

Distilling Models of Bounded-Rational Choice: A Constraint Programming Approach

Özgür Akgün

Georgios Gerasimou

July 2, 2026*

Abstract

We provide an analytical framework that allows for distilling the full explanatory and welfare-relevant content of influential yet computationally hard models of bounded-rational general choice. We do so by introducing constraint programming methods and tools from the optimization literature. We focus on the prominent “shortlisting” and “limited-attention” models. Applying our framework on imperfectly rational human choice data, we find that these models jointly account for nearly all behaviors, with limited-attention ones explaining better while being more permissive. Selection criteria that we introduce narrow down the models’ welfare-relevant predictions, considerably alleviating their indeterminacy and contributing toward their practical applicability.

Keywords: bounded rationality; general choice; constraint programming; shortlisting; limited attention; predictive success.

*Özgür Akgün: School of Computer Science, University of St Andrews, ozgur.akgun@st-andrews.ac.uk. Georgios Gerasimou (corresponding author): Adam Smith Business School, University of Glasgow, georgios.gerasimou@glasgow.ac.uk. Authors are ordered alphabetically. Georgios is grateful to Matúš Tejiščák for introducing him to Özgür and to constraint programming. We thank Paola Manzini for sharing with us the data that are analyzed in Section 5, and Thomas Demuyne for very helpful comments. The paper was presented at FUR 2026 (Alicante) and BSE 2026 (Barcelona), whose participants we thank for their feedback. Jacob Beadie, Tudor Huica and David Seruga provided excellent research assistance. Any errors are our own.

1 Introduction

Since the seminal work of Kalai, Rubinstein, and Spiegler (2002), much of the development of bounded-rational choice theory in the past two decades has taken place within the analytical framework of deterministic general choice where the analyst is assumed to observe decision makers' behavior at finite menus of discrete alternatives.¹ In addition to its generality, this environment's appealing features include the paucity of technical assumptions and interpretational clarity. Because of these features, it provides a fertile ground for the development of theories of decision making that are more general than—or form alternatives to—the textbook model of standard utility maximization, or *rational choice* (**RC**).

Arguably the two most influential theories of bounded-rational decisions in general choice environments are the intuitive “*rational shortlist method*” by Manzini and Mariotti (2007) (henceforth **RS**, or “*shortlisting*”) and the “*choice with limited attention*” (**CLA**) procedure by Masatlioglu, Nakajima, and Ozbay (2012), along with its “*overwhelming choice*” (**OC**) extension by Lleras, Masatlioglu, Nakajima, and Ozbay (2017) (henceforth “*limited-attention*” models).² The former portrays the decision maker as choosing in two stages: first, eliminating those feasible alternatives that are dominated according to some general shortlisting dominance relation; then, choosing from the set of shortlisted options by discarding those that are dominated according to a distinct second criterion. Limited-attention models are also two-stage choice procedures, but depict the agent as cognitively constrained in the sense of possibly being unable to consider *all* feasible alternatives, and rational given this constraint: specifically, the agent maximizes a stable preference ordering over the alternatives that they actually consider at each menu. Both models include **RC**, as a special case: the former when the first dominance relation is a complete ordering; the latter two when all feasible alternatives are always paid attention to.

While these seminal models have been studied at length in the theoretical literature, leading to many extensions, very few papers test their implications empirically and, to the best of our knowledge, none studies the particularly challenging task of doing so in the *full generality* that is afforded by these models, or the problem of re-

¹This is presented in the introductory chapters of major graduate-level microeconomics textbooks such as Kreps (1990), Mas-Colell, Whinston, and Green (1995), Rubinstein (2006) and Kreps (2012).

²Other prominent contributions to this literature include Masatlioglu and Ok (2005); Salant and Rubinstein (2008); Bernheim and Rangel (2009); and Ok, Ortoleva, and Riella (2015). A selection of additional contributions are discussed in the recent survey by de Clippel and Rozen (2024). Rubinstein and Salant (2008) and several other chapters in Caplin and Schotter (2008) discuss methodological implications of the behavioral/bounded-rational approach to revealed preference.

covering the complete range of the models’ behavioral primitives—by which we mean the underlying relations (preferences, attention filters)—in order to empirically assess their welfare-relevant information. This paper aims to facilitate such inquiries by introducing into the bounded-rationality and revealed-preference literatures the *Constraint (Dominance) Programming* (C(D)P) computational paradigms that achieve both these tasks simultaneously and, as highlighted above, in the most general formulations of the **RS**, **CLA** and **OC** models. In addition, to alleviate the explanatory indeterminacy that results from the well-known fact that these models are not identified, and to make them more applicable in practice, we propose two dominance-based normative (or welfare-motivated) refinements of the set of behavioral primitives that are compatible with these models on a given dataset. The criterion behind both refinements is the respective primitives’ proximity to rational choice: other things equal, the more complete either of the two rationales in **RS** is, and the more alternatives are paid attention to in a given menu, the closer the relevant primitive is to rational choice.

After recalling, in Section 2, the choice models that we focus on, in Section 3.1 and Section 4.2 we provide brief primers of CP and CDP, respectively, that are tailored to the specific optimization and model-primitive selection problems that pertain to these models. In a nutshell, CP methods—developed by researchers working at the intersection of optimization theory, computer science, artificial intelligence and operations research—allow us to exploit logical inference during search. Such inference is crucial as it avoids brute-force exploration of an exponentially large space of candidate model primitives, which we wish to allow to be either *perfectly or approximately compatible* with the choice data under some model. CDP methods, in turn, provide a low-level mechanism for enumerating only those perfectly or approximately compatible model primitives that are *undominated* according to a dominance relation of interest, thereby avoiding a separate generate-and-test post-processing step.³ Introducing CP and CDP therefore enables us to look deeper into the above prominent yet computationally demanding models of bounded-rational choice than has been possible under more conventional methods thus far.

In Section 5 we illustrate the usefulness of our C(D)P framework (which, we emphasize, can be adapted to the particular characteristics of *any* choice-theoretic model that may be defined in a general choice environment) with an empirical application. There, in addition to assessing the models’ explanatory power and permissiveness

³The specific dominance relations between shortlisting and limited-attention model primitives that we introduce and apply are dictated by which instance is closer to the model of rational choice, other things equal.

using gold-standard criteria in revealed-preference analysis, such as model-based interpretations of the Houtman and Maks (1985) proximity index, Bronar’s (1987) power test, and Selten’s (1991) measures of predictive success, we distill these models’ welfare-relevant content by enumerating the refined sets of “undominated” behavioural primitives that are compatible with a given dataset under each model. More specifically, we pin down the full explanatory, predictive and welfare-relevant content of **RS**, **CLA** and **OC** *vis-à-vis* **RC**, by applying our C(D)P methods and tool chain on two experimental datasets with riskless choices from 4 alternatives/11 menus and 5 alternatives/16 menus (Manzini, 2023).⁴ The majority of subjects’ choices in both datasets are perfectly compatible with **RC**. Focusing on those that are not, we find that the two limited-attention models (**CLA**, followed by **OC**) explain more subjects’ choices than **RS**, but are also considerably more permissive.⁵ Strikingly, our refinement criteria often filter out, respectively, thousands and hundreds of thousands of perfectly compatible but dominated model primitives in the 4- and 5-alternative datasets. In addition, the number of undominated such primitives of **CLA** and **OC** generally exceeds the corresponding number of undominated primitives in the most structured version of **RS**, which in turn exceeds the typically unique compatible preference ordering via which **RC** provides an *approximate* fit for these subjects. Our analysis, therefore, also uncovers and quantifies an interpretational trade-off that researchers of deterministic discrete choice are likely to be faced with when attempting to understand imperfectly rational behavior: perfect fit may come at the cost of poor identification.

From the extensive literature on bounded-rational choice with two-stage models since Manzini and Mariotti (2007), Masatlioglu et al. (2012) and Lleras et al. (2017), we highlight the following contributions as those that are closest in spirit to the present paper’s focus and questions: de Clippel and Rozen (2021); Inoue and Shirai (2023) on testing such models when data are limited; Freer and Nosratabadi (2026, 2024); Richter (2020) on studying non-identification in such models and formulating alternative ones that help alleviate it; Rubinstein and Salant (2012) on recovering an individual’s welfare-relevant preferences from imperfectly rational choices generated by structured procedures; and Demuynck (2011); de Clippel and Rozen (2021) on the computational complexity of testing two-stage models with possibly limited datasets. Compared to the most closely related of the studies above, the present one is more

⁴With the most efficient implementation of our CDP models, all computations on these two datasets were completed within a few hours on a regular laptop computer.

⁵As usual, permissiveness is quantified via the proportion of simulated uniform-random choice datasets that are also compatible with a model.

demanding in its assessment of the underlying models, in the sense that it performs such assessment in the models' full generality, and aims not only for the kinds of pass/fail checks that are afforded by the axiomatic approach, but also for: (i) computing the models' *proximity* to a dataset; (ii) *enumerating all* possible primitives of the model that achieve the closest such proximity; (iii) understanding which model fits could not have been achieved under random behavior; (iv) comparing models in terms of their predictive success.

Although motivated by the specific characteristics of these very influential models, and delivering their first set of in-depth empirical investigations, we view our methodological contribution primarily as forward-looking and closely related to contemporary studies on over-fitting, under-fitting and complexity trade-offs in model selection, e.g. Liang (2026); Fudenberg, Liang, and Gao (2026b); Fudenberg, Gao, and You (2026a); Fr chet te, Vespa, and Yuksel (2026); Andrews, Fudenberg, Lei, Liang, and Wu (2025); Fudenberg, Kleinberg, Liang, and Mullainathan (2022). Unlike the focus of these studies on parametric models of risk preference or belief updating, ours is a contribution to this literature on the domain of non-parametric and deterministic discrete choice models. In particular, we hope that constraint (dominance) programming will be a useful new framework in behavioral revealed preference analysis that will facilitate the development of a new generation of models in this domain that will also be informed by the—now measurable—trade-offs in their degrees of perfect or approximate explanatory power; indeterminacy; and predictive success.

2 Rational Choice, Shortlisting and Limited Attention

2.1 Definitions

The finite set that contains all choice alternatives is denoted by X . To avoid repetition, in the definitions pertaining to properties of a binary relation $\succ$ on X the relevant requirement is assumed to hold for all items in X . Specifically, such a relation $\succ$ is said to be *irreflexive* if $x \not\succ x$; *asymmetric* if $x \succ y \Rightarrow y \not\succ x$; *transitive* if $x \succ y \succ z \Rightarrow x \succ z$; and *total* if $x \neq y \Rightarrow x \succ y$ or $y \succ x$. An irreflexive relation $\succ$ is *acyclic* if $x_1 \succ x_2 \succ \dots \succ x_k \Rightarrow x_k \not\succ x_1$. A relation $\succ$ is a *strict partial order* if it is asymmetric and transitive. A total strict partial order is a *strict linear order*. Strict partial orders are acyclic relations, and the latter are asymmetric. The converse statements are false.

A single-choice dataset on X (henceforth simply *dataset*) is a tuple $\mathcal{D} = \{(A_i, C(A_i))\}_{i=1}^k$ where, for all $i \leq k$, $C(A_i) \subseteq A_i \subseteq X$ is a singleton and $A_i \neq A_j$ whenever $i \neq j$. A dataset $\mathcal{D}$ is explained by rational choice with strict preferences (**RC**) if there is a strict linear order $\succ$ on X such that, for all $i \leq k$,

$$C(A_i) = M_{\succ}(A_i), \quad (1)$$

where, for any $A \subseteq X$,

$$M_{\succ}(A) := \{x \in A : y \not\succ x \text{ for all } y \in A\}.$$

The set $M_{\succ}(A)$ is the set of $\succ$ -undominated, or $\succ$ -maximal, alternatives in A . When $\succ$ is a strict linear order, $M_{\succ}(A)$ consists of a single item, which is also the *dominant* or *greatest* element of $\succ$ in that set. If $\succ$ is acyclic and not merely asymmetric, then $M_{\succ}(A)$ is non-empty. A dataset $\mathcal{D}$ is explained by the “rational shortlist method” (**RS**) (Manzini and Mariotti, 2007) if there exist asymmetric relations $\succ_1$ and $\succ_2$ on X such that, for all $i \leq k$,

$$C(A_i) = M_{\succ_2}(M_{\succ_1}(A_i)) \quad (2)$$

A dataset $\mathcal{D}$ is explained by “choice with limited attention” (**CLA**) (Masatlioglu et al., 2012) if there is a pair $(\Gamma, \succ)$ where $\succ$ is a strict linear order on X and $\Gamma : 2^X \setminus \{\emptyset\} \rightarrow 2^X \setminus \{\emptyset\}$ an “attention filter” such that, for all $i \leq k$,

$$\Gamma(A_i) \subseteq A_i \quad (3a)$$

$$x \notin \Gamma(A_i) \implies \Gamma(A_i) = \Gamma(A_i \setminus \{x\}) \quad (3b)$$

$$C(A_i) = M_{\succ}(\Gamma(A_i)) \quad (3c)$$

Finally, a dataset $\mathcal{D}$ is explained by “overwhelming choice” (**OC**) if there is a pair $(\Gamma, \succ)$ where $\succ$ is a strict linear order on X and $\Gamma : 2^X \setminus \{\emptyset\} \rightarrow 2^X \setminus \{\emptyset\}$ a “competition filter” such that, for all $i \leq k$, (3a) and (3c) hold while (3b) is replaced by

$$x \in A_i \subset A_j \text{ and } x \in \Gamma(A_j) \implies x \in \Gamma(A_i) \quad (4)$$

Replacing (3b) with (4), and vice versa, makes **CLA** and **OC** logically distinct. The former condition is an inattention-irrelevance one: if an option is not considered at a menu, then whether or not it is feasible there does not affect the set of options that are actually paid attention to. This may be an appealing property in some situations but does not allow for thinking about cases where the only reason why some alternatives

are not considered is the menu’s off-puttingly big size. In other words, these options may well be paid attention to if they remain feasible in a subset of the original “big” menu that is now sufficiently small so as to be cognitively manageable by the agent. This is allowed by (4), though not by (3b).

Example (Part 1). Consider the following dataset from an underlying set X consisting of x , y and z :

$$\begin{aligned}\mathcal{D} &= \left\{ (A_i, C(A_i)) \right\}_{i=1}^4 \\ &= \left\{ (\{x, y\}, x), (\{y, z\}, y), (\{x, z\}, z), (\{x, y, z\}, x) \right\}\end{aligned}\tag{5}$$

Since $\mathcal{D}$ features a binary choice cycle, it is incompatible with **RC**. It is, however, explainable by **RS**, **CLA** and **OC**, as follows (we omit the simple verification details):

RS: $(\succ_1, \succ_2)$ defined by $y \succ_1 z$ and the strict linear order $z \succ_2 x \succ_2 y$.

CLA: $(\Gamma^{CLA}, \succ^{CLA})$ defined by $x \succ^{CLA} y \succ^{CLA} z$ and the attention filter Γ^{CLA} such that $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma^{CLA}} (\{x, y\}, \{y, z\}, \{z\}, \{x, y\})$.

OC: $(\Gamma^{OC}, \succ^{OC})$ defined by $y \succ^{OC} z \succ^{OC} x$ and the competition filter Γ^{OC} such that $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma^{OC}} (\{x\}, \{y, z\}, \{x, z\}, \{x\})$.

Yet these three instances of **RS**, **CLA** and **OC** are not the only ones that explain these choices. Applying our CP method we enumerate in Appendix A all 12 distinct pairs $(\succ_1, \succ_2)$ of asymmetric relations that are compatible with them under **RS**, as well as all 11 and 6 pairs $(\succ, \Gamma)$ of strict linear orders and attention/competition filters that are so under **CLA** and **OC**, respectively. $\blacklozenge$

2.2 Decomposition of Shortlisting into Three Nested Classes

Although we study three logically distinct models of bounded-rational choice, it will prove more insightful in the sequel to decompose **RS** into three nested classes, depending on the structure imposed on $\succ_1$ and $\succ_2$. The most general one clearly corresponds to both $\succ_i$ being asymmetric and possibly cyclic, i.e. the one axiomatized in Manzini and Mariotti (2007). We will refer to this as **RS_{asym}**. But special cases of natural interest here are also those where both $\succ_i$ are either acyclic (**RS_{acyc}**) or strict partial orders (**RS_{order}**). In fact, because of their relative tractability, the latter have received more attention in the choice-theoretic and revealed-preference literatures (Rubinstein and Salant, 2008; Au and Kawai, 2011; Dutta and Horan, 2015; Inoue and Shirai, 2023; de Clippel and Rozen, 2021).

Clearly, we have

$$\mathbf{RS}_{asym} \supset \mathbf{RS}_{acyc} \supset \mathbf{RS}_{order} \quad (6)$$

because $\mathbf{RS}_{asym}$ contains pairs $(\succ_1, \succ_2)$ where at least one of the two asymmetric relations may be cyclic and $\mathbf{RS}_{acyc}$ contains pairs where at least one of the two acyclic relations may not be transitive. Using this finer taxonomy, we can now clarify that out of this model’s 12 primitives that are perfectly compatible with the above example choices (Appendix A), one is in the $\mathbf{RS}_{order}$ class; nine in the $\mathbf{RS}_{acyc} \setminus \mathbf{RS}_{order}$ class (i.e. both relations in each primitive are acyclic and at least one is not transitive); and two in the $\mathbf{RS}_{asym} \setminus \mathbf{RS}_{acyc}$ class (i.e. both relations in each primitive are asymmetric and at least one is cyclic).

3 Computing Models’ Proximity to the Data

As we highlighted earlier, we are not merely interested in conducting a “pass-or-fail” test for each model, but also in assessing “how close” that model is to explaining a dataset where its fit is possibly imperfect. A standard such metric in revealed preference analysis, applicable both on discrete feasible sets and continuous/divisible ones that are determined by a budget constraint, is the so-called Houtman and Maks (1985) (HM) score (Choi, Fisman, Gale, and Kariv, 2007; Smeulders, Spieksma, Cherchye, and De Rock, 2014; Caplin, Dean, and Martin, 2011; Heufer and Hjertstrand, 2015; Apesteguia and Ballester, 2015; Dean and Martin, 2016; Costa-Gomes, Cueva, Gerasimou, and Tejiščák, 2022; Gerasimou, 2026; Demetry and Hjertstrand, 2023; Demuyne and Rehbeck, 2023). In its original formulation, this is defined as the smallest number of choices that need to be changed in (or, equivalently, dropped from) a subject’s dataset in order for them to be compatible with the model of rational choice under some preference ordering. However, this principle can be applied to any model, choice heuristic or axiom.⁶ In such cases, a model’s HM score—more descriptively, *distance score*—on a dataset similarly corresponds to the number of changes that are necessary for that model to fit those data perfectly under some of its permissible instances.

⁶“Prest” (Gerasimou and Tejiščák, 2018) is a free and open-source desktop tool that uses brute force optimization to compute this measure for several bounded-rational choice models of preference maximization that are defined only in terms of a—possibly incomplete—preference ordering. The measure itself is intuitive and natural enough that other papers in behavioral revealed preference—for example, Ellis and Freeman (2024)—report it without reference to Houtman and Maks (1985).

Formally, for each $\mathcal{D}$, the score of a model belonging to one of the **RS**_{class}, **CLA** or **OC** classes is defined, respectively, by

$$Score(\mathbf{RS}_{class}(\mathcal{D})) = \min_{(\succ_1, \succ_2) \in \mathbf{RS}_{class}} \{|\{(A_i, C(A_i) \in \mathcal{D})| : C(A_i) \neq M_{\succ_2}(M_{\succ_1}(A_i))\}| \} \quad (7)$$

$$Score(\mathbf{CLA/OC}(\mathcal{D})) = \min_{(\Gamma, \succ) \in \mathbf{CLA/OC}} \{|\{(A_i, C(A_i) \in \mathcal{D})| : C(A_i) \neq M_{\succ}(\Gamma(A_i))\}| \} \quad (8)$$

An intuitive interpretation of a model’s distance score for a given dataset is that it counts how many of the relevant decision maker’s choices might be considered “mistakes” from the perspective of that model, with a zero score reflecting a perfect model fit.

Computing the three models’ distance scores defined in (7) and (8)—let alone enumerating the full set of solutions in the respective arg min sets, which we also set out to do— is computationally hard, owing to the fact that their primitives grow exponentially in the size of the underlying set of alternatives.⁷ To give a sense of this we note that the numbers of distinct strict linear orders, strict partial orders, acyclic and asymmetric relations grow from 24, 219, 543 and 2,680 when $|X| = 4$ to 120, 4,231, 29,281 and 109,824 when $|X| = 5$.⁸ This makes brute-force search for the models’ distance scores, and recovery of the set of model instances that are compatible with these scores, an impractical approach. When $|X| = 4$, for example, using this method to test the most general version of the shortlisting model against a single subject’s dataset would involve searching $2,680^2 = 7,182,400$ pairs of asymmetric binary relations $(\succ_1, \succ_2)$ ⁹ and 630,118,440 pairs $(\succ, \Gamma)$ of strict linear orders and attention/competition filters¹⁰ toward testing either limited-attention model.

⁷Computing the HM score for **RC** is also hard. See Smeulders et al. (2014) for an analysis in the neoclassical demand-theoretic environment.

⁸Recall that the number of strict linear orders on X is $|X|!$. For the latter three types of binary relations see the A001035, A003024 and A123553 entries of the Online Encyclopedia of Integer Sequences in <https://oeis.org>.

⁹That said, we note that Demuyne (2011, Theorem 3) constructs an algorithm to establish that a pass-or-fail test of **RS** is possible in polynomial time on $|X|$.

¹⁰Indeed, there are 11 non-singleton menus when $|X| = 4$: 6 binary; 4 ternary; 1 quaternary. For binary menu $\{x, y\}$ the possible values of $\Gamma(\{x, y\})$ are 3: $\{x\}$, $\{y\}$ and $\{x, y\}$. Similarly, for ternary menu $\{x, y, z\}$ and quaternary menu $\{w, x, y, z\}$, these are 7 and 15, respectively. Thus, given that there are 24 strict linear orders on a set with 4 elements, it follows that there are $24 \times 3^6 \times 7^4 \times 15^1 = 630,118,440$ pairs $(\succ, \Gamma)$.

3.1 A Primer in Constraint Programming

The computational tasks in this paper are not standard problems where a real- or integer-valued objective function is to be minimized subject to some constraints. Had this been the case, we could have resorted to using tools for solving mixed integer linear programs (MILP).¹¹ Rather, we repeatedly need to: (i) decide whether there exist primitives—i.e. $(\succ_1, \succ_2)$ or $(\succ, \Gamma)$ —consistent with a dataset after dropping some observations, as part of the process of computing the scores in (7) and (8); and (ii) *enumerate* the set of primitives that achieve a given score for a model, possibly after applying a dominance-based refinement (we return to the latter point after introducing Definitions 1 and 2, below). This is precisely the type of setting for which constraint-based approaches are designed (Rossi, van Beek, and Walsh, 2006; Apt, 2003; Dechter, 2003).

We start by stressing that CP separates modelling from solving. Specifically, in CP one first writes a *model* of a problem in a bespoke, human-friendly programming language. This model comprises a finite collection of decision *variables* (the unknown objects to be recovered); their *domains* (the admissible values of those objects); and *constraints* (the logical conditions they must satisfy) (Apt, 2003). A CP *solver* then searches for assignments to the decision variables that satisfy all constraints. Importantly, the model is conceptually distinct from the algorithm used to solve it: the same high-level model can be compiled to different low-level solver formalisms, and solved by different backends/solving engines. In our implementation we first write models in the ESSENCE language (Frisch, Harvey, Jefferson, Martínez-Hernández, and Miguel, 2008) (full details are available in Appendix C). We then use the tools CONJURE and SAVILE ROW to translate these models into solver-optimal input (Akgün, Frisch, Gent, Jefferson, Miguel, and Nightingale, 2022; Nightingale, Özgür Akgün, Gent, Jefferson, and Miguel, 2014). Finally, this input is passed to the solver for execution and output production. We report our results with the MINION solver (Gent, Jefferson, and Miguel, 2006), which is explicitly designed for combinatorial-optimization problems such as the ones we tackle in this paper. We stress, however, that our modelling pipeline is not tied to a single solver: changing the backend is (conceptually and, in practice, often operationally) a compilation choice rather than a re-modelling exercise. In other words, the same human-friendly and choice-model specific code—i.e. in our case, those written in ESSENCE and shown in A.C.—can be

¹¹For a recent demonstration of integer-programming methods in revealed preference analysis from budget-constrained, consumer-theoretic choices we refer the reader to Demuyneck and Rehbeck (2023).

sent to different backend solvers at no cost to the researcher.

A useful way to view CP solving is as an interleaving of two components: *inference* (also called constraint propagation) and *search* (Bessiere, 2006; van Beek, 2006). This “embeds any reasoning which consists in explicitly forbidding values or combinations of values for some variables of a problem because a given subset of its constraints cannot be satisfied otherwise” (Bessiere, 2006, p.26). For example (taken from the same source), if a problem requires that integer variables x_1, x_2 must lie in the set $\{1, \dots, 10\}$ and a constraint asks that $|x_1 - x_2| > 5$, then propagating this constraint leads to the inference that $x_1, x_2 \notin \{5, 6\}$. More generally, inference reasons from the constraints to rule out variable values that cannot participate in any feasible solution, while search proposes choices (partial assignments) when inference alone does not determine a unique outcome.

In practice, solvers implement this inference via *filtering algorithms*, including specialised procedures for widely used constraint families (often called *global constraints*) (R  gin, 2005). The solver alternates between these: it makes a tentative assignment decision; propagates implications; detects contradictions early; and *backtracks* when necessary. This is the key mechanism by which CP avoids naive brute-force enumeration of all possibilities (Bessiere, 2006; van Beek, 2006; Mackworth, 1977). In the language of the CP literature, propagation *prunes* domains while backtracking organises the exploration of the remaining space.

Example (Part 2). To illustrate how inference/constraint propagation works in our setting, consider the example data introduced earlier. Under the **RS** model, constraint propagation with respect to the first rationale, $\succ_1$, of some candidate instance $(\succ_1, \succ_2)$ of the model that may explain the given choices amounts to reasoning that

$$\begin{array}{llll}
C(\{x, y\}) = x & \Rightarrow & y \not\succ_1 x & \\
C(\{y, z\}) = y & \Rightarrow & z \not\succ_1 y & \\
C(\{x, z\}) = z & \Rightarrow & x \not\succ_1 z & \\
C(\{x, y, z\}) = x & \Rightarrow & y \not\succ_1 x \ \& \ z \not\succ_1 x &
\end{array}
\quad \Longrightarrow \quad
\begin{array}{ll}
x \not\succ_1 z \not\succ_1 x & \text{(C1)} \\
y \not\succ_1 x & \text{(C2)} \\
z \not\succ_1 y & \text{(C3)}
\end{array}$$

From the exclusions imposed by constraints (C1)–(C3), then, it is inferred that the possible relations $\succ_1$ for any candidate instance $(\succ_1, \succ_2)$ are:

1. $y \succ_1 z$ (compatible with **RS**_{asym}, **RS**_{acyc}, **RS**_{order})
2. $x \succ_1 y$ (compatible with **RS**_{asym}, **RS**_{acyc}, **RS**_{order})
3. $x \succ_1 y, y \succ_1 z$ (compatible with **RS**_{asym}, **RS**_{acyc})

Constraint propagation/inference therefore leads to a meaningful computational gain by reducing the search for a compatible instance $(\succ_1, \succ_2)$ and only retaining either 2 or 3 of the 128 possible relations $\succ_1$ in this search. $\blacklozenge$

Many CP applications aim to find a single feasible solution or a single optimum. As already highlighted, our setting additionally requires *enumeration*: counting and listing all primitives that are compatible with a given model distance score, and then subjecting them to further refinement criteria. CP supports enumeration naturally: once a solution is found, the solver can be instructed to continue searching for additional solutions. Operationally, this is typically implemented by repeatedly solving while adding a *blocking constraint* that excludes the solutions already found or a region of the solution space implied by those solutions. In our application this supports the tasks of interest: “enumerate all $(\succ_1, \succ_2)$ or $(\succ, \Gamma)$ pairs that achieve the minimum score”.

Finally, a practical advantage of the particular set of free and open source CP tools that we deploy here, comprising ESSENCE together with CONJURE and SAVILE ROW, is that it lets us keep the economic decision-making structure of the problem explicit (relations, sets, and menu-indexed objects) while leaving the low-level encoding choices to the modelling tool chain. This is valuable in a study like ours where it is useful to be able to change the backend solver without rewriting the economic model, and where the same conceptual constraints are reused across: (i) model distance scoring; (ii) exact compatibility; (iii) dominance-refined enumeration.

4 Eliciting the Models’ Welfare-Relevant Content

We now use these computational tools to address the central economic question of welfare inference in the presence of imperfectly rational choice data. To this, we note that it is well-known—and already alluded to in the first part of our running example—that the **RS**, **CLA** and **OC** models are not identified: even when choices from the full domain of menus are available to the analyst, there is generally more than one instance of each model that could explain the data if *some* instance of the model does so. Aware of the possibility that the same data can be explained by multiple models, some authors have suggested the use of non-choice data in an attempt to understand which one(s) might be the most relevant for the problem of interest. But what if such additional non-choice data are not available? And, more pertinent to our study, how should an analyst decide which one(s) among the multiple instances

within the same model is/are the most appropriate for welfare analysis, irrespective of whether such additional non-choice data are available or not? We turn to this question next.

4.1 Proximity to Rational Choice as A Refinement Criterion

Multiplicity of a model's instances that provide an equally good explanation of a decision maker's choices can be either genuine, in the sense that the models' primitives across these instances are in conflict with each other, or mechanical, in the sense that those primitives differ in lesser, non-conflicting ways. To distinguish the former from the latter type of indeterminacy, we make use of the fact that all three models of bounded-rational choice in our focus include **RC** as a special case, and introduce a *dominance filtering* criterion that favors those instances of each model that are "closer" to **RC**. In the case of **RS**, this means favoring instances that contain a "more complete" relation $\succ_1$ or $\succ_2$, *ceteris paribus*. In the case of **CLA** and **OC**, it means favoring instances featuring larger consideration sets, *ceteris paribus*.

Definition 1

If $(\succ_1^a, \succ_2^a)$ and $(\succ_1^b, \succ_2^b)$ are two distinct instances of a shortlisting model $\mathbf{RS}_{class}$, $class \in \{order, acyc, asym\}$, that explain a dataset $\mathcal{D}$, then the former is said to dominate the latter in $\mathbf{RS}_{class}$ if

$$\succ_t^b \subseteq \succ_t^a \quad \text{for all } t \in \{1, 2\} \quad (9a)$$

$$\succ_t^b \subsetneq \succ_t^a \quad \text{for some } t \in \{1, 2\} \quad (9b)$$

An instance in $\mathbf{RS}_{class}$ is undominated if it is not dominated in $\mathbf{RS}_{class}$.

This notion is applicable as a filtering criterion within each of the three logically nested **RS** classes [cf. (6)], which, as was noted in Section 2.2, are defined by the consistency structure of the least consistent rationale in the relevant pair/instance of the model. Finding the undominated compatible pairs/instances within a given **RS** class, however, does not immediately suggest a way of comparing those instances across classes. For example, the pairs $(\succ_1^a, \succ_2^a)$ and $(\succ_1^b, \succ_2^b)$ defined by

$$\begin{array}{cc} \succ_1^* & \succ_2^* \\ \star = a : & (x, y) \quad (z, y), (y, x), (z, x) \\ \star = b : & (x, y), (y, z) \quad (z, y), (y, x), (z, x) \end{array}$$

correspond to shortlisting instances in the $\mathbf{RS}_{order}$ and $\mathbf{RS}_{acyc}$ classes, respectively, that feature an identical second rationale. If they were treated as belonging to the

same class, the latter instance would clearly dominate the former. A question now arises: Should an acyclic but intransitive shortlisting relation be considered closer to being rational than a transitive such relation that contains strictly fewer ranked comparisons? More generally, how should one think about trade-offs involving the hierarchical consistency of rationales—whereby *transitive* $\gg$ *acyclic* $\gg$ *asymmetric* in this order—and their plurality?

We imagine both approaches being potentially appropriate in some contexts. Prioritizing consistency may be relatively uncontroversial if one of the two relations in the debate is cyclic, not least because cyclicalities can be detrimental to any attempt to draw actionable welfare-relevant conclusions. This is less so, however, when one relation is transitive while the other—although merely acyclic—contains more pairs of comparable alternatives, as is the case in the above example. An obvious argument in favor of plurality here is that the information contained in the latter relation is not self-contradictory and, being also richer than its opponent, may provide more targeted welfare-relevant guidance. On the other hand, if it is postulated that x is somehow better than y , and y in turn is considered to be better than z , yet x and z cannot be compared, a cautious observer might be concerned about the validity (or lack thereof) of one or both of the two antecedent comparisons. For such an observer consistency might still be prioritized over plurality.

Definition 2

If $(\succ^a, \Gamma^a)$ and $(\succ^b, \Gamma^b)$ are two distinct instances of a limited-attention model that explain a dataset $\mathcal{D}$, then the former is said to dominate the latter if

$$\succ^b = \succ^a \tag{10a}$$

$$\Gamma^b(A_i) \subseteq \Gamma^a(A_i) \quad \text{for all } i \leq k \tag{10b}$$

$$\Gamma^b(A_i) \subsetneq \Gamma^a(A_i) \quad \text{for some } i \leq k \tag{10c}$$

That is, a data-compatible limited-attention model primitive dominates another such primitive if they both feature the same underlying preference relation but the former portrays the decision maker as always paying attention to at least as many choice alternatives as the latter, and in at least one case as paying attention to strictly more alternatives. One can imagine an additional filtering process being applied among the primitives that survive this refinement, where those featuring attention/competition filter mappings that are finer in the sense of (10b)–(10c) are favored even when the underlying preference relations differ. While easily implementable in our setting, such an additional filtering process is not without loss of potentially welfare-relevant

information, unlike the one put forward in Definition 2.

Example (Part 3). As noted earlier, Appendix A shows the 12, 11 and 6 distinct primitives of the **RS**, **CLA** and **OC** models that are compatible with the cyclic choices in our running example. Appendix B explains why: 3 of the 9 compatible primitives $(\succ_1, \succ_2)$ in the $\mathbf{RS}_{acyc} \setminus \mathbf{RS}_{order}$ class are undominated; 1 of the 2 in $\mathbf{RS}_{asym} \setminus \mathbf{RS}_{acyc}$; 3 of the 11 **CLA** primitives $(\Gamma, \succ)$; and 3 of the 6 **OC** primitives. $\blacklozenge$

Intuitively appealing as they may be for the empirical revealed preference analysis of bounded-rational choices, it is not obvious how the proposed dominance-based refinements of the **RS** and **CLA**, **OC** compatible model primitives can be operationalized so that they are applied systematically in practice. This is our motivation for introducing and tackling this problem with CDP.

4.2 A Primer in Constraint *Dominance* Programming

Standard constraints restrict what a *single* candidate primitive may look like (e.g. that $\succ$ is a strict linear order, or that Γ satisfies the attention-filter conditions). These are constraints *on a solution*. However, the dominance relations in Definitions 1 and 2 are not properties of a single instance in isolation: they are relations *between* instances. Such requirements are naturally expressed as constraints *among solutions*. A direct way to compute undominated instances is therefore a two-step generate-and-test procedure: 1. Enumerate *all* feasible solutions (all compatible instances); 2. Post-process the resulting list to filter out those that are dominated. This works, but has a clear computational weakness: it requires a full enumeration even when a large fraction of the enumerated instances are dominated and will ultimately be discarded.

Is there a better alternative? Indeed, CDP avoids such exhaustive generate-and-test. Instead, it is a principled way to integrate dominance filtering into the enumeration process itself (Guns, Stuckey, and Tack, 2018; Chu and Stuckey, 2015). The modeller specifies a dominance relation over solutions, such as those in Definitions 1 and 2, and the CDP mechanism ensures that once a solution is found, the subsequent search is constrained so as to exclude solutions dominated by it. Concretely, after discovering a solution s , the system posts a *dominance blocking constraint* that rules out any future solution s' such that s dominates s' . Search then continues in the remaining space, which (by construction) contains no solution dominated by an already accepted one. The net effect is that the solver spends its effort on discovering

new *undominated* solutions, rather than on discovering large numbers of dominated ones and discarding them afterwards.

CDP nevertheless does not remove the need for post-processing completely. The reason is that, during enumeration, a solution that is *non-dominated so far* may later turn out to be dominated by a solution found subsequently. In such cases, the earlier solution will have been recorded even though it is not undominated in the final set. A lightweight post-processing pass can therefore still be required to remove any solutions that are dominated *eventually* by later discoveries. CDP should thus be viewed as a way to avoid *exhaustive* generate-and-test (by steering search away from dominated regions as it progresses), rather than as a guarantee that every intermediate output is globally undominated. For the problem at hand, CDP can substantially reduce the burden of enumerating and filtering large families of dominated **RS**, **CLA** and **OC** instances, leaving only minimal post-processing to ensure that the retained set is undominated in the global sense.

5 Illustration with Choices from Two Experiments

5.1 Data

We use data from two online experiments conducted by Manzini (2023). Participants were recruited on the “Ravelry” social network which, according to Wikipedia (accessed on 26th June 2026) “*functions as an organizational tool for a variety of fiber arts, including knitting, crocheting, spinning and weaving. Members share projects, ideas, and their collection of yarn, fiber and tools via various components of the site.*” In the first experiment the grand choice set X comprised 4 knitted dresses. Subjects (primarily female senior citizens) were asked to make a single choice from all 11 menus with 2, 3 or 4 items that are derivable from that set. A well-defined, single-valued choice function was therefore elicited for every subject there. In the second experiment, X comprised 5 knitted ponchos. Here, subjects made single choices from 16 of the 26 non-trivial menus (61.5%) that are derivable from X ; specifically, those that comprised the grand set itself and those with 2 or 4 items. A partially defined choice function on the same sub-domain was therefore elicited for every subject. All subjects were rewarded with a knitting pattern, and a randomly selected subset of them also received a yarn in their chosen color to knit with it. Out of the 277 and 336 subjects with fully recorded choice responses in these two experiments, respectively, 195 (70.4%) and 169 (50.3%) were explained perfectly by rational choice with strict

preferences. Our primary focus will therefore be on the remaining 82 and 167 subjects in these “Four Dresses” (4D) and “Five Ponchos” (5P) data.

5.2 Explanatory Power and Predictive Success

Table 1 presents the results of our non-parametric model-score analysis in the sense of (7)-(8) of the non-**RC**-compliant subjects in the 4D and 5P data (panels (a) and (b), respectively). In the former data, 77 of the 82 subjects (94%) are perfectly explainable by at least one of the **RS**, **CLA** and **OC** models. Among them, 75 and 51 are perfectly compatible with **CLA** and **OC**, respectively, with all remaining ones being approximately compatible, with a score of 1. Furthermore, 4 and 53 subjects, respectively, are perfectly and approximately compatible—with a score of 1—with all three classes of the **RS** model. Similar results are obtained in the larger Five-Ponchos sample. Notably, **CLA** now explains *all* subjects perfectly, while the proportion of those explained by **RS** is doubled compared to 4D (for **OC** they are the same). Additionally, these data allow for differentiating between the explanatory power of the more general **RS**_{acyc} and **RS**_{asym} classes relative to the transitive **RS**_{order} class, with the former accounting perfectly for one additional subject’s behavior.

Table 1: Proportions of subjects who are incompatible with Rational Choice and are perfectly or imperfectly explained by the Shortlisting and Limited-Attention models.

<i>Perfect & approximate fits</i>	RS_{order}	RS_{acyc}	RS_{asym}	OC	CLA
	Four Dresses ($N = 82$)				
Score=0	4.8%	4.8%	4.8%	62.2%	91.5%
Score=1	64.6%	64.6%	64.6%	37.8%	8.5%
Score=2	20.7%	21.9%	21.9%	—	—
Score=3	9.9%	8.7%	8.7%	—	—
	Five Ponchos ($N = 167$)				
Score=0	8.9%	9.6%	9.6%	62.3%	100%
Score=1	55.7%	55.7%	55.7%	37.7%	—
Score=2	23.5%	25.1%	25.1%	—	—
Score=3	8.9%	8.4%	8.4%	—	—
Score ≥ 4	3.0%	1.2%	1.2%	—	—

Next, following the ideas put forward in Bronars (1987); Selten (1991) and later refined in Beatty and Crawford (2011), Cherchye, Demuyne, de Rock, and Lanier

(2025) and Fudenberg et al. (2026b), among many others, we assess the predictive ability of the three models in the 4D and 5P data by juxtaposing how well they explain human vs. synthetic subjects who choose uniform-randomly in the respective collections of 11 and 16 menus. In particular, for a model M in $\{\mathbf{RS}_{asym}, \mathbf{RS}_{acyc}, \mathbf{RS}_{order}, \mathbf{CLA}, \mathbf{OC}\}$ we denote by r_M (“hit rate”) and a_M (“area”), respectively, the proportions of human and synthetic subjects that are perfectly explained by M , and let

$$m_M^1 := r_M - a_M \tag{11}$$

$$m_M^2 := \frac{r_M}{a_M} \tag{12}$$

be, respectively, that model’s *linear difference* and *ratio* measures of predictive success at the given collection of choice problems (Selten, 1991). Both measures are increasing in r and decreasing in a . The linear measure—which is more frequently applied, and was axiomatized in Selten (1991)—is bounded above by 1 and below by -1. Compared to the ratio measure, it tends to favor models with better explanatory power (i.e. a higher r_M). The ratio measure on the other hand is bounded below by 0 and unbounded above. Because, by definition, it favors models with higher explanation/permissiveness ratios, models which explain few human subjects’ behavior but are also unlikely to explain random choices could be ranked above models with the opposite empirical characteristics. The two measures are therefore complementary.

Indeed, this contrast between m^1 and m^2 is manifested in the results shown in Table 2, which were computed by applying our CDP models for each of the five choice models on 10,000 simulated datasets in each of the 4D and 5P data (the latter computations were completed in less than 24 hours in a regular laptop computer). In both experimental samples, m^1 puts one of the two limited-attention models at the top whereas m^2 ranks **RS** first, with this model’s three constituent classes in reverse order of permissiveness (namely, **RS**_{order} first, followed by **RS**_{acyc} and then **RS**_{asym}). Importantly, the top-ranked model under either measure and dataset is robustly so at the 5% level of significance, with confidence intervals computed as suggested in Demuynck (2015).¹²

¹²Because all subjects in each dataset faced the same problems, the consistent estimator of the relevant variance-covariance matrix of r and a in Demuynck (2015, p. 693), with respect to some model M and given a sample of size n , has zero entries everywhere except for the variance of $r_{M,n}$. The latter is estimated by $v_{M,n} = r_{M,n}(1 - r_{M,n})$ (note: we write $v_{M,n}$ instead of $v_{M,n,m}$, where m here denotes the number of random datasets, due to the above-mentioned redundancies in the variance-covariance estimator). The 5% confidence interval of measure m^i for model M given sample size n , denoted $m_{M,n}^i$ for $i = 1, 2$, is then determined by $\left[m_{M,n}^i - c_{5\%} \sqrt{\frac{v_{M,n}}{n}}, m_{M,n}^i + c_{5\%} \sqrt{\frac{v_{M,n}}{n}} \right]$, where $c_{5\%} \approx 1.645$ is the 5% cutoff value of the standard normal cumulative density function.

Table 2: Evaluating models according to Selten’s (1991) difference and ratio measures of predictive success (in parentheses, 95% confidence intervals as in Demuyneck, 2015).

(a) Four Dresses.

	RS_{asym}	RS_{acyc}	RS_{order}	CLA	OC
<i>r</i>	0.0488	0.0488	0.0488	0.9146	0.6219
<i>a</i>	0.0119	0.0119	0.0077	0.3689	0.3471
<i>m</i> ¹	0.0369	0.0369	0.0411	0.5457	0.2748
	(−0.002, 0.076)	(−0.002, 0.076)	(0.002, 0.080)	(0.495, 0.596)	(0.187, 0.363)
<i>m</i> ²	4.1008	4.1008	6.3376	2.4792	1.7917
	(3.742, 4.460)	(3.742, 4.460)	(5.892, 6.784)	(2.396, 2.563)	(1.642, 1.941)

(b) Five Ponchos.

	RS_{asym}	RS_{acyc}	RS_{order}	CLA	OC
<i>r</i>	0.0958	0.0958	0.0898	1.0000	0.6227
<i>a</i>	0.0030	0.0027	0.0016	1.0000	0.2094
<i>m</i> ¹	0.0928	0.0931	0.0882	0.0000	0.4133
	(0.055, 0.130)	(0.056, 0.131)	(0.052, 0.125)	(0.000, 0.000)	(0.352, 0.475)
<i>m</i> ²	31.933	35.481	56.125	1.0000	2.9737
	(31.249, 32.617)	(34.760, 36.202)	(55.215, 57.035)	(1.000, 1.000)	(2.839, 3.109)

(c) Model rankings across the two datasets and measures.

Four Dresses				Five Ponchos			
<i>m</i> ¹		<i>m</i> ²		<i>m</i> ¹		<i>m</i> ²	
CLA	0.545	RS_{order}	6.34	OC	0.427	RS_{order}	56.12
OC	0.275	RS_{acyc}	4.10	RS_{acyc}	0.093	RS_{acyc}	35.48
RS_{order}	0.041	RS_{asym}	4.10	RS_{asym}	0.092	RS_{asym}	31.93
RS_{acyc}	0.037	CLA	2.48	RS_{order}	0.088	OC	2.973
RS_{asym}	0.037	OC	1.79	CLA	0.000	CLA	1.000

A result that is revealed by this analysis and may be worth highlighting is the fact that, despite the larger domain spanned by the 5P data (5 items & 16 menus, against 4 & 11 in 4D), **CLA** explains perfectly not only every human subject, but every synthetic random-behaving subject too. This is despite that, for both **OC** and **CLA**, the consideration mapping of optimal solutions properly includes the choice made at some menu, i.e. $\Gamma(A) \supsetneq C(A)$.¹³ An intuitive explanation for this finding is that the 10 three-element menus from which choices were not observed in the 5P data give disproportionately more degrees of freedom in **CLA** compared to **OC** to define a rationalizing consideration mapping Γ for any candidate matching preference relation.¹⁴

Next, taking the perfect fits of all three models and classes thereof at face value, in Table 3 and Figure 1 we summarize the results from enumerating compatible model primitives per subject after application of our CDP method, which eliminates instances that are dominated according to the refinements of Definitions 1 and 2. In the 4D data we find that, for all 4 subjects who are perfectly explained by **RS**, there are 4, 12 and 32 undominated instances that are compatible with the **RS_{order}**, **RS_{acyc}** and **RS_{asym}** classes of the model, respectively. The average number of undominated primitives for subjects who are perfectly explained by **CLA** and **OC**—both of which also explain the above 4 subjects—is 8 and 11 for **CLA** and **OC**, respectively. Importantly, however, despite these generally high numbers, for the 2 and 26 subjects who are perfectly matched *only* by **OC** and **CLA**, respectively, the compatible instances are markedly lower, down to 3.5 and 4.4 on average (Figure 1(a)). In fact, for two of

¹³The **CLA** model’s explanatory gains on both human and synthetic data gradually decrease as we expand this constraint by requiring to cover more menus. Introducing this extra analytical flexibility amounts to defining a distinct class of models which encompass the one introduced and axiomatically characterized in Masatlioglu et al. (2012). Freer and Nosratabadi (2026) study a variation of **CLA** and **OC** along those lines. There, however, “attention floors” (lower bounds on the number of alternatives that the decision maker pays attention to) apply across all menus.

¹⁴Indeed, (3b) and (4) differ meaningfully in this case: if x is chosen at $\{v, w, x, y, z\}$ or in any of this menu’s 4-element subsets, (4) implies that x must necessarily belong to the consideration set $\Gamma(A)$ of every latent 3-element menu A that contains x . By not imposing this constraint, (3b) allows for more rationalization possibilities under some Γ for any candidate linear order $\succ$. For example, suppose $C(\{v, w, x, y\}) = x$, $C(\{v, x\}) = v$ and $C(\{x, y\}) = y$ are part of a 5P choice dataset. To make the point clear while keeping other things equal, suppose $\Gamma^{OC}(\{v, w, x, y\}) = \Gamma^{CLA}(\{v, w, x, y\}) = \{w, x, y\}$. Although $C(\{v, x, y\})$ is unobserved, we know from (4) and the above data that $\Gamma^{OC}(\{v, x, y\}) \supseteq \{x, y\}$ and, therefore, $\Gamma^{OC}(\{x, y\}) = \{x, y\}$. This implies, for example, that the candidate preference order $v \succ w \succ x \succ y$ could not explain these choices under **OC**. By contrast, (3b) allows Γ^{CLA} to be agnostic on the latent menu $\{v, x, y\}$, and then to freely assign $\Gamma^{CLA}(\{v, x\}) = \{v\}$, $\Gamma^{CLA}(\{x, y\}) = \{y\}$. Such a Γ^{CLA} mapping enables the observed choices to be in principle rationalizable by **CLA** under the above preference relation.

the **CLA**-only subjects our analysis elicits the most welfare-informative possible fits by identifying a *unique* compatible undominated primitive $(\Gamma, \succ)$ per subject. This demonstrates that the CDP approach can be extremely illuminating in some cases.

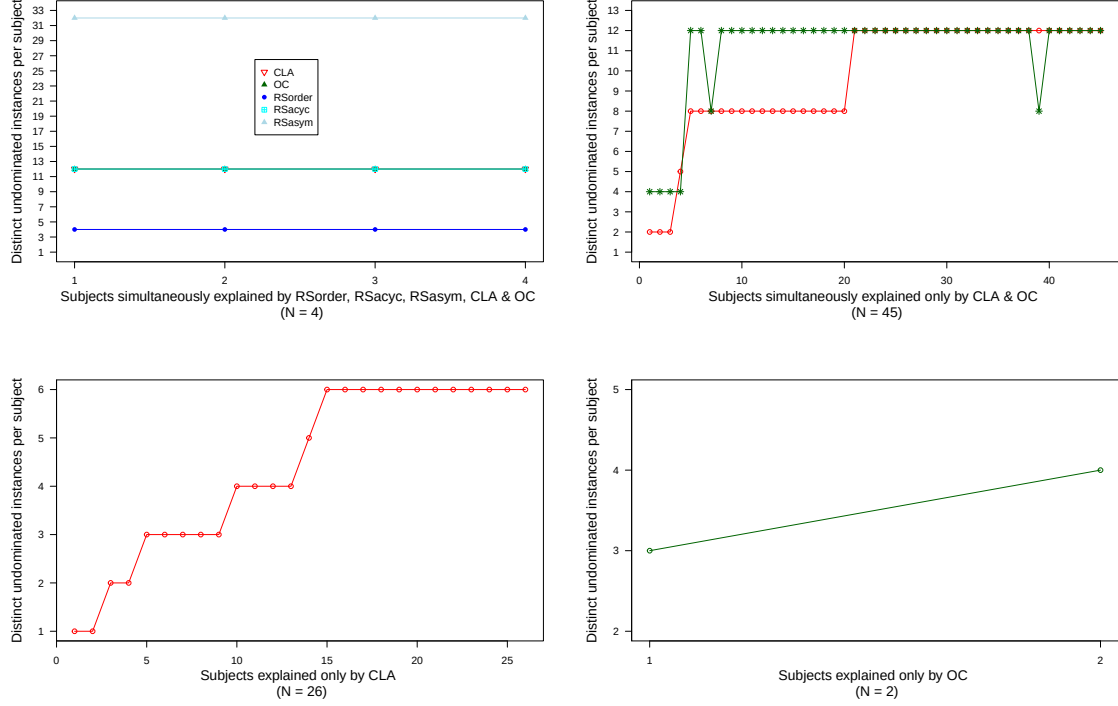
For the 4D data we also enumerated the full sets of compatible primitives (i.e. those that also include dominated ones) in order to quantify each model’s degree of non-identification prior to the application of any refinements but also—and more importantly—the degree to which a model’s compatible primitives are dominated. As Table 3(a) shows, such explanatory indeterminacy is highest in **CLA** (typically amounting to thousands of compatible primitives) and **RS_{asym}**, followed by **OC**, **RS_{acyc}** and **RS_{order}**. What is arguably most striking here is the fact that between 83.3% (**RS_{order}**) and 99.6% (**CLA**) of all instances across all models are discarded as dominated.

While we did not enumerate the full list of compatible model-specific behavioral primitives before applying our dominance refinements in the 5P data, we note that in the case of one subject this amounted to nearly 351,000 **CLA** primitives. In these data we applied our CDP models and tool chain directly, as explained in Section 4.2; in particular, we did not have to compute the full list of compatible primitives via our CP models first. We find that the average and median numbers of undominated primitives are higher for all models relative to their 4D-data counterparts, with the most indeterminate again being **RS_{asym}** (average: 512), which is now followed by **CLA** rather than **OC** (79 vs 53), and then by **RS_{acyc}** (95) and **RS_{order}** (32). The generally higher indeterminacy here is caused by the larger number of alternatives and menus in the data, as well as by the absence of choices from menus with three elements. The model-specific average computation times in 5P also increased by a varying number of factors compared to 4D (≈ 782 for **RS_{asym}** vs. ≈ 4 for **CLA**). However, all CDP computations—both for model scores and for the enumeration of undominated instances—were completed in a few hours on a regular laptop.

These results highlight the potential of the prominent shortlisting and limited-attention theories of bounded-rational choice, in conjunction with sophisticated theory-based computational methods such as the C(D)P ones that we propose, to extract non-trivial behavioral insights from human choices that deviate from the rational model. Yet the general picture that emerges from the preceding analysis generally points toward a considerable amount of indeterminacy still remaining for each of these models after application of the relevant dominance-based criteria for selecting welfare-relevant model primitives. Perhaps a natural benchmark relative to which the degree of this indeterminacy might be assessed is how **RC** performs in this respect

Figure 1: Undominated behavioral primitives that explain subjects perfectly under one or more models.

(a) Four Dresses



(b) Five Ponchos

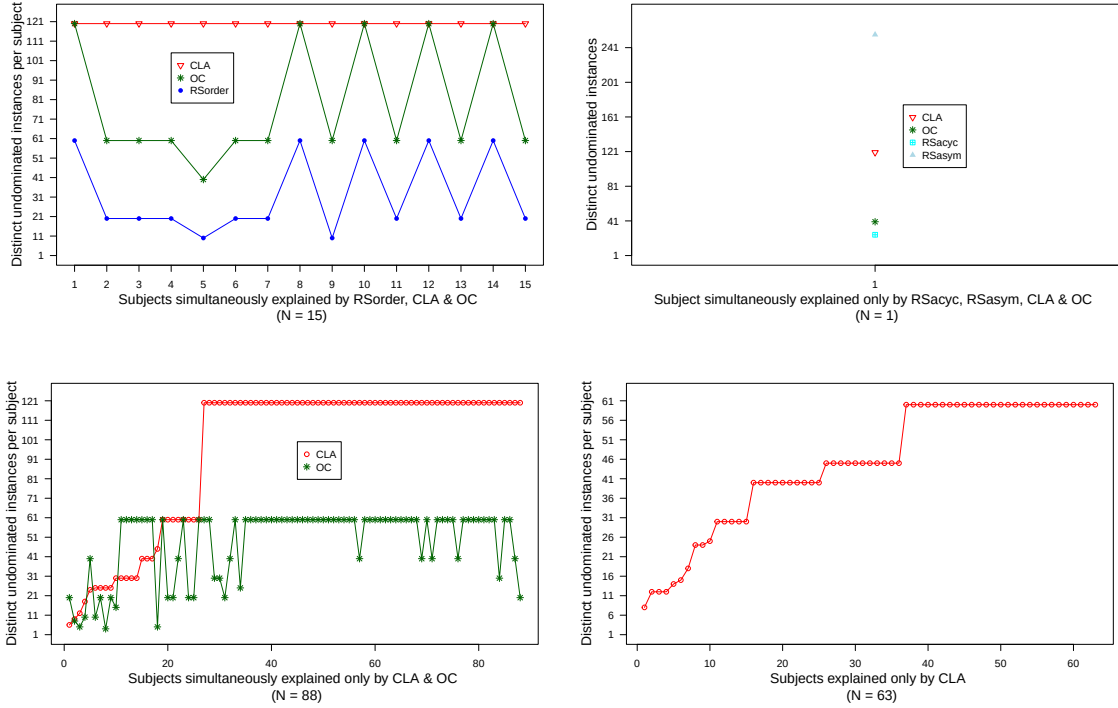


Table 3: Average and median undominated model primitives and computation times, and the corresponding statistics for approximately compatible fits under Rational Choice.

(a) Four Dresses

	RS_{order}	RS_{acyc}	RS_{asym}	CLA	OC	RC
Compatible subjects	5%	5%	5%	91.5%	62%	(100%)
Mean (median) primitives per matching model	24	480	768	1932 (1643)	664 (638)	1.51 (1)
Mean (median) undominated primitives per matching model	4 (·)	12 (·)	32 (·)	8.01 (8)	10.9 (12)	
Mean (median) total computation time, in seconds	0.0001 (·)	0.003 (·)	0.007 (·)	0.078 (0.075)	0.081 (0.079)	0.0007 (0.0002)

(b) Five Ponchos

	RS_{order}	RS_{acyc}	RS_{asym}	CLA	OC	RC
Compatible subjects	9%	9.5%	9.5%	100%	62%	(100%)
Mean (median) undominated primitives per matching model	32 (20)	94.7 (60)	640 (512)	78.8 (60)	52.8 (60)	1.57 (1)
Mean (median) total computation time, in seconds	0.053 (0.049)	0.142 (0.097)	5.48 (3.05)	0.295 (0.23)	0.747 (0.794)	0.038 (·)

Note: (·) means that the median coincides with the mean.

Table 4: Percentage of subjects whose perfect or approximate model fits were unlikely under random behavior (Bronars, 1987)*.

(a) Four Dresses

RC	68.3%	(2)
RS_{order}	5%	(1)
RS_{acyc}	5%	(1)
RS_{asym}	5%	(1)
CLA	0%	(0)
OC	0%	(0)

(b) Five Ponchos

RC	79.6%	(3)
RS_{order}	64.6%	(2)
RS_{acyc}	9.6%	(1)
RS_{asym}	9.6%	(1)
CLA	0%	(0)
OC	0%	(0)

*In parenthesis are the 2.5th percentile model scores derived from simulations.

on the same data. Indeed, although the rational-choice model is fully identified when choices from all possible menus are available (as is the case here), since our empirical analysis is limited to subjects who can only be *approximately* explained by **RC**, there is no *a priori* guarantee that any such approximate fits will also be generated by a single preference ordering per subject. Yet this is indeed what we find: an average (median) number of 1.51 (1) and 1.57 (1) strict preference relations per subject in the 4D and 5P data, respectively, almost always corresponding to approximate fits with a distance score of 1 or 2. In other words, with the exception of few subjects, the *approximate* **RC** fit generally gives more focused welfare-relevant information than the *perfect* fit of every bounded-rational model under study.

This finding, in turn, raises the question: Is it more appropriate for such human decision makers to be viewed as if they behaved perfectly in the ways described by **CLA**, **OC** and **RS**, or, instead, as “noisy” utility maximizers?¹⁵ While there may be little hope to arrive at a definitive and generally applicable answer, using synthetic data featuring random behavior in the way originally suggested in Bronars (1987) could be helpful in that direction. More specifically, we can use the model-specific distributions of distance scores on the *synthetic* datasets (which are available from the predictive-success analysis that was presented earlier) and, by analogy to the two-tailed statistical test at the 5% significance level, determine the cut-off scores that correspond to the 2.5th percentile values. Finally, once these are pinned down we can discard those *human* datasets whose model-specific scores weakly exceed the respective cut-off values, on the grounds that we cannot plausibly rule out that the (perfect or approximate) fits they give rise to could not have been achieved by a person making choices at random.

As is detailed in Table 4, the majority of approximate fits that are achieved by **RC** survive this robustness check in both the 4D and 5P experimental data. Furthermore, all perfect fits of all three **RS** model classes also survive, while the opposite is true for **CLA** and **OC** (see also the high a_M values in Table 2). Finally, this analysis suggests that the almost perfect fit achieved by the most structured version of the shortlisting model, **RS_{order}**, for 93 of 167 subjects in the 5P sample should also be considered acceptable according to that criterion. Although this method is useful in narrowing down the imperfect rational and perfect bounded-rational explanations

¹⁵A distinct related question is: How should an analyst decide which one of several bounded-rational models provides the best way to think about a choice dataset that is explained by all of them and may also be (in)compatible with **RC**? In particular, what is the role of non-choice data in such discerning attempts? We refer the reader to discussions in Caplin and Schotter (2008); Manzini and Mariotti (2014); de Clippel and Rozen (2024) for more on this.

that are simultaneously possible for a given subject, we note that it does not eliminate the need for additional scrutiny and discretion to be exercised by the analyst interested in arriving at a single best explanation.

6 Concluding Remarks

Many deterministic models of bounded-rational general choice have been developed in the past two decades with the overarching goal of explaining in introspectively intuitive ways behavioral phenomena that the textbook model of rational choice is unable to. The extensive theoretical interest in this direction, however, has not been met by an analogous interest in, and advances on, the rigorous empirical testing of these models in the kinds of general, riskless choice environments that motivated their development. The high computational complexity associated with in-depth analyses of these models is arguably one explanation for the relative absence of such work.¹⁶ In the present study we tackle this problem directly by introducing into the bounded-rationality and revealed-preference literatures constraint and constraint-dominance programming methods, and a comprehensive, free and open-source such tool chain from the frontier of research at the intersection of discrete optimization, computer science, artificial intelligence and operations research. This facilitates detailed—and in full generality—non-parametric goodness-of-fit tests of what are arguably two of the most prominent classes of such models, and assessing how informative such fits are to any analyst interested to use them toward inferring decision makers’ welfare-relevant information—namely, their shortlisting criteria or attention constraints and preferences—in practical choice situations.

The illustrative empirical application of our methods and tools on the two experimental datasets with riskless choices over four and five alternatives that we considered shows that limited-attention models explain better than their shortlisting counterparts, but are also more permissive. In addition, of all the models and special classes thereof that we analyzed, only shortlisting with transitive rationales achieves (perfect or imperfect) empirical fits that pass the “non-randomness” robustness test. This structured special case of the shortlisting model is also the one that, in both datasets, generates the most parsimonious fits —namely those with the smallest number of

¹⁶It is also one of the reasons for the increased theoretical interest in the formulation of bounded-rational *random* choice models more recently. Under distinct assumptions on the nature of the available data, and under additional identifying assumptions, these models have proved more amenable to empirical testing with econometric methods (Cattaneo, Ma, Masatlioglu, and Suleymanov, 2020; Dardanoni, Manzini, Mariotti, Petri, and Tyson, 2023; Filiz-Ozbay and Masatlioglu, 2023).

compatible behavioral primitives. Despite their higher explanatory ability, however, all bounded-rational choice models that we studied fared worse in terms of parsimony than the textbook model of rational choice, whose approximate empirical fits were typically associated with a unique compatible primitive/preference orderings in these data.

Concluding, we hope that this paper’s methodological contribution and empirical findings will aid the development of new bounded-rational models of general choice in which the—now measurable—degrees of permissiveness and explanatory indeterminacy will be evaluated alongside the models’ computational complexity and ability to explain observed behavioral phenomena.

References

- AKGÜN, Ö., A. M. FRISCH, I. P. GENT, C. JEFFERSON, I. MIGUEL, AND P. NIGHTINGALE (2022): “Conjure: Automatic Generation of Constraint Models from Problem Specifications,” *Artificial Intelligence*, 310, 103751.
- ANDREWS, I., D. FUDENBERG, L. LEI, A. LIANG, AND C. WU (2025): “The Transfer Performance of Economic Models,” *Working Paper*.
- APESTEGUIA, J. AND M. BALLESTER (2015): “A Measure of Rationality and Welfare,” *Journal of Political Economy*, 123, 1278–1310.
- APT, K. R. (2003): *Principles of Constraint Programming*, Cambridge University Press.
- AU, P. H. AND K. KAWAI (2011): “Sequentially Rationalizable Choice with Transitive Rationales,” *Games and Economic Behavior*, 73, 608–614.
- BEATTY, T. K. M. AND I. A. CRAWFORD (2011): “How Demanding is the Revealed Preference Approach to Demand?” *American Economic Review*, 101, 2782–2795.
- BERNHEIM, B. D. AND A. RANGEL (2009): “Beyond Revealed Preference: Choice-Theoretic Foundations for Behavioral Welfare Economics,” *Quarterly Journal of Economics*, 124, 51–104.
- BESSIERE, C. (2006): “Constraint Propagation,” in *Handbook of Constraint Programming*, ed. by F. Rossi, P. van Beek, and T. Walsh, Elsevier, chap. 3.
- BRONARS, S. G. (1987): “The Power of Non-parametric Tests of Preference Maximization,” *Econometrica*, 55, 693–698.

- CAPLIN, A., M. DEAN, AND D. MARTIN (2011): “Search and Satisficing,” *American Economic Review*, 101, 2899–2922.
- CAPLIN, A. AND A. SCHOTTER, eds. (2008): *The Foundations of Positive and Normative Economics: A Handbook*, New York: Oxford University Press.
- CATTANEO, M. D., X. MA, Y. MASATLIOGLU, AND E. SULEYMANOV (2020): “A Random Attention Model,” *Journal of Political Economy*, 128, 2796–2836.
- CHERCHYE, L., T. DEMUYNCK, B. DE ROCK, AND J. LANIER (2025): “Are Consumers (Approximately) Rational?: Shifting the Burden of Proof,” *The Review of Economics and Statistics*, 107, 1652–1666.
- CHOI, S., R. FISMAN, D. GALE, AND S. KARIV (2007): “Consistency and Heterogeneity of Individual Behavior under Uncertainty,” *American Economic Review*, 97, 1921–1938.
- CHU, G. AND P. J. STUCKEY (2015): “Dominance Breaking Constraints,” *Constraints*, 20, 155–182.
- COSTA-GOMES, M., C. CUEVA, G. GERASIMOU, AND M. TEJIŠČÁK (2022): “Choice, Deferral and Consistency,” *Quantitative Economics*, 13, 1297–1318.
- DARDANONI, V., P. MANZINI, M. MARIOTTI, H. PETRI, AND C. J. TYSON (2023): “Mixture Choice Data: Revealing Preferences and Cognition,” *Journal of Political Economy*, 131, 687–715.
- DE CLIPPEL, G. AND K. ROZEN (2021): “Bounded Rationality and Limited Datasets,” *Theoretical Economics*, 16, 359–380.
- (2024): “Bounded Rationality in Choice Theory: A Survey,” *Journal of Economic Literature*, 62, 995–1039.
- DEAN, M. AND D. MARTIN (2016): “Measuring Rationality with the Minimum Cost of Revealed Preference Violations,” *The Review of Economics and Statistics*, 98, 524–534.
- DECHTER, R. (2003): *Constraint Processing*, Morgan Kaufmann.
- DEMETRY, M. AND P. HJERTSTRAND (2023): “Consistent Subsets: Computing the Houtman–Maks Index in Stata,” *The Stata Journal*, 23, 578–588.
- DEMUYNCK, T. (2011): “The Computational Complexity of Rationalizing Boundedly Rational Choice Behavior,” *Journal of Mathematical Economics*, 47, 425–433.

- (2015): “Statistical Inference for Measures of Predictive Success,” *Theory and Decision*, 79, 689–699.
- DEMUYNCK, T. AND J. REHBECK (2023): “Computing Revealed Preference Goodness-of-Fit Measures with Integer Programming,” *Economic Theory*, 76, 1175–1195.
- DUTTA, R. AND S. HORAN (2015): “Inferring Rationales from Choice,” *American Economic Journal: Microeconomics*, 7, 179–201.
- ELLIS, A. AND D. J. FREEMAN (2024): “Revealing Choice Bracketing,” *American Economic Review*, 114, 2668–2700.
- FILIZ-OZBAY, E. AND Y. MASATLIOGLU (2023): “Progressive Random Choice,” *Journal of Political Economy*, 131, 716–750.
- FRÉCHETTE, G., E. VESPA, AND S. YUKSEL (2026): “People as Intuitive Modelers: How Model Complexity Adapts to Data,” *Working Paper*.
- FREER, M. AND H. NOSRATABADI (2024): “On the Welfare (Ir)Relevance of Two-Stage Models,” *arXiv:2411.08263*.
- (2026): “Revealed Preference Analysis Under Limited Attention,” *Journal of Economic Behavior & Organization*, 246, 107568.
- FRISCH, A. M., W. HARVEY, C. JEFFERSON, B. MARTÍNEZ-HERNÁNDEZ, AND I. MIGUEL (2008): “Essence: A Constraint Language for Specifying Combinatorial Problems,” *Constraints*, 13, 268–306.
- FUDENBERG, D., W. Y. GAO, AND Z. YOU (2026a): “Model Restrictiveness in Functional and Structural Settings,” *arXiv Working Paper 2602.07688*.
- FUDENBERG, D., J. KLEINBERG, A. LIANG, AND S. MULLAINATHAN (2022): “Measuring the Completeness of Economic Models,” *Journal of Political Economy*, 130, 956–990.
- FUDENBERG, D., A. LIANG, AND W. GAO (2026b): “How Flexible is that Functional Form? Quantifying the Restrictiveness of Theories,” *The Review of Economics and Statistics*, 108, 194–209.
- GENT, I. P., C. JEFFERSON, AND I. MIGUEL (2006): “Minion: A Fast, Scalable, Constraint Solver,” in *Proceedings of the 17th European Conference on Artificial Intelligence*, IOS Press.

- GERASIMOU, G. (2026): “Eliciting and Distinguishing Between Weak and Incomplete Preferences: Theory, Experiment and Computation,” *Games and Economic Behavior*, 158, 737–762.
- GERASIMOU, G. AND M. TEJIŠČÁK (2018): “Prest: Open-Source Software for Computational Revealed Preference Analysis,” *Journal of Open Source Software*, 3(30), 1015.
- GUNS, T., P. J. STUCKEY, AND G. TACK (2018): “Solution Dominance over Constraint Satisfaction Problems,” *CoRR*, abs/1812.09207.
- HEUFER, J. AND P. HJERTSTRAND (2015): “Consistent Subsets: Computationally Feasible Methods to Compute the Houtman-Maks Index,” *Economics Letters*, 128.
- HOUTMAN, M. AND J. A. MAK (1985): “Determining All Maximal Data Subsets Consistent with Revealed Preference,” *Kwantitatieve Methoden*, 19, 89–104.
- INOUE, Y. AND K. SHIRAI (2023): “On the Observable Restrictions of Limited Consideration Models: Theory and Application,” *Economic Theory*, 75, 695–715.
- KALAI, G., A. RUBINSTEIN, AND R. SPIEGLER (2002): “Rationalizing Choice Functions by Multiple Rationales,” *Econometrica*, 70, 2481–2488.
- KREPS, D. M. (1990): *A Course in Microeconomic Theory*, Princeton: Princeton University Press.
- (2012): *Microeconomic Foundations I: Choice and Competitive Markets*, Princeton: Princeton University Press.
- LIANG, A. (2026): “Using Machine Learning to Generate, Clarify, and Improve Economic Models,” *Journal of Economic Literature*, forthcoming.
- LLERAS, J. S., Y. MASATLIOGLU, D. NAKAJIMA, AND E. Y. OZBAY (2017): “When More is Less: Limited Consideration,” *Journal of Economic Theory*, 170, 70–85.
- MACKWORTH, A. K. (1977): “Consistency in Networks of Relations,” *Artificial Intelligence*, 8, 99–118.
- MANZINI, P. (2023): “Mixture Choice Function Datasets,” *author’s personal website*.
- MANZINI, P. AND M. MARIOTTI (2007): “Sequentially Rationalizable Choice,” *American Economic Review*, 97, 1824–1839.

- (2014): “Welfare Economics and Bounded Rationality: The Case for Model-Based Approaches,” *Journal of Economic Methodology*, 21, 343–360.
- MAS-COLELL, A., M. D. WHINSTON, AND J. R. GREEN (1995): *Microeconomic Theory*, New York: Oxford University Press.
- MASATLIOGLU, Y., D. NAKAJIMA, AND E. Y. OZBAY (2012): “Revealed Attention,” *American Economic Review*, 102, 2183–2205.
- MASATLIOGLU, Y. AND E. A. OK (2005): “Rational Choice with Status Quo Bias,” *Journal of Economic Theory*, 121, 1–29.
- NIGHTINGALE, P., ÖZGÜR AKGÜN, I. P. GENT, C. JEFFERSON, AND I. MIGUEL (2014): “Automatically Improving Constraint Models in Savile Row through Associative-Commutative Common Subexpression Elimination,” in *Proceedings of the 20th International Conference on Principles and Practice of Constraint Programming (CP 2014)*, Springer, 590–605.
- OK, E. A., P. ORTOLEVA, AND G. RIELLA (2015): “Revealed (P)reference Theory,” *American Economic Review*, 105, 299–321.
- RÉGIN, J.-C. (2005): “Global Constraints and Filtering Algorithms,” in *Global Constraints*, Wiley, preprint available at <https://www.constraint-programming.com/people/regin/papers/global.pdf>.
- RICHTER, M. (2020): “Choice Theory via Equivalence,” *Working Paper*.
- ROSSI, F., P. VAN BEEK, AND T. WALSH (2006): “Introduction,” in *Handbook of Constraint Programming*, ed. by F. Rossi, P. van Beek, and T. Walsh, Elsevier, chap. 1.
- RUBINSTEIN, A. (2006): *Lecture Notes in Microeconomic Theory*, Princeton, NJ: Princeton University Press.
- RUBINSTEIN, A. AND Y. SALANT (2008): “Some Thoughts on the Principle of Revealed Preference,” in *The Foundations of Positive and Normative Economics: A Handbook*, ed. by A. Caplin and A. Schotter, New York: Oxford University Press, 115–124.
- (2012): “Eliciting Welfare Preferences from Behavioural Data Sets,” *Review of Economic Studies*, 79, 375–387.

- SALANT, Y. AND A. RUBINSTEIN (2008): “ (A, f) : Choice with Frames,” *Review of Economic Studies*, 75, 1287–1296.
- SELTEN, R. (1991): “Properties of a Measure of Predictive Success,” *Mathematical Social Sciences*, 21, 153–167.
- SMEULDERS, B., F. C. R. SPIEKMA, L. CHERCHYE, AND B. DE ROCK (2014): “Goodness-of-Fit Measures for Revealed Preference Tests: Complexity Results and Algorithms,” *ACM Transactions on Economics and Computation*, 2, 3:1–3:16.
- VAN BEEK, P. (2006): “Backtracking Search Algorithms,” in *Handbook of Constraint Programming*, ed. by F. Rossi, P. van Beek, and T. Walsh, Elsevier, chap. 4.

Appendix

A All compatible primitives in the Example

RS solutions:

1. $y \succ_1 z$ & $z \succ_2 x, x \succ_2 y$ ($\mathbf{RS}_{acyc} \setminus \mathbf{RS}_{order}$)
2. $y \succ_1 z$ & $z \succ_2 x, x \succ_2 y, z \succ_2 y$ ($\mathbf{RS}_{order}$)
3. $y \succ_1 z$ & $z \succ_2 x, x \succ_2 y, y \succ_2 z$ ($\mathbf{RS}_{asym} \setminus \mathbf{RS}_{acyc}$)
4. $x \succ_1 y, y \succ_1 z$ & $z \succ_2 x$ ($\mathbf{RS}_{acyc} \setminus \mathbf{RS}_{order}$)
5. $x \succ_1 y, y \succ_1 z$ & $z \succ_2 x, z \succ_2 y$ ($\mathbf{RS}_{acyc} \setminus \mathbf{RS}_{order}$)
6. $x \succ_1 y, y \succ_1 z$ & $y \succ_2 z, z \succ_2 x$ ($\mathbf{RS}_{acyc} \setminus \mathbf{RS}_{order}$)
7. $x \succ_1 y, y \succ_1 z$ & $z \succ_2 x, y \succ_2 x$ ($\mathbf{RS}_{acyc} \setminus \mathbf{RS}_{order}$)
8. $x \succ_1 y, y \succ_1 z$ & $z \succ_2 y, y \succ_2 x, z \succ_2 x$ ($\mathbf{RS}_{acyc} \setminus \mathbf{RS}_{order}$)
9. $x \succ_1 y, y \succ_1 z$ & $y \succ_2 z, z \succ_2 x, y \succ_2 x$ ($\mathbf{RS}_{acyc} \setminus \mathbf{RS}_{order}$)
10. $x \succ_1 y, y \succ_1 z$ & $z \succ_2 x, x \succ_2 y$ ($\mathbf{RS}_{acyc} \setminus \mathbf{RS}_{order}$)
11. $x \succ_1 y, y \succ_1 z$ & $z \succ_2 x, x \succ_2 y, z \succ_2 y$ ($\mathbf{RS}_{acyc} \setminus \mathbf{RS}_{order}$)
12. $x \succ_1 y, y \succ_1 z$ & $z \succ_2 x, x \succ_2 y, y \succ_2 z$ ($\mathbf{RS}_{asym} \setminus \mathbf{RS}_{acyc}$)

CLA solutions:

1. $x \succ y \succ z$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x, y\}, \{y, z\}, \{z\}, \{x, y\})$
2. $x \succ y \succ z$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x, y\}, \{y\}, \{z\}, \{x, y\})$
3. $x \succ y \succ z$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x\}, \{y\}, \{z\}, \{x, y, z\})$
4. $x \succ y \succ z$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x\}, \{y, z\}, \{z\}, \{x, y, z\})$
5. $x \succ y \succ z$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x, y\}, \{y\}, \{z\}, \{x, y, z\})$
6. $x \succ y \succ z$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x, y\}, \{y, z\}, \{z\}, \{x, y, z\})$
7. $x \succ z \succ y$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x, y\}, \{y\}, \{z\}, \{x, y\})$
8. $x \succ z \succ y$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x\}, \{y\}, \{z\}, \{x, y, z\})$
9. $x \succ z \succ y$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x, y\}, \{y\}, \{z\}, \{x, y, z\})$
10. $z \succ x \succ y$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x, y\}, \{y\}, \{z\}, \{x, y\})$
11. $z \succ x \succ y$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x, y\}, \{y\}, \{x, z\}, \{x, y\})$

OC solutions:

1. $y \succ z \succ x$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x\}, \{y\}, \{x, z\}, \{x\})$
2. $y \succ z \succ x$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x\}, \{y, z\}, \{x, z\}, \{x\})$
3. $z \succ y \succ x$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x\}, \{y\}, \{x, z\}, \{x\})$
4. $z \succ x \succ y$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x\}, \{y\}, \{x, z\}, \{x\})$
5. $z \succ x \succ y$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x, y\}, \{y\}, \{x, z\}, \{x\})$
6. $z \succ x \succ y$ & $(A_1, A_2, A_3, A_4) \xrightarrow{\Gamma} (\{x, y\}, \{y\}, \{x, z\}, \{x, y\})$

◆

B All *undominated* model primitives in the Example

Applying the dominance notions of Definitions 1 and 2 to the list of perfectly compatible instances that are enumerated in Section A in relation to the example choices in (5) leads to the following refined elicitations per model or class thereof (the # numberings follow the respective list orders of that section):

$\text{RS}_{\text{order}}^{\text{undominated}}$: #2

$\text{RS}_{\text{acyc}}^{\text{undominated}}$: #8 #9 #11
 (dominates #5,7) (dominates #6) (dominates #1,4,10)

$\text{RS}_{\text{asym}}^{\text{undominated}}$: #12
 (dominates #3)

$\text{CLA}^{\text{undominated}}$: #6 #9 #11
 (dominates #1–5) (dominates #7,8) (dominates #10)

$\text{OC}^{\text{undominated}}$: #2 #3 #6
 (dominates #1) (dominates #4,5)

◆

C Executed CP models and example dataset

Essence encoding of a choice dataset

Saved as a `.param` file. Comments preceded by `$`. Indentation for illustration only.

```
letting dataset be sequence(  
  ({2,4,3,1},4),      $ item '4' is chosen at menu '{2,4,3,1}'  
  ({2,4,1},4),  
  ({2,3,1},3),  
  ({3,1,4},4),  
  ({4,3,2},3),  
  ({1,2},1),  
  ({4,1},4),  
  ({4,2},4),  
  ({3,2},3),  
  ({1,3},3),  
  ({3,4},4))
```

Essence model for Rational Shortlisting

Saved as an `.essence` file. Comments preceded by `$`. Indentation for illustration only.

```
given dataset : sequence of (set of int, int)  
letting n be max([ x | (-, (xs, -)) <- dataset, x <- xs ])  
letting options be domain int(1..n)  
letting nbObservations be |dataset|  
find omitList : sequence (size nbObservations) of bool  
  
$ (i, j) means i is preferred to j by asymmetric relation rel1  
find rel1 : relation (minSize 1, asymmetric) of (options * options)  
  
$ For each entry in the dataset, we will record the survivors after applying rel1.  
$ adding minSize 1 to the set here disallows Rational Choice solutions  
find survivor1 : sequence (size nbObservations) of set of options  
find rel2 : relation (minSize 1, asymmetric) of (options * options)
```

```

$ For each entry in the survivor1 set, we will record the survivors after applying
rel2.
$ size 1 - single valued choice
find survivor2 : sequence (size nbObservations) of set (size 1) of options
such that
  forAll row : int(1..nbObservations) .
    survivor1(row) subsetEq dataset(row)[1] /\
    ((omitList(row) /\ !(dataset(row)[2] in survivor1(row))) \/\
    (!omitList(row) /\ (dataset(row)[2] in survivor1(row)))) /\
    $ Element survives if it is not dominated by any other element of the set.
    forAll x in dataset(row)[1] .
      x in survivor1(row) <=> forAll y in dataset(row)[1] . !rel1(y,x)

such that
  forAll row : int(1..nbObservations) .
    survivor2(row) subsetEq survivor1(row) /\
    ((omitList(row) /\ !(dataset(row)[2] in survivor2(row))) \/\
    (!omitList(row) /\ (dataset(row)[2] in survivor2(row)))) /\
    forAll x in survivor1(row) .
      x in survivor2(row) <=> forAll y in survivor1(row) . !rel2(y,x)

$ Distance score is min number of elements to be removed from dataset
$ to comply perfectly with the model.
find score : int(0..nbObservations)
such that score = sum row : int(1..nbObservations) . toInt(omitList(row))

$ Unique solutions are defined by the preference relations, not by survivor sets.
branching on [rel1, rel2]

```

Essence model of Choice with Limited Attention

Saved as an .essence file. Comments preceded by \$. Indentation for illustration only.

```
given dataset : sequence of (set of int, int)
letting n be max([ x | (_, (xs, _)) <- dataset, x <- xs ])
letting options be domain int(1..n)

letting nbObservations be |dataset|

find omitList : sequence (size nbObservations) of bool

$ Maximal size to find all complete solutions.
find G : function (total, size 2**|options|-1) set (minSize 1) of options
--> set (minSize 1) of options

$ (i, j) means i is preferred to j under strict linear order rel
find rel : relation (aSymmetric, connex, transitive) of (options * options)

$  $\Gamma(A_i)$  is a non-empty subset of  $A_i$ 
such that forAll (menu, filter) in G .
    filter subsetEq menu /\ |filter| > 0 /\ |menu| > 0

$ If an element  $x$  is not in  $\Gamma(A_i)$  then  $\Gamma(A_i) = \Gamma(A_i \setminus \{x\})$ 
such that forAll row : int(1..nbObservations) .
    forAll x in dataset(row)[1] .
        !(x in G(dataset(row)[1])) -> (G(dataset(row)[1]) =
            G((dataset(row)[1] - x)))

$ Choice,  $x$ , is selected from  $\Gamma(A_i)$  where no other element of  $\Gamma(A_i) \succ x$ 
such that forAll row : int(1..nbObservations) .
    ((omitList(row) /\ !(dataset(row)[2] in G(dataset(row)[1]))) /\
    (!omitList(row) /\ (dataset(row)[2] in G(dataset(row)[1])))) /\
    (forAll (i, j) in rel . (i in G(dataset(row)[1]) /\ j in G(dataset(row)[1]))
    -> dataset(row)[2] != j)

find score : int(0..nbObservations)
such that score = sum row : int(1..nbObservations) . toInt(omitList(row))
```

Essence model for Overwhelming Choice

Saved as an .essence file. Comments preceded by \$. Indentation for illustration only.

```
given dataset : sequence of (set of int, int)
letting n be max([ x | (_, (xs, _)) <- dataset, x <- xs ])
letting options be domain int(1..n)
letting nbObservations be |dataset|
find omitList : sequence (size nbObservations) of bool
find G : function (total, size 2**|options|-1) set (minSize 1) of options
--> set (minSize 1) of options

$(i, j) means i is preferred to j under strict linear order rel
find rel : relation (aSymmetric, connex, transitive) of (options * options)

$  $\Gamma(A_i)$  is a non-empty subset of  $A_i$ 
such that forAll (menu, filter) in G .
    filter subsetEq menu /\ |filter| > 0 /\ |menu| > 0

$ Choice,  $x$ , is selected from  $\Gamma(A_i)$  where no other element of  $\Gamma(A_i) \succ x$ 
such that forAll row : int(1..nbObservations) .
    ((omitList(row) /\ !(dataset(row)[2] in G(dataset(row)[1]))) /\
    (!omitList(row) /\ (dataset(row)[2] in G(dataset(row)[1])))) /\
    (forAll (i, j) in rel . (i in G(dataset(row)[1]) /\ j in G(dataset(row)[1]))
    -> dataset(row)[2] != j)

$ Extra structure constraint: if  $S \subset T$ , then  $(\Gamma(T) \cap S) \subseteq \Gamma(S)$ 
such that forAll (menuT, filterT) in G .
    forAll (menuS, filterS) in G .
        menuS subset menuT -> (filterT intersect menuS) subsetEq filterS

find score : int(0..nbObservations)
such that score = sum row : int(1..nbObservations) . toInt(omitList(row))
```