



Testing probabilistic models of choice using column generation

Bart Smeulders^{a,1,*}, Clinton Davis-Stober^b, Michel Regenwetter^c, Frits C.R. Spieksma^d

^a HEC Management School, University of Liège, Liège 4000, Belgium

^b Department of Psychological Sciences, University of Missouri, 219 McAlester Hall, Columbia, MO 65211, USA

^c Department of Psychology, University of Illinois at Urbana-Champaign, 603 E. Daniel St., Champaign, IL 61820, USA

^d Department of Mathematics and Computer Science, TU Eindhoven, MB Eindhoven, 5600, The Netherlands

ARTICLE INFO

Article history:

Received 5 October 2017

Revised 21 February 2018

Accepted 7 March 2018

Available online 8 March 2018

Keywords:

Choice behavior

Probabilistic choice

Column generation

Membership problems

ABSTRACT

In so-called random preference models of probabilistic choice, a decision maker chooses according to an unspecified probability distribution over preference states. The most prominent case arises when preference states are linear orders or weak orders of the choice alternatives. The literature has documented that actually evaluating whether decision makers' observed choices are consistent with such a probabilistic model of choice poses computational difficulties. This severely limits the possible scale of empirical work in behavioral economics and related disciplines. We propose a family of column generation based algorithms for performing such tests. We evaluate our algorithms on various sets of instances. We observe substantial improvements in computation time and conclude that we can efficiently test substantially larger data sets than previously possible.

© 2018 Elsevier Ltd. All rights reserved.

1. Introduction

We consider computational challenges that arise when testing a certain type of probabilistic models of choice behavior. Imagine a decision maker who must specify a best element out of a set of distinct alternatives. In such situations, decision makers do not consistently select the same alternative as best, even when presented with the same (or nearly the same) set of alternatives (see, e.g., Tversky, 1969). Thus, assuming that a decision maker acts deterministically using a single decision rule (say, some linear order of the alternatives) is unrealistic. Probabilistic models of choice, pioneered by Block and Marschak (1960) and Luce (1959), attempt to explain uncertainty and fluctuations in behavior through probabilistic specifications. We concentrate on a class of models in which the permissible preference states are linear orders or weak orders of the alternatives. These are prominent cases in the ongoing research about rationality of preferences in behavioral economics, psychology, neuroscience and zoology (Arbuthnott et al., 2017; Brown et al., 2015; Regenwetter et al., 2011; Regenwetter and Davis-Stober, 2012). The random preference model captures the decision maker's uncertainty about preference with a probability distribution over these preference states. The

probability of choosing a particular alternative is governed by that probability distribution over preference states.

In a seminal contribution, McFadden and Richter (1990) provided several equivalent (sets of) conditions for choice probabilities to be consistent with such a probabilistic model of choice. However, actually checking these conditions on choice probabilities poses computational challenges. Indeed, straightforwardly evaluating the “axiom of revealed stochastic preference” and the “Block-Marschak polynomials” both require checking a number of conditions that is exponential in the number of choice alternatives. Likewise, the system of linear inequalities and the linear programs given in McFadden and Richter (1990) contained one variable for every preference state. The resulting number of variables grows exponentially in the number of alternatives, for most classes of preference states, including for linear orders. Even so, this linear programming model forms the basis of our column generation approach.

Most work on these probabilistic models has been on *binary choice induced by linear orders*. More precisely, the probability that a person chooses an alternative i over an alternative j , when required to choose one of the two, is the marginal probability of all linear orders in which i is preferred to j . Block and Marschak (1960) described two classes of inequalities and proved that these inequalities are necessary and sufficient conditions for consistency with the probabilistic model of choice for data sets with up to 3 choice alternatives. Dridi (1980) proved that these conditions are also necessary and sufficient for data sets with up to 5 alternatives and showed that they are no longer sufficient

* Corresponding author.

E-mail addresses: bart.smeulders@uliege.be, bart.smeulders@kuleuven.be (B. Smeulders), stoberc@missouri.edu (C. Davis-Stober), regenwet@illinois.edu (M. Regenwetter), f.c.r.spieksma@tue.nl (F.C.R. Spieksma).

¹ The author is a post-doctoral fellow the F.R.S.-FNRS.

for data sets with 6 or more alternatives. Megiddo (1977) proved that testing data sets for consistency with probabilistic choice induced by linear orders is difficult in general. He showed that the problem is equivalent to testing membership of a given point in the *linear ordering polytope*. Since optimization and separation over a particular polytope are polynomially equivalent (see Grötschel et al., 1993), it follows that testing whether a given collection of choice probabilities is consistent with a probabilistic model of choice induced by linear orders is NP-COMPLETE. In the last decade, researchers have generated extensive knowledge on the facial description of the linear ordering polytope (see Doignon et al. (2006); Fiorini (2006), the survey by Charon and Hudry (2007), and the book by Martí and Reinelt (2011), as well as the references contained therein).

When carrying out tests of probabilistic models of choice, scholars usually circumvent the computational challenges that arise when the number of alternatives grows large. Human laboratory experiments keep the number of alternatives small (see, e.g., Brown et al., 2015; Cavagnaro and Davis-Stober, 2014; Regenwetter et al., 2011; Regenwetter and Davis-Stober, 2012, who used sets of 5 alternatives). Kitamura and Stoye (2014) tested a probabilistic version of the “strong axiom of revealed preference,” using data from the U.K. Family Expenditure Survey, which they partitioned into subsets of a manageable size. One benefit of our proposed methodology is that it will allow researchers to design studies with larger numbers of choice alternatives, which will, in turn, increase their realism and generalizability. While testing probabilistic choice models is difficult in general, it becomes easy for some settings and classes of preference states. Matzkin (2007) and Hoderlein and Stoye (2014) provided conditions for a probabilistic version of the so-called “weak axiom of revealed preference.” Davis-Stober (2012) described a set of linear inequalities that are necessary and sufficient conditions for probabilistic choice induced by certain heuristic preferences. Smeulders (2018) provided necessary and sufficient conditions for a probabilistic model induced by single-peaked linear orders. Guo and Regenwetter (2014), Regenwetter et al. (2014), and Regenwetter and Robinson (2017) evaluated various sets of necessary and sufficient conditions for binary choice probabilities. For all of these settings, the conditions can be tested in polynomial time.

Here, we propose a family of algorithms based on column generation to test various probabilistic models of choice and apply it to a model induced by linear orders. Column generation is a technique to efficiently solve linear programs with a large number of variables; we come back to this technique in Section 3. Traditionally, the technique of column generation has almost always been applied to optimization problems. Here, however, we use it for a decision problem, namely, to detect whether given choice probabilities satisfy the probabilistic model of choice or not (i.e., a yes/no answer). We show how this affects the algorithm. The rest of this paper unfolds as follows. In Section 2, we lay out the notation, the definitions and the model that we use. Section 3 provides a basic description of the column generation algorithms. Section 4 discusses the implementation of a family of such algorithms and reviews results from computational experiments. In Section 5 we show that when testing the model for many similar choice probabilities, the column generation algorithm can use output from one test to speed up subsequent tests. We illustrate how this is useful for statistical analysis of probabilistic models, e.g., for calculating the Bayes factor to evaluate statistical performance on laboratory data from human subjects. We conclude in Section 6.

2. Notation and definitions

Consider a set A , consisting of n many alternatives and let $A \star A = \{(i, j) \mid i, j \in A, i \neq j\}$ denote the collection of all ordered pairs

of distinct elements of A . For each ordered pair of distinct alternatives $(i, j) \in A \star A$, we are given a nonnegative number $p_{i,j} \leq 1$. These numbers represent the probabilities that i is chosen over j for all distinct i and j in A . For now, we concentrate on *two-alternative forced choice*, that is, the case in which a person must choose one alternative or the other when offered a pair of alternatives. (We consider other cases in the appendix.) Therefore, $p_{i,j} + p_{j,i} = 1$ for each pair of $i, j \in A, i \neq j$. We refer to such a collection $\{p_{i,j} \mid (i, j) \in A \star A\}$ of binary choice probabilities as a *data set*. We denote a preference order over the alternatives by the relation $\succ$ and we use the index m to indicate a particular preference order. If, for the preference order $\succ_m$, the alternative $i \in A$ is preferred over the alternative $j \in A$, we write $i \succ_m j$. The relations $\succ_m$ are asymmetric, complete and transitive. The set of all such preference orders is O . We further consider the subsets $O_{i,j} \subset O$ for each $(i, j) \in A \star A$, where each $O_{i,j}$ contains all preference orders $\succ_m$ in which $i \succ_m j$. The particular probabilistic model of choice that we use is called the *mixture model* (also known as random preference model): this model assumes that when a decision maker is faced with a choice, each preference order has a certain probability of governing the choice. When these probabilities are consistent with the numbers $p_{i,j}$, we say that the mixture model *rationalizes* the data set.

Definition 1. Choice probabilities $\{p_{i,j} \mid (i, j) \in A \star A\}$ are rationalizable by the mixture model if and only if there exist values x_m , with $0 \leq x_m \leq 1$ for each $\succ_m \in O$, for which

$$\sum_{\succ_m \in O_{i,j}} x_m = p_{i,j}, \quad \forall (i, j) \in A \star A. \quad (1)$$

One straightforward way to find out whether a given data set is rationalizable by the mixture model is to check whether there exist nonnegative values x_m that satisfy this system of equalities (1). Similarly, a collection of empirical choice proportions (say, in a human subjects data set from a laboratory experiment) is *rationalizable* if it is consistent with having been generated by choice probabilities that are rationalizable. Determining whether this is the case is a matter of statistical inference subject to the equality constraints (1) on the generating probabilities $\{p_{i,j} \mid (i, j) \in A \star A\}$. Notice that the system of equalities (1) has a variable for every possible preference order of the alternatives, of which there exist $|O| = n!$ many. Even for a moderate number of alternatives, it is computationally prohibitive to solve this system.

Another approach is based on a result by Megiddo (1977): A collection $\{p_{i,j} \mid (i, j) \in A \star A\}$ of binary choice probabilities can be viewed as a point in a $n \times (n-1)$ -dimensional space. The collection is rationalizable if and only if that point is contained in the *linear ordering polytope*. This polytope (see Section 3.2 for its formulation) can theoretically be described by its facet-defining inequalities, which means that the data set $\{p_{i,j} \mid (i, j) \in A \star A\}$ is rationalizable by the mixture model if and only if the probabilities $p_{i,j}$ satisfy all inequalities defining the linear ordering polytope. However, the number of facet-defining inequalities needed to describe the linear ordering polytope rises very fast with the number of alternatives; a complete description is known for up to 7 alternatives only (see, e.g., Martí and Reinelt, 2011). Furthermore, the problem of establishing whether any facet-defining inequalities are violated is NP-COMPLETE for several known classes of inequalities. Here, we circumvent the need to solve a huge system of equalities (1), or to list and check all facet-defining inequalities, by moving to a different perspective: column generation.

3. Column generation

In this section, we describe an algorithm based on column generation to detect whether a given data set can be rationalized by the mixture model. Column generation is a technique dating

Download English Version:

<https://daneshyari.com/en/article/6892616>

Download Persian Version:

<https://daneshyari.com/article/6892616>

[Daneshyari.com](https://daneshyari.com)