

Working Paper Series

No. 08-2026

Identifying Behavioral Types

Christopher Kops

Maastricht University

j.kops@maastrichtuniversity.nl

Paola Manzini

University of Bristol

p.manzini@bristol.ac.uk

Marco Mariotti

Queen Mary University of London

m.mariotti@qmul.ac.uk

Illia Pasichnichenko

University of Sussex

i.pasichnichenko@sussex.ac.uk

Abstract: We study identification in models of aggregate choice generated by unobserved behavioral types. An analyst observes only aggregate choice behavior, while the population distribution of types and their type-level choice patterns are latent. Assuming only minimal and purely qualitative prior knowledge of the process generating type-level choice probabilities, we provide necessary and sufficient conditions for identifiability. Identification obtains if and only if the data exhibit sufficient cross-type behavioral heterogeneity, which we characterize equivalently through combinatorial matching conditions between types and alternatives, and through a behavioral condition that forbids certain choice patterns at type-level.

JEL codes: D81, D83, D91

Key words: revealed preferences, stochastic choice, bounded rationality, attention, characteristics, incomplete preferences.

Identifying Behavioral Types

Christopher Kops, Paola Manzini, Marco Mariotti, Illia Pasichnichenko*

July 17, 2026

Abstract

We study identification in models of aggregate choice generated by unobserved behavioral types. An analyst observes only aggregate choice behavior, while the population distribution of types and their type-level choice patterns are latent. Assuming only minimal and purely qualitative prior knowledge of the process generating type-level choice probabilities, we provide necessary and sufficient conditions for identifiability. Identification obtains if and only if the data exhibit sufficient cross-type behavioral heterogeneity, which we characterize equivalently through combinatorial matching conditions between types and alternatives, and through a behavioral condition that forbids certain choice patterns at type-level.

JEL codes: D81, D83, D91

Keywords: revealed preferences, stochastic choice, bounded rationality, attention, characteristics, incomplete preferences.

1 Introduction

Latent factors such as preferences, attention, and cognitive capacity shape choice behavior. These factors are typically heterogeneous—across individuals or within a decision maker observed over time—giving rise to distinct *behavioral types*. Even when type-level choices are unobservable, due to privacy, anonymization, or data scale, they may still generate observable aggregate choice data. When can one infer the distribution of behavioral types, or their type-specific choice patterns, from aggregate data alone? This paper characterizes when such inferences are possible, assuming minimal and purely qualitative prior knowledge of the process generating type-level choice probabilities.

*Kops: Maastricht University, the Netherlands (j.kops@maastrichtuniversity.nl); Manzini: University of Bristol, United Kingdom (p.manzini@bristol.ac.uk) and IZA Bonn; Mariotti: Queen Mary University of London, United Kingdom (m.mariotti@qmul.ac.uk) and Deakin University; Pasichnichenko: University of Sussex, United Kingdom (i.pasichnichenko@sussex.ac.uk). We thank for comments Alessandro Iaria and audiences at Catania, Sussex, Maastricht, Karlsruhe, BRIC 2025 at ITAM, the 2025 Unforeseen Contingencies Workshop at UGA, D-TEA 2026 at HEC Paris and SDM 2026 at HKUST.

To illustrate the motivation, consider a policymaker contemplating a stimulus that expands consumers’ choice sets. Assessing such a policy requires knowledge of the distribution of consumer types. Some types attend to the full range of options, others focus on a subset, and still others adhere to their current choice. The overall impact depends on the prevalence of these types and on how each responds to the expanded choice set. More generally, policies based solely on average responses may perform poorly because the “average” type need not correspond to any actual individual. Failing to identify the relevant types can therefore lead to inappropriate interventions, with some groups receiving treatments that are ineffective or harmful.

We study a general stochastic choice framework. The analyst observes choice probabilities (idealized empirical shares) over a finite set of n alternatives. These shares result from the aggregation of the unobserved stochastic choices made by a finite number of behavioral types. Formally, the model is summarized by

$$p = M\pi \tag{1}$$

where p is an observed $n \times 1$ vector of choice probabilities, and M is an $n \times r$ matrix of partially unknown type-level choice probabilities. A type t is a probability distribution over the alternatives, and $M(k, t)$ denotes the probability that type t chooses alternative k . A theory of choice specifies the entries of M . Finally, π is an unobserved $r \times 1$ vector describing the distribution of types. Because stochastic choice data are often observed across multiple choice occasions, we also consider a multi-occasion extension of Equation (1).

We impose minimal requirements on the analyst’s prior knowledge of the choice process. Only the *possibility relation* governing type-level behavior (specifying which types may choose an alternative with positive probability and which cannot) is assumed known. Such information is often implicit in the definition of behavioral types: vegan consumers do not choose animal products; risk-aversers avoid junk bonds; information avoiders omit health screenings or ignore adverse media. At a theoretical level, models may likewise restrict the support of type-level randomization, for example by allowing randomization only within a consideration set or among maximal elements of an incomplete preference relation (detailed examples are found in Section 6).

Our focus is on how the data in p identify the unknown parameters, namely the elements of π and the positive elements of M . Under the informational assumption discussed above, we derive identification results in three settings. First, we characterize identification of the type distribution π without imposing additional structure on type-level choice behavior.

Second, we specialize the framework to a type–state model, where each type’s stochastic

behavior depends on an unobserved common state. Aggregate choice probabilities then arise as weighted sums of state-dependent choice matrices:

$$p = \left(\sum_{a \in A} f(a) M^a \right) \pi$$

where A is the set of states, $f(a)$ the probability of state a , and M^a a matrix of 0 – 1 type-level choice probabilities in state a . This model yields additional identification conditions for π .

Finally, Section 5 shows that observing choices across multiple occasions strengthens identification by introducing cross-temporal variation. We combine the one-shot identification conditions with a fundamental tensor decomposition result, obtaining sufficient conditions for identifying both the type distribution and the type-level choice probabilities.

Our results can be understood as formalizing a requirement of sufficient *behavioral heterogeneity*. They fall into two classes. The first provides a combinatorial characterization based on *matchings* between types and alternatives consistent with the possibility relation, capturing heterogeneity via the *potential* association of distinct types with distinct choices. The second presents a behavioral characterization at the level of choice data: identifiability is impossible precisely when there exist certain cycles of types and corresponding choices that allow perturbations of the type distribution to leave aggregate choice frequencies unchanged.

The simplest case, a single choice occasion with two states, two types, and two alternatives, already illustrates the multifaceted nature of our results. Generic identifiability of the type distribution is equivalent here to either: (i) in at least one of the two states there exists a matching between types and alternatives; (ii) some state *separates* the two types, meaning their behaviors differ in that state.

Overall, the results yield an exhaustive characterization of identifiability. The underlying notion of behavioral heterogeneity is qualitative: the conditions depend only on the pattern of structural zeros, not on the magnitudes of type-level choice probabilities. Equivalently, the analysis depends solely on the sign pattern of M . The scope of the results is illustrated by two extended examples in Section 6. The first examines Lancaster-style choice over vectors of characteristics, while the second studies a model of incomplete preferences.

1.1 Related literature

Dardanoni *et al.* [11] has a similar motivation to ours, in that it studies generic type identification from discrete choice data generated by a single menu without additional covariate information, linking identifiability to the invertibility of an associated matrix. However their

analysis is concerned with a specific parametric model of choice with consideration sets, in which types are indexed by consideration capacities, that is by the maximum number of alternatives a decision maker can consider. In that setting, the matrix representation of the model has a known and highly structured form. A more econometrics-oriented work in the same vein, which however makes use of covariates besides choice observations from a single menu with limited consideration (or limited feasibility) is Barseghyan *et al.* [6]. By contrast, the framework developed here is deliberately abstract and qualitative. Our objective is not to analyze a particular behavioral model, but to accommodate a broad class of models as special cases and characterize identification at this level of generality and exploiting only the qualitative nature of the process of choice.

A recent strand of the literature studies the identifiability of the preference distribution from choice data within Random Utility Maximization (RUM) frameworks. Contributions in this vein include Apesteguia *et al.* [3], Chambers and Turansick [9], Doignon and Saito [13], Filiz-Ozbay and Masatlioglu [16], Lu [26], Manzini and Mariotti [28],¹ Turansick [33], and Yildiz [34] (see also the important earlier piece by Fishburn [17]). Our model shares with this literature a broad objective within a stochastic choice framework, namely the identification of underlying heterogeneity from observed choices. Our analysis departs from this literature in two main respects. First, we do not take preferences, rankings, or general choice functions as primitives; instead, we work with abstract “behavioral types.” Second, standard RUM approaches typically take as observable a random choice rule specifying a distribution of choices from each of several overlapping menus (distinct subsets of a common grand set). By contrast, the core of our analysis relies on the observed choice distribution from a single set of alternatives, although we also provide an extension to a parsimonious multi-occasion framework. These differences necessitate distinct analytical tools and arguments.

Caradonna and Turansick [8], like the papers above, also focus on RUM identification, but this is mainly for expositional purposes. Their work develops new general techniques for identification (related to the unique decomposition of flows on directed acyclic graphs into measures over paths) that extend to more general random choice models, as well as dynamic discrete choice models. Once again, a main difference from our approach lies in their use of a random choice rule as a primitive.

Progress on identification with more limited data than a full random choice rule can be credited to Azrieli and Rehbeck [5]. Their framework is one of *marginal stochastic choice*, a setting in which only the probability of choosing each alternative and the probability of each menu’s realization are observed.

¹See also the Correction in Manzini *et al.* [29].

Our multi-occasion extension of the main framework shares methodological affinities with Allman *et al.* [1] and Dardanoni *et al.* [12], which in turn build on classical tensor decomposition results by Kruskal [23], with a recent proof by Rhodes [32].

At a much broader level our work also relates to latent structure models, running from finite mixture models (McLachlan and Peel [31]) to nonparametric mixture models (Lindsay [25]), and hidden Markov models (Cappé *et al.* [7], Ephraim and Merhav [15]). The identification problem of mapping observed choice data back to latent behavioral types finds its formal roots in the seminal contributions of Koopmans and Reiersøl [20] and Koopmans [21]. Traditionally, the study of identifiability assumes a hypothetical, exact knowledge of the distribution of observables. The central question then is whether the underlying structural parameters can be uniquely recovered from this distribution. While such identification is distinct from statistical inference, it remains a prerequisite for it, as non-identifiable parameters render consistent estimation impossible. While we focus entirely on a qualitative notion of behavioral heterogeneity, a complementary approach is taken by Apesteguia and Ballester [2], who study the *measurement* of behavioral heterogeneity in environments closely related to those underlying our identification analysis.

2 Framework and definitions

Let $T = \{t_1, \dots, t_r\}$ and $X = \{x_1, \dots, x_n\}$ be finite sets of types and alternatives, respectively.² The analyst observes choice shares p for all alternatives in X . Let M be an $n \times r$ matrix with a typical column containing the choice shares for a specific type, so that the data can be described by Equation (1).

As in the RUM framework, the model admits both population and intrapersonal interpretations. Under the population interpretation, type heterogeneity reflects variation in preferences or cognitive traits across individuals. Under the intrapersonal interpretation, types represent distinct choice-relevant states — preferential or cognitive — that a decision maker may occupy over time.

Definition 1. A type t_l is *possible at alternative* x_k if $M(k, l) > 0$. In this case we write $t_l \wedge x_k$, where $\wedge$ is a binary relation over the sets T and X .

The possibility relation $\wedge$ captures the analyst’s prior knowledge about the structural zeros of the type–conditional choice matrix M , while leaving unrestricted the magnitudes of the

²The number of types r may itself be identifiable. One can begin with an upper bound $r' > r$ and iteratively aggregate types that are observationally indistinguishable until no further aggregation is possible. The resulting partition then yields the minimal number r of behaviorally distinct types.

positive entries. Thus, we identify the parameter space of the model with a full-dimensional subset of $[0, 1]^{|\wedge|-1}$, where $r - 1$ dimensions describe π and $|\wedge| - r$ dimensions describe nonzero entries of M . The model's parameterization map is defined on the parameter space and takes values in $[0, 1]^n$. The model is *globally identifiable* if its parameterization map is injective.

We focus primarily on *generic identifiability*, under which the set of parameter values where identifiability fails has measure zero. This notion is typically sufficient for empirical analysis. We next summarize the algebraic-geometric concepts used throughout the paper; broader introductions include Allman *et al.* [1] and Cox *et al.* [10].

An algebraic variety is defined as the simultaneous zero-set of a finite collection of multivariate polynomials $\{\phi_i\}$ on $\mathbb{R}^m$. A variety is all of $\mathbb{R}^m$ only when all ϕ_i are 0; otherwise, a variety is called a proper subvariety and must be of dimension less than m , and, hence, of Lebesgue measure 0 in $\mathbb{R}^m$. The same is true if we replace $\mathbb{R}^m$ by a full-dimensional subset of $[0, 1]^m$. Intersections of algebraic varieties are algebraic varieties as they are the simultaneous zero-set of the unions of the original sets of polynomials. Finite unions of varieties are also varieties, since if sets G_1 and G_2 define varieties, then $\{fg | f \in G_1, g \in G_2\}$ defines their union. Given a set $\Theta \subseteq \mathbb{R}^m$ of full dimension, we say that a property holds *generically* on Θ if it holds for all points in Θ , except possibly for those on some proper subvariety. Thus, generic identifiability of parameters means that all non-identifiable parameter choices lie within a proper subvariety, which is a set of Lebesgue measure zero.

Finally, a model with parametrization map z is *structurally non-identifiable* if no point in Θ has a singleton fiber, that is, for all $\theta \in \Theta$, $|\mathcal{F}(\theta)| > 1$, where $\mathcal{F}(\theta) = \{\theta' \in \Theta | z(\theta') = z(\theta)\}$ is the fiber of z at θ .

Often we will refer to identifiability as a property of (a subset of) the parameters in the model.

3 Characterizing identifiability of the type distribution

In this sparse observational context, our primary aim is the identification of π . This means asking whether or not different latent type populations produce the same aggregate data, for any realization of the unknown positive choice probabilities consistent with the possibility relation. We will show that identifiability is equivalent to the types' choices being sufficiently heterogeneous in a qualitative sense. Specifically, we establish a connection between the identifiability of π and the possibility relation $\wedge$.

The following notion describes the possibility for types to choose distinct alternatives.

Definition 2. A (type-to-alternative) *matching* is an injective function $m : T \rightarrow X$ such that $t \wedge m(t)$ for all $t \in T$.

In our setting, identification of π in the model defined by Equation (1) reduces to whether the structural zero pattern encoded by $\wedge$ admits parameter configurations for which the columns of M are linearly independent. The first result shows that this condition admits a simple combinatorial interpretation in terms of matchings between types and alternatives (the proofs of all results are in Appendix A).

Theorem 1. *The parameters in π are generically identifiable if and only if there exists a matching $m : T \rightarrow X$. In addition, if a matching exists and it is unique, then the parameters in π are globally identifiable. If no matching exists, then the parameters in π are structurally non-identifiable.*

Note that the existence of a matching implies $n \geq r$, i.e., there are at least as many alternatives as types. The interpretation of the matching is that r agents of distinct types can always choose r distinct alternatives. In this case, the possibility relation $\wedge$ is sufficiently rich to permit full identification of the type distribution from observed choices. Equivalently, every non-trivial perturbation of the type distribution changes the aggregate choice frequencies for generic parameter values. Theorem 1 shows that the existence of a matching restricts the behavioral theories consistent with the possibility relation exactly to those in which identification arises almost always. The theorem also clarifies that the failure of identification on measure zero sets is attributable to the multiplicity of matchings: when the matching is unique, identification occurs for all values of the parameters. As we will show below, the converse is not true.

The necessity of a matching can be understood as follows. If no matching exists between types and alternatives, then there is a subset of types whose “collectively choosable” alternatives (that is, all those at which one or more types in the subset are possible) are too few to accommodate them all. As a result, any attempt to assign distinct alternatives to these types must necessarily assign at least two types to the same alternative, or assign some type to an alternative it can never choose. Translating this into matrix terms, the columns of M corresponding to these types must be supported on a set of rows that is not rich enough to ensure generic linear independence. This issue is *structural*: it does not depend on the particular values of the positive entries of M , but follows solely from the zero pattern implied by the possibility relation. No selection of the positive parameter values in M can, in this case, permit learning of π from the data.

Remark 1. Formally, the notion of a matching used here is equivalent to that of a perfect matching in a bipartite graph whose left and right vertices correspond to columns and rows of the matrix M , respectively, with edges induced by its nonzero entries. Hence, the existence of

a matching is equivalent to the condition in Hall's marriage theorem (Hall [19]). Translated to our context, the condition requires that for any set $S \subseteq T$ of types there exist at least $|S|$ distinct alternatives such that some type in S is possible at at least one of them. This further clarifies how the existence of a matching captures behavioral heterogeneity.

Example 1. Suppose that there are three types, labeled t_1 , t_2 , and t_3 , and three alternatives, labeled x , y , and z . A behavioral theory determines the possibility relation $\Delta = \{(t_1, x), (t_2, x), (t_2, z), (t_3, x), (t_3, y), (t_3, z)\}$. In turn this induces the zero entries of M :

$$M = \begin{matrix} & t_1 & t_2 & t_3 \\ \begin{matrix} x \\ y \\ z \end{matrix} & \begin{pmatrix} \bullet & \bullet & \bullet \\ 0 & 0 & \bullet \\ 0 & \bullet & \bullet \end{pmatrix} \end{matrix},$$

where black ($\bullet$) and red ($\bullet$) dots denote parameters. This encapsulates the information on which alternatives could possibly have been chosen by which types. The parametric description of this matrix takes values in $[0, 1]^3$. The columns of M can never be linearly dependent, whatever the values assigned to the dots. To build a matching, t_1 must be matched to x because this is the only alternative for which this type is possible. Then t_2 must be matched to z and t_3 to y . This matching (red) is unique. On the other hand, if type t_1 were possible at y , with $M(2, 1)$ positive, there would exist other matchings, such as t_1 matched to y , t_2 to x and t_3 to z . The columns could be linearly dependent for some values of the parameters, but the set of such values is of measure zero.

Example 2. Let each type be a consideration capacity, as in Dardanoni et al. [11], namely the maximum number of alternatives the agent can consider. Let $T = \{t_1, t_2, \dots, t_n\}$ with type t_k having consideration capacity k .³ Each type maximizes on the set of considered alternatives a known and common linear order $\succ$ (preference) on X . Type t_k considers each subset of X of cardinality k with a probability governed by an unknown full-support probability distribution on $\{K \subseteq X \mid |K| = k\}$. This model generalizes the fixed-preference model in [11], who assume the distribution on each $\{K \subseteq X \mid |K| = k\}$ to be uniform.

To build a matching, number the alternatives $x_1, x_2, \dots, x_n$ in decreasing order of preference. Type t_n , who considers the entire set X , chooses x_1 with probability one and must be matched to x_1 . Type t_{n-1} , who considers all but one alternative, chooses with positive probability only x_1 and x_2 , so that, in view of the previous matching, must be matched with x_2 . Continuing in this way, we build a (unique) matching, with type t_k matched to x_{n-k+1} for each k . It follows

³Here, type t_n stands for a composite type including all types that can consider *at least* n alternatives.

from Theorem 1 that the type distribution is identified. Proposition 2 in [11] is a special case of this result.⁴

3.1 Global identifiability

In order to characterize global identifiability, we introduce a new property of behavioral heterogeneity. The possibility relation satisfies collective behavioral separation if, for any two disjoint nonempty groups of types, the union of alternatives possible for one group differs from the union of alternatives possible for the other. Formally:

Definition 3. The possibility relation satisfies *collective behavioral separation* if there do not exist disjoint nonempty $T', T'' \subset T$ such that

$$\bigcup_{t \in T'} \{x|t \wedge x\} = \bigcup_{t \in T''} \{x|t \wedge x\}.$$

Collective behavioral separation fully characterizes global identifiability.

Theorem 2. *The parameters in π are globally identifiable if and only if there is collective behavioral separation.*

The proof of Theorem 2 starts from the observation that global non-identification is equivalent to the existence of a nontrivial linear dependence among the columns of a matrix that is compatible with the possibility relation. Exploiting the freedom to choose compatible columns as arbitrary full-support distributions over their feasible alternatives, this linear dependence is shown to arise if and only if two disjoint groups of types possess identical collective choosable sets.

The matching condition in Theorem 1, equivalent to Hall’s marriage condition (Remark 1), guarantees that every collection of types has access to sufficiently many distinct alternatives to permit their individual behavioral separation. Collective behavioral separation strengthens this perspective. Rather than counting possible alternatives available to groups of types, it requires directly that no two disjoint groups of types have identical sets of possible alternatives. Since global identification implies generic identification, collective behavioral separation implies the Hall (or matching) condition, but the converse need not hold.

⁴What is more, Proposition 6 in [11], showing identifiability for generic preference distributions in the extension of their model to heterogeneous preferences, is also implied by Theorem 1—the uniform distribution assumption on equal-sized consideration sets makes type-level choice probabilities structurally fixed for any given preference, so the only variable parameters are related to the preference distribution.

4 The type-state framework

We now impose additional structure common in many economic settings, enabling a more refined identification analysis. A stochastic state determines the nature of the matrix M , while each type's choice is deterministic conditional on the state. Let f be a probability measure on a known finite state space A . Each type t observes a realization of a according to f and chooses according to a “choice function” $c_t : A \rightarrow X$, where $c_t(a)$ denotes the choice made by type t in state a . The distribution f is independent of type.

For a given state $a \in A$ denote by M^a the corresponding $n \times r$ matrix of choice shares for each type. A column l of this matrix is a one-hot vector such that $M^a(k, l) = 1 \iff c_{t_l}(a) = x_k$. The matrix M is a convex combination of such matrices,

$$M = \sum_{a \in A} f(a) M^a. \quad (2)$$

A parametric description of M is given by $\{f(a)\}_{a \in A}$ taking values in $[0, 1]^{|A|-1}$.

Example 3. Each state a corresponds to a linear preference order $\succ_a$ over X , with probability $f(a)$. Types differ in *attention depth*: alternatives are ranked by salience, and type t_k searches only the k most salient alternatives. The columns of M^a therefore describe type-level choices at a given attention depth. Index alternatives in decreasing order of salience. Then the choice function of type t_k is

$$c_{t_k}(a) = \arg \max_{\succ_a} \{x_1, \dots, x_k\}, \quad k = 1, \dots, n,$$

so in state a , type t_k chooses the most-preferred alternative among the first k alternatives according to $\succ_a$. The resulting state-aggregated choice probabilities in M are given by Equation (2).

4.1 A “walk through” abstract example

Identification asks whether the aggregate outcome frequencies uniquely determine the underlying distribution of behavioral types. Equivalently, starting from some type distribution π , suppose that we perturb it slightly by transferring a small amount of probability mass between types while keeping the total probability equal to one. If every non-zero perturbation necessarily changes the aggregate outcome frequencies, then the type distribution is generically identifiable. Conversely, identification fails if there exists a non-zero redistribution of probability mass across types that leaves the aggregate observations completely unchanged. The key question in the type-state framework is therefore whether the effects of changing the type

distribution survive in the observables after the state-specific choice patterns have been aggregated across all states. The following example illustrates how the structure of the latent choice mappings determines whether such perturbations are visible or can instead remain hidden in the aggregate data.

Example 4. Suppose an analyst observes aggregate choice frequencies generated by four latent behavioral types t_1, t_2, t_3, t_4 and two latent states a and b . State a occurs with probability $f(a)$ and state b with probability $1 - f(a)$. The underlying state-dependent choices are

	a	b
t_1	x	w
t_2	y	w
t_3	z	x
t_4	z	y

Although neither the types nor the realized states are observed, this latent structure determines how changes in the type distribution affect the aggregate data.

To see this, imagine moving a small amount of probability mass ε from type t_3 to type t_4 . Since both types choose z in state a , this perturbation has no effect on the contribution of state a . In state b , however, type t_3 chooses x , whereas type t_4 chooses y . Consequently, the aggregate frequency of x decreases, while that of y increases, by $(1 - f(a))\varepsilon$. Could this change be concealed from the observer of aggregate data by further redistributions of probability mass? Any such redistribution must compensate for the change in the frequencies of x and y without creating new discrepancies elsewhere. Doing so generally requires further adjustments involving other types, which in turn create additional imbalances that must themselves be offset. Whether this process can ultimately conceal the redistribution of latent probability mass from the aggregate data depends on whether the resulting chain closes in a way that exactly balances the probability-weighted contributions of the different states.

In the present example, the chain can be closed after just one further redistribution: move $\frac{1-f(a)}{f(a)}\varepsilon$ units of probability mass from type t_2 to type t_1 . Types t_1 and t_2 both choose w in state b , so this second perturbation does not affect the contribution of state b . In state a , however, type t_1 chooses x , whereas type t_2 chooses y . It therefore increases the aggregate frequency of x by $(1 - f(a))\varepsilon$ and decreases that of y by the same amount, exactly offsetting the effects of the first perturbation. The aggregate choice frequencies are therefore unchanged, even though the underlying distribution over behavioral types has changed.

The crucial observation is therefore that identification fails only if redistributions of probability mass can be propagated through the latent types so that every observable discrepancy created at one stage is exactly offset at a later stage. This requires a special form of lack of behavioral heterogeneity.

Below we establish two classes of characterizations that make this requirement precise. The first provides a compact matching-based condition, in the spirit of those developed in the previous section. The second develops the underlying behavioral intuition further, by identifying the cyclic choice structures responsible for these invisible redistributions.

4.2 Identifiability of the type distribution via matching

The next definition formalizes the idea of a state where the type is revealed through choice.

Definition 4. A state $a^* \in A$ *separates* types if, for any two types t_1 and t_2 , we have $c_{t_1}(a^*) = c_{t_2}(a^*)$ only if $t_1 = t_2$.

Thus, separation means that state a^* induces a matching between types and chosen alternatives. States that do not separate types are said to *pool* them. The following result shows that within the type-state framework the existence of a single separating state is sufficient for generic identification.

Theorem 3. *Suppose there exists a state $a^* \in A$ that separates types. Then the parameters in π are generically identifiable.*

Example 5. To illustrate the result, consider three types, three alternatives, and two states, $A = \{a, b\}$, and

$$M^a = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & 0 \end{bmatrix}, \quad M^b = \begin{bmatrix} 1 & 0 & 1 \\ 0 & 0 & 0 \\ 0 & 1 & 0 \end{bmatrix}, \quad M = \begin{bmatrix} 1 & 0 & f(b) \\ 0 & 0 & f(a) \\ 0 & 1 & 0 \end{bmatrix},$$

where $M = f(a)M^a + f(b)M^b$. State a separates types but b does not. Because $f(b) = 1 - f(a)$, the complete parameter description of M is given by $f(a) \in [0, 1]$. By Theorem 3, the parameters in π are generically identifiable. Indeed, π is not identifiable only for $f(a) = 0$.

Intuitively, the presence of one separating state reflects in choice probability changes (in that unobservable state) any perturbation in the type distribution. Except for very specific

(nongeneric) state distributions, the aggregation across states will not be able to exactly offset those changes, allowing any perturbation in the type distributions to be generically reflected in changes in the observed data.

We make further progress by specifying the definition of a matching for the current environment.

Definition 5. A *matching* is an injective function $m : T \rightarrow X$ such that for all $t \in T$, $m(t) = c_t(a)$ for some $a \in A$.

When A consists of a single state a , the existence of a matching is equivalent to the fact that a separates types. The type-state framework also allows us to record the states in which each type is matched to its corresponding alternative, motivating the following definition.

Definition 6. A *state-matching* is a pair (m, γ) where m is a matching and $\gamma : T \rightarrow A$ is such that $m(t) = c_t(\gamma(t))$ for all $t \in T$.

In words, a state-matching is defined by (i) an exclusive assignment of an alternative to each type; and (ii) a specification of the state where any given type selects the assigned alternative. For a single matching m there may be various state-matchings (m, γ) , (m, γ') , etc.

For any state-matching (m, γ) , let $\Gamma(m, \gamma) = [\nu_1 \dots \nu_{|A|}]$ be the state-usage vector, indicating the numbers of times each state a_i is used to form the matching m , i.e.,

$$\nu_i = \sum_{t \in T} \mathbf{1}[\gamma(t) = a_i].$$

To provide a full characterization of generic identifiability, we introduce some standard terminology. For any given permutation σ on the index set $\{1, 2, \dots, r\}$, an inversion is a pair (i, j) such that $i < j$ and $\sigma(i) > \sigma(j)$. A permutation is *even* (resp., *odd*) if it consists of an even (resp. odd) number of inversions. Any matching $m : T \rightarrow X$ is associated with a permutation σ_m defined by $m(t_i) = x_{\sigma_m(i)}$. The *parity* of a matching m is defined as *even* (resp., *odd*) if σ_m is an even (resp., odd) permutation.⁵ Two matchings have the same parity if and only if one can be transformed into the other by an even number of pairwise swaps of assignments between types. More precisely, a swap is a permutation that exchanges two elements and leaves all the others fixed, that is for distinct indices $i \neq j$ the swap between i and j is the permutation σ such that $\sigma(i) = j$, $\sigma(j) = i$ and $\sigma(k) = k$ for all $k \neq i, j$.

For any subset of alternatives $R \subseteq X$ with $|R| = r$, let M_R denote the $r \times r$ submatrix of M obtained by retaining the rows associated with R (and all r columns). Denoting $\text{Im}(m) =$

⁵While the parity of a matching depends on the chosen orderings of types and alternatives, only relative parity between matchings matters for our results, so the choice of ordering is immaterial.

$\{m(t) \mid t \in T\}$ the image of m , let

$$\mathcal{S}_R := \{(m, \gamma) : (m, \gamma) \text{ is a state-matching and } \text{Im}(m) = R\}.$$

Theorem 3 has identified a simple situation in which every perturbation of the type distribution leaves generically a detectable trace in the observed data thanks to a separating state. When no state separates types, every state admits type perturbations that remain locally hidden because the perturbed types are pooled. Whether identification nevertheless survives after aggregation depends on whether these locally hidden changes can be “circulated” across the remaining states without eventually producing a detectable change elsewhere. The following result gives a necessary and sufficient condition for every non-zero perturbation to leave such a detectable trace, generically.

Theorem 4. *The parameters in π are generically identifiable if and only if there exists a subset $R \subseteq X$ with $|R| = r$ such that the following two conditions hold: (i) $\mathcal{S}_R$ is nonempty. (ii) For any partition of $\mathcal{S}_R$ into disjoint pairs $((m, \gamma), (m', \gamma'))$ with $\Gamma(m, \gamma) = \Gamma(m', \gamma')$, there exists at least one pair for which m' can be obtained from m via an even number of swaps. Moreover, if no such R exists, then π is structurally non-identifiable.*

Condition (i) in Theorem 4 asserts the existence of a state-matching (m, γ) for which the image of m is the subset of alternatives R . It guarantees that every type can in principle be associated with a distinct observable alternative by selecting suitable states. Condition (ii) then ensures that perturbations which remain locally hidden within individual states cannot be combined into a perturbation that remains hidden after aggregation. It rules out exactly those pairs of state-matchings that permit a complete redistribution of probability mass across types while leaving the aggregate choice frequencies unchanged. In the proof, condition (i) guarantees the existence of a non-zero Leibniz monomial in the determinant expansion, while condition (ii) prevents all such monomials from cancelling pairwise. Theorem 3 follows from Theorem 4, since type separation by some $a^* \in A$ implies that there is a state-matching such that $\gamma(t) = a^*$ for all $t \in T$, with no other state-matchings having the same Γ . Theorem 5 below offers a behavioral interpretation.

Example 6. Consider

$$\alpha \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 0 \end{bmatrix} + \beta \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 0 \\ 1 & 0 & 1 \end{bmatrix} = \begin{bmatrix} \alpha & \beta & 0 \\ 0 & \alpha & \alpha \\ \beta & 0 & \beta \end{bmatrix},$$

where $X = \{x, y, z\}$, $A = \{a, b\}$, $\alpha = f(a)$, and $\beta = f(b)$. The columns of the right-hand side matrix are generically independent because the determinant is not zero, for example, for $\alpha = \beta = \frac{1}{2}$. There is a state-matching $(m, \gamma) = (xyz, aab)$, with $\Gamma(m, \gamma) = [2 \ 1]$. There is another state-matching $(m', \gamma') = (zxy, bba)$, with $\Gamma(m', \gamma') = [1 \ 2]$. There is no other state-matchings. Since $\Gamma(m, \gamma) \neq \Gamma(m', \gamma')$, condition (ii) is vacuously satisfied. Theorem 4 applies while Theorem 3 is silent because neither state separates types.

Example 7. The following is an instance where condition (ii) of Theorem 4 fails:

$$\alpha \begin{bmatrix} 1 & 0 & 1 & 0 \\ 0 & 1 & 0 & 1 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix} + \beta \begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \end{bmatrix} = \begin{bmatrix} \alpha & 0 & \alpha & 0 \\ 0 & \alpha & 0 & \alpha \\ \beta & \beta & 0 & 0 \\ 0 & 0 & \beta & \beta \end{bmatrix},$$

where $X = \{x, y, z, w\}$. The columns of the right-hand side matrix are linearly dependent for all α and β , because the sum of the first and fourth columns is equal to the sum of the second and third column. This case is ruled out in Theorem 4, since there is one state-matching $(m, \gamma) = (xzw y, abba)$, with $\Gamma(m, \gamma) = [2 \ 2]$, and another state-matching $(m', \gamma') = (zyxw, baab)$ with $\Gamma(m', \gamma') = [2 \ 2]$. Matching m' is obtained from m by three swaps. There are no other state-matchings. Hence, condition (ii) of Theorem 4 fails.

Example 8. Consider two types $T = \{t_1, t_2\}$, two states $A = \{a, b\}$, and two alternatives $X = \{x, y\}$. Suppose, towards a violation of the identifiability condition, that both states pool types—for instance, $c_{t_1}(a) = x = c_{t_2}(a)$ and $c_{t_1}(b) = y = c_{t_2}(b)$. Then there exist two symmetric state-matchings: $(m, \gamma) = (xy, ab)$ and $(m', \gamma') = (yx, ba)$. In both state-matchings, each state is used exactly once and m' is obtained from m by a single swap, violating condition (ii) of Theorem 4. This shows that at least one state must separate types for identifiability to hold. This condition is also sufficient, since Theorem 3 applies.

These matching statements rest on underlying determinant conditions: if state s separates types, then M^s has full rank (it is the identity or a permutation matrix) and $\det M^s \neq 0$, and if state s pools types, then M^s has rank one (two identical columns) and $\det M^s = 0$. Global identification on the interior of the simplex means that⁶ $\det M(f(a)) \neq 0$ for all $f(a) \in (0, 1)$. Since $\det M(f(a))$ is a linear function of $f(a)$ in the $2 \times 2 \times 2$ case, this is equivalent to: $\det M^a$ and $\det M^b$ have the same sign, and are not both zero. If exactly one state, say a ,

⁶Denoting $M(f(a)) = f(a)M^a + (1 - f(a))M^b$.

separates types, then $\det M(f(a)) = f(a) \det M^a$, which is nonzero for all $f(a) \in (0, 1)$. Hence the model is generically (and indeed globally) identifiable.

If both states separate types, then two cases arise. When each type chooses identically across states, $M^a = M^b = M$ and $M(f(a))$ is constant in $f(a)$ and invertible, so generic (and indeed global) identification holds. By contrast, when each type chooses differently across states, say M^a is the identity and M^b is a permutation matrix, $\det M(f(a)) = 2f(a) - 1$ which vanishes at a unique interior value of $f(a)$, ensuring generic (though not global) identification.

The reasoning in Example 8 shows as a corollary of Theorem 4 that in the $2 \times 2 \times 2$ case the parameters are generically identifiable if and only if there exists at least one state that separates types.

4.3 Identification of the type distribution via a behavioral condition

While the characterizations of generic identification in previous sections are general, they do not lay bare how choices in different states interact to enable or obstruct identification. This section zooms in on this interaction, clarifying the precise behavioral implications of such identification failures. For brevity and to obtain clear conditions, the focus here is on the $r \times r$ case and on generic identification only. Thus, we restrict attention to the situation of $p = M\pi$ with $M = \sum_{a \in A} f(a)M^a$ being a square matrix.

From Theorem 3, we know that if generic identification fails, then there is no type separating state. In other words, each state pools at least some types. Given that the M^a 's have one-hot columns, this is equivalent to each such matrix having at least two columns which are identical. But, while this condition is necessary for generic identification to fail, it is in general not sufficient.⁷ The following situation from Example 6 demonstrated this fact already

$$\alpha \underbrace{\begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 0 \end{bmatrix}}_{M^a} + \beta \underbrace{\begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 0 \\ 1 & 0 & 1 \end{bmatrix}}_{M^b} = \underbrace{\begin{bmatrix} \alpha & \beta & 0 \\ 0 & \alpha & \alpha \\ \beta & 0 & \beta \end{bmatrix}}_M$$

Generic identification holds here, even though state a pools the two types corresponding to the last two columns and state b pools the two types corresponding to the first and the last

⁷In the case of two types only, each state pooling some types boils down to each state pooling all (two) types. So, in this case the condition is both necessary and sufficient.

column. In other words, the distinctions generated by the choices in one state are not undone by the distinctions generated in the other state.

To undo such distinctions, the state-pooled choices need to be connected across states in a way that allows different underlying distributions of types to generate exactly the same aggregate outcome frequencies. This connection of choices across states takes the form of a specific cyclic sequence of an even number of types, such that the choices of any two consecutive types are equal, each type uses two different states to establish the two equalities with the neighboring types in the sequence, and the overall use of states is balanced. Formally:

Definition 7. An A -cycle is a sequence of pairwise different types $t_1, \dots, t_{2k+2}$, where $k \geq 0$, together with states $l_i, r_i \in A$, such that

$$c_i(l_i) = c_{i+1}(r_{i+1}), \quad i = 1, \dots, 2k+2$$

where indices are taken modulo $2k+2$,⁸ and

- (i) $l_i \neq r_i$ for all i .
- (ii) for every $a \in A$, the number of indices i with $l_i = a$ equals the number of indices i with $r_i = a$.

An A -cycle formalizes the idea that distinctions generated by choices in one state are exactly undone by distinctions generated in other states. To illustrate this concept and why it hinders identification, consider first the simplest such cycle. For $k = 0$, an A -cycle boils down to the same two types being pooled by all states, as in the following example.

Example 9. Let $A = \{a, b\}$ and consider

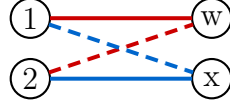
$$\underbrace{f(a) \begin{bmatrix} 1 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{bmatrix}}_{M^a} + \underbrace{f(b) \begin{bmatrix} 0 & 0 & 0 \\ 1 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}}_{M^b} = \underbrace{\begin{bmatrix} f(a) & f(a) & 0 \\ f(b) & f(b) & f(a) \\ 0 & 0 & f(b) \end{bmatrix}}_M$$

where type $t \in \{1, 2, 3\}$ corresponds to the t -th Column of M , and alternative w to Row 1, x to Row 2, and y to Row 3. The choices $c_1(a) = w = c_2(a)$ and $c_2(b) = x = c_1(b)$ describe an A -cycle with $k = 0$, $l_1 = a = r_2$, $l_2 = b = r_1$. We can depict this A -cycle as an Euler cycle⁹ in the bipartite graph whose left vertices correspond to the types in $\{1, 2\}$ and whose right

⁸I.e., $c_{t_{2k+2}}(l_{2k+2}) = c_{t_1}(r_1)$

⁹An Euler cycle is a path in a graph that visits every edge exactly once and returns to the starting vertex.

vertices to the alternatives in $\{w, x\}$. The l_i -choices in the equations above are the solid edges, and the r_j -choices the dashed edges, with state-a-edges in red and state-b-edges in blue. The balancedness of the A-cycle results in a color-balance between dashed and solid edges.



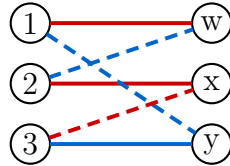
For the simple A-cycle here, it is obvious that the observables do not reveal the composition of the type population uniquely. Types 1 and 2 are pooled both in state a and in state b . Hence, any transfer of probability mass across these two types will never be revealed by aggregate choices.

Next, reconsider the situation from the start of this section where generic identification occurs.

Example 10. Let $A = \{a, b\}$, $T = \{1, 2, 3\}$, $X = \{w, x, y\}$, and consider:

$$\alpha \underbrace{\begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 0 \end{bmatrix}}_{M^a} + \beta \underbrace{\begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 0 \\ 1 & 0 & 1 \end{bmatrix}}_{M^b} = \underbrace{\begin{bmatrix} \alpha & \beta & 0 \\ 0 & \alpha & \alpha \\ \beta & 0 & \beta \end{bmatrix}}_M$$

where type $t \in T$ corresponds to the t -th Column of M , and alternative w to Row 1, x to Row 2, and y to Row 3. Note that the choice cycle $c_1(a) = w = c_2(b)$, $c_2(a) = x = c_3(a)$, $c_3(b) = y = c_1(b)$ is not balanced in the way of Condition (ii) in the definition of an A-cycle requires. We have $l_i = a$, for $i = 1, 2$, but $r_j = a$ only for $j = 3$. Indeed, depicting this choice cycle in the bipartite graph below with left vertices corresponding to types and right vertices to alternatives, the solid edges (representing l_i -choices) feature more red state-a-edges as the dashed edges (representing r_j -choices) do. There is an Euler cycle, but no color-balance between the dashed and solid edges.



Suppose we move a small amount of probability mass from type 2 to type 3. Since state a pools these two types, this redistribution leaves the contribution of state a unchanged. However, it changes the contribution of state b . To conceal this change, one might try a second redistribution, for instance from type 3 to type 1. This restores the contribution of state b , but now alters that of state a . Thus every compensating redistribution creates a new discrepancy elsewhere, and because the cycle is not balanced, this process cannot be completed without eventually changing the aggregate outcome frequencies.

Finally, consider an example with A-cycles involving more than two states.

Example 11. Let $A = \{a, b, g\}$ and consider the non-invertible matrix M given by

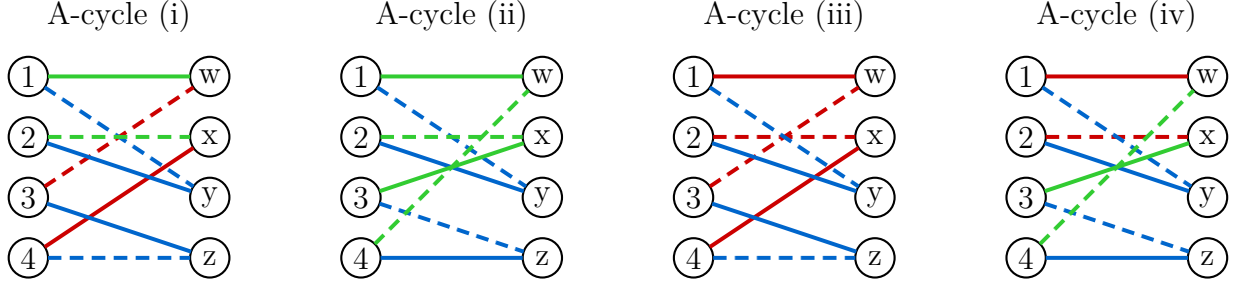
$$\alpha \underbrace{\begin{bmatrix} 1 & 0 & 1 & 0 \\ 0 & 1 & 0 & 1 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}}_{M^a} + \beta \underbrace{\begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \end{bmatrix}}_{M^b} + \gamma \underbrace{\begin{bmatrix} 1 & 0 & 0 & 1 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}}_{M^g} = \underbrace{\begin{bmatrix} \alpha + \gamma & 0 & \alpha & \gamma \\ 0 & \alpha + \gamma & \gamma & \alpha \\ \beta & \beta & 0 & 0 \\ 0 & 0 & \beta & \beta \end{bmatrix}}_M$$

where type $t \in \{1, 2, 3, 4\}$ corresponds to the t -th Column of M , and alternative w to Row 1, x to Row 2, y to Row 3, and z to Row 4. Then, there are the following four A-cycles:

- (i) $c_1(g) = w = c_3(a)$, $c_3(b) = z = c_4(b)$, $c_4(a) = x = c_2(g)$, $c_2(b) = y = c_1(b)$
- (ii) $c_1(g) = w = c_4(g)$, $c_4(b) = z = c_3(b)$, $c_3(g) = x = c_2(g)$, $c_2(b) = y = c_1(b)$
- (iii) $c_1(a) = w = c_3(a)$, $c_3(b) = z = c_4(b)$, $c_4(a) = x = c_2(a)$, $c_2(b) = y = c_1(b)$
- (iv) $c_1(a) = w = c_4(g)$, $c_4(b) = z = c_3(b)$, $c_3(g) = x = c_2(a)$, $c_2(b) = y = c_1(b)$

Again, we can depict these A-cycles as Euler cycles in a bipartite graph whose left vertices correspond to the types and whose right vertices to the alternatives. The l_i -choices in the equations above are the solid edges, and the r_j -choices are the dashed edges. State-a-edges are red, state-b-edges blue, and state-g-edges green. The color-balance between solid and dashed edges in each case indicates the balancedness of the corresponding A-cycle.

Identification fails here because the following redistribution of probability mass leaves the aggregate outcome frequencies unchanged: shifting $(\alpha - \gamma)$ from type 2 to type 1, while also shifting $(\alpha + \gamma)$ from type 3 to type 4. Indeed, state b pools type 1 and type 2, while also pooling type 3 and type 4. So, all of these shifts of frequency mass are hidden in state b alone



and all four A-cycles involve this exact pooling by state b . An observer cannot detect the move of frequency mass with the help of states a and g because it cancels out at the level of the alternatives. To see this, consider the alternative w as an illustration. The mass of $(\alpha - \gamma)$ at type 1 reaches w once as $\alpha(\alpha - \gamma)$ because state a occurs with probability α and once as $\gamma(\alpha - \gamma)$ because state g occurs with probability γ . The mass of $(\alpha + \gamma)$ at type 4 reaches w via state g 's probability of γ as $\gamma(\alpha + \gamma)$ and the mass of $-(\alpha + \gamma)$ at type 3 reaches w via state a 's probability α as $-\alpha(\alpha + \gamma)$. Finally, observe that these probability-weighted masses arriving at alternative w do cancel each other out:

$$\alpha(\alpha - \gamma) + \gamma(\alpha - \gamma) + \gamma(\alpha + \gamma) - \alpha(\alpha + \gamma) = 0.$$

Since this cancellation also occurs at all other alternatives, the observer cannot detect the move of frequency mass. Therefore, aggregate choices here do not reveal the composition of agent types uniquely.

The previous examples suggest a common mechanism. A perturbation of the type distribution can remain hidden only if the compensating redistributions required to offset its observable effects eventually close into a balanced loop. A-cycles provide a behavioral description of these balanced loops. The following theorem shows that they constitute precisely the behavioral configurations preventing generic identification. Since this section assumes the $r \times r$ square case of $|X| = |T| = r$, the set $\mathcal{S}_R$ is simply the set of all state-matchings.

Theorem 5. *The parameters in π are generically identifiable if and only if the set $\mathcal{S}_R$ is non-empty and there is no partition of $\mathcal{S}_R$ into disjoint pairs $((m, \gamma), (m', \gamma'))$ such that, for each such pair, the choices of all types $t \in T$ with $\gamma(t) \neq \gamma'(t)$ form an A-cycle.*

Additional examples that help understanding the identification logic of this section are provided in Appendix B.

5 Multiple choice occasions

We extend identification to type-level choice probabilities. Consider an analyst who observes choice data from a finite number of *occasions* $j \in \{1, \dots, J\}$. The case $J = 1$ is that of a one-shot observation, which was the focus of the previous sections. The type distribution π is assumed constant across occasions. For each occasion j , choice is from a feasible set X_j . Choices are independent across occasions conditional on the type. The notation and definitions introduced earlier extend in the natural way to all j -indexed objects of this section.

The data set is now described as a tensor¹⁰ S of order J with dimensions $n_1 \times \dots \times n_J$. A typical entry $S_{k_1, \dots, k_J}$ is the share of the population choosing alternative x_{k_j} from X_j for $j \in \{1, \dots, J\}$. Hence, S contains the joint distribution of choices for the J occasions. We can express S using M_j and π as

$$S = \sum_{l=1}^r \pi(t_l) \bigotimes_{j=1}^J M_j \mathbf{1}_l, \quad (3)$$

where $\bigotimes$ is the outer product operator and $\mathbf{1}_l$ is the unit vector for component h .¹¹

It turns out that $J = 3$ is sufficient for generic identifiability of the model in Equation (3), provided the possibility relations are sufficiently rich. Under this condition, one can apply algebraic results on tensor decompositions to identify the model. The uniqueness properties of these decompositions have been extensively studied, beginning with Kruskal's [23] fundamental theorem, later refined by Allman, Matias, and Rhodes [1] and others. Because the spirit of our analysis is that of identifying the parameters from minimal data, we restrict the agent's choices to the three occasions required. Additional occasions neither aid nor impede identification in this setting, whereas two occasions are insufficient.

5.1 General framework

In the general case, the parameter space of the model is a full dimensional subset of $[0, 1]^m$ where $m = (r - 1) + \sum_{j=1}^3 (|\wedge_j| - r)$. The model's parameterization map takes values in $[0, 1]^{n_1 \times n_2 \times n_3}$. The next result gives a sufficient condition for generic identification.

Theorem 6. *Suppose there exist natural numbers $v_j \leq r$ for $j \in \{1, 2, 3\}$ such that: (i) $v_1 + v_2 + v_3 \geq 2r + 2$; (ii) for any $j \in \{1, 2, 3\}$ and any v_j types in T , there is a matching of*

¹⁰A tensor is a multidimensional array that extends the concept of a matrix to an arbitrary number of indices, known as the tensor's order. Its dimensions specify how many values each index can take, generalizing the numbers of rows and columns in a matrix.

¹¹Thus $M_j \mathbf{1}_l$ is the l th column of M_j . The outer product operation creates a rank-1 tensor of order J . A tensor is said to be of rank 1 if it is an outer product of vectors. S is a convex combination of r rank-1 tensors.

these types with alternatives in X_j . Then the parameters in π , M_1 , M_2 , and M_3 are generically identifiable, up to label swapping.

Note that for the application of Theorem 6 we need at least $2r + 2$ alternatives across the three occasions that have positive choice shares. Additionally, because identification is unique up to a simultaneous permutation of the entries of π and columns of M_1 , M_2 , and M_3 , we do not know a priori which of the inferred values of π corresponds to a particular type. However, the correct labeling is often revealed by the pattern of 0's in M_j . The matrix in Example 1 illustrates this.

Example 12. Consider the following matrix form, with rows corresponding to four alternatives x , y , z , and w , and columns corresponding to types t_1 , t_2 , and t_3 ,

$$M_j = \begin{matrix} & t_1 & t_2 & t_3 \\ \begin{matrix} x \\ y \\ z \\ w \end{matrix} & \begin{pmatrix} \bullet & \bullet & \bullet \\ 0 & 0 & \bullet \\ \bullet & \bullet & \bullet \\ 0 & 0 & \bullet \end{pmatrix} \end{matrix}.$$

Note that $t \wedge_j x$ and $t \wedge_j z$ for all types t , but $t \wedge_j y$ and $t \wedge_j w$ only for $t = t_3$. Condition (ii) in Theorem 6 holds with $v_j = 3$ (for example, take the matching represented by the red dots). If alternative z is removed, condition (ii) holds only for $v_j = 1$, since no matching exists that covers both types t_1 and t_2 .

5.2 Type-state framework

We now consider multiple occasions of choice in the type-state framework. The states are allowed to vary with the occasions, with A_j denoting the set of states at occasion j .

Theorem 7. Consider the model described by Equations (2) and (3) with parameters π and f_j . If, for any occasion $j \in \{1, 2, 3\}$, the conditions of Theorem 4 hold, then the parameters in π and $\{M_j : j \in \{1, 2, 3\}\}$ are generically identified, up to label swapping.

Theorem 7 shows that the one-occasion type-state identification result lifts directly to the tensor setting, and provides a sufficient condition for the generic identification of π and M_j , $j \in \{1, 2, 3\}$. It is sometimes also possible to recover f_j from Equation (2). In particular, this occurs when there exists a type that selects a different alternative in every state $a \in A_j$. More generally, Equation (2) defines a system of linear equations in the unknowns $\{f(a)\}_{a \in A_j}$, which

may or may not admit a unique solution. Theorems 4 and 7 can be generalized to the case where the state distribution f_j depends on the type. In that case, instead of Equation (2), each column of M_j would be a different convex combination of the corresponding columns of $\{M^a : a \in A_j\}$.

6 Applications

6.1 Choice over vectors of characteristics

Let consumer goods be described in terms of measurable characteristics (after Lancaster [24]), or *features*. Denote $F = \{1, \dots, N\}$ the set of possible features.¹² A good is a vector $x = (x_1, \dots, x_N)$ of amounts x_i of features $i \in F$ (while we use the term “good” for simplicity, we allow features to be undesirable). Consumers choose from feasible set $Y \subseteq \mathbb{R}^N$, a nonempty closed bounded convex set of $\mathbb{R}^N$ with the boundary denoted by ∂Y . A *type* is an $I \subseteq F$ corresponding to the nonempty set of features that the consumer considers relevant for the description of the alternatives. A type I chooses an $x^* \in Y$ that is “optimal” with respect to the features I . We now define the notion of optimality.

Definition 8. Given an evaluation function $e : F \rightarrow \{-1, 0, 1\}$, a point $x \in Y$ is *e-admissible* if $y \notin Y$ whenever $e(i) y_i \geq e(i) x_i$ for all $i \in F$ and $e(i^*) y_{i^*} > e(i^*) x_{i^*}$ for some $i^* \in F$.

The consumer’s evaluation function e indicates if the consumer values a feature positively, negatively, or ignores it. For a consumer of type I , $e(i) \neq 0$ if and only if $i \in I$. The behavioral assumption is that for each consumer there exists an (unobserved) evaluation function e such that the choice $x^* \in Y$ is e -admissible. This requirement is similar to Pareto optimality; however, it is applied only to the dimensions in I , and the directions of monotonicity specified by e are not a priori known to the analyst.

Definition 9. A type $I \subseteq F$ is *possible* at $x^* \in Y$ if there exists an evaluation function e such that $e(i) \neq 0 \iff i \in I$ and x^* is e -admissible.

Let $X \subset Y$ be a finite set of goods with positive choice shares. For each type $I \subseteq F$ and each good $x \in X$, we write $I \wedge x$ if type I is possible at $x \in Y$ according to Definition 9.

As an example, let $N = 2$, so the types are $\{1\}$, $\{2\}$, and $\{1, 2\}$. The feasible set Y is shown in Figure 1. The possible types at each goods $X = \{x, y, z, w\}$ coincide with those described by M_j in Example 12. Theorem 1 guarantees that the distribution of the types $\{1\}$,

¹²[22] studies a version of this model using only *deterministic* choices.

$\{2\}$, and $\{1,2\}$ can be generically recovered from the observed choices. With three occasions, Theorem 6 further shows that both the type distribution and the matrices M_1 , M_2 , and M_3 are generically identifiable.

The type-state model can be applied to a population of consumer types that satisfy a linearity restriction. A type $I \subseteq F$ is *linear* at $x^* \in Y$ if there exists an $a \in \mathbb{R}^N$ such that $a_i \neq 0 \iff i \in I$ and $x^* \in \operatorname{argmax}_{x \in Y} \langle a, x \rangle$. It is easy to see that if a type is linear at x^* , then it is also possible at x^* . Assume the feature coefficients are drawn from a finite set of values. Thus, f_j is interpreted as a probability measure defined on a finite set $A_j \subset \mathbb{R}^N$ of feature values vectors. After a is realized according to f_j , the consumer solves

$$\max_{x \in Y_j} \langle \kappa_I(a), x \rangle \quad (4)$$

where $\kappa_I : \mathbb{R}^N \rightarrow \mathbb{R}^N$ is defined as $\kappa_I(a)_i = a_i$ if $i \in I$ and $\kappa_I(a)_i = 0$ otherwise. We exclude the non-generic case when the consumer problem (4) does not have a unique solution.¹³ Equation (4) determines a choice function in the sense previously specified, and the results of Section 4 can be applied to ensure that the distribution of consumer types can be identified.

The model admits some additional identification analysis, as M_j contains information about the evaluation functions. Consider the example of Figure 1. Recall that each consumer's choice is e -admissible with respect to the consumer's evaluation function e . The first column of M_j describes the tastes of the consumers for whom only feature x_1 matters: $M_j(1,1)$ is the share of $e(1) = 1$ (x_1 is desirable) and $M_j(3,1)$ is the share of $e(1) = -1$ (x_1 is undesirable). Likewise, the second column describes the tastes of type $\{2\}$ agents, and the third column describes the tastes of type $\{1,2\}$ agents. For example, consumers choosing alternative y (alternative w) must evaluate both features positively (resp., negatively), while this is less certain for those choosing x or z . If the population consists of only linear types, then the third column of M_j shows the shares of type $\{1,2\}$ consumers whose vectors of feature values belong to the normal cones at points x , y , z , and w .

6.2 Incomplete preferences

Suppose the set $X = \{x_1, \dots, x_n\}$ of alternatives is linearly ordered from best to worst according to a strict preference relation $\succ$, i.e., $x_1 \succ x_2 \succ \dots \succ x_n$. Agents' preferences are *incomplete* in the sense that alternatives that are sufficiently close in this ordering are treated as incomparable

¹³Alternatively, a tie-breaker assumption can be used.

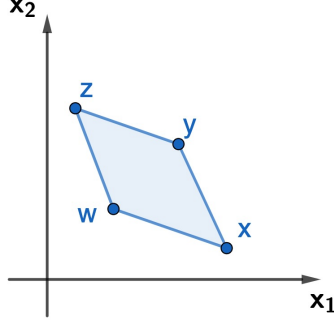


Figure 1: Feasible set Y with M_j as in Example 12

or indifferent.¹⁴ Formally, an agent of type t_l can distinguish between two alternatives only if they are at least l positions apart in the ordering induced by $\succ$. Let $\succ_l$ denote the preference relation of type t_l . It follows that each $\succ_l$ is a coarsening of $\succ$: $\succ_1 = \succ$, $\succ_{l+1} \subset \succ_l$ for all l , and $\succ_n$ is the empty relation. The set of types is $T = \{t_1, \dots, t_n\}$.

Each type never chooses a dominated alternative and randomizes among undominated ones. Thus, t_1 represents the fully rational type, whereas t_n corresponds to the fully random type. Observe that $t_l \wedge x_k$ if and only if $l \geq k$. Consequently, the parameters populate an upper-triangular type-conditional choice matrix M . This structure guarantees the existence of a unique matching, and the parameters in π are globally identifiable from the observed choice probabilities $p = M\pi$ by Theorem 1. Moreover, if data from multiple choice occasions are available, the analyst can also identify the entries of M (the randomization parameters). In particular, in the three-occasion setting discussed in Section 5, we have $v_1 = v_2 = v_3 = n$ and $r = n$. Hence, both conditions of Theorem 6 are satisfied, implying that the parameters in π , M_1 , M_2 , and M_3 are generically identifiable.

Example 13. Let $X = \{x, y, z, w\}$ and $x \succ y \succ z \succ w$. The induced preference relations for the different types are given by

$$\succ_1 = \succ, \{x \succ_2 z, x \succ_2 w, y \succ_2 w\}, \{x \succ_3 w\}, \succ_4 = \emptyset.$$

The corresponding sets of undominated alternatives are therefore $\{x\}$ for type t_1 , $\{x, y\}$ for type t_2 , the top three alternatives for type t_3 , and all four alternatives for type t_4 . As a result, the type-conditional choice matrix M takes the form

¹⁴Incomplete preferences have been studied extensively in decision theory and social choice, both as a primitive departure from classical rationality (Aumann [4], Dubra *et al.* [14]) and as an outcome of bounded rationality or limited discrimination (Masatlioglu *et al.* [30], Manzini and Mariotti [27]).

$$M = \begin{matrix} & t_1 & t_2 & t_3 & t_4 \\ \begin{matrix} x \\ y \\ z \\ w \end{matrix} & \begin{pmatrix} \bullet & \bullet & \bullet & \bullet \\ 0 & \bullet & \bullet & \bullet \\ 0 & 0 & \bullet & \bullet \\ 0 & 0 & 0 & \bullet \end{pmatrix} \end{matrix}.$$

This matrix admits a unique matching, implying that the type distribution is globally identifiable. Moreover, if data from multiple choice occasions are available, the choice matrix M itself is also generically identifiable.

The type-state specification of the framework is useful when the exact nature of the agents' randomization behavior is known. Suppose that agents resolve incomparabilities or indifferences via the salience of the alternatives. Let each state $a \in A$ specify a linear salience ordering $\triangleright_a$ over the set X of alternatives. An agent of type t_l in state a chooses the maximizer of $\triangleright_a$ over the agent's set $\{x_1, \dots, x_l\}$ of undominated alternatives. Each state thus generates a type-conditional choice matrix M^a with one-hot columns. The overall type-conditional choice matrix is obtained by aggregation, $M = \sum_{a \in A} M^a$. According to Theorem 3, the type distribution π is generically identifiable from $p = M\pi$ if the type is fully identified in at least one state $a^* \in A$.

One case when the type is fully identified given $a^* \in A$ arises when the salience ordering $\triangleright_{a^*}$ is the *reverse* of the preference ordering $\succ$. Now each type t_l selects x_l , i.e., $c_{t_l}(a^*) = x_l$ for all l , and the corresponding choice matrix M^{a^*} is the identity matrix. More generally, only diagonal state-matchings are possible in this setting, i.e., matchings satisfying $m(t_l) = x_l$ for all l . Moreover, all such state-matchings have the same parity. Hence, Theorem 4 implies that the parameters in π are generically identifiable if and only if at least one diagonal state-matching exists. This condition is equivalent to requiring that, for each type, there exists a salience ordering under which the worst undominated alternative for that type is the most salient. Formally, for every $l \in \{2, \dots, n\}$, there exists a state a^l such that $x_l \triangleright_{a^l} x_i$ for all $i \in \{1, \dots, l-1\}$.

As an illustration, consider Example 13 with two states $A = \{a, b\}$, where the salience orderings are $w \triangleright_a y \triangleright_a z \triangleright_a x$ and $z \triangleright_b w \triangleright_b y \triangleright_b x$. Let $f(a)$ and $f(b)$ denote the probabilities of states a and b , respectively. The equation $f(a)M^a + f(b)M^b = M$ then takes the form

$$f(a) \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} + f(b) \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 0 \end{bmatrix} = \begin{bmatrix} f(a) + f(b) & 0 & 0 & 0 \\ 0 & f(a) + f(b) & f(a) & 0 \\ 0 & 0 & f(b) & f(b) \\ 0 & 0 & 0 & f(a) \end{bmatrix}.$$

For a generic choice of f , the resulting matrix M is invertible, implying that the distribution π is globally identifiable.

7 Concluding remarks

We have provided necessary and sufficient conditions under which aggregate choice shares from a single set of alternatives guarantee identifiability (generic or global) of the type distribution, as well as sufficient conditions for generic identification of both the type distribution and the type-level choice probabilities. A central message of our analysis is that the informational content of aggregate choice data depends less on its volume per se than on the structure of the qualitative behavioral patterns it embeds. We have obtained strong identification results without imposing parametric structure on decision rules, and without requiring rich menu variation or many repeated observations. At the same time, our results clarify the limits of what can be learned from aggregate data: when behavioral heterogeneity is insufficient, even generic identification breaks down, regardless of the amount of data available. While the broad link between diversity in behavior and identifiability is intuitive, a contribution of our analysis is to make the required notion of “qualitative behavioral heterogeneity” precise. We do so by characterizing it combinatorially, through the notion of “matchings” between types and alternatives, and behaviorally, by specifying the cyclic choice sequences that obstruct identification.

By working with abstract behavioral types rather than more specific primitives, the framework developed here is designed to apply to a wide range of behavioral theories—such as limited consideration, salience-based choice, status quo bias, or incomplete preferences—as well as to more ad hoc behavioral hypotheses tailored to particular contexts. Moreover, the qualitative nature of the characterizing properties frees the researcher from the need for precise knowledge of type-level choice probabilities, even at a merely parametric level.

An additional question concerns whether the number of relevant types can itself be identified.

One may begin with an upper bound on the number of latent types and successively aggregate types that are observationally indistinguishable. This procedure continues until no further aggregation is possible. In this sense, identification pertains not necessarily to the full set of types, but to the minimal partition that is relevant for observed behavior.

We have focused on identification under extremely undemanding informational assumptions. In many empirical settings, the analyst will have access to additional information, such as observations from multiple menus, covariates, or partial knowledge of the cognitive processes underlying choice. An important direction for future research is to extend the framework developed here to incorporate such additional sources of information and to assess the extent to which they further strengthen identification.

References

- [1] Allman, E. S., C. Matias, and J. A. Rhodes (2009) “Identifiability of Parameters in Latent Structure Models With Many Observed Variables,” *Annals of Statistics* 37: 3099–3132.
- [2] Apesteguia, J. and M. Ballester (2025) “A Measure of Behavioral Heterogeneity,” *American Economic Journal: Microeconomics*, forthcoming.
- [3] Apesteguia, J., M. Ballester, and J. Lu (2017) “Single-Crossing Random Utility Models,” *Econometrica* 85: 661–674.
- [4] Aumann, R. J. (1962) “Utility Theory Without the Completeness Axiom,” *Econometrica*: 445-462.
- [5] Azrieli, Y. and J. Rehbeck (2022) “Marginal Stochastic Choice,” *arXiv*: 2208.08492 [econ.TH].
- [6] Barseghyan, L., M. Coughlin, F. Molinari, and J. C. Teitelbaum (2021) Heterogeneous Choice Sets and Preferences, *Econometrica* 89(5): 2015-2048.
- [7] Cappé, O., E. Moulines, and T. Rydén (2005) “Inference in Hidden Markov Models,” Springer, New York.
- [8] Caradonna, P. P. and C. Turansick (2026) “Identification in Stochastic Choice,” *arXiv*: 2408.06547 [econ.TH].
- [9] Chambers, C. P. and C. Turansick (2025) “The Limits of Identification in Discrete Choice,” *Games and Economic Behavior* 150: 537–551.

- [10] Cox, D. A., J. Little, and D. O’Shea (1997) “Ideals, Varieties, and Algorithms,” 2nd ed, Springer, New York.
- [11] Dardanoni, V., P. Manzini, M. Mariotti, and C. J. Tyson (2020) “Inferring Cognitive Heterogeneity from Aggregate Choices,” *Econometrica* 88(3): 1269-1296.
- [12] Dardanoni, V., P. Manzini, M. Mariotti, H. Petri, and C. J. Tyson (2024) “Mixture Choice Data: Revealing Preferences and Cognition,” *Journal of Political Economy* 131(3): 687-715.
- [13] Doignon, J. and K. Saito (2023) “Adjacencies on Random Ordering Polytopes and Flow Polytopes,” *Journal of Mathematical Psychology*, 114: 102768.
- [14] Dubra, J., F. Maccheroni, and E. A. Ok. (2004) “Expected Utility Theory Without the Completeness Axiom,” *Journal of Economic Theory* 115.1: 118-133.
- [15] Ephraim, Y. and N. Merhav (2002) “Hidden Markov Processes,” *IEEE Transactions on Information Theory* 48, 1518-1569.
- [16] Filiz-Ozbay, E. and Y. Masatlioglu (2023) “Progressive Random Choice,” *Journal of Political Economy* 131: 716–750.
- [17] Fishburn, P. C. (1998) “Stochastic Utility,” in *Handbook of Utility Theory*, ed. by S. Barbera, P. Hammond, and C. Seidl, Kluwer Dordrecht; 273–318.
- [18] Forney, Jr. G. D. (1975) “Minimal Bases of Rational Vector Spaces, with Applications to Multivariable Linear Systems,” *SIAM Journal on Control* 13.3: 493-520.
- [19] Hall, P. (1935) “On Representatives of Subsets,” *Journal of the London Mathematical Society* 1: 26–30.
- [20] Koopmans, T. C. and O. Reiersøl (1950) “The identification of Structural Characteristics,” *The Annals of Mathematical Statistics* 21: 165–181.
- [21] Koopmans, T. C. (1950) “Statistical Inference in Dynamic Economic Models,” Wiley, New York.
- [22] Kops, C. J., P. Manzini, M. Mariotti, and I. Pasichnichenko (2025) “Revealing Features from Pareto Optimal Choice,” Working paper.

- [23] Kruskal, J. B. (1977) “Three-Way Arrays: Rank and Uniqueness of Trilinear Decompositions, With Application to Arithmetic Complexity and Statistics,” *Linear Algebra and its Applications* 18: 95–138.
- [24] Lancaster, K. J. (1966) “A New Approach to Consumer Theory,” *Journal of Political Economy* 74(2): 132-157.
- [25] Lindsay, B. G. (1995) “Mixture Models: Theory Geometry and Applications,” *NSF-CBMS Regional Conference Series in Probability and Statistics* 5, IMS, Hayward, CA.
- [26] Lu, J. (2019) “Bayesian Identification: A Theory for State-Dependent Utilities,” *American Economic Review* 109(9): 3192-3228.
- [27] Manzini, P. and M. Mariotti (2012) “Categorize Then Choose: Boundedly Rational Choice and Welfare,” *Journal of the European Economic Association* 10.5: 1141-1165.
- [28] Manzini, P. and M. Mariotti (2018) “Dual Random Utility Maximisation,” *Journal of Economic Theory* 177: 162–182.
- [29] Manzini, P., M. Mariotti, and H. Petri (2019) “Corrigendum to “Dual Random Utility Maximisation” [J. Econ. Theory 177 (2018) 162–182],” *Journal of Economic Theory* 184: 104944.
- [30] Masatlioglu, Y., D. Nakajima, and E. Y. Ozbay (2012) “Revealed Attention,” *American Economic Review* 102.5: 2183-2205.
- [31] McLachlan, G. J. and D. Peel (2000) “Finite Mixture Models,” Wiley, New York.
- [32] Rhodes, J. A. (2010) “A Concise Proof of Kruskal’s Theorem on Tensor Decomposition,” *Linear Algebra and its Applications* 432: 1818-1824.
- [33] Turansick, C. (2022) “Identification in the Random Utility Model,” *Journal of Economic Theory* 203: 105489.
- [34] Yildiz, K. (2023) “Foundations of Self-Progressive Choice Models,” *arXiv*: 2212.13449 [econ.TH].

Appendices

A Proofs

Proof of Theorem 1 Here generic identifiability of π is equivalent to the columns of M being generically independent. If the columns of M are generically independent, then there exists an $r \times r$ submatrix M' of M such that $\det(M')$ is not identically 0. Note that $\det(M')$ is a multivariate polynomial in the parameters $\{M'(k, l) : t_l \wedge x_k\}$. From the Leibniz determinant formula

$$\det(M') = \sum_{\sigma \in S_r} \text{sgn}(\sigma) \prod_{l=1}^r M'(\sigma(l), l),$$

(where S_r denotes the set of permutations on $\{1, \dots, r\}$) it follows that there exists a permutation σ^* of the rows of M' such that the monomial

$$\prod_{l=1}^r M'(\sigma^*(l), l)$$

is not identically 0. This is possible only if all of the variables in the monomial are not identically 0. Hence, $t_l \wedge x_{\sigma^*(l)}$ for all $l \in \{1, \dots, r\}$ and there exists a matching m defined as $m(t_l) = x_{\sigma^*(l)}$. This proves one direction of the assertion on generic identifiability in the statement.

Conversely, if there exists a matching m , then we can define a permutation σ^* for all $l \in \{1, \dots, r\}$ as the unique permutation satisfying equation $m(t_l) = x_{\sigma^*(l)}$. Then, we have $t_l \wedge x_{\sigma^*(l)}$ for all $l \in \{1, \dots, r\}$. Form an $r \times r$ -submatrix M' of M using the rows corresponding to the alternatives in the image of m . Choosing the point in the parameter space defined by $M'(k, l) = 1$ if $k = \sigma^*(l)$ and $M'(k, l) = 0$ otherwise, we have $\det(M') = 1$. Hence, $\det(M')$ is not identically 0, which means that the r columns of M are generically independent. This concludes the proof of the other direction.

When the matching m is unique, for all $\sigma \in S_r$ with $\sigma \neq \sigma^*$ we have $M'(\sigma(l), l) = 0$ for some l (otherwise we would have found another matching different from m), whereas $\prod_{l=1}^r M'(\sigma^*(l), l) \neq 0$ since $M'(\sigma^*(l), l) > 0$ for all l . Equivalently, in the Leibniz expansion there is exactly one term that is nonzero, and we conclude that $\det(M') \neq 0$ for all parameter values. This shows that the columns of M are independent for all parameter values, and π is globally identified. Finally, if there exists no matching, then for any $r \times r$ -submatrix M' of M , all terms in the Leibniz expansion are zero, and therefore $\det(M') = 0$ for any such submatrix. Then the rank of M is less than r , and its columns are linearly dependent for all values of the

parameters, showing that π is structurally non-identified. $\square$

Proof of Theorem 2 *Sufficiency.* Suppose π is not globally identifiable. Then there exists an admissible matrix M and two distinct type distributions π and π' such that $M\pi = M\pi'$, or $M(\pi - \pi') = 0$. Since $\pi \neq \pi'$, there exists a vector $c \neq \mathbf{0}$ such that $c = \pi - \pi'$. Now distinguish the types according to the corresponding sign in vector c . Specifically, let

$$\begin{aligned} T^+ &:= \{t \in T \mid c_t > 0\}, \\ T^- &:= \{t \in T \mid c_t < 0\}. \end{aligned}$$

Since $\sum \pi_t = 1 = \sum \pi'_t$, we have $\sum c_t = 0$, implying that T^+ and T^- are nonempty.

For each type t , let $S(t) := \{x \mid t \wedge x\}$.

Take any alternative $x \in \bigcup_{t \in T^+} S(t)$. This means that for some $t \in T^+$, we have $M(x, t) > 0$. Hence (recalling $c_u > 0$ for all $u \in T^+$),

$$\sum_{u \in T^+} M(x, u) c_u > 0.$$

Since

$$0 = (Mc)_x = \sum_{u \in T^+} M(x, u) c_u + \sum_{v \in T^-} M(x, v) c_v,$$

it must be that

$$\sum_{v \in T^-} M(x, v) c_v < 0.$$

Hence, there exists a $v \in T^-$ such that $M(x, v) > 0$, or, $v \wedge x$. Therefore, we have shown that

$$x \in \bigcup_{t \in T^+} S(t) \implies x \in \bigcup_{t \in T^-} S(t).$$

Using a symmetric argument we obtain the reverse implication, and thus

$$\bigcup_{t \in T^+} S(t) = \bigcup_{t \in T^-} S(t),$$

proving that collective behavioral separation is violated.

Necessity. Suppose that collective behavioral separation fails, and specifically that $U := \bigcup_{t \in T^+} S(t) = \bigcup_{t \in T^-} S(t)$.

Step 1. Let $S_1, \dots, S_k \subseteq X$ with $\bigcup_{i=1}^k S_i = U$. Then for every full-support distribution z

on U (i.e. $z_x > 0$ for all $x \in U$ and $z \in \Delta(U)$)¹⁵ there exist coefficients $\lambda_1, \dots, \lambda_k > 0$ with $\sum_i \lambda_i = 1$, and full support distributions $v_1, \dots, v_k$, with each $v_i \in \Delta(S_i)$, such that

$$z = \sum_{i=1}^k \lambda_i v_i.$$

That is, a full-support distribution on U can be expressed as a mixture of full-support distributions over subsets of U that together cover the whole of U .

Proof: For $x \in U$, let $n(x) := |\{i : x \in S_i\}|$. Note that since $S_1, \dots, S_k$ cover U we have $n(x) \geq 1$. We begin by constructing auxiliary weights $W(i, x)$ which we will subsequently normalise to obtain the distributions v_i . Let $W(i, x) := z_x/n(x)$ if $x \in S_i$ and $W(i, x) := 0$ otherwise. Since $z_x > 0$ and $n(x) \geq 1$, $W(i, x) > 0$ whenever $x \in S_i$.

For each $x \in U$, the mass z_x is split uniformly across the $n(x)$ sets that contain x . Therefore, aggregating the $W(i, x)$ yields the desired mass on x :

$$\sum_{i=1}^k W(i, x) = z_x.$$

Next, define the coefficients

$$\lambda_i := \sum_{x \in S_i} W(i, x) = \sum_{x \in S_i} z_x/n(x).$$

Since S_i is nonempty and every term is positive, we have $\lambda_i > 0$. Moreover, the coefficients add up to one:

$$\sum_{i=1}^k \lambda_i = \sum_{i=1}^k \sum_{x \in S_i} \frac{z_x}{n(x)} = \sum_{x \in U} z_x \sum_{i: x \in S_i} \frac{1}{n(x)} = \sum_{x \in U} z_x \cdot \frac{n(x)}{n(x)} = \sum_{x \in U} z_x = 1.$$

Finally, define

$$v_i(x) := W(i, x) / \lambda_i$$

for $x \in S_i$. Since $v_i(x) > 0$ for all $x \in S_i$ and $\sum_{x \in S_i} v_i(x) = \lambda_i/\lambda_i = 1$, v_i is a full-support distribution on S_i . Finally, for all $x \in U$,

$$\sum_{i=1}^k \lambda_i v_i(x) = \sum_{i=1}^k W(i, x) = z_x$$

¹⁵ $\Delta(U)$ is the set of all probability distributions over U .

as required.

Step 2. Fix z as any full-support distribution over U . Apply Step 1 to both families of subsets $\{S(t)\}_{t \in T^+}$ and $\{S(t)\}_{t \in T^-}$ to obtain

$$\sum_{t \in T^+} \lambda_t v_t = z = \sum_{t \in T^-} \mu_t v'_t,$$

where $(\lambda_t)_{t \in T^+}$ and $(\mu_t)_{t \in T^-}$ are all strictly positive weights summing to 1, and the v_t and v'_t are full-support distributions on $S(t)$ for $t \in T^+$ and $t \in T^-$, respectively.

Step 3. We construct a matrix M^* that is an admissible parameter point for the type conditional choice matrix as follows: set, for each $t \in T$:

$$M^*(\cdot, t) := \begin{cases} v_t, & t \in T^+, \\ v'_t, & t \in T^-, \\ \text{any full-support distribution on } S(t), & t \notin T^+ \cup T^-. \end{cases}$$

Each column respects the possibility relation $\wedge$ in the sense that it is strictly positive exactly on $S(t)$, and its entries add up to one. Hence, M^* is an admissible parameter point as claimed.

Step 4. We construct two distinct type distributions π and π' that generate the same data with M^* . Define $c \in \mathbb{R}^r$ by $c_t := \lambda_t$ for $t \in T^+$, $c_t := -\mu_t$ for $t \in T^-$, and $c_t := 0$ otherwise. Since T^+, T^- are disjoint and nonempty by assumption, $c \neq \mathbf{0}$, and

$$\sum_{t \in T} c_t = \sum_{t \in T^+} \lambda_t - \sum_{t \in T^-} \mu_t = 1 - 1 = 0.$$

Moreover,

$$M^*c = \sum_{t \in T^+} \lambda_t v_t - \sum_{t \in T^-} \mu_t v'_t = z - z = \mathbf{0}.$$

Fix any π_0 in the interior of $\Delta(T)$ and an $\eta > 0$ small enough that $\pi := \pi_0 + \eta c$ and $\pi' := \pi_0 - \eta c$ both lie in $\Delta(T)$ (possible since $\sum_t c_t = 0$ and π_0 is interior). Then $\pi \neq \pi'$ since $c \neq \mathbf{0}$, yet

$$M^*\pi - M^*\pi' = M^*(2\eta c) = \mathbf{0},$$

so that $M^*\pi = M^*\pi'$.

We conclude that the type distribution π is not globally identifiable at the admissible parameter point M^* . $\square$

Proof of Theorem 3 We will show that the columns of M are generically independent. Suppose that $a^* \in A$ separates types. Then there exists a matching $m : T \rightarrow X$ such that, for all $t \in T$, $m(t) = c_t(a^*)$. This implies that $n \geq r$. Form an $r \times r$ -submatrix M' of M using the rows $\{k : x_k = m(t), t \in T\}$. The determinant of M' is a polynomial in the parameters $\{f(s)\}_{s \in A}$. Consider the point in the parameter space where $f(s) = 1$ if $s = a^*$ and $f(s) = 0$ otherwise. At this point, we have that M' is a permutation matrix, which implies $\det(M') \neq 0$. Hence, the zero set of $\det(M')$ is a proper subvariety and not the whole parameter space. We have just shown that there exists an $r \times r$ -minor of M such that its zero set is a proper subvariety. This implies that the columns of M are generically independent. $\square$

Proof of Theorem 4 Fix $R \subseteq X$ with $|R| = r$. For $a \in A$, it holds that

$$M_R = \sum_{a \in A} f(a) M_R^a,$$

where M_R^a is the restriction of M^a to the rows in R , i.e. $M_R^a = (M^a)_R$. By the Leibniz formula, and identifying each permutation in S_r with the corresponding matching,

$$\det(M_R) = \sum_m \delta(m) \prod_{t \in T} M_R(m(t), t).$$

where the summation is over all matchings $m : T \rightarrow R$ and $\delta(m) \in \{-1, 1\}$ is determined by the parity of the permutation associated with m . For any matching $m : T \rightarrow R$ and any $t \in T$,

$$M_R(m(t), t) = \sum_{a \in A : c_t(a) = m(t)} f(a) = \sum_{\gamma : (m, \gamma) \in \mathcal{S}_R} f(\gamma(t)).$$

Expanding the product of sums and observing that each choice of admissible states $(\gamma(t))_{t \in T}$ uniquely defines a state-matching $(m, \gamma) \in \mathcal{S}_R$, we obtain

$$\det(M_R) = \sum_{(m, \gamma) \in \mathcal{S}_R} \delta(m) \prod_{t \in T} f(\gamma(t)), \quad (5)$$

i.e. the determinant is a polynomial in $\{f(a)\}_{a \in A}$ whose monomials are indexed precisely by state-matchings with image R .

Sufficiency. Assume there exists $R \subseteq X$, $|R| = r$, satisfying (i)–(ii). By (i), $\mathcal{S}_R \neq \emptyset$, hence (5) contains at least one monomial with nonzero coefficients. Condition (ii) rules out complete cancellation of all monomials: indeed, cancellation of (5) within every Γ -class would allow one to partition $\mathcal{S}_R$ into disjoint pairs $((m, \gamma), (m', \gamma'))$ with $\Gamma(m, \gamma) = \Gamma(m', \gamma')$ and $\delta(m') = -\delta(m)$

for every pair, contradicting (ii). Therefore, $\det(M_R)$ is not the zero polynomial. Hence M has generically full column rank r , and π is generically identifiable from $p = M\pi$.

Necessity. Suppose π is generically identifiable. Then M has generically full column rank r , so there exists at least one $r \times r$ minor M_R whose determinant is not identically zero as a polynomial in $\{f(a)\}_{a \in A}$. Fix such an R . If $\mathcal{S}_R = \emptyset$, then every product in the Leibniz formula is identically zero, so $\det(M_R) \equiv 0$, a contradiction. Thus (i) holds.

If (ii) fails for this R , then there exists a partition of $\mathcal{S}_R$ into disjoint pairs $((m, \gamma), (m', \gamma'))$ such that $\Gamma(m, \gamma) = \Gamma(m', \gamma')$ and m' differs from m by an odd number of swaps, i.e. $\delta(m') = -\delta(m)$, for every pair. For each pair, $\Gamma(m, \gamma) = \Gamma(m', \gamma')$ implies

$$\prod_{t \in T} f(\gamma(t)) = \prod_{t \in T} f(\gamma'(t)),$$

while $\delta(m') = -\delta(m)$ implies the corresponding contributions in (5) cancel. Hence all terms cancel and $\det(M_R) \equiv 0$, again contradicting the choice of R . Therefore (ii) must hold for this R .

Structural non-identifiability. Suppose that no $R \subseteq X$ with $|R| = r$ satisfies (i)–(ii). Then, for every R , the polynomial $\det(M_R)$ is identically zero, hence every $r \times r$ minor of M vanishes for every parameter value. Thus, $\text{rank}(M) < r$ for all admissible parameters, and π is structurally non-identifiable. $\square$

Proof of Theorem 5 *Necessity.* Suppose π is generically identifiable. This implies that the matrix M is generically invertible and its determinant $\det(M)$ does not vanish. So, there exists a non-zero monomial in the Leibniz formula for the determinant of M which does not cancel. Since every non-zero Leibniz monomial induces a matching between types and alternatives, it follows that $\mathcal{S}_R \neq \emptyset$.

It remains to show that there is no partition of $\mathcal{S}_R$ into disjoint pairs such that some subset of the choices of each pair generates an A -cycle. Because π is generically identifiable, for any nonempty subset of types $S \subseteq T$, also the restriction of π to S , denoted $\pi|_S$, is generically identifiable. This implies that there exists a $|S| \times |S|$ square submatrix N_S of M with its columns corresponding to the types in S such that N_S is generically invertible. It follows that, for any such S , in the Leibniz formula for the determinant of N_S there is a monomial that does not cancel. Since this is also true for M itself, there exists a monomial which, for all non-empty $S \subseteq T$, does not cancel out in the Leibniz formula for the determinant of the corresponding submatrix N_S .

It follows that there exists no partition of the set of all non-zero Leibniz monomials in the

determinant of M into disjoint pairs such that, for each such pair, there exists a submatrix N_S of M , such that the two monomials in this pair cancel each other out. Since each non-zero Leibniz monomial induces a matching, there being no such partition of the set of all non-zero monomials of M implies that there also is no partition $\mathcal{S}_R$ into disjoint pairs such that some subset of the choices of each pair generates an A -cycle. To see this, suppose otherwise, i.e., towards a proof by contradiction suppose there exists such a partition of $\mathcal{S}_R$. First observe that the balancedness Condition (ii) of an A -cycle requires that there are two matchings between some subset of types, say $S \subseteq T$, and some subset of alternatives with the same overall state-usage for each of the two matchings between these subsets. This corresponds to two monomials in the Leibniz formula for the determinant of the corresponding submatrix N_S which coincide apart from a potential difference in their signs. Next, observe that the definition of an A -cycle involves an even number of different types, $2k + 2$, and describes a cyclic sequence of these types' choices, satisfying Condition (i) in the definition of an A -cycle. So, the two coinciding monomials have opposite signs, canceling each other out. Since this results in pairs of monomials canceling each other out, we have reached the desired contradiction.

Sufficiency. Suppose π is not generically identifiable. By Theorem 4, this implies that either Condition (i) or Condition (ii) of Theorem 4 does not hold. If Condition (i) of Theorem 4 does not hold, then the set $\mathcal{S}_R$ is the empty set and the condition of the present theorem does not hold.

Suppose then that Condition (ii) of Theorem 4 does not hold. That is, there exists a partition of $\mathcal{S}_R$ into disjoint pairs $((m, \gamma), (m', \gamma'))$ with $\Gamma(m, \gamma) = \Gamma(m', \gamma')$ such that m' can be obtained from m via an odd number of swaps for each pair.

For any pair $((m, \gamma), (m', \gamma'))$ of state-matchings define the state-difference set

$$\mathcal{D}^{(\gamma, \gamma')} = \{t \in T \mid \gamma(t) \neq \gamma'(t)\}.$$

When the meaning is clear, we will omit the superscripts from $\mathcal{D}^{(\gamma, \gamma')}$ and write only $\mathcal{D}$.

Also, define the successor map σ , which will be used to construct the desired sequence of types, by

$$\sigma(t) = m'^{-1}(m(t))$$

for all $t \in T$. This means that the successor $\sigma(t)$ of type t is matched by m' to the same alternative to which type t is matched by m . Since the matchings m and m' are bijections from T onto X , σ is a well-defined bijection of T onto itself. We will use its restriction to $\mathcal{D}$.

Lemma (full successor cycle): There exists a partition of $\mathcal{S}_R$ such that $\mathcal{D}^{(\gamma, \gamma')}$ is non-empty for each disjoint pair $((m, \gamma), (m', \gamma'))$, and σ is a full $|\mathcal{D}^{(\gamma, \gamma')}|$ -cycle: that is, for each

$t \in \mathcal{D}^{(\gamma, \gamma')}$, $\sigma^n(t) = t$ only for $n = |\mathcal{D}^{(\gamma, \gamma')}|$ or multiples thereof.¹⁶

Proof: Fix an arbitrary disjoint pair $((m, \gamma), (m', \gamma'))$ from the given partition of $\mathcal{S}_R$ and define the “filtered” matrix $M^{(m, m')}$ by retaining, in each entry, only the state contributions used by the two state-matchings. More precisely, for all i, j with $1 \leq i, j \leq r$, the entry in row i (corresponding to alternative x_i) and column j (corresponding to type t_j) of $M^{(m, m')}$ is defined as¹⁷

$$[M^{(m, m')}]_{ij} = \sum_{a \in \{\gamma(t_j), \gamma'(t_j)\}} \mathbf{1}(c_{t_j}(a) = x_i) f(a)$$

By construction, the only non-zero Leibniz monomials of $M^{(m, m')}$ are those induced by m and m' , and these cancel out. Hence also $M^{(m, m')}$ is singular. Let N be a singular square submatrix N of $M^{(m, m')}$, obtained by restricting to a subset of rows and columns such that the restrictions of the two monomials induced by m and m' remain nonzero Leibniz monomials of N . Choose this N to be irreducible in the following sense: for any proper subset C of columns of N , the submatrix induced by C contains a non-singular square submatrix of size $|C|$. By construction, the restrictions of the two monomials are still the only non-zero Leibniz monomials of N , and they continue to cancel.

Suppose now, towards a contradiction, that σ is not a full $|N|$ -cycle. Since σ is a permutation of a finite set, it decomposes into disjoint cycles. Hence it must decompose into at least two disjoint cycles. The cycles involve disjoint sets of rows and columns. This is immediate for columns. For rows, note that every row is $m(t)$ for a unique type t and $m'(t) = m(\sigma(t))$. Therefore the rows appearing in one successor cycle are exactly the images under m of the types in that cycle. Since the cycles partition the types, the row sets are also disjoint.

After reordering rows and columns according to these cycle components, N becomes block diagonal, with one block for each cycle. Since swapping columns may only affect the sign of the determinant and the determinant of a block diagonal matrix is the product of the determinants of each block, $\det(N) = 0$ implies that at least one block is singular. But this block is itself a proper square submatrix N' of N whose columns form a proper subset of the columns of N . Since any nonsingular square submatrix induced by these columns would have to be contained in the block itself (every row outside that block is zero in those columns), we contradict the irreducibility of N . $\diamond$

Note that the existence of a square singular submatrix in the lemma implies that the two monomials cancel each other on that submatrix. Since they permute all entries on the submatrix

¹⁶Here σ^n denotes the n^{th} iteration of the map σ .

¹⁷In the definition, note that $\{\gamma(t_j), \gamma'(t_j)\}$ is a set, not a multiset. Hence, if $\gamma(t_j) = \gamma'(t_j)$, then $\{\gamma(t_j), \gamma'(t_j)\} = \{\gamma(t_j)\}$, and the state is counted only once.

in a cycle, then cancelling must mean an odd number of swaps, so an even number of types involved.

We now complete the proof of the theorem. Consider the restrictions of the matchings m and m' to N obtained by using only the rows and columns of N . We can complete these restrictions to matchings on the level of M that coincide outside of N . Indeed, define $\tilde{m}'$ to coincide with m' on N and with m otherwise. Analogously, define $\tilde{m}$ to coincide with m on N and with m' otherwise. Clearly, $\Gamma(m, \gamma) = \Gamma(\tilde{m}', \tilde{\gamma}')$ and $\Gamma(m', \gamma') = \Gamma(\tilde{m}, \tilde{\gamma})$, where $\tilde{\gamma}'$ coincides with γ' on N and with γ otherwise and $\tilde{\gamma}$ coincides with γ on N and with γ' otherwise. Hence, (m, γ) can be paired with $(\tilde{m}', \tilde{\gamma}')$, and (m', γ') with $(\tilde{m}, \tilde{\gamma})$.

By performing this splicing for all paired matchings, we have constructed a partition of $\mathcal{S}_R$ into disjoint pairs $((m, \gamma), (m', \gamma'))$ with $\Gamma(m, \gamma) = \Gamma(m', \gamma')$ such that: m' can be obtained from m via an odd number of swaps for each pair; the set $\mathcal{D} = \{t \in T \mid \gamma(t) \neq \gamma'(t)\}$ is non-empty; and σ is a full $|\mathcal{D}|$ -cycle.

It only remains to show that any such pair $((m, \gamma), (m', \gamma'))$ generates a balanced A -cycle.

Fix any such pair. Since $\Gamma(m, \gamma) = \Gamma(m', \gamma')$, the restrictions of the two state-matchings to $\mathcal{D}$ have the same state usage. By the Lemma, $\sigma^n(t) = t$ only for $n = |\mathcal{D}|$ or multiples thereof. Clearly, $n > 1$, for otherwise $\sigma(t) = t$ and, thus, $\gamma(t) = \gamma'(\sigma(t)) = \gamma'(t)$, contradicting that $t \in \mathcal{D}$. Define

$$l_{\sigma^i(t)} = \gamma(\sigma^i(t)), \quad r_{\sigma^{i+1}(t)} = \gamma'(\sigma^{i+1}(t))$$

for $i = 0, 1, \dots, n-1$, and note that these arguments induce the cycle

$$c_t(l_t) = c_{\sigma(t)}(r_{\sigma(t)}), c_{\sigma(t)}(l_{\sigma(t)}) = c_{\sigma^2(t)}(r_{\sigma^2(t)}), \dots, c_{\sigma^{n-1}(t)}(l_{\sigma^{n-1}(t)}) = c_t(r_t)$$

Since $\sigma^n(t) = t$ only for $n = |\mathcal{D}|$, the types $t, \sigma(t), \dots, \sigma^{n-1}(t)$ are pairwise different as an A -cycle requires. Next, for any $i = 0, 1, \dots, n-1$, by definition, $\sigma^i(t) \in \mathcal{D}$ implies that $\gamma(\sigma^i(t)) \neq \gamma'(\sigma^i(t))$ and, thus, $r_{\sigma^i(t)} \neq l_{\sigma^i(t)}$ as in Condition (i) of the definition of an A -cycle. Finally, the balancedness condition in the definition of an A -cycle follows because the two state-matchings have identical state usage on $\mathcal{D}$. Hence, we have shown that the cycle above is indeed an A -cycle, which concludes the proof. $\square$

Proof of Theorem 6 As a first step, we show that the Kruskal column rank¹⁸ of M_j is at least v_j generically, for each $j \in \{1, 2, 3\}$. Take any v_j types in T and denote the set of these types by $\bar{T}$. From (ii), it follows that there is an injective function $m : \bar{T} \rightarrow X_j$ such that $t \wedge_j m(t)$ for all $t \in \bar{T}$. Let $l(t)$ be such that $x_{l(t)} = m(t)$ and $k(t)$ be such that $t_{k(t)} = t$, for all $t \in \bar{T}$. Form

¹⁸The Kruskal column rank of M is the largest number d such that every set of d columns are independent.

a $v_j \times v_j$ -submatrix of M_j using rows $\{l(t) : t \in \bar{T}\}$ and columns $\{k(t) : t \in \bar{T}\}$. Observe that the determinant of this matrix is a sum of monomials, one of which is $\prod_{t \in \bar{T}} M_j(l(t), k(t))$. The set of parameter values contains points where, for all $t \in T$, $M_j(l(t), k(t)) = 1$. At any such point the determinant of the submatrix is equal to 1. Hence, the zero set of this determinant is a proper subvariety and not the whole parameter space.

We have just shown that, for any selection of v_j columns of M_j , we can find a $v_j \times v_j$ -minor of M_j formed of these columns such that its zero set is a proper subvariety. Hence, for any v_j columns of M_j , the columns are generically independent. Denote by Θ_j the union of the zero sets of these minors. Since this is a finite union, Θ_j is a proper subvariety too. For any point in the parameter space that lies outside of Θ_j , all combinations of v_j columns are independent. This implies that the Kruskal column rank of M_j is at least v_j generically.

For the second step, observe that outside of $\Theta_0 = \cup_j \Theta_j$ the Kruskal column rank of M_j is at least v_j for all $j \in \{1, 2, 3\}$. Denote by Θ the union of Θ_0 and the set of parameter values such that some component of π is equal to 0. Clearly, Θ is a proper subvariety. Outside Θ , the sum of the Kruskal column ranks of M_1 , M_2 , and M_3 is greater than or equal to $v_1 + v_2 + v_3$ and all components of π are strictly positive. Relying on the Kruskal's result [23], Corollary 2 of [1] says that, for $J = 3$, the parameters of the model (3) are uniquely identifiable, up to label swapping, if the sum of the Kruskal column ranks of M_1 , M_2 , and M_3 is greater than or equal to $2r + 2$. From this and assumption (i), it follows that the parameters outside Θ are uniquely identifiable, up to label swapping. Hence, the parameters are generically identifiable, up to label swapping. $\square$

Proof of Theorem 7 It follows from Theorem 4 that, for each j , there is a proper subvariety Θ_j of the parameter space such that outside Θ_j the Kruskal column rank of M_j is r . Hence, outside of $\Theta_0 = \cup_j \Theta_j$ the sum of the Kruskal column ranks of the matrices M_1 , M_2 , and M_3 is $3r$. Denote by Θ the union of Θ_0 and the set of parameter values such that some component of π is equal to 0. For any point outside Θ we can apply Corollary 2 of [1]. Since Θ is a proper subvariety, the result follows. $\square$

B Additional examples for Section 4.3

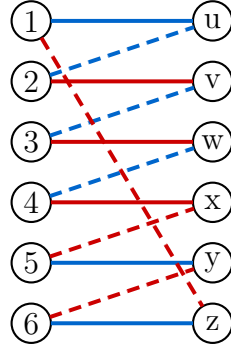
B.1 A 6×6 Example (not identifiable)

The following is an example where we have (aa), (bb), (ab), and (ba) equations in the A-cycle.

Example 14. Let $A = \{a, b\}$, $T = \{1, 2, 3, 4, 5, 6\}$, $|X| = |T|$, and consider:

$$M = f(a) \underbrace{\begin{bmatrix} 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 & 0 & 0 \end{bmatrix}}_{M^a} + f(b) \underbrace{\begin{bmatrix} 1 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix}}_{M^b}$$

where type $t \in T$ corresponds to the t -th Column of M , and alternative u to Row 1, v to Row 2, w to Row 3, x to Row 4, y to Row 5, and z to Row 6. There is the following A-cycle: $c_1(b) = u = c_2(b), c_2(a) = v = c_3(b), c_3(a) = w = c_4(b), c_4(a) = x = c_5(a), c_5(b) = y = c_6(a), c_6(b) = z = c_1(a)$. As before, we can depict this A-cycle as an Euler cycle. The l_i -choices in the equations above are the solid edges, and the r_j -choices are the dashed edges. State-a-edges are red, and state-b-edges are blue. The color-balance between solid and dashed edges indicates the balancedness of the corresponding cycle. The parameters in π are not generically identifiable since there are exactly two state-matchings, $(m, \gamma) = (uvwxyz, baaabb)$ and $(m', \gamma') = (zuvwxy, abbbba)$, for which the choices of all types form an A-cycle.



B.2 An Example similar to Example 11, but identifiable

The next example is very similar to Example 11. The two non-zero rows of M^g move down one row, resulting in an invertible matrix M .

Example 15. Let $A = \{a, b, g\}$ and consider the generically invertible matrix M

$$\alpha \underbrace{\begin{bmatrix} 1 & 0 & 1 & 0 \\ 0 & 1 & 0 & 1 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}}_{M^a} + \beta \underbrace{\begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \end{bmatrix}}_{M^b} + \gamma \underbrace{\begin{bmatrix} 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 1 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}}_{M^g} = \underbrace{\begin{bmatrix} \alpha & 0 & \alpha & 0 \\ \gamma & \alpha & 0 & \alpha + \gamma \\ \beta & \beta + \gamma & \gamma & 0 \\ 0 & 0 & \beta & \beta \end{bmatrix}}_M$$

where type $t \in \{1, 2, 3, 4\}$ corresponds to the t -th Column of M , and alternative w to Row 1, x to Row 2, y to Row 3, and z to Row 4. Then, there are the following three cycles:

- (i) $c_2(a) = x = c_4(a)$, $c_4(b) = z = c_3(b)$, $c_3(g) = y = c_2(g)$
- (ii) $c_1(a) = w = c_3(a)$, $c_3(b) = z = c_4(b)$, $c_4(g) = x = c_1(g)$
- (iii) $c_1(a) = w = c_3(a)$, $c_3(b) = z = c_4(b)$, $c_4(a) = x = c_2(a)$, $c_2(b) = y = c_1(b)$

Again, we can depict these cycles as Euler cycles in a bipartite graph whose left vertices correspond to the types and whose right vertices to the alternatives. The l_i -choices in the equations above are the solid edges, and the r_j -choices are the dashed edges. State-a-edges are red, state-b-edges blue, and state-g-edges green. While the color-balance might suggest all three are A-cycles, this is not the case. Only cycle (iii) features an even number of pairwise different types, making it the only A-cycle here. For the pair of state-matchings $(wyzx, abba)$ and $(yxwz, baab)$, the choices of the four types form this A-cycle. However, there exist other state-matchings, such as $(wxyz, aagb)$, that are not associated with the A-cycle. Therefore, the condition of Theorem 5 is satisfied, and identifiability holds.

