

NONPARAMETRIC ESTIMATES OF DEMAND IN THE CALIFORNIA HEALTH INSURANCE EXCHANGE

PIETRO TEBALDI

Department of Economics, Columbia University and NBER

ALEXANDER TORGOVITSKY

Kenneth C. Griffin Department of Economics, University of Chicago

HANBIN YANG

Harvard Business School

We develop a new nonparametric approach for discrete choice and use it to analyze the demand for health insurance in the California Affordable Care Act marketplace. The model allows for endogenous prices and instrumental variables, while avoiding parametric functional form assumptions about the unobserved components of utility. We use the approach to estimate bounds on the effects of changing premiums or subsidies on coverage choices, consumer surplus, and government spending on subsidies. We find that a \$10 decrease in monthly premium subsidies would cause a decline of between 1.8% and 6.7% in the proportion of subsidized adults with coverage. The reduction in total annual consumer surplus would be between \$62 and \$74 million, while the savings in yearly subsidy outlays would be between \$207 and \$602 million. We estimate the demand impacts of linking subsidies to age, finding that shifting subsidies from older to younger buyers would increase average consumer surplus, with potentially large impacts on enrollment. We also estimate the consumer surplus impact of removing the highly-subsidized plans in the Silver metal tier, where we find that a nonparametric model is consistent with a wide range of possibilities. We find that comparable mixed logit models tend to yield price sensitivity estimates toward the lower end of the nonparametric bounds, while producing consumer surplus impacts that can be both higher and lower than the nonparametric bounds depending on the specification of random coefficients.

KEYWORDS: Nonparametric estimation, demand estimation, discrete choice models, partial identification, linear programming, ACA, health insurance.

1. INTRODUCTION

1.1. *Motivation*

THE DESIGN OF THE PATIENT PROTECTION AND AFFORDABLE CARE ACT of 2010 (“ACA”)—and, more generally, of publicly-sponsored, private-provided health insurance—remains an object of debate among policy makers (e.g., [Einav and Levin \(2015\)](#), [Handel and Kolstad \(2022\)](#), [Handel and Ho \(2021\)](#)). Addressing key design issues, such as the structure of premium subsidies or the types of plans to offer, requires estimating

Pietro Tebaldi: p.tebaldi@columbia.edu

Alexander Torgovitsky: torgovitsky@uchicago.edu

Hanbin Yang: hyang1@g.harvard.edu

We thank Nikhil Agarwal, Stéphane Bonhomme, Steve Berry, Øystein Daljord, Michael Dinerstein, Liran Einav, Phil Haile, Kate Ho, Ali Hortaçsu, Simon Lee, Chuck Manski, Sanjog Misra, Magne Mogstad, Adam Rosen, Ariel Pakes, Áureo de Paula, Jack Porter, Bernard Salanié, Thomas Wollman, and participants in many conferences and seminars for helpful comments and feedback. Five anonymous referees provided comments that led to substantial improvements in the paper. The first author was supported by the Becker Friedman Institute Health Economics Initiative. The second author was supported by National Science Foundation grants SES-1530538 and SES-1846832.

demand. Recent research has filled this need using parametric discrete choice models, such as conditional logit (Chan and Gruber (2010), Ericson and Starc (2015)), nested logit (Saltzman (2019)), and mixed (random coefficient) logit (Tebaldi (2022), Polyakova and Ryan (2019)), following approaches that are common to analyses of market regulation, consumer welfare, and antitrust in a variety of contexts (see Berry and Haile (2016), for a review).

These alternative logit models differ in how they deal with the independence of irrelevant alternatives (e.g., Goldberg (1995), McFadden and Train (2000)), and in how they deal with the potential endogeneity of prices (e.g., Berry (1994), Berry, Levinsohn, and Pakes (1995)). However, they are all fully parametric, with the type I extreme value distribution playing a key role. This raises the concerning possibility that the counterfactual predictions and policy implications generated by such models are substantially driven by specific parametric functional forms.

1.2. Methodological Contribution

We develop a new nonparametric approach for the discrete choice model

$$Y_i = \arg \max_{j \in \mathcal{J}} V_{ij} - P_{ij}, \quad (1)$$

where Y_i is individual i 's choice from the discrete set of options $j \in \mathcal{J} = \{0, 1, \dots, J\}$ with prices P_{ij} and continuously distributed latent valuations V_{ij} . We do not require valuations to follow a specific distribution, such as normal (probit) or type I extreme value (logit), and we allow them to be arbitrarily dependent across options for each individual. We allow for prices and valuations to be correlated, and use instrumental variables to address this endogeneity.

We analyze identification by developing a strategy that builds on Manski (2007), who considered discrete choice under only the basic assumptions of rationality, without the aid of choice model (1). Manski's insight was that individuals facing a discrete choice can be characterized by a discrete set of latent types based on what their choices would be under any of a finite number of potential choice scenarios, such as different prices. The data provides some information on the relative frequency of the latent types, while assumptions such as rationality eliminate other types altogether. Manski showed how to use linear programming to combine these sources of information and produce bounds on choice probabilities.

The number of types in Manski's framework grows extremely quickly with the number of choice sets. This is because rationality is a weak assumption: the number of rational preference configurations compatible with even a small number of price vectors is astronomically large even when the number of choices (J) is small. We eliminate many of these types by assuming not only that individuals are rational, but also that preferences are consistent with choice model (1), which can be viewed as requiring quasilinearity. Under this choice model, each type can be interpreted as a subset of valuations, a point recognized by Koning and Ridder (2003, Figure 2). These subsets divide the space of valuations into a finite partition of valuation space that we call the minimal relevant partition (MRP).

We use the MRP to develop a practical approach for computing sharp bounds on various target parameters, such as choice probabilities, elasticities, and changes in consumer surplus. By definition, the sharp bounds cover the range of values for the target parameter that could be generated by a joint distribution of latent valuations (V_{ij}) that satisfies the researcher's a priori assumptions, and which could have also produced the observed

distribution of choices. But directly considering all possible distributions of valuations is difficult, because the distribution of a continuous random vector is an infinite-dimensional object.

Instead, we prove that one can find the sharp bounds by considering the overall masses placed on the finite number of sets in the MRP. We use a constructive argument to show that this strategy preserves sharpness even under additional assumptions that incorporate instrumental variables and known vertical orderings on the choice options. We show how to compute the sharp bounds by solving finite-dimensional linear programs. Then we apply our approach at scale, making use of recent advances in estimation and inference under partial identification.

1.3. *Empirical Results*

We use our approach to estimate demand counterfactuals in the California ACA marketplace (Covered California) using administrative records from 2014. We focus on the choice of coverage level (metal tier) for low-income individuals who are not covered under employer-sponsored insurance or public programs. In addition to estimates from our nonparametric approach, we also report estimates from comparable parametric logit, probit, and mixed logit models.

Our main empirical strategy leverages institutionally-induced variation in post-subsidy premiums across age and income, similar to [Polyakova and Ryan \(2019\)](#). We cast this as an instrumental variable strategy by dividing individuals into relatively homogenous groups by age and income, then using residual variation in age and income within that group to instrument for prices (post-subsidy premiums) while assuming that latent variations remain stable. For all of our results, we also conduct a sensitivity analysis that relaxes or drops either invariance to age or invariance to income (or both). The sensitivity analysis provides a transparent connection between the price variation in the data and the estimated results.

We estimate target parameters related to choice shares and welfare for several different counterfactuals. The main counterfactual we consider—which is simple but particularly policy-relevant—is a uniform increase in all post-subsidy premiums, which from the consumer’s perspective is the same as a decrease in premium subsidies. We also estimate substitution patterns by considering counterfactuals in which premiums are increased for different metal tiers separately. The other counterfactuals we consider concern regulatory design. In one, we consider the impacts of shifting premium subsidies toward younger or older consumers. In another, we evaluate the consumer surplus impact of removing Silver plans.

One theme that emerges throughout our results is that parametric models tend to produce price sensitivity estimates that are attenuated compared to the nonparametric bounds. For example, we estimate that a uniform \$120 increase in yearly premiums for all plans would cause between a 1.8% and 6.7% decline in the proportion of low-income adults who purchase insurance, off of a base of 28%. The parametric models, by contrast, produce point estimates that are all between 2.5% and 3.5%, suggesting that they could be substantially understating price sensitivity. Other logit-based estimates for Covered California have been similar, such as [Tebaldi \(2022\)](#) who estimates between a 1% and 2% decline in the proportion enrolled in response to a \$120 increase in yearly premiums, [Saltzman \(2019\)](#), who estimates a 3.3% decline, and [Saltzman \(2021\)](#), who estimates a 1.2% decline. The findings are consistent with the idea that the logit model has a “flat” functional form that attenuates the role of price at large or small utilities. [Ho and Pakes](#)

(2014) and [Compiani \(2022\)](#) reached similar conclusions, albeit using different methods in different empirical settings.

The differences between the nonparametric estimates and the comparable parametric models become more nuanced when considering consumer surplus. The nonparametric estimate of the consumer surplus impact of decreasing subsidies by \$120 annually is between \$62 and \$74 million annually when aggregated, with poorer individuals incurring the bulk of the loss. Total annual savings on premium subsidies would be between \$207 and \$602 million per year. The various mixed logit models—differing only by which coefficients are random—produce consumer surplus estimates between \$63 and \$84 million, with estimates of subsidy savings toward the lower end of our bounds. For a shift in subsidies toward younger buyers, the parametric models predict an average consumer surplus impact of less than \$1 per person, per month, while the nonparametric bounds include impacts that could be almost three times as large. The parametric models predict that removing Silver plans would create an aggregate consumer surplus impact of anywhere between \$149 to \$292 million per year, whereas the nonparametric model produces a lower bound of \$12 million, while recognizing that there is no logical upper bound. The results are consistent with long-standing concerns about the use of logit models for welfare analysis, especially when the counterfactual involves adding or removing a product (e.g., [Hausman \(1996\)](#), [Petrin \(2002\)](#), [Akerberg and Rysman \(2005\)](#), [Berry and Pakes \(2007\)](#)).

1.4. *Relationship to the Literature*

This paper contributes to both the methodological literature on discrete choice models and an empirical literature on the demand for subsidized health insurance.

The methodological contribution is most closely related to the previously-cited work by [Manski \(2007\)](#), who developed a linear programming procedure for bounding choice probabilities with exogenous prices under rationality. The most important difference is the quasilinear structure of choice model (1), which facilitates computation at scale and provides substantial identifying content. We also show how to use instruments when prices are endogenous, as well as how to further restrict the number of types to reflect known vertical orderings. In addition to bounding choice probabilities, we show how to bound elasticities, as well as changes in consumer surplus resulting from price changes or the removal of a choice option.

This paper is also related to [Chesher, Rosen, and Smolinski \(2013\)](#). Those authors demonstrate that endogenous prices can lead to partial identification in a class of discrete choice models that includes equation (1), even if one maintains the usual logit parametric distribution ([Chesher, Rosen, and Smolinski \(2013, Section 4.2\)](#)). Using random set theory, [Chesher, Rosen, and Smolinski \(2013\)](#) show that whether a candidate value of the model parameters belongs to the sharp identified set can be determined by checking a system of moment inequalities. The system has a finite number of moment inequalities when the number of prices is finite, as in our setting. Each moment inequality corresponds to a “core-determining set” of valuations, which we show can be a much larger collection of sets than the MRP (see Section 2.7.1).

[Chesher, Rosen, and Smolinski \(2013\)](#) implement their identification analysis by checking whether the moment inequalities hold *for every candidate value of the model parameters*. If the model is nonparametric, as in our setting, then one component of the model parameters is the distribution of valuations. Following the [Chesher, Rosen, and Smolinski \(2013\)](#) strategy would require checking whether *every* distribution of valuations satis-

fies a set of moment inequalities. This is computationally intractable: iterating over *every* distribution—even approximately, using a grid—is simply not possible because the space of distributions for continuous random variables is infinite-dimensional. Chesher, Rosen, and Smolinski (2013, p. 160) also derive other nonparametric outer (nonsharp) bounds that are straightforward to implement, but these do not exploit the assumption that choices were generated according to choice model (1). In Appendix S3 in the Online Supplementary Material (Tebaldi, Torgovitsky, and Yang (2023)), we replicate a simulation in Chesher, Rosen, and Smolinski (2013), demonstrating that our sharp nonparametric bounds lie strictly between their nonsharp nonparametric bounds and their sharp parametric bounds.

Our approach is computationally tractable because it focuses on specific scalar target parameters, such as choice probabilities, elasticities, and changes in consumer surplus. Instead of trying to iterate over the space of distributions, we check whether *there exists* a distribution that is consistent with the data and reproduces a candidate value of the target parameter. We prove that one can determine whether such a distribution exists by solving a finite system of linear equations, thereby providing the basis of our linear programming procedure. The strategy makes our nonparametric approach feasible to implement at scale, as showcased by our application.

Like both Manski (2007) and Chesher, Rosen, and Smolinski (2013), we embrace a partial identification view (Tamer (2010), Ho and Rosen (2017), Molinari (2020)). Nonparametric point identification arguments for discrete choice models often require a large amount of variation in prices (e.g., Manski (1975), Thompson (1989), Matzkin (1993), Fox and Gandhi (2016)). Yet variation in prices is limited in many applications, including ours. Nonparametric point identification arguments with endogenous prices and instrumental variables also typically use an additional “completeness” condition (Chiappori and Komunjer (2009), Berry and Haile (2014, 2022), Compiani (2022)), which in some cases is also necessary for point identification (Newey and Powell (2003), Santos (2012)).¹ Completeness can be difficult to interpret, and has been shown to be untestable (Canay, Santos, and Shaikh (2013)). Not assuming either large variation or completeness raises the possibility of partial identification, which we allow for as the leading case, although our approach does not rule out point identification.

There are a number of papers on semi and nonparametric discrete choice that parameterize the valuations in choice model (1) using a linear index of choice and/or individual characteristics, and then focus on identifying and estimating these index parameters, while treating the distribution of unobservables as a nuisance parameter. Examples include Manski (1975), Matzkin (1993), Lewbel (2000), Fox (2007), Pakes (2010), Ho and Pakes (2014), Pakes, Porter, Ho, and Ishii (2015), Pakes and Porter (2021), Shi, Shum, and Song (2018), and Khan, Ouyang, and Tamer (2021). In contrast, we do not treat the distribution of unobservables as a nuisance parameter because it determines the counterfactuals we want to infer. Other semi and nonparametric papers that focus on counterfactuals but do not allow for endogeneity and instruments include Thompson (1989), Manski (2007), Briesch, Chintagunta, and Matzkin (2010), Chiong, Hsieh, and Shum (2017), Allen and Rehbeck (2019), and Fosgerau and Kristensen (2021). Manski (2014), Kline and Tartari (2016), and Kamat (2021) use approaches similar to Manski (2007) for models different than equation (1).

¹Berry and Haile (2014, Section 5) also provide nonparametric identification results that do not require completeness and instead leverage both demand and supply under a shape restriction on the distribution of unobservables; see also Berry and Haile (2018).

Our empirical analysis contributes to a large literature on the demand for subsidized health insurance, including Chan and Gruber (2010), Krueger and Kuziemko (2013), Ericson and Starc (2015), Hackmann, Kolstad, and Kowalski (2015), Shepard (2022), DeLeire, Chappel, Finegold, and Gee (2017), Finkelstein, Hendren, and Shepard (2019), Saltzman (2019), Drake (2019), Polyakova and Ryan (2019), Jaffe and Shepard (2020), and Tebaldi (2022). The estimates on alternative subsidy schemes is related to the design of premium subsidies, as studied by Decarolis (2015), Einav, Finkelstein, and Tebaldi (2019), Polyakova and Ryan (2019), and Decarolis, Polyakova, and Ryan (2020). The effects of removing certain coverage options has been studied in different contexts by Dafny, Ho, and Varela (2013) and Marone and Sabety (2022). We contribute to the literature by providing nonparametric estimates that do not depend on functional form choices: our estimates depend only on the structure of the choice model together with instrumental variable assumptions and limited vertical orderings. These assumptions have straightforward economic interpretations.

The primary drawback of our empirical analysis is that we do not model supply, so all of our estimates must be interpreted as holding insurers' decisions fixed. This is a clear limitation, since one might expect potentially strong supply-side responses in the counterfactuals we consider. Integrating our nonparametric approach for estimating demand with a model of supply is an interesting avenue for further research.

2. METHODOLOGY

2.1. Nonparametric Discrete Choice Model

Individual i chooses option Y_i from a set $\mathcal{J} \equiv \{0, 1, \dots, J\}$ of $J + 1$ choices. Each choice j has a potentially endogenous characteristic called price, P_{ij} , which we model as discretely distributed, and collect into the vector $P_i \equiv (P_{i0}, P_{i1}, \dots, P_{iJ})$. Choice $j = 0$ represents the outside option of not choosing any of the inside choices $j \geq 1$, and has its price normalized to $P_{i0} = 0$. When we apply the model to Covered California, we will have five choices ($J = 4$) with options 1, 2, 3, and 4 representing Bronze, Silver, Gold, and Platinum plans.

Individual i has a vector $V_i \equiv (V_{i0}, V_{i1}, \dots, V_{iJ})$ of valuations, one for each choice, with the standard normalization that $V_{i0} = 0$. The valuations are known to the individual, but latent from the perspective of the researcher. We assume that individual i 's indirect utility from choosing j is given by $V_{ij} - P_{ij}$, so that their choice is given by

$$Y_i = \arg \max_{j \in \mathcal{J}} V_{ij} - P_{ij}. \quad (1)$$

We do not assume that the distribution of V_i follows a specific functional form such as type I extreme value (logit) or multivariate normal (probit). We also allow V_{ij} and V_{ik} to be dependent for $j \neq k$.

The main economic restriction in equation (1) is the additive separability between valuations and prices, which imposes restrictions on substitution patterns consistent with quasilinear utility. If all prices were to increase by the same amount, then an individual who chooses $j \geq 1$ before the increase will either continue to choose j after the increase, or will switch to the outside option ($j = 0$), but they will not switch to a different inside choice $k \geq 1$, $k \neq j$. This limits the role of income effects to the extensive margin of choosing any inside choice versus taking the outside option. In Appendix S1, we derive equation (1) from an insurance choice model similar to the ones in Handel (2013) and Handel, Hendel, and Whinston (2015), in which individuals have quasilinear utility and

constant absolute risk aversion. Variation in V_i across individuals arises from heterogeneity in unobserved preferences, risk factors, and risk aversion.

The limited role of income effects in equation (1) pertain to conjectured variations for a single individual i . When we take the model to the data we combine observations on many individuals, so in practice we can allow for income effects by allowing for dependence between an individual's income and their valuations. We formalize this by treating an individual's income and other observed characteristics as part of a vector, X_i , and then restricting the dependence between V_i and the various components of X_i . Throughout the paper, we model X_i as discretely distributed with finite support.

One observable characteristic that is particularly important is the individual's market. When we estimate demand we will do so conditional on a market, so that market-level characteristics—both observable and unobservables—are held fixed in the counterfactual (see, e.g., [Berry and Haile \(2022\)](#), p. 8). To emphasize this, we let M_i denote individual i 's market, and we treat M_i as separate from X_i in the notation.

Our empirical setting is different than many discrete choice applications in which prices only vary at the market level, such as [Berry, Levinsohn, and Pakes \(1995\)](#) or [Nevo \(2001\)](#). These settings would have P_i constant conditional on M_i . The approach we develop in the main text is not immediately useful for this case. In Appendix S2, we propose two ways in which our approach could be extended to handle more aggregated price variation.

2.2. Comparison With a Common Parametric Model

A common parametric specification for discrete choice demand models is

$$Y_{im} = \arg \max_{j \in \mathcal{J}} X'_{ijm} \beta_{im} - \alpha_{im} P_{ijm} + \xi_{jm} + \epsilon_{ijm}, \quad (2)$$

where i , j , and m index individuals, products, and markets, P_{ijm} is price, X_{ijm} are observed characteristics, ξ_{jm} are unobserved product-market characteristics, β_{im} and α_{im} are individual-level random coefficients, and ϵ_{ijm} are idiosyncratic unobservables. See, for example, equation (6) of [Nevo \(2011\)](#), or equation (1) of [Berry and Haile \(2016\)](#). In the influential model of [Berry, Levinsohn, and Pakes \(1995\)](#), ϵ_{ijm} are assumed to be i.i.d. logit (type I extreme value), and $(\beta_{im}, \alpha_{im})$ are assumed to be normally (or log-normally) distributed. Our motivation for considering choice equation (1) is to preserve the utility maximization structure in equation (2), while avoiding these types of parametric assumptions.

The three indices in equation (2) reflect different possible levels of data aggregation. If only market-level data is available, as in [Berry, Levinsohn, and Pakes \(1995\)](#) or [Nevo \(2001\)](#), then equation (2) is aggregated to the (j, m) level, and the data is viewed as drawn from a population of markets and/or products ([Berry, Linton, and Pakes \(2004\)](#), [Armstrong \(2016\)](#)). Our analysis presumes richer individual-level choice data as in [Berry, Levinsohn, and Pakes \(2004\)](#) or [Berry and Haile \(2022\)](#), but the number of markets we study is small and fixed. To emphasize this, we index equation (1) only over i and j , and we record individual i 's market using the random variable M_i . After subsuming m subscripts into i subscripts, equation (1) can be seen to nest equation (2) by dividing through by α_i and taking $V_{ij} \equiv \alpha_i^{-1}(X'_{ij}\beta_i + \xi_{ij} + \epsilon_{ij})$.²

²This requires the mild assumption that $\alpha_i > 0$. As suggested by a referee, an advantage of normalizing in this way is that V_{ij} can be directly interpreted as willingness-to-pay relative to the outside option (see, e.g., [Lewbel \(2019\)](#), Section 6.3).

2.3. The Density of Valuations

The primitive object in choice model (1) is the distribution of valuations, V_i , conditional on prices, P_i , market, M_i , and other covariates, X_i . We assume throughout the paper that this distribution is continuous so that ties between choices occur with zero probability. More formally, we assume that the distribution of V_i is absolutely continuously distributed with respect to Lebesgue measure on $\mathbb{R}^J$, conditional on $(P_i, M_i, X_i) = (p, m, x)$ for every (p, m, x) in the support of (P_i, M_i, X_i) . Absolute continuity means we can associate the conditional distribution of valuations with a conditional density function $f(\cdot|p, m, x)$ for each realization $P_i = p$, $M_i = m$, and $X_i = x$. Let $\mathcal{F}$ denote the set of all such conditional density functions.

2.4. Target Parameters

Common counterfactual quantities of interest can be written as integrals or sums of integrals of f ; see, for example, Section 4.2 of [Berry and Haile \(2014\)](#) or Section 3.4.1 of [Berry and Haile \(2016\)](#). For example, a natural counterfactual quantity is the proportion of individuals who would choose j at a new price vector, p^* . The proportion can be written in terms of f as

$$\int \underbrace{\mathbb{1}[v_j - p_j^* \geq v_k - p_k^* \text{ for all } k]}_{\text{choose } j \text{ if prices were } p^*} f(v|m, x) dv, \quad (3)$$

where $f(v|m, x) \equiv \mathbb{E}[f(v|P_i, m, x)|M_i = m, X_i = x]$

is the density of valuations conditional on market, m and other characteristics x , averaged over the remaining variation in observed prices (if any). Another natural counterfactual quantity is the impact on average consumer surplus caused by changing prices from p to p^* . The impact can be written as

$$\underbrace{\int \left\{ \max_{j \in \mathcal{J}} v_j - p_j^* \right\} f(v|m, x) dv}_{\text{consumer surplus under } p^*} - \underbrace{\int \left\{ \max_{j \in \mathcal{J}} v_j - p_j \right\} f(v|m, x) dv}_{\text{consumer surplus under } p}, \quad (4)$$

where again the market, m , is being held fixed in the counterfactual.

We view equations (3) and (4) as examples of *target parameters*; that is, functions $\theta : \mathcal{F} \rightarrow \mathbb{R}^{d_\theta}$ that map the collection $\mathcal{F}$ of all conditional density functions on $\mathbb{R}^J$ into real vectors. The goal is to infer the values of $\theta(f)$ that are consistent with both the observed data and the assumptions.

2.5. Assumptions

We augment the choice model (1) with two types of assumptions.

2.5.1. Instrumental Variables

Let W_i and Z_i be two subvectors (or more general functions) of prices, market and covariates, (P_i, M_i, X_i) . Define the density of valuations conditional on W_i and Z_i as

$$f(v|w, z) \equiv \mathbb{E}[f(v|P_i, M_i, X_i)|W_i = w, Z_i = z]. \quad (5)$$

ASSUMPTION IV: $f(v|w, z) = f(v|w, z')$ for all z, z', w , and v .

Assumption IV says that the distribution of valuations is invariant to Z_i , conditional on W_i . That is, Z_i are exogenous instruments and W_i are control variables. Assumption IV nests exogenous prices as a special case by taking $Z_i = P_i$ and W_i to be deterministic (e.g., $W_i = 1$). It also nests the unconditional instrumental variable assumption used by Chesher, Rosen, and Smolinski (2013, Restriction A.6) by taking W_i to be deterministic.

For Assumption IV to be useful, prices should shift with the instruments Z_i while conditioning on the controls W_i . This follows the usual intuition: if Z_i is exogenous, then changes in observed choice shares as Z_i varies reflect changes in prices, rather than changes in valuations. The more that prices vary with Z_i , the more information we have to pin down different parts of the density of valuations, f , and, therefore, the target parameter, θ . Our approach does not require the instruments to have any particular amount of variation, but greater variation produces more informative bounds.

In the application, we also use the following generalization of Assumption IV, which allows the instruments to be “imperfect” in the terminology of Nevo and Rosen (2012).

ASSUMPTION IV—Generalized: For every z, z', w , and v ,

$$(1 - \kappa(z, z', w))f(v|w, z') \leq f(v|w, z) \leq (1 + \kappa(z, z', w))f(v|w, z'),$$

where $\kappa(z, z', w) \geq 0$ is chosen by the researcher.

The generalization allows for a pointwise difference in the density of valuations between z and z' of up to $(100 \times \kappa(z, z', w))\%$. Similar approaches to relaxing assumptions involving exact equalities have also been used in different contexts by Conley, Hansen, and Rossi (2010), Manski and Pepper (2017), Torgovitsky (2019), and Rambachan and Roth (2022). Taking $\kappa(z, z', w) = 0$ returns the generalized assumption back to the baseline case of an exogenous instrument.

2.5.2. Support

The second assumption is that the support of f is contained in a known subset of $\mathbb{R}^J$.

ASSUMPTION SP: $\int_{\mathcal{V}^\bullet(w)} f(v|w, z) dv = 1$ for each w and z , where $\mathcal{V}^\bullet(w)$ is a known subset of $\mathbb{R}^J$.

In the application, we use Assumption SP to exploit the fact that certain plans in the ACA are unambiguously more attractive at equal prices than certain other plans. Assumption SP can be made nonrestrictive by taking $\mathcal{V}^\bullet(w) = \mathbb{R}^J$ for all w .

2.6. The Identified Set

Suppose that we know the observed choice shares

$$s_j(p, m, x) \equiv \mathbb{P}[Y_i = j | P_i = p, M_i = m, X_i = x]. \quad (6)$$

Each density of valuations $f \in \mathcal{F}$ also implies a set of choice shares

$$s_j(p, m, x; f) \equiv \int_{\mathcal{V}_j(p)} f(v|p, m, x) dv, \quad (7)$$

$$\text{where } \mathcal{V}_j(p) \equiv \{(v_1, \dots, v_J) \in \mathbb{R}^J : v_j - p_j \geq v_k - p_k \text{ for all } k\}. \quad (8)$$

A density f is consistent with the observed choice shares (“matches the data”) if

$$s_j(p, m, x; f) = s_j(p, m, x) \quad \text{for all } j, p, m, \text{ and } x. \quad (\text{MD})$$

The sharp identified set of valuation densities is defined as the subset of $f \in \mathcal{F}$ that both match the data and satisfy Assumptions IV and SP. We call this subset $\mathcal{F}^*$:

$$\mathcal{F}^* \equiv \{f \in \mathcal{F} : f \text{ satisfies Assumptions IV, SP, and equation (MD)}\}. \quad (9)$$

The object of interest is the sharp identified set for the target parameter, which is defined as the image of $\mathcal{F}^*$ under θ , and denoted as

$$\Theta^* \equiv \{\theta(f) : f \in \mathcal{F}^*\} \equiv \theta(\mathcal{F}^*).$$

2.7. Identification Analysis

In this section, we show how to compute Θ^* . If $\mathcal{F}$ were a parametric class of densities, then the results of Chesher, Rosen, and Smolinski (2013) could be used to approximately compute Θ^* by iterating over a grid of $f \in \mathcal{F}$ that satisfy Assumptions IV and SP, checking whether each such f satisfies a finite set of moment inequalities, and adding f to $\mathcal{F}^*$ and $\theta(f)$ to Θ^* if and only if they are satisfied. In our model, f is nonparametric, so $\mathcal{F}$ is an infinite-dimensional set. Iterating over even a grid of $f \in \mathcal{F}$ is not feasible in this case.

We develop a different strategy. Instead of trying to check a set of conditions for every $f \in \mathcal{F}$, we check whether there exists *some* $f \in \mathcal{F}^*$ that would generate a given candidate value $t = \theta(f)$ for the target parameter. As we prove ahead in Proposition 1, the existence of such an f is equivalent to the existence of a solution to a carefully-constructed finite-dimensional system of linear equations. The finite-dimensional system is constructed by first partitioning the space of valuations into the smallest number of sets needed to describe choice behavior under all relevant prices. We call this partition the minimal relevant partition (MRP).

2.7.1. The Minimal Relevant Partition

Let $\mathcal{P}$ denote a finite set of prices that contains the observed prices, as well as any additional prices relevant for evaluating the target parameter, θ . The MRP is the smallest (coarsest) partition of valuation space ($\mathbb{R}^J$) with the following property: any two valuations in the same subset produce the same choice behavior for all $p \in \mathcal{P}$, while any two valuations in different subsets produce different choice behavior for at least one $p \in \mathcal{P}$. The formal definition is as follows.

DEFINITION MRP: Let $Y(v, p) \equiv \arg \max_{j \in \mathcal{J}} v_j - p_j$ for any $(v_1, \dots, v_J), (p_1, \dots, p_J) \in \mathbb{R}^J$, where $v \equiv (v_0, v_1, \dots, v_J)$ and $p \equiv (p_0, p_1, \dots, p_J)$ with $v_0 = p_0 = 0$. The *minimal relevant partition of valuations* (MRP) is a collection $\mathbb{V}$ of sets $\mathcal{V} \subseteq \mathbb{R}^J$ for which the following property holds for almost every $v, v' \in \mathbb{R}^J$ (with respect to Lebesgue measure):

$$v, v' \in \mathcal{V} \text{ for some } \mathcal{V} \in \mathbb{V} \quad \Leftrightarrow \quad Y(v, p) = Y(v', p) \quad \text{for all } p \in \mathcal{P}. \quad (10)$$

Suppose that the data consists of a single observed price vector, p^a , and that we are concerned with behavior under a counterfactual price vector, p^* . Take $\mathcal{P} = \{p^a, p^*\}$. We illustrate the MRP with $J = 2$ in Figure 1. Panel (a) shows that considering choice behavior under price p^a divides $\mathbb{R}^2$ into three sets depending on whether an individual would

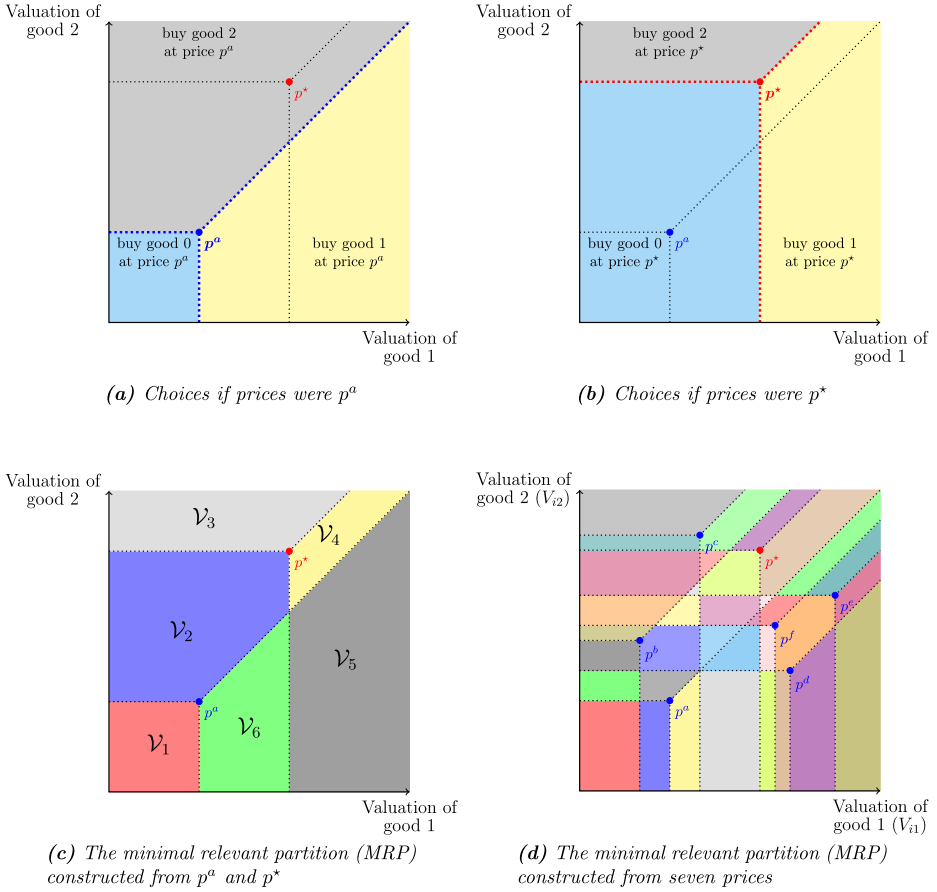


FIGURE 1.—Partitioning the space of valuations. *Notes:* Panels (a) and (b) show the partition of valuation space created by prices p^a and p^* , respectively, for the case with $J = 2$. Panel (c) shows how intersecting these two partitions yields a partition of six sets (the MRP). Panel (d) shows a richer MRP based on the prices $\mathcal{P} = \{p^a, p^b, p^c, p^d, p^e, p^f, p^*\}$.

choose options 0, 1, or 2 when faced with p^a . Panel (b) shows the analogous division under price p^* . See, for example, Thompson (1989, Figure 1), Chesher, Rosen, and Smolinski (2013, Figure 1), or Berry and Haile (2014, Figure 1) for similar diagrams. Intersecting these two three-set collections creates the collection of six sets shown in panel (c). The collection of six sets is the MRP generated by $\mathcal{P} = \{p^a, p^*\}$. In Figure 1d, we show the MRP constructed from a set $\mathcal{P}$ with seven prices. We describe an algorithm for constructing the MRP in Appendix S4.

The MRP is the coarsest division of $\mathbb{R}^J$ that captures all choice behavior under the prices in $\mathcal{P}$. The six sets in Figure 1c correspond to the six types of individuals that could exist under choice model (1), where a type is defined as a pair of choices under (p^a, p^*) . The characterization by types follows Manski (2007), but here we adapt it to the quasilinear choice model (1), as in Koning and Ridder (2003).

The MRP is related to the class of core-determining sets (CDS) derived by Chesher, Rosen, and Smolinski (2013) and Chesher and Rosen (2014), but there are fewer sets in the MRP than the CDS. With two prices and $J = 2$, the CDS consists of 12 sets (Chesher, Rosen, and Smolinski (2013, Figures 2–3)) compared to 6 sets in the MRP. With seven

prices and $J = 2$, the CDS can consist of as many as 842 sets (Chesher, Rosen, and Smolinski (2011, Table 1, p. 28)) compared to 35 sets in the MRP in Figure 1d. The MRPs in our application have between 195,000 and 900,000 sets, depending on the region and target parameter.

2.7.2. Equivalent Finite-Dimensional Problem

The MRP is formed by organizing the continuum of valuations in $\mathbb{R}^J$ into sets that yield identical choice behavior at all prices in $\mathcal{P}$. If our concern is behavior that occurs (or would occur) under these prices, then we can shift our attention from densities in $\mathbb{R}^J$ to mass functions over the sets in the MRP. For example, in Figure 1c the share of individuals who would choose 1 if prices were p^a can be written as

$$s_1(p^a, m, x; f) = \int_{\mathcal{V}_5} f(v|p^a, m, x) dv + \int_{\mathcal{V}_6} f(v|p^a, m, x) dv,$$

which depends not on f per se, but only on the mass that f places on $\mathcal{V}_5$ and $\mathcal{V}_6$. Focusing on these masses replaces the intractable infinite-dimensional problem of searching over all $f \in \mathcal{F}$ into a tractable finite-dimensional problem of searching over a finite set of non-negative numbers that sum to unity. As we now show, the replacement can be done while preserving all of the information provided by the choice model, the data, and Assumptions IV and SP.

We denote the set of conditional mass functions supported on $\mathbb{V}$ as

$$\Phi \equiv \left\{ \phi \in \mathbb{R}_+^{d_\phi} : \sum_{\mathcal{V} \in \mathbb{V}} \phi(\mathcal{V}|p, m, x) = 1 \text{ for all } (p, m, x) \in \text{supp}(P_i, M_i, X_i) \right\}, \quad (11)$$

where d_ϕ is the cardinality of $\mathbb{V} \times \text{supp}(P_i, M_i, X_i)$. Let $\mathbb{V}_j(p) \equiv \{\mathcal{V} \in \mathbb{V} : Y(v, p) = j \text{ for almost every } v \in \mathcal{V}\}$ denote the sets in the MRP associated with choosing j under price $p \in \mathcal{P}$. Then $\phi \in \Phi$ is consistent with an f that satisfies equation (MD) only if

$$\sum_{\mathcal{V} \in \mathbb{V}_j(p)} \phi(\mathcal{V}|p, m, x) = s_j(p, m, x) \quad \text{for all } j \in \mathcal{J} \text{ and } (p, m, x). \quad (\text{MD}')$$

Each ϕ generates the conditional-on- (W_i, Z_i) probability mass function

$$\phi(\mathcal{V}|w, z) \equiv \mathbb{E}[\phi(\mathcal{V}|P_i, M_i, X_i)|W_i = w, Z_i = z].$$

To be consistent with an f that satisfies Assumption IV, ϕ should satisfy

$$(1 - \kappa(z, z', w))\phi(\mathcal{V}|w, z') \leq \phi(\mathcal{V}|w, z) \leq (1 + \kappa(z, z', w))\phi(\mathcal{V}|w, z') \\ \text{for all } z, z', w, \text{ and } \mathcal{V}. \quad (\text{IV}')$$

To be consistent with an f that satisfies Assumption SP, ϕ should satisfy

$$\sum_{\mathcal{V} \in \mathbb{V}^\bullet(w)} \phi(\mathcal{V}|w, z) = 1 \quad \text{for all } w \text{ and } z, \quad (\text{SP}')$$

where $\mathbb{V}^\bullet(w)$ is the subset of $\mathbb{V}$ that intersects $\mathcal{V}^\bullet(w)$, that is, $\mathbb{V}^\bullet(w) \equiv \{\mathcal{V} \in \mathbb{V} : \lambda(\mathcal{V} \cap \mathcal{V}^\bullet(w)) > 0\}$, with λ denoting Lebesgue measure on $\mathbb{R}^J$.

Since our focus is not on f but on the target parameter, θ , we assume that it is possible to represent θ in terms of ϕ .

CONDITION TP: For any $f \in \mathcal{F}$, define $\bar{\phi}(f) \in \Phi$ as

$$\bar{\phi}(f)(\mathcal{V}|p, m, x) \equiv \int_{\mathcal{V}} f(v|p, m, x) dv \quad \text{for all } \mathcal{V} \in \mathbb{V}. \quad (12)$$

Then there exists a known function $\bar{\theta}: \Phi \rightarrow \mathbb{R}^{d_\theta}$ such that

$$\theta(f) = \bar{\theta}(\bar{\phi}(f)) \quad \text{for every } f \in \mathcal{F}.$$

Since Φ depends on the MRP, and the MRP depends on $\mathcal{P}$, Condition TP can always be satisfied by choosing $\mathcal{P}$ to include all prices relevant for evaluating the target parameter, θ . For example, if the target parameter is demand at the new premium vector p^* in Figure 1c, then Condition TP simply requires choosing $\mathcal{P}$ to include p^* .

Under Condition TP, Θ^* can be characterized in terms of the mass function ϕ without having to consider the full density f . If $t \in \Theta^*$, then (by definition) there exists an $f \in \mathcal{F}^*$ with $\theta(f) = t$. Since $f \in \mathcal{F}^*$ satisfies equation (MD), Assumption IV, and Assumption SP (again all by definition), taking the mass that any such f places on the MRP yields a $\phi \in \Phi$ that satisfies equations (MD'), (IV'), (SP'), and $\bar{\theta}(\phi) = t$; for example, take $\bar{\phi}(f)$ defined in equation (12). The opposite direction is more delicate: if $\bar{\theta}(\phi) = t$ for some $\phi \in \Phi$ that satisfies equations (MD'), (IV'), and (SP'), then is $t \in \Theta^*$? The following proposition shows that the answer is yes. The idea is to use such a ϕ to construct an $f \in \mathcal{F}^*$ that yields $\theta(f) = t$. The proof is in the Appendix.

PROPOSITION 1: Suppose that Condition TP is satisfied. Let

$$\Phi^* \equiv \{\phi \in \Phi : \phi \text{ satisfies equations (MD'), (IV'), and (SP')}\}, \quad (13)$$

and for any $t \in \mathbb{R}^{d_\theta}$, let $\Phi^*(t) \equiv \Phi^* \cap \{\phi \in \Phi : \bar{\theta}(\phi) = t\}$. Then $t \in \Theta^*$ if and only if $\Phi^*(t)$ is nonempty.

Proposition 1 shows that Θ^* is characterized by finite-dimensional systems of equations. If the target parameter is scalar, then Θ^* can be computed more directly by solving two optimization problems.

PROPOSITION 2: If $\bar{\theta}$ is continuous on Φ , then Θ^* is a compact, connected set. In particular, if $d_\theta = 1$, then $\Theta^* = [t_{lb}^*, t_{ub}^*]$, where

$$t_{lb}^* \equiv \min_{\phi \in \Phi^*} \bar{\theta}(\phi) \quad \text{and} \quad t_{ub}^* \equiv \max_{\phi \in \Phi^*} \bar{\theta}(\phi). \quad (14)$$

Each condition that defines Φ^* —equations (MD'), (IV'), and (SP')—is linear in ϕ . If $\bar{\theta}$ is also linear in ϕ , then problem (14) is a linear program. A change in choice shares, like equation (3), yields a $\bar{\theta}$ function that is linear in ϕ . A change in consumer surplus, like equation (4), also does, although constructing $\bar{\theta}$ is less obvious (see Appendix S5). A discrete approximation to an elasticity produces a $\bar{\theta}$ function that is nonlinear, but the resulting optimization problem is linear-fractional, so can be reformulated as a linear program using the Charnes and Cooper (1962) transformation; see Appendix S6 for more

detail, and [Kamat \(2021\)](#) for a similar observation. Linearity ensures that both the minimization and maximization problems can be solved reliably and relatively quickly, even when the dimension of Φ is quite large.³

2.8. Estimation and Statistical Inference

The choice shares $s_j(p, m, x)$ are in practice replaced by estimates $\hat{s}_j(p, m, x)$. Using the estimated choice shares in equation (MD') can lead to empty feasible sets in problem (14), even when the model is correctly specified so that the feasible sets are nonempty with the true choice shares. We solve this problem by applying an estimator of $[t_{lb}^*, t_{ub}^*]$ developed by [Mogstad, Santos, and Torgovitsky \(2018\)](#). The estimator has two steps.

In the first step, we find the best fit to the estimated choice shares by solving

$$\hat{Q}^* \equiv \min_{\phi \in \Phi} \hat{Q}(\phi) \text{ subject to equations (IV') and (SP'), where} \quad (15)$$

$$\hat{Q}(\phi) \equiv \sum_{j, p, m, x} \hat{\mathbb{P}}[P_i = p, M_i = m, X_i = x] \left| \hat{s}_j(p, m, x) - \sum_{\nu \in \mathbb{V}_j(p)} \phi(\nu | p, m, x) \right|,$$

and where $\hat{\mathbb{P}}[P_i = p, M_i = m, X_i = x]$ are estimated probabilities. Defining $\hat{Q}$ with absolute deviations means that problem (15) can be reformulated as a linear program by replacing terms in absolute values by the sum of their positive and negative parts (e.g., [Bertsimas and Tsitsiklis \(1997\)](#), pp. 19–20). The absolute deviations are weighted so that $\hat{Q}(\phi)$ is an estimate of the average absolute deviation in choice shares under ϕ .

In the second step, we collect values of $\bar{\theta}(\phi)$ for ϕ that come close to minimizing problem (15). That is, we construct the set

$$\hat{\Theta}^* \equiv \{\bar{\theta}(\phi) : \phi \in \Phi, \phi \text{ satisfies equations (IV'), (SP'), and } \hat{Q}(\phi) \leq (1 + \eta)\hat{Q}^*\},$$

which is never empty due to the definition of $\hat{Q}^*$. The qualifier “close” here reflects the tuning parameter η , which is a small positive constant that must converge to zero at an appropriate rate with the sample size to smooth out potential discontinuities that could conceivably arise in the convergence of a set estimator. In our empirical estimates, we set $\eta = 10^{-4}$ and found little sensitivity to values of η that were bigger or smaller by an order of magnitude. However, there are currently no theoretical results to guide the choice of this parameter.

We compute $\hat{\Theta}^*$ by solving

$$\hat{t}_{lb}^* \equiv \min_{\phi \in \Phi} \bar{\theta}(\phi) \text{ subject to equations (IV'), (SP'), and } \hat{Q}(\phi) \leq \hat{Q}^*(1 + \eta), \quad (16)$$

and an analogous maximization problem for $\hat{t}_{ub}^*$. The set estimator for Θ^* is $\hat{\Theta}^* \equiv [\hat{t}_{lb}^*, \hat{t}_{ub}^*]$, which [Mogstad, Santos, and Torgovitsky \(2018, Section S3, Theorem 1\)](#) show is consistent for $[t_{lb}^*, t_{ub}^*]$ under fairly weak conditions. When $\bar{\theta}$ is linear, problem (16) can be reformulated as a linear program, again by appropriately reformulating the absolute value terms in $\hat{Q}(\phi)$. The overall estimation procedure requires solving three linear programs: problem (15), problem (16), and the maximization problem analogous to (16).

³All of the linear (and quadratic) programs ahead were computed using Gurobi ([Gurobi Optimization \(2015\)](#)), and we checked a subset of the results using CPLEX ([IBM \(2010\)](#)).

We conduct statistical inference using the test developed by [Deb, Kitamura, Quah, and Stoye \(2021\)](#); see also [Kitamura and Stoye \(2018\)](#). The specifics of the test are notationally involved, so here we provide a high-level overview of what it entails. Appendix S7 contains a detailed discussion of how we implement the test in our application.

The null hypothesis of the test is $H_0 : t \in \Theta^*$, where t is a candidate value of the target parameter. The test statistic is the minimum weighted sum of squares between the estimated choice shares and the shares predicted by a $\phi \in \Phi$ that satisfies equations (IV'), (SP'), and $\bar{\theta}(\phi) = t$, where $\bar{\theta}$ is required to be linear in ϕ . The test statistic is thus the optimal value of a quadratic program. A critical value is found by bootstrapping the choice shares and, for each bootstrap replication, solving a tightened version of the same quadratic program. The tightening effectively provides the type of generalized moment selection used in the moment inequality literature (e.g., [Andrews and Soares \(2010\)](#)). The null hypothesis is rejected at the 5% level if the test statistic exceeds the 0.95 bootstrapped quantile of the tightened problem.

We construct a 95% confidence interval by inverting 5% tests. That is, the confidence interval contains all target parameter candidate values t for which the null hypothesis $H_0 : t \in \Theta^*$ is not rejected. The confidence interval can be constructed relatively efficiently by bisecting its endpoints. In our application, the overall procedure of building confidence intervals is still computationally demanding, since it requires solving an extremely large quadratic program for each bootstrap replication and each trial value of t . For this reason, and because our sample size is quite large, we report confidence intervals for only a few parameters.

3. COVERED CALIFORNIA

3.1. *Institutional Details*

Covered California is one of the largest state-level ACA marketplaces, accounting for more than 10% of national enrollment. The marketplace offers health insurance plans directed at individuals not covered by an employer or by a public program such as Medicaid or Medicare.

The basic structure of Covered California is determined by federal regulation common to ACA marketplaces in all states. The regulation splits states into geographic rating regions comprised of groups of contiguous counties or zip codes. In California, there are 19 rating regions. Insurers are allowed to vary premiums across (but not within) rating regions, and consumers face the premiums set for their resident region. Each year in the spring, insurers announce their intention to enter a region in the subsequent calendar year, then undergo state certification. Consumers purchase insurance for the subsequent year during an open enrollment period at the end of the year.

Covered California differs from other ACA marketplaces in important ways. An insurer who intends to participate in a rating region is required to offer a menu of four plans classified into metal tiers of increasing actuarial value: Bronze, Silver, Gold, and Platinum.⁴ Unlike other marketplaces, where insurers do not need to offer Bronze or Platinum plans, in California an insurer must provide the entire menu of four plans in any region it enters. Moreover, the actuarial features of the plans are required to have the standardized characteristics shown in Table I (among others not shown).

⁴There is a fifth coverage tier called minimum (or catastrophic) coverage. This tier is not available to the subsidized buyers we focus on (with a few, rare exceptions), so we omit it from the analysis.

TABLE I
STANDARDIZED PLAN CHARACTERISTICS IN COVERED CALIFORNIA FOR 2014.

Tier	Panel (a): Characteristics by metal tier before cost-sharing reductions (CSRs)						
	Annual deductible	Annual max out-of-pocket	Primary visit	E.R. visit	Specialist visit	Preferred drugs	Advertised AV ^(*)
Bronze	\$5000	\$6250	\$60	\$300	\$70	\$50	60%
Silver	\$2250	\$6250	\$45	\$250	\$65	\$50	70%
Gold	\$0	\$6250	\$30	\$250	\$50	\$50	79%
Platinum	\$0	\$4000	\$20	\$150	\$40	\$15	90%

Income (%FPL)	Panel (b): Silver plan characteristics after cost-sharing reductions (CSRs)						
	Annual deductible	Annual max out-of-pocket	Primary visit	E.R. visit	Specialist visit	Preferred drugs	Advertised AV ^(*)
200–250% FPL	\$1850	\$5200	\$40	\$250	\$50	\$35	74%
150–200% FPL	\$550	\$2250	\$15	\$75	\$20	\$15	88%
100–150% FPL	\$0	\$2250	\$3	\$25	\$5	\$5	95%

Note: Source: <http://www.coveredca.com/PDFs/2015-Health-Benefits-Table.pdf>.

Insurers are also regulated on how they can set premiums. Each insurer chooses a base premium for each metal tier in each rating region. The base premium is multiplied by a federally-determined adjustment factor that increases with a consumer’s age from 1 at age 21 to 3 at age 64 (see Orsini and Tebaldi (2017), for further detail). The insurer is not permitted to adjust premiums based on any nonage characteristics of the consumer. Insurer premiums are thus a deterministic function of a consumer’s age and resident rating region.

Individuals with household income below 400% of the Federal Poverty Level (FPL) receive premium subsidies. The size of the premium subsidy is set so that the subsidized premium of the second-cheapest Silver plan is lower than a so-called “maximum affordable amount” that varies with income. Post-subsidy premiums are thus a deterministic function of a consumer’s age, resident rating region, and household income.

In Figure 2, we illustrate how post-subsidy premiums vary as a function of both age and income. Holding age fixed, post-subsidy premiums increase with income equally across all plans due to lower subsidies. Holding income fixed, post-subsidy premiums increase with age differently across plans due to the adjustment factor.

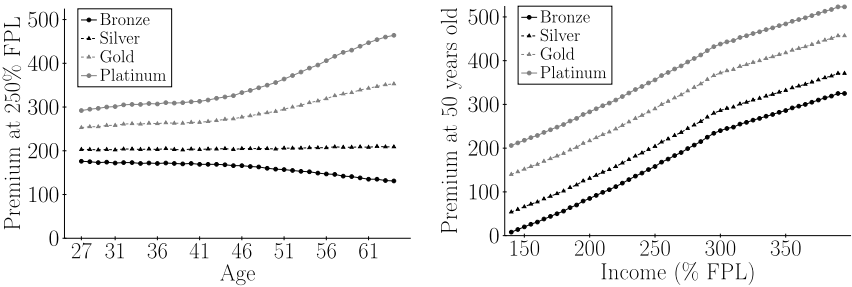


FIGURE 2.—Post-subsidy premium variation by age and income. Notes: Post-subsidy premiums shown are the median across insurers for rating region 16 (part of Los Angeles).

In addition to premium subsidies, the ACA also provided cost-sharing reductions (CSRs) for individuals with household income lower than 250% of the FPL. In Covered California, CSRs were implemented by changing the actuarial terms of Silver plans by income, with discrete changes at 150%, 200%, and 250% of the FPL; see Table I.

The ACA had a universal coverage mandate that specified an income tax penalty for remaining uninsured. We treat the tax penalty as affecting the value of the outside option of not purchasing any Covered California plan. The universal mandate was generally unenforced between 2014–2017 (Miller (2017)).

3.2. Data

Our primary data are administrative records on all individuals who purchased a plan through Covered California in 2014. The records contain unique individual and household identifiers, as well as age, income measured in percentage of the FPL, gender, zip-code of residence, choice of plan, and premium paid. Since post-subsidy premiums are a deterministic function of demographics, we can also calculate premiums for plans a consumer did not choose. We focus on adults aged 27–64 with household income between 140% and 400% of the FPL, which is 73% of the roughly 1.3 million enrollees. The remaining 27% of enrollees are either ineligible for subsidies (11%), or are younger than 26, so considered dependents under the ACA (16%).

We characterize each individual i by their resident rating region (market), M_i , and a vector X_i consisting of their age and household income. We discretize age into 38 single-year bins running from 27 to 64, and household income into 52 FPL bins that are 5% wide, running from 140% to 400%. When crossed with the 19 rating regions in Covered California, we obtain 37,544 unique rating region $\times$ age $\times$ income bins of the observable characteristics, (M_i, X_i) .

As in most demand analyses, we do not directly observe individuals who chose the outside option of not purchasing a plan through Covered California. To transform quantities purchased into shares, we estimate the number of potential buyers using data from the 2011–2013 American Community Survey public use file (via IPUMS, Ruggles, Genadek, Goeken, Grover, and Sobek, 2015). Our potential buyer estimates use a flexible linear regression, similar to Finkelstein, Hendren, and Shepard (2019) and Tebaldi (2022). More detail is provided in Appendix S8.

We combine potential buyer estimates with the administrative data to construct choice shares for each of the region $\times$ age $\times$ income bins. To avoid excessive extrapolation, we drop bins that are empty in the ACS sample. We also drop a few small bins for which we estimate fewer potential buyers than there are enrollees in the administrative data. We are left with 30,007 bins that we use as the main estimation sample. Since the number of individuals per bin varies greatly, we will report parameters that average over (M_i, X_i) and, therefore, put greater weight on larger bins.

We focus on an individual's choice of coverage level (metal tier), so that $J = 4$, with $j = 1, 2, 3, 4$ denoting Bronze, Silver, Gold, and Platinum, respectively, and $j = 0$ denoting the outside option, as usual. The implicit assumption is that the choice of coverage level is separable from the choice of insurer. We view this assumption as reasonable for Covered California because the regulations ensure that the metal tiers offered—as well as the financial characteristics of the tiers—do not vary by insurer. We define premiums for each tier in each bin as the median post-subsidy premium across insurers. Because of ACA regulations, $P_i = \pi(M_i, X_i)$ is a deterministic function of a consumer's region, M_i , and demographic characteristics, X_i . We reflect this in notation ahead by writing $f(v|p, m, x) = f(v|m, x)$.

TABLE II
SUMMARY STATISTICS.

Panel (a): Distribution of bin characteristics					
	Mean	St. Dev.	P-10	Median	P-90
Number of buyers	85.32	91.07	14	55	195
Age	43.415	10.694	29	43	59
Income (FPL%)	243.991	72.037	155	230	355
Takeup rate	0.280	0.209	0.053	0.234	0.576
Average premium paid	174.495	89.324	68	162	298
Share choosing Bronze	0.065	0.073	0	0	0
Share choosing Silver	0.188	0.173	0	0	0
Share choosing Gold	0.015	0.021	0	0	0
Share choosing Platinum	0.012	0.018	0	0	0

Panel (b): Premiums and choice shares by age and income								
	Bronze		Silver		Gold		Platinum	
	Premium	Share	Premium	Share	Premium	Share	Premium	Share
By age:								
27-34	120	0.050	174	0.122	229	0.010	272	0.010
35-49	117	0.058	181	0.175	248	0.013	299	0.011
50-64	104	0.086	207	0.259	321	0.022	409	0.016
By income (FPL%):								
140-150	5	0.011	57	0.336	133	0.005	191	0.006
150-200	28	0.046	94	0.318	170	0.008	229	0.009
200-250	86	0.084	162	0.193	241	0.018	302	0.015
250-400	196	0.074	276	0.084	357	0.019	419	0.014

Note: Panel (a) reports statistic taken across the 30,007 bins in our main estimation sample. All statistics are weighted by number of potential buyers. For income, standard deviation means the standard deviation of the within-bin medians of income and average premium paid. In panel (b), premium is the average premium paid for buyers of a given age/income group.

In Table II, we provide some summary statistics. Each (M_i, X_i) bin contains 85 potential buyers on average. The average participation rate in Covered California is 28%, but participation varies widely across demographics and rating regions. Older and poorer buyers in particular are more likely to purchase coverage. The impact of the CSRs is evident in panel (b): buyers with income below 200% face monthly premiums under \$100 for a Silver plan with actuarial value of 88% or more (see Table I). Over 30% of such consumers purchase a Silver plan, whereas among consumers with income over 250% of the FPL, fewer than 9% purchase a more expensive and less generous non-CSR Silver plan.

3.3. Identifying Assumptions

Insurers in Covered California choose the base premium for each rating region and coverage level. This choice likely depends on differences in demand and costs specific to each rating region, for example due to the underlying socioeconomic or health characteristics of the residents in a region, or due to differences in provider networks. Using variation across regions risks confounding premiums with these other unobservable differences. For this reason, our primary empirical strategy uses only within-region variation, and we

do not impose any restriction on how preferences (the density of valuations f) vary across regions, M_i .⁵

Instead, we assume—in a limited way—that valuations are locally invariant to age and income. Within a region, post-subsidy premiums vary with age and income due to ACA regulations. The variation is outside of the control of insurers, who only choose base premiums for each region. The main concern with this strategy is that valuations also change with age or income due to changes in latent risk factors. We defend against this concern by only using local variation in age and income and then conducting a sensitivity analysis using the imperfect instrument generalization of Assumption IV.

In the notation of Section 2, we let W_i denote a coarse aggregate of (M_i, X_i) . The aggregates are constructed by grouping X_i into age bins given by $\{27\text{--}30, 31\text{--}35, 36\text{--}40, \dots, 56\text{--}60, 61\text{--}64\}$ and income bins given in percentage of the FPL by $\{140\text{--}150, 150\text{--}200, 200\text{--}250, 250\text{--}300, 300\text{--}350, 350\text{--}400\}$. A value of W_i is then taken to be the region indicator M_i crossed with all possibilities of these coarser age-income bins. Conditional on W_i , we still observe variation in premiums due to variation in age and income within the W_i bin. The assumption is that the distribution of latent valuations does not change as X_i varies within this bin. In terms of Assumption IV, this means taking $Z_i = X_i$ as the instrument, while conditioning on W_i .⁶

For example, one value of $W_i = w$ corresponds to individuals in the North Coast rating region who are aged between 36 and 40 with incomes between 150% and 200% of the FPL. Within this bin there are 50 values of X_i comprised of the ages 36, 37, 38, 39, and 40 crossed with the 10 income bins between 150% and 200% in steps of 5%. We observe a different premium vector for each of these 50 values, and we assume that observed choices are generated by the same distribution of valuations.

We evaluate sensitivity to this assumption by using the imperfect instrument generalization of Assumption IV. We specify κ as

$$\kappa(z, z', w) = \begin{cases} \kappa_{\text{age}}, & \text{if } z \text{ and } z' \text{ differ only in age, and only by a single bin,} \\ \kappa_{\text{inc}}, & \text{if } z \text{ and } z' \text{ differ only in income, and only by a single bin,} \\ +\infty, & \text{otherwise,} \end{cases}$$

where $\kappa_{\text{age}}, \kappa_{\text{inc}} \geq 0$ are parameters that control how much the density of valuations is allowed to vary. For example, taking $\kappa_{\text{age}} = 0.2$ and $\kappa_{\text{inc}} = 0$ allows for a 20% difference in valuations for any two adjacent 1-year age bins with the same income, while still requiring adjacent income bins with the same age to have identical valuations.

In Appendix S9, we also implement an alternative empirical strategy that allows for valuations to change arbitrarily across age and instead uses limited cross-region variation in premiums, while still maintaining invariance across income. The demand estimates from that strategy are broadly similar to the main results that use only within-region variation in age and income.

⁵As a consequence, we can estimate bounds region-by-region, which helps with computation. For statistical inference, we need to consider all regions simultaneously, which is the main reason that inference is so much more computationally intensive; see Appendix S7 for more detail.

⁶Since $P_i = \pi(M_i, X_i)$ does not have any unobserved stochastic components, this setting is an atypical example of an instrumental variable. The simulations in Appendix S3 illustrate our method when the endogenous variable depends on an additional unobservable component.

Throughout all of the analysis, we use Assumption [SP](#) to exploit unambiguous vertical orderings between plans. We specify the support sets $\mathcal{V}^\bullet(w)$ as:

$$\mathcal{V}^\bullet(w) = \begin{cases} \{v \in \mathbb{R}^4 : v_2 \geq v_4 \geq v_3 \geq v_1\} & \text{if } w \text{ has income below 150\% FPL,} \\ \{v \in \mathbb{R}^4 : v_2 \geq v_1, v_4 \geq v_3 \geq v_1\} & \text{if } w \text{ has income in 150–200\% FPL,} \\ \{v \in \mathbb{R}^4 : v_4 \geq v_2 \geq v_1, v_4 \geq v_3 \geq v_1\} & \text{if } w \text{ has income in 200–250\% FPL,} \\ \{v \in \mathbb{R}^4 : v_4 \geq v_3 \geq v_2 \geq v_1\} & \text{if } w \text{ has income above 250\% FPL.} \end{cases}$$

This specification requires consumers to always have higher valuations for plans that dominate on all actuarial characteristics. The different cases account for the CSRs, which change at 150%, 200%, and 250% of the FPL (see [Table I](#)). We do not restrict the ordering of Silver relative to Gold and Platinum for the 150–200% FPL bracket, or relative to Gold for the 200–250% FPL bracket, because the CSRs make these plans differ in ways that may be more or less attractive for different consumers. We also do not assume that any of the plans are valued more than the outside option, that is, we allow valuations for Covered California plans to be negative.

3.4. Parametric Models

For comparison, we also consider some fully parametric models, all of which follow a specification similar to equation (2):

$$Y_i = \arg \max_{j \in \mathcal{J}} \mathbb{1}[j \neq 0](\gamma_i - \alpha_i P_{ij} + \beta_i \text{AV}_{ij}) + \epsilon_{ij}, \quad (17)$$

where γ_i is an individual-specific intercept, AV_{ij} is the actuarial value of tier j for individual i (see [Table I](#)), and α_i and β_i are individual slope coefficients. The indicator normalizes the contribution of these terms to 0 for the outside option ($j = 0$). We consider logit models in which ϵ_{ij} is assumed to follow a type I extreme value distribution, independently across j , as well as a probit model in which ϵ_{ij} follows a normal distribution. We always estimate model (17) region-by-region, so that all parameters vary by region in an unrestricted way.

The first model we estimate is a logit in which the price parameter, α_i , is constant within regions, but both γ_i and β_i vary with age and income in coarse bins. In particular, the specification allows β_i to vary freely by region with a different value in each of the age bins used to define W_i : {27–30, 31–35, 36–40, ..., 56–60, 61–64}. It allows γ_i to vary freely by region, and within each region restricts $\gamma_i = \gamma_i^{\text{inc}} + \gamma_i^{\text{age}}$, where γ_i^{inc} varies in income bins {140–150, 150–200, ..., 350–400}, and γ_i^{age} varies in the same age bins as β_i . The second model is a probit with the same specification.⁷

We then consider three mixed logit models. In each of these models, γ_i and β_i vary with observables in the same way as in the first logit model. The three models differ in whether γ_i (“mixed logit I”), α_i (“mixed logit II”), or both (“mixed logit III”) have an additional unobservable component that is normally distributed with unknown variance. In the latter case, we also assume that the two unobservable components are uncorrelated.

We use these five parametric models to contextualize the nonparametric results and demonstrate some of the benefits—and limitations—of the nonparametric approach. The

⁷The probit still has ϵ_{ij} independent across j . We had difficulty allowing for correlation across j because the likelihood is quite flat.

comparison should be qualified by pointing out that the parametric approaches also differ in other ways from the nonparametric approach. The limited verticality we impose on the nonparametric model through Assumption SP cannot be imposed in the parametric models because the distribution of ϵ_{ij} has full support. Using actuarial value (AV_{ij}) in equation (17) adds a notion of verticality that should make the comparison closer.⁸ We impose some additive separability between income and age because including more parameters created convergence problems in some regions. Perhaps most substantively, the parametric models are estimated through maximum likelihood, while our nonparametric bounds are estimated using the absolute deviations approach in Section 2.8. As a consequence of these differences, the parametric estimates need not lie inside the estimated nonparametric bounds.

3.5. Demand Only

We do not model supply, so all of the results should be interpreted as holding supply fixed. This is an important caveat to all of our counterfactual estimates, both parametric and nonparametric.

4. DEMAND

We begin by estimating demand responses to uniform changes in monthly, per-person premiums. The premium vectors in these counterfactuals take the form $\pi(M_i, X_i) + \delta$ for various choices of δ . The first class of target parameters we consider are changes in choice shares. For good j , region (market) m , and consumer characteristics x , these can be written as

$$\Delta \text{Share}_j^\delta(m, x; f) \equiv \int_{\mathcal{V}_j(\pi(m, x) + \delta)} f(v|m, x) dv - \int_{\mathcal{V}_j(\pi(m, x))} f(v|m, x) dv, \quad (18)$$

where $\mathcal{V}_j(p)$ was defined in equation (8). We aggregate these changes in choice shares into a single measure by averaging over regions and characteristics:

$$\Delta \text{Share}_j^\delta(f) \equiv \sum_{m, x} \Delta \text{Share}_j^\delta(m, x; f) \mathbb{P}[M_i = m, X_i = x]. \quad (19)$$

In Table III, we report estimated bounds for $\Delta \text{Share}_j^\delta$ across the four metal tiers together with bounds on overall participation, $\sum_{j \geq 1} \Delta \text{Share}_j^\delta = 1 - \Delta \text{Share}_0^\delta$. In the first row of each panel of Table III, δ is set to a \$10 increase for all plans together. From the consumer's perspective, this is the same as a \$10 decrease in premium subsidies, or alternatively, the same as a \$10 increase in the ACA's "maximum affordable amount." In the next four rows of each panel of Table III, δ is set to a \$10 increase in per-person monthly premiums in each of the four metal tiers alone.

4.1. Participation and Substitution

Estimated impacts on overall participation are shown in the first column of Table III. A \$10 increase in all premiums reduces the proportion of individuals who purchase coverage by between 1.8% and 6.7%. Panel (b) shows that the bounds are larger in magnitude

⁸We also estimated specifications where AV_{ij} was replaced by product-specific intercepts. These specifications produced less price sensitivity.

TABLE III
NONPARAMETRIC BOUNDS ON CHANGES IN CHOICE SHARES.

\$10/month premium increase for	Change in probability of choosing									
	Any plan		Bronze		Silver		Gold		Platinum	
	LB	UB	LB	UB	LB	UB	LB	UB	LB	UB
Panel (a): Full sample (140–400% FPL)										
All plans	−0.067	−0.018	−0.012	−0.004	−0.051	−0.011	−0.004	−0.001	−0.003	−0.001
Bronze	−0.011	−0.002	−0.047	−0.009	+0.004	+0.044	+0.000	+0.028	+0.000	+0.023
Silver	−0.050	−0.003	+0.001	+0.124	−0.165	−0.017	+0.001	+0.121	+0.000	+0.097
Gold	−0.003	−0.000	+0.000	+0.005	+0.000	+0.010	−0.013	−0.003	+0.000	+0.011
Platinum	−0.002	−0.000	+0.000	+0.003	+0.000	+0.006	+0.001	+0.009	−0.010	−0.002
Panel (b): Lower income (140–250% FPL)										
All plans	−0.091	−0.020	−0.011	−0.003	−0.077	−0.015	−0.003	−0.001	−0.003	−0.001
Bronze	−0.009	−0.001	−0.046	−0.008	+0.004	+0.044	+0.000	+0.027	+0.000	+0.023
Silver	−0.076	−0.005	+0.001	+0.178	−0.237	−0.021	+0.001	+0.173	+0.000	+0.141
Gold	−0.002	−0.000	+0.000	+0.004	+0.000	+0.009	−0.011	−0.002	+0.000	+0.010
Platinum	−0.002	−0.000	+0.000	+0.004	+0.000	+0.006	+0.001	+0.010	−0.010	−0.002
Panel (c): Higher income (250–400% FPL)										
All plans	−0.037	−0.016	−0.015	−0.006	−0.018	−0.007	−0.004	−0.001	−0.003	−0.001
Bronze	−0.013	−0.003	−0.049	−0.009	+0.003	+0.045	+0.000	+0.029	+0.000	+0.023
Silver	−0.016	−0.001	+0.001	+0.053	−0.072	−0.012	+0.001	+0.054	+0.000	+0.040
Gold	−0.003	−0.000	+0.000	+0.006	+0.000	+0.012	−0.016	−0.004	+0.000	+0.013
Platinum	−0.002	−0.000	+0.000	+0.003	+0.000	+0.005	+0.001	+0.009	−0.010	−0.003

Note: Each pair of columns contains the estimated lower and upper bound for the change in choice probability of the choice indicated in columns in response to a \$10/month premium increase for the plan(s) indicated in the rows. The column “Any plan” means any choice $j \neq 0$, and the row “All plans” means all choices $j \neq 0$.

for lower income individuals, at between 2.0% and 9.1%, and panel (c) shows that they are smaller in magnitude for higher income individuals, at between 1.6% and 3.7%. Comparing panels (b) and (c) more generally, we find a pattern consistent with higher price sensitivity for lower income enrollees, which is consistent with the literature (e.g., [Abraham, Drake, Sacks, and Simon \(2017\)](#), [Finkelstein, Hendren, and Shepard \(2019\)](#)).

Premium increases for Gold or Platinum plans would have little impact on overall enrollment, which makes sense since there are relatively few buyers for these plans to begin with, and \$10 is a proportionally smaller increase in premiums for these plans. For higher-income consumers, we estimate similar bounds on participation from increasing premiums for either Bronze or Silver. For low-income consumers, increasing Silver premiums could decrease participation by as much as 7.6%, while an increase in the Bronze premium would cause a decrease of at most 0.9%.

The estimated bounds on substitution patterns within and between coverage tiers are also informative in many cases. For example, panel (a) shows that an increase in Bronze premiums by \$10 would lead to a decrease of between 0.9% and 4.7% in the share of consumers choosing Bronze coverage, and an increase in the share choosing Silver of between 0.4% and 4.4%.⁹ The upper bound on the increase in the share choosing Gold or Platinum is significantly smaller, reflecting the closer substitutability of the Bronze and Silver plans. The change in participation from a Bronze premium increase is between 0.2% and 1.1%, which is naturally both smaller and tighter than the change when all premiums are increased together. In contrast, increasing Platinum premiums by the same amount would lead to a much smaller decline in the proportion of buyers not purchasing coverage, which we measure to be at most 0.2%.

4.2. Sensitivity

In Table IV, we report sensitivity to Assumption IV. The reported bounds are for the change in participation due to a \$10 decrease in subsidies (increase in all premiums). The three pairs of columns correspond to allowing for variation across age but not income

TABLE IV
SENSITIVITY TO ASSUMPTION IV.

κ	Change in probability of purchasing coverage if all per-person premiums increase by \$10/month					
	$\kappa_{\text{age}} = \kappa, \kappa_{\text{inc}} = 0$		$\kappa_{\text{age}} = 0, \kappa_{\text{inc}} = \kappa$		$\kappa_{\text{age}} = \kappa_{\text{inc}} = \kappa$	
	LB	UB	LB	UB	LB	UB
0	-0.0674	-0.0183	-0.0674	-0.0183	-0.0674	-0.0183
0.2	-0.0691	-0.0192	-0.1076	-0.0344	-0.1017	-0.0223
0.3	-0.0699	-0.0196	-0.1227	-0.0395	-0.1083	-0.0258
0.4	-0.0705	-0.0198	-0.1355	-0.0436	-0.1191	-0.0314
0.6	-0.0718	-0.0204	-0.1556	-0.0485	-0.1415	-0.0311
$+\infty$	-0.0865	-0.0158	-0.2602	-0.0293	-0.2798	-0.0000

Note: Each pair of columns shows estimated lower and upper bounds on the change in choosing any inside choice ($j \neq 0$). The first pair adjusts κ_{age} , while keeping $\kappa_{\text{inc}} = 0$. The second pair adjusts κ_{inc} , while keeping $\kappa_{\text{age}} = 0$. The third pair adjusts both κ_{inc} and κ_{age} simultaneously.

⁹In Appendix Figure S2, we show that the joint identified set for the two choice shares is far from rectangular. For example, if the effect on Bronze is a decrease of 4.7%, then the effect on Silver must be an increase of between roughly 2% and 4.4%, narrower than the marginal bounds of 0.4% to 4.4%.

bins ($\kappa_{\text{age}} > 0, \kappa_{\text{inc}} = 0$), allowing for variation across income but not age bins ($\kappa_{\text{age}} = 0, \kappa_{\text{inc}} > 0$), and allowing for variation across both income and age bins ($\kappa_{\text{age}} = \kappa_{\text{inc}} > 0$). The top row with $\kappa_{\text{age}} = 0$ and $\kappa_{\text{inc}} = 0$ is the same for each of the three pairs of columns, and the same as reported in Table III. At the opposite extremes, the bottom row in the first pair of columns uses only premium variation in income, the bottom row in the second pair uses only premium variation in age, and the bottom row in the third pair uses neither. Intermediate values of κ_{age} and κ_{inc} limit the amount by which adjacent bins can differ, with larger values representing weaker identifying assumptions.¹⁰

The bottom row of the first pair of columns ($\kappa_{\text{age}} = +\infty$ and $\kappa_{\text{inc}} = 0$) shows that if we completely drop age invariance, the estimated bounds are 1.6% to 8.7%, which is only modestly wider than the baseline estimates of 1.8% to 6.7%. The bottom row of the second pair ($\kappa_{\text{age}} = 0$ and $\kappa_{\text{inc}} = +\infty$) has much wider bounds, implying that the estimates depend more on the income invariance assumption. When both income and age invariance are completely removed, the bounds become completely uninformative and we cannot rule out that a \$10 decrease in subsidies causes all 28% of the population currently enrolled to stop participating. Setting $\kappa_{\text{age}} = \kappa_{\text{inc}} = 0.6$ allows for the density of valuations in adjacent age and/or income bins to increase or decrease by 60%. The bounds do widen considerably, but remain informative. For a more modest relaxation, like $\kappa_{\text{age}} = \kappa_{\text{inc}} = 0.2$, the bounds are still close to the baseline estimates.

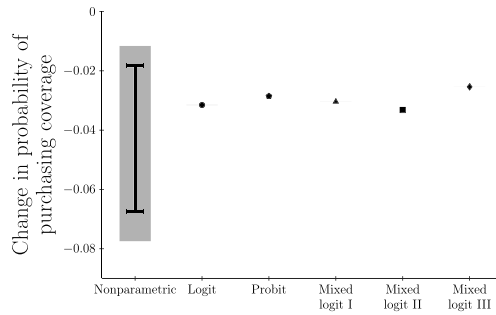
Overall, the results in Table IV show that the participation estimates primarily rely on the assumption that consumers in narrow income bins have similar preferences. This makes sense given Figure 2, since post-subsidy premiums vary more with income than with age. Assuming valuations are invariant to income is natural, since it gives the additive separability in choice model (1) the interpretation of quasilinearity (within coarse bins). The assumption could fail if either quasilinearity fails, or if higher income individuals have systematically different valuations, for example, due to lower health risk. However, Table IV shows that our results are largely robust to modest violations of income invariance. While we view local invariance of valuations to age as reasonable for the relatively homogenous groups of individuals within each 5 year bin, the assumption is unlikely to hold exactly since risk factors change with age (see, e.g., [Ericson and Starc \(2015\)](#), [Geruso \(2017\)](#), [Orsini and Tebaldi \(2017\)](#), [Tebaldi \(2022\)](#)). However, as demonstrated in Table IV, age invariance plays a secondary role in our estimates.

4.3. Comparison to Parametric Models

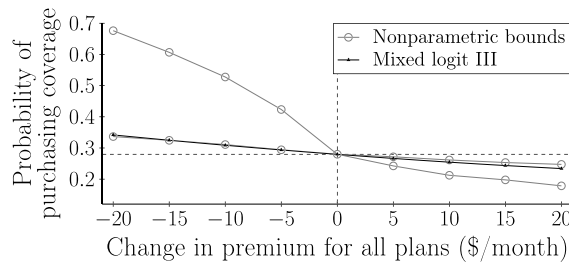
In Figure 3a, we plot the baseline bounds on the change in participation due to a \$10 decrease in subsidies (1.8% to 6.7%) together with 95% confidence intervals. For comparison, we also plot point estimates and confidence intervals for the parametric models discussed in Section 3.4. All of the point estimates lie within the nonparametric bounds, toward the upper bound, where price sensitivity is lowest.

By definition, any value within the nonparametric bounds can be produced by a distribution of valuations that matches the observed choice shares equally well. The parametric models produce a single point estimate by further requiring the distribution of valuations

¹⁰The bounds tend to be monotonic in κ_{age} and κ_{inc} , but this is not always the case. The population bounds would necessarily be monotonic, since larger values of κ_{age} and κ_{inc} correspond to weaker assumptions. However, this does not need to be the case when estimating the bounds using the procedure discussed in Section 2.8. The reason is that as κ_{age} or κ_{inc} increases, the value of $\hat{Q}^*$ always mechanically decreases, and thus the set $\hat{\Theta}^*$ over which the bounds are taken in the second step also changes, potentially excluding densities that fit the observed choice shares sufficiently well for smaller values of κ_{age} or κ_{inc} .



(a) Change in participation in response to a \$10 decrease in subsidies



(b) Change in participation in response to different sized changes in subsidies

FIGURE 3.—Comparison to parametric models. *Notes:* Top panel: Bound and point estimates are shown in solid black, and 95% confidence intervals are indicated with grey shading. The confidence interval for the logit and probit models are too narrow to be visible. Bottom panel: Nonparametric upper and lower bounds on the overall probability of purchasing coverage (choosing $j \neq 0$) for each price change are shown with light grey circles. Corresponding point estimates from mixed logit III are shown in black triangles.

to also have a particular shape. The conclusion we take from Figure 3a is therefore not that the parameterizations used in the parametric models do not matter. Instead, it is that the five parametric models we consider provide similar conclusions, and that there are other distributions of valuations that match the data equally well while leading to much different conclusions.

In Figure 3b, we explore this comparison for subsidy changes of different sizes. Larger changes reflect more ambitious counterfactuals that are more dissimilar from what was observed in the data. The nonparametric bounds widen to reflect this increasing dissimilarity. They are also wider for premium decreases than for increases, intuitively because we have less information on price sensitivity for the 72% of potential buyers who did not participate under the observed premiums. In contrast, the richest mixed logit model produces a point prediction for both increases and decreases, regardless of how dissimilar the counterfactual is to the observed data.¹¹ The point prediction hugs the lower bound for premium decreases and the upper bound for premium increases, suggesting that it could be underestimating the degree of price sensitivity, potentially by a considerable amount.

¹¹Confidence intervals for the mixed logit predictions in Figure 3b are not much different than those in Figure 3a, even for the most distant counterfactuals.

4.4. Elasticities

In Table V, we report estimated bounds on discrete approximations to average region-level elasticities. To avoid dividing by zero choice shares, we compute these elasticities at an aggregated level by grouping Silver, Gold, and Platinum into a single low deductible category, with the other option being high deductible Bronze plan. Appendix S6 contains more details, showing how we use our method to estimate bounds on discrete approximations of semielasticities, which we then turn into elasticities.

The Lerner index computed at the endpoints of the bounds implies a markup for marginal buyers of between 10% and 59% for the Bronze plan, and between 6% and 51% for the low deductible plans.¹²

Own-price elasticity estimates from the parametric models tend to be toward the non-parametric upper bounds, although there is considerable variation for the Bronze plan. The same is true for cross-price elasticities with the exception of the elasticity of the outside option to the price of Bronze, which the parametric models estimate to be at or above the nonparametric upper bound.

The sensitivity analysis in Table V aligns with the price variation in the data (Figure 2). The elasticities of choosing the outside option rely more on income variation than age variation, with the bounds widening only somewhat with $\kappa_{\text{age}} > 0$, but considerably more with $\kappa_{\text{inc}} > 0$. On the other hand, cross-price elasticities rely more on age variation, which “rotates” the tier prices in a way that income variation does not. Own-price elasticities seem to use both sources of variation in equal measure.

5. WELFARE AND MARKET DESIGN

5.1. Consumer Surplus and Government Spending

The change in consumer surplus from a change in premiums to $\pi(M_i, X_i) + \delta$ is

$$\Delta \text{CS}^\delta(m, x; f) \equiv \int \left[\max_{j \in \mathcal{J}} \{v_j - \pi_j(m, x) - \delta_j\} - \max_{j \in \mathcal{J}} \{v_j - \pi_j(m, x)\} \right] f(v|m, x) dv.$$

The associated change in government spending is

$$\begin{aligned} \Delta \text{GS}^\delta(m, x; f) &\equiv \sum_{j \geq 1} (\text{Sub}_j(m, x) - \delta_j) \times \left[\int_{\mathcal{V}_j(\pi(m, x) + \delta)} f(v|m, x) dv \right] \\ &\quad - \sum_{j \geq 1} \text{Sub}_j(m, x) \times \left[\int_{\mathcal{V}_j(\pi(m, x))} f(v|m, x) dv \right], \end{aligned}$$

where $\text{Sub}_j(m, x)$ denotes the baseline premium subsidy for purchasing plan j . We aggregate both $\Delta \text{CS}^\delta(m, x; f)$ and $\Delta \text{GS}^\delta(m, x; f)$ into single measures $\Delta \text{CS}^\delta(f)$ and $\Delta \text{GS}^\delta(f)$ by averaging over regions and demographics, the same as in equation (19).

In Figure 4, we show the estimated joint identified set for ΔCS^δ and ΔGS^δ when δ corresponds to a \$10 decrease in subsidies. The subsidy decrease would lead to a reduction in average monthly consumer surplus of between \$2.03 and \$2.40 per person, or between

¹²Markups for inframarginal buyers can be much smaller (possibly negative) if there is adverse selection and higher risk buyers are less price sensitive (e.g., Einav, Finkelstein, and Cullen (2010), Tebaldi (2022), Einav, Finkelstein, and Tebaldi (2019), Polyakova and Ryan (2019)).

TABLE V
ELASTICITIES.

1% premium increase for		% change in probability of choosing						
		Outside		High deductible		Low deductible		
		Bounds/Point estimate		Bounds/Point estimate		Bounds/Point estimate		
High deductible	Nonparametric	+0.025	+0.169	−9.797	−1.707	+0.256		+2.710
	$\kappa_{\text{age}} = 0.4 \quad \kappa_{\text{inc}} = 0$	+0.023	+0.190	−10.438	−1.787	+0.273		+3.002
	$\kappa_{\text{age}} = \infty \quad \kappa_{\text{inc}} = 0$	+0.000	+0.282	−11.369	−1.051	+0.074		+3.452
	$\kappa_{\text{age}} = 0 \quad \kappa_{\text{inc}} = 0.4$	+0.073	+0.387	−10.046	−2.632	+0.152		+2.707
	$\kappa_{\text{age}} = 0 \quad \kappa_{\text{inc}} = \infty$	+0.112	+0.898	−10.646	−2.292	+0.077		+2.727
	Logit		+0.154		−1.997		+0.154	
	Probit		+0.152		−1.902		+0.200	
	Mixed Logit I		+0.152		−1.966		+0.203	
	Mixed Logit II		+0.206		−4.411		+0.997	
	Mixed Logit III		+0.176		−4.039		+1.282	
Low deductible	Nonparametric	+0.207	+1.530	+1.364	+54.251	−15.491		−1.956
	$\kappa_{\text{age}} = 0.4 \quad \kappa_{\text{inc}} = 0$	+0.197	+1.583	+1.922	+59.219	−16.178		−2.183
	$\kappa_{\text{age}} = \infty \quad \kappa_{\text{inc}} = 0$	+0.052	+1.909	+0.235	+67.867	−18.444		−1.358
	$\kappa_{\text{age}} = 0 \quad \kappa_{\text{inc}} = 0.4$	+0.472	+3.064	+0.955	+52.746	−17.351		−3.638
	$\kappa_{\text{age}} = 0 \quad \kappa_{\text{inc}} = \infty$	+0.449	+5.851	+0.195	+63.914	−20.367		−2.288
	Logit		+0.641		+0.641		−3.549	
	Probit		+0.544		+1.200		−2.455	
	Mixed Logit I		+0.619		+0.799		−3.426	
	Mixed Logit II		+0.281		+2.876		−3.135	
	Mixed Logit III		+0.182		+3.263		−3.187	

Note: High deductible is Bronze and low deductible is a bundle consisting of Silver, Gold, and Platinum. See Appendix S6 for further details on implementation and computation.

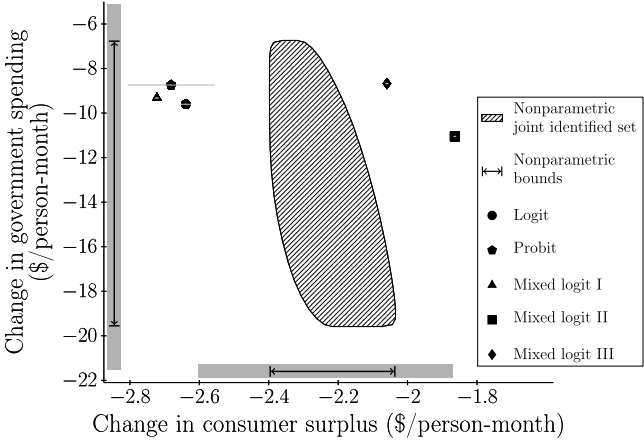


FIGURE 4.—Changes in consumer surplus and government spending from a \$10 decrease in subsidies. *Notes:* Bound and point estimates are shown in solid black. One-dimensional 95% confidence intervals are shown in grey vertical and horizontal bars.

\$62 and \$74 million yearly when aggregated, and a reduction in government spending of between \$6.74 and \$19.59 per person, for a total of between \$207 and \$602 million per year. The joint identified set is not rectangular: larger consumer surplus declines would be accompanied by smaller declines in government spending. The large spending declines are due to the marginal buyers who exit the market and relinquish their entire premium subsidy, which in most cases is considerably greater than \$10.

In Figure 4, we also plot the five-point predictions from the parametric models introduced in Section 3.4. The models yield similar predictions for the decline in government spending, all of which are toward the nonparametric upper bound. This is not surprising in light of Figure 3. Lower price sensitivity means fewer consumers leave the market when subsidies decrease, and all of the parametric models produced similarly insensitive participation responses.

More surprising is that the parametric models produce consumer surplus predictions that are both larger and smaller than the nonparametric bounds. Only the richest model (mixed logit III) makes a consumer surplus prediction within the nonparametric bounds. The usual intuition is that lower price sensitivity would lead to larger consumer surplus declines. However, the results in Figures 3a and 4 show that consumer surplus predictions in the parametric models are also driven by the functional form of unobserved heterogeneity, even among models that yield similar price sensitivity.

In Table VI, we report nonparametric consumer surplus changes by income. Changes for the lower-income sample (\$2.6–\$3.1 per person monthly) are estimated to be roughly twice as large as for the higher-income sample (\$1.3–\$1.5 per person monthly). This is likely due to the higher participation rate among the lower-income sample. In Table VI, we also report sensitivity analysis for the consumer surplus and government spending estimates. As with the participation estimates in the previous section, the results are quite robust to relaxing the age invariance assumption and depend more on the income invariance assumption.

The results indicate large differences between consumer surplus and government spending changes, consistent with a growing number of empirical analyses showing that consumers value individual health insurance significantly less than it costs in subsidies

TABLE VI
AGGREGATE IMPACTS FROM REDUCING PREMIUM SUBSIDIES BY \$10 PER MONTH.

	140–400% FPL		140–400% FPL		140–250% FPL		250–400% FPL	
	Change in government		Change in consumer		Change in consumer		Change in consumer	
	spending (\$ million/year)		surplus (\$ million/year)		surplus (\$/person-month)		surplus (\$/person-month)	
	Bounds/Point estimate		Bounds/Point estimate		Bounds/Point estimate		Bounds/Point estimate	
Nonparametric	–601.73	–207.05	–73.67	–62.49	–3.10	–2.59	–1.50	–1.32
$\kappa_{\text{age}} = 0.4 \quad \kappa_{\text{inc}} = 0$	–622.58	–217.51	–74.00	–62.17	–3.11	–2.58	–1.51	–1.31
$\kappa_{\text{age}} = \infty \quad \kappa_{\text{inc}} = 0$	–750.84	–188.13	–75.78	–56.74	–3.17	–2.39	–1.57	–1.15
$\kappa_{\text{age}} = 0 \quad \kappa_{\text{inc}} = 0.4$	–1136.26	–393.73	–72.10	–50.82	–3.02	–2.04	–1.49	–1.17
$\kappa_{\text{age}} = 0 \quad \kappa_{\text{inc}} = \infty$	–2092.82	–281.28	–74.55	–11.83	–3.08	–0.34	–1.60	–0.44
Logit		–295.13		–81.08		–3.44		–1.78
Probit		–268.52		–82.38		–3.39		–1.78
Mixed Logit I		–286.42		–83.64		–3.51		–1.88
Mixed Logit II		–339.43		–57.27		–2.32		–1.40
Mixed Logit III		–266.55		–63.27		–2.60		–1.49

Note: Each pair of columns corresponds to a different target parameter. Lower and upper bounds are shown for the nonparametric model with different sensitivity to age and income, while single point estimates are shown for the parametric models.

to induce them to purchase a plan (e.g., [Finkelstein, Hendren, and Shepard \(2019\)](#)). An important caveat when interpreting welfare estimates in this context is that they do not account for the existence of potentially large healthcare externalities such as the cost of uncompensated care, debt delinquency, or bankruptcy ([Finkelstein et al. \(2012\)](#), [Mahoney \(2015\)](#), [Garthwaite, Gross, and Notowidigdo \(2018\)](#)).

5.2. *Linking Subsidies to Age*

Regulated health insurance exchanges like Covered California rely on the participation of healthy consumers. A growing body of research considers linking subsidies to demographics as a tool for encouraging healthy consumers to participate (e.g., [Polyakova and Ryan \(2019\)](#), [Decarolis, Polyakova, and Ryan \(2020\)](#), [Tebaldi \(2022\)](#)). In Figure 5a, we diagram two counterfactual regulatory changes that directly link premium subsidies to age. In one counterfactual, subsidies are shifted from older buyers to younger buyers, while in the other the shift goes the opposite way.

We report the estimated effects of the two counterfactuals in Figures 5b and 5c. Shifting subsidies toward younger buyers could have a potentially large positive impact on both their participation and on aggregate participation, while decreasing the participation of older buyers comparatively less, even in the worst case. While the change in consumer surplus would naturally be positive for younger buyers and negative for older buyers, we find that average consumer surplus would unambiguously increase by between \$0.46 and \$2.68 per person, per month. The impacts on government spending could be either positive or negative, depending on exactly which sets of buyers change their purchase decisions. In contrast, shifting subsidies toward older buyers would increase government spending without necessarily increasing participation or average consumer surplus.

Figures 5b and 5c also include estimated effects for the five parametric models. As expected, they indicate price sensitivity toward the lower end of the nonparametric bounds. The parametric models all predict positive aggregate consumer surplus impacts from shifting subsidies toward younger buyers—like the nonparametric model—but potentially understate the magnitude of these impacts by a large amount. The parametric models predict no government spending impact from shifting subsidies toward younger buyers, whereas the nonparametric model says the data is consistent with either moderate decreases or large increases. Similarly, the parametric models unambiguously predict that the average consumer surplus impact of shifting subsidies toward older buyers would be negative, whereas the nonparametric model implies it could be either positive or negative.

5.3. *Removing Silver Plans*

The CSRs in Covered California make Silver plans especially attractive for low-income consumers (Tables I and II). Without Silver plans, subsidized consumers would face a more standard trade-off between premiums and actuarial value. If Silver plans ($j = 2$) were removed, the consumer surplus impact would be

$$\Delta RS(m, x; f) \equiv \int \left[\max_{j \in (\mathcal{J} \setminus \{2\})} \{v_j - \pi_j(m, x)\} - \max_{j \in \mathcal{J}} \{v_j - \pi_j(m, x)\} \right] f(v|m, x) dv.$$

As with the other parameters, we aggregate $\Delta RS(m, x; f)$ into $\Delta RS(f)$ by averaging over regions and demographics.

An interesting property of the nonparametric model is that the sharp lower bound on ΔRS is infinite ($-\infty$), at least unless we observe a choke price for Silver plans in the data

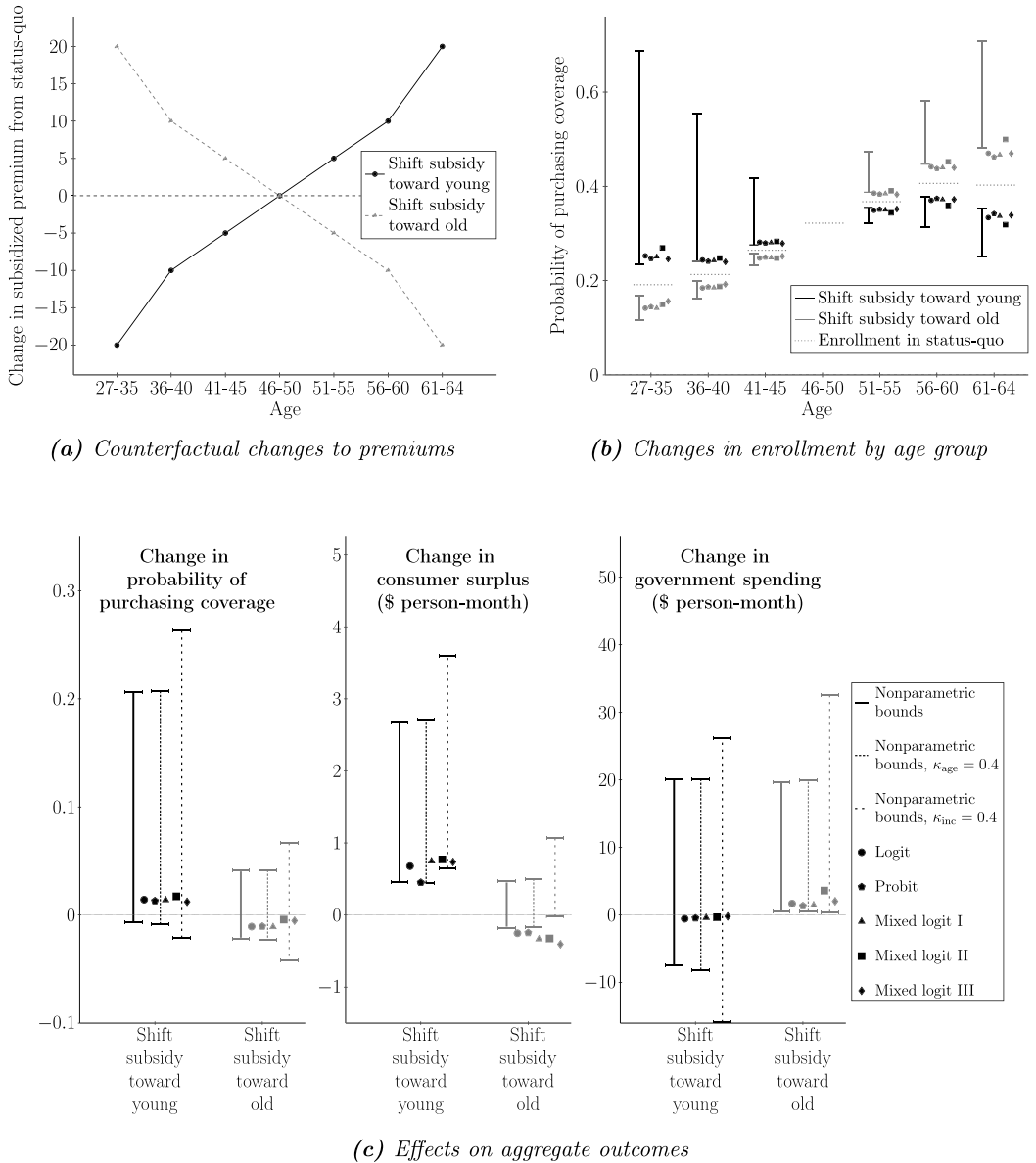


FIGURE 5.—Linking subsidies to age. *Notes:* Panel (a) illustrates the change in subsidized premiums by age under the two counterfactuals considered. Each x-axis group in panels (b) and (c) contains estimated nonparametric bounds and parametric point estimates on the indicated counterfactual, as well as the baseline value at the observed premiums.

(which we do not). The explanation is intuitive: without a parametric form for f , there is nothing to restrict the tails of the valuation for Silver, allowing for the possibility that some consumers have an unbounded preference for Silver plans. We view this property as a transparent benefit of the nonparametric model. It implies that parametric models rely on functional form to pick out a single consumer surplus impact from an unbounded set of possibilities. While an arbitrarily large consumer surplus impact is implausible, the fact

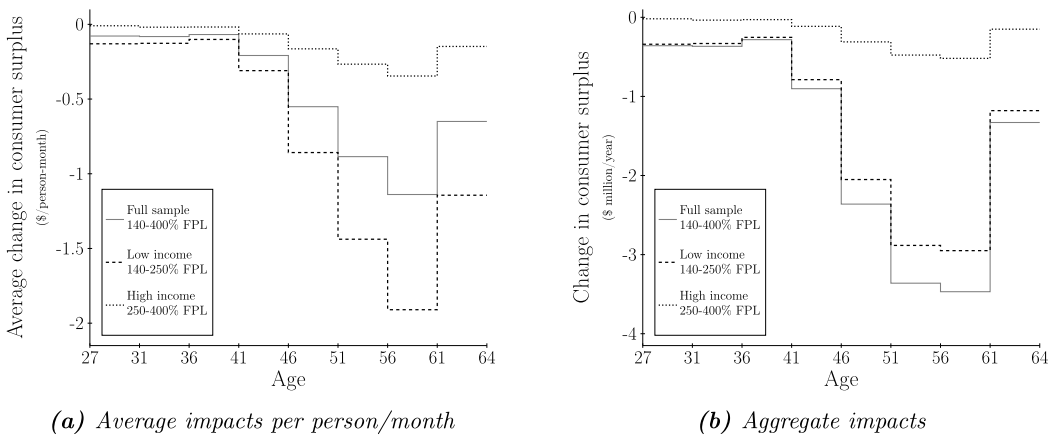


FIGURE 6.—Upper bounds on the change in consumer surplus from removing Silver plans. *Notes:* Each line indicates the estimated nonparametric upper bound on the change in consumer surplus for a different income group. The nonparametric lower bound is infinite.

that neither nonparametric assumptions nor the data rule it out highlights the role played by the specific choice of parameterization.

We focus instead on estimating the sharp upper bound on the decrease in consumer surplus from removing Silver plans. In Figure 6, we show that younger consumers would experience less surplus loss in the best-case scenario. The finding makes sense because these consumers also face a smaller premium differences between Bronze, Silver, and Gold plans. In Figure 6, we also show that low income consumers would bear the brunt of the surplus loss, which also makes sense because it is these consumers who receive the CSRs incorporated into the terms of Silver plans (Table I). In Table VII, we show that nearly \$11 million of the aggregate best-case surplus change of \$12.4 million would be borne by consumers with income less than 250% of the FPL.

TABLE VII
AGGREGATE IMPACTS FROM REMOVING SILVER PLANS.

	140–400% FPL		140–250% FPL		250–400% FPL	
	Change in consumer surplus (\$ million/year)		Change in consumer surplus (\$ million/year)		Change in consumer surplus (\$ million/year)	
	Bounds/Point estimate		Bounds/Point estimate		Bounds/Point estimate	
Nonparametric	$-\infty$	−12.43	$-\infty$	−10.78	$-\infty$	−1.66
$\kappa_{\text{age}} = 0.4 \quad \kappa_{\text{inc}} = 0$	$-\infty$	−11.63	$-\infty$	−10.25	$-\infty$	−1.38
$\kappa_{\text{age}} = \infty \quad \kappa_{\text{inc}} = 0$	$-\infty$	−0.68	$-\infty$	−0.68	$-\infty$	−0.00
$\kappa_{\text{age}} = 0 \quad \kappa_{\text{inc}} = 0.4$	$-\infty$	−9.77	$-\infty$	−8.36	$-\infty$	−1.41
$\kappa_{\text{age}} = 0 \quad \kappa_{\text{inc}} = \infty$	$-\infty$	−1.71	$-\infty$	−1.39	$-\infty$	−0.32
Logit	−281.67		−248.97		−36.24	
Probit	−290.40		−260.62		−29.79	
Mixed Logit I	−292.09		−257.40		−38.29	
Mixed Logit II	−148.95		−135.53		−14.48	
Mixed Logit III	−173.74		−162.85		−11.65	

Note: See notes for Table VI.

Estimates from the parametric models are also reported in Table VII. The parametric estimates all indicate consumer surplus impacts to be 7–10 times as large for lower income consumers, similar to the ratio of the nonparametric upper bounds. The magnitude of the estimates range from a decline of \$149 to a decline of \$292 million, which is between 2.5–3.5 times as large as the estimated impact from increasing all premiums by \$10 per month. The sensitivity of the estimates to the precise type of logit model used is reminiscent of findings by Petrin (2002). According to the nonparametric model, any figure between \$12.4 million and infinity is equally well supported by the observed choice shares.

The sensitivity estimates in Table VII show that these conclusions rely on both invariance to age and invariance to income. Fully removing either assumption renders the bounds essentially uninformative. The intuition can again be seen from Figure 6. Fixing income, post-subsidy premiums for Silver do not vary in age, making it difficult to pin down substitution between Silver and nonparticipation. Fixing age, post-subsidy premiums for Silver increase in lock-step with income, making it difficult to pin down substitution between Silver and other tiers. Information on both types of substitution are needed to infer the welfare impacts of removing Silver plans from the choice set.

6. CONCLUSION

We estimated the demand for health insurance in California’s ACA marketplace using a new nonparametric approach. The central idea of the method is to divide consumer valuations into a minimal relevant partition (MRP) of sets for which behavior remains constant under all considered prices. Using the MRP, we developed a scalable linear programming procedure for consistently estimating sharp identified sets for policy-relevant target parameters. We believe the method should be useful for other discrete choice problems as well. For example, it could be used to remove the large support assumption in the ideological voting model analyzed by Merlo and de Paula (2017).

The nonparametric estimates point to the possibility of substantially greater price sensitivity than would be recognized using comparable parametric models. This finding is consistent with the folklore that logits are “flat” models. The greater price sensitivity estimates in turn have important welfare implications for counterfactual policy changes to subsidies or plan offerings. However, we also found direct evidence that specific parametric functional forms have first-order impacts on consumer surplus estimates that operate through channels other than price sensitivity. The results provide a clear example in which functional form assumptions about the distribution of unobserved heterogeneity are not innocuous, and actually play a leading role in driving empirical conclusions.

APPENDIX: PROOFS FOR PROPOSITIONS 1 AND 2

A.1. Proposition 1

Suppose that $t \in \Theta^*$. By definition, there exists an $f \in \mathcal{F}^*$ such that $\theta(f) = t$. We will show that

$$\tilde{\phi}(\mathcal{V}|p, m, x) \equiv \int_{\mathcal{V}} f(v|p, m, x) dv \quad (20)$$

is an element of $\Phi^*(t)$.

First, note that $\tilde{\phi} \in \Phi$, because the MRP $\mathcal{V}$ is (almost surely) a partition of $\mathbb{R}^J$, and f is a conditional probability density function on $\mathbb{R}^J$. To see that $\tilde{\phi}$ satisfies equation (MD’),

observe that

$$\sum_{\mathcal{V} \in \mathbb{V}_j(p)} \tilde{\phi}(\mathcal{V}|p, m, x) \equiv \sum_{\mathcal{V} \in \mathbb{V}_j(p)} \int_{\mathcal{V}} f(v|p, m, x) dv = s_j(p, m, x; f) = s_j(p, m, x),$$

where the second equality follows from equation (7) using the definition of $\mathbb{V}_j(p)$, and the third holds for all $f \in \mathcal{F}^*$. Similarly, $\tilde{\phi}$ satisfies the upper bound in equation (IV') because

$$\begin{aligned} \tilde{\phi}(\mathcal{V}|w, z) &\equiv \mathbb{E}[\tilde{\phi}(\mathcal{V}|P_i, M_i, X_i)|W_i = w, Z_i = z] \\ &= \mathbb{E}\left[\int_{\mathcal{V}} f(v|P_i, M_i, X_i) dv | W_i = w, Z_i = z\right] \\ &= \int_{\mathcal{V}} f(v|w, z) dv \\ &\leq (1 + \kappa(z, z', w)) \int_{\mathcal{V}} f(v|w, z') dv \\ &= (1 + \kappa(z, z', w)) \tilde{\phi}(\mathcal{V}|w, z'), \end{aligned}$$

where the third equality follows by Tonelli's Theorem (e.g., [Shorack \(2000, p. 82\)](#)), the inequality uses Assumption IV, which is satisfied by any $f \in \mathcal{F}^*$, and the final equality reverses the steps of the first three. An analogous argument shows that $\tilde{\phi}$ also satisfies the lower bound of equation (IV'). That $\tilde{\phi}$ satisfies equation (SP') follows because $f \in \mathcal{F}^*$ satisfies Assumption SP, so

$$\begin{aligned} \sum_{\mathcal{V} \in \mathbb{V}^*(w)} \tilde{\phi}(\mathcal{V}|w, z) &= \sum_{\mathcal{V} \in \mathbb{V}^*(w)} \int_{\mathcal{V}} \mathbb{E}[f(v|P_i, M_i, X_i)|W_i = w, Z_i = z] dv \\ &= \int_{\cup\{\mathcal{V} : \mathcal{V} \in \mathbb{V}^*(w)\}} f(v|w, z) dv \geq \int_{\mathcal{V}^*(w)} f(v|w, z) dv = 1. \end{aligned} \quad (21)$$

The inequality in equation (21) uses the definition of $\mathbb{V}^*(w)$, which together with the fact that $\mathbb{V}$ is an a.s. partition of $\mathbb{R}^J$, implies that $\mathcal{V}^*(w)$ is contained in the union of sets in $\mathbb{V}^*(w)$. Inequality (21) implies that $\tilde{\phi}$ satisfies equation (SP'), because

$$\sum_{\mathcal{V} \in \mathbb{V}^*(w)} \tilde{\phi}(\mathcal{V}|w, z) \leq \sum_{\mathcal{V} \in \mathbb{V}} \tilde{\phi}(\mathcal{V}|w, z) = \mathbb{E}\left[\sum_{\mathcal{V} \in \mathbb{V}} \tilde{\phi}(\mathcal{V}|P_i, M_i, X_i)|W_i = w, Z_i = z\right] = 1,$$

as a result of $\tilde{\phi}$ being an element of Φ . Finally, since $\tilde{\phi} \equiv \bar{\theta}(f)$ as defined in equation (12), we immediately have that $\bar{\theta}(\tilde{\phi}) = \theta(f) = t$ due to Condition TP. We have now established that if there is an $f \in \mathcal{F}^*$ with $\theta(f) = t$, then $\tilde{\phi} \in \Phi^*(t)$.

Conversely, suppose that there exists a $\phi \in \Phi^*(t)$. Recall that W_i was assumed to be a subvector (or more generally, a function) of (P_i, M_i, X_i) , and denote this function by ω , so that $W_i = \omega(P_i, M_i, X_i)$. We will show that

$$\tilde{f}(v|p, m, x) \equiv \sum_{\mathcal{V} \in \mathbb{V}^*(\omega(p, m, x))} \frac{\mathbb{1}[v \in \mathcal{V} \cap \mathcal{V}^*(\omega(p, m, x))]}{\lambda(\mathcal{V} \cap \mathcal{V}^*(\omega(p, m, x)))} \phi(\mathcal{V}|p, m, x),$$

is an element of $\mathcal{F}^*$ that satisfies $\theta(\tilde{f}) = t$, where λ denotes Lebesgue measure on $\mathbb{R}^J$. (Note that the definition of $\mathbb{V}^\bullet(w)$ ensures that the denominator of each summand is nonzero.) Intuitively, the function $\tilde{f}(\cdot|p, m, x)$ distributes total mass of $\phi(\mathcal{V}|p, m, x)$ uniformly over each set $\mathcal{V} \in \mathbb{V}^\bullet(\omega(p, m, x))$.

For all $\mathcal{V} \in \mathbb{V}$,

$$\begin{aligned} \int_{\mathcal{V}} \tilde{f}(v|p, m, x) dv &\equiv \sum_{\mathcal{V}' \in \mathbb{V}^\bullet(\omega(p, m, x))} \int_{\mathcal{V}} \frac{\mathbb{1}[v \in \mathcal{V}' \cap \mathcal{V}^\bullet(\omega(p, m, x))]}{\lambda(\mathcal{V}' \cap \mathcal{V}^\bullet(\omega(p, m, x)))} \phi(\mathcal{V}'|p, m, x) dv \\ &= \mathbb{1}[\mathcal{V} \in \mathbb{V}^\bullet(\omega(p, m, x))] \phi(\mathcal{V}|p, m, x), \end{aligned} \quad (22)$$

since the sets in $\mathbb{V}$, and thus $\mathbb{V}^\bullet(\omega(p, m, x))$ are disjoint (almost surely). Using equation (22), we have that

$$\begin{aligned} \int_{\mathbb{R}^J} \tilde{f}(v|p, m, x) dv &= \sum_{\mathcal{V} \in \mathbb{V}} \int_{\mathcal{V}} \tilde{f}(v|p, m, x) dv \\ &= \sum_{\mathcal{V} \in \mathbb{V}^\bullet(\omega(p, m, x))} \phi(\mathcal{V}|p, m, x) = 1, \end{aligned} \quad (23)$$

where the first equality uses the fact that $\mathbb{V}$ is an (a.s.) partition of $\mathbb{R}^J$, and the final equality is implied by the hypothesis that ϕ satisfies equation (SP'), since

$$1 = \sum_{\mathcal{V} \in \mathbb{V}^\bullet(w)} \phi(\mathcal{V}|w, z) = \mathbb{E} \left[\sum_{\mathcal{V} \in \mathbb{V}^\bullet(\omega(P_i, M_i, X_i))} \phi(\mathcal{V}|P_i, M_i, X_i) | W_i = w, Z_i = z \right],$$

and every $\phi \in \Phi$ satisfies

$$\sum_{\mathcal{V} \in \mathbb{V}^\bullet(\omega(p, m, x))} \phi(\mathcal{V}|p, m, x) \leq \sum_{\mathcal{V} \in \mathbb{V}} \phi(\mathcal{V}|p, m, x) = 1.$$

Since $\tilde{f}$ inherits nonnegativity from $\phi \in \Phi^* \subseteq \Phi$, we conclude from equation (23) that $\tilde{f}$ is a conditional density, that is, $\tilde{f} \in \mathcal{F}$.

To see that $\tilde{f}$ satisfies the upper bound of Assumption IV, notice that

$$\begin{aligned} \tilde{f}(v|w, z) &\equiv \mathbb{E}[\tilde{f}(v|P_i, M_i, X_i) | W_i = w, Z_i = z] \\ &\equiv \mathbb{E} \left[\sum_{\mathcal{V} \in \mathbb{V}^\bullet(w)} \frac{\mathbb{1}[v \in \mathcal{V} \cap \mathcal{V}^\bullet(w)]}{\lambda(\mathcal{V} \cap \mathcal{V}^\bullet(w))} \phi(\mathcal{V}|P_i, M_i, X_i) | W_i = w, Z_i = z \right] \\ &= \sum_{\mathcal{V} \in \mathbb{V}^\bullet(w)} \frac{\mathbb{1}[v \in \mathcal{V} \cap \mathcal{V}^\bullet(w)]}{\lambda(\mathcal{V} \cap \mathcal{V}^\bullet(w))} \phi(\mathcal{V}|w, z) \\ &\leq (1 + \kappa(z, z', w)) \sum_{\mathcal{V} \in \mathbb{V}^\bullet(w)} \frac{\mathbb{1}[v \in \mathcal{V} \cap \mathcal{V}^\bullet(w)]}{\lambda(\mathcal{V} \cap \mathcal{V}^\bullet(w))} \phi(\mathcal{V}|w, z') \\ &= (1 + \kappa(z, z', w)) \tilde{f}(v|w, z'), \end{aligned}$$

where the inequality uses condition (IV'), and the final equality reverses the steps of the first three. An analogous argument can be used to show that $\tilde{f}$ also satisfies the lower bound of Assumption IV. Assumption SP is satisfied because

$$\begin{aligned} \int_{\mathcal{V}^\bullet(w)} \tilde{f}(v|w, z) dv &\equiv \int_{\mathcal{V}^\bullet(w)} \mathbb{E}[\tilde{f}(v|P_i, M_i, X_i)|W_i = w, Z_i = z] dv \\ &= \mathbb{E}\left[\sum_{\mathcal{V} \in \mathbb{V}^\bullet(w)} \phi(\mathcal{V}|P_i, M_i, X_i)|W_i = w, Z_i = z\right] \\ &= \sum_{\mathcal{V} \in \mathbb{V}^\bullet(w)} \phi(\mathcal{V}|w, z) = 1, \end{aligned}$$

where the second equality uses Tonelli's theorem with equation (22). That $\tilde{f}$ satisfies equation (MD) follows from the definition of $\mathbb{V}_j(p)$, equation (22), and equation (MD') via

$$\begin{aligned} s_j(p, m, x; \tilde{f}) &\equiv \sum_{\mathcal{V} \in \mathbb{V}_j(p)} \int_{\mathcal{V}} \tilde{f}(v|p, m, x) dv \\ &= \sum_{\mathcal{V} \in \mathbb{V}_j(p)} \mathbb{1}[\mathcal{V} \in \mathbb{V}^\bullet(\omega(p, m, x))] \phi(\mathcal{V}|p, m, x) \\ &= \sum_{\mathcal{V} \in \mathbb{V}_j(p)} \phi(\mathcal{V}|p, m, x) = s_j(p, m, x), \end{aligned}$$

for all $j \in \mathcal{J}$ and $(p, m, x) \in \text{supp}(P_i, M_i, X_i)$. The second-to-last equality here used the implication of equation (23) that $\phi(\mathcal{V}|p, m, x) = 0$ for any $\mathcal{V} \notin \mathbb{V}^\bullet(\omega(p, m, x))$. Finally, note that in the notation of equation (12), equation (22) says

$$\overline{\phi}(\tilde{f})(\mathcal{V}|p, m, x) = \int_{\mathcal{V}} \tilde{f}(v|p, m, x) dv = \mathbb{1}[\mathcal{V} \in \mathbb{V}^\bullet(\omega(p, m, x))] \phi(\mathcal{V}|p, m, x).$$

This equality implies that $\overline{\phi}(\tilde{f})(\mathcal{V}|p, m, x) = \phi(\mathcal{V}|p, m, x)$ for all $\mathcal{V}$, since equation (23) implies that $\phi(\mathcal{V}|p, m, x) = 0$ for $\mathcal{V} \notin \mathbb{V}^\bullet(\omega(p, m, x))$. As a consequence,

$$\theta(\tilde{f}) = \bar{\theta} \circ \overline{\phi}(\tilde{f}) = \bar{\theta}(\phi) = t,$$

and therefore $t \in \Theta^*$.

A.2. Proof of Proposition 2

Observe that Φ is a compact and connected subset of $\mathbb{R}^{d_\phi}$. Since equations (MD'), (IV'), and (SP') are linear (in)equalities, the subset of Φ that satisfies them is also compact and connected. If $\bar{\theta}$ is continuous on this subset, then its image over it—which Proposition 1 established to be Θ^* —is compact and connected as well. If $d_\theta = 1$, then Θ^* is a compact interval, so by definition its endpoints are t_{lb}^* and t_{ub}^* .

REFERENCES

ABRAHAM, JEAN, COLEMAN DRAKE, DANIEL W. SACKS, AND KOSALI SIMON (2017): "Demand for Health Insurance Marketplace Plans Was Highly Elastic in 2014–2015," *Economics Letters*, 159, 69–73. [129]

- ACKERBERG, DANIEL A., AND MARC RYSMAN (2005): "Unobserved Product Differentiation in Discrete-Choice Models: Estimating Price Elasticities and Welfare Effects," *The RAND Journal of Economics*, 36, 771–788. [110]
- ALLEN, ROY, AND JOHN REHBECK (2019): "Identification With Additively Separable Heterogeneity," *Econometrica*, 87, 1021–1054. [111]
- ANDREWS, DONALD W. K., AND GUSTAVO SOARES (2010): "Inference for Parameters Defined by Moment Inequalities Using Generalized Moment Selection," *Econometrica*, 78, 119–157. [121]
- ARMSTRONG, TIMOTHY B. (2016): "Large Market Asymptotics for Differentiated Product Demand Estimators With Economic Models of Supply," *Econometrica*, 84, 1961–1980. [113]
- BERRY, STEVEN (1994): "Estimating Discrete-Choice Models of Product Differentiation," *The RAND Journal of Economics*, 25, 242–262. [108]
- BERRY, STEVEN, AND PHILIP A. HAILE (2014): "Identification in Differentiated Products Markets Using Market Level Data," *Econometrica*, 82, 1749–1797. [111,114,117]
- (2016): "Identification in Differentiated Products Markets," *Annual Review of Economics*, 8, 27–52. [108,113,114]
- (2018): "Identification of Nonparametric Simultaneous Equations Models With a Residual Index Structure," *Econometrica*, 86, 289–315. [111]
- (2022): "Nonparametric Identification of Differentiated Products Demand Using Micro Data," arXiv:2204.06637 [econ]. [111,113]
- BERRY, STEVEN, JAMES LEVINSOHN, AND ARIEL PAKES (1995): "Automobile Prices in Market Equilibrium," *Econometrica*, 63, 841–890. [108,113]
- (2004): "Differentiated Products Demand Systems From a Combination of Micro and Macro Data: The New Car Market," *Journal of Political Economy*, 112, 68–105. [113]
- BERRY, STEVEN, OLIVER B. LINTON, AND ARIEL PAKES (2004): "Limit Theorems for Estimating the Parameters of Differentiated Product Demand Systems," *Review of Economic Studies*, 71, 613–654. [113]
- BERRY, STEVEN, AND ARIEL PAKES (2007): "The Pure Characteristics Demand Model," *International Economic Review*, 48, 1193–1225. [110]
- BERTSIMAS, DIMITRIS, AND JOHN N. TSITSIKLIS (1997): *Introduction to Linear Optimization*, Vol. 6. Belmont, MA: Athena Scientific. [120]
- BRIESCH, RICHARD A., PRADEEP K. CHINTAGUNTA, AND ROSA L. MATZKIN (2010): "Nonparametric Discrete Choice Models With Unobserved Heterogeneity," *Journal of Business & Economic Statistics*, 28, 291–307. [111]
- CANAY, IVAN A., ANDRES SANTOS, AND AZEEM M. SHAIKH (2013): "On the Testability of Identification in Some Nonparametric Models With Endogeneity," *Econometrica*, 81, 2535–2559. [111]
- CHAN, DAVID, AND JONATHAN GRUBER (2010): "How Sensitive Are Low Income Families to Health Plan Prices?" *American Economic Review*, 100, 292–296. [108,112]
- CHARNES, ABRAHAM, AND WILLIAM W. COOPER (1962): "Programming With Linear Fractional Functionals," *Naval Research Logistics Quarterly*, 9, 181–186. [119]
- CHESHER, ANDREW, AND ADAM ROSEN (2014): "An Instrumental Variable Random Coefficients Model for Binary Outcomes," *The Econometrics Journal*, 17, S1–S19. [117]
- CHESHER, ANDREW, ADAM ROSEN, AND KONRAD SMOLINSKI (2011): "An Instrumental Variable Model of Multiple Discrete Choice," Working paper. [118]
- (2013): "An Instrumental Variable Model of Multiple Discrete Choice," *Quantitative Economics*, 4, 157–196. [110,111,115–117]
- CHIAPPORI, PIERRE-ANDRE, AND IVANA KOMUNJER (2009): "On the Nonparametric Identification of Multiple Choice Models," Working paper. [111]
- CHIONG, XIANG CHIONG, YU-WEI HSIEH, AND MATTHEW SHUM (2017): "Counterfactual Estimation in Semiparametric Multinomial Choice Models," Working paper. [111]
- COMPIANI, GIOVANNI (2022): "Market Counterfactuals and the Specification of Multi-Product Demand: A Nonparametric Approach," *Quantitative Economics*, 13, 545–591. [110,111]
- CONLEY, TIMOTHY G., CHRISTIAN B. HANSEN, AND PETER E. ROSSI (2010): "Plausibly Exogenous," *Review of Economics and Statistics*, 94, 260–272. [115]
- DAFNY, LEEMORE, KATE HO, AND MAURICIO VARELA (2013): "Let Them Have Choice: Gains From Shifting Away From Employer-Sponsored Health Insurance and Toward an Individual Exchange," *American Economic Journal: Economic Policy*, 5, 32–58. [112]
- DEB, RAHUL, YUICHI KITAMURA, JOHN K.-H. QUAH, AND JÖRG STOYE (2021): "Revealed Price Preference: Theory and Empirical Analysis," arXiv:1801.02702 [econ]. [121]
- DECAROLIS, FRANCESCO (2015): "Medicare Part d: Are Insurers Gaming the Low Income Subsidy Design?" *American Economic Review*, 105, 1547–1580. [112]

- DECAROLIS, FRANCESCO, MARIA POLYAKOVA, AND STEPHEN P. RYAN (2020): "Subsidy Design in Privately Provided Social Insurance: Lessons From Medicare Part d," *Journal of Political Economy*, 128, 1712–1752. [112,136]
- DELEIRE, THOMAS, ANDRE CHAPPEL, KENNETH FINEGOLD, AND EMILY GEE (2017): "Do Individuals Respond to Cost-Sharing Subsidies in Their Selections of Marketplace Health Insurance Plans?" *Journal of Health Economics*, 56, 71–86. [112]
- DRAKE, COLEMAN (2019): "What Are Consumers Willing to Pay for a Broad Network Health Plan?: Evidence From Covered California," *Journal of Health Economics*, 65, 63–77. [112]
- EINAV, LIRAN, AMY FINKELSTEIN, AND MARK R. CULLEN (2010): "Estimating Welfare in Insurance Markets Using Variation in Prices," *Quarterly Journal of Economics*, 125, 877–921. [132]
- EINAV, LIRAN, AMY FINKELSTEIN, AND PIETRO TEBALDI (2019): "Market Design in Regulated Health Insurance Markets: Risk Adjustment vs. Subsidies," Unpublished mimeo, Stanford University, MIT, and University of Chicago. [112,132]
- EINAV, LIRAN, AND JONATHAN LEVIN (2015): "Managed Competition in Health Insurance," *Journal of the European Economic Association*, 13, 998–1021. [107]
- ERICSON, KEITH M., AND AMANDA STARC (2015): "Pricing Regulation and Imperfect Competition on the Massachusetts Health Insurance Exchange," *The Review of Economics and Statistics*, 97, 667–682. [108,112,130]
- FINKELSTEIN, AMY, NATHANIEL HENDREN, AND MARK SHEPARD (2019): "Subsidizing Health Insurance for Low-Income Adults: Evidence From Massachusetts," *American Economic Review*, 109, 1530–1567. [112,123,129,136]
- FINKELSTEIN, AMY, SARAH TAUBMAN, BILL WRIGHT, MIRA BERNSTEIN, JONATHAN GRUBER, JOSEPH P. NEWHOUSE, HEIDI ALLEN, KATHERINE BAICKER, AND OREGON HEALTH STUDY GROUP (2012): "The Oregon Health Insurance Experiment: Evidence From the First Year," *Quarterly Journal of Economics*, 127, 1057–1106. [136]
- FOSGERAU, MOGENS, AND DENNIS KRISTENSEN (2021): "Identification of a Class of Index Models: A Topological Approach," *The Econometrics Journal*, 24, 121–133. [111]
- FOX, JEREMY T. (2007): "Semiparametric Estimation of Multinomial Discrete-Choice Models Using a Subset of Choices," *The RAND Journal of Economics*, 38, 1002–1019. [111]
- FOX, JEREMY T., AND AMIT GANDHI (2016): "Nonparametric Identification and Estimation of Random Coefficients in Multinomial Choice Models," *The RAND Journal of Economics*, 47, 118–139. [111]
- GARTHWAITE, CRAIG, TAL GROSS, AND MATTHEW J. NOTOWIDIGDO (2018): "Hospitals as Insurers of Last Resort," *American Economic Journal: Applied Economics*, 10, 1–39. [136]
- GERUSO, MICHAEL (2017): "Demand Heterogeneity in Insurance Markets: Implications for Equity and Efficiency," *Quantitative Economics*, 8, 929–975. [130]
- GOLDBERG, PINELOPI KOUJIANOU (1995): "Product Differentiation and Oligopoly in International Markets: The Case of the U.S. Automobile Industry," *Econometrica*, 63, 891–951. [108]
- GUROBI OPTIMIZATION, INC. (2015): "Gurobi Optimizer Reference Manual." [120]
- HACKMANN, MARTIN B., JONATHAN T. KOLSTAD, AND AMANDA E. KOWALSKI (2015): "Adverse Selection and an Individual Mandate: When Theory Meets Practice," *American Economic Review*, 105, 1030–1066. [112]
- HANDEL, BEN, AND KATE HO (2021): "The Industrial Organization of Health Care Markets," in *Handbook of Industrial Organization*, Vol. 5. Elsevier, 521–614. [107]
- HANDEL, BEN, IGAL HENDEL, AND MICHAEL D. WHINSTON (2015): "Equilibria in Health Exchanges: Adverse Selection versus Reclassification Risk," *Econometrica*, 83, 1261–1313. [112]
- HANDEL, BENJAMIN R. (2013): "Adverse Selection and Inertia in Health Insurance Markets: When Nudging Hurts," *American Economic Review*, 103, 2643–2682. [112]
- HANDEL, BENJAMIN R., AND JONATHAN T. KOLSTAD (2022): "The Affordable Care Act After a Decade: Industrial Organization of the Insurance Exchanges," *Annual Review of Economics*, 14, 287–312. [107]
- HAUSMAN, JERRY A. (1996): "Valuation of New Goods Under Perfect and Imperfect Competition," in *The economics of new goods*, University of Chicago Press, 207–248. [110]
- HO, KATE, AND ARIEL PAKES (2014): "Hospital Choices, Hospital Prices, and Financial Incentives to Physicians," *American Economic Review*, 104, 3841–3884. [109–111]
- HO, KATE, AND ADAM M. ROSEN (2017): "Partial Identification in Applied Research: Benefits and Challenges," in *Advances in Economics and Econometrics*, ed. by Bo Honore, Ariel Pakes, Monika Piazzesi, and Larry Samuelson. Cambridge University Press, 307–359. [111]
- IBM (2010): "IBM ILOG CPLEX Version 12.8," International Business Machines Corporation. [120]
- JAFFE, SONIA, AND MARK SHEPARD (2020): "Price-Linked Subsidies and Imperfect Competition in Health Insurance," *American Economic Journal: Economic Policy*, 12, 279–311. [112]

- KAMAT, VISHAL (2021): "Identification of Program Access Effects With an Application to Head Start," arXiv:1711.02048 [econ]. [111,120]
- KHAN, SHAKEEB, FU OUYANG, AND ELIE TAMER (2021): "Inference on Semiparametric Multinomial Response Models," *Quantitative Economics*, 12, 743–777. [111]
- KITAMURA, YUICHI, AND JÖRG STOYE (2018): "Nonparametric Analysis of Random Utility Models," *Econometrica*, 86, 1883–1909. [121]
- KLINE, PATRICK, AND MELISSA TARTARI (2016): "Bounding the Labor Supply Responses to a Randomized Welfare Experiment: A Revealed Preference Approach," *American Economic Review*, 106, 972–1014. [111]
- KONING, RUUD H., AND GEERT RIDDER (2003): "Discrete Choice and Stochastic Utility Maximization," *The Econometrics Journal*, 6, 1–27. [108,117]
- KRUEGER, ALAN B., AND ILYANA KUZIEMKO (2013): "The Demand for Health Insurance Among Uninsured Americans: Results of a Survey Experiment and Implications for Policy," *Journal of health economics*, 32, 780–793. [112]
- LEWBEL, ARTHUR (2000): "Semiparametric Qualitative Response Model Estimation With Unknown Heteroscedasticity or Instrumental Variables," *Journal of Econometrics*, 97, 145–177. [111]
- (2019): "The Identification Zoo: Meanings of Identification in Econometrics," *Journal of Economic Literature*, 57, 835–903. [113]
- MAHONEY, NEALE (2015): "Bankruptcy as Implicit Health Insurance," *American Economic Review*, 105, 710–746. [136]
- MANSKI, CHARLES F. (1975): "Maximum Score Estimation of the Stochastic Utility Model of Choice," *Journal of Econometrics*, 3, 205–228. [111]
- (2007): "Partial Identification of Counterfactual Choice Probabilities," *International Economic Review*, 48, 1393–1410. [108,110,111,117]
- (2014): "Identification of Income-Leisure Preferences and Evaluation of Income Tax Policy," *Quantitative Economics*, 5, 145–174. [111]
- MANSKI, CHARLES F., AND JOHN V. PEPPER (2018): "How Do Right-to-Carry Laws Affect Crime Rates? Coping With Ambiguity Using Bounded-Variation Assumptions," *The Review of Economics and Statistics*, 100, 232–244. [115]
- MARONE, VICTORIA R., AND ADRIENNE SABETY (2022): "Should There Be Vertical Choice in Health Insurance Markets?" *American Economic Review*, 112, 304–342. [112]
- MATZKIN, ROSA L. (1993): "Nonparametric Identification and Estimation of Polychotomous Choice Models," *Journal of Econometrics*, 58, 137–168. [111]
- MCFADDEN, DANIEL, AND KENNETH TRAIN (2000): "Mixed MNL Models for Discrete Response," *Journal of Applied Econometrics*, 15, 447–470. [108]
- MERLO, ANTONIO, AND ÁUREO DE PAULA (2017): "Identification and Estimation of Preference Distributions When Voters Are Ideological," *Review of Economic Studies*, 84, 1238–1263. [139]
- MILLER, THOMAS P. (2017): "Examining the Effectiveness of the Individual Mandate Under the Affordable Care Act," American Enterprise Institute, statement before the House Committee on Ways and Means Subcommittee on Oversight. [123]
- MOGSTAD, MAGNE, ANDRES SANTOS, AND ALEXANDER TORGOVITSKY (2018): "Using Instrumental Variables for Inference About Policy Relevant Treatment Parameters," *Econometrica*, 86, 1589–1619. [120]
- MOLINARI, FRANCESCA (2020): "Microeconometrics With Partial Identification," in *Handbook of Econometrics*. Handbook of Econometrics, Vol. 7, ed. by Steven N. Durlauf, Lars Peter Hansen, James J. Heckman, and Rosa L. Matzkin 355–486. [111]
- NEVO, AVIV (2001): "Measuring Market Power in the Ready-to-Eat Cereal Industry," *Econometrica*, 69, 307–342. [113]
- (2011): "Empirical Models of Consumer Behavior," *Annual Review of Economics*, 3, 51–75. [113]
- NEVO, AVIV, AND ADAM M. ROSEN (2012): "Identification With Imperfect Instruments," *Review of Economics and Statistics*, 94, 659–671. [115]
- NEWWEY, WHITNEY K., AND JAMES L. POWELL (2003): "Instrumental Variable Estimation of Nonparametric Models," *Econometrica*, 71, 1565–1578. [111]
- ORSINI, JOE, AND PIETRO TEBALDI (2017): "Regulated Age-Based Pricing in Subsidized Health Insurance: Evidence From the Affordable Care Act," Available at SSRN 2464725. [122,130]
- PAKES, ARIEL (2010): "Alternative Models for Moment Inequalities," *Econometrica*, 78, 1783–1822. [111]
- PAKES, ARIEL, AND JACK PORTER (2021): "Moment Inequalities for Multinomial Choice With Fixed Effects," *Quantitative Economics*, Forthcoming. [111]
- PAKES, ARIEL, JACK PORTER, KATE HO, AND JOY ISHII (2015): "Moment Inequalities and Their Application," *Econometrica*, 83, 315–334. [111]

- PETRIN, AMIL (2002): “Quantifying the Benefits of New Products: The Case of the Minivan,” *Journal of political Economy*, 110, 705–729. [110,139]
- POLYAKOVA, MARIA, AND STEPHEN P. RYAN (2019): “Subsidy Targeting With Market Power,” Tech. rep, National Bureau of Economic Research. [108,109,112,132,136]
- RAMBACHAN, ASHESH, AND JONATHAN ROTH (2022): “A More Credible Approach to Parallel Trends,” Working paper. [115]
- RUGGLES, STEVEN, KATIE GENADEK, RONALD GOEKEN, JOSIAH GROVER, AND MATTHEW SOBEK (2015): “Integrated Public Use Microdata Series: Version 6.0 [Dataset],” University of Minnesota, Minneapolis. [123]
- SALTZMAN, EVAN (2019): “Demand for Health Insurance: Evidence From the California and Washington ACA Exchanges,” *Journal of Health Economics*, 63, 197–222. [108,109,112]
- (2021): “Managing Adverse Selection: Underinsurance versus Underenrollment,” *The RAND Journal of Economics*, 52, 359–381. [109]
- SANTOS, ANDRES (2012): “Inference in Nonparametric Instrumental Variables With Partial Identification,” *Econometrica*, 80, 213–275. [111]
- SHEPARD, MARK (2022): “Hospital Network Competition and Adverse Selection: Evidence From the Massachusetts Health Insurance Exchange,” *American Economic Review*, 112, 578–615. [112]
- SHI, XIAOXIA, MATTHEW SHUM, AND WEI SONG (2018): “Estimating Semi-Parametric Panel Multinomial Choice Models Using Cyclic Monotonicity,” *Econometrica*, 86, 737–761. [111]
- SHORACK, GALEN R. (2000): *Probability for Statisticians*. Springer-Verlag. [140]
- TAMER, ELIE (2010): “Partial Identification in Econometrics,” *Annual Review of Economics*, 2, 167–195. [111]
- TEBALDI, PIETRO (2022): “Estimating Equilibrium in Health Insurance Exchanges: Price Competition and Subsidy Design Under the ACA,” Tech. rep, National Bureau of Economic Research. [108,109,112,123,130,132,136]
- TEBALDI, PIETRO, ALEXANDER TORGOVITSKY, AND HANBIN YANG (2023): “Supplement to ‘Nonparametric Estimates of Demand in the California Health Insurance Exchange’,” *Econometrica Supplemental Material*, 91, <https://doi.org/10.3982/ECTA17215>. [111]
- THOMPSON, T. SCOTT (1989): “Identification of Semiparametric Discrete Choice Models,” Discussion paper 249, Center for Economics Research, Department of Economics, University of Minnesota. [111,117]
- TORGOVITSKY, ALEXANDER (2019): “Nonparametric Inference on State Dependence in Unemployment,” *Econometrica*, 87, 1475–1505. [115]

Co-editor Aviv Nevo handled this manuscript.

Manuscript received 9 April, 2019; final version accepted 26 July, 2022; available online 29 August, 2022.