

Recommender systems and the value of user data^{*}

Gunhaeng Lee[†] Julian Wright[‡]

October 21, 2025

Abstract

Artificial intelligence (AI) increasingly shapes economic decisions through recommender systems, which personalize recommendations from user data. Yet, how the value from personalization scales with data remains unclear. We introduce a probability-based model designed to flexibly accommodate varying degrees of recommendation personalization. Using this model, we derive the value of recommender systems at different personalization levels, and decompose the value into customization, selection, and screening components. We also find a condition under which increased customization strictly benefits users. Our analysis reveals that deeper customization can yield non-diminishing returns and super-additivity. Empirical analyses using Bayesian predictions support our theoretical insights.

Keywords: recommender system, collaborative filtering, personalized recommendations, big data.

^{*}We thank Maximilian Schäfer, Juwon Seo and Tat-How Teh, as well as seminar participants at NUS, National Taiwan University, Academia Sinica, National Taipei University, SAET, APIOC, Baptist-Kyoto-NTU-Osaka-Sinica Workshop on Economic Theory. All errors are ours.

[†]Institute of Economics, Academia Sinica, glee@econ.sinica.edu.tw

[‡]Department of Economics, National University of Singapore, jwright@nus.edu.sg

1 Introduction

Artificial intelligence (AI) increasingly shapes economic decision making, with recommender systems emerging as a key example of how AI adds value to users. These AI-driven tools convert high dimensional user-item interaction data into personalized recommendations for products, media content, and services. Rather than peripheral features, recommender systems represent core engines of value creation for major platforms such as YouTube, TikTok, Instagram, Netflix, Spotify, and Amazon, directly influencing user choices and platform revenues. Despite their widespread adoption and economic significance, our understanding of how value is generated from and scales with the volume of data and customization remains limited, especially in settings with both within-user and across-user learning from large-scale data.

This paper addresses these gaps by proposing an economic model that captures how increases in data scale and algorithmic capacity translate into deeper customization, and how this personalization in turn generates value for users. The evolution of recommender systems illustrates a clear trajectory toward this goal. Early systems, limited by data processing capabilities, relied on simple metrics like popularity or aggregate ratings, which offered limited personalization. Over the last few decades, we have witnessed a steady, and at times radical, progression in recommender system technology. The move from naïve, popularity-based systems to more sophisticated methods like collaborative filtering marked a significant leap in personalization. More recently, the advances in generative AI models represent another step in this evolution, enabling even deeper levels of customization.¹ This expanded capability allows recommender systems to integrate diverse user behaviors and contextual features, creating significant value for users. For instance, music platforms can help a user discover novel songs matching their niche tastes, while online marketplaces can identify complementary products to complete a purchasing goal.²

In our stylized model, a platform aggregates historical user feedback to generate personalized probability estimates. While our framework can incorporate any observable user data, we focus specifically on user ratings for simplicity, as the model generalizes to other data forms through straightforward relabeling. The platform uses historical data to identify correlations across items and predicts the likelihood that a new user will like each item conditional on their past ratings on other items. Our recommendation framework closely mirrors the practical stages of modern

¹For an extensive review comparing these paradigms of recommender systems, see Ayemowa et al. (2024) and Wu et al. (2024).

²For detailed insights, see Subbiah and Aggarwal (2024).

systems: item retrieval, ranking based on predicted probabilities, and threshold-based filtering. Recommended items yield positive payoffs if tried and liked, negative payoffs if disliked, and zero otherwise.

Our study provides three main contributions. First, we introduce a tractable model that explicitly accommodates different degrees of customizations. This flexible model allows for both theoretical analysis and empirical implications. Within this framework, we quantify the economic value created by personalization in recommendations, decomposing it into three distinct components: *customization*, *selection* and *screening*. Customization segments users based on their historical data; selection identifies the best item(s) to recommend from a pool of candidate items; screening determines if the selected item(s) are good enough to be recommended, filtering out suggestions that fall below a threshold level. We show that the user-optimal threshold that perfectly aligns with user’s interest by recommending items only if the user’s expected payoff is positive, ensures additional customization always weakly increases user welfare. Conversely, we show that only the user-optimal threshold provides this guarantee. For any misaligned threshold, we demonstrate that there always exist patterns in how items are correlated for which deeper customization strictly harms certain user segments. We then derive a necessary and sufficient condition under which deeper customization always strictly increases user welfare.

Second, we characterize the learning curve as data grows along two distinct dimensions: breadth and depth. Typically, analyses of increased data focus primarily on breadth—expanding the number of users from whom data is collected. However, in AI-driven recommender systems, the depth of data—detailed individual user information—is equally critical for precision. We investigate how user welfare evolves as customization deepens. Unlike the breadth dimension, which consistently yields positive but diminishing returns, increasing data depth can produce non-diminishing and even increasing marginal returns. Moreover, the depth dimension can display super-additivity, where the value of recommendations conditioned on multiple items surpasses the sum of values derived from conditioning on each item individually. We formalize the conditions under which this synergy occurs. These two dimensions jointly create a learning curve composed of cascading, overlapping diminishing return curves. Increasing the breadth of data raises user welfare at a diminishing rate, but simultaneously enables deeper customization, which initiates a new, higher potential welfare curve. Consequently, as data expands along both dimensions, user welfare progresses upward through a series of these curves.

Our third contribution empirically validates these theoretical insights using a Bayesian proba-

bility estimation model applied to a dataset of over four million anonymous joke ratings from 73,421 users. Counterfactual analyses demonstrate substantial welfare improvements: a rating system that does not offer customization increases user welfare by 30.4% over one with no rating system, and customizing recommendations yields an additional 33.7% increase. After demonstrating that both the marginal value of additional users and deeper customization exhibit diminishing returns in our data, we assess the extent of complementarity or substitutability between these two dimensions of learning. We find the number of previous users exhibits strong complementarity with the degree of customization when both the number of users and the degree of customization is low, but this complementarity disappears quickly as the system accumulates data from more users.

The remainder of the article proceeds as follows. Section 2 surveys the related literature. Section 3 formally defines the recommender system, develops the theoretical framework for analyzing its economic value, explores its theoretical implications, and presents key empirical results that illustrate our findings. Section 4 concludes. All proofs appear in the Appendix, with supplementary analyses relegated to the Online Appendix.

2 Related Literature

Our work contributes to a rapidly growing economic literature on artificial intelligence and recommender systems. It is arguably closest to Calvano et al. (2025), who study how collaborative filtering based recommendations on a platform affect various outcomes including the match between products and consumers, and the additional value generated for consumers. A key difference is that they incorporate how a platform’s recommendations can alter consumer search, and as a result, how competing firms price their products. This means market outcomes such as market concentration and prices, which are their main focus, are endogenous in their setting. We abstract from endogenous market outcomes and focus just on the matching aspect of a recommender system. Methodologically, whereas Calvano et al. (2025) adopts a parametric taste-feature representation, we develop a non-parametric, probability-based model that works directly with the correlation structure between items and user histories. This lets us study theoretically how value scales with data depth as well as breadth, and it delivers conditions for non-diminishing returns and super-additivity from deeper customization.

By decomposing the roles of across-user and within-user learning, our work relates to Hagiu and Wright (2023) who introduce these concepts (and terminology) albeit in a quite different

environment — they study how competitive advantage evolves in the face of each type of learning. Our focus is on a setting where both types of learning co-exist, something Hagiu and Wright (2023) only briefly consider.

More broadly, our research connects to a considerable body of work examining how user data and firms’ algorithms influence market outcomes. Calvano et al. (2020) show in controlled simulations that Q-learning pricing agents can sustain supracompetitive prices via collusive strategies. Experiments on algorithmic pricing show that the design of learning protocol matters for market performance. Asker et al. (2023) demonstrate that different AI learning rules can lead to sharply different market prices. In contrast to these pricing-centric studies, our paper abstracts from prices and competition to isolate the intrinsic welfare value generated by deeper customization from user data. Beyond pricing, papers on learning consumer preferences and rankings study how algorithms infer heterogeneity from feedback signals. e.g., Chen et al. (2018) and Feng et al. (2022). On the other hand, contributions by Biglaiser et al. (2019), de Cornière and Taylor (2024), Fainmesser et al. (2022), and Aridor and Gonçalves (2022) consider equilibrium outcomes when firms use consumer data to enhance their offerings, with Aridor and Gonçalves (2022) explicitly addressing welfare impacts arising from platforms steering consumers toward proprietary products.

Finally, several recent empirical contributions explore the role and economic significance of recommender systems. Aridor et al. (2024) conduct a field experiment on a movie-recommendation platform, demonstrating that recommendations directly affect user beliefs about product quality and drive subsequent information acquisition. Relatedly, studies by Bajari et al. (2019), Peukert et al. (2023), and Yoganarasimhan (2020) quantify economic returns to data primarily through improvements in forecast accuracy. Particularly relevant to the distinction between data breadth and depth, Schäfer and Sapi (2023) investigates the impact of data size on search quality and finds that both across-user and within-user learning significantly enhance search outcomes. Our paper complements these empirical studies by proposing a microfounded model that provides a theoretical basis for their findings. Specifically, our framework shows how learning from user data translates into tangible user surplus, offering a mechanism that links the improvements in predictive accuracy documented in the empirical literature to welfare gains.

3 Value of data

3.1 The model

A platform has $I \geq 1$ items to consider recommending to a target user and chooses an item to recommend to the user who can have either a positive or a negative experience with the item. Each of the I items is called a target item and we denote the set of target items by $\mathcal{I}$. The user has tried $C \geq 0$ items before and we call them conditioning items as the platform can customize its recommendation based on the user's reported experiences with the C items. Although our analysis focuses on the scenario where conditioning items are past experiences, they can more generally represent any observable user features, such as demographics, geographic location, or browsing history. Thus, our model closely mirrors modern recommender systems that leverage diverse historical and contextual data for highly personalized recommendations.

In our framework, a user obtains a payoff of $v_1 > 0$ from a positive experience and $v_0 \leq 0$ from a negative experience with a target item. Similarly, the platform obtains w_1 and w_0 corresponding to the user's positive or negative experiences. If the user does not try the item, both parties receive a payoff of zero. The probability of either experience with a target item is initially unknown to the platform or to the target user. Therefore, the platform aims to estimate these probabilities based on accumulated user data before making a recommendation.

Data

The platform's data is represented by an $M \times N$ matrix, X , containing ratings from $N \geq 1$ users for M items. The first C rows represent ratings on conditioning items, and the remaining I rows correspond to target items. Thus, the set of conditioning items is $\{1, \dots, C\}$ and that of target items is $\mathcal{I} = \{C + 1, \dots, M\}$. Assume each item $i \in \{1, 2, \dots, C, C + 1, \dots, M\}$ has a finite number of possible ratings. For analytical convenience, we assume that the ratings for target items are binary (positive or negative). Some users may not have ratings on some of the M items. If a user has a missing rating of item i , we denote this by $\emptyset$. Thus, in the data X , a typical element $x_{ij} \in \{\emptyset, 0, 1\}$ records user j 's rating of item i . Without loss of generality, we reserve the first column of X to denote the target user's ratings. Naturally, we have $x_{i1} = \emptyset$, $\forall i \in \mathcal{I}$ given the target user's ratings of target items must be predicted.

Data X is said to be *complete* if we have $x_{ij} \neq \emptyset$, $\forall i, \forall j$ s.t. $i \leq C$ or $j \neq 1$. In other words, the platform has complete data if the only missing ratings are the target user's rating of the target

items. On the other hand, if the condition for complete data is not satisfied, the data is said to be *partial*. To reduce the notational burden, we focus here on the case in which the existing user data is complete, and we explain the extension to partial data in the Online Appendix C and D.³

Statistical model

The platform estimates the likelihood of a target user having a positive user experience with each target item based on historical data. Formally, we define outcomes for each combination of conditioning items and a target item. An outcome is a vector of length $C+1$ that records the user’s ratings on C conditioning items and on a specific target item, i . The set of all possible outcomes corresponding to target item i is denoted by R^i .

Although the interpretation of an outcome vector depends on the target item i , the set of possible vectors is structurally identical across all target items. For notational simplicity, we can therefore refer to this common set of outcomes as simply R , letting the relevant target item be clear from the context. For example, consider a platform that conditions on the reported (binary) rating of item 1 in making predictions about item 2 and item 3, i.e., $C = 1$ and $I = 2$. The main example we use throughout the article is represented in Table 1. The set of outcomes R in this case is given by $R = \{(1, 1), (1, 0), (0, 1), (0, 0)\}$.

	user 1	user 2	user 3	user 4	user 5
Rating on item 1	1	1	0	1	1
Rating on item 2	$\emptyset$	1	0	1	0
Rating on item 3	$\emptyset$	0	1	1	0

Table 1: Main example

The target user’s ratings on the C conditioning items can be represented by r' , which is a subvector of $r \in R$ whose length is C . We refer to r' as a history. The set of all such histories of length C is denoted by R' . When $C = 0$, we take $R' = \{\emptyset\}$. By slightly abusing notation, we denote the vector of dimension $C+1$ created by adding 1 to the end of r' by $(r', 1)$. Similarly, $(r', 0)$ denotes the vector created by appending 0 to the end of r' . Any $r \in R$ one-to-one corresponds to either $(r', 1)$ or $(r', 0)$ for some r' . For example, $r = (1, 0)$ is equivalent to $(r', 0)$ for $r' = (1)$.

For each target item i , the outcome is governed by an unobserved probability vector $p^i =$

³We focus on the complete data case as it serves as a clear theoretical benchmark, allowing us to derive a closed-form solution for the value of data. The framework extends to partial data, although the analysis becomes more complex as it requires integrating over the distributions of missing ratings. The extensions in the Online Appendix detail how our recommender system makes predictions in this more complex setting and provides the corresponding welfare analysis.

$(p_r^i)_{r \in R}$, where p_r^i denotes the joint probability of a specific outcome r , representing the vector of ratings across the C conditioning items and target item i . We assume $p_r^i > 0$, for all i and r . This probability vector p^i is referred to as a *correlation structure* as it reveals how the user experience with the target item i and conditioning items are correlated each other. Naturally, these probabilities satisfy $\sum_{r \in R} p_r^i = 1$. Additionally, the true probability associated with a given user history r' is denoted by $p_{r'}^i$, which equals the sum of probabilities for the corresponding positive and negative outcomes: $p_{r'}^i = p_{(r',1)}^i + p_{(r',0)}^i$.

Recommender system

Given the correlation structure, the platform estimates the probability $z_i(r')$ that the target user will have a positive experience with target item i , conditional on their history r' .

$$z_i(r') = \frac{p_{(r',1)}^i}{p_{r'}^i}. \quad (1)$$

The fact that the same probability $z_i(r')$ applies to all the users with the same history r' does not imply that these users have identical preferences. Instead, the platform treats these users equivalently due to the limitations imposed by the degree of customization it uses. As richer and more detailed user data become available, enabling a higher degree of customization, the platform can better differentiate between users and provide increasingly precise, personalized recommendations.

Let $\hat{z}_i^N(r')$ be a pointwise estimator of $z_i(r')$ for a target user whose history is r' , based on data from $N - 1$ previous users. Additionally, the platform sets a threshold $\tau(r') \in [0, 1]$ that applies to users with history r' . If the estimated probability $\hat{z}_i^N(r')$ is below this threshold, the item i is deemed not suitable for the user with history r' , and therefore, the system does not recommend this item to the user. Formally, we define a recommender system as follows:

Definition 1 *A recommender system is a collection of functions*

$$\left\{ \left\{ \hat{z}_i^N(r') \right\}_{i \in \mathcal{I}}, \tau(r') \right\}_{r' \in R}.$$

Given data X and the target user's history r' , target item i is recommended to the target user if and only if $\hat{z}_i^N(r')(X) \geq \max_{j \in \mathcal{I}} \{ \hat{z}_j^N(r')(X), \tau(r') \}$.

Note that the value derived from data is inherently linked to the specific recommender system

employed. Different estimators within recommender systems can lead to varied inferences, user decisions, and subsequent learnings from identical data. Consequently, no single recommender system can capture the universal value that data might offer. Nevertheless, as the number of previous users N becomes very large, we can evaluate the value offered by a broad class of recommender systems termed *consistent recommender systems*. A recommender system is said to be consistent if, for each r' , the estimator $\hat{z}_i^N(r')$ is statistically consistent for $z_i(r')$. Formally, a consistent recommender system can be defined as follows.

Definition 2 *A recommender system is said to be consistent if $\hat{z}_i^N(r')$ converges in probability to $z_i(r') \in [0, 1]$ for each $r' \in R'$.*

The existence of consistent recommender systems is guaranteed.⁴ Although consistency is an asymptotic property, it represents a fundamental requirement for reliable recommendation performance, as violating it would imply persistent prediction errors. In the following sections, we focus exclusively on consistent recommender systems, examining both the asymptotic value that such systems create for users and the incremental benefit arising from deeper customization. While user value from any consistent recommender system approaches its asymptotic limit with more data, finite-sample values are dependent on the specific system chosen. We detail the finite-sample implications for user value in Section 3.5, where we impose an additional unbiasedness condition on the estimator.

User behavior

We model the target user’s decision with a cutoff rule. Suppose, hypothetically, the correlation structure is fully known to the platform. Because the user receives $v_1 > 0$ from a positive experience and $v_0 \leq 0$ from a negative experience, the user-optimal threshold is $\tau^u = \frac{-v_0}{v_1 - v_0}$. It is therefore optimal for the target user if platform recommends the item with the highest probability of liking only when that probability exceeds τ^u .

On the other hand, the platform receives w_1 and w_0 from the user’s positive and negative experiences (with $w_1 > w_0$), which yields the platform-optimal threshold $\tau^p = \max\{0, -w_0/(w_1 - w_0)\}$.

Prior to any recommendation, for a target user with history r' , let the user’s belief about $z_j(r')$ for each target item j be represented by some distribution G_j . Upon receiving a recommendation

⁴We explicitly construct a Bayesian recommender system satisfying this consistency criterion in Online Appendix A.

of item i , the user updates to the implied posterior.⁵ Using this posterior, define $\underline{\tau}$ as the obedience lower bound: if the platform sets its per-history threshold below $\underline{\tau}$, a user declines to try the item even when it is recommended. For example, consider the simplest case with a single target item, namely item i , where $z_i(r')$ is uniformly distributed over $[0, 1]$. With $(v_1, v_0) = (1, -2)$, if the platform sets $\tau = 1/4$, then posterior mean success probability conditional on a recommendation is $(\tau + 1)/2 = 5/8$, so the user's expected utility from trying is $-1/8 < 0$. Equivalently, the obedience lower bound here is $\underline{\tau} = 1/3$. Thus, when $\tau < \underline{\tau}$, the user does not follow the recommendation.

We discard the trivial case $\tau^p < \underline{\tau}$, in which users do not follow recommendations and the value of the system is trivially zero. In the relevant case that we will focus on, $\tau^p \geq \underline{\tau}$, and users optimally follow recommendations in the asymptotic regime.⁶

3.2 Value of data and its decomposition

We characterize and decompose the value created by personalized recommender systems. To isolate the gains from personalization, we contrast these systems with a simpler, non-personalized benchmark, which we refer to a generic recommender system. This system relies solely on aggregate user data without personalizing recommendations based on individual user histories. Formally, the generic recommender system selects recommendations based on an estimated probability, $\hat{z}_i^G$ of the true probability $z_i^G = \sum_{r' \in R'} p_{(r', 1)}^i$, that an average user will positively experience an item, independent of any specific history.

Definition 3 *A generic recommender system is a recommender system $\{\{\hat{z}_i^G\}_{i \in \mathcal{I}}, \tau^G\}$ such that $\hat{z}_i^G$ is a consistent estimator for z_i^G .*

The generic system corresponds to early popularity or average-based rules that rely solely on across-user data and serve as a transparent benchmark for isolating the incremental value of conditioning. In contrast, the type of personalized recommender system we study represents a recommendation framework that leverages both across-user and within-user learning, providing a higher degree of customization. We analyze the ex-ante value created by these more sophisticated personalized systems, averaging their expected benefits across all possible data realizations weighted by true probabilities.

⁵E.g., if $\{z_j(r')\}_{j \in \mathcal{J}}$ is i.i.d., the recommended item would be the largest order statistic among those above the platform's threshold.

⁶If instead we assume user behavior such that users always follow the platform's recommendation regardless of τ^p , then we can drop the obedience lower bound condition we impose, that $\tau^p \geq \underline{\tau}$.

Proposition 1 quantifies this value, separating it into two distinct components: (1) generic across-user learning and (2) additional gains from personalized recommendations. All the proofs are given in the Appendix.

Proposition 1 1. For the user-optimal threshold, in any consistent recommender system, as $N \rightarrow \infty$, the expected utility to a user converges to

$$\sum_{r' \in R'} \max_{i \in \mathcal{I}} \left\{ v_1 p_{(r',1)}^i + v_0 p_{(r',0)}^i, 0 \right\}.$$

2. For the user-optimal threshold, the expected utility created from a generic recommendation in the limit as $N \rightarrow \infty$ is

$$\max_{i \in \mathcal{I}} \left\{ \sum_{r' \in R'} \left(v_1 p_{(r',1)}^i + v_0 p_{(r',0)}^i \right), 0 \right\}.$$

Customization can be regarded as a finer segmentation of users according to observable features. To see why a finer segmentation necessarily leads to a higher user welfare under the user-optimal threshold level, consider a history r' group of users whose average probability of liking a (unique) target item is given by z . Under the user-optimal threshold, the average utility of the history r' users is $\max\{v_1 z + v_0(1 - z), 0\}$. If a further segmentation of r' users into subgroups r'_a and r'_b is available to the system, the system makes two separate recommendations according to the two new average probabilities, z_a and z_b , of the two subgroups. Let p_a portion of r' users is assigned to r'_a subgroup. The expected utility of users under segmentation is then $p_a \max\{v_1 z_a + v_0(1 - z_a), 0\} + (1 - p_a) \max\{v_1 z_b + v_0(1 - z_b), 0\}$, where $z = p_a z_a + (1 - p_a) z_b$. The customization reveals the subgroup that is on-average better off trying (or not trying) the item, which in turn increases the overall expected utility. In Figure 1, the expected utility without customization is zero while with customization it is $(1 - p_a)(v_1 z_b + v_0(1 - z_b))$.

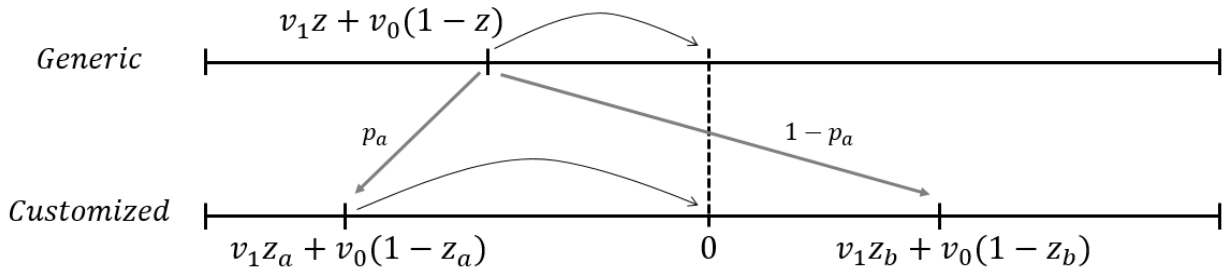


Figure 1: Customization under the user-optimal threshold

The difference between the two expressions in the proposition is due to the different degree of customization each recommender system offers. A consistent recommender system makes history-dependent recommendations by selecting the best item for each group of history r' users. Instead, a generic recommender system only selects the item on average best for every user. Therefore, the value from generic recommendations in Proposition 1 can be attributed to pure across-user learning. On the other hand, the difference between the value from consistent recommendations and the value from generic recommendations can be attributed to the customization benefits that the consistent recommendation provides. Note that, under the user-optimal threshold, the expected utility of the target user created from a consistent recommender system is always positive, and it is at least as great as the expected utility generated from the generic recommendation, which is also always positive. That is, both the pure across-user learning and the customization component always add value to users when the recommender system is user-centric.

Note that, although the two systems differ in their degrees of customizations, they share the same value creation process: a recommender system selects the most suitable item and at the same time, screens items not suitable for the user. Obviously, the more items in the target item pool, the more value is created from both of the systems. On the other hand, the benefit from screening is maximized when the threshold is properly chosen for the user.

Example

To illustrate and better understand how a recommender system adds value through learning and customization, consider our main example in Table 1 and assume $(v_1, v_0) = (1, -1)$. In the example, the platform makes a recommendation between item 2 and item 3, or none based on the target user's rating on item 1. According to Proposition 1, when the threshold level is τ^u , the expected utility of a user under a consistent mechanism in the limit as $N \rightarrow \infty$ is given by

$$\underbrace{\max \left\{ p_{(1,1)}^2 - p_{(1,0)}^2, p_{(1,1)}^3 - p_{(1,0)}^3, 0 \right\}}_{\text{Recommendation for history 1 user}} + \underbrace{\max \left\{ p_{(0,1)}^2 - p_{(0,0)}^2, p_{(0,1)}^3 - p_{(0,0)}^3, 0 \right\}}_{\text{Recommendation for history 0 user}}. \quad (2)$$

Depending on the user's history, different items can be recommended. For example, if $p_{(1,1)}^2 - p_{(1,0)}^2 > \max\{p_{(1,1)}^3 - p_{(1,0)}^3, 0\}$ and $0 > \max\{p_{(0,1)}^2 - p_{(0,0)}^2, p_{(0,1)}^3 - p_{(0,0)}^3\}$, item 2 is recommended to users who had a positive experience with item 1 but no item is recommended to users with a negative experience with item 1. On the other hand, under a generic recommender system, which only makes

use of across-user learning, users receive the same recommendation regardless of their history. The value from a generic recommendation with the same threshold being applied is given by

$$\max \left\{ \underbrace{p_{(1,1)}^2 + p_{(0,1)}^2 - p_{(1,0)}^2 - p_{(0,0)}^2, p_{(1,1)}^3 + p_{(0,1)}^3 - p_{(1,0)}^3 - p_{(0,0)}^3}_\text{No customization in recommendations}, 0 \right\}. \quad (3)$$

It is clear that the value from customization is always weakly positive. On the other hand, it is strictly positive *if and only if* users who had different experiences with item 1 receive different recommendations.

Asymptotic lower bound of prediction error

Our ex-ante analytical framework and explicitly introduced theoretical model of correlation structure enable us to derive a concise closed-form representation of the achievable asymptotic lower bound on prediction error for consistent recommender systems. It can also be shown that this lower bound of the error is weakly decreasing in the degree of customization. For a consistent recommender system $\{\{\hat{z}_i^N(r')\}_{i \in \mathcal{I}}, \tau(r')\}_{r' \in R'}$, let ϵ_i^m be the number of wrong predictions, both the false positives (negative experience from a recommended item) and false negatives (unrealized positive experience when no item is recommended), out of the total of m predictions made about the target items. The prediction error is defined to be $\epsilon_i^m(X)/m$. The lower bound of this asymptotic error can be achieved when the threshold is user-optimal, and it is characterized as follows for the general $I \geq 1$ cases.

Corollary 1 (Corollary to Proposition 1) *For $C \geq 0$ and $I \geq 1$,*

1. *The asymptotic lower bound of prediction error in a consistent recommender system is*

$$L_C = \sum_{r' \in R'} \min_{i \in \mathcal{I}} \min \left\{ p_{(r',1)}^i, p_{(r',0)}^i \right\}. \quad (4)$$

2. *L_C weakly decreases in C .*

Thus, personalized recommender systems not only significantly boost user welfare through customization but also consistently achieve higher prediction accuracy compared to traditional generic systems.

3.2.1 Empirical illustration: Decomposing the value of data

To provide an empirical illustration of our theoretical decomposition above, we conduct a simulation study using the Jester dataset, which contains over four million anonymous joke ratings from 73,421 users. We convert the continuous ratings into a binary format (positive/negative) to fit our framework. We estimate correlation structure with a Bayesian Dirichlet-Multinomial model with a uniform prior.⁷ Using the estimates, we run counterfactual simulations to quantify the value created by the three key functions of a personalized recommender system: customization, selection, and screening.

To do so, we measure user utility under different system configurations. We compare a generic recommender system ($C = 0$) with a customized recommender system that conditions on user history ($C = 9$). We also vary the number of target items available for selection (I) and the screening threshold (τ). Table 2 summarizes the results. Full details of the dataset and our empirical methodology are available in Online Appendix B.

		<i>Average Utility</i>	<i>Standard Error</i>	<i>Min / Max</i>
$I = 1, \tau = 0$	<i>Generic RS ($C = 0$)</i>	0.148	(0.0077)	-0.567/0.616
	<i>Customized RS ($C = 9$)</i>	0.148	(0.0077)	-0.567/0.616
$I = 1, \tau = 1/2$ (<i>scr only</i>)	<i>Generic RS</i>	0.193	(0.0057)	-0.159/0.616
	<i>Customized RS</i>	0.258	(0.0049)	-0.051/0.622
$I = 3, \tau = 0$ (<i>sel only</i>)	<i>Generic RS</i>	0.342	(0.0048)	-0.284/0.617
	<i>Customized RS</i>	0.357	(0.0049)	-0.295/0.651
$I = 3, \tau = 1/2$ (<i>scr and sel</i>)	<i>Generic RS</i>	0.344	(0.0046)	-0.133/0.617
	<i>Customized RS</i>	0.371	(0.0042)	-0.013/0.634

Table 2: User surplus from recommender systems
scr: screening, *sel*: selection

Table 2 presents the estimated average utility that users receive from recommender systems with different parameter settings. Without selection and screening ($I = 1, \tau = 0$), the average utility is 0.148, which corresponds to the baseline where users try randomly selected items. On top of this, screening ($\tau = 1/2$) adds 0.045 to average utility with a generic recommender system, while customization ($C = 9$) adds a further 0.065, suggesting significant value from within-user learning. The value of customization is less pronounced when selection is wider (comparing $I = 3$

⁷The full detail of this model and its properties are relegated to Online Appendix A.

to $I = 1$ with $\tau = 0$), which is consistent with our theoretical predictions for items that are popular across many user types. When all three functions coexist ($C = 9, I = 3, \tau = 1/2$), the personalized recommender system adds 0.223 more value to users compared to the situation without any recommender system.

3.3 Harmful customization

While Proposition 1 shows customization always (weakly) benefits users under the user-optimal threshold, it can also be shown that such a user-optimal threshold is the only threshold level for which customization benefits the user regardless of the correlation structure and user history.

Proposition 2 *For any $\tau(r') \neq \tau^u$ and for any $I \geq 1$, there exists a collection of I correlation structures $q = \{q^i\}_{i \in \mathcal{I}}$ such that the history r' target user is strictly worse-off from customization when the true correlation structures are q .*

The customization in predictions is advantageous to the target user only when the threshold level is properly chosen. Put differently, when $\tau^u \neq \tau^p$, either the user or the platform will be strictly worse off under some correlation structures when customization is used. Thus, customization has the scope to hurt users if the platforms' interests cause its threshold level to diverge from the user-centric threshold. The proof is by construction: We identify q for each $\tau(r')$. Conversely, for each realized correlation structure p , we can also find $\tau(r') \neq \tau^u$ under which history r' user strictly worse off from customization. The exact value of such threshold levels is presented in Online Appendix E.

3.4 Marginal value of customization

As highlighted in the previous section, recommender systems learn about target items not only from the ratings left on the target items by other users, but also from the ratings left by the target user on other items so as to better customize the recommendation. Hence, the quality of a recommendation and the resulting user surplus depend also on the degree of customization the platform provides. In this section, we study how customization affects user value through the recommender system. Specifically, we are interested in the effect of changes in the degree of customization to the value generated by each target item under a recommender system. To do so, we isolate the target item by considering the case $I = 1$, and investigate the sources behind the value creation process of a recommender system.

Suppose that a consistent recommender system takes one more conditioning item into account in making a recommendation on the target item. One can equally think of a situation where the target user tries an item recommended by the platform, and the platform now can provide a customized recommendation that conditions also on the rating of the newly tried item. We denote the new item and the target item respectively by item $C + 1$ and item $C + 2$. As the recommender system can condition its predictions on one more item, the predictive accuracy of the target item improves. However, the improvement in accuracy does not necessarily lead to higher user welfare because of the misalignment of interests between the user and the platform which is captured by the threshold level. Here, we study this marginal benefit or harm of customization. Because there is only one target item, we save notation by using p for p^{C+2} , the correlation structure associated with the $C + 1$ conditioning items and the target item.

Definition 4 For $r' \in R'$ over items $\{1, \dots, C\}$, item $C + 1$ and item $C + 2$ are positively correlated conditional on r' if $v_1 p_{(r',1,1)} + v_0 p_{(r',1,0)} \geq 0$ and $v_1 p_{(r',0,1)} + v_0 p_{(r',0,0)} \leq 0$, with at least one inequality holding strictly. If both inequalities are strict, we say that they are strictly positively correlated conditional on r' .

When the target item and item $C + 1$ are positively correlated conditional on r' , users whose history is r' over the other C conditioning items are more likely to have net positive utility from the target item if they liked item $C + 1$. On the other hand, if they did not like the item, it is more likely that they have net negative utility from the target item. When $(v_1, v_0) = (1, -1)$, the condition is satisfied if we have $p_{(r',1,1)} \geq p_{(r',1,0)}$ and $p_{(r',0,1)} \leq p_{(r',0,0)}$ with at least one inequality holding strictly. In a parallel way we can define a (strict) negative correlation between item $C + 1$ and the target item. If item $C + 1$ and the target item are positively or negatively correlated conditional on $r' \in R'$, they are said to be correlated conditional on r' . For strict inequalities, they are said to be strictly correlated conditional on r' .

We focus on a universal threshold level $\tau = \tau(r')$, $\forall r'$. For each r' , there are two trivial cases in which an extra degree in customization yields only zero marginal customization effect: when the threshold level is too high or too low as stated below.

$$\tau \leq \min \left\{ \frac{p_{(r',1,1)}}{p_{(r',1)}}, \frac{p_{(r',0,1)}}{p_{(r',0)}} \right\} \text{ or } \tau \geq \max \left\{ \frac{p_{(r',1,1)}}{p_{(r',1)}}, \frac{p_{(r',0,1)}}{p_{(r',0)}} \right\}. \quad (5)$$

If the former is the case, the item is recommended to the user regardless of the history and the

introduction of the extra customization. This follows from the following equality.

$$\frac{p(r',1,1) + p(r',0,1)}{p_{r'}} = \frac{p(r',1)}{p_{r'}} \frac{p(r',1,1)}{p(r',1)} + \frac{p(r',0)}{p_{r'}} \frac{p(r',0,1)}{p(r',0)}.$$

Similarly, in the latter case, the item is not recommended regardless of the introduction of the extra degree in customization. As a result, the extra customization cannot affect user welfare if (5) holds. In what follows, we focus on the remaining case, i.e.,

$$\min \left\{ \frac{p(r',1,1)}{p(r',1)}, \frac{p(r',0,1)}{p(r',0)} \right\} < \tau < \max \left\{ \frac{p(r',1,1)}{p(r',1)}, \frac{p(r',0,1)}{p(r',0)} \right\}.$$

In the next proposition, we find that a strict correlation between the newly added item and the target item is a necessary and sufficient condition under which an extra degree in customization is (on average) beneficial to users regardless of the misalignment of interests between the platform and users.

Proposition 3 *The marginal customization strictly benefits r' user for any threshold level if and only if item $C + 1$ and the target item (item $C + 2$) are strictly correlated conditional on r' .*

When strict correlation exists, the system can use the item to segment r' users into subgroups based on their experiences with item $C + 1$. This segmentation enables the system to accurately identify and screen out subgroups unlikely to benefit from the target item, enhancing overall user welfare.

An important insight is that marginal value from additional customization does not always diminish; it can increase or decrease. Moreover, the marginal value conditioning on multiple items simultaneously can even surpass the sum of individual values from conditioning on the items separately, demonstrating super-additivity in the value of customization.

To see this, let $(v_1, v_0) = (1, -1)$ and consider a platform that has three items under the user-optimal threshold level. The target item is fixed at item 3 and we compare user surplus under two cases depending on whether we condition on item 1 first or item 2 first.

The correlation structure p is given as follows:⁸

$$p = \left(\frac{2}{20}, \frac{1}{20}, \frac{1}{20}, \frac{7}{20}, \frac{1}{20}, \frac{2}{20}, \frac{5}{20}, \frac{1}{20} \right).$$

⁸Here $p = (p_{(1,1,1)}, p_{(1,1,0)}, p_{(1,0,1)}, p_{(1,0,0)}, p_{(0,1,1)}, p_{(0,1,0)}, p_{(0,0,1)}, p_{(0,0,0)})$.

We can find the expected user utility under different degrees of customizations. First, when the platform initially does not customize its prediction, the expected utility from trying the target item is zero as the item is not recommended to try:

$$p_{(1,1,1)} + p_{(1,0,1)} + p_{(0,1,1)} + p_{(0,0,1)} = \frac{9}{20} < \frac{1}{2}.$$

On the other hand, if the platform conditions on both items in making a prediction about the target item, we can check that the expected utility to a user is $\frac{1}{4}$. In this case, users of history $(1, 1)$ and $(0, 0)$ try the item. Now, when the platform conditions only on item 1, the expected utility is

$$\max\{p_{(1,1,1)} + p_{(1,0,1)} - p_{(1,1,0)} - p_{(1,0,0)}, 0\} + \max\{p_{(0,1,1)} + p_{(0,0,1)} - p_{(0,1,0)} - p_{(0,0,0)}, 0\} = \frac{3}{20}.$$

Thus, in this case, the marginal value from customization is diminishing. It is $\frac{3}{20}$ from degree zero to degree one, and then $\frac{2}{20}$ from degree one to degree two.

On the other hand, if we customize the recommendation conditioning on item 2 first, the resulting expected utility is given as

$$\max\{p_{(1,1,1)} + p_{(0,1,1)} - p_{(1,1,0)} - p_{(0,1,0)}, 0\} + \max\{p_{(1,0,1)} + p_{(0,0,1)} - p_{(1,0,0)} - p_{(0,0,0)}, 0\} = 0.$$

Thus, in this case, the marginal value from customization is increasing: it is zero from degree zero to degree one, and then $\frac{1}{4}$ from degree one to degree two.

Another notable observation from this example is the super-additivity of values generated by customization. Specifically, the value created by simultaneously conditioning recommendations on multiple items can exceed the sum of the values generated by individually conditioning on each item separately. In the example, the value from conditioning on both item 1 and 2 simultaneously is $\frac{1}{4}$, while conditioning individually on either item 1 or item 2 yields values of $\frac{3}{20}$ or 0, respectively. The combined individual values $\frac{3}{20}$ is thus strictly lower than the value from combined conditioning, which is $\frac{1}{4}$.

At the user-optimal threshold, super-additivity occurs whenever a particular user history generates positive expected utility if considered alone, but this value is canceled out when aggregated with other histories under a lower degree of customization. In the example, consider the history $r' = (1, 1)$. Individually, this segment's expected utility is positive and fully realized when customization involves both items. However, when this segment is combined with the $(1, 0)$ history

(by conditioning on item 1 only), the aggregate expected utility becomes negative, completely masking the existence of valuable information. This illustrates that more precise segmentation can unlock substantial value that is otherwise hidden.

This intuition can be formalized in the following proposition, which provides a necessary and a sufficient condition for super-additivity. Consider two conditioning items, 1 and 2, and the user-optimal threshold level.

Proposition 4 *The value of customization exhibits super-additivity only if there exists i, j such that*

$$\begin{aligned} (v_1 p_{(i,1,1)} + v_0 p_{(i,1,0)}) (v_1 p_{(i,0,1)} + v_0 p_{(i,0,0)}) &< 0 \text{ and} \\ (v_1 p_{(1,j,1)} + v_0 p_{(1,j,0)}) (v_1 p_{(0,j,1)} + v_0 p_{(0,j,0)}) &< 0. \end{aligned}$$

Furthermore, super-additivity is guaranteed if i, j satisfies $v_1 p_{(ij1)} + v_0 p_{(ij0)} > 0$,

$$\begin{aligned} v_1 (p_{(i,1,1)} + p_{(i,0,1)}) + v_0 (p_{(i,1,0)} + p_{(i,0,0)}) &\leq 0 \text{ and} \\ v_1 (p_{(1,j,1)} + p_{(0,j,1)}) + v_0 (p_{(1,j,0)} + p_{(0,j,0)}) &\leq 0. \end{aligned}$$

This proposition formalizes the mechanism observed in the numerical example. The first part of the proposition establishes that a fundamental level of informational relevance—strict correlation—is required for super-additivity to be possible. The second part identifies a specific data structure—the existence of valuable information that is only achievable when multiple conditions are met—that is sufficient to guarantee a super-additive outcome. Recognizing that the value of data can scale in these complex, synergistic ways, rather than only with simple diminishing returns, is crucial for understanding the full potential of modern AI-driven recommender systems, which excel at uncovering precisely these types of multi-feature interaction effects.

3.5 Value of an additional user: the breadth and depth of data

In this section, we study the value from finite data to study how a marginal data point (i.e. on another user) adds value to other users. Specifically, we analyze how the value of this additional data scales along the two data dimensions: breadth and depth. Contrary to the asymptotic value we have focused in the previous section, the value from finite data varies depending on the estimator that a recommender system adopts as different consistent estimators can induce different predictions

from finite data. Thus, we impose a minimal structure that a consistent estimator should satisfy, and study the value from marginal data and how it behaves as the system accumulates more data as it adds users.

To keep the analysis as simple as possible, we consider the case when $(v_1, v_0) = (1, -1)$ and look at the case when there is only one target item and the history is given as r' over C conditioning items. The platform learns the correlation structure over $C + 1$ items from the previous N history r' users and the target user's experience with the conditioning items.

In this case, the data X from N users can be summarized by $y = (y_{(r',1)}, y_{(r',0)})$, where y_r records the occurrences of outcome $r \in R = \{(r', 1), (r', 0)\}$ that appear in X . The true but unknown probability of a positive experience with the target item is $\frac{p_{(r',1)}}{p_{(r',1)} + p_{(r',0)}}$, which we denote by z . We focus on the recommender systems that satisfy the following condition.

Definition 5 *A recommender system is said to be unbiased if the target item is recommended to history r' target user if and only if $y_{(r',1)} \geq y_{(r',0)}$.*

In any unbiased recommender system, the target item is recommended to try if and only if the target user is more likely to have a positive experience with the target item in the sense that the previous users who share the same history with the target item have reported more positive ratings than negative ratings.

As will be apparent from the proof, to avoid a mathematical complexity that is involved with binomial probability and the ceiling function, assume that N is only an odd number by considering the case of $y_{(r',1)} + y_{(r',0)} = 2m - 1$, $m \in \mathbb{N}$. The same results apply to the even-number cases.

Let $v(m|r', p)$ denote the expected value of the target user in terms of the number of previous users $N = 2m - 1$ and the correlation structure is given as p . Formally, we have

$$\begin{aligned} v(m|r', p) &= \frac{p_{(r',1)}}{p_{(r',1)} + p_{(r',0)}} P[y_{(r',1)} \geq y_{(r',0)}] - \left(1 - \frac{p_{(r',1)}}{p_{(r',1)} + p_{(r',0)}}\right) P[y_{(r',1)} \geq y_{(r',0)}] \\ &= zP[y_{(r',1)} \geq y_{(r',0)}] - (1 - z)P[y_{(r',1)} \geq y_{(r',0)}]. \end{aligned}$$

Note that this value can take a negative value. For example, for any $z \in (0, \frac{1}{2})$, it is possible when the occurrences of $y_{(r',1)}$ exceed that of $y_{(r',0)}$. Therefore, the possibility of a wrong recommendation always exists even though the platform is committed to giving the recommendation that it expects to be best for the target user. Although the addition of a data point from an additional user always

leads to a more informative data structure, it can be shown that the incremental value diminishes. The following results characterize the nature of data under the recommender system.

Proposition 5 *Let $v(m|r', p)$ denote the value to the target user and $\Delta v(m|r', p)$ be the marginal value that the target user contributes to the next user.*

1. *The value is positive and increases in the number of previous users. i.e., $v(m|r', p)$ increases in m , $\forall p$.*
2. *The marginal value diminishes. i.e., $\Delta v(m|r', p)$ decreases in m , $\forall p$.*
3. *For $m = 1$, the expected value decreases in C .*

In words, the value to users increases in the number of available data points (1), but the value increment diminishes (2). Finally, in contrast to our observation in Proposition 1 where we show the value increases in the degree of customization when we have enough data points, it decreases in the degree of customization when the number of data points is not large enough to accommodate the degree of customization (3). Consider a correlation structure q over the initial C conditioning items, a new conditioning item and the target item. Thus, it satisfies $q_{(r', 1, 1)} + q_{(r', 0, 1)} = p_{(r', 1)}$, $\forall r' \in R'$.

The reason behind the diminishing return to a data point has a clear connection to how an unbiased prediction is formulated. An unbiased prediction always involves a tradeoff between using prior information and new data. As we increase the data size, the relative contribution of each data point to the posterior becomes smaller. Thus, any particular data point that has the same history as the target user's becomes less influential in forming the expected utility of the target user. On the other hand, when there is not enough number of data, an excessive customization hurts the prediction precision and decrease the value to users.

Taken together, our findings imply a distinctive characterization of the learning curve associated with recommender systems. As more data accumulates, both dimensions, the breadth and the depth, expand concurrently. This combined expansion generates cascading waves of diminishing returns. Initially, increasing the breadth of data improves user welfare, though at a diminishing rate. Once sufficient breadth is reached to make a higher degree of customization feasible, further benefits arise by increasing that depth, thereby initiating a new diminishing return curve but with a higher potential utility, as established in Proposition 1. Importantly, as illustrated by our examples, this magnitude of the incremental benefit from a higher degree in customization

does not necessarily diminish and can, in fact, sometimes increase, showcasing super-additivity. Consequently, the overall learning curve formed by simultaneous expansions along both dimensions resembles the upper envelope of these diminishing return curves.

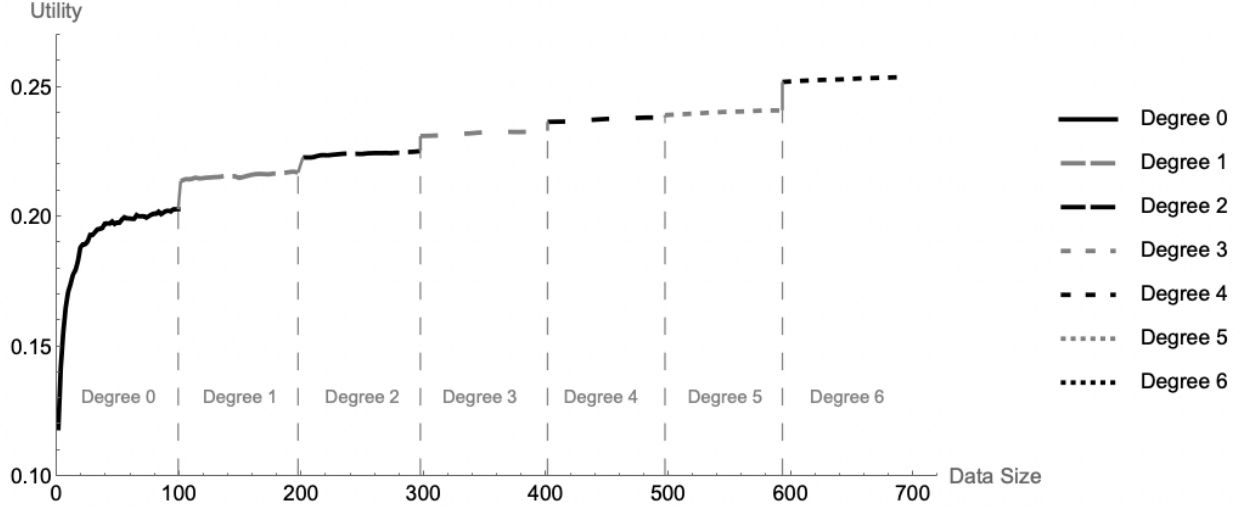


Figure 2: Cascading learning curves

To empirically illustrate this theoretical characterization of the learning curve, we conduct a further simulation exercise using the same Jester dataset introduced in Section 3.2. In this exercise, we progressively expand the breadth and depth of data: we expand the breadth of data one user at a time, and for every 100 data-points from users, we also increase the depth by one. The resulting learning curve, plotted in Figure 2, closely aligns with our theory. It clearly showcases an upper envelope formed by cascading waves of diminishing returns as data breadth and customization depth expand concurrently. Further details of this simulation are presented in Online Appendix B.

Lastly, we assess whether the two dimensions of data, breadth and depth, are complements or substitutes in creating user value. We use our benchmark simulation setting with $(I, \tau) = (1, 1/2)$ to run 1,000 simulations where we adjust the degree of customization C and the number of training users N . Let $V(C, N)$ be the average utility from a recommender system that learns from N previous user's data about $C + 1$ number of items. To investigate how the two relevant dimensions interact, we measure the discrete cross-partial of V with respect to C and N :

$$\Delta_{C,N} = (V(C + \delta_C, N + \delta_N) - V(C + \delta_C, N)) - (V(C, N + \delta_N) - V(C, N)).$$

Here, δ_C and δ_N are increments in the degree of customization and in the number of previous users. Depending on whether this cross partial difference is positive or negative, it can be evi-

dence that the two dimensions are complements or substitutes. For instance, $\Delta_{C,N} > 0$ indicates complementarity, meaning that a higher degree of customization makes additional user data more valuable. In our analysis, we take $\delta_C = 1$ and $\delta_N = 500$. Figure 3 plots the value of this measure.

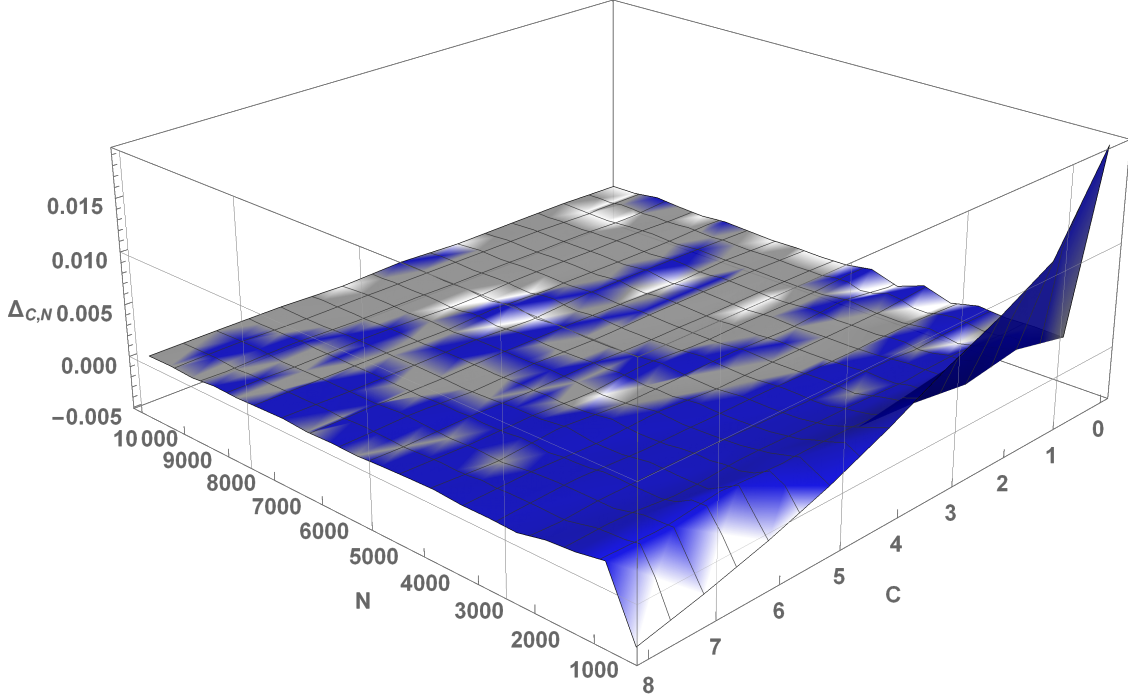


Figure 3: Complementarity between the two dimensions

In Figure 3, the area depicted in dark blue represents combinations of (C, N) which generate positive values of the cross-partial derivative (above 0.0005), so in which the two dimensions are complements. On the other hand, the area in white represents combinations of (C, N) which generate negative values of the cross-partial derivative (below -0.0005), so in which the two dimensions are substitutes. Lastly, the area with light gray represents combinations of (C, N) around zero, within the error bound of $(-0.0005, 0.0005)$.

As shown in the figure, the two dimensions exhibit strong complementarity when both the degree of customization and the number of users are low (the peak in the dark blue area). This finding provides a direct rationale for the cascading effect illustrated in Figure 2. The simulation that generates the cascading curve, covering data sizes up to 700 users and customization degrees up to 6, operates within this region of strong complementarity. Because data breadth and depth are complements in this range, an increase in the number of users enhances the value of deeper customization. However, this complementarity diminishes as the user base grows, and the two dimensions can even become substitutes when customization is high but the number of users is

very small. For instance, at $N = 0$, the cross marginal effects are $\Delta_{6,0} = -0.0011$, $\Delta_{7,0} = -0.0026$, and $\Delta_{8,0} = -0.0035$ ⁹. This reflects the lack of degrees of freedom when N is small and C is large. This substitutability also shapes the structure of the cascading learning curve. It explains why a higher degree of customization is only viable after a sufficient number of users have been accumulated. Increasing C too early, when N is still low, would actually harm prediction accuracy and lead to lower expected utility. Finally, we note that regardless of the degree of customization, the cross-effects largely disappear once there are enough users to learn from. In our exercise, once $N \geq 5000$, the magnitude of any cross marginal effect is less than 0.0005.

4 Conclusion

In this paper, we propose a tractable economic model to understand and quantify the value created by recommender systems, highlighting the critical role of personalization. Our framework identifies three fundamental components underlying value creation: customization, which segments users based on their observable features (history); selection, which then identifies the best potential item for each group; and screening, which finally determines if that best item is good enough to be recommended. We provide theoretical conditions ensuring that additional customization always enhances user welfare. Empirically, using a large dataset of over four million anonymous joke ratings from 73,421 users, we confirm significant value creation attributable to these three elements.

Our findings reveal distinctive scaling properties of value creation along two key dimensions: the breadth of data (number of users) and its depth (level of detailed user information). While broader user bases consistently generate positive but diminishing returns, deeper customization, which is enabled by more capable learning models, can deliver non-diminishing and even super-additive returns. Empirically, customization constantly delivers incremental value, generating a learning frontier characterized by cascading waves of diminishing value creation instead of a single concave curve as may typically be assumed.

Our flexible model, suitable for both theoretical and empirical analyses, can serve as a starting point for further exploration as we establish a theoretical upper bound on the surplus an ideal recommender system can confer to its users. Potential extensions include analyzing firms' optimal pricing strategies in the presence of recommender systems, exploring the strategic manipulation of information to steer consumers toward more profitable items, and modeling platform competition

⁹As in the definition of $\Delta_{C,N}$, these values are obtained when we change N from 0 to 500. A similar result holds if we increase the size of the training set from $N = 1$ to $N = 500$.

employing different recommender technologies to analyze the resulting market outcomes and long-run market structure. Furthermore, our formalization of the conditions for super-additivity opens new avenues for research into optimal data acquisition strategies, as platforms could prioritize collecting data on features known to have strong, synergistic interaction effects.

References

- G. Aridor and D. Gonçalves. Recommenders’ originals: The welfare effects of the dual role of platforms as producers and recommender systems. *International Journal of Industrial Organization*, 83:102845, 2022.
- G. Aridor, D. Gonçalves, D. Kluver, R. Kong, and J. Konstan. The informational role of online recommendations: Evidence from a field experiment. Working paper, 2024.
- J. Asker, C. Fershtman, and A. Pakes. The impact of artificial intelligence design on pricing. *Journal of Economics & Management Strategy*, 33(2), 2023.
- M. O. Ayemowa, R. Ibrahim, and M. M. Khan. Analysis of recommender system using generative artificial intelligence: A systematic literature review. Number 10.1109/ACCESS.2024.3416962, pages 87742–87766. IEEE Access, 2024.
- P. Bajari, V. Chernozhukov, A. Hortaçsu, and J. Suzuki. The impact of big data on firm performance: An empirical investigation. In *AEA Papers and Proceedings*, volume 109, pages 33–37, 2019.
- G. Biglaiser, E. Calvano, and J. Crémer. Incumbency advantage and its value. *Journal of Economics & Management Strategy*, 28(1):41–48, 2019.
- E. Calvano, G. Calzolari, V. Denicolò, and S. Pastorello. Artificial intelligence, algorithmic pricing, and collusion. *American Economic Review*, 110(10):3267–97, October 2020.
- E. Calvano, G. Calzolari, V. Denicolò, and S. Pastorello. Artificial intelligence, algorithmic recommendations and competition. Working paper, 2025.
- X. Chen, Y. Li, and J. Mao. A nearly instance optimal algorithm for top k ranking under the multinomial logit model. Proc. 29th Annual ACM-SIAM Symposium on Discrete Algorithms, Society for Industrial and Applied Mathematics, Philadelphia, 2018.
- A. de Cornière and G. Taylor. Data and competition: A simple framework. *RAND Journal of Economics*, 2024. Forthcoming.
- I. P. Fainmesser, A. Galeotti, and R. Momot. Digital privacy. *Management Science*, 69(6):3157–3173, June 2022.
- Y. Feng, R. Caldentey, and C. T. Ryan. Robust learning of consumer preferences. *Operations Research*, 70(2):918–962, 2022.
- A. Hagiwara and J. Wright. Data-enabled learning, network effect and competitive advantage. *RAND Journal of Economics*, 54(4):638–667, 2023.

- C. Peukert, A. Sen, and J. Claussen. The editor and the algorithm: Recommendation technology in online news. *Management Science*, 70(9):5816–5831, 2023.
- M. Schäfer and G. Sapi. Complementarities in learning from data: Insights from general search. *Information Economics and Policy*, 65(101063), December 2023.
- A. Subbiah and V. Aggarwal. Transformers in music recommendation. Technical report, Google Research, <https://research.google/blog/transformers-in-music-recommendation/>, August 2024.
- L. Wu, Z. Zheng, Z. Qiu, T. Shen, C. Qin, C. Zhu, and E. Chen. A survey on large language models for recommendation. *World Wide Web*, 27(60), 2024.
- H. Yoganarasimhan. Search personalization using machine learning. *Management Science*, 66(3):1045–1070, 2020.

A Appendix - Proofs of Propositions

A.1 Proof of Proposition 1

Proof. Consider the user-optimal threshold $\tau^u = \frac{-v_0}{v_1 - v_0}$. By the definition of a consistent recommender system, i is recommended to history r' user *if and only if* the following inequality holds:

$$\frac{p_{(r',1)}^i}{p_{r'}^i} \geq \max_{j \in \mathcal{I}} \left\{ \frac{p_{(r',1)}^j}{p_{r'}^j}, \frac{-v_0}{v_1 - v_0} \right\}.$$

As $p_{(r',1)}^i + p_{(r',0)}^i = p_{r'}^i = p_{r'}^j$, $\forall i, j \in I$, the condition is equivalent to the following expressions

$$\begin{aligned} \frac{p_{(r',1)}^i}{p_{r'}^i} &\geq \max_{j \in \mathcal{I}} \left\{ \frac{p_{(r',1)}^j}{p_{r'}^j}, \frac{-v_0}{v_1 - v_0} \right\} \\ \Leftrightarrow (v_1 - v_0)p_{(r',1)}^i &\geq \max_{j \in \mathcal{I}} \left\{ (v_1 - v_0)p_{(r',1)}^j, -v_0(p_{(r',1)}^i + p_{(r',0)}^i) \right\} \\ \Leftrightarrow v_1p_{(r',1)}^i + v_0p_{(r',0)}^i &\geq \max_{j \in \mathcal{I}} \left\{ v_1p_{(r',1)}^j + v_0p_{(r',0)}^j, 0 \right\}. \end{aligned}$$

As this holds for any $r' \in R'$ and r' happens with probability $p_{r'}^i$, $\forall i \in I$, the result in Proposition 1 follows immediately. ■

A.2 Proof of Corollary 1

Proof. Consider the situation that there is only one target item. In the limit, the prediction error of a consistent recommender system with the threshold level $\{\tau(r')\}_{r' \in R'}$ when there is only one

target item ($i = C + 1$) converges to

$$\sum_{r' \in R'} p_{r'}^i \left(\frac{p_{(r',0)}^i}{p_{r'}^i} \mathbf{1} \left\{ \frac{p_{(r',1)}^i}{p_{r'}^i} \geq \tau(r') \right\} + \frac{p_{(r',1)}^i}{p_{r'}^i} \mathbf{1} \left\{ \frac{p_{(r',1)}^i}{p_{r'}^i} < \tau(r') \right\} \right). \quad (6)$$

That is, history r' occurs with probability $p_{r'}^i$, a false positive recommendation is made with probability $\frac{p_{(r',0)}^i}{p_{r'}^i}$ once the item is recommended to history r' target user, and a false negative event happens with probability $\frac{p_{(r',1)}^i}{p_{r'}^i}$ when the item is not recommended to the user. For each history r' , the expression (6) inside of the summation is

$$p_{(r',0)} \mathbf{1} \{ (1 - \tau)p_{(r',1)} \geq \tau p_{(r',0)} \} + p_{(r',1)} \mathbf{1} \{ (1 - \tau)p_{(r',1)} < \tau p_{(r',0)} \}.$$

Regardless of r' , the expression is minimized at $\tau = \frac{1}{2}$. Thus, the error bound in (6) is minimized at $\tau^u(r') = \frac{1}{2}$. Taking $\tau^u(r') = \frac{1}{2}$, the expression (6) is equivalent to

$$\begin{aligned} & \sum_{r' \in R'} p_{r'} \left(\frac{p_{(r',0)}}{p_{r'}} \mathbf{1} \{ p_{(r',1)} \geq p_{(r',0)} \} + \frac{p_{(r',1)}}{p_{r'}} \mathbf{1} \{ p_{(r',1)} < p_{(r',0)} \} \right) \\ &= \sum_{r' \in R'} p_{(r',0)} \mathbf{1} \{ p_{(r',1)} \geq p_{(r',0)} \} + p_{(r',1)} \mathbf{1} \{ p_{(r',1)} < p_{(r',0)} \} \\ &= \sum_{r' \in R'} \min \{ p_{(r',1)}, p_{(r',0)} \}, \end{aligned}$$

For the situations in which there are multiple target items, the expression above can immediately be extended to (4).

■

A.3 Proof of Proposition 2

Proof. To begin, we introduce a lemma which characterizes the limit utility that the target user expect from a consistent recommender system and a generic recommender system when there is only one item available to the user and a general threshold level is used.

Lemma 1 *Suppose there is only one item available to the target user, i.e., $i = C + 1$.*

1. *For any consistent mechanism, the expected utility to the target user converges to*

$$\sum_{r' \in R'} \mathbf{1} \{ p_{(r',1)}^i \geq \tau(r') p_{r'}^i \} (v_1 p_{(r',1)}^i + v_0 p_{(r',0)}^i).$$

2. The expected utility to the target user converges to

$$\mathbf{1}\left\{\sum_{r' \in R'} p_{(r',1)}^i \geq \tau^G\right\} \sum_{r' \in R'} (v_1 p_{(r',1)}^i + v_0 p_{(r',0)}^i).$$

(Proof of Lemma 1.) The target item is recommended to a user with history r' if and only if $\frac{p_{(r',1)}^i}{p_{r'}^i} \geq \tau(r')$. Once the item is recommended and tried by the user, the user receives $v_1 \frac{p_{(r',1)}^i}{p_{r'}^i} + v_0 \frac{p_{(r',0)}^i}{p_{r'}^i}$. The ex-ante utility is derived taking into account that the target user has a history r' with probability $p_{r'}^i$. The same logic can be applied in deriving the other utility characterizations. $\square$

Now, we prove Proposition 2. For each history of a target user r' , we construct a set of I correlation structures $q = \{q^i\}_{i \in \mathcal{I}}$ under which an extra degree in customization strictly hurts the target user. Because the proof is done by construction, it is without loss of generality to assume $I = 1$. For any cases with $I > 1$, we can simply let q^j , $j \neq C + I = M$, satisfies the following inequality and focus on q^M only:

$$z_j(r') = \frac{q_{(r',1)}^j}{q_{r'}^j} < \tau(r'), \quad \forall j \neq M \text{ and } \forall r' \in R'.$$

Let $I = 1$ and $\tau \neq \tau^u$, where τ is the threshold level the system adopts. Similar to the construction of R' , we denote the set of all outcomes that can be generated by the first $C - 1$ conditioning items by R'' . Similarly, $(r'', k, 1)$ and $(r'', k, 0)$ represents the positive and negative user experience with the target item of a history (r'', k) user, $k \in \{1, \dots, n_C\}$.

We will show that for any $\tau \neq \tau^u$, there exists q^M such that the target user with history $r'' \in R''$ receive strictly lower utility when the recommendation is customized based on C conditioning items and the target item than when it is customized based on first $C - 1$ conditioning items and the target item. By an induction argument, this will prove the existence of a correlation structure that users are strictly worse off from an extra degree in customization.

Firstly, if $\tau < \tau^u$, then consider the following correlation structure q^M . For each $r'' \in R''$ and $k \in \{1, \dots, n_C\}$, the following equalities and inequality hold:

$$\begin{cases} \frac{q_{(r'',k,1)}^M}{q_{(r'',k)}^M} = \tau & \text{if } k \neq n_C \\ \frac{q_{(r'',k,1)}^M}{q_{(r'',k)}^M} < \tau & \text{if } k = n_C. \end{cases}$$

Thus, we have $\frac{\sum_{k \in \{1, \dots, n_C\}} q_{(r'', k, 1)}^M}{\sum_{k \in \{1, \dots, n_C\}} q_{(r'', k)}^M} < \tau$. Under this correlation structure, when the recommendation is fully customized, item M is recommended to all users except the history (r'', n_C) user, whereas when the system omits conditioning item C in making recommendations, no user is recommended to try item M .

The user's utility under fully customized recommendations can be represented as follows using Lemma 1:

$$\begin{aligned} & \sum_{r'' \in R''} \sum_{k \in \{1, \dots, n_C\}} q_{(r'', k)}^M \mathbf{1} \left\{ \frac{q_{(r'', k, 1)}^M}{q_{(r'', k)}^M} \geq \tau \right\} \left(v_1 \frac{q_{(r'', k, 1)}^M}{q_{(r'', k)}^M} + v_0 \frac{q_{(r'', k, 0)}^M}{q_{(r'', k)}^M} \right) \\ &= \sum_{r'' \in R''} \sum_{k \in \{1, \dots, n_C\}} q_{(r'', k)}^M \mathbf{1} \left\{ \frac{q_{(r'', k, 1)}^M}{q_{(r'', k)}^M} \geq \tau \right\} (v_1 - v_0) \left(\frac{q_{(r'', k, 1)}^M}{q_{(r'', k)}^M} - \tau^u \right). \end{aligned}$$

As $\frac{q_{(r'', k, 1)}^M}{q_{(r'', k)}^M} = \tau < \tau^u$ for $k \neq n_C$, the user utility is strictly negative for users whose history is not $(r'', n_C, 1)$. On the other hand, the history $(r'', n_C, 1)$ user does not try the item. By construction, the item is recommended to no users when the system omits C , and the resulting user utility is zero.

Secondly, suppose we have $\tau > \tau^u$. Again, for each $r'' \in R''$ and $k \in \{1, \dots, n_C\}$, consider a correlation structure q^M that satisfies the following equality and inequality:

$$\begin{cases} \frac{q_{(r'', k, 1)}^M}{q_{(r'', k)}^M} > \tau & \text{if } k = 1 \\ \frac{q_{(r'', k, 1)}^M}{q_{(r'', k)}^M} = \tau & \text{if } k \neq 1 \text{ or } n_C \\ \frac{q_{(r'', 1, 1)}^M + q_{(r'', n_C, 1)}^M}{q_{(r'', 1)}^M + q_{(r'', n_C)}^M} = \tau. \end{cases}$$

Under this correlation structure, item M is recommended to all users when the system omits conditioning item C in making predictions. However, if it takes all conditioning items into account, the history (r'', n_C) user does not receive a recommendation. Because trying M is actually beneficial to all users, there is a missing opportunity if recommendations are fully customized. ■

A.4 Proof of Proposition 3

Proof. Without loss of generality, we only look at the case of $\frac{p_{(r', 1, 1)}}{p_{(r', 1)}} \geq \frac{p_{(r', 0, 1)}}{p_{(r', 0)}}$. The opposite case can be shown using the exact same logic.

To begin, note that the expected utility of r' user before the extra degree in customization is

$$\mathbf{1}\left\{\frac{p(r',1,1) + p(r',0,1)}{p(r',1) + p(r',0)} \geq \tau\right\}\left(v_1 \frac{p(r',1,1) + p(r',0,1)}{p_{r'}} + v_0 \frac{p(r',1,0) + p(r',0,0)}{p_{r'}}\right).$$

On the other hand, the expected utility of the user after the extra degree in customization is

$$\frac{p(r',1)}{p_{r'}} \mathbf{1}\left\{\frac{p(r',1,1)}{p(r',1)} \geq \tau\right\}\left(v_1 \frac{p(r',1,1)}{p(r',1)} + v_0 \frac{p(r',1,0)}{p(r',1)}\right) + \frac{p(r',0)}{p_{r'}} \mathbf{1}\left\{\frac{p(r',0,1)}{p(r',0)} \geq \tau\right\}\left(v_1 \frac{p(r',0,1)}{p(r',0)} + v_0 \frac{p(r',0,0)}{p(r',0)}\right).$$

Therefore, for each $r' \in R'$, the benefit of the marginal customization is

$$\begin{cases} 0 & \text{if } \tau \leq \frac{p(r',0,1)}{p(r',0)} \text{ or } \tau \geq \frac{p(r',1,1)}{p(r',1)} \\ -(v_1 p(r',0,1) + v_0 p(r',0,0))/p_{r'} & \text{if } \frac{p(r',0,1)}{p(r',0)} \leq \tau < \frac{p(r',1,1) + p(r',0,1)}{p(r',1) + p(r',0)} \\ (v_1 p(r',1,1) + v_0 p(r',1,0))/p_{r'} & \text{if } \frac{p(r',1,1) + p(r',0,1)}{p(r',1) + p(r',0)} \leq \tau < \frac{p(r',1,1)}{p(r',1)}. \end{cases} \quad (\text{Diff})$$

The if-condition in the first case utilizes our assumption $\frac{p(r',1,1)}{p(r',1)} \geq \frac{p(r',0,1)}{p(r',0)}$.

First, let item $C + 1$ and the target item are strictly correlated conditional on r' . Because we focus on $\frac{p(r',1,1)}{p(r',1)} \geq \frac{p(r',0,1)}{p(r',0)}$, the two items are strictly positively correlated. That is, we have

$$v_1 p(r',0,1) + v_0 p(r',0,0) < 0 \text{ and } v_1 p(r',1,1) + v_0 p(r',1,0) > 0.$$

Thus, for any r' that does not satisfies (5), the expected utility related to r' is strictly positive.

Conversely, suppose that the two items are not correlated conditional on r' . By definition of correlation, it is either both $v_1 p(r',1,1) + v_0 p(r',1,0)$ and $v_1 p(r',0,1) + v_0 p(r',0,0)$ are positive or both are negative. Thus, there always exists τ such that the utility represented in (Diff) is negative. ■

A.5 Proof of Proposition 4

Proof. For brevity, define x_{ij} , the expected payoff of the history (i, j) users under τ^u , by

$$x_{ij} = v_1 p_{ij1} + v_0 p_{ij0}.$$

The marginal benefit from conditioning on two item simultaneously is given as

$$\sum_{i,j \in \{0,1\}} \max\{x_{ij}, 0\} - \max\left\{\sum_{i,j \in \{0,1\}} x_{ij}, 0\right\}.$$

On the other hand, the sum of marginal benefits from conditioning on one item each is

$$\sum_{i \in \{0,1\}} \max \{x_{i1} + x_{i0}, 0\} + \sum_{j \in \{0,1\}} \max \{x_{1j} + x_{0j}, 0\} - 2 \max \left\{ \sum_{i,j \in \{0,1\}} x_{ij}, 0 \right\}.$$

Defining the difference between the two approaches by Δ and denoting $\max\{a, 0\}$ by a_+ , we have

$$\Delta = \sum_{i,j \in \{0,1\}} (x_{ij})_+ + \left(\sum_{i,j \in \{0,1\}} x_{ij} \right)_+ - \sum_{i \in \{0,1\}} (x_{i1} + x_{i0})_+ - \sum_{j \in \{0,1\}} (x_{1j} + x_{0j})_+.$$

Here, note that $a_+ = (a + |a|)/2$. Using this, we can write Δ as

$$\Delta = \frac{1}{2} \left(\sum_{i,j \in \{0,1\}} |x_{ij}| + \left| \sum_{i,j \in \{0,1\}} x_{ij} \right| - \sum_{i \in \{0,1\}} |x_{i1} + x_{i0}| - \sum_{j \in \{0,1\}} |x_{1j} + x_{0j}| \right).$$

This is because the non-absolute value terms cancel out to zero. Now, the super-additivity is equivalent to Δ be strictly positive.

We first show the necessity. Suppose, for a contrary that, for all $i, j \in \{0, 1\}$, we have either $x_{i1}x_{i0} \geq 0$ or $x_{1j}x_{0j} \geq 0$. First, if $x_{i1}x_{i0} \geq 0$ for all $i \in \{0, 1\}$, we have $|x_{i1}| + |x_{i0}| = |x_{i1} + x_{i0}|$, $\forall i$. Thus,

$$\Delta = \frac{1}{2} \left(\left| \sum_{i,j \in \{0,1\}} x_{ij} \right| - \sum_{j \in \{0,1\}} |x_{1j} + x_{0j}| \right) \leq 0,$$

which is a contradiction. Similarly, if $x_{1j}x_{0j} \geq 0$ for all $j \in \{0, 1\}$ we have

$$\Delta = \frac{1}{2} \left(\left| \sum_{i,j \in \{0,1\}} x_{ij} \right| - \sum_{i \in \{0,1\}} |x_{i1} + x_{i0}| \right) \leq 0.$$

The two remaining cases are 1) $x_{i1}x_{i0} \geq 0$ and $x_{(1-i)1}x_{(1-i)0} < 0$ and 2) 1) $x_{1j}x_{0j} \geq 0$ and $x_{1(1-j)}x_{0(1-j)} < 0$ for $i, j \in \{0, 1\}$. In these cases, if we have at least one product that is strictly smaller than zero in one dimension, say i -dimension, then it should be the case that all the inequalities on the other dimension, i.e., j -dimension, should be positive by our hypothesis. For instance, if $x_{i1}x_{i0} < 0$, it must be $x_{1j}x_{0j} \geq 0$ for all $j \in \{0, 1\}$. Thus, we go back to the previous case we already have shown to have a contradiction.

Now, to show the sufficiency, suppose that there is i, j such that $x_{ij} > 0$, $x_{i1} + x_{i0} \leq 0$ and $x_{1j} + x_{0j} \leq 0$. Then, it should be the case $x_{i(1-j)} < 0$ and $x_{(1-i)j} < 0$.

Under these assumptions, we have $(x_{ij})_+ = x_{ij}$, $x_{i(1-j)} = 0$, $x_{(1-i)j} = 0$, $(x_{i1} + x_{i0})_+ = 0$ and $(x_{1j} + x_{0j})_+ = 0$. Thus,

$$\begin{aligned} \Delta = & x_{ij} + (x_{(1-i)(1-j)})_+ + (x_{ij} + x_{i(1-j)} + x_{(1-i)j} + x_{(1-i)(1-j)})_+ \\ & - (x_{(1-i)j} + x_{(1-i)(1-j)})_+ - (x_{i(1-j)} + x_{(1-i)(1-j)})_+. \end{aligned} \quad (\star)$$

Now, compare $x_{(1-i)(1-j)}$ to two thresholds:

$$\theta_0 = \max\{-x_{i(1-j)}, -x_{(1-i)j}\} \text{ and } \theta_1 = -x_{ij} - x_{i(1-j)} - x_{(1-i)j}.$$

From $x_{i(1-j)} \leq -x_{ij}$ and $x_{(1-i)j} \leq -x_{ij}$, we have $\theta_0 \geq x_{ij}$ and $\theta_1 \geq \theta_0$. Thus, consider the following three cases:

Case 1: $x_{(1-i)(1-j)} \leq \theta_0$. Then

$$(x_{(1-i)(1-j)} + x_{(1-i)j})_+ = 0, \quad (x_{i(1-j)} + x_{(1-i)(1-j)})_+ = 0, \quad (x_{ij} + x_{i(1-j)} + x_{(1-i)j} + x_{(1-i)(1-j)})_+ = 0.$$

Hence by $(\star)$,

$$\Delta = x_{ij} + (x_{(1-i)(1-j)})_+ \geq x_{ij} > 0.$$

Case 2: $\theta_0 < x_{(1-i)(1-j)} \leq \theta_1$. Exactly one of the two sums is positive. Without loss of generality, suppose $\theta_0 = -x_{i(1-j)}$, so

$$(x_{i(1-j)} + x_{(1-i)(1-j)})_+ = x_{(1-i)(1-j)} + x_{i(1-j)},$$

while the other two positive-part terms in $(\star)$ vanish. Also $x_{(1-i)(1-j)} > 0$, so $(x_{(1-i)(1-j)})_+ = x_{(1-i)(1-j)}$. Therefore

$$\Delta = x_{ij} + x_{(1-i)(1-j)} - (x_{(1-i)(1-j)} + x_{i(1-j)}) = x_{ij} - x_{i(1-j)} \geq 2x_{ij} > 0,$$

because $x_{i(1-j)} \leq -x_{ij}$.

Case 3: $x_{(1-i)(1-j)} > \theta_1$. All three positive-part terms in $(\star)$ are active, giving

$$\begin{aligned}\Delta &= x_{ij} + x_{(1-i)(1-j)} + (x_{ij} + x_{i(1-j)} + x_{(1-i)j} + x_{(1-i)(1-j)}) \\ &\quad - (x_{(1-i)(1-j)} + x_{(1-i)j}) - (x_{i(1-j)} + x_{(1-i)(1-j)}) = 2x_{ij} > 0.\end{aligned}$$

In every case $\Delta > 0$. This proves strict super-additivity under τ^u . ■

A.6 Proof of Proposition 5

Proof. For notational simplicity, let $P[y_{(r',1)} \geq N_1 - y_{(r',1)}] = T(m, 2m - 1, z)$. That is,

$$T(m, 2m - 1, z) = \sum_{k=m}^{2m-1} \binom{2m-1}{k} z^k (1-z)^{2m-1-k}.$$

We first need to show that $v(m|r', p)$ increases in m . Let $X_n \sim \text{Binomial}(n, z)$. Note that the associated cumulative distribution function of the binomial random variable is

$$P[X_n \leq k] = \sum_{i=0}^k \binom{n}{i} z^i (1-z)^{n-i} = 1 - \sum_{i=k+1}^n \binom{n}{i} z^i (1-z)^{n-i} = 1 - T(k+1, n, z).$$

Using this, we have

$$\begin{aligned}T(m+1, 2m+1, z) &= 1 - P[X_{2m+1} \leq m] \\ &= 1 - P[X_{2m+1} \leq m | X_{2m-1} \leq m-2] P[X_{2m-1} \leq m-2] \\ &\quad - P[X_{2m+1} \leq m | X_{2m-1} = m-1] P[X_{2m-1} = m-1] \\ &\quad - P[X_{2m+1} \leq m | X_{2m-1} = m] P[X_{2m-1} = m] \\ &= 1 - P[X_{2m-1} \leq m-2] - (1-z^2) P[X_{2m-1} = m-1] - (1-z)^2 P[X_{2m-1} = m] \\ &= T(m-1, 2m-1, z) - (1-z^2) \binom{2m-1}{m-1} z^{m-1} (1-z)^m \\ &\quad - (1-z)^2 \binom{2m-1}{m} z^m (1-z)^{m-1}.\end{aligned}$$

Here, note that $\binom{2m-1}{m-1} = \binom{2m-1}{m}$ and $T(m-1, 2m-1, z) = T(m, 2m-1, z) + \binom{2m-1}{m-1} z^{m-1} (1-z)^m$.

The last expression can be simplified to

$$T(m+1, 2m+1, z) = T(m, 2m-1, z) - (1-2z) \binom{2m-1}{m} z^m (1-z)^m.$$

Thus, we have

$$\begin{aligned} v(m+1|r', p) - v(m|r', p) &= - (1-2z)(T(m+1, 2m+1, z) - T(m, 2m-1, z)) \\ &= (1-2z)^2 \binom{2m-1}{m} z^m (1-z)^m \geq 0. \end{aligned}$$

That is, for any realization of z , the user always expects weakly higher utility when there are more accumulated data points.

To show the diminishing marginal return property, let $\Delta v(m|r', p)$ be the marginal externality that the $(m+1)^{th}$ user contributes to the subsequent user. That is,

$$\Delta v(m|r', p) = v(m+1|r', p) - v(m|r', p).$$

Using the derivation in the above proposition, it has a closed form representation of

$$\Delta v(m|r', p) = (1-2z)^2 \binom{2m-1}{m} (z(1-z))^m.$$

Consider the ratio between the following two increments:

$$\begin{aligned} \frac{\Delta v(m+1|r', p)}{\Delta v(m|r', p)} &= \frac{v(m+2|r', p) - v(m+1|r', p)}{v(m+1|r', p) - v(m|r', p)} \\ &= \frac{\binom{2m+1}{m+1} z^{m+1} (1-z)^{m+1}}{\binom{2m-1}{m} z^m (1-z)^m} \\ &= \frac{2(2m+1)}{m+1} z(1-z) < 4z(1-z) \leq 1. \end{aligned}$$

That is, the increment diminishes.

Lastly, to show the third point, we want to show $v(1|r', p) > \frac{p(r',1)}{p_{r'}} v(1|(r', 1), q) + \frac{p(r',0)}{p_{r'}} v(1|(r', 0), q)$. Here, the right hand side is the expected value when the degree of customization is $C+1$. For $m=1$, the probability that the target item is recommended is $p_{(r',1)}/(p_{(r',1)} + p_{(r',0)})$ when the degree of customization is C and it is $p_{(r',1,1)}/(p_{(r',1,1)} + p_{(r',1,0)})$ when we have one more conditioning item. Thus, the expected utility for the case of customization of degree C is

$$\begin{aligned} v(1|r', p) &= \left(2 \frac{p(r',1)}{p_{(r',1)} + p_{(r',0)}} - 1 \right) \frac{p(r',1)}{p_{(r',1)} + p_{(r',0)}} \\ &= \left(2 \frac{p(r',1,1)}{p_{r'}} + 2 \frac{p(r',0,1)}{p_{r'}} - 1 \right) \frac{p(r',1,1) + p(r',0,1)}{p_{r'}} \end{aligned}$$

On the other hand, for the customization of degree $C + 1$, we have

$$v(1|(r', i), q) = \left(2 \frac{p(r', i, 1)}{p(r', i)} - 1 \right) \frac{p(r', i, 1)}{p(r', i)} \text{ for } i \in \{0, 1\}$$

Thus, we have

$$\begin{aligned} & v(1|r', p) - \left(\frac{p(r', 1)}{p_{r'}} v(1|(r', 1), q) + \frac{p(r', 0)}{p_{r'}} v(1|(r', 0), q) \right) \\ &= \left(2 \frac{p(r', 1, 1)}{p_{r'}} + 2 \frac{p(r', 0, 1)}{p_{r'}} - 1 \right) \frac{p(r', 1, 1) + p(r', 0, 1)}{p_{r'}} - \left(2 \frac{p(r', 1, 1)}{p(r', 1)} - 1 \right) \frac{p(r', 1, 1)}{p_{r'}} - \left(2 \frac{p(r', 0, 1)}{p(r', 0)} - 1 \right) \frac{p(r', 0, 1)}{p_{r'}} \\ &= \left(2 \frac{p(r', 1, 1)}{p_{r'}} + 2 \frac{p(r', 0, 1)}{p_{r'}} \right) \frac{p(r', 1, 1) + p(r', 0, 1)}{p_{r'}} - 2 \frac{p(r', 1, 1)}{p(r', 1)} \frac{p(r', 1, 1)}{p_{r'}} - 2 \frac{p(r', 0, 1)}{p(r', 0)} \frac{p(r', 0, 1)}{p_{r'}} \\ &= \left(2 \frac{p(r', 1, 1)}{p_{r'}} + 2 \frac{p(r', 0, 1)}{p_{r'}} \right) \frac{p(r', 1)}{p_{r'}} - 2 \frac{p(r', 1, 1)}{p(r', 1)} \frac{p(r', 1, 1)}{p_{r'}} - 2 \frac{p(r', 0, 1)}{p(r', 0)} \frac{p(r', 0, 1)}{p_{r'}} \\ &= \frac{2}{(p_{r'})^2} (p(r', 1, 1)p(r', 1) + p(r', 0, 1)p(r', 1) - (p(r', 1, 1))^2 - (p(r', 0, 1))^2) \\ &= \frac{4p(r', 1, 1)p(r', 0, 1)}{(p_{r'})^2} \end{aligned}$$

To derive the third and the fourth equations, we used the property that $p_{(r', 1)} = p(r', 1, 1) + p(r', 0, 1)$.

As $p_r > 0$, for all $r \in R$, we conclude that the difference is strictly positive. ■