

Valuing Policy Characteristics and New Products using a Simple Linear Program:

The Pure Characteristics Model Revisited

H. Spencer Banzhaf

Center for Environmental and Resource Economic Policy
Working Paper Series: No. 26-001
May 2026

Suggested citation: Banzhaf, Spencer, "Valuing Policy Characteristics and New Products using a Simple Linear Program: The Pure Characteristics Model Revisited," CEnREP Working Paper No. 26-001, May 2026, <https://go.ncsu.edu/cenrep.wp.26.001>.



VALUING POLICY CHARACTERISTICS AND NEW PRODUCTS USING A SIMPLE LINEAR PROGRAM: THE PURE CHARACTERISTICS MODEL REVISITED

H. Spencer Banzhaf*

May 2026

Abstract

The Random Utility Model (RUM) is a way to value new products or changes in public goods. But RUMs have been faulted along two lines. First, for including idiosyncratic errors that imply unreasonably high values for new alternatives and unrealistic substitution patterns. Second, for involving strong functional form restrictions. This paper shows how, starting with a revealed preference framework, one can partially identify nonparametrically the answers to policy questions about discrete alternatives. When the Generalized Axiom of Revealed Preference (GARP) is satisfied, the approach weakly identifies a pure characteristics model. When GARP is violated, it recasts the RUM errors as departures from GARP (critical cost efficiency), to be minimized using a minimum-distance criterion, thus coming as close to the pure characteristics model as the data allows. The paper illustrates the approach by estimating bounds on the values of ecological improvements using survey data.

Key words: Discrete choice, random utility, nonparametric estimation, revealed preference, ecosystem services.

JEL Codes: C14, C25, C61, D12, Q51

*Professor, Department of Agricultural and Resource Economics, North Carolina State University, PERC, and NBER. spencer_banzhaf@ncsu.edu.

I thank seminar participants at Cornell U., Georgetown U., Georgia State U., LSE, NC State U., the Triangle Econometrics Conference, U.C. Santa Barbara, the U. of Chicago, and the U. of Virginia for comments on earlier drafts.

Valuing Policy Characteristics and New Products Using a Simple Linear Program: The Pure Characteristics Model Revisited

1 Introduction

Valuing policy-induced changes in the attributes of public goods is a common goal in public economics and related fields like health and environmental economics. Examples include hospitals and medical plans, local public goods, and ecosystem services. (Ho and Pakes 2014; Keane et al. 2021; Bayer et al. 2016; Banzhaf et al. 2016; English et al. 2018). Similarly, estimating consumers' values for new products is a common goal in industrial organization and marketing. Examples include new automobiles, consumer appliances, and computers (Berry et al. 2004; Bajari and Benkard 2005).

For some 40 years, the standard approach to this problem has been to use a discrete choice random utility model (RUM), such as a multinomial logit. However, these models raise two difficulties. First, they give a large role to additive idiosyncratic errors, which in turn can imply unrealistic substitution patterns. They also imply that expected utility always increases in the number of alternatives, so adding an alternative—even one dominated in the characteristic space—always has value. Consequently, products always have market power, so markups above marginal cost can never be competed away. Accordingly, Bajari and Benkard (2005) and Berry and Pakes (2007) earlier introduced a pure characteristics model with no idiosyncratic errors at all, only unobserved tastes and product characteristics. However, these models proved computationally challenging and can be quite sensitive to outlying observations (Athey and Imbens 2007; Berry and Pakes 2007; Gandhi and Nevo 2021).

A second criticism of discrete choice models is that, in practice, most estimators are highly parametric. Typically, they specify distributions for the additive errors (e.g. logit) and the distribution of tastes for characteristics, as well as a functional form for the utility function. Much of the identification from discrete choice models comes from these assumptions.

This paper shows how to overcome both problems by relating the RUM framework to the Generalized Axiom of Revealed Preference (GARP). Whereas the RUM

framework looks for a utility function with patterns of heterogeneity and errors that best explain the observed choices, the GARP framework tests whether there is any utility function that can perfectly explain the choices without errors. If so, it can weakly identify indifference curves and the answers to policy questions—in the sense that it can point identify these objects as the sample goes to infinity, but for finite samples the estimator can only recover sets of parameters and bound answers to policy questions (Lewbel 2019). For discrete choices, this approach assumes only that utility is strictly monotonic, continuous, and quasi-concave.¹

To see its connection to the RUM approach, two long-standing features of the GARP approach are especially salient. First, economists can test GARP using a linear program (LP), which also provides a way to construct a candidate utility index nonparametrically (Afriat 1967). Second, if the tests reject GARP, economists can quantify the degree of the rejection using indices known as “critical cost efficiency” (Afriat 1972; Varian 1985, 1990). Such measures typically ask how much expenditure would have to be wasted to make the data consistent with GARP, and typically appear in an additively separable form. These two features of the GARP literature are well established for competitive budget sets. However, their implications for RUMs with discrete alternatives has not previously been noted.

With those features in mind, my argument can be summarized in the following steps. First, consider micro-level data with a panel of discrete choices, which allows for non-trivial revealed preference tests. The “panel” structure of the data may comprise repeated choice observations for each individual or a cross section of data but with the individuals considered as part of a panel at a group level (Berry and Haile 2024). Second, I argue that if a set of discrete choices is consistent with GARP in characteristics space, then we can weakly identify the pure characteristics demand model and bound answers to policy questions in the spirit of Manski (2007, 2010). Moreover, the bounds can conveniently be obtained by “searching” over a domain defined by the values which are consistent with GARP, directly within the LP. However, it is likely that GARP will be violated in most data sets big enough to be of practical use, in which case the pure characteristics model is incoherent. Accordingly, the third

¹The revealed preference approach, begun by Samuelson and extended by Houthakker, Afriat, Varian, and others, is now quite large and has been extended to firm behavior, Nash equilibria, and other contexts. See e.g. Chambers and Echenique (2016) and Hands (2014) for overviews and discussion.

step of the logic is to generalize the model by introducing the adjustments associated with critical cost efficiency. Crucially, I show that these adjustments are equivalent to the additively separable, idiosyncratic errors of a RUM. That is, optimization mistakes in a model with no idiosyncratic shocks are observationally equivalent to a model without mistakes, but with idiosyncratic tastes, a “moment’s fancy,” or quality shocks.

Unfortunately, introducing such errors nonparametrically has the potential to make the model radically incomplete. For any candidate utility functions, there are *always* some errors that could rationalize the observed data. In other words, the identified sets are unbounded. However, this problem can be overcome by assuming that the additive idiosyncratic errors of the RUM are, like other econometric errors, objects to be minimized. They are only there to explain why we cannot otherwise perfectly satisfy GARP in characteristics space. As with critical cost efficiency, they are the minimum adjustment to utility (potentially measured in dollars) that would rationalize the data. Completing the link, the model can then be estimated with an LP, with the objective to minimize some norm of the errors. Rather than fitting market shares conditional on an error structure, this way of thinking about the problem inverts the logic and finds the pure characteristics model that comes closest to the data, with the residuals being the unexplained gaps. The resulting error distribution is selected nonparametrically, imposing only additive separability and mean independence. In other words, rather than point identifying a model from this unbounded set using functional form and/or distributional assumptions like logit, I propose selecting it by a minimum distance criterion.

This approach has four advantages over others. First, by construction, it minimizes a norm of the errors, rather than imposing distributional assumptions that put more weight on random utility than necessary. Second, consequently, it weakly identifies the pure characteristics model (when GARP is satisfied) or comes as close to it as the data allows (when GARP is violated), thus overcoming a long-standing estimation challenge for these models (Athey and Imbens 2007; Berry and Pakes 2007; Gandhi and Nevo 2021). Third, answers to economic and policy questions can be directly introduced into the objective of the LP, while GARP is satisfied through inequality constraints whenever possible. Finally, it is nonparametric, imposing only monotonicity, continuity, and quasi-concavity on the function of observables and additivity and

mean independence on the unobservables.

I illustrate this approach with an application to the valuation of ecosystem services in the Southern Appalachian Mountain area of the southeastern United States, previously analyzed using traditional logit models (Banzhaf et al. 2016). The ecosystems in that region have been damaged by years of acid precipitation. A survey of the area’s residents elicited preferences for various investments to improve different dimensions of the ecosystem (including the health of streams, forests, and bird populations). The survey was set up as a series of choice experiments, over which respondents chose their preferred option from a menu. Analyzing these data using the nonparametric approach posed here, and allowing for heterogeneity across households, I find the average household’s WTP for a given set of improvements can be bounded between \$44 and \$210, compared to \$49 for a mixed logit, which imposes stronger assumptions. The average absolute value of the errors in the nonparametric model, when monetized, is about \$0.75, compared to about \$75 in the mixed logit—a difference of two orders of magnitude. Moreover, 97% of observations require no error at all, and 88% of households require no error for any of the alternatives they face. Thus, in this application, the pure characteristics model is much closer to the data than standard RUMs give it credit for.

This paper is related to the existing literature in at least four ways. First, it relates to work that extends GARP to discrete goods. Blow et al. (2008) showed how GARP can be extended from goods space to Gorman-Lancaster characteristics space. They considered the value of new products, but still assumed choice sets were continuous. Polisson and Quah (2013) and Cosaert and Demuyne (2015) extend GARP from the usual competitive budget sets, which are convex, to discrete choice sets. I rely on their results, but also extend them by introducing additive errors and by showing how their results can help bound economic questions in characteristics space.

Second, this paper is related to recent work on identification of discrete choice models. Berry and Haile (2016, 2024) emphasize that identification of demand in such models is built on nonparametric foundations, and several results have pushed in that direction. Fox et al. (2012) consider semi-parametric identification of the distribution of tastes for characteristics, within the mixed logit model. Shi et al. (2018) show how to semi-parametrically identify average tastes for characteristics with minimal

assumptions on the distribution of idiosyncratic errors, except for additive separability and a continuous distribution function. However, they impose a linear form for the utility function and restrict the admissible distribution of observed characteristics. Allen and Rehbeck (2019) nonparametrically identify average demands via a representative consumer. Their approach bears some similarity to mine insofar as it imposes the integrability conditions on this consumer, which are comparable to GARP. However, for welfare analysis, it requires restrictions to guarantee that the representative consumer is normative (Gorman 1953; Jackson and Yariv 2024). In contrast to these papers, my approach makes no assumption on the distribution or continuity of the distribution of idiosyncratic errors (which, in practice, will have a large probability mass at zero), makes no assumption about the functional form for utility, and recovers information on the distribution of tastes. However, in general it only weakly identifies demands. My approach also relates to Ho and Pakes (2014), who bound answers to policy questions based on revealed preference inequalities. But whereas they average out the errors and derive moment conditions, my approach incorporates the errors into an LP.

One limitation of my approach is that, although it can account for an index of unobserved product characteristics, it assumes they are independent of observed characteristics. This assumption is consistent with McFadden’s early work on logit models and with more recent work by Lewbel et al. (2011) and Allen and Rehbeck (2019). It is applicable to cases where price is plausibly determined exogenously by travel costs (e.g. Beckert et al. 2012; English et al. 2018) or experimentally controlled in a survey, which I use in my empirical application. However, it will be problematic in cases where price or other characteristics are endogenous. Future work might consider imposing additional exclusion restrictions, as emphasized by Berry and Haile (2016, 2024) and in the spirit of Blundell et al. (2003).

Third, this paper relates to other work on non-market valuation in environmental and other settings. Blow and Blundell (2018) use revealed preference methods to gauge the plausibility of benefits transfer exercises. Of especial relevance, Kuminoff (2009) partially identifies parameters within a structural model with functional form restrictions. I extend this body of work by introducing RUM errors as critical cost efficiency indices of GARP violations. I also extend it by considering non-parametric bounds on WTP, without functional form assumptions.

Finally, this paper relates to the nonparametric testing of RUMs using an axiom of revealed stochastic preference (McFadden 2006; Kitamura and Stoye 2018). However, in some ways it takes the opposite approach. Rather than test axiomatically whether an entire sample of data could be generated by a random utility function, it tests whether individual responses could be generated from non-random utility functions, and constructs a random utility function through least deviations from such functions.

The paper proceeds along the following outline. Section 2 introduces the basic intuition with a simple linear-in-parameters structural model, weakly identified through GARP rather than point identified with distributional assumptions on the errors. It begins by assuming the observed choice behavior does satisfy GARP, and hence that there is no need for additive errors. It then derives bounds on the policy questions of interest. Next, it introduces unobserved product characteristics into the model, as in Bajari and Benkard (2005) and Berry and Pakes (2007).

Section 3 introduces idiosyncratic errors into the linear model, to rationalize the data when GARP is violated, using the logic of critical cost efficiency. Completing the steps in the theoretical argument, Section 4 relaxes the linearity assumption, making the model fully nonparametric, except for mild regularity conditions. Section 5 offers an application, estimating WTP for improvements to ecosystem service in the southern Adirondack mountains in the United States. It compares the GARP-based approach to more traditional logit models. Section 6 concludes.

2 The Linear Pure Characteristics Model

This section introduces and develops the basic intuition using the canonical linear RUM, before eventually relaxing any functional form restrictions in Section 4.

2.1 Weak Identification

Consider the following utility function:

$$v_{ijt} = \alpha(y_{it} - p_{jt}) + \bar{\beta}'\mathbf{x}_{jt} + u_{ijt}, \quad (1)$$

where v_{ijt} is the utility for individual (or household) i conditional on choosing alternative j in period t , y_{it} is income in period t and p_{jt} is the price of alternative j

in period t (so $y_{it} - p_{jt}$ is numeraire consumption) and $\mathbf{x}_{jt}$ is a vector of M other observed characteristics $(x_1, \dots, x_m, \dots, x_M)$. The j alternatives could be products, residential locations, or simply policy scenarios. The utility parameters include $\alpha > 0$ and $\bar{\beta} \geq \mathbf{0}$. These sign restrictions imply utility is strictly increasing in the numeraire and non-decreasing in all $\mathbf{x}$. (This is a harmless normalization as we can always model bads as absence of goods.) Note additionally that, though this is a very simple linear model, it easily can be extended to any linear-in-parameters model in the usual way. Given the sign restrictions and the linearity assumption, utility is strictly monotonic and concave.² Finally, an outside option ($p_{0t} = 0$, $\mathbf{x}_{0t} = \mathbf{0}$) is normalized to have $v_{i0t} = 0 \forall i, t$.

As written in the general form of Equation (1), the errors u_{ijt} may have at least three interpretations. First, they could be random idiosyncratic errors: $u_{ijt} = \varepsilon_{ijt}$, as for example in the multinomial logit model. Second, they could reflect unobserved product attributes: $u_{ijt} = \mu_j$ (Berry et al. 1995, 2004; Berry and Pakes 2007; Bajari and Benkard 2005). Third, they could reflect heterogeneity in tastes for $\mathbf{x}$, as in the mixed logit model (Berry et al. 1995; McFadden and Train 2000). In that case, $u_{ijt} = \tilde{\beta}_i' \mathbf{x}_{jt}$, where individual i 's taste vector is $\beta_i = (\bar{\beta} + \tilde{\beta}_i)$ with $\bar{\beta}$ interpreted as the population average taste and $\tilde{\beta}_i$ the departure from the average. Or any combination of these. From the perspective of this paper, the first of the three interpretations (idiosyncratic errors) differs from the other two in that there are *always* values for such errors that guarantee we can rationalize the data.

To develop the basic argument in the simplest setting, in this section I initially assume the data are rationalizable without needing either idiosyncratic errors or unobserved product characteristics to explain observed choices, and thus temporarily omit them from the model. I thus assume the errors reflect only heterogeneity in β_i . This is a reasonable starting point, as much of the previous work on revealed preference in characteristics space also has taken this approach, including Manski (2007); Blow et al. (2008); Polisson and Quah (2013); and Cosaert and Demuyne (2015). Later sections relax these restrictions.

²With competitive budget sets and a finite set of observations, monotonicity and concavity are observationally equivalent (Afriat 1967). However, this equivalence no longer holds with discrete budget sets, because an open ball around some choice alternative might not be in the choice set (Cosaert and Demuyne 2015). Accordingly, I impose strict monotonicity (utility is non-decreasing in all goods and increasing in income) and quasi-concavity throughout this paper.

With these initial assumptions, we can write the model as:

$$v_{ijt} = \alpha_i(y_{it} - p_{jt}) + \beta'_i \mathbf{x}_{jt},$$

with $\alpha_i > 0$ and $\beta_i \geq \mathbf{0}$. Furthermore, we can take an affine transformation of this utility function by dividing by α_i and subtracting y_{it} :

$$v_{ijt} = -p_{jt} + \beta'_i \mathbf{x}_{jt}. \quad (2)$$

With this renormalization, the coefficients β_i can be interpreted as individual-specific marginal WTP for the attributes $\mathbf{x}$, as they are now relative to the marginal utility of money. Additionally, we've conveniently eliminated the strict inequality constraint on α_i .³

Note this model nests several special cases, including the simplest case of total homogeneity ($\beta_i = \beta \forall i$) or observed heterogeneity ($\beta_i = \beta'_z \mathbf{z}_i$ for some vector of observables $\mathbf{z}_i$). Thus, although revealed preference approaches always require repeated observations, the model allows for several kinds of data sets, each with its respective interpretation of the panel. If only a cross section of data (or repeated cross sections) are observed, one could take one of these two special cases. Or, one could define a group by some set of observables, and think of i as the group, with individuals within the group being treated like repeated observations from the same preference ordering. The most general case is individual-level panel data, which allows for unobserved heterogeneity (i.e., β_i as random coefficients).⁴ In my application to ecosystem services, these data are available. Accordingly, for concreteness, I will use the language of this kind of setup, though the others are also possible.

Each individual $i = 1 \dots I$ faces $T(i)$ choice occasions indexed by $t(i) = 1 \dots T(i)$, chooses from a choice set $\mathcal{J}(i, t)$ at each choice occasion t , and actually chooses

³Technically, the β_i in Equation (2) is now a new β_i/α_i , but for simplicity I abuse the notation.

⁴Panel data are necessary to identify random coefficients using the nonparametric GARP approach of this paper, whereas the conventional MLE approach to RUMs can identify random coefficients models even in a cross section. Similar issues arise in other contexts for testing GARP and bounding policy questions, motivating Blundell et al. (2003, 2008) to nonparametrically adjust budget sets for income differences, thereby making different choice occasions comparable. While motivated by a similar data problem, their solution is not directly relevant to my application, as I focus on a single differentiated product, for which budget-balancedness need not hold.

alternative $j'(i, t) \in \mathcal{J}(i, t)$ at choice occasion t . The individual chooses alternative $j'(i, t)$ if and only if $v_{i,j'(t)} \geq v_{i,j} \forall j \in \mathcal{J}(i, t)$. Substituting Equation (2) into this condition, re-arranging, and considering all choice occasions in the data, we have:

$$-p_{j'(i,t)t} + \beta'_i \mathbf{x}_{j'(i,t)t} \geq -p_{jt} + \beta'_i \mathbf{x}_{jt}, \forall i, \forall t, \forall j \in \mathcal{J}(i, t). \quad (3)$$

That is, individual i 's utility for the alternative chosen at any choice occasion must be at least as great as the utility of all other alternatives in that occasion's choice set.

We will say the data are “rationalizable by a linear utility function,” or “linearly rationalizable” for short, if there is a linear utility function that explains the choice patterns observed in the data, as well as patterns of monotonicity (j' is always preferred to j if it dominates in every dimension). Following the revealed preference approach, utility maximization is equivalent to satisfying a set of choice axioms, to the effect that if an individual chooses a bundle A on one occasion when B is available, thus revealing that A is at least as good as B , then there is no sequence of choices such that it ever indirectly reveals B to be strictly better than A . More formally, using the notation $A \succeq_R B$ to mean “ A revealed at least as good as B ” and $A \succ_R B$ to mean “ A revealed strictly preferred to B ,” then $A \succeq_R B$ implies there is no sequence $B \succeq_R C \succeq_R \dots \succeq_R D \succeq_R A$, with one of the relations being strict in the sequence. Thus, GARP assures that there is a no-cycling condition on choice patterns. Simplifying the problem, Afriat (1967) showed that the no-cycling condition is equivalent to the existence of a solution to an LP. His approach relies on the equivalence of GARP and, for competitive budget sets, the existence of utilities for bundles and marginal utilities of income that explain choice patterns. Polisson and Quah (2013) and Cosaert and Demuyne (2015) extend Afriat's approach to discrete choices.

The linearity assumption in this section greatly simplifies the problem: The data are linearly rationalizable if and only if $\exists \beta_i \geq 0$ satisfying Expression (3). Equivalently, there must be some solution (possibly a continuum of solutions) to an LP to minimize the degenerate objective $\beta'_i \mathbf{0}$, subject to the inequalities given by Expression (3) and the non-negativity constraints. In other words, Expression (3) represents the Afriat inequalities for this model. Without additional distributional assumptions about the random coefficients, this is all we can say about the β_i . That is, the set $\{\beta_i | \beta_i \geq 0\}$ consistent with utility maximization are all those satisfying

Expression (3). However, intuitively, the sets shrink as we observe more data. This intuition introduces an alternative argument for identification of the discrete choice model.

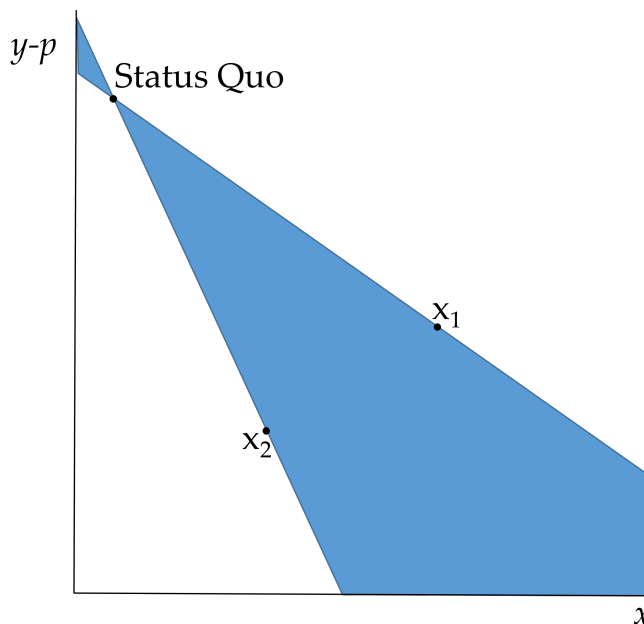
This logic can be summarized more formally by the following theorem:

Theorem 1 (weak identification). If the data are generated from utility-maximizing individuals with linear utility, as in Equation (2), then the data are linearly rationalizable for each individual i , and the distribution of β_i is partially identified by $\{\beta_i | \beta_i \geq \mathbf{0} \text{ and } \beta'_i(\mathbf{x}_{j'(i,t)t} - \mathbf{x}_{jt}) \geq p_{j'(i,t)t} - p_{jt}\}$. Moreover, if $\mathbf{x}$ is distributed over continuous support, then as the set of observed alternatives in the data increases and becomes dense, each set $\{\beta_i\}$ converges to its true value. Thus, the individual random coefficients are point identified in the limit.

Proof: The first part of the theorem is just an application of Afriat's Theorem (Afriat 1967) and follows immediately from Expression (3). To prove the second part, on convergence, let $\{\mathbf{x}_{j'(i,t)t}, \mathcal{J}(i, t)\} \rightarrow \{\beta_i\}$ be the correspondence mapping the data to the set of parameters that satisfy GARP (which by the premise of the theorem is non-empty). Next, consider some true $\tilde{\beta}_i$, some $\mathbf{x}_j \in \mathcal{J}(i, t)$, the set $\sim \mathbf{x}_j = \{\mathbf{x} | \tilde{\beta}'_i \mathbf{x} = \tilde{\beta}'_i \mathbf{x}_j\}$ and the set $\{\hat{\mathbf{x}}_j\} = \{\mathbf{x} | \beta' \mathbf{x} = \beta' \mathbf{x}_j \forall \beta \in \{\beta_i\}\}$. The first of these sets is the true indifference set for $\mathbf{x}_j$; the second comprises the candidates for the indifference set given the data. Now, consider a sequence l which adds additional data points to the choice set $\mathcal{J}(i, t)$. Given the continuous support for $\mathbf{x}$, as $l \rightarrow \infty$, at some step l there will be $\mathbf{x} \in \{\hat{\mathbf{x}}_j\}_{l-1}, \mathcal{J}(i, t)_l$ which, by completeness, is either in the no-worse-than set or no-better-than set for $\mathbf{x}_j$, thus removing regions from $\{\beta_i\}$. As the minimum distance in $\mathbb{R}^k$ between the hyperplane defined by $\sim \mathbf{x}_j$ and an element in $\mathcal{J}(i, t)_l$ shrinks to zero, the set $\{\beta_i\} \rightarrow \tilde{\beta}_i$. In the limit, as $\mathcal{J}(i, t)_l$ becomes dense, there are $M - 1$ linearly independent elements from $\sim \mathbf{x}_j$ and $\{\beta_i\} = \tilde{\beta}_i$.

The basic intuition of the proof follows the logic first introduced by Mas-Colell (1978) for competitive budget sets, adapted to the case of linear utility and discrete choices in characteristics space. The proof uses large- J asymptotics, but a similar argument could be constructed using large- T (i.e., repeated choice occasions), so long as each choice occasion involves an independent draw from the space of possible choice

Figure 1: Partial Identification of the Linear Pure Characteristics Model



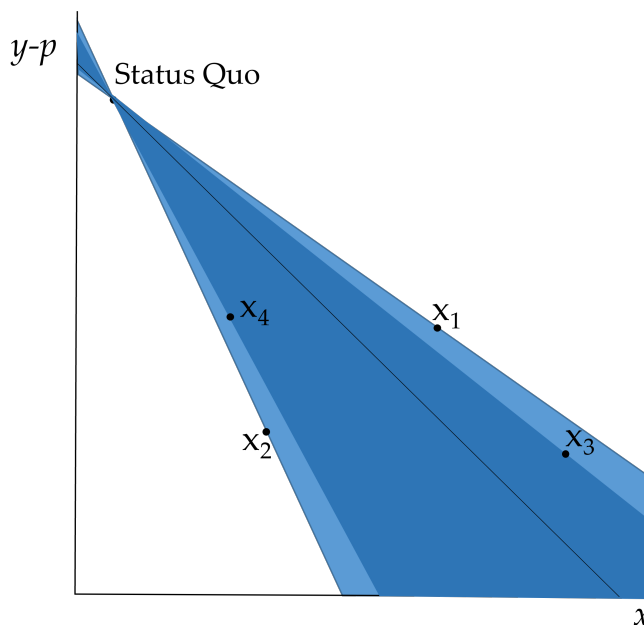
In this example, we observe $\{x_1, p_1\} \succ_R \text{SQ} \succ_R \{x_2, p_2\}$. The shaded area depicts the set $\{\beta_i\}$ of linear taste parameters consistent with these data.

sets. Additionally, if the “panel” data treats individuals as repeated draws from a group with identical preferences, the proof could use large- I .

As noted above, if the data are rationalizable, all we can accomplish with finite data and without further structural assumptions is to partially identify each β_i . Figure 1 illustrates this point with an example with $M = 1$ and three alternatives, a status quo option (SQ), $\{x_1, p_1\}$, and $\{x_2, p_2\}$. In this example, we observe $\{x_1, p_1\} \succ_R \text{SQ} \succ_R \{x_2, p_2\}$. The shaded area depicts the set $\{\beta_i\}$ consistent with these data. The linear indifference curve going through the SQ option could lie in any part of this shaded area. The absolute values of the slopes of such indifference curves reflect the possible WTP for attribute x .

Nevertheless, it is reassuring that, as the data set becomes large for each i , the identified set converges. This second part of the proposition is illustrated in Figure 2. In this example, x_3 and x_4 now enter $\mathcal{J}$, and we observe $\{x_3, p_3\} \succ_R \text{SQ} \succ_R \{x_4, p_4\}$. Now the light shaded areas can be eliminated, and the set $\{\beta_i\}$ shrinks to the darker shaded area, closing in on the true indifference curve shown. However, the rate of convergence is not known *a priori*.

Figure 2: Asymptotic Point Identification of the Linear Pure Characteristics Model



In this example, we observe $\{x_1, p_1\} \succ_R \{x_3, p_3\} \succ_R \text{SQ} \succ_R \{x_4, p_4\} \succ_R \{x_2, p_2\}$. The dark shaded area depicts the set $\{\beta_i\}$ of linear taste parameters consistent with these data, which has shrunk from the light shaded area of Fig. 1.

To get some insight into this question, I construct simulations, described in Appendix A1. In summary, I consider an individual with a linear utility function and randomly sample choice alternatives with characteristics distributed symmetrically around a policy of interest. Based on the bounds on the indifference curves, I estimate the range of WTP for a change in x consistent with the data. Even when the data are generated completely at random (a case I call “naïve sampling”), the simulations illustrate weak convergence, with the width of the estimated interval for WTP (i.e., the gap between the max and min of conceivable WTP consistent with GARP) shrinking faster than root- n . However, the simulations also highlight the practical importance of optimal sampling, as might be possible with updated survey questions or as might be generated from market data from profit-maximizing firms and price discovery. When the sample is taken sequentially and prices are dynamically set based on a guess from previous samples, the gaps between max and min WTP contract quite fast, essentially obtaining point identification at 50 choice occasions. Appendix Figure A1 depicts a graphical summary of these simulations and Table A1

gives numerical results.

2.2 Bounding Willingness to Pay

Even when indifference curves are only partially identified, we can construct bounds on policy questions defined by scalar-valued functions of β Manski (2010). To do so, we need only change the objective function of the above LP, while satisfying the same Afriat inequalities. In this sub-section, I consider two such policy questions: (i) the WTP to switch to a new policy scenario, such as a change in the vector of public goods; and (ii) the highest price that can be charged for a product competing against a set of substitutes.

Value of a new policy scenario

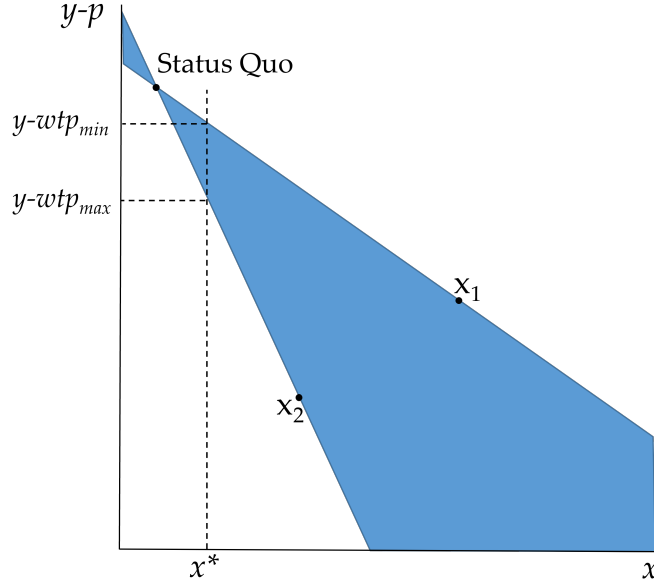
Suppose we are interested in the WTP for an alternative with characteristics $\mathbf{x}^*$ when the only other alternative is the outside option. For example, suppose $\mathbf{x}$ is a vector of public goods and we wish to know individuals' WTP for a policy moving $\mathbf{x}$ from the status quo to some $\mathbf{x}^*$. Note that in this linear model the status quo can always be renormalized to the outside option by subtracting the status quo values of $\mathbf{x}$ from each alternative, so characteristics now represent changes.

Since each candidate utility function implies a unique WTP, it is easy to see that because utility is weakly identified in this model so is WTP. Moreover, the following theorem shows how we can construct bounds on WTP for such a policy question.

Theorem 2 If the data are linearly rationalizable and the set of $\{\beta_i\}$ rationalizing the data is bounded, then the WTP for some counterfactual scenario characterized by $\mathbf{x}^*$, relative to the status quo, can be bounded by solving the two linear programs: $Max_{\beta_i} \sum_i \beta'_i \mathbf{x}^*$ and resp. $Min_{\beta_i} \sum_i \beta'_i \mathbf{x}^*$ subject to $\beta_i \geq \mathbf{0}$ and Inequalities (3).

Proof. Note first that the set $\{\beta_i \mid \beta_i \geq \mathbf{0} \text{ and } \beta'_i(\mathbf{x}_{j'(i,t)t} - \mathbf{x}_{jt}) \geq p_{j'(i,t)t} - p_{jt} \forall t, \forall j \in \mathcal{J}(i,t)\}$ is the intersection of half spaces, so it is a convex polytope. Because it is bounded as well, the set of $\{\beta_i\}$ satisfying GARP is compact. Then the objective described in the proposition is to maximize (resp. minimize) a continuous function defined on a compact set. By the extreme value theorem, there is a solution to this program, and it is the arg-max (resp. arg-min) of the function. Finally, note trivially that this

Figure 3: Bounding Willingness to Pay Relative to the Status Quo



The figure illustrates the set of plausible WTP values for moving from the status quo to a bundle $\mathbf{x}^*$.

maximized (resp. minimized) function $\sum_i \beta'_i \mathbf{x}$ is the WTP function.

Figure 3 illustrates Theorem 2, assuming as in Figure 1 that we observe $\{x_1, p_1\} \succ_R SQ \succ_R \{x_2, p_2\}$. Steeper candidate indifference curves in the shaded region represent higher marginal WTP for $\mathbf{x}$; flatter curves represent lower marginal WTP. Therefore, we simply “search” over all indifference curves satisfying GARP to find the highest and lowest. The figure shows the highest and lowest WTP, consistent with the data, that would make the individual indifferent between $\mathbf{x}^*$ and the status quo—that is, the range of feasible WTP.

To implement this LP, we need to assure that the set of $\{\beta_i\}$ satisfying GARP is bounded. Without additional restrictions, it may not be. For example, the (finite) data for a particular individual may show it always chooses the alternative with the highest value of some attribute m , regardless of price. In that case the set of $\{\beta_{ik}\}$ satisfying GARP would be unbounded from above. A practical solution is to set some arbitrary (but defensible) upper bound on the support for β through additional inequalities. Equivalently, one could set an upper bound through one additional constraint on the maximal permissible WTP, i.e. $\beta' \mathbf{x}^* \leq WTP^{max}$. When the $\{\beta_i\}$

are actually unbounded but such a constraint is binding, the above procedure will then return the upper bound of the support. At the researchers' discretion, values at the boundary could be flagged.

Pricing a product

Another policy question of interest is the highest price a consumer would pay for a discrete product with characteristics $\mathbf{x}^*$ when facing some choice set $\mathcal{J}^*$. This is still a WTP, but in the presence of more substitutes. Such a question might still be useful in the context of public policy, when an agency is offering a program into which individuals can select, or in the context of a local government changing a vector of local public goods when individuals can sort into that jurisdiction or others. But it also pertains to industrial organization, where we are interested in the price the market will bear for a product. The following two-part theorem shows how we can construct bounds for such questions.

Theorem 3A. If the data are linearly rationalizable and the set of $\{\beta_i\}$ rationalizing the data is bounded, then the upper bound on the WTP for a product $\mathbf{x}^*$ is determined by the linear program:

$$\begin{aligned} p_i^* &= \text{Max}_{\{\beta_i, p_i\}} \sum_i p_i, \text{ subject to:} \\ \beta_i &\geq \mathbf{0}, \\ \text{Inequalities (3), and} \\ -p_i^* + \beta_i' \mathbf{x}^* &\geq -p_j + \beta_j' \mathbf{x}_j, \forall i, \forall j \in \mathcal{J}^*. \end{aligned} \tag{4}$$

That is, we are looking for the highest price p_i^* that would satisfy the participation constraint (4) that individual i would choose product $\mathbf{x}^*$ from the choice set $\mathcal{J}^*$. At the same time, we search for the β_i that makes the product as valuable as possible, so long as the β_i are consistent with the observed choice patterns in the sense of satisfying GARP, i.e. consistent with Expressions (3). Mechanically, Expression (4) is just like (3), only it is not conditioned by the observed data, but rather by a hypothetical condition requiring the new product to be weakly preferred when priced at WTP. Accordingly, we can just add an additional “observation” to the data that states the individual chooses $\mathbf{x}^*$ at a price p_i^* out of choice set $\mathcal{J}^*$, and then find the highest value of p_i^* that still satisfies GARP. (The proof, which is omitted, follows

directly from the proof of Theorem 2.)

To find the lower bound, we cannot just minimize the same objective, because we are not looking for the lowest possible price we could charge for a product such that consumers buy it (which might be low, indeed!). We are looking for the lower bound, consistent with the data, on the *highest* price we can charge. However, we can obtain this value by reversing the sign in the participation constraint.

Theorem 3B. If the data are linearly rationalizable and the set of β_i rationalizing the data is bounded, then the lower bound on the WTP for a product $\mathbf{x}^*$ is determined by:

$$p_i^* = \min_{j \in \mathcal{J}^*} \left\{ \min_{\beta_i \geq 0} \{ \beta_i'(\mathbf{x}^* - \mathbf{x}_{jt^*}) + p_{jt^*} \} \right\}, \text{ subject to:}$$

$$\beta_i \geq \mathbf{0} \text{ and}$$

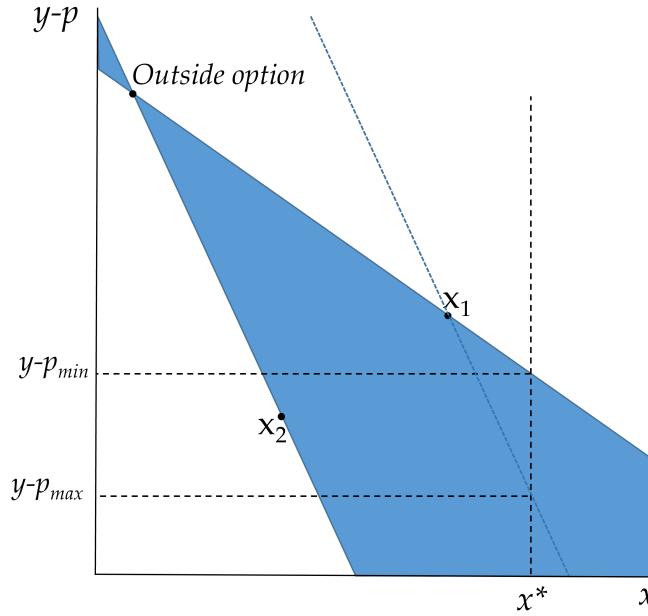
Inequalities (3).

Proof. For a given β_i , the highest price we can charge i for product $\mathbf{x}^*$ in period t^* is $\beta_i' \mathbf{x}^* - \max_{j \in \mathcal{J}^*} \{-p_{jt^*} + \beta_i' \mathbf{x}_{jt^*}\}$, which is the consumer surplus for $\mathbf{x}^*$ relative to the best alternative. This price is equivalent to $\min_{j \in \mathcal{J}^*} \{ \beta_i'(\mathbf{x}^* - \mathbf{x}_{jt^*}) + p_{jt^*} \}$. By continuity and monotonicity of v , the greatest lower bound on the individual's maximum willingness to pay is the lowest possible value of this expression, subject to GARP. Thus, we seek $\min_{\beta_i \geq 0} \left\{ \min_{j \in \mathcal{J}^*} \{ \beta_i'(\mathbf{x}^* - \mathbf{x}_{jt^*}) + p_{jt^*} \} \right\} = \min_{j \in \mathcal{J}^*} \left\{ \min_{\beta_i \geq 0} \{ \beta_i'(\mathbf{x}^* - \mathbf{x}_{jt^*}) + p_{jt^*} \} \right\}$ such that Inequalities (3) are satisfied.

While the double $\min\{ \}$ function may at first look daunting, note that, in practice, we can simply solve the LP $\min_{\beta_i \geq 0} \{ \beta_i'(\mathbf{x}^* - \mathbf{x}_{jt^*}) + p_{jt^*} \}$ separately for each alternative $j \in \mathcal{J}^*$ and then find the minimal value of the solution over $\mathcal{J}^*$. Potentially, the value may be negative, but this is a completely valid finding: we might have to pay people to induce them to use a new product instead of their usual favorite.

Figure 4 illustrates the procedure, for the same data as Figure 3 (but relabeling the status quo as the “outside option”). A new product is shown at x^* with higher x than any other existing product. Because $\{x_1, p_1\}$ is the customer's current choice, that is the one the new product must outcompete. Because the new product has higher x than x_1 , we know for sure we can charge a higher price than p_1 . In fact, we can charge at least p_{min} shown in the figure, because that is on the linear indifference

Figure 4: Bounding the Highest Price for a New Product



The figure illustrates the plausible range of the highest price for an unobserved new good when, in this example, we observe that $\mathbf{x}_1$ is currently the preferred good and the outside option is preferred to $\mathbf{x}_2$.

curve associated with the lowest value x_1 consistent with the data and GARP. We can charge no more than p_{max} , as that is associated with the highest WTP for x consistent with the data. How do we know this? The slope of the line going through x_2 represents the highest value of x consistent with the data, but we need to outcompete x_1 , not x_2 . So we translate that marginal willingness to pay through x_1 as shown by the dashed blue line in the figure.

Aggregating across individuals, the solutions bound the demand curve, and hence the revenue that a firm could obtain under perfect price discrimination. Bounding revenue for a single-price monopolist (or oligopolist in Bertrand competition) would involve taking the solutions as bounds on discrete demand functions and searching for the revenue-maximizing prices.

2.3 Introducing Unobserved Quality into the Model

So far, we have assumed the econometrician observes all relevant information about the choice alternatives faced by individuals. However, if they are reacting to un-

observed characteristics, individuals may appear to violate GARP even when they are behaving rationally, and the sets of β will be misidentified. In their application to computers, for example, Bajari and Benkard (2005) found that, even with 19 characteristics included in $\mathbf{x}$, over half of the data could not be rationalized by a linear utility function with no errors (see also Kuminoff 2009, for discussion). One response to this problem is to introduce idiosyncratic errors, as in most RUMs, which I will consider in Section 3. First, however, consider the possibility of adding only a vertically differentiated unobserved quality attribute.

In particular, consider the following pure characteristics model from Bajari and Benkard (2005) and Berry and Pakes (2007):

$$v_{ijt} = -p_{jt} + \beta'_i \mathbf{x}_{jt} + \lambda_i \xi_j, \quad (5)$$

where ξ_j is an unobserved time-invariant characteristic and $\lambda_i \geq 0$ is an individual-specific weight.⁵ The unobserved characteristic ξ can be thought of as an index of multiple characteristics, so long as the underlying index is the same for everybody. Additionally, utility is monotonic in ξ , so the unobservables differentiate products vertically. Nevertheless, it is easy to see how the model can be interpreted as a RUM with $u_{ijt} = \lambda_i \xi_j$.

We can still bound the value of existing products or policy scenarios using this model. For example, following the outline of the previous sub-section, we can bound the value of some product $(\mathbf{x}^*, \xi^*)$ relative to an outside option by solving the following program:

$$\begin{aligned} & \max_{\{\beta_i, \lambda_i, \xi_j\}} \text{ (resp. } \min_{\{\beta_i, \lambda_i, \xi_j\}}) \sum_i \beta'_i \mathbf{x}^* \text{ subject to:} \\ & \beta_i \geq \mathbf{0}, \lambda_i \geq 0, \text{ and} \\ & -p_{j'(i,t)t} + \beta'_i \mathbf{x}_{j'(i,t)t} + \lambda_i \xi_{j'(i,t)} \geq -p_{jt} + \beta'_i \mathbf{x}_{jt} + \lambda_i \xi_j, \quad \forall i, \forall t, \forall j \in \mathcal{J}(i, t). \end{aligned} \quad (6)$$

Here, the Inequalities (6) replace (3) as the Afriat inequalities, simply adding $\lambda_i \xi_j$. It would be straightforward to answer other economic questions using a similar strategy.

Three remarks are worth making. First, in taking $\sum_i \beta'_i \mathbf{x}^*$ as the objective, I am

⁵Note we can normalize $\lambda_i = 1$ for one individual and $\xi_j = 0$ for one product (such as the outside option).

implicitly assuming researchers will want to set $\xi^* = 0$ for the new product, but this is not necessary. One could incorporate an arbitrary value of ξ^* or place it at some percentile of the distribution. Second, the bilinear terms $\lambda_i \xi_j$ in Expression (6) require quadratic (rather than linear) programming. However, the objective remains convex. Third, we now must impose an assumption of mean independence between the ξ_j and $\mathbf{x}$. Normally, in econometric practices, such an assumption is simply assumed. Here, because the ξ_j are solved for, we need to impose it in the program. In practice, we can do this by binning the joint distribution of $\mathbf{x}$ into cells $c = \{1 \dots C\}$, and imposing the additional constraint that $\frac{1}{\sum_{j \in c} 1} \sum_{j \in c} \xi_j$ is equal for all c (i.e., the mean ξ_j is equal across all cells of $\mathbf{x}$).

Somewhat speculatively, it may be possible to generalize the model further. Using the notation $\beta_i = (\bar{\beta} + \tilde{\beta}_i)$ and $\mu_{ij} = f_i(\xi_j) + \bar{\beta}'\mathbf{x}_j$, where $f_i(\cdot)$ is any increasing function, we can re-write Equation (5) as:

$$v_{ijt} = -p_{jt} + \tilde{\beta}'\mathbf{x}_j + \mu_{ij},$$

with the restriction

$$(\mu_{ij} - \mu_{ij'}) (\mu_{i'j} - \mu_{i'j'}) \geq \forall i, i', \forall j, j'.$$

That is, individuals rank the ξ_j in the same order: For any two individuals i, i' and any two products j, j' , if $f_i(\xi_j) + \bar{\beta}'\mathbf{x}_j \geq f_i(\xi_{j'}) + \bar{\beta}'\mathbf{x}_{j'}$ then $f_{i'}(\xi_j) + \bar{\beta}'\mathbf{x}_j \geq f_{i'}(\xi_{j'}) + \bar{\beta}'\mathbf{x}_{j'}$. However, in general this quadratic program is non-convex and may be difficult to solve.⁶

3 When GARP is Not Satisfied: Rationalizing the Data with Minimal Idiosyncratic Errors

If the data do not satisfy GARP, the model to this point will be incoherent: no values of β will satisfy the constraints. Consequently, the model needs to be generalized to allow for otherwise unexplained departures from GARP. We can do this by introducing additive idiosyncratic errors, as is common practice in discrete choice modeling. We

⁶Too, without imposing additional identifying restrictions on the joint distribution of $\{\mathbf{x}, \xi\}$, we would no longer identify the mean tastes for the product $\bar{\beta}$ separately from $f_i(\xi_j)$. However, recovering the parameter μ_{ij} is sufficient information for many policy questions.

now have the utility function:

$$v_{ijt} = -p_{jt} + \beta'_i \mathbf{x}_{jt} + \lambda_i \xi_j + \varepsilon_{ijt}.$$

The Afriat inequalities become:

$$\begin{aligned} -p_{j'(i,t)t} + \beta'_{i'} \mathbf{x}_{j'(i,t)t} + \lambda_{i'} \xi_{j'(i,t)} + \varepsilon_{ij'(i,t)t} &\geq -p_{jt} + \beta'_i \mathbf{x}_{jt} + \lambda_i \xi_j + \varepsilon_{ijt}, \\ \forall i, \forall t, \forall j \in \mathcal{J}(i, t). \end{aligned} \quad (7)$$

If we view these idiosyncratic errors through the lens of the revealed preference literature (and estimate the model with an LP), the ε_{ijt} are now taken as parameters (in econometrics parlance, “variables” in programming parlance) which allow us to satisfy these inequalities for every choice occasion.⁷ That is, we are not just fitting market shares or the probabilities that a consumer chooses a product given some distribution of ε , we are jointly estimating permissible sets of β and ε that perfectly explain each observed choice.⁸ As Kuminoff (2009) discusses, this makes such models radically incomplete. The set of $\{\beta_i\}$ consistent with GARP are now unbounded: For any β_i , there are always idiosyncratic errors that could explain a choice. Table 1 illustrates this fact. It depicts a stylized choice context with 5 choice occasions, each with a status quo option and two other alternatives characterized by a scalar-valued characteristic x and a cost p . Model I can rationalize these data with $\beta = 5$ and with errors that were drawn from an i.i.d. logit distribution. The three columns show the expected utility (i.e. $\beta' \mathbf{x}$), the logit errors, and realized random utility. It can readily be verified that the utility is highest for the chosen alternatives. Model II, summarized by a similar set of three columns, also can rationalize these data with $\beta = 4.85$ and an alternative set of errors, an alternative set with a much smaller norm ($\sum |\varepsilon| = 9.1$ vs 3,661). However, the errors are not mean independent of $\mathbf{x}$. Model III imposes mean independence, returning to the estimate that $\beta = 5$, with errors almost as small as those of Model II ($\sum |\varepsilon| = 10.0$).

Accordingly, we need some rule for resolving, or at least reducing, this incom-

⁷In this paper, I use the language of econometrics: “variables” are given and the objects to be estimated are “parameters.” Confusingly, in the language of programming, “parameter” are given and the objects to be estimated are “variables.”

⁸In this sense, the “solved” errors are comparable to the conditional distribution of parametric RUM errors, where one conditions on observed choices (Von Haefen 2003).

Table 1: Hypothetical Example Illustrating the Incompleteness of Random Utility Models

Choice Occasion	Alternative	x	p	Choice	Model I: $-\beta x/\alpha = 5$			Model II: $-\beta x/\alpha = 4.85$			Model III: $-\beta x/\alpha = 5$		
					Expected Utility	Error	Utility	Expected Utility	Error	Utility	Expected Utility	Error	Utility
1	1	0	0	1	0	137.1	137.1	0	2.49	2.49	0	5	5
1	2	3	20	0	-5	-223.2	-228.2	-5.44	0	-5.44	-5	0	-5
1	3	17	80	0	5	-488.3	-483.3	2.49	0	2.49	5	0	5
2	1	0	0	0	0	-110.0	-110.0	0	0	0	0	0	0
2	2	11	70	0	-15	49.3	34.3	-16.62	-0.83	-17.45	-15	0	-15
2	3	6	10	1	20	469.1	489.1	19.11	0.66	19.78	20	0	20
3	1	0	0	0	0	-7.0	-7.0	0	0	0	0	0	0
3	2	6	40	0	-10	154.4	144.4	-10.89	0	-10.89	-10	0	-10
3	3	4	10	1	10	343.2	353.2	9.41	0	9.41	10	0	10
4	1	0	0	1	0	75.8	75.8	0	0	0	0	0	0
4	2	7	40	0	-5	-521.2	-526.2	-6.03	0	-6.03	-5	0	-5
4	3	17	90	0	-5	-56.0	-61.0	-7.51	0	-7.51	-5	0	-5
5	1	0	0	0	0	-96.0	-96.0	0	0	0	0	-5	-5
5	2	1	10	1	-5	591.4	586.4	-5.15	5.15	0	-5	0	-5
5	3	10	60	0	-10	-318.7	-328.7	-11.48	0	-11.48	-10	0	-10

This table illustrates that the same choice data can be rationalized by three structural models. The illustration assumes we observe five choice occasions each with three alternatives. We observe a scalar-valued amenity x , a tax (or cost) p , and the chosen alternative at each occasion. Model I illustrates that we can rationalize these data with a random utility model with iid logistic errors and a marginal rate of substitution (MRS) between x and money equal to 5. The model shows three columns: Expected utility is the utility ignoring the errors, the errors shown are drawn from the logit distribution, and the utility is the sum of expected utility and the error. Utility is highest for the chosen alternative, indicating rationalization. Model II illustrates that we can also rationalize the data with MRS=4.85 and a distribution of errors with a much smaller norm. The three columns are analogous to those in Model I. Model III imposes mean independence, again recovering MRS=5 as in Model I with a norm of errors almost as small as Model II.

pleteness, ideally a rule with economic and/or statistical foundations (Lewbel 2019). A natural choice is to select the subset of permissible $\{\beta_i, \lambda_i, \xi_j, \varepsilon_{ijt}\}$ that minimizes some norm of ε , not unlike a least squares problem. In fact, the revealed preference literature has a long tradition of doing just that, to find what is known as “critical cost efficiency” (Afriat 1972; Varian 1985, 1990; Chambers and Echenique 2016). The basic idea is to find the smallest adjustment to the data needed to make them consistent with GARP. To understand these adjustments, consider as a simple example a violation of Samuelson’s Weak Axiom of Revealed Preference (WARP). If WARP is violated, the individual chooses bundle A when B is available but on another occasion chooses B when A is available. Equivalently, the individual behaves as if either B was not available the first time or A was not available the second time. With competitive (convex) budget sets the adjustments are equivalent to adjustments to the individual’s income: Either it wasted expenditure in the first case (by buying A when it could have done at least as well with the cheaper B) or vice versa. With discrete alternatives in characteristic space, we can adapt the idea of the efficiency index by considering how much money an individual must seemingly throw away when it makes a particular choice. In other words, how much does it understate the cost of products it buys, or overstate those it does not buy?

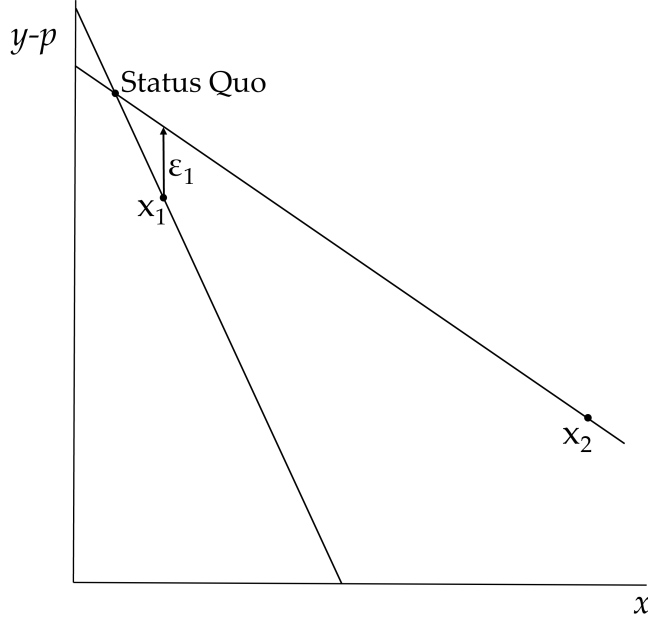
That is, suppose we simply re-write Expression (7) by regrouping some terms:

$$-(p_{j'(i,t)t} - \varepsilon_{ij'(i,t)t}) + \beta'_i \mathbf{x}_{j'(i,t)t} + \lambda_i \xi_{j'(i,t)} \geq -(p_{jt} - \varepsilon_{ijt}) + \beta'_i \mathbf{x}_{jt} + \lambda_i \xi_j, \\ \forall i, \forall t, \forall j \in \mathcal{J}(i, t).$$

Written this way, the relationship between the errors and money becomes transparent. The errors either lower the effective price of the chosen product or raise it for the products not chosen (or both). The interpretation is similar to that of Ho and Pakes (2014), who view errors as measurement error in prices. But whereas they “average out” the errors and derive moment conditions, I incorporate the errors into an LP.

Figure 5 illustrates the idea. Again, $\{x_1, p_1\} \succ_R SQ \succ_R \{x_2, p_2\}$, but the alternatives are relocated in the characteristics space. Now these choices are not linearly rationalizable. The indifference curve going through SQ must pass below $\{x_1, p_1\}$, which puts $\{x_2, p_2\}$ in the strictly-better than set for SQ . We can explain this apparent contradiction with a positive error on the first bundle, as if it had a cheaper price,

Figure 5: Rationalizing the Data with Idiosyncratic Errors



In this example, we observe $\{x_1, p_1\} \succ_R \text{SQ} \succ_R \{x_2, p_2\}$, which is not linearly rationalizable. The vertical distance ε_1 represents the minimum adjustment to p_1 that would rationalize the data.

$\{x_1, p_1 - \varepsilon_1\}$, which shows up as a vertical displacement in the figure, until it appears in the upper contour of the lowest candidate indifference curve going through x_2 . Minimizing the errors in this way can be thought of as bringing the model as close as possible to the pure characteristics model, introducing idiosyncratic errors only when that model violates GARP. Thus, this generalization nests the pure characteristics model, but also address the concerns raised by Athey and Imbens (2007) about the difficulty of forcing the pure characteristics model on data that violate it.

In the application below I minimize the sum of absolute values of the errors.⁹ I also impose a mean independence assumption, which I incorporate as a constraint.¹⁰

⁹Varian (1990) suggested this is computationally difficult when using his proposed algorithms. However, as I show in the following section, it can easily be built into the objective functions for LP tests of discrete choices proposed by Cosaert and Demuynck (2015). Although the usual absolute value function is non-linear, it is well known that it can be reformulated as an LP by replacing ε_i with $(\varepsilon_i^+ - \varepsilon_i^-)$ in the constraints, where ε_i^+ and ε_i^- are both defined as positive variables, and minimizing $(\varepsilon_i^+ + \varepsilon_i^-)$ in the objective.

¹⁰In practice, all the models estimated here also satisfy median independence, which arguably is more fitting for a model minimizing absolute errors. Defining the median of ε as $\inf(t : \text{prob}(\varepsilon \leq t) \geq 0.5)$, median independence is also satisfied whenever there is a large enough probability mass

Model III in Table 1 is a simple example of such a solution. Note also that the distribution of ε can—and typically will—have a large probability mass at zero. (Indeed, it has a degenerate distribution with all mass at zero when GARP is satisfied.) Thus, the recovered distribution function of ε will not be continuous. (These issues are further discussed below in the empirical application.)

As McFadden (2014) has observed, in a cross section, the interpretation of idiosyncratic errors as fixed inter-agent heterogeneity or as intra-agent heterogeneity across choice alternatives (based on a “moment’s fancy”) are observationally equivalent. By the same token, so too in a panel are the interpretation of the errors as either transient tastes and/or shifting quality, or alternatively as irrational behavior in the face of unchanging tastes and quality. That is, ε could be a measure of the error, either by the agents or the analyst in observing prices, or it could be a measure of how wrong the pure characteristics model is in the first place, how mistaken we are to ignore idiosyncratic tastes.¹¹ We can never tell the difference between a model where people do not violate GARP, but where there are transient shocks to tastes and/or transient unobserved shocks to product quality, and one where there are no such shocks and people do violate GARP. This observational equivalence is why we can interpret the errors through either the lens of critical cost efficiency or through the lens of random utility. To my knowledge, this connection between critical efficiency in the GARP framework and errors in the RUM framework has not previously been pointed out.

Even starting from the RUM framework rather than the GARP framework, arguably the errors still are best interpreted as just a retrofit. They are a way to improve the fit of an empirical model when tastes for characteristics alone cannot fully explain observed choice patterns. Thus, minimizing a metric of the errors makes sense regardless of the interpretation. If they are interpreted as actual violations of GARP, they should be minimized following the logic of critical cost efficiency. Alternatively, if they are interpreted as idiosyncratic tastes, they still should be minimized so that the pure characteristics portion of the model, $\beta_i' \mathbf{x} + \lambda_i \xi_j$, gives the best fit to the data

at $\varepsilon = 0$, which is the case for all models I have considered.

¹¹Of course, as the argument has been developed so far, they also could be a measure of how wrong the linearity assumption is. However, because I have only introduced linearity for expositional simplicity, I ignore this possibility here. Section 4 relaxes linearity, taking a nonparametric approach.

possible. However, when adopting the GARP framework, “the best fit to the data” becomes equivalent to “satisfying GARP,” or finding the pure characteristics model that comes as close to the data as possible. This is different than the standard RUM approach of integrated over the errors to fit market shares.

To incorporate the policy question into this model, one can minimize a weighted sum of the norm of ε and the policy objective, with sufficient weight on the ε ’s to guarantee they are minimized, with the range of answers to the policy limited to the set of parameters where the norm of ε is tied. For example, to maximize WTP for a bundle $\mathbf{x}^*$ relative to the status quo, one could find

$$\underset{\{\beta_i, \varepsilon_{ijt}\}}{Min} \quad \omega \sum_i \sum_j \sum_t |\varepsilon_{ijt}| - (1 - \omega) \sum_i \beta_i' \mathbf{x}^*$$

subject to $\beta_i \geq \mathbf{0}$, Inequalities (3), and mean independence, with arbitrarily high weight ω . To minimize WTP, one would find

$$\underset{\{\beta_i, \varepsilon_{ijt}\}}{Min} \quad \omega \sum_i \sum_j \sum_t |\varepsilon_{ijt}| + (1 - \omega) \sum_i \beta_i' \mathbf{x}^*$$

subject to the same conditions.

In practice, this procedure sometimes will point identify β_i for any individual (or group) that violates GARP, as the focus is on limiting the norm of ε . Returning to the policy questions, this means that, say, the sum of WTP may involve the sum over point-identified WTP values (for individuals that violate GARP) and over partially identified WTP values (for individuals that don’t), which of course still results in partially identified aggregate WTP. In theory, we cannot rule out the possibility that β_i is only partially identified even for individuals who do violate GARP, if two solutions are tied on the criteria of minimizing the norm of ε but result in different WTP values. These ties might happen, for example, in the following nonparametric model when violations of GARP happen away from the policy-relevant region in characteristics space.

To assess the performance of the estimator, I run simulations as described in Appendix A2. I consider four data-generating processes (DGPs), generate observed choices, and then estimate both mixed logit models and the GARP-based models proposed in this paper. One DGP is in fact the logit model. The others consider

successively smaller norms of the error distribution. See the appendix for details. Appendix Table A2 displays the results. Interestingly, both models perform well at estimating *marginal* values for all DGPs. The key issues arise when we integrate over the idiosyncratic errors, which comes into play when evaluating the value of new products. The results illustrate a key point of this paper: RUMs are incomplete (Kuminoff 2009; Lewbel 2019). When choice data is generated from an underlying logit structure, the logit model naturally can recover that structure. But the GARP approach recovers an alternative structure—also consistent with the observed choice data—that is much closer to the pure characteristics model, with a lower value for including new products.

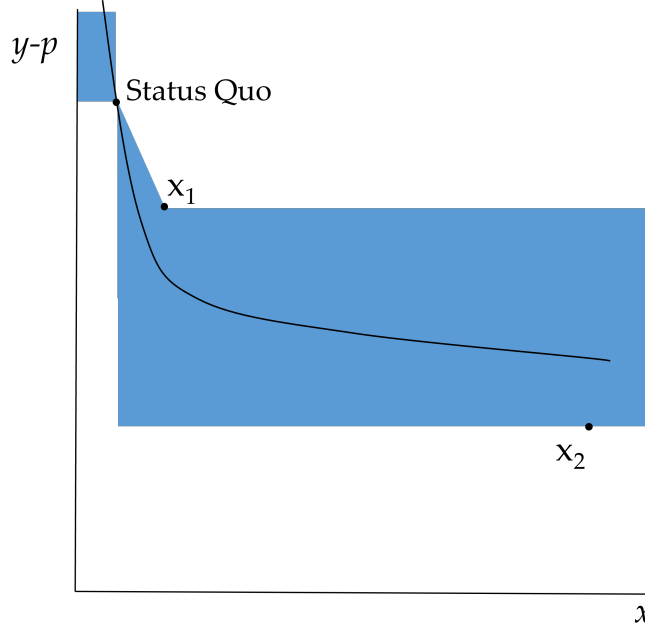
When the choice data is generated from a DGP that is in fact closer to the pure characteristics model, the logit model tends to overstate the value of new products, a point previously emphasized by Berry and Pakes (2007) in motivating the pure characteristics model. In contrast, the GARP approach is more accurate.

4 A Nonparametric Approach

So far, to develop the argument on familiar ground, I have relied on a linear utility function. But eschewing such functional form restrictions has been a central feature of the revealed preference approach since Samuelson first introduced it. Figure 6 illustrates the potential importance of relaxing linearity. It repeats the example from Figure 5, with $\{x_1, p_1\} \succ_R \text{SQ} \succ_R \{x_2, p_2\}$. As previously noted, this choice behavior is not linearly rationalizable. However, it is easily rationalizable by a quasi-concave utility function, such as the indifference curve shown in the figure. Thus, these data are actually consistent with the most general expression of GARP; only the linearity assumption is violated. With competitive budget sets, Afriat’s theorem allows us to easily test whether there is any locally non-satiated utility function that rationalizes the data, using LP. With discrete choices, we can still test for the existence of any strictly monotonic, continuous, quasi-concave utility function that rationalizes the data (Cosaert and Demuynck 2015). In Figure 6, the shaded area represents the valid region for the indifference curve passing through SQ, given the choice behavior, monotonicity, and quasi-concavity.

Extending Afriat’s basic logic, Cosaert and Demuynck (2015, Thm. 3) formally demonstrate that, for an individual i , discrete choice data satisfy GARP, and can be

Figure 6: Rationalizable but not Linearly Rationalizable



This figure repeats the example of Fig. 5. The blue shaded area shows the permissible area for indifference curves going through the status quo.

represented by a weakly monotonic concave utility function, if and only if there exist numbers $v_{ij'(i,t)}, v_{ij}, \hat{\alpha}_{ij}, \hat{\pi}_{ij}$ such that

$$v_{ij'(i,t)} \geq v_{ij}, \forall i, \forall t, \forall j \in \mathcal{J}(i, t). \quad (8)$$

$$v_{ij} - v_{ik} \geq -\hat{\alpha}_{ij} + \hat{\pi}'_{ij}(\mathbf{x}_j - \mathbf{x}_k), \forall i, \forall t, \forall j, k \in \mathcal{J}(i, t). \quad (9)$$

Expression (8) simply states that we must find utilities for chosen bundles j' that exceed those for other bundles available on that choice occasion. Expression (9) states that for any alternative j , there is a supporting hyperplane, defined by $\hat{\alpha}_{ij}$ and $\hat{\pi}_{ij}$, to the indifference curve going through bundle j . This hyperplane represents a vector of shadow prices on the characteristics (times the marginal utility of income). Although the set of alternatives is discrete, the shadow prices represent the budget constraint that would have induced a purchase of j over any other alternative k , if, hypothetically, we were in a competitive budget setting. Importantly, they include the marginal utilities of income, explaining why a person might choose a bundle A on one choice occasion even though it is dominated by a bundle B available on another

occasion. These conditions differ from the usual Afriat inequalities (with competitive budget sets) in that in this case the $\hat{\alpha}_{ij}, \hat{\pi}_{ij}$ are shadow values that must be solved for rather than observed prices. However, whereas in the familiar competitive case we can solve for marginal utilities of income given the observed prices, here we solve for shadow prices that are not identified to scale separately from the marginal utilities of income.

I extend Cosaert and Demuyne's model in two ways. First, I introduce additive errors representing critical efficiency, so Expression (8) now becomes

$$v_{ij'(i,t)} + \varepsilon_{ij'(i,t)} \geq v_{ij} + \varepsilon_{ij}, \quad \forall i, \forall t, \forall j \in \mathcal{J}(i, t),$$

which are the Afriat inequalities accounting for the errors.

Second, I introduce the policy as an artificial alternative j^* . Letting 0 represents the baseline or status-quo state for comparison, maximum WTP can be found by solving for the maximum value of WTP such that

$$v_{ij^*} \geq v_{i0}, \tag{10}$$

$$v_{ij^*} - v_{i0} \geq -\hat{\alpha}_{ij^*}[(p_{j^*} + WTP) - p_0] + \hat{\pi}'_{ij^*}(\mathbf{x}_{j^*} - \mathbf{x}_0),$$

and

$$v_{i0} - v_{ij^*} \geq -\hat{\alpha}_{i0}[p_0 - (p_{j^*} + WTP)] + \hat{\pi}'_{i0}(\mathbf{x}_0 - \mathbf{x}_{j^*}),$$

The intuition is that we are looking for the highest possible price that would be associated with the policy and allow us to choose it over the status quo, while still satisfying GARP. Minimum WTP can be found by maximizing the value of WTP with the inequality in Expression (10) reversed. In that case, we are looking for the lowest possible price that would allow us to choose the status quo over the policy while satisfying GARP. The two values bracket WTP . As written, this problem is non-linear because it involves the product of two variables, $\hat{\alpha}$ and WTP . In my application, the NLP converged quickly to a solution using modern non-linear solvers like Baron (Khajavirad and Sahinidis 2018). However, if analysts find it difficult in their applications, the problem can also be solved by nesting an LP within a line

search over WTP .¹²

5 Application to Valuing Ecosystem Services

I apply the methods proposed in this paper to the valuation of improvements in ecosystem services in the Southern Appalachian Mountains of the United States. The data are based on a “conjoint” survey, or choice experiment, of households in the region. Choice experiments ask respondents to choose an alternative from a hypothetical menu of options, with each option described by a set of attributes and a price (see e.g., Ben-Akiva et al. 2019; Mühlbacher and Johnson 2016). Such data make for an appropriate illustration of the proposed methods, for three reasons. First, they involve a panel of choices at the individual level, making it possible to partially identify each β_i . Second, the data involve experimentally controlled characteristics $\mathbf{x}$ as well as prices, sidestepping the need for instruments in this case. Finally, data from choice experiments are widely used in marketing studies for new products, as well as in applications to environmental policy (e.g., Banzhaf et al. 2016; Lewbel et al. 2011) and health (Mühlbacher and Johnson 2016).

The area studied includes portions of West Virginia, Virginia, North Carolina, Tennessee, and Georgia. The survey was distributed to residents of those states both online and by mail from October 2009 to December 2010. It introduces the state of ecosystem services in the region, with quantitative information in three dimensions: streams and fish, forests, and birds. Based on scientific models, it describes 20 percent of small streams in high-elevation areas as being affected by acid precipitation (about 60,000 streams). Additionally, the survey describes six species of fish affected in the streams: brook trout, rainbow trout, mottled sculpin, longnose dace, rosyside dace, and fantail darter. It also describes three percent of the forested area as affected by acidification. Finally, it describes three bird species (waterthrush, crossbill, and

¹²Specifically, for the case of maximizing WTP, first bound WTP and begin with an initial guess. Then, simply test for GARP using the Cosaert-Demuynck LP, based on the observed data plus the augmented “observation” that the individual votes for the policy (or buys the product) at the trial WTP value. If these augmented data satisfy GARP, this WTP value is consistent with the observed behavior. The initial guess can be set as the lower bound and a new, higher guess made. If the data are not consistent with GARP, the guess is too high and can be set as the upper bound. This process can be repeated until the WTP range converges to a tolerable limit. A similar process can be used for minimizing WTP. When the initial data do not satisfy GARP, the same procedure can be used, constraining the sum of absolute errors to its value without the augmented policy “observation.”

ovenbird) as harmed by acidification, with their populations now at 65 percent of what they once were.

The survey next tells respondents that a liming program has been proposed to improve the streams, forests, and affected bird populations in the region. After describing the program, the survey explains that the respondent's state government will fund its share of the program with a revenue bond that will be paid off by additional state income tax payments for the next 10 years. Next, the survey describes a set of improvements to streams, forests, and bird populations that the program could achieve over a ten year period. It provides a sequence of six choice sets, each including the status quo (or "no program") option plus one to two variants of the program. The choice alternatives range from no improvement in each dimension, to a maximal improvement of 45,000 streams (15% of all streams), 625,000 acres of forest improvement (2.5% of all forests) or a 30 pp improvement in the three bird populations (from 65% of what they once were to 95%). Respondents indicated whether they would be willing to support the program and the indicated annual cost to them over ten years. Collectively, 1,512 respondents evaluated 293 unique bundles of the three attributes on at least one choice occasion, or 1,651 bundles when also considering different costs.¹³ The alternatives were selected randomly using a Bayesian experimental design with updating, to improve the statistical efficiency of the variable space. See Banzhaf et al. (2016) for additional details about the survey and its implementation.

Based on these data, I evaluate the mean WTP of one bundle of possible improvements to ecosystem services from a reduction in acidification from air pollution, in particular improving 15,000 streams in the region, 250,000 acres of forest, and 15 pp of bird populations. These improvements are within the range of the data considered by the survey. I assume they take 50 years to reach full maturity, and that they then continue indefinitely. I report annualized WTP for this full stream of services (i.e., as of year 50) in present value terms, using a discount rate of 3 percent.

For comparison purposes, I first compute values using a simple parametric RUM, a linear logit model imposing homogeneity in taste parameters, with heterogeneity

¹³An additional 1,678 respondents faced one or two choice occasions in a traditional dichotomous choice contingent valuation format. These respondents were included in the original parametric analysis of the data in Banzhaf et al. (2016) but were excluded here.

entering only through the additive errors. Parameters include the marginal utility of money α and a vector β for streams, forest, and birds, plus a dummy variable for the status quo option. This dummy can be interpreted as representing the effect of a reference point, the taste for “doing something” if positive or “not interfering” if negative. Alternatively, it can simply be interpreted as adding a kink to an otherwise linear indifference curve, to improve the fit to the data. Other than this dummy, I do not include product-specific unobservables ξ_j , as in this application there are no obvious “products” to model. After the logit, I estimate the linear GARP alternative, with no distributional assumptions on the errors, based on the logic of Theorem 2. This model finds the vector β that minimizes the additive errors needed to satisfy GARP. In this case, the panel structure of the data can be thought of as individuals representing separate observations of a representative consumer.

Table 2 displays the results. Because the policy relevance of tastes for the status quo is controversial (Banzhaf et al. 2016), the table shows WTP estimates respectively ignoring or including this parameter. The first column shows the results for the standard multinomial model. Mean WTP is \$52 if we ignore the effect of the status quo dummy when calculating WTP, with a 95% confidence interval (based on standard errors clustered by individual respondent) spanning \$40 to \$64. If we allow for the disutility of interfering with the status quo when evaluating the policy, mean WTP falls to \$21, with a confidence interval of \$7 to \$35. The average absolute difference in logit errors between any two options, normalized by the marginal utility of income, is \$146, indicating substantial randomness in this RUM relative to the fit of the parameters.¹⁴

The second column shows the results for the GARP approach proposed in this paper, imposing homogeneity in β . Unlike the logit model, the GARP approach imposes no distributional assumption on the ε ’s, except for satisfying mean independence. It imposes this independence through constraints, requiring that the ε ’s are mean zero in each of 33 cells, as well as globally orthogonal to p and $\mathbf{x}$. Because this model imposes homogeneity in the face of what is likely substantial heterogeneity, all the burden of the model is on satisfying GARP as close to possible, with no weight

¹⁴The logit model is only identified to scale. Although the standard procedure is to normalize the scale parameter of the idiosyncratic errors to 1, there is nothing to prevent us from instead normalizing the marginal utility of income to 1, as in Section 2. This puts the errors in dollar terms and makes them more comparable to the errors from the GARP model.

Table 2: Results Imposing Homogeneous Preferences

	Multinomial Logit	Linear GARP
Mean WTP without Status Quo	\$52.04 (\$40.32 - \$63.76)	\$25.31 (\$21.31 - \$29.30)
Mean WTP with Status Quo	\$20.55 (\$6.53 - \$34.57)	\$40.12 (\$36.86 - \$43.38)
Average $ \varepsilon $	\$146.04	\$11.83
Pct $\varepsilon = 0$	0%	80.3%
Pct of hholds $\Sigma \varepsilon = 0$	0%	6.7%
Pct Correctly Predicted, Pure Characteristics	50.1%	50.7%
Pct Correctly Predicted, Conditioning on $F(\varepsilon)$	41.7%	49.4%

This table gives the results from estimating a standard logit model and the proposed LP, assuming homogeneous preference functions across households. The numbers in parentheses represent 95% confidence intervals. All standard errors clustered by household. For Linear GARP, the standard errors are based on a clustered bootstrap with 500 draws.

given to altering the parameters to minimize or maximize WTP. Thus, this model is point identified. Mean WTP is estimated to be \$25 when omitting the status quo and \$40 when including it. The confidence intervals, based on cluster-bootstrapping the data by respondent and recalculating the LP for each bootstrapped sample, are fairly tight, at about +/- \$4. These estimates are in the same range as the linear logit model, but they flip the relative order of the models with and without the status quo.

As we would expect, the GARP-based LP approach drives the errors much lower. In other words, it comes much closer to the pure characteristics model. The average $|\varepsilon|$ is only \$12 versus \$146 for the logit, and 80% of observations and 7% of households require no ε at all. However, it has a similar record as the logit at predicting choices, with each predicting about 50% correctly based only on the characteristics. We can also look at the prediction rate when integrating over the ε 's. For the logit model, this is given by the closed-form likelihood, $\frac{e^{\beta' \mathbf{x}_{j'} - p_{j'}}}{\sum_j e^{\beta' \mathbf{x}_j - p_j}}$. For the GARP model, I simulate repeated draws from the empirical distribution of ε 's and calculate the percentage of occasions when $\beta \mathbf{x}_{j'} - p_{j'} + \varepsilon_{ijt} \geq \beta \mathbf{x}_j - p_j + \varepsilon_{ijt}$. The GARP model still outperforms

the logit model by this metric, at 49% correct vs. 42%.

The assumption of homogeneity imposed in these models is a strong one, but it can be relaxed both with parametric methods (e.g. random coefficients, or mixed, logit) as well as with the GARP-based methods suggested in this paper. Table 3 presents results allowing for preference heterogeneity. Again for comparison purposes, the first column shows results using a mixed logit model, with random coefficients on the status quo dummy distributed normal and random coefficients on streams and forests distributed log-normal (to impose positive values). (The coefficient on birds is kept homogeneous, as attempts to estimate heterogeneity were fragile.) Heterogeneity has little effect on the estimate when ignoring the status quo dummy, which is now \$49 (CI \$35 to \$64), compared to \$52 in the model with homogeneity. However, it now increases it substantially when including the status quo dummy, from \$21 to \$77 (CI \$55 to \$100). Including heterogeneity in the random parameters allows the scale of the additive logit errors to fall, with the average absolute value of the difference in errors now at \$75 instead of \$146.

The next pair of columns shows the estimates using the linear GARP-based approach with heterogeneity, as in Theorem 2. As most individuals satisfy GARP, the model is now only partially identified, with a continuum of β_i satisfying GARP but having different implications for WTP. The two columns in the pair show the results for the argmin and argmax in $\{\beta_i\}$ for the WTP objective. Plausible mean WTP ranges from \$36 to \$259 when ignoring the status quo dummy, and -\$130 to \$190 when including it. Both values bracket the respective range from the logit model, although this result is not guaranteed.¹⁵ The WTP range reflects model uncertainty, which is large, but sampling variance remains small, with tight confidence intervals. (To compute standard errors, we can now simply take the sampling distribution across households rather than bootstrap.)

Again, the model shrinks the additive errors much more than the logit model assumes, with average errors now \$2 rather than \$75 for logit and \$10 for the GARP approach without heterogeneity, and with about 97% of observations having no need

¹⁵There is nothing that guarantees the GARP-based approach will always bound the logit estimates. Consider for example a case where all the data satisfy GARP. The GARP-approach will partially identify a range of WTP using a pure characteristics model (with no idiosyncratic errors). The logit estimator conditions on its distributional assumption, which can take it outside the range of the GARP model.

Table 3: Results Allowing Heterogeneous Preferences

	Mixed Logit	Linear GARP		Nonparametric GARP (without status quo)		Nonparametric GARP (with status quo)	
		Min	Max	Min	Max	Min	Max
Mean WTP without Status Quo	\$49.07 (34.56 - 63.58)	\$36.48 (33.33 - 39.62)	\$259.40 (249.36 - 269.43)	\$44.44 (41.03 - 47.85)	\$210.09 (199.22 - 220.96)	N/A	N/A
Mean WTP with Status Quo	\$77.40 (55.08 - 99.71)	-\$129.90 (-143.34 - -116.46)	\$189.95 (178.75 - 201.15)	N/A	N/A	-\$281.51 (-295.61 - -267.40)	\$223.40 (212.48 - 234.32)
Avg $ \varepsilon $	\$74.96	\$2.00	\$2.00	\$0.74	\$0.79	\$0.56	\$0.57
Pct $\varepsilon = 0$	0%	97.2%	97.2%	97.3%	97.1%	98.0%	98.1%
Pct of hholds $\Sigma \varepsilon = 0$	0%	74.0%	73.6%	88.0%	88.0%	92.3%	92.3%
Pct Correctly Predicted, Pure Characteristics	N/A	92.8%	92.8%	95.2%	95.2%	96.9%	96.9%
Pct Correctly Predicted, Conditioning on $F(\varepsilon)$	38.5%	82.3%	81.3%	92.5%	92.5%	95.2%	95.2%

This table gives the results from estimating a mixed logit model, the proposed LP model with heterogeneous linear preference functions across households, and the heterogeneous nonparametric model. The numbers in parentheses represent 95% confidence intervals. All standard errors clustered by household.

of errors to fit the data. Seventy-four percent of households have choice patterns that are linearly rationalizable, so require no error at all to fit the data. Finally, the GARP model can fit 93% of choices correctly focusing only on the characteristics. This prediction rate falls to 81-82% when integrating over the error distribution, but this rate still is a substantial improvement over the mixed logit model, which correctly predicts only 39% of choices.¹⁶

The remaining columns use the nonparametric GARP-based approach of Section 4. In this case, separate models were estimated with and without the status quo dummy.¹⁷ The plausible range of mean WTP actually is now \$44 to \$210 for the model without status quo, which is not much different than the linear model. However, the range grows rapidly for the model with the status quo, at $-\$282$ to $\$223$. Because, in this application, so much of the mass of data is at the status quo, it is apparently difficult to nonparametrically identify WTP for policy improvements (without any linear interpolations between data points) while controlling for the status quo nonparametrically as well. Nevertheless, it should be noted that these values do continue to bracket the point estimates from the mixed logit model.

Not surprisingly, these models are able to reduce GARP violations even more, with mean errors falling to about \$0.56 to \$0.79 depending on the model, and with 97–98% of observations and 88–92% of households having no error. Finally, in the bottom rows we see that the models’ predictions are better than the linear version, and remain much better than the mixed logit.

For policy analysis, the nonparametric model without the status quo dummy is probably the most appropriate, as it is the most general but does not allow an arbitrary reference point to affect policy. With 13.3 million households in study area, and using a discount rate of 3 percent, even the most conservative estimate of \$44 implies an aggregate present value of \$10.4 billion for ecosystem service improvements

¹⁶Note it is not possible to compute the predictions in pure characteristics space for the mixed logit, except at the “average tastes,” since the value of the coefficients for any one given respondent is unknown. Thus, the last row of the table treats the random coefficients and additive errors as part of the error distribution $F(\cdot)$, for the mixed logit. In contrast, for the GARP model, the β_i are treated as fixed and only the additive errors are in $F(\cdot)$.

¹⁷It may seem surprising that a dummy variable can be included in characteristics space when the function over them is already non-parametric. The dummy essentially relaxes the constraint of continuity and quasi-concavity in the neighborhood of the status quo, i.e. $\{\mathbf{x} = \mathbf{0}, p = 0\}$.

that might be realized from a reduction in air pollution. When added to the substantial health benefits of most air pollution control policies, these benefits are likely to outweigh the costs of most policies.

6 Conclusions

This paper shows how, starting with a RUM framework, one can nonparametrically partially identify the answers to policy questions using only the Generalized Axiom of Revealed Preference (GARP). This approach recasts the RUM errors as departures from GARP, to be minimized using a minimum-distance criterion, and provides another avenue for nonparametric identification of the RUM model. When GARP is satisfied, this approach easily estimates the pure characteristics model. When GARP is violated, it gets as close to that model as possible.

This approach has four advantages over other approaches. First, by construction, it minimizes a norm of the errors, rather than imposing distributional assumptions that put more weight on random utility than necessary. Second, consequently, it weakly identifies the pure characteristics model (when GARP is satisfied) or comes as close to it as the data allows (when GARP is violated), thus overcoming a long-standing estimation challenge for such models. Third, answers to economic and policy questions can be directly introduced into the objective of the LP, while GARP is satisfied through inequality constraints. Finally, as discussed in Section 4, it is nonparametric, imposing only monotonicity, continuity, and quasi-concavity on the function of observables and additivity and mean independence on the unobservables.

Despite these advantages, the model might still be further generalized. First, given my use of survey data, it was plausible to treat all variables as exogenous. As Berry and Haile (2016, 2024) have emphasized, valid instruments are the core foundation of RUMs. Future work might consider how to introduce additional exclusion restrictions. For example, if prices p are endogenous one might estimate $E[p|\mathbf{x}, \mathbf{z}]$ for additional instruments $\mathbf{z}$ and then minimize errors to satisfy GARP over the conditional expectation of p rather than p itself, not unlike the approach taken by Blundell et al. (2003). The potential to apply generalized instrumental variables (Chesher and Rosen 2017) might be another fruitful avenue to pursue.

Future work might also consider the implications of this approach for designing

experiments. When using surveys or conducting laboratory experiments, researchers have considered algorithms to update the bundles of characteristics offered, to provide the best information for estimating a parametric utility function (e.g. Sándor and Wedel 2005). Further work might consider how the logic of weak identification using GARP might be applied to this problem. As highlighted by the simulations, with experimental data the characteristics presented to respondents can be updated to more efficiently converge to point identification. A related issue is the possibility that tastes may change between observations, especially in naturally occurring data. Blundell et al.’s (2003; 2008) suggestions to adjust budget sets for changes in income might also be combined with the approaches suggested in this paper.

The paper illustrates the approach by estimating bounds on the values of ecological improvements in the Southern Appalachian Mountains. In the most general model, WTP is bounded between \$44 and \$210 using the GARP-based approach, and point identified at \$49 using a mixed logit model. While such bounds may seem less satisfying than a model that is point identified through additional restrictions, they represent an honest appraisal of model uncertainty and can help take the “con” out of econometrics (Leamer 1983).

References

- Afriat, S. N. (1967). The construction of utility functions from expenditure data. *International Economic Review* 8(1), 67–77.
- Afriat, S. N. (1972). Efficiency estimation of production functions. *International Economic Review*, 568–598.
- Allen, R. and J. Rehbeck (2019). Identification with additively separable heterogeneity. *Econometrica* 87(3), 1021–1054.
- Athey, S. and G. W. Imbens (2007). Discrete choice models with multiple unobserved choice characteristics. *International Economic Review* 48(4), 1159–1192.
- Bajari, P. and C. L. Benkard (2005). Demand estimation with heterogeneous consumers and unobserved product characteristics: A hedonic approach. *Journal of Political Economy* 113(6), 1239–1276.
- Banzhaf, H. S., D. Burtraw, S. C. Criscimagna, B. J. Cosby, D. A. Evans, A. J. Krupnick, and J. V. Siikamaki (2016). Valuation of ecosystem services in the Southern Appalachian mountains. *Environmental Science & Technology* 50(6), 2830–2836.
- Bayer, P., R. McMillan, A. Murphy, and C. Timmins (2016). A dynamic model of demand for houses and neighborhoods. *Econometrica* 84(3), 893–942.
- Beckert, W., M. Christensen, and K. Collyer (2012). Choice of NHS-funded hospital services in England. *Economic Journal* 122(560), 400–417.
- Ben-Akiva, M., D. McFadden, and K. Train (2019). Foundations of stated preference elicitation: Consumer behavior and choice-based conjoint analysis. *Foundations and Trends in Econometrics* 10(1-2), 1–144.
- Berry, S. and P. Haile (2016). Identification in differentiated products markets. *Annual Review of Economics* 8(1), 27–52.
- Berry, S. and P. A. Haile (2024). Nonparametric identification of differentiated products demand using micro data. *Econometrica* 92(4), 1135–1162.
- Berry, S., J. Levinsohn, and A. Pakes (1995). Automobile prices in market equilibrium. *Econometrica* 63(4), 841–890.
- Berry, S., J. Levinsohn, and A. Pakes (2004). Differentiated products demand systems from a combination of micro and macro data: The new car market. *Journal of*

- Political Economy* 112(1), 68–105.
- Berry, S. and A. Pakes (2007). The pure characteristics demand model. *International Economic Review* 48(4), 1193–1225.
- Blow, L. and R. Blundell (2018). A nonparametric revealed preference approach to measuring the value of environmental quality. *Environmental and Resource Economics* 69, 503–527.
- Blow, L., M. Browning, and I. Crawford (2008). Revealed preference analysis of characteristics models. *The Review of Economic Studies* 75(2), 371–389.
- Blundell, R., M. Browning, and I. Crawford (2008). Best nonparametric bounds on demand responses. *Econometrica* 76(6), 1227–1262.
- Blundell, R. W., M. Browning, and I. A. Crawford (2003). Nonparametric Engel curves and revealed preference. *Econometrica* 71(1), 205–240.
- Chambers, C. P. and F. Echenique (2016). *Revealed Preference Theory*. Cambridge University Press.
- Chesher, A. and A. M. Rosen (2017). Generalized instrumental variable models. *Econometrica* 85(3), 959–989.
- Cosaert, S. and T. Demuynck (2015). Revealed preference theory for finite choice sets. *Economic Theory* 59, 169–200.
- English, E., R. H. von Haefen, J. Herriges, C. Leggett, F. Lupi, K. McConnell, M. Welsh, A. Domanski, and N. Meade (2018). Estimating the value of lost recreation days from the Deepwater Horizon oil spill. *Journal of Environmental Economics and Management* 91, 26–45.
- Fox, J. T., K. il Kim, S. P. Ryan, and P. Bajari (2012). The random coefficients logit model is identified. *Journal of Econometrics* 166(2), 204–212.
- Gandhi, A. and A. Nevo (2021). Empirical models of demand and supply in differentiated products industries. In K. Ho, A. Hortaçsu, and A. Lizzeri (Eds.), *Handbook of Industrial Organization*, Volume 4, pp. 63–139. Elsevier.
- Gorman, W. M. (1953). Community preference fields. *Econometrica* 21(1), 63–80.
- Hands, D. W. (2014). Paul Samuelson and revealed preference theory. *History of Political Economy* 46(1), 85–116.

- Ho, K. and A. Pakes (2014). Hospital choices, hospital prices, and financial incentives to physicians. *American Economic Review* 104(12), 3841–3884.
- Jackson, M. O. and L. Yariv (2024). The non-existence of representative agents. Unpublished manuscript.
- Keane, M., J. Ketcham, N. Kuminoff, and T. Neal (2021). Evaluating consumers’ choices of Medicare Part D plans: A study in behavioral welfare economics. *Journal of Econometrics* 222(1), 107–140.
- Khajavirad, A. and N. V. Sahinidis (2018). A hybrid LP/NLP paradigm for global optimization relaxations. *Mathematical Programming Computation* 10(3), 383–421.
- Kitamura, Y. and J. Stoye (2018). Nonparametric analysis of random utility models. *Econometrica* 86(6), 1883–1909.
- Kuminoff, N. V. (2009). Decomposing the structural identification of non-market values. *Journal of Environmental Economics and Management* 57(2), 123–139.
- Leamer, E. E. (1983). Let’s take the con out of econometrics. *American Economic Review* 73(1), 31–43.
- Lewbel, A. (2019). The identification zoo: Meanings of identification in econometrics. *Journal of Economic Literature* 57(4), 835–903.
- Lewbel, A., D. McFadden, and O. Linton (2011). Estimating features of a distribution from binomial data. *Journal of Econometrics* 162(2), 170–188.
- Manski, C. F. (2007). Partial identification of counterfactual choice probabilities. *International Economic Review* 48(4), 1393–1410.
- Manski, C. F. (2010). Partial identification in econometrics. In S. N. Durlauf and L. E. Blume (Eds.), *Microeconometrics*, pp. 178–188. Macmillan.
- Mas-Colell, A. (1978). On revealed preference analysis. *Review of Economic Studies* 45(1), 121–131.
- McFadden, D. (2006). Revealed stochastic preference: A synthesis. In *Rationality and Equilibrium: A Symposium in Honor of Marcel K. Richter*, pp. 1–20.
- McFadden, D. (2014). The new science of pleasure: Consumer choice behavior and the measurement of well-being. In S. Hess and A. Daly (Eds.), *Handbook of Choice modelling*, pp. 7–48. Edward Elgar.

- McFadden, D. and K. Train (2000). Mixed MNL models for discrete response. *Journal of Applied Econometrics* 15(5), 447–470.
- Mühlbacher, A. and F. R. Johnson (2016). Choice experiments to quantify preferences for health and healthcare: State of the practice. *Applied Health Economics and Health Policy* 14, 253–266.
- Polisson, M. and J. K.-H. Quah (2013). Revealed preference in a discrete consumption space. *American Economic Journal: Microeconomics* 5(1), 28–34.
- Sándor, Z. and M. Wedel (2005). Heterogeneous conjoint choice designs. *Journal of Marketing Research* 42(2), 210–218.
- Shi, X., M. Shum, and W. Song (2018). Estimating semi-parametric panel multinomial choice models using cyclic monotonicity. *Econometrica* 86(2), 737–761.
- Varian, H. R. (1985). Non-parametric analysis of optimizing behavior with measurement error. *Journal of Econometrics* 30(1-2), 445–458.
- Varian, H. R. (1990). Goodness-of-fit in optimizing models. *Journal of Econometrics* 46(1-2), 125–140.
- Von Haefen, R. H. (2003). Incorporating observed choice into the construction of welfare measures from random utility models. *Journal of Environmental Economics and Management* 45(2), 145–165.

Appendix: Simulations

A1 Weak Identification

To gauge the weak identification of the GARP-based approach, I employ simulations and estimate the rate at which the identified bounds on WTP converge to the known true value.

Following the premise of Theorem 1, I assume that utility takes a linear form and that individuals' choices are rationalizable. Specifically, I assume they have tastes for ecosystem improvements comparable to those estimated in Section 5, with tastes for improvements to streams, forests, and bird populations close to those estimated in the simplest model of that section, but slightly adjusted so that the true WTP for the policy in the Southern Adirondacks is \$30. (As discussed in Section 5, the policy involves an improvement in 15,000 streams, 250,000 acres of forest, and 15 pp of bird populations.)

I randomly draw choice alternatives with these three characteristics, on a uniform distribution between 0 and 2x the policy-level improvements. In one version of the simulations, I also randomly draw prices on a uniform distribution between zero and \$180. I then create choice occasions with three choice alternatives each, two of these draws plus a status quo option of no change (at no cost). Next, using the true utility parameters, I compute which alternative the individuals would choose at each choice occasion. This constitutes all the data for the simulations. Next, I estimate the model with the following total numbers of choice occasions: 5, 10, 15, ..., 50, 60, 70, ..., 150, 200, 250, ..., 500, 600, 700, ..., 1000, 2000, 3000, ..., 10,000. Last, I repeat this process 100 times.

One potential criticism of this simulation is that the prices are naïvely drawn at random, independently of the other attributes. In an experimental setting, such as that of Section 5, we would normally compute a better guess at the price, e.g. with higher-quality alternatives being assigned higher prices. Too, in a market setting we would expect the equilibrium to have converged on prices that are closer to WTP. Accordingly, in a second version of the simulations, at each addition of choice occasions to the data, I use the previous iteration's estimates to take a guess at the WTP for the new alternatives coming in, confronting the individuals with those prices in

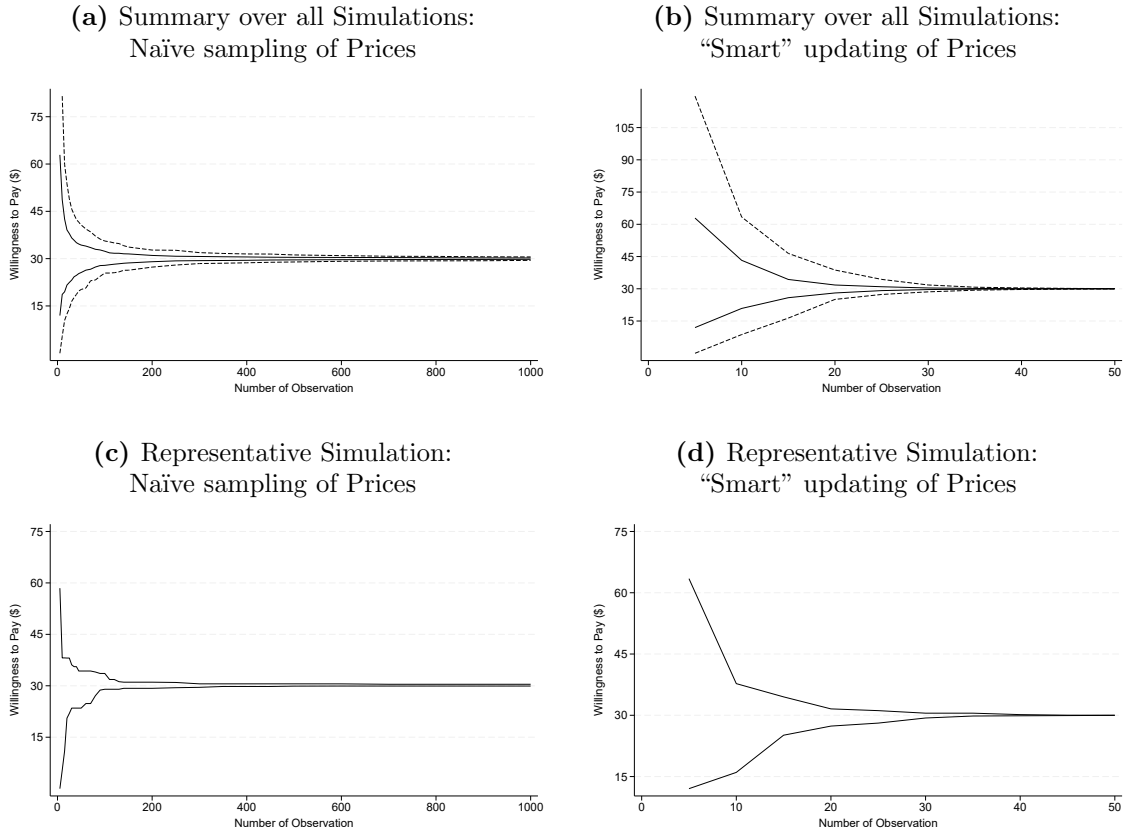
their new choice sets.

Figure A1 shows the results. Panel (a) summarizes the upper and lower bounds across simulations, when using naïve sampling of observed prices. The inner (solid) lines show the median upper and lower bounds respectively, across simulations. The outer (dashed) lines show the 95th percentile of all upper bounds across simulations and the 5th percentile of all lower bounds, at each observed number of choice occasions. (Note that each data point in this figure could come from any simulation.) As shown in the figure, the spreads are very wide at first and do converge rapidly, but improvements begin to stall. The probability of drawing the needed information becomes very low. Panel (b) shows a similar figure but with “smart” sampling of prices (i.e. updated based on previous rounds within the iteration). We now get much more rapid convergence, with, for all practical purposes, point identification at about 50 choice occasions.

The lower panels show the respective results for one representative simulation. For the naively sampled prices (panel c), we can see flat portions in the curves, where the new data does not help at all. Moreover, we still have not reached perfect point identification at 1,000 choice occasions. In contrast, in the “smart” sample (panel d), there is steady convergence, with point identification at about 50 choice occasions.

Table A1 presents the simulation results in tabular form. Each row of the table is a given number of choice occasions. The columns represent the 5th, 50th, and 95th percentiles (across simulations) in the gap between the estimated max and min of WTP. With only five choice occasions, the median range in the identified set of WTP is \$52, but it can be over \$100 in some simulations. By 50 choice occasions, the range falls to a typical level of \$9 with naïve sampling, which may be as good as many applications in the literature today when accounting for model uncertainty. Moreover, this rate of convergence is faster than root- n . For example, from $n = 5$ to $n = 50$ the median gap falls from \$52.28 to \$9.05 or to a factor of 0.17 its former level, vs. $\sqrt{5}/\sqrt{50} = 0.32$, and this remains true throughout the range of sample sizes.

Figure A1: Convergence to Point Identification in Simulations.



The figure illustrates how the gap between the maximum and minimum WTP, consistent with GARP, converges as the sample grows. The simulated true value is \$30. The left column shows results where price data are naïvely sampled from the support regardless of the quality of the alternative. The right column shows results when prices are sampled based on updated estimates from previous rounds. The top row summarizes results across simulations. The inner (solid) lines show the median of all upper and lower bounds, respectively, across simulations. The outer (dashed) lines show the 95th percentile of all upper bounds across simulations and the 5th percentile of all lower bounds, at each observed number of choice occasions. (Note that each data point in this figure could come from any simulation.) The lower panels show the respective results for one representative simulation.

Table A1: Convergence to Point Identification in Simulations.

# choice occasions	Naïve sampling of prices				“Smart” sampling of prices			
	5th percentile	Median	95th percentile		5th percentile	Median	95th percentile	
5	\$19.21	\$52.28	\$101.88		\$19.21	\$52.28	\$101.88	
10	\$14.78	\$32.70	\$65.18		\$8.58	\$21.35	\$50.94	
25	\$7.40	\$16.19	\$30.43		\$0.69	\$1.88	\$5.88	
50	\$3.95	\$9.05	\$16.88		\$0.01	\$0.04	\$0.13	
100	\$1.84	\$4.74	\$9.63		\$0	\$0	\$0	
200	\$1.00	\$2.26	\$4.58		\$0	\$0	\$0	
300	\$0.78	\$1.46	\$3.02		\$0	\$0	\$0	
400	\$0.50	\$1.08	\$2.20		\$0	\$0	\$0	
500	\$0.36	\$0.87	\$1.67		\$0	\$0	\$0	
1000	\$0.17	\$0.45	\$0.90		\$0	\$0	\$0	
2000	\$0.07	\$0.22	\$0.53		\$0	\$0	\$0	
3000	\$0.06	\$0.16	\$0.33		\$0	\$0	\$0	
4000	\$0.05	\$0.11	\$0.26		\$0	\$0	\$0	
5000	\$0.03	\$0.08	\$0.20		\$0	\$0	\$0	
10,000	\$0.02	\$0.04	\$0.12		\$0	\$0	\$0	

This table shows the gap between the upper and lower bounds on estimated WTP, by number of choice occasions included, at three percentiles across the distribution of simulations. It shows results for two strategies for sampling prices: prices that are naïvely sampled from the support regardless of the quality of the alternative and prices that are sampled based on updated WTP estimates from previous rounds.

However, we can do even better with “smart” sampling, where by $n = 50$ the median gap falls to \$0.04. From there, it quickly falls to zero with “smart” sampling, but with random sampling it is still \$0.04 with as many as 10,000 choice occasions. These simulations illustrate weak convergence in principle but also highlight the practical importance of optimal sampling.

A2 Welfare Estimates

To assess how well the LP-based GARP model proposed in this paper performs at estimating WTP, I consider several simulations where the data-generating process (DGP) is known, and, hence, so is WTP. I consider four choice-generating models. In all four, the heterogeneous utility functions take the simple form:

$$u_{ij} = -p_j + \beta_i x_j + \varepsilon_{ij}. \quad (\text{A1})$$

In all simulations, β_i is drawn from a normal distribution with $\mu = 10$ and $\sigma = 2$, so the average marginal WTP for the single attribute x is \$10. In the first DGP, the errors ε_{ij} are drawn from a logit distribution. In the second model, the errors are still from a logit distribution, except half the errors are set to zero. In the third, 98% are set to zero. The idea is to see how the models perform as the norm of the errors approaches zero—i.e., as they approach the pure characteristics model. In this spirit, the fourth model uses the empirical error distribution solved for by the GARP model for basic logit distribution (with about 90% of errors equal to zero).

There are 100 simulations, each with 100 households, with each household facing 10 choice occasions of three choice alternatives. One choice alternative is always the “status quo,” with $\{x = 0, p = 0\}$. The second alternative is randomly drawn with $x \sim u(1, 6)$ and $p/x \sim (0, 20)$. The third alternative is higher quality, with x doubled and the incremental price increase uniformly distributed on $(0, 20)$ per unit of additional x . Thus, within choice occasions the price is necessarily increasing in x , while across choice occasions higher x ’s have higher prices on average. For each choice-generating process, both a mixed logit and the GARP model were estimated on the observed data.

Table A2 displays the results. The columns represent the four “true” DGPs. The rows are the true values of various parameters or the estimates from the logit

model or the GARP model, the latter divided into the minimization and maximization versions, generating lower and upper bound estimates respectively. In all cells, the values in parentheses represent the range from the 5th to 95th percentiles across simulations. The values are organized into three panels. Panel A displays the results for the marginal WTP for the characteristic. Both the logit and the GARP models do well in this respect, accurately recovering the true value of \$10 for all DGPs.

Panel B displays the results for the WTP to include the highest-characteristic (and most expensive) alternative in the choice set. Across the DGPs (columns), the true value for including this alternative declines slightly as the norm of the errors shrink, as the component value due to a high idiosyncratic error draw for the alternative shrinks. Under the pure characteristics model, or if we want to ignore the errors because, say, they represent mistakes that should not influence welfare, the true value is somewhat lower. The two values converge as we move to the right, with fewer idiosyncratic errors in the DGP. Panel C displays the results for the value of having the same alternative replicated again as a fourth option. This change to the choice set can *only* have value in the event that the new alternative gets a sufficiently high draw for the idiosyncratic error. Thus, the values again shrink as the norm of the errors shrinks, this time converging to zero. Under the assumptions of the pure characteristics model, it necessarily has zero value.

When the choice data is generated from the logit model (col. 1), the logit model unsurprisingly does well, accurately capturing the value for the high alternative (panel B) and for the replicated alternative (panel C). By the same token, when the choice data is generated from the GARP model (col. 4), the GARP model does very well, with a range of values for the high alternative that covers the value according to the utilities that generated the choice data. It also accurately captures the true value of zero for the replicated alternative.

The remaining results are situations where one model is fit to choice data that was actually generated by a different DGP. In these cases, the logit model continues to do well at estimating the value for the high alternative in all models (panel b, cols. 2-4). However, it does not do as well at estimating the value of a replicated alternative. As the errors shrink in the DGP (moving across columns to the right), the logit model begins to overstate the value of this alternative, attributing too much to idiosyncratic error draws that are not truly present. This bug of the logit is part of

the original motivation for estimating pure characteristics models (Berry and Pakes 2007).

In contrast, when the choice data are generated from the logit distribution (col. 1), the GARP-based approach does not recover the same model. By construction, it *does* find parameters and error distributions that are consistent with the observed choices. But it finds a different set, one much closer to the pure characteristics model. Accordingly, it finds a smaller value for additional alternatives. On the other hand, as the underlying DGP approaches the pure characteristics model (moving to the right), the logit model is simply wrong while the GARP-based approach recovers the true values.

In summary, RUMs as previously emphasized are incomplete (Kuminoff 2009; Lewbel 2019). When choice data is generated from an underlying logit structure, the logit model naturally can recover that structure. The GARP approach recovers an alternative structure—also consistent with the observed choice data—that is much closer to the pure characteristics model, with a lower value for including new products. When the choice data is generated from an alternative structure, one in fact closer to the pure characteristics model, the logit model tends to overstate the value of new products, while the GARP approach is more accurate.

Table A2: Results from Simulation

Scenario:	Logit errors, all non-zero	Logit errors, 50% zero	Logit errors, 98% zero	GARP errors
A. Marginal WTP				
True value	\$10.00	\$10.00	\$10.00	\$10.00
Mixed Logit	\$9.98 (9.41–10.36)	\$10.02 (9.65–10.40)	\$9.99 (9.60–10.36)	\$9.97 (9.63–10.35)
GARP LP-Min	\$10.11 (9.68–10.51)	\$9.81 (9.49–10.22)	\$8.98 (8.66–9.39)	\$8.94 (8.56–9.32)
GARP LP-Max	\$10.28 (9.78–10.64)	\$10.45 (10.09–10.86)	\$11.00 (10.63–11.43)	\$11.02 (10.72–11.46)
B. WTP for high alt.				
True value	\$7.65 (6.83–8.60)	\$7.12 (6.41–8.02)	\$6.45 (5.57–7.20)	\$6.42 (5.70–7.31)
True value if pure char.	\$6.42 (5.70–7.15)	\$6.48 (5.76–7.32)	\$6.42 (5.54–7.21)	\$6.42 (5.70–7.31)
Mixed Logit	\$7.67 (6.86–8.71)	\$7.05 (6.20–7.87)	\$6.43 (5.57–7.27)	\$6.40 (5.59–7.32)
GARP LP-Min	\$6.64 (5.95–7.45)	\$6.18 (5.53–7.02)	\$4.84 (4.09–5.37)	\$4.80 (4.11–5.55)
GARP LP-Max	\$6.95 (6.14–7.82)	\$7.30 (6.49–8.34)	\$8.31 (7.31–9.47)	\$8.36 (7.39–9.76)
C. WTP to replicate high alternative				
True value	\$2.82 (2.41–3.27)	\$1.40 (1.16–1.74)	\$0.06 (0.00–0.13)	\$0.00 (0–0.00)
True value if pure char.	\$0	\$0	\$0	\$0
Mixed Logit	\$2.81 (2.50–3.20)	\$1.67 (1.39–1.92)	\$0.31 (0.05–0.55)	\$0.08 (0.05–0.23)
GARP LP-Min	\$0.30 (0.26–0.35)	\$0.14 (0.12–0.18)	\$0.01 (0.00–0.02)	\$0.00 (0.00–0.00)
GARP LP-Max	\$0.31 (0.27–0.36)	\$0.16 (0.13–0.19)	\$0.01 (0.00–0.02)	\$0.00 (0.00–0.00)

This table reports results from simulations comparing the GARP approach to a mixed logit model. The columns represent data-generating processes and the rows represent “true” welfare values and estimates. See text for further explanation.