data from July 2016 to January 2017 from Booking, Expedia, and the meta-search site Kayak for hotels in 250 cities (most of them within Europe), featuring more than 18,000 hotels. They find that for a given price on a hotel booking platform, a lower price on the other platform or on the hotel’s website leads to a worse position on the list. This suggests that hotel booking platforms bias their recommendations.

The interaction between organic and sponsored links can provide another reason that search engines opt to bias their search results—this insight is relevant

²⁶ If the platform is allowed to impose a most-favored nation (MFN) clause that does not allow sellers to offer lower prices elsewhere, it no longer has the incentive to bias search results in that way. However, such MFN clauses have been declared illegal in several jurisdictions on competition grounds.

not only for general search engines, but also platforms such as Booking, which offers advertising opportunities in addition to providing organic search results.²⁷ As Xu, Chen, and Whinston (2012), Taylor (2013), and White (2013) point out, organic links give producers a free substitute to sponsored links on the search engine. Hence, if the search engine provides high quality in its organic links, it cannibalizes its revenue from sponsored links (if it is not able to fully recoup them through higher charges on its sponsored links). At the same time, providing better (*i.e.*, more reliable) organic search results makes the search engine more attractive. If consumers have search costs, a more attractive search engine obtains a larger demand. However, if the latter effect is (partially) dominated by self-cannibalization, a search engine optimally distorts its organic search results.

Finding 11. Profit-maximizing platforms may degrade the quality of their recommender systems or provide biased recommendations. This tends to reduce the size of within-group external effects among buyers.

— *Search bias due to partial integration*

A misalignment of buyer and platform incentives may also be the result of partial vertical integration. In particular, this may be alleged to give rise to or exacerbate search engine bias—an issue that received prominence in the Google Shopping case in the European Union. Does partial vertical integration lead to additional worries about *search engine bias*, or can integration possibly reduce search engine bias? In what follows, we present the models of de Cornière and Taylor (2014) and Burguet, Caminal, and Ellman (2015) to systematically analyze the costs and benefits of search engine integration.

De Cornière and Taylor (2014) analyze a market with a monopoly search engine, two websites, sellers and users. The websites offer horizontally differentiated content. This is formalized by the Hotelling line, with platform 1 located at point 0 and platform 2 at point 1, and users uniformly distributed on the unit interval. Prior to search, users are not aware of their preferred content. This implies that without searching, a user cannot identify which website has the content that interests her the most. A user incurs a user-specific search cost when engaging in search on the search engine (specifically, the search cost is drawn from some cumulative distribution function).

Websites and the search engine obtain revenues exclusively from advertising posted by sellers, which users are assumed to dislike. The search engine works

²⁷ Our discussion of search engine bias closely follows the exposition in Peitz and Reisinger (2016).

as follows: if a user decides to use the search engine, she enters a query. The search engine then directs the user to one of the websites. The search engine's decision rule is a threshold rule such that all users to the left of the threshold are directed to platform 1 and those to the right are directed to platform 2. A key assumption is that ads on the search engine and those on the media platforms are imperfect substitutes. That is, the marginal value of an ad on one outlet decreases as the number of advertisements on the other outlet increases. This implies that the advertising revenue generated by a website falls if the amount of advertising on the search engine rises (which is treated as exogenous).

The timing of the game is as follows. First, websites choose their advertising levels and the search engine chooses the threshold. Second, the advertising market clears. Third, users decide whether or not to rely on the search engine. Finally, those users who rely on the search engine type in a query and visit the website suggested by the search engine. When deciding whether or not to rely on the search engine, a user knows the threshold and has an expectation about the websites' advertising levels. The search engine is said to be biased if its chosen threshold differs from the one that maximizes the expected user utility (and, thus, the users participation rate).

The search engine faces the following trade-off. On the one hand, it is interested in high user participation. Other things equal, a larger number of search engine users leads to higher profits because advertisers are willing to pay more to the search engine. Therefore, the search engine cares about relevance to users. In addition, since users dislike advertising, they prefer to be directed to a site that shows few ads. These considerations align the incentives of the search engine with those of users. On the other hand, the search engine obtains profits from advertisers and, thus, aims to maintain a high price for its own links. Therefore, if ads on website i are particularly good substitutes for ads on the search engine, the search engine prefers to bias results against this website.

De Cornière and Taylor (2014) then analyze the effects of integration of the search engine with one of the websites—say, website 1. Suppose that there is partial integration without control of ad levels—that is, website 1 shares a fraction ρ_1 of its profit with the search engine but retains full control with respect to its ad level (this corresponds to partial ownership, but no control rights for the search engine). Then, the search engine has an incentive to bias its result in favor of website 1 because it benefits directly from this website's revenues. However, it also benefits more from higher user participation, implying that the search engine wants to implement higher quality (i.e., less-biased results). Because of these two potentially countervailing forces, partial integration can increase or decrease the level of bias. In particular, if the search engine were biased to the detriment of website 1 without integration, partial integration

might mitigate this bias. Even if the search engine is biased in favor of media outlet 1 without integration, partial integration can lead to a reduction in the bias. If the websites are symmetric, partial (or full) integration always leads to an increase in bias. However, users may be better off because of lower ad levels.

Burguet, Caminal, and Ellman (2015) propose a different setup to analyze the problem of search engine bias and integration. They do not account for ad nuisance but explicitly model consumer search for sellers' products. User i is interested in the content of one of the N websites only—this website is denoted by $n(i)$ —while any other content generates a net utility of zero. Each website's content interests the same fraction of users, $1/N$.

Users do not know which website matches their interests and need the help of a search engine. Suppose that the search engine can perfectly identify the relevant website $n(i)$ once a user i has typed in the search query. When using the search engine, a user incurs a search cost.²⁸ The search engine displays a link to a website after a user has typed in the query. The search engine chooses the probability that the link leads to the content matching the user's interest. Since the links to websites are non-paid, this corresponds to organic search.

The search engine also features sponsored search in which it advertises the sellers products. This is the source of profits for the search engine and websites. Sellers belong to one of J different product categories, indexed by j . User i values only one category $j(i)$. Each category's products interest the same fraction of users, $1/J$. There are two sellers in each category. Seller 1 provides the best match to a user, leading to a net utility of v_1 . Producer 2 provides a worse match such that $0 < v_2 < v_1$. The sellers' margins are m_1 and m_2 . Users' and sellers' interests are assumed to be misaligned, and, thus, $m_2 > m_1$. In addition, it is assumed that buyer preferences dominate for the welfare ranking—i.e., $v_1 + m_1 > v_2 + m_2$. The monopoly search engine provides a single link after a user has typed in a query for product search in a particular category.²⁹ Then, the search engine sets a pay-per-click price. The search engine chooses to display the link of producer 1 with some probability and the link of producer 2 with the remaining probability.³⁰

²⁸ The search cost is heterogeneous across consumers and drawn from some cumulative distribution function.

²⁹ Both models described here (Burguet, Caminal and Ellman, 2014, and de Cornière and Taylor, 2014) assume that users visit only a single website after typing in a query. However, in reality users may click on multiple search results (in sequential order). They can be expected to broadly follow the respective ranking of the results. In such a situation, advertisers exert negative externalities on each other when bidding for more prominent placement. Athey and Ellison (2011) and Kempe and Mahdian (2008) study the question of how the optimal selling mechanism of the search engine takes these externalities into account.

³⁰ This is a simplified version of the model of Burguet, Caminal, and Ellman (2015), which is developed in Peitz and Reisinger (2016).

Absent vertical integration, search results are distorted because websites compete for advertisers. As Burguet, Caminal, and Ellman (2015) show, generically, the search engine will distort, at most, one type of search—product search or content search—setting the other at the optimal value. If the search engine was integrated with all websites, it would internalize the externality exerted by one websites on others and, as a result, improve its reliability. This is an unambiguously positive effect. However, in case the search engine is integrated only with a fraction of the websites, it has an incentive to divert search from non-affiliated websites to affiliated ones. Here, partial integration may lead to a lower consumer surplus compared to no integration.

The findings from the theoretical literature suggest that search engine bias may arise due to (partial) integration. However, partial integration sometimes is a remedy for search engine bias prior to integration, and, in any case, its consumer-welfare implications are ambiguous. So, to ascertain whether recommender systems work better or worse under (partial) integration, a detailed understanding of the specific case is needed. What is clear is that when (partial) integration reduces bias and increases buyer participation, integration tends to improve the recommender system.

Finding 12. Partial integration of a platform with sellers or content providers may increase or decrease the bias of its recommender system. Even if partial integration increases bias, it may increase buyer participation and buyer surplus.

— *Search discouragement to reduce sellers' competition*

Finally, a platform may want to make its recommender system less informative so as to discourage search. Chen and He (2011) and Eliaz and Spiegler (2011) provide a reason that a search engine may bias its recommendations or search results if it takes a cut from the transaction between buyer and seller—this is a situation with sponsored links. In this case, it is in the search engine's best interest for sellers' revenues from sponsored links to be high. Because revenues increase if product market competition between sellers becomes softer, the search engine may distort search results so as to relax product market competition. As formalized in Chen and He (2011) and Eliaz and Spiegler (2011), a monopoly search engine has an incentive to decrease the relevance of its search results, thereby discouraging users from searching extensively. This quality degradation leads to less competition between sellers and, thus, to higher seller revenues, which can be partly extracted by the search engine.

IV. CONCLUSION

It is our contention that one cannot understand the functioning of prominent digital platforms such as Airbnb, Amazon, Booking, Expedia, Ebay, Google Shopping and Uber without taking proper account of their rating and recommender systems.

Such systems are crucial for the performance of digital platforms for the following simple reason: potential buyers incur an opportunity cost in evaluating how products and services fare in terms of quality and how they fit their tastes; thus, they appreciate ratings, reviews and recommendations because knowing what other buyers did in the past helps them to make better-informed decisions. Rating and reviews are particularly useful for product characteristics that everyone appreciates (in terms of value for money—these characteristics may be observable prior to purchase or only after purchase, possibly only by a fraction of buyers). In the presence of taste heterogeneity, buyers benefit from personalized recommendations, which help them find their way in selecting products.

When two-sidedness is an essential feature of a digital platform, users are often keen to infer information about the reliability of the counterparties to the transactions that they may conduct on the platform. Here, rating systems can possibly steer buyers away from low-quality sellers and can discourage sellers from misbehaving. Conversely, thanks to rating systems, sellers can stay clear of problematic buyers, and buyers may have a stronger incentive to behave properly.

In this [article](#), we have analyzed the economic roles that rating and recommender systems play. In particular, we have shed light on how the effectiveness of these systems depends on the joint actions of their users and designers: not only can buyers and sellers take actions that damage the functioning of rating systems, but for-profit platforms also may have an incentive to manipulate their rating and recommender systems. Finally, throughout our analysis, we have argued that rating and recommender systems are the source of positive within-group and cross-group external effects. They are, thus, in many cases, a key driver allowing a platform to attract many buyers (and, if applicable, sellers), which is an undeniable source of competitive advantage in markets with competing platforms.

BIBLIOGRAPHY

ALBANO, G., and A. LIZZERI (2001), "Strategic certification and the provision of quality," *International Economic Review*, 42: 267-283.

ANDERSON, C. (2006), *The Long Tail: Why the Future of Business is Selling Less of More*, Hyperion Press, New York.

ARAL, S. (2014), "The problem with online ratings," *MIT Sloan Management Review*, 55: 45-52.

ATHEY, S., and G. ELLISON (2011), "Position auctions with consumer search," *Quarterly Journal of Economics*, 126: 1213-1270.

BAJARI, P., and A. HORTACSU (2004), "Economic insights from internet auctions," *Journal of Economic Literature*, 42: 457-486.

BANERJEE, A. V. (1992), "A simple model of herd behavior," *Quarterly Journal of Economics*, 107: 797-817.

BELLEFLAMME, P., and M. PEITZ (2015), *Industrial Organization: Markets and Strategies*, 2nd edition, Cambridge University Press, Cambridge.

— (2018a), *The Economics of Platforms*, book manuscript, work in progress.

— (2018b), "Platforms and network effects," in L. CORCHON and M. A. MARINI (eds.), *Handbook of Game Theory and Industrial Organization*, vol. II, Edward Elgar Publisher.

BOLTON, G.; GREINER, B., and A. OCKENFELS (2013), "Engineering trust: Reciprocity in the production of reputation information," *Management Science*, 59(2): 265-285.

BOUVARD, M., and R. LEVY (2016), "Two-sided reputation in certification markets," *Management Science*, forthcoming.

BRYNJOLFSSON, E.; HU, Y., and D. SIMESTER (2011), "Goodbye Pareto Principle, hello long tail: The effect of search cost on the concentration of product sales," *Management Science*, 57: 1373-1386.

BURGUET, R.; CAMINAL, R., and M. ELLMAN (2015), "In Google we trust?," *International Journal of Industrial Organization*, 39: 44-55.

CABRAL, L., and A. HORTACSU (2010), "The dynamics of seller reputation: Evidence from Ebay," *Journal of Industrial Economics*, 58: 54-78.

CAI, H.; JIN, G. Z.; LIU, C., and L.-A. ZHOU (2014), "Seller reputation: From word-of-mouth to centralized feedback," *International Journal of Industrial Organization*, 34: 51-65.

CHEN, Y., and C. HE (2011), "Paid placement: Advertising and search on the internet," *Economic Journal*, 121: F309-F328.

CHEVALIER, J. A., and D. MAYZLIN (2006), "The effect of word of mouth on sales: Online book reviews," *Journal of Marketing Research*, 43: 345-354.

CORNIÈRE, A. DE, and G. TAYLOR (2014), "Integration and search engine bias," *Rand Journal of Economics*, 45: 576-597.

ELBERSE, A., and F. OBERHOLZER-GEE (2007), Superstars and underdogs: An examination of the long tail phenomenon in video sales, *MSI Reports: Working Paper Series*, 4: 49–72.

ELIAZ, K., and R. SPIEGLER (2011), "A simple model of search engine pricing," *Economic Journal*, 121: F329-F339.

FLEDER, D., and K. HOSANAGAR (2009), "Blockbuster culture's next rise or fall: The impact of recommender systems on sales diversity," *Management Science*, 55: 697-712.

FRADKIN, A.; GREWAL, E., and D. HOLTZ (2017), *The determinants of online review informativeness: Evidence from field experiments on Airbnb*, unpublished manuscript.

HERVAS-DRANE, A. (2015), "Recommended for you: The effect of word of mouth on sales concentration," *International Journal of Research in Marketing*, 32: 207-218.

HUI, X.-A.; SAEEDI, M.; SUNDARESAN, N., and Z. SHEN (2016), "Reputation and regulations: Evidence from Ebay," *Management Science*, 62: 3604-3616.

HUNOLD, M.; KESLER, R., and U. LAITENBERGER (2017), *Hotel rankings of online travel agents, channel pricing and consumer protection*, unpublished manuscript.

JIN, G. Z., and A. KATO (2006), "Price, quality, and reputation: Evidence from an online field experiment," *Rand Journal of Economics*, 37: 983-1005.

KEMPE, D., and M. MAHDIAN (2008), "A cascade model for advertising in sponsored search", *Proceedings of the 4th International Workshop on Internet and Network Economics (WINE)*.

KLEIN, T.; LAMBERTZ, C., and K. STAHL (2016), "Adverse selection and moral hazard in anonymous markets," *Journal of Political Economy*, 124: 1677-1713.

LI, L. I.; TADELIS, S., and X. ZHOU (2016), Buying reputation as a signal: Evidence from online marketplace, *NBER Working Paper*, No. 22584.

LIVINGSTON, J. A. (2005), "How valuable is a good reputation? A sample selection model of internet auctions," *Review of Economics and Statistics*, 87(3): 453-465.

LIZZERI, A. (1999), "Information revelation and certification intermediaries," *Rand Journal of Economics*, 30: 214-231.

LUCA, M., and G. ZERVAS (2016), "Fake it till you make it: Reputation, competition, and Yelp review fraud," *Management Science*, 62: 3412-3427.

MAYZLIN, D.; DOVER, Y., and J. CHEVALIER (2014), "Promotional reviews: An empirical investigation of online review manipulation," *American Economic Review*, 104: 2421-2455.

MCDONALD, C. G., and V. C. SLAWSON (2002), "Reputation in an internet auction market," *Economic Inquiry*, 40: 633-650.

MELNIK, M. I., and J. ALM (2002), "Does a seller's ecommerce reputation matter? Evidence from ebay auctions," *Journal of Industrial Economics*, 50: 337-349.

MILGROM, P., and J. ROBERTS (1986), "Price and advertising signals of product quality," *Journal of Political Economy*, 94: 796-821.

MUCHNIK, L.; ANAL, S., and S. J. TAYLOR (2013), "Social influence bias," *Science*, 341: 647-651.

NOSKO, C., and S. TADELIS (2015), The limits of reputation in platform markets: An empirical analysis and field experiment, *NBER Working paper*, No. 20830.

OESTREICHER-SINGER, G., and A. SUNDARARAJAN (2012a), "Recommendation networks and the long tail of electronic commerce," *MIS Quarterly*, 36(1): 65-83.

— (2012b), "The visible hand? Demand effects of recommendation networks in electronic markets," *Management Science*, 58: 1963-1981.

OTT, M.; CARDIE, C., and J. HANCOCK (2012), "Estimating the prevalence of deception in online review communities," *Proceedings of the 21st international conference on World Wide Web*: 201-210.

PEITZ, M., and M. REISINGER (2016), "Media economics of the internet," in S. ANDERSON, D. STROMBERG and J. WALDFOGEL (eds.), *Handbook of Media Economics*, vol. 1A, North Holland (2016): 445-530.

RESNICK, P.; ZECKHAUSER, R.; SWANSON, J., and K. LOCKWOOD (2006), "The value of reputation on Ebay: A controlled experiment," *Experimental Economics*, 9: 79-101.

TADELIS, S. (2016), "Reputation and feedback systems in online platform markets," *Annual Review of Economics*, 8: 321-340.

TAYLOR, G. (2013), "Search quality and revenue cannibalisation by competing search engines," *Journal of Economics and Management Strategy*, 22: 445-467.

TUCKER, C., and J. ZHANG (2011), "How does popularity information affect choices? Theory and a field experiment," *Management Science*, 57: 828-842.

VANA, P., and A. LAMBRECHT (2018), *Online reviews: Star ratings, position effects and purchase likelihood*, unpublished working paper.

WHITE, A. (2013), "Search engines: Left side quality versus right side profits," *International Journal of Industrial Organization*, 31: 690-701.

XU, H.; LIU, D.; WANG, H., and A. STAVROU (2015), "E-commerce reputation manipulation: The emergence of reputation-escalation-as-a-service," *Proceedings of 24th World Wide Web Conference (WWW 2015)*: 1296-1306.

XU, L.; CHEN, J., and A. WINSTON (2012), "Effects of the presence of organic listing in search advertising," *Information System Research*, 23: 1284-1302.

ZERVAS, G.; PROSERPIO, D., and B. JOHN (2015), A first look at online reputation on Airbnb, where every stay is above average, *Working paper*, Boston University.