

Optimal Place-Based Transfers: Theory, Identification, and a New Survey for Measurement¹

Md Moshir Ul Alam Morris A. Davis Jesse Gregory
Clark University Rutgers University University of Wisconsin
and NBER

July 28, 2026

Abstract

The optimal redistribution of income across locations is determined by two sufficient statistics: location by location, the average marginal utility of income and the average tax paid by residents in substitute locations that receive individuals migrating in response to a marginal tax change. We show that location-choice data do not identify average marginal utilities without strong restrictions on preferences and that standard data sets are insufficient to identify the cross-location elasticities that pin down the average tax of substitute locations. We design and field a unique nationally representative survey that allows for heterogeneity in preferences, and develop a method to aggregate survey responses to recover both sufficient statistics for 270 U.S. locations. We estimate large variation in the average marginal utility of income and strong regional substitution patterns. Optimal place-based taxes and transfers involve substantially more redistribution across locations than implied by policy rules resulting from common specifications of preferences.

¹We are grateful to the Rutgers Survey Research Center for their help in fielding the survey. We also thank David Albouy, Nate Baum-Snow, Leonardo D’Amico, Roberto Chang, Jonathan Dingel, Hector Blanco Fernandez, Cecile Gaubert, Andra Ghent, Edward Glaeser, John Kennan, Patrick Kline, Yuhei Miyauchi, Fabrizio Perri, Esteban Rossi-Hansberg, Gianluca Violante, Chris Taber, and Matt Wiswall for helpful conversations and seminar participants at the 2026 ASSA meetings, the 2026 Econometric Society North American Summer Meetings, the University of Illinois at Urbana-Champaign, the Indian Institute of Management-Calcutta, Johns Hopkins University, and Princeton University for helpful comments and suggestions. All errors are our own.

1 Introduction

Place-based transfers are taxes and subsidies that people pay or receive based only on where they live. Several recent papers characterize optimal place-based transfers in the presence of agglomeration externalities (Fajgelbaum and Gaubert, 2020, 2025; Donald et al., 2025, 2026), knowledge spillovers (Rossi-Hansberg et al., 2026), and redistributive motives (Gaubert et al., 2025). Across these papers, optimal place-based taxes and transfers depend on the elasticity of location choices with respect to after-tax income and the size of externalities when they are present. An extensive empirical literature has produced credible estimates of both of these factors.² Importantly, optimal place-based transfers in these papers are also driven by differences across locations in the average marginal utility of income and we have almost no *direct* evidence on how the marginal utility of income varies across locations. In this paper we take a first step at providing this and other key information that is required by theory to determine optimal place-based taxes and transfers.

We begin by solving a planning problem to derive a general formula for optimal place-based taxes and transfers that has as a central ingredient endogenous location choice. Like Gaubert et al. (2025), we set aside the possibility of agglomeration or other externalities to focus entirely on optimal redistribution. We show in an environment with an arbitrary number of locations that a planner needs two pieces of information (“sufficient statistics”) to determine optimal transfers across these locations. The first is average marginal utility of income in each location, which determines where an extra dollar delivers the largest gain in welfare. The second is, for each location, the average tax of substitute locations, i.e. the average tax in a new location paid by an outmigrant who leaves their existing location in response to a change in taxes. This average tax is determined by a location demand system, a matrix of location-choice elasticities that describes how people re-sort when taxes change, and therefore how costly it is to redistribute income. The theory thus makes explicit what must be identified by data.

We then show why these two sufficient statistics are difficult to recover from standard data sets without very strong assumptions. Under additively separable random-utility models, relative marginal utilities can in principle be identified from ratios of cross-location responses to location-specific income shocks (Allen and Rehbeck, 2019; Donald et al., 2026). We prove that if the idiosyncratic, location-specific shocks that affect location choice are allowed to also affect the marginal utility of income, full-support location-specific exogenous shifters of income only identify marginal utilities for households at the margin of moving, and at perturbed income levels. The solution to the planning problem requires knowledge of the

²See, for example, Davis et al. (2014) and Fajgelbaum et al. (2019).

average marginal utility in each location of *all* residents at baseline income levels.

Similarly, estimating the location demand system from standard location-choice data is difficult because the object required for optimal policy is a full matrix of own- and cross-location responses to changes in taxes and transfers, not a single migration elasticity. The problem is analogous to estimating substitution patterns across differentiated products: Identification of flexible substitution patterns requires rich variation in relative prices across alternatives.³ In spatial applications, researchers typically observe one national system of locations, or a short panel of that system, with limited plausibly exogenous variation in location-specific net incomes. As a result, existing work usually imposes structure on the location demand system – for example logit, nested logit, Fréchet, or other low-dimensional elasticity restrictions – rather than separately estimating own- and cross-location responses.

The main innovation of this paper is that we develop a new, novel survey to directly measure both required sufficient statistics to compute optimal place based policy. We field the survey on a nationally representative sample in the United States. Each survey respondent identifies four metropolitan areas (“locations”) where they could most plausibly live in several years: their hometown (pre-selected) and three additional destinations.⁴ We refer to each survey respondent’s individualized choice set of potential future locations as a “menu.” Respondents report their expected income in each of the locations in the menu, describe contextual factors like the presence of family and friends and how well the local environment fits their interests, and state the probability with which they would choose each location.

The survey then elicits estimates of the additional benefit that hypothetical extra permanent income would provide in each location. Explaining, respondents first indicate the location in which a future version of themselves would experience the largest increase in well-being from a permanent \$5,000 increase in income; call this location as “A.” For each of the three remaining locations in the menu, call them “B”, “C,” and “D,” respondents then report the required change to permanent income, greater than \$5,000, that delivers the same change in well-being in B, C, or D (as appropriate) as the \$5,000 increase in permanent income in location A. We show this set of 4 questions identifies the respondent-specific marginal utility of income in each of location of the menu up to a single, respondent-specific scale factor. We repeat this full set of questions across experimentally varied scenarios that shift incomes and contextual factors. Allowing for unobserved differences between locations

³The analogy is to the differentiated-products demand literature, where flexible substitution patterns are identified by variation in relative prices or incentives across products, often across many markets or periods, together with instruments for endogenous prices (McFadden, 1974; Berry et al., 1995; Nevo, 2000; Train, 2009; Berry and Haile, 2014).

⁴Throughout the paper, when we refer to a “metropolitan area” we consider a large set of metropolitan and micropolitan areas. Specifically we use all Core Based Statistical Areas (CBSAs) obtained from the US Census Bureau March 2020 CBSA Delineation file (U.S. Census Bureau, 2020).

and exploiting within-location variation across scenarios allows us to trace out how each respondent’s marginal utility of income differs, across locations, over a wide range of possible income levels and the only assumption is that all unobserved characteristics of each location otherwise remain constant across scenarios.

Experimental variation in our survey also lets us recover the location-demand system. We ask each respondent to report their baseline expected probability of living in each of the locations in their menu. This reveals which locations are commonly considered together and therefore likely to receive one another’s outflows if relative net incomes change and everything else remains unchanged. Across scenarios we shift incomes in specific locations and track how each respondent’s choice probabilities move, yielding individual-level substitution patterns across locations within each menu. In other words, the survey identifies not just how sensitive demand for a location is to its own after-tax income, but where people tend to go when they leave a given location. Aggregating the survey responses lets us recover the location-demand responses required for optimal place-based taxes and transfers.

We highlight four sets of results from our survey. First, using only data from our survey, we estimate a location-choice elasticity with respect to the wage of 1.69. This falls within the range of 1.2 to 3.1 from other papers that use revealed-preference methods to generate estimates (Bound and Holzer, 2000; Suárez Serrato and Wingender, 2014; Suárez Serrato and Zidar, 2015; Notowidigdo, 2020; Diamond, 2016). Second, our estimates of average marginal utility by location exhibit substantial variation that is not explained by differences in average income. Third, our menu-based design, which allows for individual-specific location choice sets, reveals regional substitution patterns across space which cannot be recovered by using region-specific nested logit models. Fourth, the survey data allow us to reject two assumptions frequently made in the literature: that the marginal utility of recent movers to a location is the same as the marginal utility of the existing population of that location (separability) and that when people move from an area they move in proportion to the existing population (aggregate IIA). We find instead that new movers to a location tend to have much different marginal utility than the existing residents, and that when people move from an area they tend to move to locations that are similar in some dimension, for example geographically.

Next, we provide a procedure for aggregating these individual-level survey objects into the location-level sufficient statistics that enter the planning problem. We propose an aggregation method for marginal utilities that is invariant to arbitrary within-person rescaling, preserves the possibility that some respondents’ entire menus consist of high- or low-marginal-utility locations, and avoids creating spurious ex-ante redistribution among individuals with the same menu, baseline probabilities, and Pareto weight. The procedure uses

overlap across menus to place normalized within-person comparisons on a common scale, aggregates them using choice probabilities and Pareto weights, and applies a precision adjustment so that locations with imprecisely estimated average marginal utilities are not assigned excessive redistribution by a planner.

Using these survey-based measures and our aggregation procedure, we compute optimal location-specific taxes and transfers for 270 locations (mostly CBSAs as commonly defined) in the United States. Relative to the location-weighted standard deviation of optimally-computed taxes that results from the standard assumptions of log utility and aggregate IIA migration, the location-weighted standard deviation of optimal place-based taxes implied by our survey is larger by more than a factor of four. Additionally, we show that our survey implies that the optimal tax policy that emerges after assuming log utility and aggregate IIA migration reduces welfare.

1.1 Related Literature

We build directly on the recent literature characterizing optimal spatial and place-based policy in quantitative spatial models (Fajgelbaum and Gaubert, 2020, 2025; Gaubert et al., 2025; Donald et al., 2025, 2026; Rossi-Hansberg et al., 2026).⁵ These papers show how optimal transfers across locations depend on equilibrium sorting, mobility responses, local fiscal externalities, housing and labor markets, agglomeration or knowledge spillovers, and distributional concerns. Closest to our theoretical analysis is Gaubert et al. (2025), who ask when place-based taxation can improve welfare relative to place-blind income taxation. In their two-location lump-sum case, the optimal place-based transfer trades off the difference across the two locations in average social marginal welfare weights against the fiscal cost of the migration it induces, the same two categories of objects that are the sufficient statistics in our framework, and their broader analysis embeds this tradeoff in a model with labor supply, nonlinear income taxation, housing markets, and redistribution between households and landlords. Rather than adding a new externality or fully calibrating a spatial equilibrium model, we exactly derive, for a general system of J locations, the two sufficient statistics that characterize optimal place-based taxes and transfers in the absence of agglomeration externalities. Our J -location formulation makes explicit that the planner needs the full matrix of own- and cross-location demand responses, not a single migration elasticity. We then focus on identifying and measuring those objects.

We also contribute to measurement in this literature. As explained above, neither the

⁵For broader discussions and empirical evaluations of place-based policies – including enterprise zones, local economic-development programs, and policies targeted to distressed regions – see Busso et al. (2013), Kline and Moretti (2014), Neumark and Simpson (2015), and Austin et al. (2018).

resident-average marginal utility of income in each location nor a flexible matrix of own- and cross-location demand responses can be recovered from standard location-choice data without strong assumptions. The first is a marginal-versus-average problem: additively separable random-utility models link marginal utilities to choice responses (McFadden, 1974; Train, 2009), and ratios of opposing cross-location responses identify relative marginal utilities (Allen and Rehbeck, 2019; Donald et al., 2026), but without separability, choice-shifting variation identifies only households at the margin of moving, a local-versus-average distinction familiar from the IV and marginal-treatment-effect literatures (Imbens and Angrist, 1994; Heckman and Vytlacil, 2005, 2007). The second is a high-dimensional demand problem: unrestricted cross-location substitution is the spatial analogue of flexible differentiated-products demand (Berry et al., 1995; Nevo, 2000; Berry and Haile, 2014), and standard spatial variation – including Bartik and shift-share designs – rarely supplies the independent variation in the full vector of location-specific net incomes that it requires (Bartik, 1991; Goldsmith-Pinkham et al., 2020; Borusyak et al., 2022). Our survey is designed to address both limitations directly.

Methodologically, we contribute to the elicited-choice-probability literature (Blass et al., 2010; Delavande and Manski, 2015; Wiswall and Zafar, 2015, 2018; Bleemer and Zafar, 2018; Koşar et al., 2022). Respondents in our survey construct individualized menus of feasible future locations, which preserves idiosyncratic attachments such as family ties, reveals which locations individuals view as substitutes, and avoids forcing all respondents to evaluate the same (possibly artificial) set of alternatives. Respondents also report probabilities over each choice and make location-by-location comparisons involving hypothetical changes to permanent income. Our approach therefore builds on this literature by allowing for unobserved differences across locations, rather than requiring alternatives to differ only in researcher-specified attributes. Closest to our setting, we build on Koşar et al. (2022) by recovering aggregate own- and cross-location demand responses, rather than the single average migration elasticity common in empirical migration work (e.g., Diamond, 2016; Notowidigdo, 2020). We also build on Alam et al. (forthcoming), who use compensating differentials to identify preferences in separable models, by recovering location-specific marginal utility without imposing separability.

Finally, we develop an aggregation and computation framework that turns survey-elicited individual objects into the aggregates required for welfare analysis (Chetty, 2009). The aggregation problem is nontrivial in part because marginal utilities are identified only up to a single person-specific scale. We use overlap across heterogeneous choice sets to discipline location-level aggregates, analogous in spirit to connected-set identification in high-dimensional fixed-effect models (Abowd et al., 1999, 2002), apply precision adjustments to

these aggregates using empirical-Bayes shrinkage (Efron and Morris, 1975; Morris, 1983), and combine these objects with demand responses in the sufficient-statistics formula for optimal transfers. To our knowledge, we are the first to do all this while also simultaneously allowing for heterogeneity across choice sets, marginal-utility profiles, and demand elasticities. Because marginal utilities are endogenous to the policy under evaluation (Saez and Stantcheva, 2016), in computation we allow each location’s marginal utility to respond to policy-induced changes to net income through an explicit, empirically estimated channel.

2 Optimal Redistribution across Locations

We now characterize optimal place-based redistribution when people can freely choose the metropolitan area (“location”) in which to reside and a federal government can impose taxes or transfers based *only* on where people live. The main message is that optimal policy in this environment depends on two sufficient statistics. The first is the marginal value of an additional dollar of income for the people who live in each location, after accounting for the planner’s Pareto weights. This object determines where a marginal dollar delivers the largest welfare gain. The second is the location demand system: how sensitive populations in each location are to location-specific taxes, including where movers go when a given location becomes more or less attractive. These mobility responses determine how costly it is to raise revenue from a particular location, because changing a tax shifts the tax base both locally and in other locations.

There is a unit mass of individuals indexed by i .⁶ Individuals belong to finitely many observable types $t(i) \in \{1, \dots, T\}$. Let N_t denote the population share of type t , so $\sum_t N_t = 1$. Locations are indexed by $j \in \{1, \dots, J\}$ and each individual chooses to live in exactly one location. Importantly, locations can vary in the income a person can earn (w_{ij}), in head taxes that are paid (τ_j) by each individual choosing to live in j , and in other factors that may determine a person’s overall enjoyment (ϵ_{ij}).

For a given vector of location-specific head taxes (or subsidies when negative), $\tau = (\tau_1, \dots, \tau_J)$, type t individuals face the indirect utility vector $\{u_{ij}^t(w_{ij} - \tau_j, \epsilon_{ij})\}_{j=1}^J$ over locations and choose their location optimally. Let $P_j^t(\tau) = \Pr(\text{type } t \text{ chooses location } j \mid \tau)$ denote the resulting choice probability. Since each individual chooses exactly one of the

⁶Throughout the theory, an “individual” denotes the residential decision-making unit. In our empirical implementation, each survey respondent is treated as representing a household decision-making unit, and the survey elicits location choices and marginal valuations using annual household income. Accordingly, location-specific incomes, location-specific taxes or transfers, and the dollar-valued welfare effects in the empirical analysis are interpreted as annual amounts per household. The theoretical characterization is unchanged when the decision-making unit is a household rather than a person.

J locations, $\sum_j P_j^t(\tau) = 1$. Define the total population share in location j as $L_j(\tau) = \sum_t N_t P_j^t(\tau)$. Since $\sum_t N_t = 1$ we have $\sum_j L_j(\tau) = 1$.

A marginal change in τ_j changes consumption one-for-one for all individuals who choose location j , so the relevant object for welfare is the expected marginal utility of income in that location, averaged over the individuals who actually live there. For type t , let μ_j^t denote the expected marginal utility of income conditional on choosing location j . Specifically, $\mu_j^t \equiv \mathbb{E}[u_w^t(w_{ij} - \tau_j, \epsilon_{ij}) \mid j \text{ chosen}]$. Let Π_t denote the Pareto weight that the planner puts on expected utility for type t individuals. The Pareto-weighted average marginal utility of income in location j is

$$\mu_j \equiv \frac{\sum_{t=1}^T N_t P_j^t(\tau) \Pi_t \mu_j^t}{\sum_{t=1}^T N_t P_j^t(\tau)}. \quad (1)$$

Given this, the population-average Pareto-weighted marginal utility of income is

$$\bar{\mu}(\tau) \equiv \sum_{j=1}^J L_j(\tau) \mu_j(\tau).$$

For type t , let $U_t(\tau)$ denote expected indirect utility given tax vector τ , i.e.

$$U_t(\tau) = \mathbb{E} \left[\max_{j \in \{1, \dots, J\}} u_{ij}^t(w_{ij} - \tau_j, \epsilon_{ij}) \mid t(i) = t \right]$$

and define the planner's objective as $W(\tau) = \sum_{t=1}^T \Pi_t N_t U_t(\tau)$. The planner maximizes the objective by choosing τ subject to a balanced budget constraint

$$\max_{\tau} W(\tau) \text{ s.t. } \sum_{j=1}^J \tau_j L_j(\tau) = G,$$

where G is exogenous net spending, possibly zero.

Theorem 2.1 (Sufficient statistics for optimal location-specific head taxes). *Suppose utility functions are continuously differentiable and the ϵ follow a continuous multivariate distribution (made precise in the regularity conditions in online Appendix A). Then any interior optimum τ^* that satisfies the balanced budget constraint $\sum_{j=1}^J \tau_j^* L_j(\tau^*) = G$, must satisfy,*

for each location j ,

$$\frac{\mu_j(\tau^*) - \frac{\bar{\mu}(\tau^*)}{M(\tau^*)}}{\frac{\bar{\mu}(\tau^*)}{M(\tau^*)}} = \frac{1}{L_j(\tau^*)} \sum_{j'=1}^J \tau_{j'}^* \frac{\partial L_{j'}(\tau^*)}{\partial \tau_j} \quad (2)$$

where $M(\tau)$ denotes the marginal revenue from a common increase in all location head taxes,

$$M(\tau) \equiv \frac{d}{d\delta} \sum_{j=1}^J (\tau_j + \delta) L_j(\tau + \delta \mathbf{1}) \Big|_{\delta=0} = 1 + \sum_{j'=1}^J \tau_{j'} \sum_{j=1}^J \frac{\partial L_{j'}(\tau)}{\partial \tau_j}.$$

and optimal taxes $\tau^* = (\tau_1^*, \tau_2^*, \dots, \tau_J^*)$ satisfy:

$$\tau_j^* = \underbrace{\frac{\sum_{j' \neq j} \frac{\partial L_{j'}(\tau^*)}{\partial \tau_j} \tau_{j'}^*}{\sum_{j'' \neq j} \frac{\partial L_{j''}(\tau^*)}{\partial \tau_j}}}_{\text{Average tax of substitute locations}} + \underbrace{\frac{L_j(\tau^*)}{\frac{\partial L_j(\tau^*)}{\partial \tau_j}} \left[\frac{\mu_j(\tau^*) - \frac{\bar{\mu}(\tau^*)}{M(\tau^*)}}{\frac{\bar{\mu}(\tau^*)}{M(\tau^*)}} \right]}_{\text{Marginal Utility based adjustment}} \quad (3)$$

See proof in online Appendix A.

The factor $M(\tau)$, defined in Theorem 2.1, is the marginal revenue of a uniform increase in all location head taxes;⁷ equivalently, the threshold marginal utility against which $\mu_j(\tau)$ is compared is $\bar{\mu}(\tau)/M(\tau)$ rather than $\bar{\mu}(\tau)$. It is exactly equal to 1.0 when a common head-tax increase does not re-sort households, and when we compute optimal taxes and transfers later on in the paper it is within 0.4% of 1.0. For this reason, the discussion below which treats the adjustment as centered on $\bar{\mu}$ is unaffected; we nonetheless impose the exact value of $M(\tau)$ in the computation (see online Appendix B for details on computation).

Since $\partial L_j / \partial \tau_j$ is assumed to be negative, the second term in (3) captures the intuition that locations with above-average Pareto-weighted marginal utility of income should be net recipients of transfers (i.e., face lower taxes), while locations with below-average marginal utility should be net contributors (i.e., face higher taxes). The magnitude of this adjustment is governed by the responsiveness of the location's population share to changes in its tax, $\partial L_j(\tau) / \partial \tau_j$: when population flows are more elastic, redistribution through location-specific taxes is more distortionary, so the marginal-utility adjustment is scaled by the inverse of this

⁷ As people migrate in response to this uniform tax, they move into locations with different baseline taxes, which either generates extra revenue or causes fiscal leakage. If this re-sorting causes revenue to leak (i.e., $M < 1$), then public funds are relatively expensive to raise. The planner becomes more fiscally conservative, so a location must have an exceptionally high marginal utility to justify receiving a subsidy. Conversely, if $M > 1$, funds are cheaper to raise, and the subsidy threshold drops.

elasticity.

The first term in (3) captures the fiscal-externality channel created by migration. When location j is taxed more heavily or subsidized less, marginal residents who are induced to leave substitute toward other locations and pay the taxes prevailing in those destinations. Because residents leaving a given location may substitute more toward some locations than others, the average tax paid by those exiting j will in general differ across j . The weighted average of those destination taxes is exactly the first term, and can be interpreted as the average tax faced by the marginal outflow from j .

Together, the formula shows that optimal policy has a starting point that is determined by the taxes in substitute locations and then shifts taxes up or down depending on whether μ_j is below or above the population average. In particular, if $\mu_j = \bar{\mu}$, the optimal policy generally does *not* set $\tau_j = 0$; instead, it sets τ_j equal to the average tax faced by the marginal movers who would leave j in response to a marginal increase in τ_j .

Equations (2) and (3) highlight two takeaways. First, (2) makes clear the objects the planner needs to set optimal location-based taxes and transfers: the Pareto-weighted marginal utilities $\{\mu_j\}$ and the matrix of migration responses $\{\partial L_{j'}/\partial \tau_j\}$. Second, (3) rewrites the same condition in a form that is closer to a policy prescription: a location's optimal tax equals the average tax faced by the marginal outflow from that location, plus an adjustment that depends on whether μ_j is above or below the population average, scaled by the own migration response. In the multi-location setting these conditions define a system in which each τ_j depends on the full vector τ , so computing the optimal taxes requires solving a fixed point.

3 Non-identification

In this section we explain why $\{\mu_j\}$ is not identified from location choice data alone without strong assumptions on preferences, even in a hypothetical environment with rich or even full-support location-specific wage shifters.

As a starting point for discussion, consider a standard separable random-utility model (McFadden, 1974; Train, 2009) (a “benchmark example”) in which utility in location j takes the form $U_{ij} = \alpha_j \log(w_{ij}) + a_j + \epsilon_{ij}$, where w_{ij} in this case is after-tax income, $\alpha_j > 0$ governs the marginal utility of income in location j , a_j is an amenity to living in j that is common to all i and enters additively, and ϵ_{ij} is an amenity to living in j that is specific to i and that also enters additively. In this model, cross-location choice responses encode information about marginal utilities. In particular, under a logit structure and in equivalent canonical

Fréchet-based formulations (e.g., [Eaton and Kortum, 2002](#)), cross derivatives take the form

$$\frac{\partial P_j}{\partial w_{j'}} \propto -P_j P_{j'} \frac{\alpha_{j'}}{w_{j'}}, \quad j \neq j'.$$

That is, the effect of increasing income in j' on the probability of choosing j is proportional to the marginal utility of income in j' . Crucially, relative marginal utilities are obtained from *cross derivatives in both directions*. For any pair of distinct locations j and j' ,

$$\frac{\partial P_j}{\partial w_{j'}} \propto -P_j P_{j'} \frac{\alpha_{j'}}{w_{j'}}, \quad \frac{\partial P_{j'}}{\partial w_j} \propto -P_{j'} P_j \frac{\alpha_j}{w_j}, \quad \implies \quad \frac{\partial P_j / \partial w_{j'}}{\partial P_{j'} / \partial w_j} = \frac{\alpha_{j'} / w_{j'}}{\alpha_j / w_j}.$$

This example illustrates the revealed-preference logic: comparing how the probability of choosing j responds to income changes in j' with the reverse cross response identifies the relative marginal utility of income across the two locations (up to the mechanical $1/w$ scaling under log income utility).⁸ This cross-derivative identification argument follows [Allen and Rehbeck \(2019\)](#), whose nonparametric results cover latent-utility models with additively separable unobserved heterogeneity; [Donald et al. \(2026\)](#) apply the result directly to spatial location choice, deriving exactly the ratio-of-cross-effects expression in the display above.

Two points are important for our purposes. First, this identification argument relies on *cross* responses – how shocks to income in one location shift choice probabilities in other locations – not on the own-location elasticity that applied work typically targets. In practice, available quasi-experimental variation is typically used to estimate a single average elasticity with respect to a location’s own income, and even that is often at the limit of what the data can support.⁹

Second, this example relies on separability as an identifying assumption: conditional on income, the marginal utility of income in a location does not depend on the idiosyncratic location shifter ϵ_{ij} that drives selection into that location. Location-choice responses reveal how exogenous changes to income shift a location’s attractiveness for people who are close to switching locations. This therefore speaks most directly to marginal utilities for households that are marginally attached to a location, i.e. most willing to move in response to a change in income. The assumption of separability is what allows researchers to treat the marginal

⁸The same mapping arises in the two canonical discrete-choice models. Additive/logit utility $U_{ij} = u(w_j) + a_j + \sigma \epsilon_{ij}$, with ϵ_{ij} i.i.d. Gumbel, gives $\partial P_j / \partial w_{j'} = -\frac{1}{\sigma} P_j P_{j'} u'(w_{j'})$; multiplicative/Fréchet utility $U_{ij} = w_j^\alpha A_j \epsilon_{ij}$, with ϵ_{ij} i.i.d. Fréchet of shape θ , gives $\partial P_j / \partial w_{j'} = -\theta \alpha P_j P_{j'} / w_{j'}$. In both, the ratio of opposing cross-responses $(\partial P_j / \partial w_{j'}) / (\partial P_{j'} / \partial w_j)$ recovers the marginal-utility ratio $MU_{j'} / MU_j$, equal to $u'(w_{j'}) / u'(w_j)$, and to $w_j / w_{j'}$ under log-income utility.

⁹For example, traditional Bartik-style instruments implemented with decennial Census data yield one predicted shifter per location per decade, which is far from the variation needed to recover the full matrix of cross-location responses required by this identification argument.

utilities of marginally attached residents as equal to the average marginal utility of income of *all* residents.

The more fundamental problem is that once one relaxes separability between income and idiosyncratic location shifters, entertaining the possibility that idiosyncratic factors that affect location choices may also shift the marginal utility of income, location-choice data do not identify the policy-relevant objects $\{\mu_j\}$. [Gaubert et al. \(2025\)](#) and [Fajgelbaum and Gaubert \(2025\)](#) construct an observational-equivalence example in which non-additive idiosyncratic location tastes generate identical sorting behavior but imply opposite welfare rankings across locations.¹⁰ Below we sharpen this concern by providing a proof for how location data do not identify marginal utilities, *even with a set of full-support instruments*. Consider two locations 1 and 2 and a population of individuals with wages (w_{i1}, w_{i2}) , and idiosyncratic shocks (e_{i1}, e_{i2}) . Suppose the researcher has access to a valid full support location-specific instrument vector $z = (z_1, z_2) \in \mathcal{Z} \subseteq \mathbb{R}^2$ that exogenously shifts location-specific incomes. Define the scalar utility difference $\Delta_i(z) \equiv u(w_{i1} + z_1, e_{i1}) - v(w_{i2} + z_2, e_{i2})$.

Theorem 3.1 (Non-identification of μ_j without separability assumptions). *Suppose instruments are location-specific and have full support (made precise in online Appendix C) and regularity conditions in online Appendix A hold. For each $z \in \mathcal{Z}$, the data identify*

$$\left. \frac{\partial P_2 / \partial z_1}{\partial P_1 / \partial z_2} \right|_z = \frac{\mathbb{E}[u_w(w_{i1} + z_1, e_{i1}) \mid \Delta_i(z) = 0]}{\mathbb{E}[v_w(w_{i2} + z_2, e_{i2}) \mid \Delta_i(z) = 0]}, \quad (4)$$

and thus do not identify the ratios relevant for the solution to the planning problem,

$$\frac{\mu_1}{\mu_2} \equiv \frac{\mathbb{E}[u_w(w_{i1}, e_{i1}) \mid \Delta_i(0) \geq 0]}{\mathbb{E}[v_w(w_{i2}, e_{i2}) \mid \Delta_i(0) < 0]}.$$

See online Appendix C for the proof. The non-identification intuition of Theorem 3.1 rests on two arguments. First, with full-support location-specific instruments, the data only identify ratios of marginal utilities of those at the *margin* in terms of location choice i.e., $\{i : \Delta_i(z) = 0\}$. Second, the cross responses identify objects tied to the set of individuals who are indifferent at a given configuration of perturbed wages $(w_{i1} + z_1, w_{i2} + z_2)$, not at baseline wages (w_{i1}, w_{i2}) . Hence, even with full-support instruments, which under weak assumptions allow the researcher to observe circumstances where residents who are inframarginal when $z = (0, 0)$ are made to be marginal, this delivers information about *their* marginal utilities of income at *perturbed* wages, not the average marginal utilities μ_1 and μ_2 among inframarginal

¹⁰In online Appendix D we provide two simple, intuitive examples for how marginal utilities are not identified using location choice data. In the first example, idiosyncratic location-specific preferences that determine location choice are complementary to consumption and in the second these preferences are substitutes.

and marginal residents at *baseline* wages.

In summary, even under the assumption of separability of w_j and ϵ_{ij} , the revealed-preference link from variation in wages to marginal utilities requires a full set of *cross*-location responses that is rarely feasible to estimate in practice. Once separability is relaxed, exogenous income shifters identify how wages change location attractiveness for individuals at the margin of moving, but they do not identify the conditional averages $\{\mu_j\}$ among *all* residents, including the inframarginal, that enter the planner’s first-order conditions. The observational-equivalence examples shown in online Appendix D show that this is not merely a power or instrument-availability issue, as without additional restrictions, distinct preference specifications can fit the same location-choice responses while implying different $\{\mu_j\}$.

4 Sufficient Statistics from Individual Survey Data

Motivated by the non-identification results discussed in Section 3, in this section we develop a formal measurement approach to estimate the required inputs to the planning solution using a nationwide hypothetical-choice survey. Section 4.1 shows how compensating-differential questions recover within-person ratios of state-contingent expected marginal utilities. Section 4.2 then shows how we estimate location-demand elasticities by proposing hypothetical future scenarios in which expected income and other variables change relative to a baseline (holding all else constant), and asking survey respondents how their expected behavior changes in response.

4.1 Marginal utility

How can a survey elicit cross-location differences in marginal utility? Our approach is to propose hypothetical changes to income in different locations and ask the respondent which has a larger impact on well-being for the “future version of themselves” who ends up living there. We start with the following warm-up question to get the respondent thinking about comparisons of changes to well-being in a hypothetical future where they may live in one of four locations – their hometown and three locations that they choose:

Imagine you have chosen to live in each of these locations. In which location would you benefit the most from an unexpected and permanent increase of \$5,000 in household income?

Call the selected location j^* ; this is the location where, at baseline expectations about the future, additional income would provide the most benefit. Then, for each remaining location

j , we ask the respondent to report the increase in permanent income in that location that would deliver the same gain in well-being as \$5,000 in j^* :

You said that \$5,000 of extra income would help you the most in $[j^]$. Now think about the version of yourself living in $[j]$, given the circumstances that led you to live there. What permanent increase in income in $[j]$ would improve your well-being by the same amount as \$5,000 in $[j^*]$?*

Let $\mu_{ij} \equiv \mathbb{E}[u_{w,ij} \mid j \text{ chosen}]$ denote the expected marginal utility of income for person i conditional on having chosen to live in location j . Denote the answer to the second question, the “compensating differential,” as b_{ij} . A first-order approximation gives

$$\mu_{ij} \cdot b_{ij} \approx \mu_{ij^*} \cdot \$5,000.$$

Defining $m_{ij} \equiv 1/b_{ij}$, it follows that

$$m_{ij} \approx \underbrace{\frac{1}{\mu_{ij^*} \times \$5,000}}_{s_i} \times \mu_{ij},$$

where s_i is a person-specific scale factor. The key observation is that s_i cancels in any across-location, within-person comparison:

$$\frac{m_{ij}}{m_{ij'}} \approx \frac{\mu_{ij}}{\mu_{ij'}} \quad \forall j, j' \in \mathcal{M}_i.$$

where we use the notation $\mathcal{M}_i$ to denote the menu of four possible locations respondent i may live in the future. Thus, for each person, the compensating-differential design recovers the full set of marginal-utility *ratios* across locations up to a first-order approximation, *without imposing any functional-form restriction on utility*.

4.2 Location demand elasticities

Estimating the migration-response matrix requires observing how choice probabilities respond to exogenous changes in location-specific income. We designed the survey to generate this variation by presenting each respondent with multiple hypothetical scenarios that vary income across locations while explicitly instructing respondents to hold all other location-specific attributes fixed across scenarios. Note that this is weaker than the typical instruction in the elicited-choice-probability literature to hold all attributes *across* choices (locations, in our case) to be the same. Since the menu is individual-specific and well-defined, we can

exploit within-location variation across scenarios for each individual and thus allow for unobserved differences between locations. As in [Koşar et al. \(2022\)](#), stated choice probabilities under experimentally varied scenarios provide the exogenous within-person variation needed to estimate migration elasticities.

We estimate elasticities using a two-way fixed-effect regression in logs:

$$\log(p_{ijk}) = \nu_i \log(w_{ijk}) + \phi_{ij} + \phi_{ik} + u_{ijk}, \quad (5)$$

where i indexes survey respondents, j indexes locations, k indexes the scenario (where $k = 0$ reflects baseline expectations and $k = 1, \dots, 8$ are hypothetical scenarios we discuss later), w_{ijk} is the expected annual household income for respondent i in location j in scenario k , ϕ_{ij} are person-location fixed effects, ϕ_{ik} are person-scenario fixed effects, and ν_i is person i 's elasticity of location choice with respect to income. The person-location effects absorb all time-invariant differences in the attractiveness of each location to each person (amenities, family ties, idiosyncratic attachment), while the person-scenario effects absorb any scenario-level shifters common to all locations in the menu. The survey's hold-all-else-constant instruction ensures that only income varies across scenarios within a person-location cell, so ν_i is identified from the survey design.

Note that

$$p_{ijk} = \exp(\nu_i \log(w_{ijk}) + \phi_{ij} + \phi_{ik}) . \quad (6)$$

We can solve for ϕ_{ik} :

$$\begin{aligned} 1 &= \sum_{j \in \mathcal{M}_i} p_{ijk} = \exp(\phi_{ik}) \sum_{j \in \mathcal{M}_i} \exp(\nu_i \log(w_{ijk}) + \phi_{ij}) \\ \rightarrow \phi_{ik} &= -\log \left(\sum_{j \in \mathcal{M}_i} \exp(\nu_i \log(w_{ijk}) + \phi_{ij}) \right) \end{aligned}$$

This allows us to rewrite (6) as

$$p_{ijk} = \frac{\exp(\nu_i \log(w_{ijk}) + \phi_{ij})}{\sum_{j' \in \mathcal{M}_i} \exp(\nu_i \log(w_{ij'k}) + \phi_{ij'})}$$

Suppressing the scenario subscript k in the discussion that follows, we assume each person assigns $p_{ij} = 0$ to every location not included in their menu $\mathcal{M}_i$, so all location-choice probability mass is allocated within the menu. The population share in location j is $L_j =$

$\sum_i p_{ij}$, and the aggregate response of location j 's share to a change in location j' 's tax is:

$$\frac{\partial L_j}{\partial \tau_{j'}} = \sum_i \frac{\partial p_{ij}}{\partial \tau_{j'}}. \quad (7)$$

Person i contributes to this sum only if both j and j' are in $\mathcal{M}_i$. In the survey, w_{ij} denotes respondent i 's expected annual household income in location j . In mapping these responses to the planner's problem, we treat a \$1 increase in τ_j as equivalent to a \$1 reduction in w_{ij} , holding all other location attributes fixed. Equation (5) therefore implies

$$\frac{\partial p_{ij}}{\partial \tau_j} = - \frac{\nu_i p_{ij} (1 - p_{ij})}{w_{ij}}.$$

For cross effects ($j \neq j'$), the adding-up constraint $\sum_{j \in \mathcal{M}_i} p_{ij} = 1$ requires that probability mass leaving j' be redistributed across the other menu locations. Adopting within-menu IIA as a simplifying assumption, outflows from j' are allocated in proportion to baseline choice probabilities:

$$\frac{\partial p_{ij}}{\partial \tau_{j'}} = \frac{\nu_i p_{ij} p_{ij'}}{w_{ij'}} \quad j \neq j'.$$

Substituting into (7):

$$\frac{\partial L_j}{\partial \tau_{j'}} = \sum_i \frac{\nu_i p_{ij} p_{ij'}}{w_{ij'}}, \quad j \neq j'. \quad (8)$$

Equation (8) combines a within-menu behavioral restriction with observed heterogeneity across menus. At the respondent level, the logit specification allocates probability mass leaving one location across the other three locations in that respondent's menu in proportion to their baseline choice probabilities. Aggregate substitution need not satisfy IIA, however, because menu composition, baseline probabilities, expected incomes, and elasticities vary across respondents. Thus, within-menu logit governs how a given respondent's outflow is divided, whereas the respondent's menu determines which destinations can receive that outflow.

4.3 Survey overview and structure

Table A1 in online Appendix E provides a schematic overview of our survey, but we summarize the design logic here. Survey respondents are nationally representative (discussed later in section 5). Our survey begins by collecting geographic and economic background,

current and hometown metropolitan areas, and individual and household income. After collecting this preliminary data, each respondent constructs a personalized menu $\mathcal{M}_i$ of four metropolitan areas they would realistically consider living in the future: their hometown (i.e., where they grew up) plus three additional destinations.¹¹ We construct this respondent-led menu of four locations – as compared to pre-selected set of locations that are the same for all respondents – so that the respondent answers questions related to choice probabilities and marginal-utility comparisons over locations the respondent actually views as feasible.¹² For each survey respondent i and for each location $j \in \mathcal{M}_i$, we record the following variables

1. Expected household income five years ahead, w_{ij0} .
2. An indicator for whether family or close friends would be present, $f_{ij0} \in \{0, 1\}$.
3. Stated choice probabilities over the menu, $\{p_{ij0}\}_{j \in \mathcal{M}_i}$, with $\sum_{j \in \mathcal{M}_i} p_{ij0} = 1$.
4. Subjective rankings of each location by expensiveness, fun / hobbies fit, and support from family/friends.

These variables provide covariates for empirical analyses and anchor baseline expectations for the hypothetical block of scenarios described next, ensuring that the experimentally varied scenarios start from each respondent’s own realistic expectations. In a last step, we ask the set of marginal-utility questions described in Section 4 that identify the location where they would (under these baseline expectations) benefit the most from an extra \$5,000 in income and then the compensating differentials for the remaining three locations.¹³

Next, respondents complete questions under eight additional scenarios of assumptions, where scenarios are indexed by $k = 1, \dots, 8$.¹⁴ In each scenario, respondents are shown shocked values of income (w_{ijk}) and presence of family/friends (f_{ijk}) for each location $j \in \mathcal{M}_i$, and are instructed to hold all other aspects of each location fixed relative to baseline.¹⁵ For each scenario $k = 1, \dots, 8$, respondents report updated choice probabilities $\{p_{ijk}\}_{j \in \mathcal{M}_i}$ given they hypothetical new values of w_{ijk} and f_{ijk} . At the end of each scenario, respondents

¹¹63 percent of the survey respondents currently live in their hometown; for these respondents their menu of future locations includes their current location. For the other 37 percent, the menu of future locations includes their hometown but does not include the location in which they currently reside.

¹²Later in the paper, we show that limiting the cardinality of the menu to four locations does not change the pattern of our estimates of marginal utility across locations.

¹³Our use of quantitative beliefs and stated probabilities builds on the subjective-expectations literature, which shows that survey respondents can provide meaningful probabilistic beliefs about future outcomes (Dominitz and Manski, 1997; Manski, 2004).

¹⁴Throughout the rest of the paper, we let $k = 0$ denote the baseline.

¹⁵ w_{ijk} is generated for $k > 0$ by applying multiplicative shocks to w_{ij0} and f_{ijk} is generated for $k > 0$ by quasi-randomly toggling the presence of family and friends across locations and scenarios. We use quasi-random assignment rules to avoid scenarios with strictly dominated options, following the stated-preference design literature (Louviere et al., 2000; Carson and Louviere, 2011). Online Appendix E provides the full details of scenario construction and randomization.

answer the two marginal-utility questions discussed earlier. Let b_{ijk} denote the reported compensating differential for location j in scenario k (with $b_{ijk}=\$5,000$ for the location chosen in the first question in each scenario). Using the notation introduced in section 4.1, these questions identify

$$m_{ijk} \equiv \frac{1}{b_{ijk}}. \quad (9)$$

m_{ijk} is proportional to the respondent’s marginal utility of income in location j in scenario k , up to a respondent-scenario-specific scale factor.

Together with the baseline, the survey yields nine sets of choice probabilities and nine sets of compensating differentials per respondent, each over four locations. The survey concludes with a standard demographics module.

4.4 Discussion: relationship to hypothetical-choice literature

The closest paper methodologically to ours is [Koşar et al. \(2022\)](#), who elicit counterfactual choice probabilities across locations to measure migration aversion. Our design differs in two ways. First, our identification exploits variation within each location between scenarios for each individual (and not within scenarios). Specifically, we ask individuals to assume that given a location, any unobserved factor remains the same across scenarios. Thus, we can allow for unobserved differences across the locations which are held constant across scenarios. This is weaker than assuming that everything else between the locations within each scenario is held fixed.¹⁶ We are able to do this because the choice sets typically involve large metro areas which are well-defined, unlike other surveys that use hypothetical choice experiments to elicit preferences for objects in choice-sets which can have varied definitions and thus can be more difficult to define and compare.¹⁷ We also differ in the way we induce experimental variation in our survey. Instead of providing quasi-randomly-drawn incomes, we use the distribution of income documented in [Blundell et al. \(2008\)](#) to exogenously vary expected location-specific income across different scenarios. We also use quasi-randomly drawn zeros and ones for each of the four locations to exogenously vary the expected presence of family and friends in those locations.¹⁸ Finally, we extend the usage of compensating differentials

¹⁶Additionally, by allowing for unobserved differences across locations, we reduce the cognitive load on respondents, and make the identification robust to variation in unobserved factors across locations.

¹⁷For example [Wiswall and Zafar \(2018\)](#) use hypothetical choice experiments to identify preferences over hypothetical job attributes, [Alam et al. \(forthcoming\)](#) embed information experiments within hypothetical choice surveys to separately identify beliefs on gender gaps in manager quality, and [Andrew and Adams \(2025\)](#) identify beliefs on future marriage market matches.

¹⁸Following [Wiswall and Zafar \(2018\)](#) and [Alam et al. \(forthcoming\)](#), we use quasi-random (instead of random) variation to avoid scenarios in which one or more choices are strictly dominated in all attributes.

to identify marginal utilities without imposing any functional form assumptions, extending [Alam et al. \(forthcoming\)](#) who show validity of using compensating differentials to identify preferences in semi-parametric (separable) models.

5 Survey Results

Combining data from our second pilot in July, 2024 and final data collection in November, 2024, we collected responses from 2,193 individuals with non-missing data on key questions. In online Appendix [F](#), we compare the composition of our analysis sample to the 2022 5-year American Community Survey (ACS) ([U.S. Census Bureau, 2022](#)) along several core dimensions. The survey sample is broadly similar to the ACS along Census divisions and age bins, with modest differences in gender and race shares. Applying survey weights that balance the joint distribution of our sample to the ACS benchmark has negligible effects on our estimates. Online Appendix [F](#) also contains a figure that shows the spatial distribution of future destinations in the menus of sample respondents, reflecting both the geographic dispersion of our sample and expected migration patterns.

5.1 Validation of survey responses

We now conduct several exercises to validate that respondents provided substantively meaningful information and thoughtfully responded to survey questions. First, we examine the relationship between reported location-specific expected income and metro level wages. Specifically, we estimate

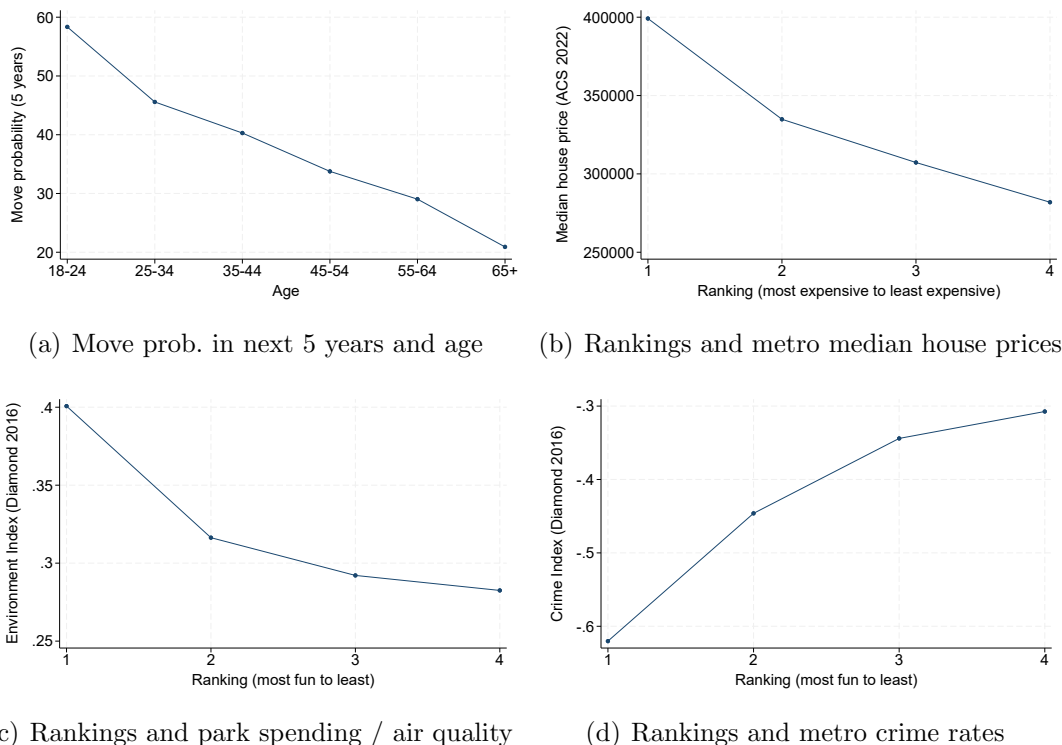
$$\log(w_{ij}) = \alpha_0 + \alpha_1 \overline{\log w_j} + \alpha_i + e_{ij}$$

where w_{ij} is respondent i 's baseline expected income in location j , α_i is a respondent fixed effect, and $\overline{\log w_j}$ is a skill-adjusted log earnings index for location j that we estimate from the 2022 5-year ACS ([U.S. Census Bureau, 2022](#)).¹⁹ We estimate $\alpha_1 = 0.421$ with a standard error of 0.017, suggesting respondents anticipate systematic income differentials across metropolitan areas in patterns that align with observed wage variation.

Next, in [Figure 1](#) we show that survey-elicited probabilities and rankings align with external data. Panel (a) of that figure shows a binned scatter plot of moving probabilities and age. The panel shows that young adults (18-24) report approximately 60% probability of moving within five years, declining systematically to approximately 20% for respondents over

¹⁹See online Appendix [G](#) for details on how we construct $\overline{\log w_j}$.

Figure 1: Further Validation of Survey Responses



65. This pattern is consistent with existing literature on migration patterns (Molloy et al., 2011; Kennan and Walker, 2011). Respondents’ rankings of metropolitan area characteristics also correlate highly with objective measures. The binned scatterplot in panel (b) shows a strong correlation of locations ranked by “expensiveness” (most to least) and median house prices from ACS data (U.S. Census Bureau, 2022).²⁰ Similarly, binned scatterplots of rankings of metropolitan “fun-ness” are positively correlated with indices of environmental quality obtained from Diamond (2016) that incorporate per-capita park spending and air quality measures (panel c) and negatively correlated with indices of metropolitan crime (panel d) that are also obtained from Diamond (2016).²¹

²⁰We ask “Given the sorts of goods and services you like to buy and given the kind of housing unit you wish to live in, which locations do you think would be most expensive for you?” and then ask the respondent to rank the likely future locations from 1 (most expensive) to 4 (least).

²¹We ask “Thinking about your hobbies and things you like to do for entertainment, where do you think you will have the most fun?” and then ask the respondent to rank the likely future locations from 1 (most fun) to 4 (least).

5.2 Survey Results: Location Choices

5.2.1 Elasticities of location choice

We start with a version of equation (5) that includes the presence of family and friends as a regressor and assumes the elasticity of location choice with respect to the wage is identical for all survey respondents:

$$\log(p_{ijk}) = \nu_1 \log(w_{ijk}) + \nu_2 f_{ijk} + \phi_{ik} + \phi_{ij} + u_{ijk} \quad (10)$$

Since we include controls for individual-location fixed effects (absorbing all factors that are time-invariant that affect the attractiveness of that location for each person) and individual-scenario fixed effects (absorbing scenario-level common shocks), the coefficient on log income ν_1 is identified purely from the quasi-randomly induced variation in income across scenarios. Due to the fact that choice probabilities are typically reported in multiples of five, we estimate equation (10) with a least-absolute-deviation quantile estimator, using the two-step approach of Canay (2011), and block-bootstrap the standard errors at the individual level.²²

We estimate an average location choice elasticity with respect to income (ν_1) of 1.69 (SE = 0.004), indicating substantial responsiveness of location decisions to income differences.²³ Our estimate is squarely in line with existing estimates from the literature of 1.2 to 3.1 as surveyed by Fajgelbaum et al. (2019), with an average estimate of 1.9 and variation in estimates depending on context and methodology: see Bound and Holzer (2000), Suárez Serrato and Wingender (2014), Suárez Serrato and Zidar (2015), Notowidigdo (2020), and Diamond (2016).

Next, we relax the assumption that all survey respondents have the same value of ν_1 and estimate equation (10) separately for each individual i using the two-step procedure of Canay (2011). As scenarios for each individual approach infinity, the individual location choice elasticities are identified. However, since estimation uses 8 scenarios giving us 24 linearly independent observations for each individual, the standard errors will be very high. Thus, we implement a parametric empirical Bayes (also known as a shrinkage estimator)

²²We use a least-absolute-deviation estimator (Blass et al., 2010; Wiswall and Zafar, 2018); this is insensitive to small perturbations of the extreme probabilities 0 and 1, which we set to 0.005 and 0.995 in estimation. Identification comes from the median-independence assumption $\mathbb{M}[u_{ijk} | w_{ijk}, f_{ijk}, \phi_{ik}, \phi_{ij}] = 0$. Because this conditions on the fixed effects, we follow the two-step procedure of Canay (2011): we first residualize $\log(p_{ijk})$ for the estimated person-scenario (ϕ_{ik}) and person-location (ϕ_{ij}) fixed effects, and then estimate a quantile regression of the residuals on the experimentally induced variation in income and the presence of family and friends. The key assumption is that the fixed effects shift all quantiles by the same amount; as the number of scenarios grows, this second step is consistent for the elasticities.

²³Our estimate of the location semi-elasticity with respect to the presence of family and friends is $\nu_2 = 0.20$ with a standard error of 0.007, quantifying the non-monetary value of social connections in location choice.

approach which shrinks the individual estimates towards the mean of the distribution based on estimated signal-to-noise ratio. The shrunk distribution is centered around 1.73, with a standard deviation of 0.55, yielding a standard error for the mean equal to 0.011.

5.2.2 Geographic substitution and diversion ratios

The optimal tax formula (3) shows that the tax on location j depends on the taxes of its substitutes, weighted by the share of marginal movers from j that each substitute absorbs. We now use the estimated individual-level elasticities, together with the structure of our survey, to characterize these substitution patterns.

The aggregate response of the population share in location j to a change in the tax in location j' is

$$\frac{\partial L_{j'}}{\partial \tau_j} = \sum_i \frac{\nu_i p_{ij} p_{ij'}}{w_{ij}}, \quad j \neq j',$$

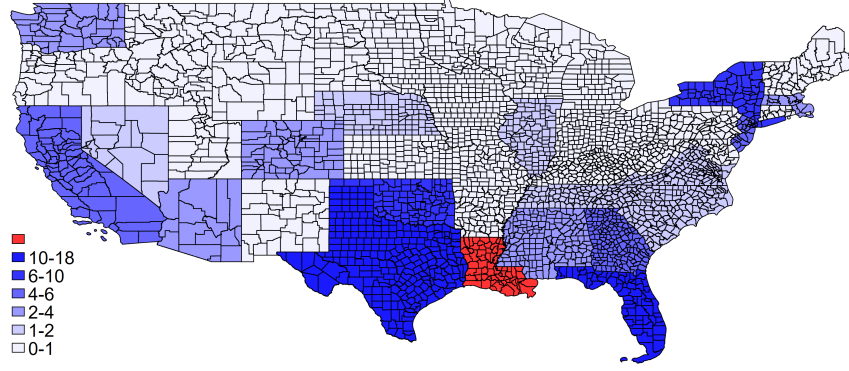
where ν_i is the location-choice elasticity estimated for respondent i , p_{ij} and $p_{ij'}$ are baseline choice probabilities for locations j and j' , and w_{ij} is the income respondent i expects in location j at baseline.²⁴ We set $p_{ij} = 0$ whenever $j \notin M_i$. Among respondents whose menus contain j , those who also list j' contribute $(\nu_i/w_{ij})p_{ij}p_{ij'}$, while those who do not list j' contribute zero. The frequency with which j and j' are co-listed is therefore informative about their aggregate substitutability (together with respondents' baseline probabilities, expected incomes, and elasticities).

To illustrate, Figure 2 presents the substitution patterns when metro areas in Louisiana are subject to an increase in taxes. The map shows which states would receive marginal movers, with darker shading indicating areas receiving more of the population leaving Louisiana. The pattern reveals strong regional clustering: Texas, Mississippi, and Florida receive the bulk of potential out-migrants, while there is minimal substitution to Northern or Western states. This regional clustering is consistent with evidence that distance, local labor-market information, and nonpecuniary attachments shape migration decisions (Dahl, 2002; Kennan and Walker, 2011; Kaplan and Schulhofer-Wohl, 2017).

To quantify these patterns systematically, we normalize the cross-derivatives into *diversion ratios*, defined as the share of marginal outflows from one location absorbed by another.

²⁴From equation (5), the change in the probability of living in location j following a one-dollar increase in its tax is $-(\nu_i/w_{ij})p_{ij}(1 - p_{ij})$.

Figure 2: Substitution Patterns from Louisiana Metros



Notes: Shading indicates the share of marginal movers from Louisiana metros that would relocate to each destination state following a hypothetical tax increase. Darker shading represents higher substitution elasticities. Calculated as: $\frac{dL_{state'}/dw_{state}}{-dL_{state}/dw_{state}}$ where $state = \text{Louisiana}$.

For an origin location j and a destination location $j' \neq j$, the diversion ratio is

$$d_{jj'} = \frac{\sum_i \left(\frac{\nu_i}{w_{ij}} \right) p_{ij} p_{ij'}}{\sum_{i'} \left(\frac{\nu_{i'}}{w_{i'j}} \right) p_{i'j} (1 - p_{i'j})},$$

By construction, $\sum_{j' \neq j} d_{jj'} = 1$, so the diversion ratios for any origin j form a probability distribution over destinations. The first term of the optimal-tax formula, call it $\bar{\tau}_j$ for convenience, is the diversion-weighted average tax of the substitutes of j ,

$$\bar{\tau}_j = \sum_{j' \neq j} d_{jj'} \tau_{j'}.$$

Because the tax vector $\{\tau_{j'}\}$ is common to all origins, any variation in $\bar{\tau}_j$ must come entirely from the diversion weights $d_{jj'}$. The remainder of this subsection characterizes how those diversion weights are structured geographically and how far they depart from the IIA benchmark under which every origin would face nearly the same $\bar{\tau}_j$. We return to the distribution of $\bar{\tau}_j$ itself once optimal taxes have been computed.

Before turning to these empirical patterns, it is useful to distinguish how precisely the pairwise diversion ratios and the average substitute tax can be estimated in finite samples. When j and j' are rarely co-listed, the numerator of $d_{jj'}$ contains relatively few nonzero contributions, so a given pairwise diversion ratio may be estimated imprecisely. By contrast, the average substitute tax need not be constructed by first treating all $J-1$ pairwise diversion

ratios as separately estimated intermediate parameters. For a fixed candidate tax vector, the average substitute tax can instead be written directly as the origin-level ratio

$$\bar{\tau}_j = \frac{\sum_i \left(\frac{\nu_i}{w_{ij}} \right) p_{ij} \sum_{j' \neq j} p_{ij'} \tau_{j'}}{\sum_{i'} \left(\frac{\nu_{i'}}{w_{i'j}} \right) p_{i'j} (1 - p_{i'j})},$$

where probabilities for locations j' outside a respondent's menu are set to zero. Conditional on a fixed candidate tax vector and under standard regularity conditions for the respondent-level inputs, the sample analogue is a smooth ratio of averages over the number of sample respondents whose menus contain j (call this number n_j). The delta method therefore implies that the sample analogue of the model-implied $\bar{\tau}_j$ is $\sqrt{n_j}$ -consistent and asymptotically normal.

As an important point of comparison, we define the *IIA diversion ratio* as

$$d_{jj'}^{\text{IIA}} = \hat{n}_{j'} / \sum_{j'' \neq j} \hat{n}_{j''}$$

where $\hat{n}_j$ is the survey-implied population of location j . It is the diversion implied by any IIA model, including a standard logit: every origin location would substitute toward the same national distribution, implying the first term of the optimal tax formula would be nearly constant across locations.

The geographic structure of diversion: Table 1 presents the diversion matrix aggregated to eight categories defined by Census region crossed with population size (above or below 250,000). Each cell reports the share of diversion from the origin category (row) absorbed by the destination category (column); rows sum to one. The bold diagonal blocks show that substitution is overwhelmingly within-region: the within-region diversion share ranges from 44% (Northeast Large) to 75% (South small). There is also meaningful asymmetry in the off-diagonal elements. Large locations in every region draw substantial diversion from small locations in distant regions, while the reverse flows are much smaller. For the optimal tax formula, this means that each location's tax is anchored to a geographically distinct set of neighbors, and not the national average.

The similarity of origins and substitutes: The concentration visible in Table 1 raises a natural follow-up: along which dimensions are substitution pairs (i.e. the locations sending migrants to the locations receiving those migrants) similar? Table 2 answers this by comparing origin-destination pair characteristics under survey-implied diversion weights ($\hat{n}_j \times d_{jj'}$) and under the IIA diversion ratio ($\hat{n}_j \times d_{jj'}^{\text{IIA}}$). In summary, survey-implied substitution pairs

Table 1: Diversion matrix by region and location size

	Northeast		Midwest		South		West		
	Small	Large	Small	Large	Small	Large	Small	Large	Total
NE, small/pooled	0.129	0.416	0.018	0.046	0.050	0.221	0.037	0.081	1.000
NE, large	0.128	0.311	0.029	0.043	0.053	0.262	0.035	0.139	1.000
MW, small/pooled	0.011	0.036	0.189	0.365	0.086	0.097	0.127	0.090	1.000
MW, large	0.016	0.038	0.223	0.255	0.061	0.240	0.044	0.122	1.000
S, small/pooled	0.015	0.037	0.052	0.043	0.241	0.509	0.035	0.068	1.000
S, large	0.022	0.082	0.016	0.064	0.183	0.490	0.044	0.098	1.000
W, small/pooled	0.017	0.037	0.055	0.045	0.060	0.131	0.214	0.440	1.000
W, large	0.015	0.077	0.026	0.059	0.048	0.154	0.186	0.434	1.000
<i>Pools</i>	14	23	19	29	42	75	36	32	270
<i>Pop. share</i>	0.023	0.161	0.050	0.144	0.069	0.306	0.047	0.201	1.000

Notes: Each cell reports the share of survey-implied diversion from the origin category (row) absorbed by the destination category (column). Rows sum to one. Categories are defined by Census region and size of location. “Large” = standalone single CBSA with Census population $\geq 250,000$; “Small” = standalone single CBSA below that threshold or any multi-CBSA pool of locations ($N < 10$ respondents at the CBSA level). Weights are $\hat{n}_j \times d_{jk}$.

are geographically much closer than the IIA diversion ratio would imply and the absolute changes in house prices, wages, demographics, and amenities between origin and destination are smaller under survey-implied weights for every characteristic but one (the environmental-quality index), confirming that people substitute toward locations that resemble their current one along multiple dimensions.

5.3 Survey Results: Average Marginal Utilities by Location

In this next section, we use our survey data to estimate the marginal utility of income across locations in the United States. Section 5.3.1 examines the cross-sectional relationship between marginal utility and location characteristics at baseline wages. Section 5.3.2 uses the experimental variation in location-income offers across scenarios to test for selection and non-separability.

5.3.1 Cross-Sectional Analysis of Marginal Utility at Baseline Wages

We estimate the variation in marginal utility of income across locations using survey data from the baseline. Specifically, we non-parametrically estimate location-specific marginal

Table 2: How similar are origins and their substitutes?

	Survey-implied	IIA
Same state	0.288	0.035
Same Census division	0.460	0.132
Distance (miles)	657.2	1114.2
<i>Absolute change in pool characteristics (origin to destination):</i>		
<i>Census:</i>		
log HPI	0.432	0.533
log population	1.460	1.530
Wage index	0.118	0.136
Hispanic share	0.122	0.158
NH White share	0.164	0.203
NH Black share	0.082	0.100
NH Asian share	0.043	0.051
<i>Diamond (2016) amenity indices (z-scores):</i>		
School quality	0.896	1.081
Retail	0.846	0.931
Crime	0.990	1.120
Road/transit	0.960	0.997
Environment	0.846	0.837
Jobs	0.886	0.995

Notes: Each row reports the weighted mean of the pair-level variable. Survey-implied weights are $\hat{n}_j \times d_{jj'}$; IIA weights are $\hat{n}_j \times d_{jj'}^{\text{IIA}}$. Amenity indices (school quality through jobs) are from [Diamond \(2016\)](#) and are available for a subset of metro pairs; weighted means for these rows are computed over non-missing pairs only.

utility values with a two-way fixed-effect-style framework that uses as a sampling weight the reported survey probability that respondent i lives in location j in the baseline, p_{ij0} :

$$\log(m_{ij}) = \alpha_i + \psi_j + \varepsilon_{ij} \quad (11)$$

where m_{ij} is as defined in equation (9), α_i are individual fixed effects that control for respondent heterogeneity, ψ_j are location fixed effects (our parameters of interest) and ε_{ij} is the error in this equation with standard deviation σ_ε .

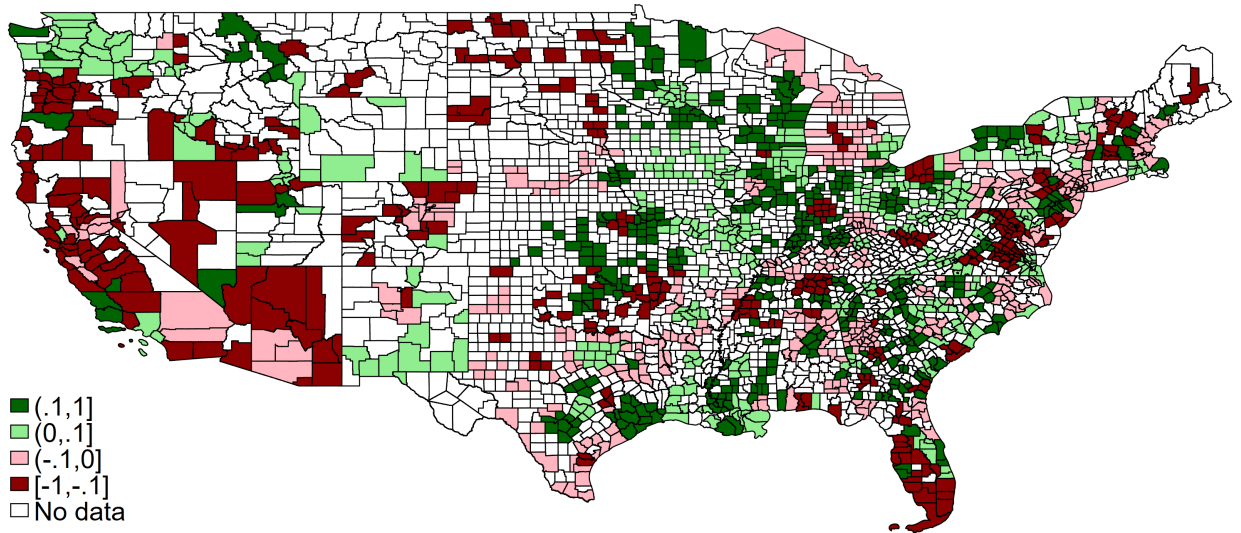
The estimates $\hat{\psi}_j$ provide a descriptive index of the common location component of reported marginal utility.²⁵ They are useful for summarizing spatial patterns in the baseline

²⁵ In AKM-style regressions of wages ([Abowd et al., 1999](#)) where ψ_j has the interpretation of a firm effect on wages, researchers explicitly adjust for selection of workers to firms such that estimates of ψ_j have the interpretation of a firm effect on wages for a worker selected at random. Here, we explicitly do not make such an adjustment, as we wish to measure average marginal utility in a location only among households

responses of marginal utility and for linking partially overlapping menus. That said, they are not the final location-level welfare weights that we use when we solve the planning problem to compute optimal taxes and transfers across locations. The object that enters the planning problem is the choice-probability-weighted location average μ_j , and we describe how we construct this object in Section 6.1.

Figure 3 below graphs our estimates of $\hat{\psi}_j$ for the 270 locations used in the policy analysis. These locations are CBSAs or, in some cases, collections of nearby CBSAs such that the probability-weighted sample size of the group is at least 10 people. The figure provides a first look at the spatial pattern in reported marginal utility. Coastal locations, particularly in the Northeast and California, tend to have low estimated marginal utility, while locations in the South, Southeast, and parts of the Midwest tend to have high marginal utility. Many states contain both high- and low-MU locations in close proximity.

Figure 3: Estimates of location marginal-utility fixed effects ($\hat{\psi}_j$)



As a check on whether limiting respondents to four locations on their location-choice menu was so restrictive as to affect results (as compared to menus including more locations), we recomputed $\hat{\psi}_j$ after eliminating the location in each person's menu with the lowest probability and proportionately rescaling the remaining probabilities in the menu to sum to 1.0. We then compared the resulting restricted estimates to the ones shown in Figure 3. The R^2 of a regression of the estimates shown in Figure 3 on the restricted estimates is 0.89 and the slope coefficient is 1.10, giving us some confidence that our results would not materially change if we were to have allowed for a larger number of locations in each respondent's menu.

that have selected into that location.

Next, we study how reported marginal utilities at the baseline correlate with observable characteristics. Specifically, we estimate the following probability-weighted regression using only baseline data:

$$\log(m_{ij}) = X'_{ij}\beta + \gamma_i + e_{ij},$$

where X_{ij} is a vector of location characteristics evaluated at baseline and γ_i is an individual fixed effect. Coefficient estimates and standard errors are reported in Table 3. Marginal utility is negatively correlated with local housing costs, positively correlated with the presence of family and friends and positively correlated with the “fun”-ness ranking of the area. Importantly, marginal utility is weakly but *positively* correlated with location-specific income at the baseline w_{ij} . Specifically, we find a 10 percent increase in income in a location is associated with an increase in marginal utility in that location of about 0.5 percent.

Table 3: Regression of $\log m_{ij}$ on location characteristics X_{ij}

$\log(w_{ij})$	0.053* (0.023)	Fun rank: 2	-0.173*** (0.019)
$\log(h_j)$ (house price index)	-0.142*** (0.019)	Fun rank: 3	-0.251*** (0.018)
f_{ij} (family/friends)	0.134*** (0.018)	Fun rank: 4	-0.327*** (0.018)
$\log(d_{ij})$ (dist. from hometown)	-0.016*** (0.003)	Constant	-7.482*** (0.251)
Observations	8,350		
Individual fixed effects	Yes		

Notes: The table reports coefficients from a within-person regression of log marginal utility on the listed location characteristics. Individual fixed effects absorb time-invariant heterogeneity. Standard errors in parentheses. * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$.

We explore this last result further by leveraging our experimental design where respondents evaluate multiple hypothetical scenarios with varying income levels and presence of family/friends across locations. The exogenous variation across scenarios removes endogeneity concerns, and the panel structure enables us to account for unobserved heterogeneity. We estimate the following regression:

$$\log(m_{ijk}) = \eta_1 \log w_{ijk} + \eta_2 f_{ijk} + \phi_{ik} + \phi_{ij} + e_{ijk} \quad (12)$$

where ϕ_{ik} and ϕ_{ij} are individual-by-scenario and individual-by-location fixed effects, respectively. Again, we estimate a *positive* elasticity of marginal utility with respect to expected

income of η_1 equal to 0.45 with a standard error of 0.01. As a check, we re-estimate equation (12) after allowing for each respondent to have their own coefficients on $\log w_{ijk}$ and f_{ijk} , η_{i1} and η_{i2} . After implementing an empirical Bayes shrinkage procedure to adjust for sampling variance, we find the mean value for η_{i1} is 0.43 with a standard deviation of 0.16.

5.3.2 Separability and Selection

We now discuss the implications of our finding of a positive correlation of income and marginal utility. Denote by $u_w(w_j, e_i)$ the marginal utility of income in location j for an individual with idiosyncratic shock e_i , and consider the average marginal utility of income conditional on choosing j . The derivative with respect to the wage in location j can be decomposed into a direct and a selection effect:

$$\frac{d}{dw_j} E[u_w | j \text{ chosen}] = \underbrace{\int \frac{\partial u_w(w_j, e_i)}{\partial w_j} f(e_i | j \text{ chosen}) de_i}_{\text{Direct effect}} + \underbrace{\int u_w(w_j, e_i) \frac{\partial f(e_i | j \text{ chosen})}{\partial w_j} de_i}_{\text{Selection effect}}$$

Furthermore, the selection effect can be re-written as:

$$\frac{1}{P_j(w_j)} \int \frac{\partial}{\partial w_j} \Pr(U_j \geq U_k | e_i) \left[u_w(w_j, e_i) - E(u_w | j \text{ chosen}) \right] f(e_i) de_i$$

The direct effect captures how marginal utility changes for a *fixed individual* in a given location when income in that location changes. The selection effect captures how the *composition of residents* changes as the location becomes more attractive and marginal movers enter or exit, which in turn changes the average marginal utility among those who choose the location. This pattern is consistent with selection effects dominating the direct effects that are typically negative in standard models. In other words, marginal entrants to higher-income locations have systematically higher marginal utility than inframarginal residents. The possibility of this kind of selection effect is not present in a model with separable shocks. This selection operates even within our within-respondent design. Because the marginal-utility question is conditional on the respondent having chosen the location, a scenario involving a change in income can change the latent circumstances under which that respondent imagines choosing a location. The within-respondent scenario response therefore combines a direct effect of income on marginal utility with a selection effect over the latent future states.

More broadly, several of our main empirical findings are difficult or impossible to rationalize with a discrete location choice model that specifies that utility from consumption is additively separable from location-specific shocks or attachment draws, such as $\log(w_{ij} - \tau_j) + \epsilon_{ij}$. First, in this model marginal utility is a decreasing function of net income. As a result, high-

wage locations must have lower average marginal utility. Second, within a given location individuals with the same wage will have the same marginal utility regardless of the strength of their attachment to the location.²⁶ Our empirical findings reject both of these restrictions.

6 Aggregating Survey Data

In this section, we document how to convert survey responses into the sufficient statistics that enter the planner’s solution. The survey identifies within-person ratios of marginal utilities, while the solution requires the average marginal utility of income of residents in every location. Section 6.1 constructs this object by normalizing individual marginal utilities, linking limited location-choice menus for all survey respondents using estimates of $\hat{\psi}_j$, appropriately aggregating to construct an average for each location, and then applying one empirical-Bayes precision adjustment to this average. Section 6.2 then describes how we update this average when taxes change after-tax incomes. We show that to appropriately update requires an auxiliary parameter that determines how an individual’s marginal utility changes in each location when after-tax income uniformly increases across all locations in their menu.

6.1 Aggregating individual MU

We face two challenges in using our survey data to construct the average marginal utility of income in every location. First, the optimal tax formula requires the average marginal utility of all residents in each location, μ_j , but the survey design yields individual-level measures m_{ij} that are only comparable within person. Second, and related, each survey respondent is asked questions about a small set of locations, their “menu,” such that constructing the average marginal utility of each location in *all* locations requires linking partially overlapping menus onto a common scale. We develop a procedure to overcome these challenges that satisfies the following three properties:

- (P1) **Scale invariance.** The procedure is invariant to rescaling $m_{ij} \mapsto s_i m_{ij}$, important because the survey design only identifies within-person ratios of marginal utilities across locations.
- (P2) **Preserve high-marginal-utility menus.** Some people may have a menu of locations that, when viewed relative to all possible locations, all or almost all have

²⁶These predictions do not change when utility is augmented to allow non-wage factors to affect location choice, for example by including additive linear terms in observable location characteristics, or including an additive location amenity parameter.

relatively high or relatively low average marginal utilities. Because menus vary across respondents, the procedure does not mechanically force everyone to have the same average marginal utility.

- (P3) **No spurious ex ante redistribution.** If two people have the same menu, baseline choice probabilities, and Pareto weight, our procedure ensures the planner values an additional dollar in their hands equally before any future location choice is determined.

We start by estimating equation (11) (repeated here for convenience) using the baseline survey data to obtain estimates of ψ_j : $\log(m_{ij}) = \alpha_i + \psi_j + \varepsilon_{ij}$. Note that because our survey design identifies ratios, taking logs converts within-person comparisons into additive differences. α_i absorbs the person-specific scale s_i since $\log(s_i \cdot \mu_{ij}) = \log s_i + \log \mu_{ij}$ and ψ_j captures the systematic metro component of marginal utility common across people. Even though the set of locations in each respondent's menu varies, the overlap of locations across menus of all respondents allows us to estimate the full set of values of ψ_j .

After estimating the full set of ψ_j , we (i) normalize within person to remove the arbitrary scale s_i , (ii) reattach a common level to each person in the sample using the estimates of ψ_j , and (iii) average across people with choice-probability and Pareto weights. The following proposition shows this procedure satisfies P1-P3.

Proposition 6.1 (Aggregation to location-specific welfare weights). *We make two assumptions:*

- **(A1) Overlap / connectedness.** *The network of locations, with an edge between any two locations that co-occur in some respondent menus, is connected.*
- **(A2) Additive error structure.** $\mathbb{E}[\varepsilon_{ij} \mid \alpha_i, \psi_j] = 0$ for all (i, j) with $j \in \mathcal{M}_i$.²⁷

Define:

$$m_{ij}^{\text{rel}} \equiv \frac{m_{ij}}{\sum_{\ell \in \mathcal{M}_i} p_{i\ell} m_{i\ell}}, \quad \kappa_i \equiv \sum_{\ell \in \mathcal{M}_i} p_{i\ell} \exp(\psi_\ell), \quad g_{ij} \equiv \Pi_{t(i)} m_{ij}^{\text{rel}} \kappa_i,$$

where $\Pi_{t(i)}$ is the planner's Pareto weight for person i 's type. Then g_{ij} satisfies Properties

²⁷This assumption only applies to pairs where $j \in \mathcal{M}_i$. It places no restriction on pairs where $j \notin \mathcal{M}_i$, because those have $p_{ij} = 0$ and enter neither estimation nor policy objects. As discussed earlier in footnote 25, this is different from the standard exogeneity concern in the labor literature, where each worker is observed at only one of many possible firms, so sorting on match-specific productivity biases firm effects. Here, the researcher would observe m_{ij} at every menu location – the complete set with positive probability.

P1-P3, and the location-average welfare weight that enters the optimal tax formula is

$$\mu_j = \frac{\sum_i p_{ij} g_{ij}}{\sum_{i'} p_{i'j}}.$$

The proof, which verifies properties P1-P3, is in online Appendix H.

We implement this aggregation procedure using baseline survey responses, i.e. $k = 0$. For each respondent i and menu location $j \in \mathcal{M}_i$, we compute the within-person normalized marginal utility m_{ij}^{rel} , the scale κ_i using our estimates $\hat{\psi}_j$, and the Pareto-weighted welfare weight $g_{ij} = \Pi_{g(i)} m_{ij}^{\text{rel}} \kappa_i$, where $\Pi_{g(i)}$ is the planner's weight for respondent i 's skill group. A starting estimate of location-level average marginal utility is then

$$\hat{\mu}_j = \frac{\sum_i p_{ij0} g_{ij}}{\sum_{i'} p_{i'j0}},$$

where p_{ij0} is respondent i 's baseline stated probability of choosing j , serving as the analogue of the population share residing in location j .

6.1.1 Precision adjustment of aggregated location marginal utilities

This initial estimate of $\hat{\mu}_j$ is measured with unequal precision because some locations appear in many respondent menus and others in relatively few. To avoid spurious redistribution across locations arising from imprecise estimates of $\hat{\mu}_j$, we therefore apply a precision adjustment that follows the logic of empirical-Bayes shrinkage and small-area estimation (Efron and Morris, 1975; Morris, 1983; Fay and Herriot, 1979). This adjustment, described next, is *not* applied to the individual values of m_{ij} , the normalized ratios m_{ij}^{rel} , or the index $\hat{\psi}_j$.

Let L_j denote a location's effective headcount, $\bar{\mu}$ as the population-weighted average value of μ_j , and $\hat{x}_j$ as the log difference of μ_j and $\bar{\mu}$, such that

$$L_j \equiv \sum_i p_{ij0}, \quad \bar{\mu} \equiv \frac{\sum_j L_j \hat{\mu}_j}{\sum_j L_j}, \quad \hat{x}_j \equiv \log \left(\frac{\hat{\mu}_j}{\bar{\mu}} \right).$$

We model $\hat{x}_j$ as a noisy estimate of the true log relative marginal utility of location j , x_j : $\hat{x}_j = x_j + \varepsilon_j$. We assume that x_j is Normally distributed with mean 0 and variance σ_x^2 . The

posterior mean is

$$\hat{x}_j^{EB} = \omega_j^\mu \hat{x}_j \quad \text{where} \quad \omega_j^\mu = \frac{\sigma_x^2}{\sigma_x^2 + (\sigma_\varepsilon^2/n_j^e)}.$$

$n_j^e = (\sum_i p_{ij0})^2 / \sum_i p_{ij0}^2$ is the Kish effective sample size for location j (Kish, 1965). If we define $\varepsilon_{ij} \sim N(0, \sigma_\varepsilon^2)$ as the respondent-level deviations from x_j , then ε_j can be written as $\sum_{i \in j} p_{ij0} \varepsilon_{ij} / \sum_{i \in j} p_{ij0}$, such that $\sigma_\varepsilon^2/n_j^e$ is equal to the variance of ε_j . We estimate $\sigma_x^2 = 0.048$ and $\sigma_\varepsilon^2 = 0.19$, such that the shrinkage factor is 0.5 at $n_j^e = \sigma_\varepsilon^2/\sigma_x^2 \approx 4.0$, below the median effective sample size of about 9.5.

The precision-adjusted location marginal utility is $\hat{\mu}_j^{EB} = \bar{\mu} \exp(\hat{x}_j^{EB})$. We implement the adjustment as a location-specific multiplier on the individual weights:

$$g_{ij}^{EB} \equiv s_j^\mu g_{ij} \quad \text{where} \quad s_j^\mu \equiv \hat{\mu}_j^{EB} / \hat{\mu}_j$$

and this implementation preserves aggregation exactly:

$$\frac{\sum_i p_{ij0} g_{ij}^{EB}}{\sum_{i'} p_{i'j0}} = \hat{\mu}_j^{EB}$$

All policy calculations below use g_{ij}^{EB} and $\hat{\mu}_j^{EB}$. This is the only shrinkage of marginal utility used in the policy analysis, and we describe in section 7.2 how this shrinkage affects our results.

6.2 Updating marginal utility under various policies

When taxes change after-tax incomes, marginal utility may change as well, making the welfare weights entering the optimal-tax formula endogenous to the policy being evaluated (Saez and Stantcheva, 2016). Taxes affect the relative profile of marginal utility across a respondent's menu of locations, which our survey identifies, and the common level across the menu, which is not identified and which we calibrate from the risk-aversion literature. For each individual i and given a candidate tax vector τ , define after-tax income as $y_{ij}(\tau) = w_{ij0} - \tau_j$. Given this, define average pre-policy income at baseline y_i^0 and average post-policy income evaluated at baseline choice probabilities $y_i(\tau)$ for individual i as

$$\bar{y}_i^0 \equiv \sum_{\ell \in M_i} p_{i\ell 0} w_{i\ell 0} \quad \bar{y}_i(\tau) \equiv \sum_{\ell \in M_i} p_{i\ell 0} y_{i\ell}(\tau).$$

Now decompose the change in log after-tax income for individual i in location j into a component common to all menu locations and a location-specific relative component, adding and subtracting $\log \bar{y}_i(\tau)$ and $\log \bar{y}_i^0$:

$$\log y_{ij}(\tau) - \log w_{ij0} = [\log y_{ij}(\tau) - \log \bar{y}_i(\tau)] + [\log \bar{y}_i(\tau) - \log \bar{y}_i^0] - [\log w_{ij0} - \log \bar{y}_i^0].$$

Collecting terms,

$$\log y_{ij}(\tau) - \log w_{ij0} = \underbrace{\log \bar{y}_i(\tau) - \log \bar{y}_i^0}_{\text{menu-common}} + \underbrace{\left[\log \left(\frac{y_{ij}(\tau)}{\bar{y}_i(\tau)} \right) - \log \left(\frac{w_{ij0}}{\bar{y}_i^0} \right) \right]}_{\text{relative}}.$$

The first term captures how taxes shift the overall level of after-tax income across the respondent's menu; the second captures how taxes change the relative position of location j within that menu. Exponentiate and define the menu-common and within-menu relative income ratios,

$$\mathcal{C}_i(\tau) = \frac{\bar{y}_i(\tau)}{\bar{y}_i^0} \quad \mathcal{R}_{ij}(\tau) = \frac{y_{ij}(\tau)/\bar{y}_i(\tau)}{w_{ij0}/\bar{y}_i^0}.$$

Now define the relative policy-updated marginal-utility profile as

$$\tilde{m}_{ij}^{rel}(\tau) = \frac{m_{ij}^{rel,0} \mathcal{R}_{ij}(\tau)^{\eta_i}}{\sum_{\ell \in M_i} p_{i\ell 0} m_{i\ell}^{rel,0} \mathcal{R}_{i\ell}(\tau)^{\eta_i}},$$

where $m_{ij}^{rel,0}$ is as defined earlier. Then (including the precision factor s_j^μ and letting the superscript EB denote the Empirical Bayes precision-adjusted estimates) we compute the updated weights and the location marginal utility that enters the optimal-tax formula as:

$$g_{ij}^{EB}(\tau) = s_j^\mu \Pi_{t(i)} \kappa_i^0 [\mathcal{C}_i(\tau)]^{-\varphi_i} \tilde{m}_{ij}^{rel}(\tau) \quad \mu_j^{EB}(\tau) = \frac{\sum_i p_{ij}(\tau) g_{ij}^{EB}(\tau)}{\sum_{i'} p_{i'j}(\tau)}.$$

The normalization in the relative profile $\tilde{m}_{ij}^{rel}(\tau)$ ensures it captures only the relative tilt implied by differential income changes across locations; the common-direction shift is absorbed by $[\mathcal{C}_i(\tau)]^{-\varphi_i}$ in the counterfactual welfare weights. The two key parameters here are η_i , which determines the impact of a within-menu shift of income, and φ_i , which determines the impact of a common shift of income. Our survey identifies η_i ; we set it to its median shrunk estimate, 0.43. The level parameter φ_i is not identified by our survey, but we set $\varphi_i = 1$ in the baseline based on estimates in the risk-aversion literature (e.g., [Chetty, 2006](#);

Barsky et al., 1997) and examine sensitivity of our results to this assumption in Section 7.2.

Note that if taxes change all locations in a respondent’s menu by a similar amount, $\mathcal{R}_{ij}(\tau)$ is close to 1 for every j , so the relative profile $\tilde{m}_{ij}^{rel}(\tau)$ changes little, while the common shift $\mathcal{C}_i(\tau)$ is large and, through $[\mathcal{C}_i(\tau)]^{-\varphi_i}$, scales all of the respondent’s welfare weights in the same direction. Conversely, when taxes differ sharply across menu locations, the relative channel tilts the profile toward locations where after-tax income has fallen relative to the menu average. Because updating operates at the individual level and is then aggregated into $\mu_j^{EB}(\tau)$, a location that appears on menus where all options are taxed is updated mainly through the common channel, whereas one appearing on menus with a mix of taxed and subsidized locations is updated mainly through the relative channel.

7 Survey-Implied Optimal Policy and Welfare

We now characterize optimal place-based policy using the results of our survey. In computing optimal taxes we specify that post-tax income must exceed \$10,000 for all households.²⁸ In our baseline specification of parameters, we set the planner’s Pareto weight on all people $\Pi_i = 1$; the parameters governing the relationship of marginal utility to income $\varphi = 1$ and $\eta_i = 0.43$ for everyone; and for the elasticity of location choice with respect to the wage, for each person in the survey we use that person’s shrunk estimate of ν_i from section 5.2.1. In section 7.2 we examine the sensitivity of our results to these choices. Note, importantly, that we hold each respondent’s set of future possible locations fixed while computing optimal policy; where this might matter, for example, is that respondents might wish to add a location to their menu if the planner allocates very large subsidies to that location.

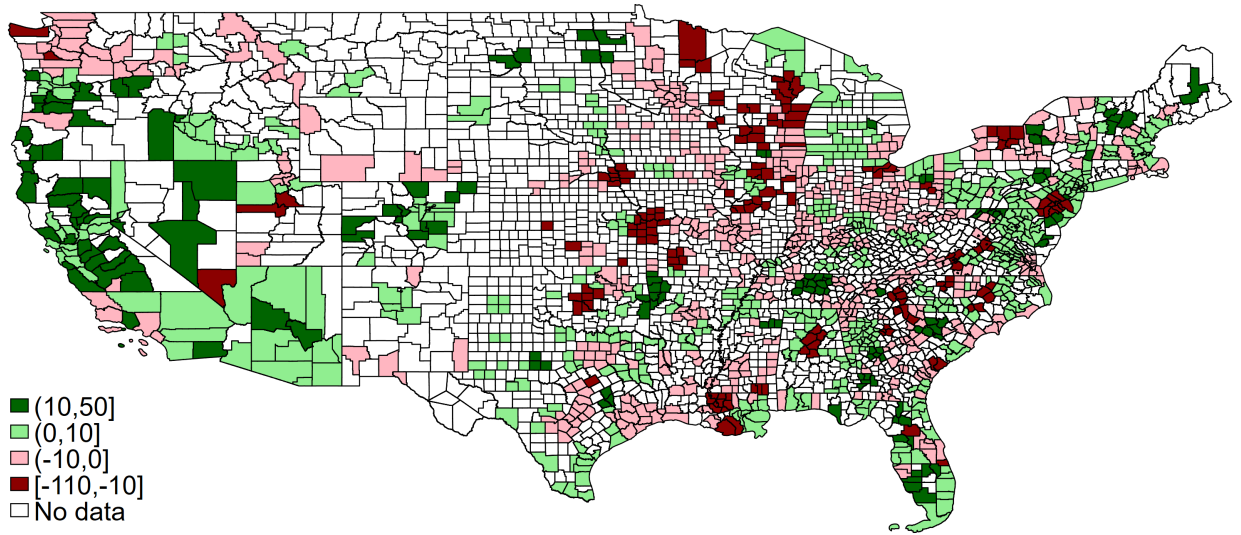
Figure 4 maps the geographic distribution of optimal location-level taxes across the continental United States. We have freely made available online a file containing the precision-adjusted location marginal utility entering the planner’s problem in log deviations from the mean, $\log(\hat{\mu}_j^{EB}/\bar{\mu})$, the optimal policy τ_j , the average value of τ_j for substitute locations, the log wage index we create using 2022 ACS data, and the number of survey respondents that have this location in their menu for all 270 locations we consider.²⁹ In summary, the distribution of optimal taxes has a standard deviation of \$11.5 thousand; the median tax is \$0.9 thousand and the interquartile range is from a subsidy of \$5.6 thousand to a tax of \$7.7 thousand. The map shows clear regional patterns. Coastal locations along the East Coast, California, and the Pacific Northwest tend to face higher optimal taxes (dark green). Many

²⁸Additionally, for the purposes of computing optimal policy, we winsorize (only from below) survey respondents’ direct responses on expected future income in each location in their menu. Specifically, we specify that pre-tax income must exceed \$30,000 for non-college graduates and \$60,000 for college graduates.

²⁹The file is available at <http://morris.marginalq.com/files/ADG-Location-Data.xlsx>

locations in the South, coastal Southeast, and portions of the Midwest receive subsidies (dark red). The spatial clustering visible in the map – neighboring locations tend to have similar tax assignments – reflects the fact that nearby locations serve as close substitutes and therefore have correlated optimal taxes. Also note that several large, high-wage locations in the South receive subsidies, while some smaller Northeast locations face positive taxes. This divergence from a naive “tax the rich” prescription reflects the fact that optimal transfers depend on marginal utility of income conditional on location choice, not merely on income levels.

Figure 4: Geographical variation in optimal taxes



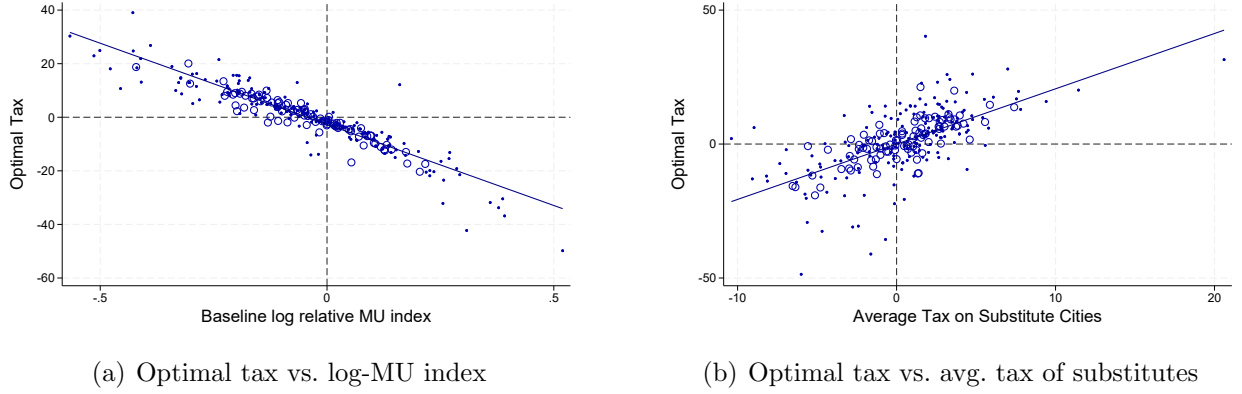
Notes: The color scheme distinguishes locations receiving subsidies (dark red: $[-110, -10]$ thousand; light red: $[-10, 0]$ thousand) from those facing positive taxes (light green: $[0, 10]$ thousand; dark green: $[10, 50]$ thousand). White areas indicate locations for which we lack sufficient survey coverage to compute optimal taxes.

Figure 5 below shows how optimal taxes vary with the two key ingredients of the formula determining optimal taxes. Panel (a) plots optimal taxes against the baseline log relative location marginal utility, $\log(\hat{\mu}_j^{EB,0}/\bar{\mu})$ – the EB-adjusted aggregated location MU that enters the planner’s problem (abstracting from $M(\tau)$, which is almost exactly 1.0). The scatter reveals a clear negative relationship. Since $\partial L_j / \partial \tau_j < 0$, the MU deviation term pushes optimal taxes below the substitute-tax average for high-MU locations and above it for low-MU locations.

Panel (b) plots each location’s optimal tax against the migration-weighted average tax of its substitutes, where the weights reflect how individuals leaving location j would reallocate across destinations. A key question is whether this term varies meaningfully across origins.

Under a standard logit, marginal movers from any origin substitute toward destinations in proportion to aggregate shares, so the average substitute tax would be nearly identical for every origin. In contrast, our menu-based design reveals strong geographic clustering in substitution patterns, so the implied destination mix differs sharply across origins. The figure confirms that a location’s optimal tax comoves with its average substitute tax. The average substitute tax varies enough across locations for fiscal externalities through migration to be an important source of cross-location differences in optimal taxes.

Figure 5: Optimal taxes and the two ingredients of the optimal-tax formula



Notes: Each point is a location (CBSA or group of CBSAs); the optimal tax is computed using the baseline specification of parameters. Panel (a) plots each unit’s optimal tax against its baseline log relative location marginal utility, $\log(\hat{\mu}_j^{EB,0}/\bar{\mu})$, the “log-MU index”; panel (b) plots it against the choice-probability-weighted average optimal tax of substitute locations. Larger markers denote locations with at least 30 respondents and smaller markers denote locations with 10-29 respondents; the line is a respondent-count-weighted linear fit. Taxes are in \$000s.

7.1 Analysis of Optimal Policy

We start our analysis by asking how much optimal place-based policy can improve average welfare. Throughout this section we keep $\Pi_i = 1$ and $\varphi = 1$ for everyone and set the value of ν_i for each survey respondent to its shrunk value from section 5.2.1. Our baseline specification of parameters, in column (1) of Table 4, combines the EB-adjusted aggregated location marginal utilities constructed in Section 6.1 with the survey-revealed location-choice and substitution patterns. The other columns impose the conventional homogeneous log consumption utility (“Log-C”) benchmark in columns (3)-(4); and the even-numbered columns (2 and 4) replace the survey-revealed substitution patterns with proportional (“IIA”) diversion.

The headline result, shown in Panel A of column (1), is that the optimal policy under the

baseline specification of parameters increases average expected utility by \$343 per household, corresponding to 0.35% of average income of our sample respondents (\$99 thousand). This net benefit is the difference of a gross welfare gain from transfers evaluated at the optimal policy (marked “INS” in the table) of \$698 per household and a deadweight cost from policy-induced migration (marked “EDC” in the table) of \$355.³⁰ Panel C shows the amount of redistribution required to achieve these welfare gains: the standard deviation of transfers is \$11.5 thousand across locations and \$8.8 thousand across people.

Table 4: Welfare from optimal place-based policy under various assumptions

	Shrunk MU (preferred)		Log-C	
	Survey subst. (1)	IIA (2)	Survey subst. (3)	IIA (4)
<i>Panel A. What the planner thinks they are getting</i> (welfare evaluated under each policy’s own optimal criterion)				
ΔW (\$/household)	+\$343 (+0.35%)	+\$527 (+0.53%)	+\$83 (+0.08%)	+\$111 (+0.11%)
Insurance value (INS)	+\$698	+\$1,058	+\$168	+\$213
Deadweight cost (EDC)	+\$355	+\$530	+\$85	+\$102
<i>Panel B. What they actually get</i> (welfare evaluated assuming the baseline specification of parameters)				
ΔW (\$/household)	+\$343 (+0.35%)	+\$40 (+0.04%)	-\$212 (-0.21%)	-\$51 (-0.05%)
Insurance value (INS)	+\$698	+\$207	-\$127	-\$18
Deadweight cost (EDC)	+\$355	+\$167	+\$85	+\$33
<i>Panel C. Optimal-tax dispersion</i>				
SD(τ) across locations	\$11.5k	\$6.9k	\$4.8k	\$2.7k
SD(τ), pop.-weighted	\$8.8k	\$6.2k	\$5.4k	\$3.1k

Notes: Each column is a planner defined by a marginal-utility assumption and a demand system; the entries are the per-household welfare gain $\Delta W = \text{INS} - \text{EDC}$ (\$/household, with % of mean baseline income in parentheses) from that planner’s optimal location-tax schedule. *Panel A* (“*what the planner thinks they are getting*”) evaluates each schedule under the same criterion used to derive it – the “diagonal” of a cross-evaluation grid – so the columns are not graded on a common metric. *Panel B* (“*what they actually get*”) re-scores all four schedules using the baseline specification of parameters. *Panel C* reports the cross-location dispersion of the optimal taxes. Columns (1)-(2) (“Shrunk MU”) use the EB-adjusted aggregated location marginal utility $\hat{\mu}_j$ (see Section 6.1.1); columns (3)-(4) (“Log-C”) impose homogeneous log consumption utility, $1/e^{\log w_j}$. Within each pair the first column uses survey-revealed substitution and the second imposes IIA choices. For IIA choices, income at every location is set to the respondent’s location-choice-probability weighted average income, adjusted by $\log w_j$. SD(τ) is computed across the policy locations; the population-weighted version uses Census population at the location level. Calculations use the 2,070 respondents with complete wage-index coverage over a 270-location choice set. Mean baseline income is about \$99k.

Panel A reports what each planner *believes* it achieves, scoring its policy under its own assumptions; Panel B re-scores all four policies under the single *true* baseline metric, so

³⁰See online Appendix I for a discussion on how we compute INS and EDC.

values reported in Panel A and Panel B coincide in column (1). The value of the survey is most apparent when comparing column (1) to column (4). Column (4) shows results when a planner sets optimal place-based policy assuming homogeneous log consumption utility (“Log-C”) and substitution patterns consistent with IIA diversion.³¹ These are the assumptions commonly made in calibrated urban models, so column 4 illustrates what researchers would typically report for the benefits and scope of place-based policy after abstracting from agglomeration or other externalities. Panel A shows that the expected researcher-reported benefits of tax policy would be *much* lower, \$111 per person as compared to \$343 in column 1 (a 68% relative decline), and corresponding amount of optimal redistribution across locations also much lower, a standard deviation of \$2.7 thousand as compared to \$11.5 thousand in column 1 (76% relative decline). Panel B shows the “real-world” consequence of implementing the column-4 policy: evaluated using our survey-based estimates of marginal utilities and substitution patterns, average welfare would *decline* by \$51 per person. The cross-location Pearson correlation of the optimal policies from columns 1 and 4 is only 0.20. The two policies agree on the broad sign for many locations but disagree on magnitudes throughout, and the survey-based policy assigns substantially larger subsidies and modestly larger taxes than a planner when assuming log utility and IIA.

Columns 2 and 3 help evaluate what is causing the gap in results between 1 and 4. Column 2 evaluates policy if the planner uses our shrunk estimates of marginal utility but assumes IIA substitution patterns, and column 3 evaluates policy when the planner assumes Log-C preferences but assumes the substitution patterns implied by our survey. Shown by comparing columns (1), (2) and (3) of Panel A, most of the planner perceived welfare gains are attributable to planner exploiting the estimated variation in marginal utility across locations implied by the survey. That said, replacing survey substitution with IIA cuts the standard deviation of optimal transfers by about 40 percent, from \$11.5 thousand (column 1) to \$6.9 thousand (column 2): the planner assuming IIA chooses far less redistribution. This is because our survey reveals strong geographic clustering in where people would relocate – the destinations a taxed location loses migrants to are themselves typically taxed – so co-taxing close substitutes generates smaller deadweight losses from migration than the diffuse, proportional diversion that IIA assumes, allowing the planner to redistribute more aggressively for the same efficiency cost.³²

³¹Under the assumptions in column 4, optimal policy collapses to the textbook case of a single constant tax *rate* on income rebated lump-sum: with $\nu_1 = 1.73$ that rate is $1/(1+\nu_1) = 36.6\%$. Under IIA the planner must attach an after-tax income and an after-tax marginal utility to every respondent-location pair; online Appendix J describes how we do this.

³²With IIA mobility, when a location is taxed there are lots of locations a person will consider moving to that are not being taxed. When location decisions are clustered, and when taxes and transfers reflect that clustering, if a person is considering leaving a location because it is being taxed it is likely the locations the

7.2 Sensitivity of Results

We now briefly discuss the sensitivity of our main findings to choices we have made about various parameters; online Appendix K contains a more fulsome discussion. We first study what would happen to predicted optimal taxes if (a) each respondent’s location-choice elasticity was twice the value we estimated or (b) if the sensitivity of marginal utility to the average level of income was double what we assumed. In both of these scenarios, the sensitivity of optimal taxes to differences in the average level of marginal utility falls, and the standard deviation of taxes and transfers falls. This occurs in scenario (a) because location choices are assumed more responsive to taxes than in the baseline, so deadweight losses increase more rapidly than in the baseline as redistribution rises. In scenario (b), this occurs because marginal utilities are more sensitive to income than assumed in the baseline, so less redistribution is required to equalize marginal utilities.

We then study what would happen if the planner uses the *non-shrunk* respondent-by-location weights and aggregated location marginal utilities – that is, without the empirical-Bayes precision adjustment of Section 6.1.1 – holding all other ingredients at baseline values. Because the precision adjustment pulls imprecisely estimated locations toward the mean, omitting this adjustment inflates the dispersion of measured marginal utility and hence of optimal taxes: the standard deviation of optimal taxes nearly doubles and the optimal tax becomes substantially more sensitive to measured marginal utility. This is the potentially excess redistribution the precision adjustment is designed to attenuate.

We next consider a case where the planner assigns a Pareto weight of 4 to non-college graduates (holding fixed the Pareto weight of 1 for college graduates).³³ This scenario is designed to capture a planner wishing redistribution to more heavily factor imperfectly signaled permanent income potential. In this scenario, the distribution of taxes changes (depending on the prevalence of college graduates in each location) and the amount of redistribution nearly doubles relative to the baseline.

We finally study optimal redistribution when we replace the balanced-budget requirement with a net-revenue target of \$3,000 per household ($G = 3$). Because the optimal-tax formula pins down only the *relative* taxes across locations and the budget condition sets their common level, raising G shifts every location’s tax upward by an almost exactly uniform \$3,010 and leaves the cross-location *pattern* unchanged. Even with the pattern of taxes and transfers unchanged, the higher tax level raises the migration deadweight cost and lowers the insurance value, so the welfare gain the planner expects from optimal redistribution declines.

person is considering moving to are also being taxed.

³³We label the household as a college graduate or not based on the education status of the survey respondent.

8 Conclusion

This paper develops and implements a new approach to measuring a central input to optimal place-based policy: how the marginal utility of income varies across locations. We show that this object is not identified from location choice data alone without strong assumptions on preference separability, and the implied methods are practically infeasible to implement. Our solution is to measure marginal utilities directly through a survey that elicits within-person comparisons of the value of additional income across locations where each respondent could plausibly live. Our survey also recovers a rich matrix of substitution elasticities across locations needed to characterize the efficiency costs of spatial redistribution. The survey data allow us to reject two assumptions frequently made in the literature on place-based transfers: that the marginal utility of recent movers to a location is the same as the marginal utility of the existing population of that location (separability) and that when people move out from an area, they move to locations in proportion to the existing population (aggregate IIA). The optimal place-based policy implied by our survey results in an increase in the standard deviation of place-based taxes by more than a factor of four relative to the standard deviation of optimally-computed taxes arising from standard assumptions.

References

- Abowd, J.M., Creedy, R.H., Kramarz, F., 2002. Computing Person and Firm Effects Using Linked Longitudinal Employer-Employee Data. Technical Paper TP-2002-06. U.S. Census Bureau, Center for Economic Studies.
- Abowd, J.M., Kramarz, F., Margolis, D.N., 1999. High wage workers and high wage firms. *Econometrica* 67, 251–333.
- Alam, M.M.U., Mookerjee, M., Roy, S., forthcoming. Worker-side discrimination: Beliefs and preferences. *Quantitative Economics* .
- Allen, R., Rehbeck, J., 2019. Identification with additively separable heterogeneity. *Econometrica* 87, 1021–1054. doi:[10.3982/ECTA15867](https://doi.org/10.3982/ECTA15867).
- Andrew, A., Adams, A., 2025. Revealed beliefs and the marriage market return to education. *The Quarterly Journal of Economics* , qjaf020.
- Austin, B., Glaeser, E., Summers, L., 2018. Jobs for the heartland: Place-based policies in 21st-century america. *Brookings Papers on Economic Activity* , 151–232Spring.

- Barsky, R.B., Juster, F.T., Kimball, M.S., Shapiro, M.D., 1997. Preference parameters and behavioral heterogeneity: An experimental approach in the health and retirement study. *Quarterly Journal of Economics* 112, 537–579.
- Bartik, T.J., 1991. Who Benefits from State and Local Economic Development Policies? W.E. Upjohn Institute for Employment Research, Kalamazoo, MI.
- Berry, S., Levinsohn, J., Pakes, A., 1995. Automobile prices in market equilibrium. *Econometrica* 63, 841–890.
- Berry, S.T., Haile, P.A., 2014. Identification in differentiated products markets using market level data. *Econometrica* 82, 1749–1797.
- Blass, A.A., Lach, S., Manski, C.F., 2010. Using elicited choice probabilities to estimate random utility models: Preferences for electricity reliability. *International Economic Review* 51, 421–440.
- Bleemer, Z., Zafar, B., 2018. Intended college attendance: Evidence from an experiment on college returns and costs. *Journal of Public Economics* 157, 184–211.
- Blundell, R., Pistaferri, L., Preston, I., 2008. Consumption inequality and partial insurance. *American Economic Review* 98, 1887–1921.
- Borusyak, K., Hull, P., Jaravel, X., 2022. Quasi-experimental shift-share research designs. *The Review of Economic Studies* 89, 181–213.
- Bound, J., Holzer, H.J., 2000. Demand shifts, population adjustments, and labor market outcomes during the 1980s. *Journal of Labor Economics* 18, 20–54.
- Busso, M., Gregory, J., Kline, P., 2013. Assessing the incidence and efficiency of a prominent place based policy. *American Economic Review* 103, 897–947.
- Canay, I.A., 2011. A simple approach to quantile regression for panel data. *The econometrics journal* 14, 368–386.
- Carson, R.T., Louviere, J.J., 2011. A common nomenclature for stated preference elicitation approaches. *Environmental and Resource Economics* 49, 539–559.
- Chetty, R., 2006. A new method of estimating risk aversion. *American Economic Review* 96, 1821–1834.
- Chetty, R., 2009. Sufficient statistics for welfare analysis: A bridge between structural and reduced-form methods. *Annual Review of Economics* 1, 451–488.
- Dahl, G.B., 2002. Mobility and the return to education: Testing a Roy model with multiple markets. *Econometrica* 70, 2367–2420.

- Davis, M.A., Fisher, J.D., Whited, T.M., 2014. Macroeconomic implications of agglomeration. *Econometrica* 82, 731–764.
- Delavande, A., Manski, C.F., 2015. Using elicited choice probabilities in hypothetical elections to study decisions to vote. *Electoral Studies* 38, 28–37.
- Diamond, R., 2016. The determinants and welfare implications of u.s. workers’ diverging location choices by skill: 1980-2000. *The American Economic Review* 106, 479–524.
- Dominitz, J., Manski, C.F., 1997. Using expectations data to study subjective income expectations. *Journal of the American Statistical Association* 92, 855–867.
- Donald, E., Fukui, M., Miyauchi, Y., 2025. Optimal dynamic spatial policy. Technical Report. National Bureau of Economic Research.
- Donald, E., Fukui, M., Miyauchi, Y., 2026. Unpacking aggregate welfare in a spatial economy. *Review of Economic Studies* doi:[10.1093/restud/rdag021](https://doi.org/10.1093/restud/rdag021). advance access.
- Eaton, J., Kortum, S., 2002. Technology, geography, and trade. *Econometrica* 70, 1741–1779.
- Efron, B., Morris, C., 1975. Data analysis using stein’s estimator and its generalizations. *Journal of the American Statistical Association* 70, 311–319.
- Fajgelbaum, P.D., Gaubert, C., 2020. Optimal spatial policies, geography and sorting. *Quarterly Journal of Economics* 135, 959–1036.
- Fajgelbaum, P.D., Gaubert, C., 2025. Optimal spatial policies, in: Donaldson, D. (Ed.), *Handbook of Regional and Urban Economics*. Elsevier. volume 6, pp. 143–223.
- Fajgelbaum, P.D., Morales, E., Suárez Serrato, J.C., Zidar, O., 2019. State taxes and spatial misallocation. *The Review of Economic Studies* 86, 333–376.
- Fay, R.E., Herriot, R.A., 1979. Estimates of income for small places: An application of james-stein procedures to census data. *Journal of the American Statistical Association* 74, 269–277.
- Gaubert, C., Kline, P., Vergara, D., Yagan, D., 2025. Place-based redistribution. *American Economic Review* 115, 3415–3450.
- Goldsmith-Pinkham, P., Sorkin, I., Swift, H., 2020. Bartik instruments: What, when, why, and how. *American Economic Review* 110, 2586–2624.
- Heckman, J.J., Vytlacil, E., 2005. Structural equations, treatment effects, and econometric policy evaluation. *Econometrica* 73, 669–738.

- Heckman, J.J., Vytlacil, E.J., 2007. Econometric evaluation of social programs, part I: Causal models, structural models and econometric policy evaluation, in: Heckman, J.J., Leamer, E.E. (Eds.), *Handbook of Econometrics*. Elsevier. volume 6B, pp. 4779–4874.
- Imbens, G.W., Angrist, J.D., 1994. Identification and estimation of local average treatment effects. *Econometrica* 62, 467–475.
- Kaplan, G., Schulhofer-Wohl, S., 2017. Understanding the long-run decline in interstate migration. *International Economic Review* 58, 57–94.
- Kennan, J., Walker, J.R., 2011. The effect of expected income on individual migration decisions. *Econometrica* 79, 211–251.
- Kish, L., 1965. *Survey Sampling*. John Wiley & Sons, New York.
- Kline, P., Moretti, E., 2014. Local economic development, agglomeration economies and the big push: 100 years of evidence from the tennessee valley authority. *Quarterly Journal of Economics* 129, 275–331.
- Koşar, G., Ransom, T., van der Klaauw, W., 2022. Understanding migration aversion using elicited counterfactual choice probabilities. *Journal of Econometrics* 231, 123–147.
- Kuziemko, I., Norton, M., Stantcheva, S., Saez, E., 2015. How elastic are preferences for redistribution: Evidence from randomized survey experiments. *American Economic Review* 105, 1478–1508.
- Louviere, J.J., Hensher, D.A., Swait, J.D., 2000. *Stated Choice Methods: Analysis and Applications*. Cambridge University Press, Cambridge.
- Manski, C.F., 2004. Measuring expectations. *Econometrica* 72, 1329–1376.
- McFadden, D., 1974. Conditional logit analysis of qualitative choice behavior, in: Zarembka, P. (Ed.), *Frontiers in Econometrics*. Academic Press, New York, pp. 105–142.
- Molloy, R., Smith, C.L., Wozniak, A., 2011. Internal migration in the united states. *Journal of Economic Perspectives* 25, 173–196.
- Morris, C.N., 1983. Parametric empirical bayes inference: Theory and applications. *Journal of the American Statistical Association* 78, 47–55.
- Neumark, D., Simpson, H., 2015. Place-based policies, in: Duranton, G., Henderson, J.V., Strange, W.C. (Eds.), *Handbook of Regional and Urban Economics*. Elsevier. volume 5, pp. 1197–1287.
- Nevo, A., 2000. A practitioner’s guide to estimation of random-coefficients logit models of demand. *Journal of Economics & Management Strategy* 9, 513–548.

- Notowidigdo, M.J., 2020. The incidence of local labor demand shocks. *Journal of Labor Economics* 38, 687–725.
- Rossi-Hansberg, E., Sarte, P.D., Schwartzman, F., 2026. Cognitive hubs and spatial redistribution. *American Economic Journal: Macroeconomics* 18, 72–111. doi:[10.1257/mac.20230013](https://doi.org/10.1257/mac.20230013).
- Saez, E., Stantcheva, S., 2016. Generalized social marginal welfare weights for optimal tax theory. *American Economic Review* 106, 24–45.
- Suárez Serrato, J.C., Wingender, P., 2014. Estimating the Incidence of Government Spending. Technical Report. Duke University and the International Monetary Fund.
- Suárez Serrato, J.C., Zidar, O.M., 2015. Who Benefits from State Corporate Tax Cuts? A Local Labor Markets Approach with Heterogeneous Firms. Technical Report. Duke University and The University of Chicago Booth School of Business.
- Train, K.E., 2009. *Discrete Choice Methods with Simulation*. 2nd ed., Cambridge University Press, Cambridge.
- U.S. Census Bureau, 2020. March 2020 delineation files: Core based statistical areas, metropolitan divisions, and combined statistical areas. <https://www.census.gov/geographies/reference-files/time-series/demo/metro-micro/delineation-files.html>. U.S. Census Bureau.
- U.S. Census Bureau, 2022. American community survey, 2018–2022 five-year estimates. <https://www.census.gov/programs-surveys/acs>. U.S. Census Bureau.
- Wiswall, M., Zafar, B., 2015. Determinants of college major choice: Identification using an information experiment. *The Review of Economic Studies* 82, 791–824.
- Wiswall, M., Zafar, B., 2018. Preference for the workplace, investment in human capital, and gender. *The Quarterly Journal of Economics* 133, 457–507.

Online Appendix

A Derivation of Solution to the Planner's problem

This appendix derives the sufficient-statistics formula for optimal place based redistribution in absence of agglomeration externalities.

A.1 Regularity conditions

We impose the following regularity conditions. We first state a set of primitive conditions on preferences and choice that are shared by the optimal-tax characterization here and by the non-identification result of Appendix C.

Assumption 1 (Utility primitives). We assume the following standard primitives on the utility function

- Each individual's utility from a location is continuously differentiable in own income (net of the location-specific policy)
- Marginal utility of income is strictly positive and uniformly bounded for every realization of ϵ .
- The individual location-choice problem has a unique maximizer almost surely; equivalently, ties occur with probability zero.

The optimal-tax characterization additionally imposes the following regularity conditions.

Assumption 2 (Regularity). In addition to Assumption 1, the following conditions hold in an open neighborhood of the candidate optimum τ^* .

1. The distribution of potential incomes and location shifters, $\{w_{ij}, \epsilon_{ij}\}_{j=1}^J$, is fixed with respect to τ .
2. The derivatives of the individual indirect utility function are locally dominated by an integrable random variable, so that differentiation and expectation can be interchanged.
3. The expected indirect utility $U_t(\tau)$ is differentiable in τ , and the envelope theorem applies to the individual's location-choice problem.
4. The location choice probabilities $P_j^t(\tau)$ are continuously differentiable in τ . Hence $L_j(\tau) = \sum_{t=1}^T N_t P_j^t(\tau)$ is continuously differentiable in τ .
5. At the candidate optimum, each location has positive population: $L_j(\tau^*) > 0$ for all j

6. The planner's optimum is interior with respect to the tax domain, and the balanced-budget constraint satisfies a constraint qualification:

$$\nabla_{\tau} \left[\sum_{j=1}^J \tau_j L_j(\tau) \right]_{\tau=\tau^*} \neq 0. \quad (13)$$

A.2 Proof of Theorem 2.1

Given these regularity conditions, the planner chooses τ to maximize

$$W(\tau) = \sum_{t=1}^T \Pi_t N_t U_t(\tau)$$

subject to $\sum_{j=1}^J \tau_j L_j(\tau) = G$, where $L_j(\tau) = \sum_t N_t P_j^t(\tau)$ and $U_t(\tau)$ is expected utility for a type- t individual under taxes τ . For each type t , define expected indirect utility as

$$U_t(\tau) = \mathbb{E} \left[\max_{j \in \{1, \dots, J\}} u_{ij}^t(w_{ij} - \tau_j, \epsilon_{ij}) \mid t(i) = t \right].$$

Assume that, for each τ in a neighborhood of the planner's optimum, the individual location-choice problem has a unique maximizer almost surely, the utility functions are differentiable in after-tax income, and the relevant derivatives are integrable so that differentiation and expectation can be interchanged. Then the envelope theorem for the individual's location-choice problem implies

$$\frac{\partial U_t(\tau)}{\partial \tau_j} = -\mathbb{E} \left[u_w^t(w_{ij} - \tau_j, \epsilon_{ij}) \mathbf{1}\{j \text{ chosen}\} \mid t(i) = t \right].$$

Equivalently,

$$\frac{\partial U_t(\tau)}{\partial \tau_j} = -P_j^t(\tau) \mu_j^t(\tau),$$

where

$$\mu_j^t(\tau) = \mathbb{E} \left[u_w^t(w_{ij} - \tau_j, \epsilon_{ij}) \mid j \text{ chosen}, t(i) = t \right].$$

Let $R(\tau) \equiv \sum_{j=1}^J \tau_j L_j(\tau)$ denote tax revenue. The planner's Lagrangian is

$$\mathcal{L}(\tau, \Lambda) = W(\tau) + \Lambda (R(\tau) - G).$$

By the constraint qualification in Assumption 2, there exists a multiplier Λ such that the first-order

condition with respect to τ_j is

$$0 = \frac{\partial W(\tau^*)}{\partial \tau_j} + \Lambda \frac{\partial R(\tau^*)}{\partial \tau_j}. \quad (14)$$

By the envelope theorem for the individual location-choice problem,

$$\frac{\partial V_i^t(\tau)}{\partial \tau_j} = -u_{w,ij}^t(w_{ij} - \tau_j, \epsilon_{ij}) \mathbf{1}\{j \text{ chosen}\} \quad \text{a.s.}$$

The local domination condition allows us to interchange differentiation and expectation by dominated convergence. By the envelope theorem for the individual's location-choice problem,

$$\frac{\partial U_t(\tau)}{\partial \tau_j} = -\mathbb{E} [u_{w,ij}^t(c_{ij}(\tau), \epsilon_{ij}) \mathbf{1}\{j \text{ chosen}\} \mid t(i) = t] = -P_j^t(\tau) \mu_j^t$$

Therefore,

$$\frac{\partial W(\tau)}{\partial \tau_j} = -\sum_{t=1}^T \Pi_t N_t P_j^t(\tau) \mu_j^t(\tau) = -L_j(\tau) \mu_j(\tau).$$

The derivative of revenue is

$$\frac{\partial R(\tau)}{\partial \tau_j} = L_j(\tau) + \sum_{j'=1}^J \tau_{j'} \frac{\partial L_{j'}(\tau)}{\partial \tau_j}.$$

Substituting these two expressions into the planner's first-order condition (14) gives

$$0 = -L_j(\tau^*) \mu_j(\tau^*) + \Lambda L_j(\tau^*) + \Lambda \sum_{j'=1}^J \tau_{j'}^* \frac{\partial L_{j'}(\tau^*)}{\partial \tau_j}.$$

Define $D_{j'j}(\tau) \equiv \frac{\partial L_{j'}(\tau)}{\partial \tau_j}$ as how the share of population in j' changes as a result of increase in tax on location j . Dividing by $L_j(\tau^*)\Lambda$ when $\Lambda \neq 0$ and simplifying gives:

$$L_j(\tau^*) \frac{\mu_j(\tau^*) - \Lambda}{\Lambda} = D_{jj}(\tau^*) \tau_j^* + \sum_{j' \neq j} D_{j'j}(\tau^*) \tau_{j'}^*.$$

Since the choice set is closed and the economy has unit mass, $\sum_{j'=1}^J L_{j'}(\tau) = 1$ for all τ , and therefore $\sum_{j'=1}^J D_{j'j}(\tau) = 0$ for all j . Hence, we have $-D_{jj}(\tau^*) = \sum_{j' \neq j} D_{j'j}(\tau^*)$. If $D_{jj}(\tau^*) \neq 0$, then the first-order condition can be rearranged to solve for τ_j^* :

$$\tau_j^* = \frac{\sum_{j' \neq j} D_{j'j}(\tau^*) \tau_{j'}^*}{\sum_{j' \neq j} D_{j'j}(\tau^*)} + \frac{L_j(\tau^*)}{D_{jj}(\tau^*)} \left[\frac{\mu_j(\tau^*) - \Lambda}{\Lambda} \right] \quad (15)$$

If, in addition, $D_{jj}(\tau^*) < 0$ and $D_{j'j}(\tau^*) \geq 0$ for all $j' \neq j$, then the first term is a weighted average of the taxes in locations that gain residents when location j 's tax increases.

At the optimal solution, the revenue term gives

$$\sum_{j=1}^J \left[L_j(\tau^*) + \sum_{j'=1}^J \tau_{j'}^* \frac{\partial L_{j'}(\tau^*)}{\partial \tau_j} \right] = 1 + \sum_{j'=1}^J \tau_{j'}^* \sum_{j=1}^J \frac{\partial L_{j'}(\tau^*)}{\partial \tau_j} = M(\tau^*).$$

Thus the summed first-order condition is

$$0 = -\bar{\mu}(\tau^*) + \Lambda M(\tau^*).$$

If $M(\tau^*) \neq 0$, then

$$\Lambda = \frac{\bar{\mu}(\tau^*)}{M(\tau^*)}.$$

Substituting this expression for Λ into equation (15) gives the result.

B Solution Method and Computational Details

Each location's optimal tax depends on the taxes of its substitutes, which in turn depend on the original location's tax. We solve for the fixed point iteratively. The system can be written in matrix form as:

$$\mathbf{A}(\boldsymbol{\tau})\boldsymbol{\tau} = \mathbf{b} \tag{16}$$

where $\mathbf{A}$ depends on migration elasticities at the current tax vector and $\mathbf{b}$ contains the marginal-utility adjustments. Define the outflow share from location j to location j' as

$$s_{jj'} \equiv \frac{\partial L_{j'}/\partial \tau_j}{\sum_{j'' \neq j} \partial L_{j''}/\partial \tau_j},$$

which satisfies $s_{jj'} \geq 0$ and $\sum_{j' \neq j} s_{jj'} = 1$. Then $A_{jj} = 1$, off-diagonal elements are

$$A_{jj'} = -s_{jj'} \text{ for } j' \neq j,$$

and the vector $\mathbf{b}$ has elements

$$b_j = \frac{L_j}{\frac{\partial L_j}{\partial \tau_j}} \cdot \frac{M(\boldsymbol{\tau})\mu_j - \bar{\mu}}{\bar{\mu}},$$

Here $M(\boldsymbol{\tau})$ is the marginal-revenue factor from Theorem 2.1: the revenue raised by a uniform increase in all location head taxes once the re-sorting it induces is accounted for, so that the

threshold marginal utility is $\Lambda(\boldsymbol{\tau}) = \bar{\mu}/M(\boldsymbol{\tau})$. In our data $M(\boldsymbol{\tau})$ is within 0.4% of one at every optimum, so the correction is quantitatively small, but we impose it exactly rather than assume $M \equiv 1$.

Because $A_{jj} = 1$ and the off-diagonal diversion shares satisfy $\sum_{j' \neq j} s_{jj'} = 1$, every row of $\mathbf{A}$ sums to zero, so $\mathbf{A}$ is singular: the first-order conditions determine only *relative* taxes across locations, not their common level. The level is pinned down by the balanced-budget constraint $\sum_j \tau_j L_j(\boldsymbol{\tau}) = G$ (with $G = 0$ in our baseline), imposed at the *endogenous* populations the policy itself induces. We handle this in two parts. Within the fixed-point iteration we hold the level fixed by recentering each sweep against the baseline populations, which keeps the relative taxes well-conditioned. At convergence we solve for the common level shift c that balances the budget at realized populations,

$$\sum_j (\tau_j + c) L_j(\boldsymbol{\tau} + c\mathbf{1}) = G,$$

by Newton's method. The derivative of the left-hand side with respect to c is exactly $(\sum_j L_j) M(\boldsymbol{\tau})$ – the same marginal-revenue object that enters $\mathbf{b}$ – so a single statistic governs both the optimal-tax threshold and the budget level. Shifting every τ_j by the common constant c moves $\boldsymbol{\tau}$ along $\mathbf{A}$'s null space and, because a common change in marginal utilities cancels in the ratio $\mu_j/\bar{\mu}$, leaves the relative pattern and $\mathbf{b}$ essentially unchanged; it selects the budget-balanced level without disturbing the optimized differences.

Our iterative algorithm proceeds as follows:

- **Initialize:** Set $k = 0$ and $\boldsymbol{\tau}^{(0)} = \mathbf{0}$.
- **Repeat until convergence:**
 - Compute after-tax incomes $y_j^{(k)} = w_j - \tau_j^{(k)}$.
 - Compute location choice probabilities $p_{ij}^{(k)}$ and population shares $L_j^{(k)} = \sum_i p_{ij}^{(k)}$
 - Evaluate migration elasticities $\frac{\partial \mathbf{L}}{\partial \boldsymbol{\tau}} \big|_{\boldsymbol{\tau}^{(k)}}$
 - Update taxes by solving $\mathbf{A}(\boldsymbol{\tau}^{(k)}) \boldsymbol{\tau}^{(k+1)} = \mathbf{b}$ using Gauss-Seidel:

$$\tau_j^{(k+1)} = \frac{1}{A_{jj}} \left(b_j - \sum_{m < j} A_{jm} \tau_m^{(k+1)} - \sum_{m > j} A_{jm} \tau_m^{(k)} \right).$$

- Set $k \leftarrow k + 1$.

- **Stop** when $\|\boldsymbol{\tau}^{(k)} - \boldsymbol{\tau}^{(k-1)}\| < \epsilon$
- **Balance the budget at realized populations:** Recenter $\boldsymbol{\tau}^{(k+1)}$ by subtracting the post-policy population-weighted mean tax. This holds the common level fixed while the relative taxes converge; the exact balanced budget at the policy's realized populations is imposed once, after convergence.

As mentioned in the text, we impose income floors to prevent implausible outcomes: pre-tax income must exceed \$30,000 for non-college graduates and \$60,000 for college graduates, and post-tax income must exceed \$10,000 for all individuals.

C Non-identification without separability

In this section we formally show non-identification of policy relevant μ_j under non-separability of utility functions even with full support location specific instruments.

Environment. Individual i chooses location 1 iff $\Delta_i(z) \geq 0$, where $\Delta_i(z) = u(w_{i1} + z_1, e_{i1}) - v(w_{i2} + z_2, e_{i2})$; the location shares are $P_1(z) = \Pr(\Delta_i(z) \geq 0)$ and $P_2 = 1 - P_1$, and the planner-relevant averages, evaluated at baseline income $z = 0$, are

$$\mu_1 = \mathbb{E}[u_w(w_{i1}, e_{i1}) \mid \Delta_i(0) \geq 0], \quad \mu_2 = \mathbb{E}[v_w(w_{i2}, e_{i2}) \mid \Delta_i(0) < 0].$$

We maintain Assumption 1 ($u(\cdot, e), v(\cdot, e) \in C^1$ in income, with u_w, v_w positive and bounded); we assume that for each $z \in \mathcal{Z}$ the scalar $\Delta_i(z)$ admits a density $f_{\Delta(z)}$ that is continuous and strictly positive at 0; and we take $\mathcal{Z}$ open, with z independent of (e_{i1}, e_{i2}) given wages and rich enough that every baseline-inframarginal individual is rendered marginal for some $z \in \mathcal{Z}$ (full support). The econometrician observes $\{P_1, P_2, \nabla P_1, \nabla P_2\}_{\mathcal{Z}}$ and the wages w , but not the idiosyncratic tastes e .

Lemma C.1 (Marginal choice-probability response). *Under these assumptions $P_1 \in C^1(\mathcal{Z})$, and for every $z \in \mathcal{Z}$*

$$\frac{\partial P_1}{\partial z_1} = f_{\Delta(z)}(0) \mathbb{E}[u_w(w_{i1} + z_1, e_{i1}) \mid \Delta_i(z) = 0], \quad \frac{\partial P_1}{\partial z_2} = -f_{\Delta(z)}(0) \mathbb{E}[v_w(w_{i2} + z_2, e_{i2}) \mid \Delta_i(z) = 0],$$

and $\nabla P_2 = -\nabla P_1$.

Proof. Since $P_1(z) = 1 - F_{\Delta(z)}(0)$, perturbing z_2 by h changes each individual's index by $\Delta_i(z_1, z_2 + h) - \Delta_i(z) = -v_w(w_{i2} + z_2, e_{i2})h + o(h)$, so the individuals who leave location 1 are exactly those with $\Delta_i(z) \in [0, h v_w + o(h))$. By continuity of $f_{\Delta(z)}$ at 0 this set has probability $h f_{\Delta(z)}(0) \mathbb{E}[v_w(w_{i2} + z_2, e_{i2}) \mid \Delta_i(z) = 0] + o(h)$; dividing by h and letting $h \rightarrow 0$ gives $\partial P_1 / \partial z_2$, and z_1 is symmetric with the opposite sign. Continuity in z of the right-hand sides gives $P_1 \in C^1(\mathcal{Z})$. This is the standard Leibniz rule for a threshold-crossing probability: a marginal income change relocates only the indifferent individuals, weighted by the threshold density $f_{\Delta(z)}(0)$. $\square$

Proof of Theorem 3.1. Taking the ratio of the two cross-responses in Lemma C.1 cancels the common density $f_{\Delta(z)}(0)$ – which the assumptions leave otherwise unrestricted, and which no other linear functional of ∇P eliminates – and yields (4). The identified content of the data is therefore, for each $z \in \mathcal{Z}$, a ratio of marginal utilities *among individuals indifferent at the perturbed income $w + z$* . Writing $g(t) = \mathbb{E}[u_w(w_{i1}, e_{i1}) \mid \Delta_i(0) = t]$, at $z = 0$ the numerator of (4) is the boundary value $g(0)$.

By iterated expectations the planner’s target is

$$\mu_1 = \frac{\int_0^\infty g(t) f_{\Delta(0)}(t) dt}{\int_0^\infty f_{\Delta(0)}(t) dt},$$

which depends on $g(t)$ for all $t > 0$. At baseline the derivative data determine only the boundary value $g(0)$. Full support does not supply the rest: rendering a baseline-inframarginal individual marginal requires perturbing that individual’s own income, so each z identifies $\mathbb{E}[u_w(w_{i1} + z_1, e_{i1}) \mid \Delta_i(z) = 0]$ – a conditional mean at the *perturbed* income over a re-sorted threshold set – not the baseline-income means $g(t)$, $t > 0$, that the integral requires. Without separability $u_w(w, e)$ depends on e , and the event $\{\Delta_i(0) = t\}$ selects systematically on e , so g is non-constant and $g(0)$ does not determine $\{g(t) : t > 0\}$. Hence the map $\{P, \nabla P\}_{\mathcal{Z}} \mapsto \mu_1$ is not injective: two structures satisfying the assumptions agree on every observable yet differ in μ_1 (Appendix D constructs such pairs). The argument for μ_2 is symmetric. Under separability u_w is independent of e , so g is constant, $\mu_1 = g(0)$, and (4) identifies μ_1/μ_2 . $\square$

D Examples illustrating observational equivalence

We give two simple examples in which, absent separability, distinct utility specifications generate identical location-choice behavior – and hence identical responses of choice probabilities to financial incentives – yet imply different values of the policy-relevant averages μ_j (in a two-location case, μ_1 and μ_2). Related observational-equivalence logic appears in Gaubert et al. (2025, Appendix B.3.2) and Fajgelbaum and Gaubert (2025), who also note, citing an earlier version of this paper, that unidentified preference heterogeneity can affect the marginal utility of consumption and therefore social preferences.

D.1 Example 1: Complementarity between income and shocks

Consider J locations and a utility index non-separable in income w_j and an idiosyncratic shock ϵ_{ij} :

$$\underbrace{u_{ij} = w_j \epsilon_{ij}}_{\text{Model 1}} \quad \text{versus} \quad \underbrace{u_{ij} = w_j \tilde{\epsilon}_{ij}}_{\text{Model 2}}, \quad \text{where} \quad \tilde{\epsilon}_{ij} = \epsilon_{ij} \times D(\epsilon_{i1}, \dots, \epsilon_{iJ}),$$

with $D(\epsilon_{i1}, \dots, \epsilon_{iJ}) > 0$ an arbitrary strictly positive function of the individual’s entire shock vector. Because the realized D does not vary across locations j , *it does not affect any location choice* but it does shift marginal utilities.

Interpret ϵ_{ij} as idiosyncratic factors – proximity to family, a spouse’s job prospects, health or schooling needs – that make location j especially appealing. In Model 1, $u_{ij} = w_j \epsilon_{ij}$, these same factors mechanically scale the marginal utility of income there, since $\partial u_{ij} / \partial w_j = \epsilon_{ij}$: the very

forces that make a location more likely to be chosen also raise the marginal utility of income *in that location*. The factor $D(\cdot)$ in Model 2 breaks this link, letting the drivers of location choice correlate with a separate, non-location-specific shifter of marginal utility. Under the parameterization $D = \prod_{k=1}^J \epsilon_{ik}^{\phi_k}$, for example, $\phi_k < 0$ makes high- ϵ_{ik} realizations coincide with low marginal utility and $\phi_k > 0$ the reverse – capturing that the circumstances drawing someone to a location may coincide with periods when a dollar is especially valuable regardless of where they live (financial stress, caregiving, liquidity needs), or especially not. Conditioning on choosing j can then select systematically higher- or lower-marginal-utility individuals even though choice behavior is unchanged.

The second model scales each individual’s entire utility vector by a positive factor, so it implies the same argmax and hence the same choices; marginal utilities, however, differ. In Model 1,

$$\mu_j^{(1)} \equiv E \left[\frac{\partial u_{ij}}{\partial w_j} \middle| j \text{ chosen} \right] = E \left[\epsilon_{ij} \middle| j \text{ chosen} \right],$$

while in Model 2,

$$\mu_j^{(2)} \equiv E \left[\frac{\partial u_{ij}}{\partial w_j} \middle| j \text{ chosen} \right] = E \left[\tilde{\epsilon}_{ij} \middle| j \text{ chosen} \right] = E \left[D(\epsilon_{i1}, \dots, \epsilon_{iJ}) \epsilon_{ij} \middle| j \text{ chosen} \right].$$

Unless $D(\cdot)$ is almost surely constant, $\mu_j^{(2)}$ generically differs from $\mu_j^{(1)}$: exogenous wage variation identifies location-demand responses but not the shape of $D(\cdot)$ – how the forces that shift location choice correlate with the marginal utility of income.

D.2 Example 2: Shocks that are perfect substitutes for income

As shown by [Gaubert et al. \(2025\)](#) and [Fajgelbaum and Gaubert \(2025\)](#), a second class treats idiosyncratic location factors and income as perfect substitutes, combining them in a common “income-equivalent” index, so choice is governed by which location delivers the higher value of that index. Even when observed choices pin the index down, they need not reveal how its idiosyncratic component maps into the policy-relevant marginal utilities μ_j . Let both locations have the same deterministic income w , and consider:

$$\begin{aligned} \text{Model A: } u_{i1} &= u(w + \epsilon_{i1} - \epsilon_{i2}), & u_{i2} &= u(w), \\ \text{Model B: } u_{i1} &= u(w), & u_{i2} &= u(w + \epsilon_{i2} - \epsilon_{i1}), \end{aligned}$$

with $u(\cdot)$ strictly increasing. Here $\epsilon_{i1} - \epsilon_{i2}$ is the individual’s net, income-equivalent advantage of location 1 (e.g., childcare from nearby family, commute-time savings, school-quality or spousal-job considerations, provider access, or local tax and insurance differences), assumed to enter utility as an additive dollar-equivalent term inside the same index as income.

In both models location 1 is chosen iff $\epsilon_{i1} \geq \epsilon_{i2}$ – in Model A because $u(w + \epsilon_{i1} - \epsilon_{i2}) \geq u(w)$,

in Model B because $u(w) \geq u(w + \epsilon_{i2} - \epsilon_{i1})$, each equivalence using only that $u(\cdot)$ is increasing. The two models therefore generate identical choices and identical choice probabilities for every distribution of $(\epsilon_{i1}, \epsilon_{i2})$: what matters for choice is only whether the income-equivalent advantage of location 1 is positive.

Yet the models imply different marginal utilities, depending on where the idiosyncratic component of the index is “loaded.” In Model A the index varies in location 1 while location 2 delivers the fixed $u(w)$; in Model B the reverse. Whenever $u(\cdot)$ is concave, so $u'(\cdot)$ is decreasing, the location whose index is shifted upward for its chosen residents has *lower* marginal utility of income. Formally, in Model A

$$\mu_1^A = E[u'(w + \epsilon_{i1} - \epsilon_{i2}) \mid 1 \text{ chosen}], \quad \mu_2^A = u'(w),$$

so $\mu_1^A < \mu_2^A$ (since $\epsilon_{i1} - \epsilon_{i2} > 0$ for everyone choosing location 1 and u is concave), while in Model B the roles reverse,

$$\mu_1^B = u'(w), \quad \mu_2^B = E[u'(w + \epsilon_{i2} - \epsilon_{i1}) \mid 2 \text{ chosen}],$$

giving $\mu_1^B > \mu_2^B$. The same location-choice behavior is thus consistent with opposite rankings of marginal utility across locations: choices reveal which location tends to deliver the higher income-equivalent index, but not how that index’s idiosyncratic component maps into the conditional marginal utilities μ_j that enter optimal policy.

E Survey design

E.1 Details on design choices

Definition of a location: A location in our survey is defined as a metropolitan statistical area (MSA). We obtain the list of all states and MSAs within each state from the U.S. Census Bureau.³⁴ We provide a drop-down menu of states followed by a drop-down menu of MSAs within the chosen state. At the beginning of the survey, we define a location as a metropolitan area: *Note: A metropolitan area is a large region comprising one or more counties, which includes a city with a significant population and its surrounding areas that are economically interconnected. It is defined by factors such as population size and the extent of linkage among nearby communities in terms of jobs, commuting, and transportation.*

Additionally, the three locations other than the respondent’s hometown are required to be

³⁴MSA county crosswalk obtained from the Delineation files of Core Based Statistical Areas (CBSAs), metropolitan divisions, and combined statistical areas (CSAs): <https://www.census.gov/geographies/reference-files/time-series/demo/metro-micro/delineation-files.html>. Connecticut had changed counties to Planning Regions. The file linking old FIPS codes to new FIPS codes are obtained from the Associated Press at https://developer.ap.org/ap-elections-api/docs/CT_FIPS_Codes_forPlanningRegions.html.

unique and different from the current location of the respondent and the location where they grew up.

Choice of language: We administered the survey in English and in Spanish. The Spanish translation was done by experts at the Rutgers Survey Center. We also conducted a pilot survey and a focus group discussion with a small sample of Spanish-speaking individuals to ensure that the translation was accurate and that the survey was understandable.

Displaying the exogenous location-specific attributes: We color coded the direction of the shocks to highlight the incomes to reduce cognitive load on the respondents. This decision was based on tests of whether the color coding had any effect on the responses in two simultaneous rounds of pilot data collection where we found no difference in the responses between the color-coded and non-color-coded versions of the pilot survey, except for shorter times to completion in the color-coded version. Thus, we decided to use the color-coded version in the final survey to reduce the cognitive load on the respondents. In order to avoid the possibility of any parallel being drawn with political affiliation, we specifically avoided using the colors blue and red. To represent increment (reduction) in income we used green (yellow). We did not color code the presence or absence of family and friends. We could have used the same color coding for the presence of family and friends, but we decided against it to avoid assuming whether the presence of family and friends increases or decreases utility, which in the case of income is more straightforward. In the same spirit of minimizing cognitive load, we showed respondent the shocked income instead of showing both the percentage change and the new income.

Ordering: Since we are identified within individuals, we are not concerned about the order of the locations. However, we randomized the order of the scenarios to avoid any ordering effects.

Facilitating ease in responses: We made several choices in order to facilitate ease in responding to the survey. Here are some of the most important ones:

- We provided a slider for the respondents to report their expected income and the compensating differential. This also avoids the possibility of different respondents using different units of income.
- In each exogenously varied scenario, respondents were shown their baseline probabilities and compensating differentials to facilitate comparison.
- In order to emphasize that their response in the compensating differential questions is only conditional on that location being their choice in 5 years, each option of location was followed with the phrase “*had I chosen to live there in 5 years*”

E.2 Background

The Survey Research Center at Rutgers University has helped us design and implement the survey. Our survey was administered in the July and November of 2024 to a nationally representative sample of adults. The survey was available in English and Spanish, and participants that

completed the survey were given small payments in line with typical compensation given to on-line survey participants. The survey respondents themselves were anonymous, and no personally identifiable data was collected.

The survey started with a small pilot in April of 2024. Based on responses and feedback from the pilot, we adjusted survey questions and then administered the survey to a larger sample as a second pilot (July). Once we reviewed the results from that pilot and decided not to implement any changes, we then launched the full survey in November.³⁵

E.3 Survey module overview

Table A1: Survey flow and what each module provides

Module	What respondents do	What it provides / how used
Geographic & economic background	• Current MSA; hometown MSA • Current income (indiv/HH)	• Descriptives + validation inputs • Baseline income measures for the choice menu
Future location menu (4 MSAs)	• Choose 4 feasible MSAs: – hometown – 3 distinct destinations (not current)	• Individual menu $\mathcal{M}_i$ • Which locations are considered together (substitution)
Baseline attributes by MSA	• For each $j \in \mathcal{M}_i$: – expected HH income (5y) – family/friends present (0/1)	• Baseline covariates (validation + correlates) • Inputs to scenario perturbations
Rankings by MSA	• Rank the 4 MSAs by: – support from family/friends – fun / hobbies fit – expensiveness (COL + desired housing)	• Validation vs external data: – expensiveness vs housing prices/costs – fun vs amenity indices (and crime)
Baseline choice probabilities	• Report p_{ij0} over 4 MSAs (sum to 1)	• Baseline choice weights for aggregation • Menu-based substitution information
Baseline MU elicitation	• Pick MSA where \$5k is most valuable • Set compensating differentials in other MSAs	• Within-person MU comparisons across MSAs • Relative MU (up to person scale)
Scenario block (8 scenarios)	• Instruction: hold all else fixed within each MSA • Randomly vary within menu: – income by MSA – family/friends by MSA • In each scenario, collect: – choice probabilities p_{ijk} – MU block (location of \$5k, compensating differentials)	• Location-demand responsiveness: – income elasticity – family/friends effect • Own/cross substitution patterns across locations • Panel MU responses to income and family/friends shifts
Demographics module	• Age, gender, race/ethnicity, education, household, employment	• Sample description; ACS comparison; Pareto weighting by skill group

E.4 Locations and Income

We start the survey by asking basic questions about current and previous location and current and expected future incomes.

³⁵The second pilot tested whether responses varied depending on whether certain phrases were highlighted or not. The responses did not seem to depend on highlighting. For the full survey, we used the highlighted version.

- We first ask the state, metropolitan area or nearest city, and zip code that the respondent considers his or her home town, and then the same data for where the respondent currently lives.
- We next ask the respondent their total pre-tax income over the previous 12 months. If the respondent lives with one or more adults, the respondent is asked to only report income from his or her job.
- Then we ask the respondent to report his or her expected individual annual income 5 years from now.
- Then we ask the respondent their current total annual *household* income. Respondents are instructed to report income from all sources and include incomes from all adults in the household.

E.5 Baseline Data

The survey next asks respondents to choose the three most likely state and metropolitan areas where they may live 5 years from now. The first question is the most likely, the second is the second-most likely, and the third is the least likely.

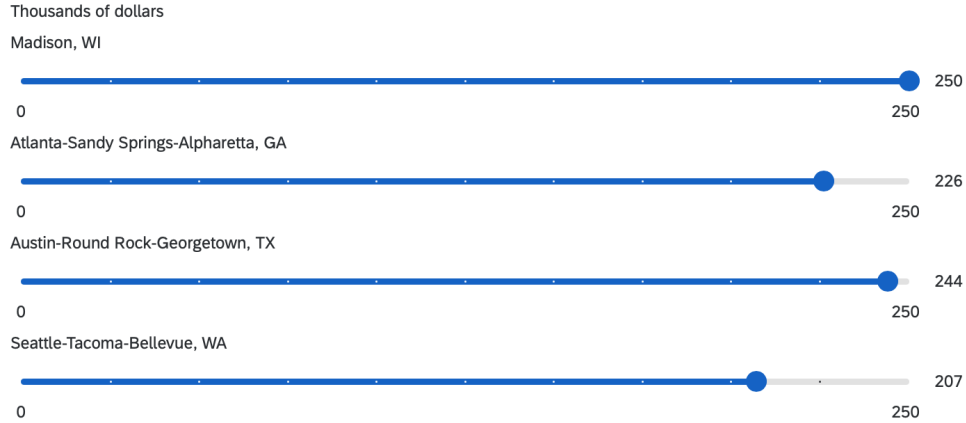
To help fix ideas, going forward we will consider an example where the most likely location is Atlanta, GA, the second most likely location is Austin, TX, the third most likely location is Seattle, WA. In this example, the respondent’s hometown is Madison, WI and current location is Washington, DC.

After the respondent reports the three places he or she is most likely to live, the respondent is asked to report the probabilities he or she will be “living in these locations 5 years from now.” The set of locations includes the respondent’s hometown, and the respondent is asked to make sure the probabilities sum to 100 percent. A sample response is below.

Madison, WI	10	%
Atlanta-Sandy Springs-Alpharetta, GA	20	%
Austin-Round Rock-Georgetown, TX	40	%
Seattle-Tacoma-Bellevue, WA	30	%
Total	100	%

We will call the collection of the three most likely locations respondents expect to live 5 years from now and the home town as the respondent’s “likely future locations.”

The respondent is then asked to report his or her expected annual total income in 5 years in each of the likely future locations. The respondent answers these questions by moving a slider, shown below.



The next questions ask about the presence of savings and friends and family in each of the likely future locations as well as expected future costs of living.

- We first ask respondents if they have enough cash on hand to “cover expenses for a few months in the event of a job loss or a relatively small event causing some unanticipated expenses.” A sample response is shown below.

	Definitely not	Probably not	Probably	Definitely
Madison, WI	☐	☐	☒	☐
Atlanta-Sandy Springs-Alpharetta, GA	☐	☐	☐	☒
Austin-Round Rock-Georgetown, TX	☐	☐	☐	☒
Seattle-Tacoma-Bellevue, WA	☐	☐	☒	☐

- We next ask if respondents have a social network of family or friends in each of the likely future locations. In each location, the answer is restricted to be “yes” or “no.”
- The respondent is then asked, *In the future, if your life circumstances become very difficult, in which location(s) do you expect to get the most support from your social network of family and friends?* The respondent is asked to rank the likely future locations from 1 (most likely to get support) to 4 (least likely).
- Next, the respondent is asked, *Thinking about your hobbies and things you like to do for entertainment, where do you think you will have the most fun?* The respondent is asked to rank the likely future locations from 1 (most fun) to 4 (least fun).
- Finally, the respondent is asked, *Given the sorts of goods and services you like to buy and given the kind of housing unit you wish to live in, which locations do you think would be most expensive for you?* The respondent is asked to rank the likely future locations from 1 (most expensive) to 4 (least expensive).

E.6 Experimental Variation in Hypothetical Scenarios

The next two blocks of questions are key for estimating the marginal utility of consumption. First, we ask how probabilities over location choices change as assumptions about income and the presence of family and friends vary. Instructions leading up to this block of questions are as follows:

Earlier you told us four places you could imagine living 5 years from now. You also gave us your best guess about what your annual income would be in each location, whether family would be present, and your probability of living in each place.

Now we are interested in how the probability of living in these locations would be affected if your anticipated income or family ties were changed, holding all other aspects of the location the same.

For your convenience, we will also show your reported expectation of income and presence or not of family/friends in parenthesis below the changed income and family ties.

Respondents are shown new values of income and family and friends in each location, with the baseline values shown in parentheses. We select the new values quasi-randomly to span a large support of possible variation. The variation in incomes and presence of family and friends for each MSA, by scenario, is shown below

Table A2: Survey-induced exogenous variation across scenarios

Scenario	Option shown to respondent							
	Hometown	Alt. 1	Alt. 2	Alt. 3	Hometown	Alt. 1	Alt. 2	Alt. 3
	<i>Income growth factor (multiplicative)</i>				<i>Family and friends indicator</i>			
1	1.2520	0.9162	1.3590	0.8913	0	1	1	0
2	0.7849	1.1480	0.8662	1.2660	0	1	1	0
3	1.2090	0.8617	0.6198	1.1050	1	0	0	1
4	0.8848	1.4440	1.2390	0.7793	1	0	0	1
5	0.7254	0.7256	1.3950	1.1870	1	0	0	1
6	1.0650	1.2130	0.9453	0.7797	1	1	0	0
7	0.8673	0.8899	0.7091	1.4020	0	0	1	0
8	1.4440	0.9491	1.1100	1.2090	0	1	1	1

Notes: Each row is a scenario used in the survey.

Continuing with our example, in a first scenario respondents are shown the following:

Please report your revised chances of living in each location if, in 5 years, your expected household income in these locations and presence or not of family/friends are different than what you expected as shown below, while holding all other aspects of the locations the same.

Location	Future Income	Family/Friends
Madison, WI	313k (was 250k)	No (was Yes)
Atlanta-Sandy Springs-Alpharetta, GA	206k (was 226k)	Yes (was No)
Austin-Round Rock-Georgetown, TX	332k (was 244k)	Yes (was Yes)
Seattle-Tacoma-Bellevue, WA	184k (was 207k)	No (was No)

Given the new assumptions, respondents are asked to recompute the probabilities they live in each of the likely future locations in 5 years:³⁶

*Your reported % chances of living in each location is indicated in parentheses.
Please make sure that your new % chances sum to 100.*

Madison, WI (10%)	New%	15
Atlanta-Sandy Springs-Alpharetta, GA (20%)	New%	15
Austin-Round Rock-Georgetown, TX (40%)	New%	45
Seattle-Tacoma-Bellevue, WA (30%)	New%	25
Total	New%	100

Respondents are shown 9 scenarios in total: baseline plus 8. Call this entire set of 9 scenarios as the “probability scenarios” which help us identify the elasticity of location choice with respect to expected income and the presence of friends and family.

The next block of questions allows for estimation of the marginal utility of consumption in each of the likely future locations, as we show later. Instructions for this block of questions are as follows:

The following questions explore a circumstance in which your income unexpectedly and permanently jumps after you have chosen your future location. We will then ask the same two questions under a number of different scenarios where your expected household income and presence or not of family/friends might differ from your stated expectations.

³⁶Note the probabilities at the baseline assumptions are shown in parentheses.

1. The first question will ask you to choose the location in which you would benefit most from the permanent jump in income, had you chosen to live in that location. Call the benefit arising from this permanent jump in income in the location you choose your “change in happiness.” Notes that this location may not necessarily be the one where you are most likely to move. Simply think of where the permanent jump in income would be most helpful had you chosen to live there.
2. Next, we will ask you to imagine living in the other locations 5 years from now, and how much of a jump in income would be needed for your change in happiness at each place to be equal to the change in happiness at the location in the first question.

Respondents will be asked the same two questions nine times. The first time, they will be shown a scenario corresponding to their baseline assumptions. The next eight times, they will be shown scenarios that are the same as the scenarios shown in the probability questions. In each of these nine scenarios, respondents will be instructed to hold all factors other than income and the presence of family and friends the same.

- The first question is as follows: *Fast forward 5 years. Out of many possible future versions of yourself, think about 4 versions A, B, C, and D each of whom has chosen to live in one of the locations you have specified and with the incomes and presence of family/friends as given below.*

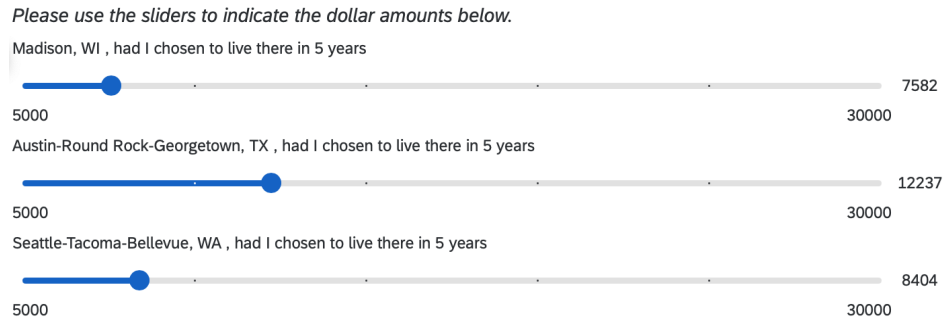
In the first scenario, respondents are shown baseline assumptions

Location	Future Income	Family/Friends
Madison, WI	250k	Yes
Atlanta-Sandy Springs-Alpharetta, GA	226k	No
Austin-Round Rock-Georgetown, TX	244k	Yes
Seattle-Tacoma-Bellevue, WA	207k	No

The question continues: *Fast forward 5 years and imagine yourself in each of the locations below. In which location would you benefit the most from an unexpected and permanent increase of \$5,000 in household income, had you chosen to live there? Note that this location may not necessarily be the one where you are most likely to move. Call the benefit arising from this permanent jump in income in the location you choose your “change in happiness.”* Respondents choose a location from the list of likely future locations:

- ☐ Madison, WI , had I chosen to live there in 5 years
- ☒ Atlanta-Sandy Springs-Alpharetta, GA , had I chosen to live there in 5 years
- ☐ Austin-Round Rock-Georgetown, TX , had I chosen to live there in 5 years
- ☐ Seattle-Tacoma-Bellevue, WA , had I chosen to live there in 5 years

- The second question is: *Now, for each of the other versions living in the other locations, indicate the size of the jump in household income that they would need living in those locations 5 years from now such that their change in happiness after receiving this jump in household income is the same as the change in happiness from the \$5,000 jump in household income given to the version living in [Atlanta-Sandy Springs-Alpharetta, GA].*³⁷ Respondents then adjust a slider for each of the other three likely future locations to answer that question:



The respondents then answer the same 2 questions 8 more times, but with the assumptions about income and friends and family in each location changed, for example:

Location	Future Income	Family/Friends
Madison, WI	313k (was 250k)	No (was Yes)
Atlanta-Sandy Springs-Alpharetta, GA	206k (was 226k)	Yes (was No)
Austin-Round Rock-Georgetown, TX	332k (was 244k)	Yes (was Yes)
Seattle-Tacoma-Bellevue, WA	184k (was 207k)	No (was No)

E.6.1 Economic Policy and Political Leanings

After we ask these 18 questions (2 questions under 9 sets of assumptions), we then ask a block of 7 questions about economic policy and political leanings that we borrow from [Kuziemko et al. \(2015\)](#). We are interested in whether political preferences on views on redistribution across agents are correlated with expected variation in their own marginal utility of income across potential likely future places.

³⁷Note that the location in brackets will be the location chosen as the answer to the previous question.

E.6.2 Demographic Questions

In the final 9 questions of the survey, we ask each respondent to list (1) his or her age (restricted to one of 7 intervals), (2) gender (male, female, or other), (3) if the respondent is of Hispanic, Latino, or Spanish origin (yes/no), (4) ethnicity or race (one of six labels including “Other,”), (5) current marital status, (6) number of adults in the household, (7) children under 18 in the household, (8) highest level of education (one of nine categories), and (9) employment status over the past 12 months (one of seven categories).

F Comparison of the Sample with 2022 5-Year ACS

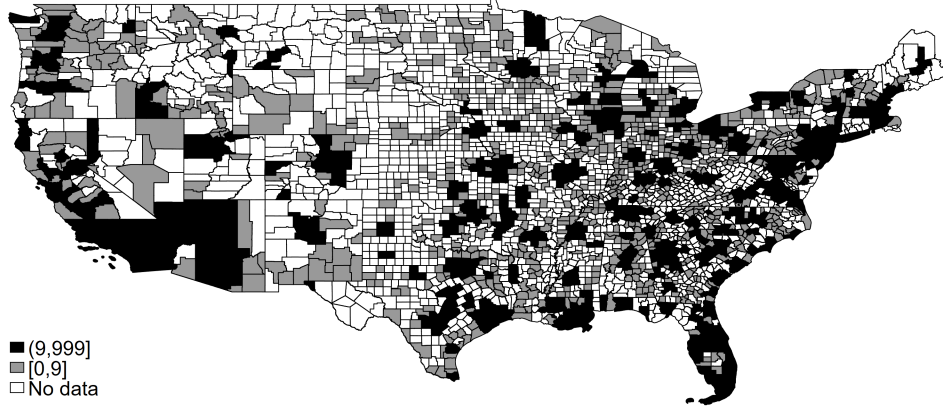
Table A3 compares our sample to key data from the 2022 5-year American Community Survey (ACS) (U.S. Census Bureau, 2022) and Figure A1 shows the raw number of times that survey respondents refer to locations as potentially living in those locations in the future. In the figure, darker-shaded MSAs are included more frequently in the menu of possible future locations than locations that are shaded more lightly. The spatial distribution of these destinations reflects both the geographic dispersion of our sample and expected migration patterns.

Table A3: Sample composition relative to the ACS

	ACS	Survey		ACS	Survey
Census division			Gender		
Pacific Division	0.159	0.136	Male	0.510	0.433
Mountain Division	0.076	0.051	Female	0.490	0.562
West South Central Div.	0.122	0.121			
East South Central Div.	0.058	0.057	Hispanic origin		
South Atlantic Division	0.206	0.196	Hispanic	0.176	0.156
West North Central Div.	0.064	0.066	Not Hispanic	0.824	0.844
East North Central Div.	0.141	0.172			
Middle Atlantic Division	0.127	0.158	Race		
New England Division	0.047	0.043	Other	0.069	0.036
			Mixed race	0.111	0.024
Education			Native American	0.010	0.019
Doctoral Degree	0.015	0.022	Asian / Pacific Islander	0.064	0.057
Professional Degree (JD, MD, MBA)	0.021	0.018	Black	0.118	0.146
Master’s Degree	0.092	0.075	White	0.629	0.719
4-year College Degree	0.207	0.281			
2-year College Degree	0.085	0.072	Age		
Some College	0.205	0.147	65+ years old	0.226	0.209
High School degree / GED	0.271	0.264	55–64 years old	0.159	0.192
Some High School	0.060	0.112	45–54 years old	0.154	0.187
Primary education or less	0.043	0.008	35–44 years old	0.171	0.191
			25–34 years old	0.173	0.113
			18–24 years old	0.116	0.107

Notes: Entries report shares. The ACS column reports shares from the ACS. The Survey column reports shares in the survey. Shares within each block sum to 1.

Figure A1: Potential Future Locations



G Construction of the wage index

The wage index $\overline{\log w_j}$ is a skill-adjusted metro log-earnings index: the metro fixed effect from a regression of individual log annual wage income on demographics (sex, race/ethnicity, age and age², an education category, and marital status), estimated on employed, full-time workers in the 2022 ACS five-year sample (U.S. Census Bureau, 2022), and recorded as one value per Census CBSA. Some metros in respondents' menus have no directly estimated ACS index, because too few ACS workers identify them. For these we impute $\overline{\log w_j}$ from the survey's own income data, in two steps. First, we construct a survey-based location income index – the metro fixed effect from a two-way regression of respondents' reported expected income on respondent and metro fixed effects – which plays the same role as $\overline{\log w_j}$ but is identified from the survey rather than the ACS. Second, we regress the ACS index $\overline{\log w_j}$ on this survey income index across the well-sampled metros where both are available (those appearing in at least fifty respondents' menus), and use the fitted relationship to predict $\overline{\log w_j}$ for the metros that lack the ACS measure. A metro for which neither source is available – it lacks even the survey income index – causes that respondent's menu to be dropped, preserving the menu structure. How this index enters the tax base and, in the log-C specifications, the marginal utility of consumption is described in Appendix J and summarized in Table A4.³⁸

H Aggregation of survey marginal utilities

This appendix proves Proposition 6.1 from Section 6.1, which establishes the three aggregation properties (P1–P3) used to construct the location-level welfare weights.

³⁸Two computational details specific to the IIA log-C specification: pool baseline choice probabilities are the within-menu baseline probability mass falling in each pool, averaged across respondents; and the single pool with no survey marginal-utility estimate has its relative marginal utility set to the cross-pool mean.

Proof of Proposition 6.1. P1 (Scale invariance). Rescaling $m_{ij} \mapsto s_i m_{ij}$ multiplies both the numerator and denominator of m_{ij}^{rel} by s_i , so m_{ij}^{rel} is unchanged. The location effects ψ_j from (11) are also invariant, since $\log s_i$ is absorbed by α_i . Therefore g_{ij} is invariant to respondent-specific rescaling.

P2 (Preserve high-MU menus). The anchor $\kappa_i = \sum_{\ell} p_{i\ell} \exp(\psi_{\ell})$ is larger for people whose menus concentrate on high- ψ locations. This cross-person variation in κ_i is preserved in g_{ij} , so the construction does not mechanically equalize menu-average welfare weights across people.

P3 (No spurious redistribution). If two people share the same menu, baseline probabilities, and Pareto weight, they have the same κ_i . By construction $\sum_j p_{ij} m_{ij}^{\text{rel}} = 1$, so $\sum_j p_{ij} g_{ij} = \Pi_{t(i)} \kappa_i$ regardless of the within-menu MU profile. A planner therefore values an additional dollar in their hands equally before location choice is realized. $\square$

I Computation: Deadweight Cost and Insurance Value

The deadweight cost is computed as the sum over all pairs of cities of the net migration between cities times the tax differential of the cities multiplied by one half. This cost is a migration-weighted sum of squared tax *differences* between substitute locations, so a spatially uniform tax incurs no deadweight loss.

Technically, we set $\text{EDC} = \frac{1}{2} \boldsymbol{\tau}' (D - \Omega) \boldsymbol{\tau} = \frac{1}{2} \sum_{j < k} \Omega_{jk} (\tau_j - \tau_k)^2$, where $\Omega_{jk} \equiv \frac{1}{2} (\partial L_j / \partial \tau_k + \partial L_k / \partial \tau_j) \geq 0$ is the *average* of the two directional migration responses between locations j and k , and $D = \text{diag}(W)$ with $W_j = \sum_{k \neq j} \Omega_{jk}$. We are not assuming these responses are symmetric: a tax on k pulls people into j at rate $\partial L_j / \partial \tau_k$, while a tax on j pulls them into k at rate $\partial L_k / \partial \tau_j$, and the two differ – the flow is larger when the tax falls on the location with the lower after-tax income. Replacing the two directional responses by their average does not affect the deadweight cost, because that cost is a quadratic form in the tax vector and the asymmetric part of any matrix drops out of such a form; each pair simply contributes its squared tax difference weighted by the average of its two responses. The factor $\frac{1}{2}$ is the usual Harberger-triangle term.

The matching insurance value is the first-order gross gain $\text{INS} = -\boldsymbol{\tau}' \mathbf{B} = \sum_j \tau_j L_j (\bar{\mu} - \mu_j) / \bar{\mu}$, where $B_j = L_j (\mu_j / \bar{\mu} - 1)$ is location j 's population-weighted marginal utility relative to the mean $\bar{\mu}$. The two are tightly linked: because the gain is linear in $\boldsymbol{\tau}$ while the cost is quadratic, the optimality condition $(D - \Omega) \boldsymbol{\tau}^* = -\mathbf{B}$ gives $\boldsymbol{\tau}' (D - \Omega) \boldsymbol{\tau}^* = \text{INS}$; hence the insurance value is exactly twice the deadweight cost, $\text{INS} = 2 \text{EDC}$, and net welfare retains half the gross gain, $\Delta W = \text{INS} - \text{EDC} = \frac{1}{2} \text{INS}$. Our imposition of income floors may cause INS to not exactly equal twice EDC.

J Construction and use of location income

This appendix describes how the location-specific income measures are built and where they enter the optimal-policy calculation, summarized in Table A4. For each respondent i and location

j the survey elicits a hypothetical income income_{ij} ; the skill-adjusted log earnings index $\overline{\log w_j}$ (2022 ACS, constructed in Appendix G) summarizes local earnings. These objects are put to two distinct uses.

Use 1: the tax base. The planner’s instrument is a location-specific tax/transfer τ_j levied on pre-tax income, so after-tax income at j is $y_{ij}(\tau) = w_{ij}^0 - \tau_j$; location choices respond to it through each respondent’s elasticity ν_i , generating the migration responses and deadweight costs in the optimal-tax problem. We impose income floors: pre-tax income is winsorized from below at \$30,000 (non-college) and \$60,000 (college), and post-tax income at \$10,000. The pre-tax base w_{ij}^0 is formed in one of two ways.

(1a) Base case (own menu). The choice set is the respondent’s survey menu, and the pre-tax base is the survey-elicited income, used directly: $w_{ij}^0 = \text{income}_{ij}$. The wage index does not enter the tax base.

(1b) IIA case (all locations). The choice set expands to all 270 pools, so an income must be attached to every respondent-location pair. We strip the wage index from each respondent’s menu incomes to form a wage-neutral “real” income level,

$$\alpha_i = \sum_{j \in \text{menu}} p_{ij0} \frac{\text{income}_{ij}}{\exp(\overline{\log w_j})},$$

and re-inflate at every location by its pool wage index $\overline{\log w_{\text{pool}}}$:

$$w_{i,\text{pool}}^0 = \alpha_i \exp(\overline{\log w_{\text{pool}}}).$$

Under IIA the wage index directly determines the tax base.

Use 2: marginal utility in the log-C specification. In our survey-based specification marginal utility comes directly from the elicitation, and income is used only as the tax base. In the log-consumption (“log-C”) benchmark, marginal utility is $1/\text{consumption}$, so income determines the welfare weights. This takes two forms.

(2a) Survey-substitution log-C. Consumption at j is the local wage $\exp(\overline{\log w_j})$, so $\mu_{ij} = 1/\exp(\overline{\log w_j})$. The tax base is the survey income as in case (1a); the wage index enters only through marginal utility.

(2b) IIA log-C. Consumption is the reconstructed income from case (1b), so $\mu_{i,\text{pool}} = 1/w_{i,\text{pool}}^0 = 1/(\alpha_i \exp(\overline{\log w_{\text{pool}}}))$. Here the wage index determines both the tax base and marginal utility.

Table A4: How location income is constructed and used, by specification

Specification	Choice set	Pre-tax base w_{ij}^0	Marginal utility μ_{ij}
Survey MU, menu (1a)	own menu	income_{ij} (survey)	survey elicitation
Survey MU, IIA (1b)	all pools	$\alpha_i \exp(\log \overline{w}_{\text{pool}})$	EB-shrunk location MU $\hat{\mu}_j^{\text{EB}}$
Log-C, menu (2a)	own menu	income_{ij} (survey)	$1 / \exp(\log \overline{w}_j)$
Log-C, IIA (2b)	all pools	$\alpha_i \exp(\log \overline{w}_{\text{pool}})$	$1 / (\alpha_i \exp(\log \overline{w}_{\text{pool}}))$

Notes: income_{ij} is the survey-elicited hypothetical income; $\log \overline{w}_j$ is the skill-adjusted log earnings index (Appendix G), and $\log \overline{w}_{\text{pool}}$ its population-weighted pool aggregate;

$\alpha_i = \sum_{j \in \text{menu}} p_{ij0} \text{income}_{ij} / \exp(\log \overline{w}_j)$ is respondent i 's wage-neutral real income. In every case the tax acts as $y_{ij}(\tau) = w_{ij}^0 - \tau_j$. Appendix G gives further detail on the log-C and IIA income constructions.

K Sensitivity of Results

This appendix assesses the sensitivity of our results to parameter choices. Table A5 varies one ingredient at a time from the baseline specification (column 1). Panel B regresses each location's optimal tax on its shrunk average marginal utility, the average tax of its substitute locations, and its college-graduate share.

Columns (2) and (4) double, respectively, each respondent's location-choice elasticity and the sensitivity of marginal utility to average income. In both, the Panel B sensitivity of optimal taxes to marginal utility falls – from -56.7 at baseline to -35.5 (column 2, a 37% decline) and -39.7 (column 4, a 30% decline) – and the dispersion of optimal taxes drops from \$11.5 thousand to \$7.9 and \$7.6 thousand. In column (2), more elastic location choice raises deadweight losses faster as redistribution rises; in column (4), more income-sensitive marginal utilities require less redistribution to equalize. Column (3) instead reduces the sensitivity of marginal utility to average income by half ($\varphi = 0.5$), so the planner redistributes *more* aggressively – the coefficient steepens to -66.8 and dispersion rises to \$15.2 thousand.

Column (5) assigns a Pareto weight of 4 to non-college graduates (keeping 1 for college graduates), capturing a planner who wants redistribution to reflect imperfectly signaled permanent income but whose only instrument is place-based policy. Panel B implies that a 10 percentage point higher college share raises a location's per-household optimal tax by \$4,802; Panel C shows dispersion rising from \$11.5 thousand to \$20.1 thousand. Column (6) uses the *non-shrunk* respondent-by-location weights and aggregated marginal utilities – without the empirical-Bayes precision adjustment of Section 6.1.1 – holding all else at baseline. Omitting the adjustment inflates the dispersion of measured marginal utility and hence of optimal taxes: dispersion rises to \$20.6 thousand and the Panel B coefficient steepens from -56.7 to -89.0 – the excess redistribution the adjustment is designed to attenuate.

Column (7) replaces the balanced budget with a net-revenue target of \$3,000 per household ($G = 3$). Because the formula pins down only *relative* taxes and the budget sets their common level, raising G shifts every tax up by a near-uniform \$3,010 and leaves the cross-location *pattern*

unchanged: the Panel B slopes, the dispersion, and the correlation with the baseline (1.00) are indistinguishable from column (1). The requirement shows up in the *level* – the regression intercept rises by about \$1,900, with the rest of the \$3,010 shift absorbed by the average-substitute-tax term. It is not free: the higher level is more distortionary (it raises the migration deadweight cost and lowers the insurance value, since both depend on the level of taxes, not only their differences), so the expected welfare gain falls from \$343 to \$270 per household (the ΔW row of Panel C). Finally, the Panel C correlation row measures how much each variation re-ranks *which* locations are taxed: it stays near one for the elasticity, φ , and revenue variations (which rescale or relevel the same pattern) but falls to 0.63 under the higher non-college Pareto weight – the only variation that materially re-ranks locations.

Table A5: Optimal-tax sensitivity across specifications

	(1)	(2)	(3)	(4)	(5)	(6)	(7)
	Baseline	Higher elast.	Lower φ	Higher φ	Different PW	Non-shrunk MU	Net revenue
<i>Panel A. Specification</i>							
Marginal-utility estimate	EB-shrunk μ_j	EB-shrunk μ_j	EB-shrunk μ_j	EB-shrunk μ_j	EB-shrunk μ_j	Non-shrunk μ_j	EB-shrunk μ_j
Pareto weight (non-college)	1	1	1	1	4	1	1
MU-updating parameter φ	1	1	0.5	2	1	1	1
Location-choice elasticities ν_i	1 $\times$	2 $\times$	1 $\times$	1 $\times$	1 $\times$	1 $\times$	1 $\times$
Net revenue requirement G (\$000s)	0	0	0	0	0	0	3
<i>Panel B. Regression of Optimal Tax on Location Attributes</i>							
Location MU, $\log(\hat{\mu}_j^{EB,0}/\bar{\mu})$	-56.70*** (1.867)	-35.50*** (1.231)	-66.81*** (2.290)	-39.74*** (1.233)	-52.37*** (4.402)	-88.97*** (4.535)	-56.73*** (1.863)
Avg. Tax of Substitutes	0.393*** (0.091)	0.673*** (0.077)	0.722*** (0.072)	0.088 (0.103)	0.685*** (0.120)	0.624*** (0.138)	0.392*** (0.091)
College Share	-0.497 (1.527)	-0.414 (1.017)	-1.897 (1.896)	0.288 (1.008)	48.02*** (4.837)	-4.820 (3.810)	-0.547 (1.528)
Constant	-1.483* (0.629)	-0.957* (0.418)	-1.426 (0.780)	-1.085** (0.415)	-20.34*** (1.966)	-1.929 (1.565)	0.392 (0.658)
R^2	0.866	0.874	0.881	0.863	0.684	0.736	0.866
<i>Panel C. Welfare and optimal-tax dispersion</i>							
ΔW , own criterion (\$/household)	+\$343	+\$278	+\$633	+\$162	+\$820	+\$822	+\$270
SD(τ) across locations	\$11.5k	\$7.9k	\$15.2k	\$7.6k	\$20.1k	\$20.6k	\$11.5k
SD(τ), pop.-weighted	\$8.8k	\$6.1k	\$12.1k	\$5.5k	\$16.4k	\$14.6k	\$8.8k
Corr. with baseline τ^* (pop.-wtd)	1.00	0.99	0.99	0.99	0.63	0.93	1.00

Standard errors in parentheses

* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

Notes: Column (1) is the baseline; columns (2)–(7) each vary one ingredient at a time, as described in the appendix text. Panel A gives the parameter configuration. Panel B regresses each location's optimal tax τ on its baseline log relative marginal utility $\log(\hat{\mu}_j^{EB,0}/\bar{\mu})$ (the EB-adjusted MU used as the x -axis of Figure 5, identical across columns), the choice-probability-weighted average tax of its substitute locations, and its college-graduate share, with standard errors in parentheses. Panel C reports the own-criterion welfare gain ΔW (per-household dollars), the dispersion of optimal τ (in \$k), and its population-weighted correlation with the column-(1) taxes.