

Online Appendix

Recommender systems and the value of user data

by Gunhaeng Lee and Julian Wright

This Online Appendix is organized as follows. Appendix A details the construction of our Bayesian recommender system, deriving the posterior predictive distribution for complete data. Appendix B provides the full methodology and detailed results of the empirical analysis summarized in the main text. Appendix C extends our Bayesian model to handle incomplete (or partial) data. Appendix D analyzes the welfare implications of our model with finite data. Finally, Appendix E explores the possibility of harmful customization.

A Bayesian model for a recommender system

In order to quantify the results implied by our theory with data, in this section we construct a Bayesian model of a recommender system that learns the correlation structure and makes customized predictions based on the target user’s history. We focus on *complete* data, although we do not require that the target user has given ratings for all the C items. A general Bayesian recommender system that handles *incomplete* data is explained in the section C.

A.1 Bayesian model of a recommender system

Let item $C + 1$ be the target item. The platform’s objective is to learn p^{C+1} , the correlation structure, so as to make predictions about a target user’s preference for items that the user has not yet tried. Again, to save notation, we write p for p^{C+1} . We use Bayesian parametric inference to model the learning. To specify the Bayesian model, we first set our prior distribution for p , and then update the prior distribution using the collected data. The posterior distribution for p immediately follows from the prior distribution and the likelihood function that generates the data. From the posterior distribution, we take our point estimator as the posterior mean.

The prior.

Initially, p is known only to the extent of a prior belief, which captures the platform’s knowledge about p . This includes any information about the items’ intrinsic values and relationships between the values of items. It is expressed in our model through the Dirichlet distribution:

$$q^0 = (q_r^0)_{r \in R} \sim \text{Dir}(\alpha^0),$$

where $\alpha^0 = (\alpha_r^0)_{r \in R}$ represents the concentration parameters. Note q^0 itself is a random vector, and there is no restriction imposed on the concentration parameters α_r^0 , $r \in R$ as long as they

are positive scalars. The Dirichlet distribution is a generalization of the Beta distribution to the multivariate case. The shape of the distribution is determined by the concentration parameters α^0 , and different concentration parameters can be used to accommodate different prior information. For example, $\alpha_r = 1, \forall r \in R$, corresponds to the uniform prior. Jeffreys prior, a commonly used non-informative prior, also can be accommodated by letting $\alpha_r = \frac{1}{2}, \forall r \in R$. In our various examples, we focus on cases in which all items have binary ratings. In such cases, the prior distribution is represented as

$$q^0 = (q_{(1,\dots,1)}^0, \dots, q_{(0,\dots,0)}^0) \sim \text{Dir}(\alpha_{(1,\dots,1)}^0, \dots, \alpha_{(0,\dots,0)}^0).$$

The Dirichlet distribution is a widely accepted way to describe prior knowledge in settings like this. Most importantly from our perspective, it is a conjugate prior for the multinomial distribution, so it is analytically and computationally tractable.

Data and likelihood function.

For X , let X' refer to the collection of all ratings excluding user 1's history and other items irrelevant in making a prediction, i.e., item $C+2$ to item $C+I$. That is, X' is the $(C+1) \times (N-1)$ submatrix of X which is obtained from X by removing the $C+2$ to $C+I$ rows of the first column. When the condition for complete data is satisfied, data from the $N-1$ previous users is simply a collection of outcomes each of which is independently generated according to the unknown probability vector p . We use y_r to denote the occurrences of r in X' . Let $y = (y_r)_{r \in R}$. The likelihood function associated with the data is

$$y|q^0 \sim \text{Multinomial}(N-1, q^0).$$

The posterior.

Finally, it can be verified that the posterior distribution induced from the prior distribution and the likelihood is a Dirichlet distribution with a concentration parameter $y + \alpha^0$:

$$q^1 \stackrel{d}{=} q^0|X' \sim \text{Dir}(y + \alpha^0).$$

In our main example, Table 1 in the main text, we have $X' = \{(1,1), (0,0), (1,1), (1,0)\}$ and $y = (2,1,0,1)$. Furthermore, if we take the uniform prior distribution, the resulting posterior distribution is $\text{Dir}(3,2,1,2)$.

A.2 Predicting what users like

From the posterior in the previous section, we are now ready to make a prediction, which we define as the conditional probability of the target user liking the target item i given the data. Let x denote the target user's ratings of items, i.e., x is the first column of X . For x , we define two disjoint sets

$R_i^+(x)$ and $R_i^-(x)$ as follows:

$$\begin{aligned} R_i^+(x) &= \{r \in R | r_i = 1 \text{ and } r_k = x_k, \forall k \text{ s.t. } x_k \neq \emptyset\} \\ R_i^-(x) &= \{r \in R | r_i = 0 \text{ and } r_k = x_k, \forall k \text{ s.t. } x_k \neq \emptyset\}. \end{aligned}$$

That is, $R_i^+(x)$ is the subset of R that satisfies (1) the rating of the target item i is positive, and (2) the rating of the conditioning item k , $k \neq i$, is the same as the target user's rating of the conditioning item k if the user has left a rating on it. The sets are specific to the target user's ratings and also to the target item. For example, in the case of Table 1 in the main text, when user 1 is the target user, we have

$$R_2^+((1, \emptyset)) = \{(1, 1)\} \text{ and } R_2^-((1, \emptyset)) = \{(1, 0)\}.$$

On the other hand, if a new target user, user 6, participates in the platform and when item 1 is the target item, the sets now become

$$R_1^+((\emptyset, \emptyset)) = \{(1, 1), (1, 0)\} \text{ and } R_1^-((\emptyset, \emptyset)) = \{(0, 1), (0, 0)\}.$$

Let the posterior belief q^1 follow $Dir(\alpha^1)$ after learning from data X' , where $\alpha^1 = y + \alpha^0$ and y is derived from the data X' . The true probability that the target user likes the target item i is denoted by

$$z_i(x) = \frac{\sum_{r \in R_i^+(x)} p_r}{\sum_{r \in R_i^+(x) \cup R_i^-(x)} p_r}.$$

Let Y_i be a random variable which takes value 1 with probability $z_i(x)$ and 0 otherwise (i.e. $Y_i \sim Bernoulli(z_i(x))$). Using q^1 and x , we define the pointwise estimator $\hat{z}_i^N(x)$ of $z_i(x)$ by its conditional probability

$$\hat{z}_i^N(x) \equiv P[X_i = 1 | X] = \frac{P[Y_i = 1, x | X']}{P[x | X']}.$$

We denote the associated posterior predictive distribution by $\hat{\mathbf{z}}_i^N(x)$. Note that an equivalent representation for the estimator is the expected value of Y_i conditional on X .

The following proposition characterizes the estimator and the predictive distribution for the parameters, showing they have concise expressions in terms of the concentration parameters. Moreover, the process of Bayesian learning and prediction is *computationally efficient*, in the sense that the new information can be updated by counting.

Proposition OA 1 *With complete data, let $\alpha_r^1 = y_r + \alpha_r^0$, where y_r denotes the occurrences r in the data and $\{\alpha_r^0\}_{r \in R}$ denotes the Dirichlet model parameters. The recommender system implies:*

1. The probability of the target user liking item i is

$$\hat{z}_i^N(x) = \frac{\sum_{r \in R_i^+(x)} \alpha_r^1}{\sum_{r \in R_i^+(x) \cup R_i^-(x)} \alpha_r^1}.$$

2. The associated distribution for the target user liking item i is

$$\hat{\mathbf{z}}_i^N(x) \sim \text{Beta} \left(\sum_{r \in R_i^+(x)} \alpha_r^1, \sum_{r \in R_i^-(x)} \alpha_r^1 \right).$$

3. For any parameter α^0 for the prior distribution and for all x , $\hat{\mathbf{z}}_i^N(x) \rightarrow z_i(x)$, in mean-square. Hence, $\hat{z}_i^N(x) \rightarrow z_i(x)$.

Proof. To begin, let x be the history of the target user. Note that the Dirichlet distribution is stable with respect to any aggregation.¹ Given this property, we have

$$\left(\sum_{r \in R_i^+(x)} q_r^1, \sum_{r \in R_i^-(x)} q_r^1, \sum_{r \notin R_i^+(x) \cup R_i^-(x)} q_r^1 \right) \sim \text{Dir} \left(\sum_{r \in R_i^+(x)} \alpha_r^1, \sum_{r \in R_i^-(x)} \alpha_r^1, \sum_{r \notin R_i^+(x) \cup R_i^-(x)} \alpha_r^1 \right).$$

For the sake of notation, we denote $\sum_{r \in R_i^+(x)} q_r^1$ and $\sum_{r \in R_i^-(x)} q_r^1$ by b_1 and b_2 respectively. Similarly, let β_1 , β_2 and β_3 denote $\sum_{r \in R_i^+(x)} \alpha_r^1$, $\sum_{r \in R_i^-(x)} \alpha_r^1$ and $\sum_{r \notin R_i^+(x) \cup R_i^-(x)} \alpha_r^1$ respectively. The conditional probability of positive experience with i can be represented as $\frac{b_1}{b_1+b_2}$ when $(b_1, b_2, 1 - b_1 - b_2) \sim \text{Dir}(\beta_1, \beta_2, \beta_3)$.

For $k \in \{1, 2, 3\}$, define an independent set of random variables γ_k each of which follows a gamma distribution with a shape parameter β_k and a rate parameter θ for some $\theta > 0$; i.e., $\gamma_k \sim \text{Gamma}(\beta_k, \theta)$.

It is well known that $\left(\frac{\gamma_1}{\gamma_1 + \gamma_2 + \gamma_3}, \frac{\gamma_2}{\gamma_1 + \gamma_2 + \gamma_3}, \frac{\gamma_3}{\gamma_1 + \gamma_2 + \gamma_3} \right) \sim \text{Dir}(\beta_1, \beta_2, \beta_3)$. Thus, we have the following equality in distribution:

$$(b_1, b_2) \stackrel{d}{=} \left(\frac{\gamma_1}{\gamma_1 + \gamma_2 + \gamma_3}, \frac{\gamma_2}{\gamma_1 + \gamma_2 + \gamma_3} \right).$$

Here, note that if $X \stackrel{d}{=} Y$, then $h(X) \stackrel{d}{=} h(Y)$ for any deterministic function h . Letting $h_1(x_1, x_2) = \frac{x_1}{x_1 + x_2}$ and $h_2(x_1, x_2) = \frac{x_2}{x_1 + x_2}$, we have

$$\frac{b_1}{b_1 + b_2} \stackrel{d}{=} \frac{\gamma_1}{\gamma_1 + \gamma_2} \text{ and } \frac{b_2}{b_1 + b_2} \stackrel{d}{=} \frac{\gamma_2}{\gamma_1 + \gamma_2}.$$

Furthermore, using the relationship between the gamma distribution and the Dirichlet distribution

¹If $(p_1, \dots, p_n)$ is Dirichlet with parameter $\alpha_1, \dots, \alpha_n$, then a collection of sums of elements also follows Dirichlet. For example, when $n = 4$, $(p_1 + p_4, p_2 + p_3) \sim \text{Dir}(\alpha_1 + \alpha_4, \alpha_2 + \alpha_3)$.

once again, we conclude that

$$\left(\frac{b_1}{b_1 + b_2}, \frac{b_2}{b_1 + b_2} \right) \sim \text{Beta}(\beta_1, \beta_2).$$

That is, the predictive experience of the target user with i takes the form of the following random variable: $\hat{\mathbf{z}}_i^{\mathbf{N}} \sim \text{Beta}(\beta_1, \beta_2)$. The first statement in the proposition can be obtained by simply taking an expectation of $\hat{\mathbf{z}}_i^{\mathbf{N}}$:

$$\hat{z}_i^N = \frac{\beta_1}{\beta_1 + \beta_2} = \frac{\sum_{r \in R_i^+(x)} \alpha_r^1}{\sum_{r \in R_i^+(x) \cup R_i^-(x)} \alpha_r^1}.$$

Lastly, recall that we already have shown that

$$\frac{\sum_{r \in R_i^+(x)} q_r^1}{\sum_{r \in R_i^+(x) \cup R_i^-(x)} q_r^1} \sim \text{Beta} \left(\sum_{r \in R_i^+(x)} \alpha_r^1, \sum_{r \in R_i^-(x)} \alpha_r^1 \right),$$

for given data X and user rating x . Consider now that the data is collected from K users. Let K_1 be the number of ratings whose associated outcomes are in $R_i^+(x)$, K_2 be the number of ratings whose associated outcomes are in $R_i^-(x)$, and M be the total number of trials whose outcome is consistent with x , i.e., $K = K_1 + K_2$. We have

$$\lim_{K \rightarrow \infty} \mathbb{E}[\hat{\mathbf{z}}_i^{\mathbf{K}} | X_K, x] = \lim_{K \rightarrow \infty} \frac{K_1 + \sum_{r \in R_i^+(x)} \alpha_r^0}{K_1 + K_2 + \sum_{r \in R_i^+(x) \cup R_i^-(x)} \alpha_r^0} \stackrel{a.s.}{=} \lim_{K \rightarrow \infty} \frac{K_1}{K}.$$

Here, X_K denotes the data from K users. By the law of large numbers, the last term is the same as z_i . On the other hand,

$$\begin{aligned} \lim_{K \rightarrow \infty} \mathbb{V}ar[\hat{\mathbf{z}}_i^{\mathbf{K}} | X_K, x] &= \lim_{K \rightarrow \infty} \frac{(K_1 + \sum_{r \in R_i^+(x)} \alpha_r^0)(K_2 + \sum_{r \in R_i^-(x)} \alpha_r^0)}{(K + \sum_{r \in R_i^+(x) \cup R_i^-(x)} \alpha_r^0)^2 (K + 1 + \sum_{r \in R_i^+(x) \cup R_i^-(x)} \alpha_r^0)} \\ &\leq \lim_{K \rightarrow \infty} \frac{1}{K + 1 + \sum_{r \in R_i^+(x) \cup R_i^-(x)} \alpha_r^0} = 0, \end{aligned}$$

which gives the convergence in mean-square to z_i . ■

The proposition implies asymptotic learning occurs in our model, so that the prediction becomes more and more precise as we use more data. The Bayesian approach we present is consistent and robust in the sense that the predictions converge to the true probabilities regardless of the choice of the prior distribution within the Dirichlet family. Furthermore, for properly chosen concentration parameters α^0 , such as parameters of the uniform prior or the Jeffreys prior, they also satisfy the *unbiasedness* condition presented in Definition 5 in the main text.

B Evidence from data

This section provides the full methodology and detailed results of the empirical analysis summarized in the main text. We first describe the dataset and the benchmark simulation setup. We then present the detailed results for the value decomposition, the basis for Table 2 in the main text, and the marginal value of customization, selection, screening, and additional user data.

The Jester dataset² contains anonymous ratings of 100 jokes from 73,421 users, collected over the period from April 1999 to May 2003. Participants choose their ratings via a rating bar over the interval $[-10, 10]$. The dataset contains their recorded ratings, which are rounded to two decimal places. To fit the data into our environment, we convert the data to a binary rating: positive ratings and negative ratings.⁴ We call this the Jester binary dataset. Not all users rate all jokes, with around 40% of the ratings out of the total 7,342,100 being missing. Although the sparsity is 40%, there are 14,116 users who have completed rating on all 100 jokes.

Throughout the empirical analysis, we employ the Bayesian learning and prediction model in the previous section as the estimator of the recommender system.⁵ There are three key parameters, C, I, τ , which define such a recommender system (recall these are the number of target items, the threshold, and the number of non-target items). As a benchmark setting, we take $M = 10$ and assume $C = 9$ and $I = 1$ with $\tau = \frac{1}{2}$, as well as a fixed way of running simulations, as we will now explain.

For a given counterfactual experiment, we run 1,000 simulations. In each simulation, we randomly select 10,000 users (the training group) from the 14,116 users who have completed rating on all 100 jokes. We then randomly select another 10,000 users (the test group) from the remaining 63,421 users. In this benchmark setting, we assume $I = 1$ and choose nine conditioning items (i.e. jokes) at random, i.e., $C = 9$.⁶ We use the uniform prior, i.e., a Dirichlet distribution with the parameter value being a vector of ones. The Bayesian mechanism learns the correlation structure about the ten items using the training group’s data. As we will see, ten items turns out to be sufficient to learn the value that the recommender system creates.⁷ After learning the correlation structure from the training group, we make predictions about the test group’s experiences with the target item. We assume $(v_1, v_0) = (1, -1)$ and apply the user-optimal threshold $\tau^u = \frac{1}{2}$. Thus, a try-recommendation is made to the target user only if the target item is more likely to induce a positive experience with the user (i.e., if the predicted probability of a positive rating is above $1/2$). The simulation assumes the recommended item is tried. And we take the user’s actual rating

²Jester dataset released by AUTOLAB.³

⁴There are 4,116 zero ratings and they are converted to negative ratings, reflecting that out of the total of 4,136,360 ratings, 2,418,393 are positive ratings and 1,717,967 are negative ratings, with a zero rating being below both the average and the median rating.

⁵We will use the uniform distribution for the prior distribution.

⁶All the random samplings within each simulation are done without replacement.

⁷According to our theoretical results in the main text, the user surplus (weakly) increases in the number of items that the prediction mechanism conditions on. However, more items require more data points to ensure asymptotic learning. With our dataset, it can be checked that ten items is sufficient to approximate maximal learning.

(either one or negative one) as the resulting utility to the user. If the predicted rating is strictly below $1/2$, the item is assumed to be not recommended to the user and so not tried. The resulting utility is recorded as zero in this case. We consider 1,000 such simulations, each time with a different random selection of the training group and test group, and a different random selection of the target item and nine other items. Note that for any $C > 0$, since target users often do not have ratings for all C items, the actual degree of customization in effect is smaller than C .

B.1 Value of data and its decomposition: Evidence

To quantify our theoretical results in the main text, we first investigate how data increases user welfare through the recommender system for the Jester binary dataset.

Our theoretical findings imply that recommender systems add value via three key functions: customization, selection and screening. Table B1 summarizes the corresponding empirical finding. For details, refer to the main text.

		<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 B1: User surplus from recommender systems
scr: screening, sel: selection

Table B1 presents⁸ the estimated average utility that users receive from recommender systems with different parameter settings. Without selection and screening, i.e., $I = 1$ and $\tau = 0$, the simulation results show that the average utility of users is 0.148 without any customization. This corresponds to the utility users should expect when there is no recommender system available and instead they try each of the randomly selected items. On top of this baseline value, screening ($\tau = 1/2$) adds 0.045 additional average utility when there is no customization ($C = 0$) and 0.110 under customization ($C = 9$). This is for the baseline case without selection ($I = 1$). This shows, without selection, out of the total 0.110 increase in average utility from learning, generic learning contributes 0.045 to the increase in average utility, whereas 0.065 is created from customization, suggesting more value comes from within-user customization than from across-user learning about

⁸Standard Error=sample standard deviation/ $\sqrt{\text{number of simulations}}$

the target item.⁹ On the other hand, adding selection without screening (so focusing on the case with $\tau = 0$ but comparing $I = 3$ with $I = 1$), adds 0.194 without customization and 0.209 with customization. The additional value from customization is less pronounced when there are multiple jokes to select from. This reflects that in the jokes database, even with a small number of jokes, there is a high likelihood of there being at least one joke which most people like.¹⁰ This is consistent with our Proposition 3, since items that have positive (or negative) ratings regardless of history cannot add significant amounts of value through customization. The recommender system adds even more value to users when selection, screening and customization coexist. When $I = 3$, $\tau = \frac{1}{2}$, the consistent recommender system adds 0.223 more value to users compared to the situation without a recommender system.

Marginal value of customization

Proposition 3 predicts additional customization increases user welfare under the user-optimal threshold level. To measure how much user surplus increases from additional customization, we run simulations following our benchmark setting ($I = 1$, $\tau = \frac{1}{2}$) but changing the degree of customization C from one to nine. Thus, for a given simulation with 10,000 training group users, 10,000 test group users, one target item and nine conditioning items, all randomly drawn, we first let the Bayesian recommender system learn from the training group users' ratings on the target item only, and make predictions about the test group users. This corresponds to the situation in which zero degrees of customization have taken place ($C = 0$ case). Next, we increase C by one ($C = 1$), letting the system learn from the same training group user data about the two-item correlation structure associated with the target item and one of the items from the nine-item selection before it makes customized predictions about the test group. We repeat this until $C = 9$. In short, we increase the degree of customization in recommendations from zero to nine and compare the resulting average utilities, keeping everything else fixed according to the benchmark setting. The final user utility is calculated as the average utility over 1,000 such simulations.

Figure B1 summarizes the average utility of users from different degrees of customization and the corresponding 95% confidence intervals from running the simulation 1,000 times. For comparison purposes, the average utility users receive in the absence of a recommender system is also presented in the plot ("w/o RS" — the average utility under $I = 1$ and $\tau = 0$).

Consistent with our theoretical predictions, the user value increases as the recommender system provides more customized predictions. It is also observed that the average utility exhibits a diminishing return to customization when averaged over the 1,000 simulations. However, for any particular set of items that the increment can sometimes increase rather than decrease in the degree

⁹However, note that customization creates no value unless it is combined with across-user learning since without data from multiple users there is no way to learn the correlation structure required for customization. So in this sense, the value created from customization augments the value created from having data on many users and can really be thought of as value created from combining customization with across-user learning.

¹⁰More than 50% of all 100 jokes have above 60% positive ratings.

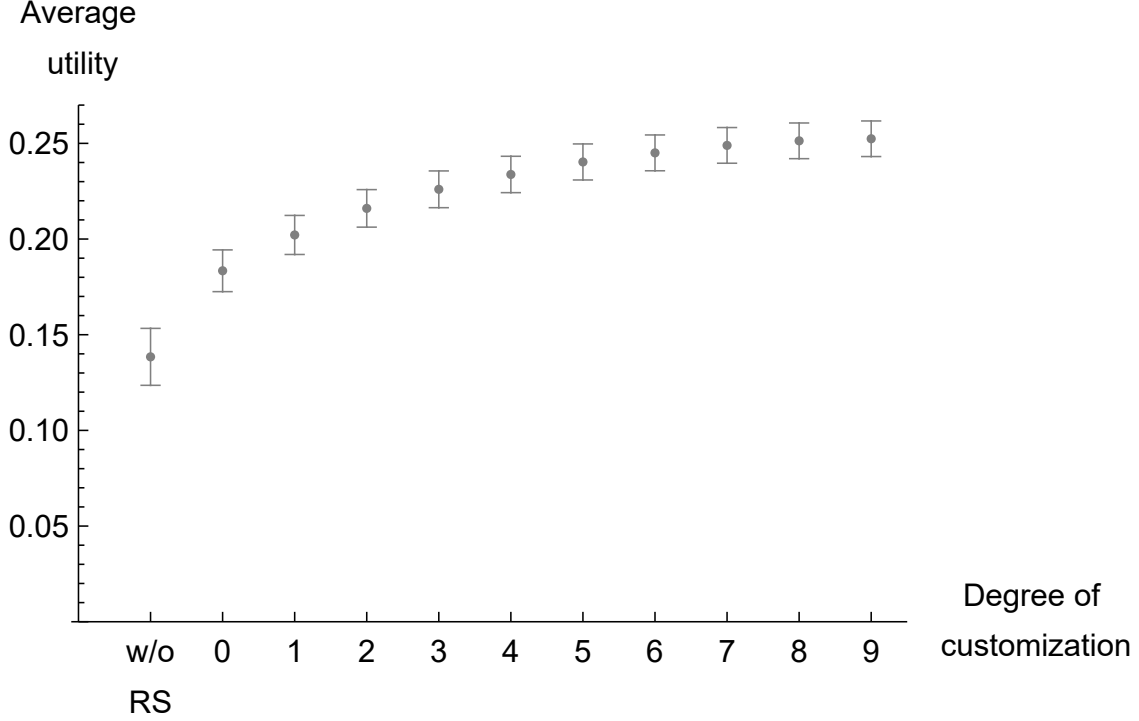


Figure B1: Average utility for different degrees of customization

of customization. We confirmed this is true in our data by considering a single draw of ten items and inspecting how the increment of average utility changes in the degree of customization.

Marginal value of selection

We turn to measuring the marginal value to users of additional target items. As predicted in Proposition 1, wider availability of selection makes it easier for the system to find a better item for users. To quantify the marginal value, based on our benchmark setting, we change the number of target items (I) from one to fifteen, and evaluate the resulting average utility of users. Figure B2 depicts the average utility of users in terms of I under $C = 9$ and $\tau = \frac{1}{2}$ and the corresponding 95% confidence intervals.

Figure B2 also shows that the marginal benefit of having an additional target item to choose from diminishes as more items become available for selection. This is related to the diminishing marginal increment property of the mean of the largest order statistic. When I items are drawn at random according to a fixed distribution, it can be shown that the mean of the largest order statistic increases in the number of items, I , at a diminishing rate. Although we do not explicitly assume any distribution behind the item's selections, we believe a similar logic also applies to our case.

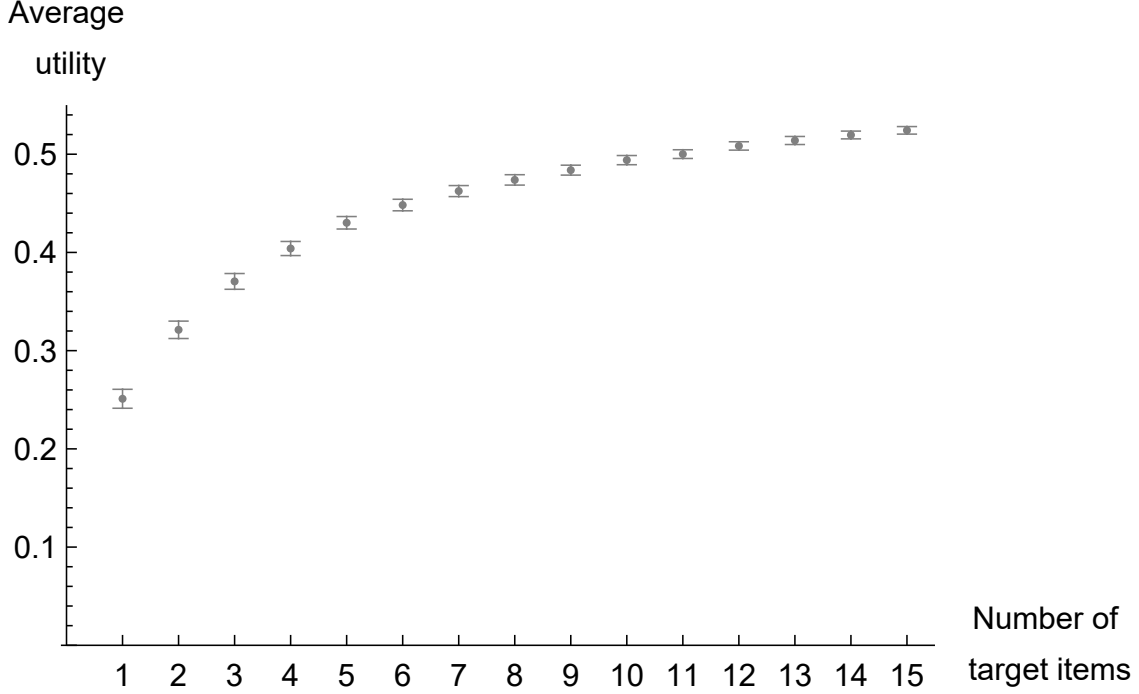


Figure B2: Average utility in terms of the number of target items

Marginal value of screening

In the model, the recommender system adjusts its level of screening by changing the threshold, which in our framework is the only channel through which a possible misalignment of interest between the platform and the users can arise. We explore what happens when the platform-optimal threshold differs from that which is best for a user ($\tau^p \neq \tau^u$) at each degree of customization. We evaluate the value to users under different platform-optimal thresholds, $\tau^p = 0.1, 0.2, \dots, 0.9$. For all other parameters, we stick to our benchmark setting. Note a low value of τ^p could capture a platform that is biased towards usage because it is compensated based on users trying the item rather than whether they like it or not, while a high value of τ^p could capture an overly conservative platform that does not want the user to try the item unless it is very confident the user will like it.

For each threshold τ^p , we repeat 1,000 simulations exactly as in our benchmark setting, changing the degree of customization from zero to nine. We find that the reduction in average utility associated with the misalignment becomes more significant and the customization becomes less valuable as the level of misalignment increases. The simulation results are summarized in Figure B3. At the maximum degree of customization we test, $C = 9$, the average utility is 0.148, 0.191, 0.204, 0.031 when $\tau^p = 0.1, 0.3, 0.7, 0.9$, which are 0.110, 0.057, 0.053, 0.221 lower than the average utility the optimal threshold $\tau^p = 0.5$ creates. Overall, we observe around 10.9% of reduction in average utility if we lower τ^p by 0.1 from the optimal level $\tau^p = 0.5$. On the other hand, if we increase τ^p

by the same margin, there is around a 22.0% reduction in average utility.¹¹

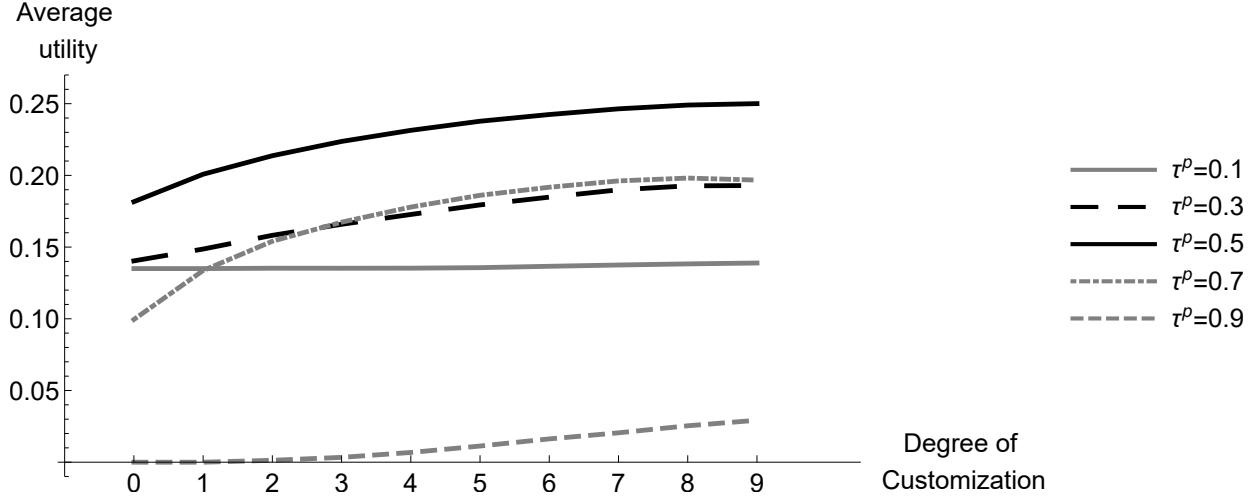


Figure B3: Average utility under different threshold levels τ^p

Additionally, a higher degree in customization leads to a higher average utility of users, as can be seen from the increasing nature of the curves in Figure B3. However, the increased utility from customization is less when the platform’s threshold is misaligned with users. Especially, when $\tau^p < 1/2$, the gain from customization is much more limited due to the misalignment: Customization contributes 0.06 to the increase in average utility when $\tau^p = 1/2$, whereas it only contributes 0.003 to the increase when $\tau^p = 0.1$ and 0.048 when $\tau^p = 0.3$. Although our theoretical analysis suggests that when τ^p is far away from the user-optimal level of $1/2$ it is possible for customization to harm consumers, that situation doesn’t arise in our data.

For the full detail of this experiment, Figure B4 presents the 95% confidence intervals related to the experiment of Figure B3.

Marginal value of additional users

We study how the average utility of users changes as the number of previous users that the platform can learn from increases. For one simulation, we start as usual by randomly selecting the training group (10,000 users) and test group (10,000 users) and ten items according to our benchmark setting, i.e., $(C, I, \tau) = (9, 1, 1/2)$. Then for each given simulation, we run 1,000 rounds of predictions as follows. We start by taking only one user from the training group, to form the training *set*. We run our usual prediction exercise but with the training group replaced by this training set. In each subsequent round, we pick ten new users at random from the training group and add those users’ data to the existing training set to form a larger training set. We repeat our usual prediction

¹¹The findings that the average utility stays positive even with $\tau^p = 0.1$ and that there is a larger utility loss from a high threshold than a low threshold reflect that the overall average utility from our rating data is positive. Since we normalize the outside option of not trying a joke to zero, the average utility from trying all of any set of jokes recommended will tend to be strictly positive.

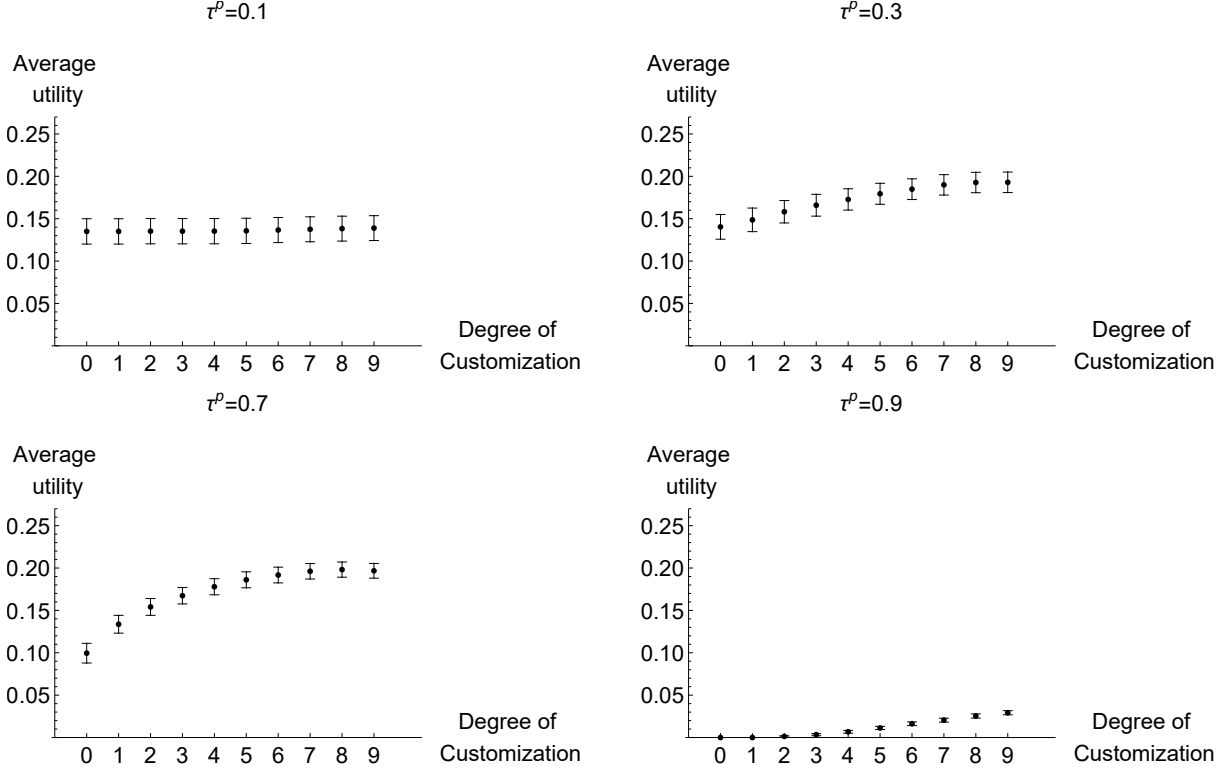


Figure B4: Different shapes of learning curves

exercise with the corresponding training set in each round. This is repeated until all 10,000 users have been added. This way we can measure how the size of the training set affects the resulting user welfare from the test group. Note as usual, the test group and the set of ten items remain fixed for a given simulation. All other aspects remain the same as in our benchmark setting, and we run 1,000 such simulations.

Figure B5 depicts average user utility in terms of the number of previous user data points, with the corresponding 95% confidence interval being represented by the shaded area. After all 10,000 users' data have been used, the average target user utility is around 0.260, which is around 78.87% higher than the average target user utility when no learning has taken place (i.e., when only one user's data has been used to train the model). Furthermore, it is clear that the marginal increment in the average user utility diminishes as we increase the data size as is predicted in Section 3.5.

C Finite data analysis: Making predictions from a partial dataset

The Bayesian approach with complete data uses the conjugacy relationship of the multinomial distribution and the Dirichlet distribution. Since each user's review is an outcome of an i.i.d. multinomial distribution, the posterior distribution remains in the Dirichlet family. However, this conjugacy relationship no longer holds when some ratings are missing (i.e., *partial data*) reflecting that the outcomes are not drawn from an identical distribution. For example, suppose there are

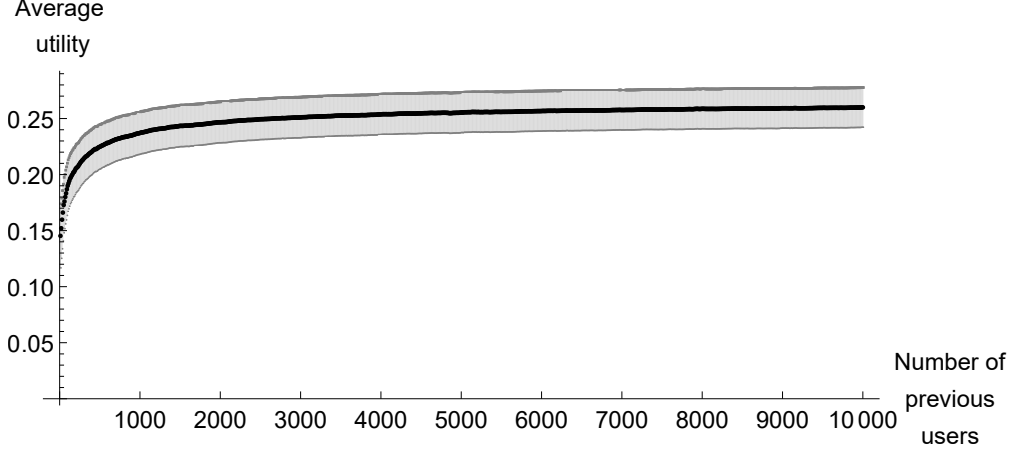


Figure B5: Average utility in terms of number of previous users

two items. The associated multinomial trial has the outcomes $R = \{(1, 1), (1, 0), (0, 1), (0, 0)\}$. If a user has an empty rating on item 2 while she has a positive rating on item 1, this outcome corresponds to an outcome of a binomial trial with success probability $(p_{(1,1)} + p_{(1,0)}, p_{(0,1)} + p_{(0,0)})$. It is theoretically possible to find a posterior distribution, which is a combination of 2^E Dirichlet distributions in this case, where E refers to the number of missing ratings. Thus, computationally, it becomes more and more intractable to use the conditional probability for a prediction as E grows. However, in what follows, we show that the Bayesian prediction can still be derived using a relatively simple formula.

Let γ_j denote the j^{th} user's ratings over the $C + 1$ items. Note that with *partial* data, even a non-target user may not have ratings on all $C + 1$ items. Thus, unlike $r \in R$, γ_j may contain $\emptyset$ in its elements.

To begin, we first define some notations that are used to express the prediction. As in the main model, suppose that we have N users and the number of conditioning items is C . Since we are interested in making a prediction about each target item, we let $I = 1$ for simplicity. The corresponding set of outcomes is again denoted by R . A typical element r of R is a length $C + 1$ vector whose i^{th} element records a non-empty rating of item i , $i \in \{1, \dots, C + 1\}$. For notational simplicity, we denote $\{1, \dots, N\}$ by $\bar{N}$. For $r' \in R$ and for a given collection of ratings $\{r_j\}_{j \in \bar{N}}$, define $N(r' | \{r_j\}_{j \in \bar{N}})$ to be the number of r' in $\{r_j\}_{j \in \bar{N}}$. For example, when the collection of ratings is given by $\{r_1, r_0, r_1\}$, we have $N(r_1 | \{r_1, r_0, r_1\}) = 2$ and $N(r_0 | \{r_1, r_0, r_1\}) = 1$.

Now, for partial data X , which is a collection $\{x_j\}_{j \in \bar{N}}$, we want to estimate the probability that the target user likes the target item, $C + 1$. Let $\bar{X}$ be an augmented data that has a positive rating of the target user 1's item $C + 1$ while all the other ratings including the empty ratings stay the same as in X . That is, the $(C + 1) \times 1^{th}$ element in X is $\emptyset$ and it is 1 in $\bar{X}$. Using this notation, the prediction can be denoted by $P[\bar{X} | X]$. To describe the process of the prediction, let x_j be the collection of x_{ij} for $i \in \{1, \dots, C + 1\}$ which is the set of ratings by user j , and $R(x_j)$ be the set

of outcomes consistent with the ratings left by the user. That is,

$$R(x_j) = \{r \in R | r_i = x_{ij}, \forall i \text{ s.t. } x_{ij} \neq \emptyset\}.$$

Recall that in the Bayesian model in Section A, we define $R_i^+(x_j)$ as follows

$$R_i^+(x_j) = \{r \in R | r_{ij} = 1 \text{ and } r_{kj} = x_{kj}, \forall k \text{ s.t. } x_{kj} \neq \emptyset\}.$$

Proposition OA 2 *The prediction with partial data can be characterized as follows:*

$$P[\bar{X}|X] = \frac{\sum_{\gamma_1 \in R_{C+1}^+(x_1), \gamma_j \in R(x_j), j \neq 1} \prod_{r \in R} \Gamma(\alpha_r^0 + N(r|\{\gamma_j\}_{j \in \bar{N}}))}{\sum_{\gamma_j \in R(x_j), j \in \bar{N}} \prod_{r \in R} \Gamma(\alpha_r^0 + N(r|\{\gamma_j\}_{j \in \bar{N}}))}.$$

When the prior parameters α^0 is a vector of natural numbers, the expression for the prediction can be represented as follows.

Corollary 1

$$P[\bar{X}|X] = \frac{\sum_{\gamma_1 \in R_{C+1}^+(x_1), \gamma_j \in R(x_j), j \neq 1} \prod_{r \in R} (\alpha_r^0 - 1 + N(r|\{\gamma_j\}_{j \in \bar{N}}))!}{\sum_{\gamma_j \in R(x_j), j \in \bar{N}} \prod_{r \in R} (\alpha_r^0 - 1 + N(r|\{\gamma_j\}_{j \in \bar{N}}))!}.$$

We can illustrate this prediction mechanism using a simple example. The ratings are summarized by the following table when user 1 is the target user and the target item is item 2. The set of outcomes is $R = \{(1,1), (1,0), (0,1), (0,0)\}$ and the associated probability is $p =$

	user 1	user 2	user 3
item 1	1	$\emptyset$	1
item 2	$\emptyset$	1	0

Table C1: Example 2

$(p_{(1,1)}, p_{(1,0)}, p_{(0,1)}, p_{(0,0)})$. Suppose that we begin with a uniform prior, i.e., $p^0 \sim Dir(1, 1, 1, 1)$ whose associated joint density is denoted by f^0 .

Instead of calculating the prediction directly out of the expression in the proposition, we will compute the prediction step by step following the steps in the proof of the proposition. From the data, we first construct the sets of outcomes that are consistent with the data: $R(x_1) = \{(1,1), (1,0)\}$, $R(x_2) = \{(1,1), (0,1)\}$, $R(x_3) = \{(1,0)\}$. Based on the prior q^0 , we know that $R(x_1)$ happens with probability $p_{(1,1)}^0 + p_{(1,0)}^0$, $R(x_2)$ happens with probability $p_{(1,1)}^0 + p_{(0,1)}^0$ and

$R(x_3)$ happens with probability $p_{(1,0)}^0$. Thus, the probability that this data is generated is

$$\begin{aligned}
P[X] &= \int_{p \in [0,1]^3} (p_{(1,0)}(p_{(1,1)} + p_{(0,1)})(p_{(1,1)} + p_{(1,0)})) f^0(p) dp \\
&= \mathbb{E}_{f^0} [p_{(1,0)} p_{(1,1)}^2 + p_{(1,0)}^2 p_{(1,1)} + p_{(1,0)} p_{(0,1)} p_{(1,1)} + p_{(1,0)}^2 p_{(0,1)}] \\
&= \mathbb{E}_{f^0} [p_{(1,0)} p_{(1,1)}^2] + \mathbb{E}_{f^0} [p_{(1,0)}^2 p_{(1,1)}] + \mathbb{E}_{f^0} [p_{(1,0)} p_{(0,1)} p_{(1,1)}] + \mathbb{E}_{f^0} [p_{(1,0)}^2 p_{(0,1)}] \\
&= \frac{\Gamma(4)}{\Gamma(7)\Gamma(1)} (\Gamma(2)\Gamma(3) + \Gamma(3)\Gamma(2) + \Gamma(2)\Gamma(2)\Gamma(2) + \Gamma(3)\Gamma(2)) \\
&= \frac{7}{120}.
\end{aligned}$$

Similarly, letting $R_2(x_1) = \{(1, 1)\}$,

$$\begin{aligned}
P[\bar{X}] &= \int_{p \in [0,1]^3} (p_{(1,0)}(p_{(1,1)} + p_{(0,1)})p_{(1,1)}) f^0(p) dp \\
&= \mathbb{E}_{f^0} [p_{(1,0)} p_{(1,1)}^2] + \mathbb{E}_{f^0} [p_{(1,0)} p_{(0,1)} p_{(1,1)}] \\
&= \frac{\Gamma(4)}{\Gamma(7)\Gamma(1)} (\Gamma(2)\Gamma(3) + \Gamma(2)\Gamma(2)\Gamma(2)) \\
&= \frac{3}{120}.
\end{aligned}$$

Thus, the predicted probability of user 3's liking item 2 is given by $\frac{3}{7}$.

Table C2 demonstrates predictions on x_{21} based on different historical ratings. The uniform prior is assumed to derive the predictions.

	user 1	user 2	user 3		user 1	user 2	user 3
item 1	1	$\emptyset$	1	item 1	1	$\emptyset$	0
item 2	3/11	0	0	item 2	3/7	0	0

Table C2: Predictions with partial data

Proof. We first find the probability that the data X happens given the prior distribution. Consider a review vector (possibly incomplete) x_j . For a given prior q^0 , we have

$$P[x_j] = \sum_{r \in R(x_j)} q_r^0.$$

Since the review of j is independent of j' conditional on the true parameter, the probability of X can be represented as

$$P[X] = \int \prod_{j=1}^N \left(\sum_{r \in R(x_j)} q_r^0 \right) f^0(q^0) dq^0,$$

where f^0 is a joint distribution of q^0 according to the Dirichlet prior distribution. Changing the

order of the multiplication and the summation,

$$\begin{aligned}
P[X] &= \int \prod_{j=1}^N \left(\sum_{r \in R(x_j)} q_r^0 \right) f^0(q^0) dq^0 \\
&= \sum_{\gamma_j \in R(x_j), j \in \bar{N}} \int \prod_{k=1}^N q_{\gamma_k}^0 f^0(q^0) dq^0 \\
&= \sum_{\gamma_j \in R(x_j), j \in \bar{N}} \mathbb{E}_{f^0} \left[\prod_{k=1}^N q_{\gamma_k}^0 \right].
\end{aligned}$$

Note that the product moments of a Dirichlet random variable has a representation using Gamma functions.¹² Using the representation, we have

$$\begin{aligned}
P[X] &= \sum_{\gamma_j \in R(x_j), j \in \bar{N}} \frac{\Gamma(\sum_{r \in R} \alpha_r^0)}{\Gamma(\sum_{r \in R} (\alpha_r^0 + N(r|\{\gamma_j\}_{j \in \bar{N}})))} \prod_{r \in R} \frac{\Gamma(\alpha_r^0 + N(r|\{\gamma_j\}_{j \in \bar{N}}))}{\Gamma(\alpha_r^0)} \\
&= \sum_{\gamma_j \in R(x_j), j \in \bar{N}} \frac{\Gamma(\sum_{r \in R} \alpha_r^0)}{\Gamma(N + \sum_{r \in R} \alpha_r^0)} \prod_{r \in R} \frac{\Gamma(\alpha_r^0 + N(r|\{\gamma_j\}_{j \in \bar{N}}))}{\Gamma(\alpha_r^0)} \\
&= \frac{\Gamma(\sum_{r \in R} \alpha_r^0)}{\Gamma(N + \sum_{r \in R} \alpha_r^0)} \prod_{r \in R} \Gamma(\alpha_r^0) \sum_{\gamma_j \in R(x_j), j \in \bar{N}} \prod_{r \in R} \Gamma(\alpha_r^0 + N(r|\{\gamma_j\}_{j \in \bar{N}})).
\end{aligned}$$

Similarly, we can compute $P[\bar{X}]$ and it is given by

$$P[\bar{X}] = \frac{\Gamma(\sum_{r \in R} \alpha_r^0)}{\Gamma(N + \sum_{r \in R} \alpha_r^0)} \prod_{r \in R} \Gamma(\alpha_r^0) \sum_{\gamma_1 \in R_{C+1}^+(x_1), \gamma_j \in R(x_j), j \neq n} \prod_{r \in R} \Gamma(\alpha_r^0 + N(r|\{\gamma_j\}_{j \in \bar{N}})).$$

Now, by the conditional probability formula, we have

$$P[\bar{X}|X] = \frac{\sum_{\gamma_1 \in R_{C+1}^+(x_1), \gamma_j \in R(x_j), j \neq 1} \prod_{r \in R} \Gamma(\alpha_r^0 + N(r|\{\gamma_j\}_{j \in \bar{N}}))}{\sum_{\gamma_j \in R(x_j), j \in \bar{N}} \prod_{r \in R} \Gamma(\alpha_r^0 + N(r|\{\gamma_j\}_{j \in \bar{N}}))}.$$

■

D Finite data analysis: Welfare analysis

In the paper, we mainly focused on the asymptotic value that a recommender system creates. The value when there are a finite number of data points available depends on the recommender system that is in use. In this section, building upon the Bayesian recommender system we proposed in Section A, we study the value a recommender system that learns from a finite dataset offers to users.

¹²If $Z \sim Dir(w_1, \dots, w_m)$, $\mathbb{E}[\prod_{i=1}^m Z_i^{n_i}] = \frac{\Gamma(w_1 + \dots + w_m)}{\Gamma(w_1 + n_1 + \dots + w_m + n_m)} \prod_{i=1}^m \frac{\Gamma(w_i + n_i)}{\Gamma(w_i)}$.

The following proposition presents the expected utility to users that the Bayesian recommender system offers under the user-optimal threshold. The user welfare under a different threshold can immediately be derived using the proposition. For simplicity, we focus on the case of $I = 1$. Consider complete data X . We can find y and r' such that y records the occurrences of r in X' as defined in the paper and r' corresponds to the target user's history. As the learning model is invariant to the order of the elements in X' , y and r' characterize X .

Proposition OA 3 (Finite sample data) *When there are $N - 1$ previous users, the value to the target user from item $C + 1$ can be characterized as follows:*

1. *Ex-post value: Given y and r' , the expected utility to a history r' user is*

$$\left(v_1 \frac{p(r',1)}{p_{r'}} + v_0 \frac{p(r',0)}{p_{r'}} \right) \mathbf{1}[v_1(y_{(r',1)} + \alpha_{(r',1)}^0) + v_0(y_{(r',0)} + \alpha_{(r',0)}^0) \geq 0].$$

2. *Interim value: The expected utility to a history r' user is*

$$\left(v_1 \frac{p(r',1)}{p_{r'}} + v_0 \frac{p(r',0)}{p_{r'}} \right) \sum_{k=0}^{N-1} \binom{N-1}{k} (1 - p_{r'})^{N-1-k} p_{r'}^k \sum_{j=\lfloor \tau \rfloor + 1}^k \binom{k}{j} \left(\frac{p(r',1)}{p_{r'}} \right)^j \left(\frac{p(r',0)}{p_{r'}} \right)^{k-j}.$$

3. *Ex-ante value: The expected utility to a user is*

$$\sum_{r' \in R'} (v_1 p(r',1) + v_0 p(r',0)) \sum_{k=0}^{N-1} \binom{N-1}{k} (1 - p_{r'})^{N-1-k} p_{r'}^k \sum_{j=\lfloor \tau \rfloor + 1}^k \binom{k}{j} \left(\frac{p(r',1)}{p_{r'}} \right)^j \left(\frac{p(r',0)}{p_{r'}} \right)^{k-j},$$

$$\text{where } \tau = \frac{-v_1 \alpha_{(r',1)} - v_0 (k + \alpha_{(r',0)})}{v_1 - v_0}.$$

The derivation of the above expressions proceeds as follows. The target item is recommended to a user with history r' if and only if the consumption of the item is more likely to induce positive utility than not. That is, when the prediction is $\hat{x}$, the recommendation is made if and only if

$$v_1 \hat{x} + v_0 (1 - \hat{x}) \geq 0.$$

This condition is equivalent to requiring $v_1(y_{(r',1)} + \alpha_{(r',1)}^0) + v_0(y_{(r',0)} + \alpha_{(r',0)}^0) \geq 0$. For example, when $(v_1, v_0) = (1, -1)$ and $(\alpha_{(r',1)}^0, \alpha_{(r',0)}^0) = (1, 1)$, the item is recommended if and only if $y_{(r',1)} \geq y_{(r',0)}$ is observed in the data. Since the recommendation is unbiased, the user tries the item whenever she receives the recommendation. Once she tries the item, she receives $v_1 \frac{p(r',1)}{p_{r'}} + v_0 \frac{p(r',0)}{p_{r'}}$. The ex-ante value is derived from the interim value taking into account that the new user has a history r' with probability $p_{r'}$. From the ex-ante perspective, the probability that the item is

recommended to the user with r' history is

$$\sum_{k=0}^{N-1} \binom{N-1}{k} (1 - p_{r'})^{N-1-k} p_{r'}^k \sum_{j=\lfloor \tau \rfloor + 1}^k \binom{k}{j} \left(\frac{p(r',1)}{p_{r'}} \right)^j \left(\frac{p(r',0)}{p_{r'}} \right)^{k-j}.$$

It is the sum of probabilities that k out of $N - 1$ previous users have history r' , and among the k users, j users have positive experiences with the item. Here, we sum only the cases in which j exceeds the threshold level for a recommendation, namely τ .

E Harmful customization

In relation to Proposition 2 of the main text, we identify two situations under which an extra degree in customization harms the history r' user's surplus.

Proposition OA 4 *Let τ be the threshold. A history r' target user is strictly worse off from an additional degree in customization if and only if one of the following is true*

$$\tau^u < \frac{p(r',0,1)}{p(r',0)} < \tau \leq \frac{p(r',1,1) + p(r',0,1)}{p(r',1) + p(r',0)} \text{ or}$$

$$\tau^u > \frac{p(r',1,1)}{p(r',1)} \geq \tau > \frac{p(r',1,1) + p(r',0,1)}{p(r',1) + p(r',0)}.$$

Proof. We compare user utilities before and after item $C + 1$ is added to the system. As before, let r' denote the ratings over C conditioning items. Using Lemma 1 in the Appendix of the main paper, the history r' user's utility before item $C + 1$ is added is

$$\mathbf{1} \left\{ \frac{p(r',1,1) + p(r',0,1)}{p(r',1) + p(r',0)} \geq \frac{-w_0}{w_1 - w_0} \right\} \left(v_1 \frac{p(r',1,1)}{p_{r'}} + v_0 \frac{p(r',1,0)}{p_{r'}} + v_1 \frac{p(r',0,1)}{p_{r'}} + v_0 \frac{p(r',0,0)}{p_{r'}} \right).$$

On the other hand, the user utility after the addition of the item is

$$\begin{aligned} & \frac{p(r',1)}{p_{r'}} \mathbf{1} \left\{ \frac{p(r',1,1)}{p(r',1)} \geq \frac{-w_0}{w_1 - w_0} \right\} \left(v_1 \frac{p(r',1,1)}{p(r',1)} + v_0 \frac{p(r',1,0)}{p(r',1)} \right) \\ & + \frac{p(r',0)}{p_{r'}} \mathbf{1} \left\{ \frac{p(r',0,1)}{p(r',0)} \geq \frac{-w_0}{w_1 - w_0} \right\} \left(v_1 \frac{p(r',0,1)}{p(r',0)} + v_0 \frac{p(r',0,0)}{p(r',0)} \right). \end{aligned}$$

We will only deal with the case when $p(r',1,1) \geq p(r',0,1)$ as the exact same logic can be applied to the other case. In this case, we have

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

Thus, if the item is recommended to the user in the system without item $C + 1$, it will also be recommended to the history $(r', 1)$ user in the system with item $C + 1$. Conversely, if the item is

not recommended to the user in the system without item $C + 1$, it will not be recommended to the user with history $(r', 0)$ in the system with item $C + 1$. Therefore, there are two cases when the extra degree in customization strictly hurts the user. The first is

$$\tau^u < \frac{p_{(r',0,1)}}{p_{(r',0)}} < \tau^p \leq \frac{p_{(r',1,1)} + p_{(r',0,1)}}{p_{(r',1)} + p_{(r',0)}}.$$

That is, the item is recommended to the target user in the system without item $C + 1$, but in the system with item $C + 1$ it is only recommended to the users whose history is $(r', 1)$ even though it is expected to generate positive utility to history $(r', 0)$ users.

On the other hand, there is also a case that the item is recommended to history $(r', 1)$ users under the system with item $C + 1$ even though it is not expected to generate positive utility to the users and the system without item $C + 1$ does not recommend the item to users. This case arises when

$$\tau^u > \frac{p_{(r',1,1)}}{p_{(r',1)}} \geq \tau^p > \frac{p_{(r',1,1)} + p_{(r',0,1)}}{p_{(r',1)} + p_{(r',0)}}.$$

■