

Online discrete choice models

Applications in personalized recommendations

Danaf, Mazen; Becker, Felix; Song, Xiang; Atasoy, Bilge; Ben-Akiva, Moshe

DOI

[10.1016/j.dss.2019.02.003](https://doi.org/10.1016/j.dss.2019.02.003)

Publication date

2019

Document Version

Final published version

Published in

Decision Support Systems

Citation (APA)

Danaf, M., Becker, F., Song, X., Atasoy, B., & Ben-Akiva, M. (2019). Online discrete choice models: Applications in personalized recommendations. *Decision Support Systems*, 119, 35-45.
<https://doi.org/10.1016/j.dss.2019.02.003>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

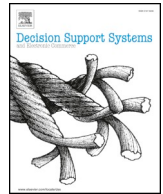
Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Green Open Access added to TU Delft Institutional Repository

'You share, we take care!' - Taverne project

<https://www.openaccess.nl/en/you-share-we-take-care>

Otherwise as indicated in the copyright section: the publisher is the copyright holder of this work and the author uses the Dutch legislation to make this work public.



Online discrete choice models: Applications in personalized recommendations

Mazen Danaf^{a,*}, Felix Becker^b, Xiang Song^b, Bilge Atasoy^c, Moshe Ben-Akiva^d

^a Massachusetts Institute of Technology, Department of Civil and Environmental Engineering, 77 Massachusetts Avenue, Room 1-181, Cambridge, MA 02139, United States of America

^b Massachusetts Institute of Technology, Department of Civil and Environmental Engineering, United States of America

^c Delft University of Technology, Department of Maritime and Transport Technology, Netherlands

^d Massachusetts Institute of Technology, Department of Civil and Environmental Engineering, Room 181, United States of America

ARTICLE INFO

Keywords:

Personalization
Intra-consumer heterogeneity
Hierarchical Bayes
Preference updates, recommender systems

ABSTRACT

This paper presents a framework for estimating and updating user preferences in the context of app-based recommender systems. We specifically consider recommender systems which provide personalized menus of options to users. A Hierarchical Bayes procedure is applied in order to account for inter- and intra-consumer heterogeneity, representing random taste variations among individuals and among choice situations (menus) for a given individual, respectively. Three levels of preference parameters are estimated: population-level, individual-level and menu-specific. In the context of a recommender system, the estimation of these parameters is repeated periodically in an offline process in order to account for trends, such as changing market conditions. Furthermore, the individual-level parameters are updated in real-time as users make choices in order to incorporate the latest information from the users. This online update is computationally efficient which makes it feasible to embed it in a real-time recommender system. The estimated individual-level preferences are stored for each user and retrieved as inputs to a menu optimization model in order to provide recommendations. The proposed methodology is applied to both Monte-Carlo and real data. It is observed that the online update of the parameters is successful in improving the parameter estimates in real-time. This framework is relevant to various recommender systems that generate personalized recommendations ranging from transportation to e-commerce and online marketing, but is particularly useful when the attributes of the alternatives vary over time.

1. Introduction

Personalization has gained increasing interest among researchers and practitioners in the past two decades. Greater ease in data collection about users has made it possible for service providers to recommend items, services, and content in a non-intrusive way [1] through online recommendations. The conventional recommendation techniques, which mainly rely on item and user profiles, produce ratings that do not take full advantage of the available data. On the other hand, discrete choice models, which have been rarely used in online recommendations, integrate item specific, user specific, and contextual data in a single model [11].

According to Jiang et al. [23], the use of discrete choice models in recommender systems can address some limitations associated with the standard recommendation approaches. The first limitation is the tradeoff between relevancy and diversity [23,38]. The second limitation is

that both metrics (relevancy and diversity), which are commonly used to measure the degree of matching, do not necessarily explain user preferences. On the other hand, discrete choice models directly measure the individual-specific utility of an alternative (or a set of alternatives) as a function of its attributes (without the need to measure relevancy or diversity separately). Finally, and unlike most standard recommendation techniques, discrete choice models can be applied even when the universal set from which alternatives are recommended and the alternative attributes vary over time. This is because the utility of each alternative is represented as a function of its attributes. For example, in travel recommendations, the travel time, cost, and availability of the different alternatives might vary over time.

This paper presents a methodology for estimating discrete choice models online, which can be used in updating user preferences continuously in an app-based setting such as recommender systems. The framework presented in this paper utilizes the Hierarchical Bayes (HB)

* Corresponding author.

E-mail addresses: mdanaf@mit.edu (M. Danaf), fbecker@mit.edu (F. Becker), bensong@mit.edu (X. Song), b.atasoy@tudelft.nl (B. Atasoy), mba@mit.edu (M. Ben-Akiva).

<https://doi.org/10.1016/j.dss.2019.02.003>

Received 1 June 2018; Received in revised form 11 February 2019; Accepted 14 February 2019

Available online 22 February 2019

0167-9236/ © 2019 Elsevier B.V. All rights reserved.

estimator proposed by Becker et al. [6] and Ben-Akiva et al. [7] which accounts for inter- and intra-consumer heterogeneity. An offline-online procedure is proposed in which individual-specific parameters are updated after each choice without the need to re-estimate the whole model. Periodically, data from multiple individuals are pooled, and population level parameters are updated by re-estimating the model with the new data.

This paper addresses important gaps associated with using discrete choice models in recommender systems:

- **Online estimation:** Although discrete choice models have been used in some recommender systems [11,23], the applications were mostly offline because updating individual-level preferences requires re-estimating the entire model (which becomes computationally burdensome as the sample size, number of attributes, or number of alternatives increases). On the other hand, many online applications (such as recommender systems, personalized advertisement, etc.) require updating individual preferences in real-time. With infrequent preference updates, users' most recent choices might not be taken into consideration in generating personalized recommendations. The proposed online methodology therefore enables us to use the most up-to-date preferences for recommendations without computational constraints.
- **Advanced level of heterogeneity:** The existing few online applications of discrete choice models in recommender systems were based on multinomial or nested logit/probit models, which do not account for preference heterogeneity. Such models can only be used in non-personalized recommendations. On the other hand, logit mixture models (which account for heterogeneity) cannot be estimated in real-time because estimation requires integration over multi-dimensional distributions (in Maximum Likelihood Estimation), or drawing from complex posteriors (in Hierarchical Bayes methods). Applications of logit mixture models were also limited to inter-consumer heterogeneity, and assumed that preferences are stable over time. The proposed methodology accounts for more complex patterns of heterogeneity (inter- and intra-consumer heterogeneity), which improves the quality of predictions and recommendations [7,33].
- **Identification of individual preferences:** Other studies have calibrated choice models on the individual level. However, this method also has limitations since it requires a sufficiently large number of observations per user. When limited data per user is available, a good prior on the individual-specific parameters is needed. To the best of our knowledge, specifying a good prior has not been adequately addressed in the literature. In our proposed methodology, the online estimation procedure overcomes this issue since it is comparable to estimating models at the individual level, but with good priors which are obtained from the offline HB estimator.

In order to validate this methodology, Monte Carlo data on the choice of grapes and real stated preferences (SP) data on the choice of transport mode in Switzerland [9] are used. Individual preferences are estimated and updated using repeated observations, and then used in predicting the next choice and generating personalized recommendations. While our applications focus on personalized recommendations, this methodology allows discrete choice models to be applied online in various real-time applications and decision support systems such as personalized advertisement, real-time forecasting, personalized tolling, and others.

The remainder of this paper is organized as follows: [Section 2](#) presents an overview of online recommendations and recent applications of discrete choice models in this domain. [Section 3](#) presents the proposed methodology for estimating and updating user preferences online. [Section 4](#) presents an application of this methodology to Monte Carlo data. [Section 5](#) presents a similar application to real SP data. [Section 6](#) presents a discussion of the modeling approach and its applications in online recommendations, and [Section 7](#) concludes.

2. Background

2.1. Online recommendations

The goal of online recommendations is to suggest items of interest to a user from a much larger set in order to handle information overload [11,23,28]. Personalized recommender systems must deliver relevant and precise recommendations based on each user's tastes and preferences, which should be determined with minimal involvement from the user. Recommendations must also be delivered in real-time so users are able to act immediately [11].

According to Ansari et al. [5], online recommendations can make use of several information sources including the individual's expressed preferences or choices among different alternatives, preferences for product attributes, other people's preferences or choices, expert judgments, and individual characteristics that may predict these preferences and choices.

Collaborative filtering and content-based filtering are the two most popular recommendation techniques. Other techniques include knowledge-based and context-aware methods [28]. Collaborative filtering [14] provides recommendations to an individual based on overlapping interests with other individuals. In other words, it mimics 'word-of-mouth' recommendations. Content-based techniques match the attributes of the user profile against the attributes of an item [28]. These techniques make recommendations similar to those a given user has liked in the past [11]. Knowledge- or utility-based recommender systems base their recommendations on the computation of the utility of each item for the user [21]. These systems utilize previous knowledge about users, items, and the utility function [28]. Context-aware recommender systems (CARS) account for contextual information such as the user's knowledge level (e.g. expert user or beginner), the time a recommendation is requested, and the external context (e.g. proximity of restaurants to the user) [28]. Other techniques have been proposed that utilize traditional machine learning techniques such as support vector machines and latent class models [12] and multi-armed bandit methods [24,30].

Despite the significant advances in online recommendations, several theoretical and practical challenges have been identified. For example, Ziegler et al. [38] showed that the commonly used top-N lists do not necessarily map user satisfaction and utility. In some cases, measuring the (expected) utility of recommendations may be more important than measuring the accuracy of recommendations [16,38]. Another major challenge is that the commonly used recommendation techniques are designed to consider different configurations as different items [28]. Therefore, very few of these techniques can be applied when the universal set from which items are recommended varies over time (as in the case of travel advisers, where the attributes of alternatives such as time and cost vary over time).

Using discrete choice models in personalized recommendations overcomes many of the limitations mentioned above. First, these models represent utility as a function of the attributes of items (or alternatives), and the individual preferences towards each of these attributes. Therefore, utility is not inferred from measures of similarity obtained from item or user profiling. Second, since utility is modeled as a function of attributes, this method is able to handle cases where new items (with known attributes) could be recommended (e.g. items that have not been chosen or rated before), and cases where the attributes vary over time. The researcher decides on the specification of the utility functions, which may include the attributes, the individual preferences for attributes, contextual variables, and individual characteristics, thus making use of all the available data. Third, since the users' preferences are inferred from their previous choices, this reduces the burden on users because they are not required to rate or evaluate any items. Finally, this method is able to deal with diversification and the exploration-exploitation problem using simple extensions described in [Section 6.2](#).

2.2. Econometric and discrete choice models in recommender systems

Discrete choice models are often used to predict choices on an aggregate level. More recently these models have been utilized in recommender systems due to their ability to predict individual choices [11,23,27].

Chaptini [11] utilized discrete choice models to predict choices on the individual level and provide personalized recommendations. He developed an online academic advisor for MIT students that recommends academic courses based on observed and latent attributes of the courses (e.g. difficulty, workload, overall impression, etc.). These attributes were expressed as functions of students' characteristics (such as gender, degree program, etc.). The model was estimated using maximum likelihood estimation with data collected via an online revealed preferences (RP)/stated preferences (SP) survey. He then conditioned on the individual choices to find individual-level parameters that were used in generating course recommendations. In this study, preferences were estimated offline for each student and not updated as more choices were observed. In addition, the behavioral model accounted for inter-consumer heterogeneity only (and ignored intra-consumer heterogeneity).

Jiang et al. [23] used discrete choice models to measure users' preferences towards an entire recommendation list. The goal was to identify a recommendation list with the highest choice probability. A multi-level nested multinomial logit model was proposed, and the recommendation problem was formulated as a nonlinear binary integer programming problem. The authors noted that unlike typical recommender systems, discrete choice models introduce product diversity in the proposed recommendations. The main limitation of this approach was the lack of personalization, since a nested logit model was used (this model can be estimated at the individual level only if a large number of choices per individual is available).

Rubin and Steyvers [29] introduced a probabilistic model of the process by which an individual selects and later rates an item. This model was applied to movie rating data collected by Netflix. A Latent Dirichlet Allocation (LDA) model was used to model the probability of selecting a movie given a set of movies classified by topic. An ordered logit model was used to model movie ratings. This model included an individual-specific bias term which determines the general tendency of a user to give favorable ratings, however, it did not account for heterogeneity in the other parameters (preferences). In addition, since all parameters are learnt through Markov-Chain Monte Carlo methods, this model can only be run offline.

Ansari et al. [5] proposed a Hierarchical Bayes approach for a recommender system that accounts for unobserved heterogeneity in user preferences, unobserved product heterogeneity and attributes (such as holistic consumer judgements and product appeal structures), and expert judgements. Customer ratings were modeled as a function of product attributes, customer characteristics, and expert evaluations. User preferences were expressed as a function of fixed effects (i.e. observed customer and movie variables and their interactions) and random effects pertaining to the customer. The model was applied to movie recommendations on the internet and estimated using MCMC. The main advantage of this paper was accounting for various sources of information (i.e. movie genres, expert evaluations, and socio-demographic characteristics). The authors also suggested different extensions which are included in this paper. First, preferences can be learnt from implicit rather than explicit information (i.e. revealed preferences or actual choices). Second, more complex forms of heterogeneity can be considered. This model also cannot be estimated online because of the excessive running times.

Our methodology extends the abovementioned studies by estimating models that account for personalization, and yet can be estimated online after each choice. Our model also accounts for complex forms of user heterogeneity, and uses only implicit data (observed choices) in order to estimate and update user preferences.

2.3. User heterogeneity and personalization

According to Castells et al. [10], user preferences are complex, dynamic, context-dependent, heterogeneous, and even contradictory. Therefore, accounting for consumer heterogeneity is crucial in recommender systems. Most of the methods mentioned earlier account, either directly or indirectly, for inter-consumer heterogeneity. On the other hand, limited research has been done on intra-consumer taste heterogeneity, representing taste variation among different choices done by the same individual. For example, in travel recommendations (such as the one presented in Section 5), the same user might be more or less sensitive to travel time depending on various unobserved factors specific to the particular choice situation, such as his/her schedule, the trip purpose, weather conditions, etc.

According to Ben-Akiva et al. [7], ignoring intra-consumer heterogeneity assumes a nearly neoclassical consumer with “permanent” individual preferences. Perturbations in these preferences are treated as nuisance factors. In the presence of multiple observations from each individual, it is possible to identify inter- and intra-consumer heterogeneity. In the context of discrete choice models, excluding intra-consumer heterogeneity when its effect is significant will result in biases due to a greater degree of unobserved effects [7].

Models with inter- and intra-consumer heterogeneity have been estimated by Hess and Train [20], Yáñez et al. [37] and Hess and Rose [19] using maximum simulated likelihood (MSL). These studies investigated taste variations among different choices done by the same individual, and demonstrated that accounting for such effects results in better estimates. However, these studies were mainly exploratory and limited to offline applications, and estimation was computationally burdensome. Becker et al. [6] and Ben-Akiva et al. [7] introduced a Hierarchical Bayes (HB) estimator for such models by extending the standard HB procedure for logit mixture [2,36]. This model significantly reduces the computation time compared to the previously used MSL estimators. All of these studies presented methods to estimate choice models with inter- and intra-consumer heterogeneity offline, and did not address online applications.

In the following sections, a novel framework is proposed that utilizes discrete choice models in estimating and updating individual level preferences in an online setting, building on the HB estimator proposed by Becker et al. [6] and Ben-Akiva et al. [7]. This framework can be used in various applications, but is particularly useful in recommender systems. The estimated preferences account for both inter- and intra-consumer heterogeneity, and are updated in real-time after each choice. They can serve as input to an assortment optimization algorithm, which recommends personalized menus to users by maximizing an objective function (e.g. the probability of choosing an alternative from the menu, the expected revenue of the menu, etc.).

3. Methodology

This section explains the methodology for estimating and continuously updating population-level and individual-level preferences. The Hierarchical Bayes estimator of a logit mixture model with inter- and intra-consumer taste and scale heterogeneity proposed by Becker et al. [6] is used in order to estimate these preferences. Inter- and intra-consumer heterogeneity are used to improve the estimation results, and thus the predictive capabilities of the choice models [7]. Individual-specific coefficients, which can be extracted from the estimation procedure, are used for personalization.

3.1. Estimating preferences

We consider the case whereby individual n ($n = 1, 2, \dots, N$) is presented with a menu m ($m = 1, 2, \dots, M_n$) and makes a choice among a set of alternatives ($j = 1, 2, \dots, J_{nm}$). Thus, each menu refers to a *choice situation*. The total number of individuals is N and the total

number of menus presented to each individual is M_n . The number of parameters to be estimated is denoted by T .

In order to estimate user preferences, we use the HB estimator proposed by Becker et al. [6], which extends the widely used 3-step HB estimator of logit mixture [36] to a 5-step estimator in order to account for intra-consumer heterogeneity.

We assume the utility specification of choice j in menu m presented in Eq. (1):

$$U_{jmn} = \frac{1}{\exp(\alpha_{mn})} (-P_{jmn} + X_{jmn}\eta_{mn}) + \epsilon_{jmn} \quad (1)$$

where P_{jmn} is the price of alternative j in menu m faced by individual n (with its coefficient fixed to -1), U_{jmn} is individual n 's unobserved utility of alternative j in menu m , X_{jmn} represents a vector of individual characteristics and alternative attributes, η_{mn} represents a vector of coefficients/preferences, α_{mn} is a scale parameter for individual n in menu m , and ϵ_{jmn} is an error term following the extreme value distribution. The subscripts (mn) in η_{mn} and α_{mn} indicate that these coefficients might vary among individuals and among choice situations of the same individual respectively.

The model uses the Willingness-to-Pay space notation defined by Ben-Akiva et al. [7], whereby the price coefficient is fixed to -1 . Therefore, all other coefficients represent the willingness-to-pay for the corresponding attributes. Since the price coefficient is fixed, the scale parameter α_{mn} can be estimated.

We start by defining three levels of parameters needed to account for both inter- and intra-consumer heterogeneity as proposed by Ben-Akiva et al. [7]:

1. Population-level parameters μ and Ω_b : represent the average tastes/preferences in the population and the inter-consumer covariance matrix respectively.
2. Individual-level parameters ζ_n and Ω_w : represent the average tastes/preferences of a specific individual and the intra-consumer covariance matrix respectively.
3. Menu-specific parameters η_{mn} : reflect the tastes/preferences specific to each choice situation.

We assume that ζ_n and η_{mn} are normally distributed:

$$\eta_{mn} \sim \mathcal{N}_T(\zeta_n, \Omega_w) \quad (2)$$

$$\zeta_n \sim \mathcal{N}_T(\mu, \Omega_b) \quad (3)$$

The probability of a sequence of choices (d_n) made by individual n can be expressed as:

$$P(d_n | \mu, \Omega_b, \Omega_w) = \int_{\zeta_n} \prod_{m=1}^{M_n} \left[\int_{\eta_{mn}} \prod_{j=1}^{J_{mn}} P_j(\eta_{mn})^{d_{jmn}} H(d\eta_{mn} | \zeta_n, \Omega_w) \right] F(d\zeta_n | \mu, \Omega_b) \quad (4)$$

where d_{jmn} is equal to one if individual n chooses alternative j in menu m and zero otherwise, and:

$$P_j(\eta_{mn}) = \frac{\exp(V_{jmn}(\eta_{mn}))}{\sum_{j'=1}^{J_{mn}} \exp(V_{j'mn}(\eta_{mn}))} \quad (5)$$

$$H(d\eta_{mn} | \zeta_n, \Omega_w) \sim \mathcal{N}_T(\zeta_n, \Omega_w) \quad (6)$$

$$F(d\zeta_n | \mu, \Omega_b) \sim \mathcal{N}_T(\mu, \Omega_b) \quad (7)$$

The posterior distribution is presented in Eq. (8):

$$K(\mu, \zeta_n \forall n, \eta_{mn} \forall mn, \Omega_b, \Omega_w | d)$$

$$\propto \prod_{n=1}^N \left[\prod_{m=1}^{M_n} \left[\prod_{j=1}^{J_{mn}} [P_j(\eta_{mn})^{d_{jmn}}] h(\eta_{mn} | \zeta_n, \Omega_w) \right] f(\zeta_n | \mu, \Omega_b) \right] k(\Omega_w) k(\mu) \quad (8)$$

where:

$$k(\mu) \sim \mathcal{N}_T(\mu_0, A) \quad (9)$$

$$k(\Omega_b) \sim \text{IW}(T, I) \quad (10)$$

$$k(\Omega_w) \sim \text{IW}(T, I) \quad (11)$$

μ_0 represents a vector of prior means, A is a diagonal covariance matrix with diagonal values $\rightarrow \infty$ (uninformative prior), T is the number of unknown parameters, I is the T -dimensional identity matrix, and $\text{IW}(T, I)$ represent an Inverse Wishart distribution with T degrees of freedom and parameter I .

The model is estimated using the five-step Gibbs sampling procedure proposed by Ben-Akiva et al. [7] and Becker et al. [6]. This procedure is explained below:

Step I: drawing from the population means by drawing from the conditional posterior:

$$K(\mu | \zeta_n \forall n, \eta_{mn} \forall mn, \Omega_w, \Omega_b) \propto f(\zeta_n \forall n | \mu, \Omega_b) k(\mu) \quad (12)$$

The conditional posterior on μ is $\mathcal{N}\left(\bar{\zeta}^{i-1}, \frac{\Omega_b^{i-1}}{N}\right)$ (where i is an iteration index) where:

$$\bar{\zeta} = \frac{1}{N} \sum_n \zeta_n^{i-1} \quad (13)$$

Step II: drawing from the population-level covariance matrix by drawing from the conditional posterior:

$$K(\Omega_b | \mu, \zeta_n \forall n, \eta_{mn} \forall mn, \Omega_w) \propto f(\zeta_n \forall n | \mu, \Omega_b) k(\Omega_b) \quad (14)$$

The conditional posterior on Ω_b is Inverted Wishart with $T + N$ degrees of freedom and parameter $TI + N\bar{V}_b$, where T is the number of unknown parameters, I is the T -dimensional identity matrix, and:

$$\bar{V}_b = \frac{1}{N} \sum_{n=1}^N (\zeta_n^{i-1} - \mu^i)(\zeta_n^{i-1} - \mu^i)' \quad (15)$$

Step III: drawing from the individual-level covariance matrix by drawing from the conditional posterior:

$$K(\Omega_w | \mu, \zeta_n \forall n, \eta_{mn} \forall mn, \Omega_b) \propto h(\eta_{mn} \forall mn | \zeta_n \forall n, \Omega_w) k(\Omega_w) \quad (16)$$

Given η_{mn}^{i-1} and ζ_n^{i-1} for all n , the conditional posterior on Ω_w is Inverted Wishart with degrees of freedom $T + M_t$ and parameter $\frac{TI + M_t \bar{V}_w}{T + M_t}$, where M_t represents the total number of menus faced by all individuals, and:

$$\bar{V}_w = \frac{1}{M_t} \sum_{n=1}^N \sum_{m=1}^{M_n} (\eta_{mn}^{i-1} - \zeta_n^{i-1})(\eta_{mn}^{i-1} - \zeta_n^{i-1})' \quad (17)$$

In this step, we assume a single covariance matrix for all individuals. Due to the potentially small number choice situations faced by each individual in a typical recommender system, it might not be possible to estimate an individual-specific covariance matrix.

Step IV: drawing from the individual-level means by drawing from the conditional posterior:

$$K(\zeta_n | \mu, \eta_{mn} \forall mn, \Omega_b, \Omega_w) \propto h(\eta_{mn} \forall mn | \zeta_n \forall n, \Omega_w) f(\zeta_n | \mu, \Omega_b) \quad (18)$$

Using $N(\mu, \Omega_b)$ as a prior for ζ_n , the conditional posterior is $N(\zeta_n, \Sigma_{\zeta_n})$ where:

$$\bar{\zeta}_n = ([\Omega_b^i]^{-1} + M_n [\Omega_w^i]^{-1})^{-1} \left([\Omega_b^i]^{-1} \mu_i + M_n [\Omega_w^i]^{-1} \frac{1}{M_n} \sum_{m=1}^{M_n} \eta_{mn}^{i-1} \right) \quad (19)$$

and:

$$\Sigma_{\zeta_n} = ([\Omega_{i+1}^b]^{-1} + M_n [\Omega_{i+1}^w]^{-1})^{-1} \quad (20)$$

Step V: drawing from the individual- and menu-specific coefficients by drawing from the conditional posterior:

$$K(\eta_{mn} | \mu, \zeta_n, \Omega_b, \Omega_w) \propto \prod_{j=0}^{j_{mn}} [P_j(\eta_{mn})^{d_{jmn}}] h(\eta_{mn} \forall mn | \zeta_n, \Omega_w), \quad n = 1, 2, \dots, N, m = 1, 2, \dots, M_n \quad (21)$$

A draw of η_{mn}^i is obtained by the Metropolis-Hastings procedure.

This five-step procedure assumes that all coefficients have inter- and intra-consumer distributions. However, the estimator can account for coefficients with only inter-consumer heterogeneity, or coefficients without any heterogeneity by including two additional MH steps [6,7].

3.2. Updating preferences

These parameters are estimated and updated through two interacting and repeated steps: offline and online estimation procedures.

Offline Estimation: The offline estimation procedure updates all the parameters across three levels. Namely, data are pooled and all coefficients (μ , Ω_b , ζ_n , Ω_w , and η_{mn}) are updated to reflect the effects of all choices made by all individuals since the last update. This is performed periodically (e.g. overnight or once a week) as it is computationally expensive. Updating population-level coefficients accounts for population trends when estimating individual-level coefficients.

Online Estimation: The online estimation procedure updates users' preferences in real time as they make choices. The individual specific parameters (ζ_n and η_{mn}) are updated after every choice, assuming that the population parameters μ and Ω_b and the intra-consumer covariance matrix Ω_w are fixed. This update is computationally inexpensive, and it can be done for each individual at a time, i.e., when a choice is observed for a given individual, his/her parameters are updated only. The online procedure is executed by iterating steps IV and V of the 5-step Gibbs sampler only for all choices available after the last offline update.

Ideally, if we ignore the computational constraints, the 5-step offline procedure would be used to update individual preferences after each choice. In this procedure, Steps IV and V update the individual- and menu-specific preferences for each individual using the intra-consumer covariance matrix Ω_w^i and the inter-consumer distribution $\mathcal{N}_T(\mu^i, \Omega_b^i)$ as a prior as shown in Eq. (22). Conditional on the population-level parameters (μ , Ω_b , and Ω_w), obtaining draws from ζ_n and η_{mn} for each individual is done independently from all other individuals. Therefore, if draws from the population level parameters were available, the individual- and menu-specific coefficients could be updated separately for each individual by iterating steps IV and V.

Consequently, if we use a prior that is close to $\mathcal{N}_T(\mu^i, \Omega_b^i)$ and a covariance matrix that is close to Ω_w^i in the online procedure, we would obtain results that are similar to those obtained from the offline procedure. Since population level parameters are not expected to vary significantly between successive offline estimations (which is the key assumption in this methodology), these values can be obtained from the

last offline estimation and used as fixed values in the online estimation. Sections 4.3.1 and 5.2.2 illustrate that this method is able to provide very close results compared to the offline procedure as a benchmark.

Additionally, since an informative prior is used on the individual-specific parameters, the Markov Chains converge faster; stationarity is achieved quickly and a fewer number of draws is required in the online procedure (compared to the offline procedure). This procedure can also be implemented on the users' mobile phone in app-based settings.

The key assumption in this procedure is that the population level preferences μ , Ω_b , and Ω_w do not vary significantly between successive offline estimations. The frequency of offline estimations depends on how fast the population level preferences change over time. This might vary from one application to another, and even between different attributes within the same application. This can be mitigated by observing the population level parameters obtained from successive offline estimations and deciding on the frequency of these estimations accordingly.

In addition, since μ and Ω_b are used as priors in Step IV, their effect diminishes as more observations per individual are observed, as the individual specific means ζ_n get closer to their true values. With few choice observations from each individual, the model suffers from “shrinkage”, whereby individual-level preferences are shrunk towards population means. HB is defined as a “data borrowing” technique that stabilizes individual-level preferences for each individual using information not only from his/her past choices, but also from other individuals within the same data set [26]. Therefore, if the number of observations per individual is large, then deviations in the population-level parameters from their true values will have smaller effects on the individual-level preferences.

The frequency of offline estimations results in a tradeoff between the computational complexity of this estimation and the enhanced accuracy of the online procedure. Section 5.2.4 demonstrates that if the population level parameters are misspecified, then the predictions obtained from the online procedure will be inferior to those obtained from the offline procedure.

3.3. Personalized menu generation

The offline and online procedures result in updated individual- and population-level parameters. These parameters are used as inputs to an online optimizer that performs menu optimization to present the user with a personalized list of alternatives to choose from. The system architecture is presented in Fig. 1, which demonstrates how the online procedure uses the individual choices and the population level parameters obtained from the offline estimation (μ , Ω_b , and Ω_w) in order to update user preferences.

Personalized recommendations are generated using the menu optimization model proposed by Song et al. [32,33]. This model maximizes hit rate or consumer-surplus (CS) in the form of log-sum, subject to constraints specifying the maximum number of alternatives to be shown in a menu. Binary decision variables are defined for each alternative representing whether or not it is shown in the recommended menu. In the latter study, a Monte-Carlo experiment representing a smart mobility service showed that models with intra-consumer heterogeneity provide better menus (i.e., achieve higher hit-rates) compared to models with only inter-consumer heterogeneity.

4. Monte Carlo application

4.1. Data and model structure

The procedure described above is applied to Monte-Carlo CBC Grapes data [6,7]. The data assumes that 10,000 individuals are presented with eight menus, each including three different alternatives which are bunches of grapes with varying prices and attributes (presented in Table 1) and an opt-out alternative. The eight menus are

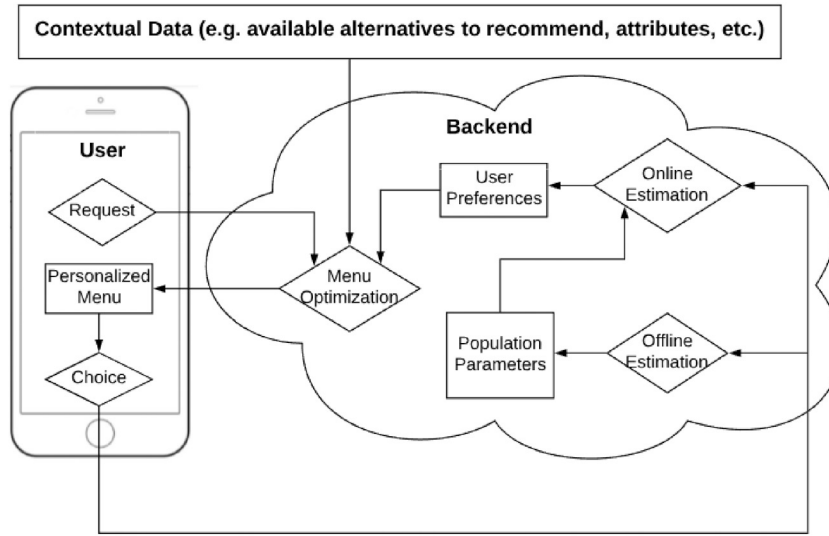


Fig. 1. System architecture.

Table 1
Grape CBC attributes and levels [7].

Attribute	Symbol	Levels
Price	P	\$1.00 to \$4.00
Sweetness	S	Sweet (1) or Tart (0)
Crispness	C	Crisp (1) or Soft (0)

assumed to be divided into three old choices (menus 1–3), four new choices (menus 4–7), and a test menu (menu 8). The goal is to update individual preferences in order to account for the new choices. The dependent variable is the choice between the three different bunches or not buying grapes at all. Both the data and the model are simplified (compared to [7]); only four coefficients are used, two of which are fixed and two have inter- and intra-consumer heterogeneity.

The utility equations (normalized to the opt-out alternative) are presented in Eq. (22):

$$U_{jmn} \equiv \frac{1}{\exp(\alpha)} (-P_{jmn} + S_{jmn}\beta_{s_{mn}} + C_{jmn}\beta_{c_{mn}} + B_{jmn}\beta_q) + \varepsilon_{jmn} \quad (22)$$

where:

- U_{jmn} represents the utility of alternative j in menu m faced by individual n .
- P_{jmn} is the price of bunch j in menu m faced by individual n , with its coefficient normalized to -1 .
- S_{jmn} and C_{jmn} represent sweetness and crispness of bunch j as indicated in Table 1, with coefficients $\beta_{s_{mn}}$ and $\beta_{c_{mn}}$ respectively. The subscript mn indicates that these coefficients have inter- and intra-consumer heterogeneity.
- B_{jmn} is a binary variable equal to one for all three bunches of grapes and zero for the opt-out alternative with coefficient β_q . This coefficient is fixed across all menus and individuals.
- α is a scale parameter, which is fixed across all menus and individuals.

Details on the data generation process, assumptions, true values, and estimates are included in Ben-Akiva et al. [7].

The true values of the population means for $\beta_{s_{mn}}$ and $\beta_{c_{mn}}$ are 1.0 and 0.3. The true values of the fixed coefficients β_q and α are 2.0 and -0.5 . Intra-consumer heterogeneity in the data is in the same order of magnitude as inter-consumer heterogeneity (all inter- and intra-consumer standard deviations are equal to 1.0 for sweetness and crispness).

The model is estimated for menus 1–7 using two procedures (the full offline procedure and the offline-online procedure), and the eighth menu is used for testing. In the full offline procedure, we iterate Steps I through V of the Gibbs sampling procedure in Section 3.1 on all 7 menus. In the offline - online procedure, we iterate steps I through V (offline procedure) for menus 1–3 and then iterate steps IV and V for the remaining menus (online procedure).

This experiment mimics a scenario in which three observations are initially observed from each individual. Individual- and population-level preferences are already estimated using these three observations (menus 1–3) by applying the five-step Gibbs sampler offline. Afterwards, four new observations are made by each individual. In order to update individual preferences to account for the new observations, either the full offline procedure or the online procedure can be used.

4.2. Analysis methods

In order to avoid overfitting, all the analyses are done using test data, which include the eighth choice. The analyses are based on the posterior predictive distribution (PPD) given by Eq. (23) and the conditional log-likelihood of the estimated parameters.

$$P(d_{jmn} = 1 | d_{m^*}) = \int_{\eta_{mn}} P_j(\eta_{mn}) K(d_{jmn} | d_{m^*}) \quad (23)$$

where d_{m^*} denotes choices from recent menus and $K(d_{jmn} | d_{m^*})$ is the posterior (marginal) distribution of menu-specific parameters. The predicted probability of the chosen alternative is defined as the mean of the posterior predictive distribution across all individuals and draws. In addition, 95% confidence intervals of the predicted probabilities are presented. The conditional log-likelihood of the test data is calculated using individual-specific parameters and distributions, and therefore is conditioned on the choices made by individuals.

On the other hand, in order to test the effect of personalization, the results are compared to those obtained by the standard “random coefficients” procedure (which does not allow for any personalization). This is done by generating draws from the unconditional distributions $\zeta_n \sim \mathcal{N}(\hat{\mu}, \hat{\Omega}_b)$ and $\eta_{mn} \sim \mathcal{N}(\zeta_n, \hat{\Omega}_w)$ respectively (where $\hat{\mu}$, $\hat{\Omega}_b$, and $\hat{\Omega}_w$ are estimates of μ , Ω_b , and Ω_w).

The predicted probability of the observed choice of individual n in the test data using the random coefficients approach can be calculated as shown in Eq. (24).

$$P(d_n^* | \mu, \Omega_b, \Omega_w) = \int_{\zeta_n} \int_{\eta_{mn}} P(d_n^* | \eta_{mn}) H(d\eta_{mn} | \zeta_n, \Omega_w) F(d\zeta_n | \mu, \Omega_b) \quad (24)$$

Since the five-step Gibbs sampler uses uninformative priors on μ , Ω_b , and Ω_w , the estimates and log-likelihood values obtained using this estimator are the same as those obtained using maximum simulated likelihood (MSL) since the posterior will be dominated by the likelihood [7,22,39]. Since we are replicating the MSL estimates using a Bayesian approach, we use the conditional log-likelihood on the test data as a measure of performance.

4.3. Estimation results

4.3.1. Application of the offline - online procedure

The model is estimated using 200,000 Gibbs iterations, 100,000 of which are used as burn-in draws while the remaining 100,000 are used for sampling from the posterior distributions. Individual-level and menu-level parameters are obtained directly from the MCMC (ζ_n and η_{mn} draws respectively).

The stationarity of the Markov chains obtained from two offline procedures (using menus 1–7 and 1–3) is verified using the Heidelberg-Welch test [40] and Gelman and Rubin's convergence diagnostic [41]. All Markov chains pass the tests at the 95% level of confidence. The estimation results are presented in Table 2.

The results presented in Table 3 indicate that the full offline procedure achieves the highest final log-likelihood values and probabilities of the chosen alternatives on the test data (menu 8) as expected. However, this procedure would be infeasible in real-time. Updating the sample level estimates is computationally expensive; for this Monte-Carlo experiment, the run time is approximately 12 h.

Alternatively, the results of the partial offline procedure (using menus 1–3) have lower log-likelihood values and predicted probabilities. However, the subsequent application of the online procedure increases the probability of the chosen alternative by approximately 1.5% and yields results that are very close to those obtained by the full offline procedure. The offline-online procedure is also feasible and efficient in real time because it can be applied to the individual making the choice only rather than the whole sample.

4.3.2. Benefits of individual-level parameters

As shown in Table 4, the non-personalized (unconditional) log-likelihood values and the predicted probabilities of the chosen alternative are inferior to the respective conditional values calculated using the posterior draws. In this example, using individual-level parameters improves the predicted probabilities of the observed choices by about 4–6% compared to the random coefficients procedure.

Table 2
Estimation results.

		Full offline (menus 1–7)		Offline on menus 1–3	
True value		Posterior mean	Std. dev	Posterior mean	Std. dev
Population mean					
Constant	2	2.004	0.011	2.015	0.016
Log(scale)	−0.5	−0.508	0.009	−0.510	0.014
Sweetness	1	1.004	0.015	0.989	0.020
Crispness	0.3	0.313	0.015	0.294	0.020
Inter-consumer std. dev					
Sweetness	1	0.961	0.017	0.954	0.029
Crispness	1	1.036	0.017	1.054	0.030
Intra-consumer std. dev					
Sweetness	1	0.997	0.024	0.972	0.039
Crispness	1	0.989	0.025	0.991	0.043

Table 3

Estimation results for the offline-online procedure.

Estimation procedure and menus	PPD mean	PPD confidence interval	Log-likelihood
Full offline (menus 1–7)	0.458	[0.456, 0.461]	−9955.2
Offline (menus 1–3)	0.442	[0.439, 0.446]	−10,335.6
Offline (menus 1–3) Online (menus 4–7)	0.458	[0.456, 0.461]	−9962.7

Table 4

Comparison between individual-specific and random coefficients.

Estimation procedure and menus	Random coefficients		Individual-specific coefficients	
	PPD mean	Log-likelihood	PPD mean	Log-likelihood
Full offline (menus 1–7)	0.397	−11,090.7	0.458	−9955.2
Offline (menus 1–3)	0.397 ^a	−10,959.5 ^a	0.442	−10,335.6
Offline (menus 1–3) Online (menus 4–7)			0.458	−9962.7

^a Using the non-personalized approach, the probabilities predicted for the test data using the offline-online procedure would be similar to those predicted using the partial offline procedure.

4.3.3. Applications in personalized recommendations

In this section, we assume that users are offered only one alternative from the test menu. In order to maximize consumer surplus, the alternative with the highest predicted probability is chosen. The choice between the recommended alternative and opting-out is then simulated. The hit-rate is defined as the probability of accepting the recommendation instead of choosing the opt-out alternative.

The simulated hit-rate with individual-specific parameters obtained from the full offline procedure is 67.0%. On the other hand, the simulated hit-rates obtained using the offline procedure with menus 1–3 only is 65.7%. However, accounting for the new choices using the online procedure raises the hit-rates back to 67.0%. On the other hand, using population-level parameters instead of individual-level parameters results in a simulated hit-rate of 62.5% even when we consider all 7 choices.

5. Real application: Swissmetro data

5.1. Data and model

The procedure described in Section 3 is also applied to the Swissmetro data set [9], with the dependent variable being the transportation mode choice. The data was collected in Switzerland on the trains between St. Gallen and Geneva in 1998. Each survey respondent was presented with 9 hypothetical choice tasks, each having three alternatives (private car, Swissmetro (SM), and train). The attributes of these modes include the travel cost (fuel and parking costs for private car and fares for Swissmetro and train), travel time for all three modes, and Swissmetro and train headway. Since multiple observations are available from each respondent, we can use the offline-online procedure to demonstrate how preferences are learnt as more choices are observed.

In this application, we consider the simplified utility equations presented in Eqs. (25)–(27). Since the cost coefficient is fixed to −1, all the estimated coefficients represent the willingness to pay for the corresponding attributes (i.e. the time coefficient represents the value of time). Consequently, a scale parameter (α_{mn}) is estimated.

Table 5
Estimation results.

Coefficient	Population mean		Inter-consumer Standard deviation		Intra-consumer Standard deviation	
	Posterior mean	Std. Dev.	Posterior mean	Std. Dev.	Posterior mean	Std. Dev.
ASC_{SM}	0.321	0.064	0.748	0.058	0.255	0.036
ASC_{Car}	0.574	0.071	1.286	0.050	0.091	0.045
Scale	−2.019	0.067	1.053	0.075	0.163	0.101
Travel time	0.179	0.037	0.912	0.029	0.024	0.016

$$U_{Car, nm} = (ASC_{Car, mn} - \exp(\eta_{mn}) \times Time_{Car, mn} - Cost_{Car, mn}) / \exp(\alpha_{mn}) + \epsilon_{Car, mn} \quad (25)$$

$$U_{SM, nm} = (ASC_{SM, mn} - \exp(\eta_{mn}) \times Time_{SM, mn} - Cost_{SM, mn}) / \exp(\alpha_{mn}) + \epsilon_{SM, mn} \quad (26)$$

$$U_{Train, nm} = (-\exp(\eta_{mn}) \times Time_{Train, mn} - Cost_{Train, mn}) / \exp(\alpha_{mn}) + \epsilon_{Train, mn} \quad (27)$$

where:

- $U_{Car, mn}$, $U_{SM, mn}$, and $U_{Train, mn}$ represent the utilities for car, Swissmetro, and train in menu m for individual n , respectively.
- $Time_{j, mn}$ and $Cost_{j, mn}$ represent the total (door-to-door) travel time and travel cost of alternative j in menu m presented to individual n , respectively. The cost coefficient is fixed to -1 .
- $ASC_{Car, mn}$ and $ASC_{SM, mn}$ represent alternative specific constants for car and train, respectively. The standard deviation of the train constant has been normalized to zero since it has the lowest value among all three alternatives.
- $\exp(\eta_{mn})$ and $\exp(\alpha_{mn})$ represent the coefficient of travel time and the scale parameter, respectively. Exponentiation is used in order model the log-normal distribution, which ensures that the travel time coefficient and the scale parameter are positive (and thus travel time and cost have a negative effect on utility to all individuals).
- $\epsilon_{Car, mn}$, $\epsilon_{SM, mn}$, and $\epsilon_{Train, mn}$ are error terms independently and identically distributed as extreme value type I.

5.2. Results

The model is estimated for menus 1–8 and the ninth menu is used for testing. In the following sections, we explore the estimation results with regards to inter- and intra-consumer heterogeneity and personalization. Afterwards, we estimate models using fewer choices done by each individual (2 or 5 choices out of 8), then apply the online procedure to the remaining choices up to the eighth choice.

5.2.1. Estimation with inter- and intra-consumer heterogeneity

The model is estimated using 400,000 Gibbs iterations, 200,000 of which are used as burn-in draws while the remaining 200,000 are used for sampling from the posterior distributions. The estimation results with menus 1–8 show significant inter-consumer heterogeneity for all coefficients. In addition, we find significant intra-consumer heterogeneity in the car and Swissmetro constants as shown in Table 5.

5.2.2. Predicting the next choice

The model estimated above is based on eight choices done by each individual. It utilizes all of the available training data. In this section, we estimate similar models using fewer menus (e.g. 2 or 5) and then perform the online procedure to all individuals.

Table 6 shows the log-likelihood and the predicted probability of the chosen alternative for the test menu (9th choice) using different estimation procedures. We first present the results for the full offline

procedure (8 menus for each individual). This procedure achieves an average predicted probability of 0.717 and a log-likelihood of -400.9 .

The following rows present the results with a subset of the data (2 choices and 5 choices per individual respectively). The results indicate that with fewer observations, we estimate models with lower average probabilities and log-likelihood values on the test data. However, the subsequent application of the online-procedure to the remaining menus recovers the drop in prediction accuracy as shown in the last two rows.

It can also be observed that the confidence intervals of the predicted probabilities (calculated empirically using the posterior distribution) of the full offline and the partial offline estimations (menus 1–2 and 1–5) do not overlap, indicating that the differences are statistically significant, and thus estimation with a fewer number of menus results in inferior predictions.

When more menus are included in the offline estimation, the predicted probabilities are higher because better priors are used (i.e. population level parameters μ , Ω_b , and Ω_w). Therefore, it is critical that the estimates of these parameters (which are obtained from the offline procedure) are accurate and up to date. To demonstrate the effects of using bad population level parameters, we perform online estimation using all 8 choices, but with the sample level means (μ) set to zeroes, and the inter- and intra-consumer covariance matrices set to identity matrices. The results indicate significantly worse predictions, with the mean of the posterior predictive distribution being 0.630, with the 95th percentile confidence interval $[0.622, 0.637]$. In addition, the likelihood of the test data is -463 , which is substantially worse than the values in Table 6.

5.2.3. Generating personalized recommendations

In order to demonstrate the accuracy and robustness of the proposed method in personalized recommendations, it is compared to two different approaches: a simple content-based method (in which the most chosen alternative in the previous menus, 1–7, is recommended), and non-personalized discrete choice models (flat logit and double mixture model with inter- and intra-consumer heterogeneity). The first

Table 6

Prediction results with the full offline, partial offline, and offline-online procedures.

Estimation procedure	Non-personalized		Personalized	
	Log-likelihood	Probability	Log-likelihood	Probability
Full offline (1–8)	−657	0.509 [0.487, 0.531]	−401	0.717 [0.709, 0.725]
Partial offline (1–2)	−666	0.514 [0.492, 0.536]	−551	0.668 [0.652, 0.683]
Partial offline (1–5)	−656	0.504 [0.480, 0.528]	−437	0.699 [0.689, 0.709]
Online (3–8)	–	–	−410	0.700 [0.692, 0.708]
Online (6–8)	–	–	−403	0.716 [0.708, 0.724]

Numbers in brackets indicate the 95th percentile confidence intervals of the mean predicted probability.

Table 7

Prediction results with the full offline, partial offline, and offline-online procedures.

Menu size	Content-based (most chosen)		Flat logit		Double mixture		Double mixture - personalized	
	1	2	1	2	1	2	1	2
Full offline (1–8)	0.763	0.954	0.636	0.910	0.609	0.914	0.770	0.977
Partial offline (1–2)	0.713	0.912	0.588	0.910	0.608	0.912	0.725	0.941
Partial offline (1–5)	0.745	0.947	0.626	0.911	0.609	0.912	0.757	0.968
Online (3–8)	–	–	–	–	–	–	0.767	0.952
Online (6–8)	–	–	–	–	–	–	0.777	0.975

approach accounts for personalization by considering the choice history of each individual, however, it cannot account for the impact of changes in the attributes (travel cost and travel time) as these vary among different choices. On the other hand, the non-personalized logit models account for attributes, but do not make use of the choice history of each individual. The offline-online estimation methodology presented in Section 3 accounts for both the individual choice history and alternative attributes.

Personalized menu optimization is performed with the objective of maximizing the expected hit rate [32,33] on the 9th choice. We simulate recommended menus that have either one or two out of the three original alternatives. The hit rate is defined as the fraction of individuals who choose an alternative that is included in the recommended menu. As shown in Table 7, the offline-online procedure can also approximate the full offline procedure in terms of hit rate, and the observed effect of personalization is substantial. (In this table, the online procedure using menus 6–8 achieves the highest hit rate. We would expect the full offline procedure to perform better, but predictions are based on the testing data; the sample level coefficients obtained from the offline estimation with the first 5 menus might fit the test data better than those obtained from the full estimation).

Table 7 indicates that the personalized double mixture model outperforms the content-based recommendation in all cases (by a margin of 1–2%), even when the online procedure is used. In addition, it is substantially better than the non-personalized flat logit and double mixture models.

6. Discussion

6.1. Model estimation

In this section, we discuss practical issues related to the effect of priors, identifiability, and applications in recommender systems.

6.1.1. Effect of priors

The basic HB procedure utilizes the Inverse Wishart (IW) prior, which has some undesirable properties, and thus can lead to biased estimates of standard deviations [4]. Particularly, this prior tends to inflate standard errors if their true values are small since it has a low density near zero. Although this issue did not impose any problems in our Monte-Carlo examples (because the standard errors are substantially distinguishable from zero), other priors can be used to avoid these biases such as the Hierarchical Inverse Wishart (HIW), Scaled Inverse Wishart (SIW), and Separation Strategy (or BMM) [34]. It should also be noted that the effect of priors decreases with increasing the sample size. With sufficient data, and with the “infinitely” diffuse priors, the posterior distribution is completely determined by the data, and therefore replicates the estimates obtained by maximum simulated likelihood.

6.1.2. Identifiability of individual-level preferences and accounting for uncertainty

The model with inter- and intra-consumer heterogeneity is only identifiable if multiple choice situations from each individual are

available. In addition, with few choice observations from each individual, the model suffers from “shrinkage”, whereby individual-level preferences are shrunk towards population means. HB is defined as a “data borrowing” technique that stabilizes individual-level preferences for each individual using information not only from his/her past choices, but also from other individuals within the same data set [26].

While Allenby and Rossi [2] state that this procedure allows us to estimate the distributions of the population level parameters (μ and Ω_b) and yields exact finite-sample estimates of the posterior distribution of individual-level parameters, Greene [15] argues that these estimates are only “exact” for the assumed priors and the data used, and up to simulation variance. To account for uncertainty in individual-level estimates, Allenby and Rossi [3] indicate that these estimates are not precisely estimated, and the use of point-estimates leads to over-confident predictions of effect-sizes. To avoid this over-confidence, Allenby and Rossi [3] suggest using all the posterior draws to make predictions instead of the point estimates (which is applied in our results in Sections 4 and 5).

Despite the fact that individual-level preferences are not accurately estimated, the results show that we achieve significantly better predictions compared to those without any personalization. These preferences are “learned” with more choices, which makes this procedure suitable for application in recommender systems.

Alternatively, the posterior distributions of the individual- and menu-specific parameters can be used in Multi-armed bandit methods such as Thompson sampling and Upper Confidence Bounds as described in Teo et al. [35] and Song [31]. For example, Thompson sampling uses individual draws from the posterior distributions. A distribution with large variance indicates high uncertainty in the estimated parameter. Therefore, attributes with uncertain parameter distributions become more likely to be recommended, which allows for learning these distributions more efficiently.

6.2. Application in recommender systems

6.2.1. Sample size and scalability

The sample size considered in our Monte-Carlo application is 10,000, which is sufficient for demonstrating the methodology for estimating and updating preferences. However, in app-based systems, the number of users can be potentially greater than tens of thousands. The five-step Gibbs sampler scales well with increasing sample size, as the estimates become closer to their true values and the required number of burn-in iterations decreases. Since the offline procedure is only performed periodically, long computational times can be tolerated. Individual preferences, on the other hand, are updated after each choice using the online procedure which can be performed in a few seconds to minutes, and can even be implemented on mobile devices.

6.2.2. Application to new users

The online procedure can be applied to new users with known choices with estimates of μ , Ω_b and Ω_w obtained from the offline procedure. For instance, these users may have joined the system and made choices after the last offline update. On the other hand, the random coefficients procedure described in Section 4.2 can be applied to

individuals with no previous choice history, and thus, these users will be first presented with non-personalized menus, i.e., population-level parameters will be used for menu optimization.

6.2.3. Data collection and endogeneity

In applications to recommender systems, the estimated models must account for endogeneity; the choice set presented to the user in each menu is based on this user's preferences, which are estimated based on his/her previous choices. Extensive research has been done on endogeneity corrections in discrete choice models, most of which falls into two categories: the BLP method [8], and the control-function method [17,18].

Endogeneity is not a concern in the models presented in this paper (since all the attributes used in the estimation of preferences were generated exogenously). In addition, Danaf et al. [13] show that endogeneity bias in recommender systems is ignorable if all the relevant data are used in estimation. This can be achieved by initializing the system with exogenous recommendations, and including all the available data in subsequent offline estimations.

7. Conclusions

This paper presented a methodology for estimating and updating consumer preferences online in the context of app-based recommender systems. We proposed an offline estimator, which estimates and updates individual and population level parameters periodically using a five-step Gibbs sampling procedure, and a real-time online estimator, which updates individual-specific parameters in real-time as more choices are made and assumes that population level parameters are fixed until the next offline estimation.

The proposed online estimator enables the use of discrete choice models in online decision support systems because it is (1) computationally efficient, (2) empirically accurate, and (3) theoretically justified. It is computationally efficient because it uses the data of the individual making the choice only, without the need to use data from other users. It is empirically accurate as it can achieve the same level of prediction accuracy as the offline estimator (which is computationally expensive and infeasible in real-time) as we have shown using real and Monte Carlo data. Finally, it is theoretically justified since it is equivalent to calibrating the model at the individual level, but with good priors representing the distribution of preferences in the population.

Our methodology subsumes the utility-based advantages of discrete choice models and the personalization capabilities of standard recommendation techniques by making use of all the available data including user-specific characteristics and preferences, alternative-specific attributes, and contextual variables. In our formulation of the utility equations, the estimated distributions can be interpreted as the individual's "willingness-to-pay" for different features, which can be used in pricing, designing, and recommending new alternatives. In addition, our models are able to account for complex patterns of preference heterogeneity, namely intra-consumer heterogeneity which represents variations in preferences across different choices of the same individual. Therefore, we avoid the unrealistic assumption that preferences are stable over time. This has also been shown to improve the accuracy of recommendations and predictions [7,33].

Several limitations arise in the application of our proposed methodology. The Monte-Carlo results indicate that sample level parameters (μ , Ω_b , and Ω_w) are recovered using the five-step Gibbs sampler. However, as in most Hierarchical models, individual- and menu-specific parameters might not be estimated precisely due to shrinkage [25]. These preferences are "learned" gradually from repeated choices. Nevertheless, using these preferences results in substantially better predictions compared to using an "average individual" (or unconditional distributions) even with a few number of choice situations.

The results presented in this paper are static and mimic SP

experiments. Consumer behavior may differ significantly between SP experiments and app-based choices. For instance, the time intervals between successive choices may vary considerably between the app-based systems and SP experiments.

Finally, there is a tradeoff between the model complexity (which results in high computational times) and the accuracy of predictions and recommendations. The complexity is determined by the utility equations which are specified by the researcher. The offline estimation results can be used to identify the significant predictors of choices, and adjust the utility equations accordingly. In addition, the model structure can be simplified by accounting for inter-consumer heterogeneity only (if intra-consumer heterogeneity does not appear to be significant). This would reduce the running time of the online procedure from a few seconds to less than 1 s.

This framework is implemented in the app-based travel adviser Tripod (Sustainable Travel Incentives with Prediction, Optimization and Personalization) [32,33] which incentivizes travelers to shift towards more sustainable alternatives (e.g. changing mode, route, or departure time choice behavior). Once more data from Tripod becomes available, the proposed methodology will be further validated, especially that users will be presented with real-life situations rather than SP experiments, and will have longer time intervals as well as contextual differences between successive choices, which allows for a higher level of intra-consumer heterogeneity. In addition, ongoing research is focused on modeling extensions to allow for flexible mixing distributions of inter- and intra-consumer heterogeneity, and incorporating socio-demographic and contextual information in order to partially explain inter- and intra-consumer heterogeneity respectively.

Acknowledgements

This work is funded by TRIPOD: Sustainable Travel Incentives with Prediction, Optimization and Personalization research project sponsored by the U.S. Department of Energy Advanced Research Projects Agency-Energy (ARPA-E). It was awarded through the ARPA-E Traveler Response Architecture using Novel Signaling for Network Efficiency in Transportation (TRANSNET) program.

References

- [1] C.C. Aggarwal, *Recommender systems*, Springer International Publishing, 2016.
- [2] G.M. Allenby, P.E. Rossi, Marketing models of consumer heterogeneity, *J. Econ.* 89 (1–2) (1999) 57–78.
- [3] G.M. Allenby, P.E. Rossi, Hierarchical Bayes models, *The Handbook of Marketing Research: Uses, Misuses, and Future Advances*, 2006, pp. 418–440.
- [4] I. Alvarez, J. Niemi, M. Simpson, Bayesian Inference for a Covariance Matrix, (2014) (arXiv preprint arXiv:1408.4050).
- [5] A. Ansari, S. Essegai, R. Kohli, Internet recommendation systems, *J. Mark. Res.* XXXVII (2000) 363–375.
- [6] F. Becker, M. Danaf, X. Song, B. Atasoy, M. Ben-Akiva, Hierarchical Bayes estimator of a logit mixture with inter- and intra-consumer heterogeneity, *Transp. Res. B Methodol.* 117 (2018) 1–17.
- [7] M. Ben-Akiva, D. McFadden, K. Train, Foundations of stated preference elicitation: consumer behavior and choice-based conjoint analysis, *Found. Trends Econ.* 10 (1–2) (2019) 1–144.
- [8] S. Berry, J. Levinsohn, A. Pakes, Automobile prices in market equilibrium, *Econometrica* (1995) 841–890.
- [9] M. Bierlaire, K. Axhausen, G. Abay, Acceptance of Modal Innovation: The Case of the Swissmetro, *Proceedings of the 1st Swiss Transportation Research Conference*, Ascona, Switzerland, 2001.
- [10] Castells, P., Hurley, N. J., & Vargas, S. (2015). Novelty and diversity in recommender systems. In Ricci, F., Rokach, L., and Shapira, B., *Recommender Systems Handbook*, Second Edition, Springer Science + Business Media, New York, 2015.
- [11] Chaptini, B. (2005). Use of discrete choice models with recommender systems. PhD dissertation, Department of Civil and Environmental Engineering, MIT.
- [12] K.W. Cheung, J.T. Kwok, M.H. Law, K.C. Tsui, Mining customer product ratings for personalized marketing, *Decis. Support. Syst.* 35 (2) (2003) 231–243.
- [13] M. Danaf, A. Guevara, B. Atasoy, X. Song, F. Becker, M. Ben-Akiva, Endogeneity Bias in Adaptive Choice Contexts: Choice-based Recommender Systems and Adaptive Stated Preferences Surveys, *Transportation Research Board Annual Meeting*, Washington D.C., 2019.
- [14] D. Goldberg, D. Nichols, B.M. Oki, D. Terry, Using collaborative filtering to weave

- an information tapestry, *Commun. ACM* 35 (12) (1992) 61–70.
- [15] W.H. Greene, Interpreting estimated parameters and measuring individual heterogeneity in random coefficient models, Department of Economics, Stern School of Business, New York University, 2003.
- [16] Gunawardana, A., & Shani, G. (2015). Evaluating recommender systems. In Ricci, F., Rokach, L., and Shapira, B., *Recommender Systems Handbook*, Second Edition, Springer Science + Business Media, New York, 2015.
- [17] J. Hausman, Specification tests in econometrics, *Econometrica* 46 (6) (1978) 1251–1272.
- [18] J. Heckman, Dummy endogenous variables in a simultaneous equation system, *Econometrica* 46 (4) (1978) 931–959.
- [19] S. Hess, J. Rose, Allowing for intra-respondent variations in coefficients estimated on repeated choice data, *Transp. Res. B Methodol.* 43 (6) (2009) 708–719.
- [20] S. Hess, K. Train, Recovery of inter- and intra-consumer heterogeneity using mixed logit models, *Transp. Res. B Methodol.* 45 (7) (2011) 973–990.
- [21] S. Huang, Designing utility-based recommender systems for e-commerce: evaluation of preference-elicitation methods, *Electron. Commer. Res. Appl.* 10 (4) (2011) 398–407.
- [22] J. Huber, K. Train, On the similarity of classical and Bayesian estimates of individual mean partworths, *Mark. Lett.* 12 (3) (2001) 259–269.
- [23] H. Jiang, X. Qi, H. Sun, Choice-based recommender systems: a unified approach to achieving relevancy and diversity, *Oper. Res.* 62 (5) (2014) 973–993.
- [24] Li, L., Chu, W., Langford, J., & Schapire, R. E. (2010). A contextual-bandit approach to personalized news article recommendation. In *Proceedings of the 19th International Conference on World Wide Web* (pp. 661–670). ACM.
- [25] Q. Liu, T. Otter, G.M. Allenby, Investigating endogeneity bias in marketing, *Mark. Sci.* 26 (5) (2007) 642–650.
- [26] B. Orme, G. Baker, Comparing Hierarchical Bayes Draws and Randomized First Choice for Conjoint Simulations, (2000) (Sawtooth Software Research Paper Series).
- [27] A. Polydoropoulou, M. Lambrou, Development of an e-Learning Recommender System Using Discrete Choice Models and Bayesian Theory: A Pilot Case in the Shipping Industry, (2012) In *Security Enhanced Applications for Information Systems*. (978-953-51-0643).
- [28] Ricci, F., Rokach, L., & Shapira, B. (2015). Recommender systems: introduction and challenges. In Ricci, F., Rokach, L., and Shapira, B., *Recommender Systems Handbook*, Second Edition, Springer Science + Business Media, New York, 2015.
- [29] Rubin, T., & Steyvers, M. (2009). A topic model for movie choices and ratings. In *Proceedings of the 9th International Conference on Cognitive Modeling*, Manchester.
- [30] Song, X. (2016). A Bayesian bandit approach to personalized online coupon recommendation. Master's Thesis, MIT.
- [31] X. Song, Personalization of future urban mobility. PhD dissertation, Department of Civil and Environmental Engineering, MIT, 2018.
- [32] X. Song, B. Atasoy, M. Ben-Akiva, Smart Mobility Through Personalized Menu Optimization. Transportation Research Board Annual Meeting, Washington, DC, USA, 2017.
- [33] Song, X., Danaf, M., Atasoy, B., & Ben-Akiva, M. (2018). Personalized menu optimization with preference updater: a Boston case study. Transportation Research Record: Journal of the Transportation Research Board.
- [34] Song, X., Becker, F., Danaf, M., Atasoy, B., & Ben-Akiva, M. (2019). Enhancement of Hierarchical Bayes for Logit Mixture. Working paper, MIT.
- [35] Teo, C. H., Nassif, H., Hill, D., Srinivasan, S., Goodman, M., Mohan, V., & Vishwanathan, S. V. N. (2016, September). Adaptive, personalized diversity for visual discovery. In *Proceedings of the 10th ACM Conference on Recommender Systems* (pp. 35–38). ACM.
- [36] K. Train, *Discrete Choice Methods with Simulation*, Cambridge University Press, 2009 (Chapter 12).
- [37] M. Yáñez, E. Cherchi, B. Heydecker, J. Dios-Ortuzar, On the treatment of repeated observations in panel data: efficiency of mixed logit parameter estimates, *Netw. Spat. Econ.* 11 (3) (2011) 393–418.
- [38] Ziegler, C.N., McNee, S.M., Konstan, J.A., & Lausen, G. (2005). Improving recommendation lists through topic diversification. In *Proceedings of WWW '05*.
- [39] W.H. Greene, Interpreting estimated parameters and measuring individual heterogeneity in random coefficient models, NYU Working Paper, EC-04-08, (2004).
- [40] P. Heidelberger, P.D. Welch, Simulation run length control in the presence of an initial transient, *Oper. Res.* 31 (6) (1983) 1109–1144.
- [41] A. Gelman, D.B. Rubin, Inference from iterative simulation using multiple sequences, *Stat. Sci.* 7 (4) (1992) 457–472.

Mazen Danaf is a PhD candidate in the Intelligent Transportation Systems Lab at MIT. His research focuses on smart mobility applications and personalization using behavioral models. Prior to that, he worked on activity-based travel demand models, and smart-phone applications used to collect data used in calibrating these models. He also holds BE and ME degrees in Civil and Environmental Engineering from AUB, where he has done research on students' travel behavior, drivers' behavior, and pedestrian-vehicular interactions.

Felix Becker studied a combined degree of computer science and business at Freie Universität Berlin. He is currently a graduate research assistant at ETH Zurich. His research interests include modeling people's preferences for automated vehicles and Bayesian analysis for discrete choice models.

Xiang Song is a PhD candidate in the Intelligent Transportation Systems Lab at MIT. His research focuses on smart mobility systems, Bayesian econometrics, and the interface between statistical learning and operations management. Prior to that, he has also worked on simulation-based decision support systems and market research.

Bilge Atasoy is an assistant professor in the Transport Engineering and Logistics section of the Department of Maritime and Transport Technology at TU Delft. Prior to joining TU Delft, she was a research scientist at MIT, in the ITS Lab, where she managed research projects in the areas of real-time optimization, travel behavior and choice-based optimization. She received her PhD from EPFL in 2013 where she was working on choice-based airline optimization models. Her main research interests lie at the intersection of optimization and behavioral modeling in order to achieve more efficient, innovative and demand-responsive transportation systems.

Moshe Ben-Akiva is the Edmund K. Turner Professor of Civil and Environmental Engineering and Director of the Intelligent Transportation Systems (ITS) Lab at the Massachusetts Institute of Technology (MIT) and at the Singapore-MIT Alliance for Research and Technology. He holds a PhD degree in Transportation Systems from MIT and honorary degrees from the University of the Aegean, the Université Lumière Lyon, the Royal Institute of Technology (KTH), and the University of Antwerp. His awards include the Robert Herman Lifetime Achievement Award in Transportation Science given by the Institute for Operations Research and the Management Sciences (INFORMS) Transportation Science and Logistics (TSL) Society, the Lifetime Achievement Award of the International Association for Travel Behavior Research, the Jules Dupuit prize from the World Conference on Transport Research Society, and the Institute of Electrical and Electronics Engineers ITS Society Outstanding Application Award for DynaMIT, a system for dynamic network management. He has worked as a consultant in industries such as transportation, energy, telecommunications, financial services and marketing for a number of private and public organizations.