

UC Berkeley

UC Berkeley Previously Published Works

Title

Testing theories of financial decision making

Permalink

<https://escholarship.org/uc/item/87f2z6cx>

Journal

Proceedings of the National Academy of Sciences of the United States of America,
113(15)

ISSN

0027-8424

Authors

Chambers, Christopher P
Echenique, Federico
Saito, Kota

Publication Date

2016-04-12

DOI

10.1073/pnas.1517760113

Peer reviewed

Testing theories of financial decision making

Christopher P. Chambers^a, Federico Echenique^b, and Kota Saito^{b,1}

^aDepartment of Economics, University of California, San Diego, La Jolla, CA 92093; and ^bDivision of the Humanities and Social Sciences, California Institute of Technology, Pasadena, CA 91125

Edited by Matthew O. Jackson, Stanford University, Stanford, CA, and approved February 22, 2016 (received for review September 8, 2015)

We describe the observable content of some of the most widely used models of decision under uncertainty: models of translation invariant preferences. In particular, we characterize the models of variational, maxmin, constant absolute risk aversion, and constant relative risk aversion utilities. In each case we present a revealed preference axiom that is satisfied by a dataset if and only if the dataset is consistent with the corresponding utility representation. We test our axioms using data from an experiment on financial decisions.

revealed preference | translation invariance | homotheticity | uncertainty

This paper is an investigation of the testable implications of models of decision under uncertainty. We carry out this investigation in financial markets, one of the most common environments in which human subjects face uncertainty.

Risk is uncertainty that can be objectively quantified probabilistically. A gambler in a casino faces risk: He may calculate the probability that a roulette wheel stops on the number 7, or that a die lands on 5. Most scientists, in contrast, face the more general concept of uncertainty and study subjects who face uncertainty. Scientists conduct or analyze experiments with outcomes that they do not know, and for which no probabilities are objectively given.

Of course, the scientist or the subject may have a subjective judgment of how likely different events are. Such judgments may even have a probabilistic expression, but the uncertainty is not resolved by means of a mechanical device for which probabilities can be objectively calculated. Moreover, this fact may cause subjects to display uncertainty aversion, a tendency to prefer risky bets over uncertain ones. Uncertainty aversion was famously documented by David Ellsberg (1), and the theories we treat in our paper are in part designed to describe uncertainty aversion.

In uncertain situations, human subjects choose among uncertain prospects. These are functions specifying an “outcome” for each element of a given set of “states of the world.” Think of an insurance contract that pays off a given sum only if some accident occurs. The set of states of the world is the binary set that codifies whether an accident has occurred, and the outcome is the payoff. In financial markets, the uncertain prospects correspond to financial assets, whereas the state of the worlds describe the relevant economic fundamentals, and the outcomes monetary payoffs.*

A long tradition in decision theory develops models of how humans make decisions under uncertainty. A crucial idea in this development is that of translation invariance. Translation invariance means that if two uncertain prospects are transformed in the same way, by adding to each prospect a given, fixed, monetary payment, then the subject’s preference between the two prospects should be preserved. For example, if the subject prefers insurance contract A over B, then the preference should be maintained after the price of each insurance contract has been raised by the same amount. A related idea is homotheticity, where scaling the payoffs of the two contracts should not affect how they are ranked. Translation invariance and homotheticity give rise to different theories of decision under uncertainty.

Theories demand to be tested, and our contribution lies in working out the testable implications of theories of homothetic and translation invariant behavior under uncertainty. We focus on financial markets because these are some of the most familiar and common uncertain environments for human subjects. If one is to test a theory, it makes sense to study it in the subjects’ most familiar

environments. It is plausible that agents do not know how to behave in an artificial environment, but that they have learned how to deal with uncertainty in familiar environments. For human subjects, few uncertain environments are as familiar as financial markets. Most existing experimental environments are artificial: They involve human subjects choosing among bets on extractions of colored balls from urns of uncertain composition [Ellsberg’s thought experiments are the best known of these (1)]. Our contribution is instead to focus on designs based on financial markets.

Our main results characterize the financial datasets that are consistent with the theories. Given is a finite collection of data on purchases of financial assets. The question is, When are such data consistent with a theory of choice under uncertainty? We provide answers for some of the most commonly encountered theories, those based on translation invariance and homotheticity.

We show that our results are applicable to the analysis and design of experiments by using a recent experiment by Hey and Pace (2). Hey and Pace have subjects decide on purchases of financial assets. We use the data they collect to test for consistency with maxmin expected utility, a theory of decision under uncertainty based on translation invariance and homotheticity. The conclusion of our analysis is that Hey and Pace’s data reject the maxmin theory. The finding is preliminary and meant mainly as an illustration of our methods, but if confirmed it would mean that some of the best-known theories of choice under uncertainty, theories that are thought of as weak, and accommodating of diverse behavioral and psychological phenomena, do not in fact stand up to empirical scrutiny on data from financial experiments.

The theories covered by our results include risk-neutral variational preferences (3), risk-neutral maxmin preferences (4), and subjective expected utility preferences with constant absolute risk aversion: so-called CARA preferences. Analogously to the CARA case, we also work out the testable implications of subjective expected utility preferences with constant relative risk aversion: so-called CRRA preferences. The theories have been used for different purposes. Variational and maxmin preferences are the most commonly used models of uncertainty aversion (3–6). They are also used to capture model robustness (7). CARA and CRRA preferences are extremely common in applied work in macroeconomics and finance, among other fields.

Significance

This paper uncovers the empirical content of many behavioral models of decisions under uncertainty. Studies of global financial markets often ignore a very important piece of the puzzle: individual behavior. We provide tests (and as such, predictions) about how individuals behave when facing uncertainty, such as that faced in financial or asset markets. Behavior at the individual level must be understood before the behavior of the economy at large can even begin to be understood.

Author contributions: C.P.C., F.E., and K.S. designed research, performed research, analyzed data, and wrote the paper.

The authors declare no conflict of interest.

This article is a PNAS Direct Submission.

¹To whom correspondence should be addressed. Email: saito@caltech.edu.

*In probability theory, an uncertain prospect together with an underlying probability over the states of the world is termed a “random variable.”

The authors of refs. 8, 9, and 10 carried out exercises similar to ours, also focusing on financial market experiments, but in a context of risk, not uncertainty. The closest papers to ours are those given in refs. 11, 12, and 13. Echenique and Saito (11) studied the case of subjective expected utility but did not address the more general theories studied here, and that have been proposed to address the empirical shortcomings of subjective expected utility. Bayer et al. (12) and Polisson and Quah (13) treat some of the same theories we do but give a characterization in terms of the solution of a system of inequalities. We give a revealed preference axiom (a characterization that references only observable data) that has to be satisfied for the data to be rationalizable. It can be written in the UNCAF (universal negation of concatenations of atomic formulas) form, which is the kind of axiom that characterizes the empirical content of a theory (14). A system of (nonlinear) inequalities may not give an economic interpretation to the characterization, and it may not be computationally feasible.[†]

Definitions

Let S be a finite set of states of the world. An act is a function from S into $\mathbf{R}$; $\mathbf{R}^S$ is the set of acts. An act can be interpreted as a state-contingent monetary payment. Define $\|x\|_1 = \sum_s x_s$. $\Delta(S)$ represents the set of probability distributions on S , that is, $\Delta(S) = \{\pi \in \mathbf{R}_+^S : \sum_s \pi_s = 1\}$.

A preference relation on $\mathbf{R}^S$ is a complete and transitive binary relation $\geq$; we denote by $>$ the strict part of $\geq$. A function $u : \mathbf{R}^S \rightarrow \mathbf{R}$ defines a preference relation $\geq$ by $x \geq y$ if and only if $u(x) \geq u(y)$. We say that u represents $\geq$, or that it is a utility function for $\geq$. A preference relation $\geq$ on $\mathbf{R}^S$ is locally nonsatiated if for every x and every $\epsilon > 0$ there is y such that $\|x - y\| < \epsilon$ and $y > x$.

Preferences, Utilities, and Data

A dataset D is a finite collection $\{(p^k, x^k)\}_{k=1}^K$, where each $p^k \in \mathbf{R}_{++}^S$ is a vector of strictly positive (so-called Arrow–Debreu) prices, and each $x^k \in \mathbf{R}^S$ is an act. The interpretation of a dataset is that each pair (p^k, x^k) consists of an act x^k chosen from the budget $\{x \in \mathbf{R}^S : p^k \cdot x \leq p^k \cdot x^k\}$ of affordable acts. Such datasets are common in financial markets experiments (2, 12, 15).

A dataset $\{(p^k, x^k)\}_{k=1}^K$ is rationalizable by a preference relation $\geq$ if $x^k \geq x$ whenever $p^k \cdot x^k \geq p^k \cdot x$. So, a dataset is rationalizable by a preference relation when the choices in the dataset would have been optimal for that preference relation. A dataset $\{(p^k, x^k)\}_{k=1}^K$ is rationalizable by a utility function u if it is rationalizable by the preference relation represented by u . So, a dataset is rationalizable by a utility function when the choices in the dataset would have maximized that utility function in the relevant budget set.

A preference relation $\geq$ is translation invariant if for all $x, y \in \mathbf{R}^S$ and all $c \in \mathbf{R}$, we have $x \geq y$ if and only if $x + (c, \dots, c) \geq y + (c, \dots, c)$.

A preference relation $\geq$ is homothetic if for all $x, y \in \mathbf{R}^S$ and all $\alpha > 0$, we have $x \geq y$ if and only if $\alpha x \geq \alpha y$.

A preference relation $\geq$ is a risk-neutral variational preference if there is a convex and lower semicontinuous function $c : \Delta(S) \rightarrow \mathbf{R} \cup \{+\infty\}$ such that the function

$$\inf_{\pi \in \Delta(S)} \pi \cdot x + c(\pi)$$

represents $\geq$. If a dataset is rationalizable by a risk-neutral variational preference relation, we will say that the dataset set is risk-neutral variational-rationalizable.

A special case of variational preference is maxmin: A preference relation is risk-neutral maxmin if there is a closed and convex set $\Pi \subseteq \Delta(S)$ such that the utility function

$$\inf_{\pi \in \Pi} \pi \cdot x$$

represents $\geq$. If a dataset is rationalizable by a risk-neutral maxmin preference relation, we will say that the dataset set is risk-neutral maxmin-rationalizable. More generally, a preference relation is risk-averse maxmin if there is a closed and convex set $\Pi \subseteq \Delta(S)$, where for each $\pi \in \Pi$ and each $s \in S$, $\pi_s > 0$, and a concave utility $u : \mathbf{R} \rightarrow \mathbf{R}$ such that the utility function

$$\inf_{\pi \in \Pi} \sum_{s=1,2} \pi_s u(x_s)$$

represents $\geq$. If a dataset is rationalizable by a maxmin preference relation, we will say that the dataset set is maxmin-rationalizable.

A utility $u : \mathbf{R}^S \rightarrow \mathbf{R}$ is CARA if there is $\alpha > 0$ and $\pi \in \Delta(S)$ for which for all $s \in S$, $\pi_s > 0$, and

$$u(x) = \sum_{s \in S} \pi_s (-\exp(-\alpha x_s)).$$

Note that CARA is a special case of subjective expected utility.[‡]

A utility $u : \mathbf{R}^S \rightarrow \mathbf{R}$ is CRRA if there is $\alpha \in (0, 1)$ and $\pi \in \Delta(S)$ for which for all $s \in S$, $\pi_s > 0$, and

$$u(x) = \sum_{s \in S} \pi_s \left(\frac{x_s^{1-\alpha}}{1-\alpha} \right).$$

If a dataset is rationalizable by a CARA (CRRA) utility, we will say that the dataset set is CARA (CRRA) rationalizable.

Variational and Maxmin Preferences

We present the results on variational and maxmin rationalizability as *Theorems 1* and 2. These models satisfy the hypothesis that for any x, y , $x \sim y \Rightarrow \frac{1}{2}x + \frac{1}{2}y \geq y$. This hypothesis is known as convexity of preference. Convexity is related to uncertainty aversion in the sense of ref. 4. In fact, given the assumptions of monotonicity found in that paper, together with the assumption that the preference is risk-neutral (i.e., lotteries are evaluated according to their expected value), it is equivalent to uncertainty aversion. Uncertainty aversion is the idea that an agent dislikes uncertainty and suffers from his or her ignorance of the possible probability distribution that governs outcomes.

One important conclusion that emerges from our analysis is that convexity is not testable with market data. This therefore means that under the maintained hypothesis of risk neutrality (and monotonicity), uncertainty aversion cannot be detected with financial data.

Theorem 1. *The following statements are equivalent:*

- Dataset D is rationalizable by a locally nonsatiated, translation invariant preference.
- Dataset D is rationalizable by a continuous, strictly increasing, concave utility function satisfying the property $u(x + (c, \dots, c)) = u(x) + c$.
- Dataset D is risk-neutral variational-rationalizable.
- For every $l = 1, \dots, M$, and every sequence $\{k_l\} \subseteq \{1, \dots, K\}$,

$$\sum_{l=1}^M \frac{p^{k_l}}{\|p^{k_l}\|_1} \cdot (x^{k_{l+1}} - x^{k_l}) \geq 0,$$

where addition is modulo M , as usual.

Note that the equivalence between *ii* and *iii* is due to ref. 3.

[†]The paper by Bayer et al. (12) is a case in point, where the solution to the system of inequalities is implemented by a grid search. A conclusive test is not possible since they results depend on the assumed granularity of the grid.

[‡]In fact, it is also a special case of a risk neutral variational preference, a fact exploited by ref. 16.

Remark: The fact that i implies ii and iii implies that if a dataset is rationalizable by a translation invariant preference, the dataset is also rationalizable by a risk-neutral variational preference (which automatically satisfies convexity).

Remark: The preceding result can be generalized. Suppose we were interested in the testable implications of preferences that are β -translation invariant, for some $\beta \geq 0$, $\beta \neq 0$. That is, we want to know whether for all x, y , we have $x \succ y$ if and only if for all t , $x + t\beta \succ y + t\beta$. Define the seminorm $\|x\|_1^\beta = \sum_i |\beta x_i|$. Then it is an easy exercise to verify that the testable implications of β -translation invariance are given by Eq. 4, replacing $\|\cdot\|_1$ with $\|\cdot\|_1^\beta$.

Remark: The test in ref. 4 is related to cyclic monotonicity. This is similar to the test given in ref. 17 for quasilinear preferences (and to a result in ref. 18).

We now turn our attention to maxmin preferences.

We say that a function $u: \mathbf{R}^S \rightarrow \mathbf{R}$ is linearly homogeneous if for all $x \in \mathbf{R}^S$ and all $\alpha > 0$, we have $u(\alpha x) = \alpha u(x)$.

Theorem 2. The following statements are equivalent:

- Dataset D is rationalizable by a locally nonsatiated, homothetic and translation invariant preference.
- Dataset D is rationalizable by a continuous, strictly increasing, linearly homogeneous and concave utility function satisfying the property that $u(x + (c, \dots, c)) = u(x) + c$.
- Dataset D is risk-neutral maxmin-rationalizable.
- For every k and l ,

$$\frac{p^k}{\|p^k\|_1} \cdot x^k \leq \frac{p^l}{\|p^l\|_1} \cdot x^k.$$

The equivalence between ii and iii is due to ref. 4.

It is interesting to note that, just as in *Theorem 1*, under the maintained hypotheses of risk aversion and monotonicity, uncertainty aversion has no content for behavior.

Remark: The rationalizing variational and maxmin preferences can be taken to imply “full support” priors. In the proof of *Theorem 1*, we shown that there is $\pi \in \Delta(S)$ satisfying $c(\pi) < +\infty$, which implies for all $s \in S$, $\pi_s > 0$. In the proof of *Theorem 2* we show that for each $\pi \in \Pi$ and all $s \in S$, $\pi_s > 0$.

CARA and CRRA

The previous section considers translation invariance and homotheticity as general properties of preferences in choice under uncertainty. Here we focus on the case of subjective expected utility. So, we consider models in which the agent has a single prior over states, and maximizes his or her expected utility. The prior is unknown and must be inferred from his or her choices. Translation invariance gives rise to CARA preferences, and homotheticity to CRRA.

Theorem 3. A dataset D is CARA rationalizable if and only if there is $\alpha^* > 0$ such that Eq. 1 holds for all $k, k' \in K$ and $s, t \in S$ and CRRA rationalizable if and only if there is $\alpha^* \in (0, 1)$ such that Eq. 2 holds for all $k, k' \in K$ and $s, t \in S$.

$$\alpha^* (x_t^k - x_s^k + x_s^{k'} - x_t^{k'}) = \log \left(\frac{p_s^k p_t^{k'}}{p_t^k p_s^{k'}} \right) \quad [1]$$

$$\alpha^* \log \left(\frac{x_t^k x_s^{k'}}{x_s^k x_t^{k'}} \right) = \log \left(\frac{p_s^k p_t^{k'}}{p_t^k p_s^{k'}} \right) \quad [2]$$

The conditions in *Theorem 3* may look like existential conditions: essentially Afriat inequalities. Afriat inequalities are indeed the source of Eqs. 1 and 2, as evidenced by the proof of *Theorem 3*, but note that the statements are equivalent to nonexistential statements: Eq. 1 says that when $(x_t^k - x_s^k + x_s^{k'} - x_t^{k'}) \neq 0$,

$$\frac{\log \left(\frac{p_s^k p_t^{k'}}{p_t^k p_s^{k'}} \right)}{(x_t^k - x_s^k + x_s^{k'} - x_t^{k'})}$$

is independent of k, t, k' and s , and that when $(x_t^k - x_s^k + x_s^{k'} - x_t^{k'}) = 0$ then $\log \left(\frac{p_s^k p_t^{k'}}{p_t^k p_s^{k'}} \right) = 0$. Similarly for Eq. 2.

It is worth pointing out that, except in the case when for all observations, all prices are equal, and consumption of all goods are equal, Eq. 1 can have only one solution. Hence, risk preferences are uniquely identified.

The next corollary also shows that beliefs are identified. Recall that a CARA utility is defined by a pair (a, π) , with $a > 0$ and $\pi \in \Delta(S)$.

Corollary. If (a, π) and (a', π') define CARA utilities that rationalize D , then $(a, \pi) = (a', \pi')$. Furthermore, $a = a'$ coincide with the unique solution to Eq. 1. Similarly for CRRA rationalizability and Eq. 2.

Risk-Averse Maxmin with Two States

Theorem 2 is about risk-neutral maxmin. Here we turn to maxmin with risk aversion. In this section, we assume that there are two states (i.e., $S = \{1, 2\}$). A preference relation is maxmin if there is a closed and convex set $\Pi \subseteq \Delta(S)$, where for each $\pi \in \Pi$ and each $s \in S$, $\pi_s > 0$, and a concave utility $u: \mathbf{R} \rightarrow \mathbf{R}$ such that the utility function

$$\inf_{\pi \in \Pi} \sum_{s=1,2} \pi_s u(x_s)$$

represents $\succ$. If a dataset is rationalizable by a maxmin preference relation, we will say that the dataset set is maxmin-rationalizable.

Let K_0 be the set of all k such that $x_1^k = x_2^k$. Let K_1 be the set of all k such that $x_1^k < x_2^k$, and K_2 be the set of all k such that $x_1^k > x_2^k$. Note that $K = K_0 \cup K_1 \cup K_2$.

Say that a sequence of pairs $(x_{s_i}^{k_i}, x_{s'_i}^{k'_i})_{i=1}^n$ is balanced if each k appears as k_i (on the left of the pair) the same number of times it appears as k'_i (on the right).

Given a sequence of pairs $(x_{s_i}^{k_i}, x_{s'_i}^{k'_i})_{i=1}^n$, consider the following notation: Let $I_{l,s} = \{i: k_i \in K_l \text{ and } s_i = s\}$, $I'_{l,s} = \{i: k'_i \in K_l \text{ and } s'_i = s\}$, for $l = 0, 1, 2$ and $s = 1, 2$.

Axiom: Strong Axiom of Revealed Maxmin Expected Utility. For any balanced sequence of pairs $(x_{s_i}^{k_i}, x_{s'_i}^{k'_i})_{i=1}^n$ in which

- $x_s^{k_i} > x_{s'_i}^{k'_i}$ for all i ;
- $\#I_{0,1} + \#I_{1,1} - \#I'_{1,1} = \#I'_{0,1} + \#I_{2,1} - \#I_{2,1} \leq 0$

The product of price ratios satisfies

$$\prod_{i=1}^n \frac{p_{s_i}^{k_i}}{p_{s'_i}^{k'_i}} \leq 1. \quad [3]$$

Theorem 4. A dataset is maxmin rationalizable if and only if it satisfies strong axiom of revealed maxmin expected utility (SARMEU).

Echenique and Saito (11) show that a stronger axiom, strong axiom of revealed subjective expected utility (SARSEU), characterizes rationalizability by subjective expected utility. Instead of condition ii of SARMEU, SARSEU requires

$$\#I_{0,1} + \#I_{1,1} + \#I_{2,1} = \#I'_{0,1} + \#I'_{1,1} + \#I'_{2,1}. \quad [4]$$

Theorem 4 is useful because it makes explicit what one would need to see in an experiment (with two states, a common setup in laboratory experiments) in order for choices to be consistent with maxmin utility, but inconsistent with subjective expected utility. For

$$\sum_{j=1}^{M'} \frac{p^{k_j}}{\|p^{k_j}\|_1} \cdot (x^{k_{j+1}} - x^{k_j}) \leq \sum_{l=1}^M \frac{p^{k_l}}{\|p^{k_l}\|_1} \cdot (x^{k_{l+1}} - x^{k_l}).$$

Therefore, $u(x) = \inf \left\{ \sum_{l=1}^M \frac{p^{k_l}}{\|p^{k_l}\|_1} \cdot (x^{k_{l+1}} - x^{k_l}) : \{k_l\}_{l=1}^M \in \Sigma_x \right\}$ is well defined, because the infimum can be taken over a finite set.

That $u: \mathbf{R}^S \rightarrow \mathbf{R}$ defined in this fashion is concave, strictly increasing and continuous is immediate. To see that it rationalizes the data, suppose that $p^k \cdot x^l \leq p^k \cdot x^k$. Then $\frac{p^k}{\|p^k\|_1} \cdot x^l \leq \frac{p^k}{\|p^k\|_1} \cdot x^k$. It is clear then by definition that $u(x^l) \leq u(x^k) + \frac{p^k}{\|p^k\|_1} \cdot (x^l - x^k) \leq u(x^k)$.

Finally, to show that $u(x + (c, \dots, c)) = u(x) + c$, note that for any p^k , we have $\frac{p^k}{\|p^k\|_1} \cdot (x + (c, \dots, c)) = c + \frac{p^k}{\|p^k\|_1} \cdot x$. The result then follows by construction.

We end the proof by showing that $ii \Rightarrow iii$. Let $u: \mathbf{R}^S \rightarrow \mathbf{R}$ be as in the statement of ii . Define the concave conjugate of u by

$$\begin{aligned} f(\pi) &= \inf \{ \pi \cdot x - u(x) : x \in \mathbf{R}^S \} \\ &= \inf \{ \pi \cdot x + c \pi \cdot \mathbf{1} - u(x) - c : x \in \mathbf{R}^S, c \in \mathbf{R} \} \\ &= \inf \{ \pi \cdot x - c(1 - \pi \cdot \mathbf{1}) - u(x) : x \in \mathbf{R}^S, c \in \mathbf{R} \}, \end{aligned}$$

where the second equality uses that $u(x + (c, \dots, c)) = u(x) + c$. Now note that $f(\pi) = -\infty$ if $(1 - \pi \cdot \mathbf{1}) \neq 0$. Note also that the monotonicity of u implies that $f(\pi) = -\infty$ if there is s such that $\pi_s < 0$. One can also show that there is $\pi \in \Delta(S)$ for which $f(\pi) \in \mathbf{R}$.¹¹ Finally, observe that by strict monotonicity, if there is $s \in S$ for which $\pi_s = 0$, then $f(\pi) = -\infty$. Hence, we can consider the domain of f to be a subset of $\Delta(S)$. Moreover, $f(\pi) < +\infty$ implies for all $s \in S$, $\pi_s > 0$.

Now, since u is continuous, it is a standard application of the separating hyperplane theorem to establish that $u(x) = \inf_{\pi \in \Delta(S)} \pi \cdot x - f(\pi)$. Because u rationalizes the dataset, the dataset is variational rationalizable.

Proof of Theorem 2: It is obvious that $iii \Rightarrow ii$ and that $ii \Rightarrow i$. Hence, to show the theorem, it suffices to show that iv implies iii and that i implies iv .

For a dataset D , let $\pi^k = \frac{p^k}{\|p^k\|_1}$. It is easy to see that $iv \Rightarrow iii$. Let Π be the convex hull of $\{\pi^k : k = 1, \dots, K\}$. Then, it is immediate that $u(x) = \min_{\pi \in \Pi} \pi \cdot x$ rationalizes D . Moreover, for each $\pi \in \Pi$ and all $s \in S$, $\pi_s > 0$ because $\pi_s^k > 0$ for all $s \in S$ and $k \in K$.

We prove that $i \Rightarrow iv$. Suppose that D satisfies i but not iv . Then, there are k and l for which $\pi^l \cdot x^k < \pi^k \cdot x^k$. Let $\succsim$ be a preference relation as stated in i . By homotheticity of $\succsim$, for any scalar $\theta > 0$, $\succsim$ rationalizes the data $D' \equiv \{(x^j, \pi^j) : j = 1, \dots, K\} \cup \{(\theta x^l, \pi^l)\}$. To see this, observe that if $\pi^l \cdot x \leq \pi^l \cdot \theta x^l$, then $\pi^l \cdot \theta^{-1} x \leq \pi^l \cdot x^l$, so that $x^l \succsim \theta^{-1} \cdot x$, and by homogeneity, $\theta x^l \succsim x$. Now, for $\theta > 0$ sufficiently small, $\pi^l \cdot x^k < \pi^k \cdot x^k$ implies that

$$x^k \cdot (\pi^l - \pi^k) + \theta x^l \cdot (\pi^k - \pi^l) < 0.$$

So either $x^k \cdot (\pi^l - \pi^k) < 0$ or $\theta x^l \cdot (\pi^k - \pi^l) < 0$. Then the dataset D' violates iv in Theorem 1, contradicting the fact that it is rationalized by $\succsim$, which is assumed to be translation invariant.

Proof of Theorem 3: The idea in the proof is to solve the first-order conditions for the unknown terms. Consider first the case of CARA. Let $\pi \in \Delta(S)$ and $\alpha > 0$ rationalize D . Then, we know that x^k maximizes $\sum_s \pi_s (-\exp(-\alpha x_s))$ subject to $p^k \cdot x \leq p^k \cdot x^k$. By considering the Lagrangean and the first-order conditions, we may conclude that for every $s, t \in S$ and every $k \in \{1, \dots, K\}$, we have

¹¹For example, take π to support $\{z \in \mathbf{R}^S : u(z) \geq 0\}$ at 0. We claim that $f(\pi) \geq 0$. Suppose by means of contradiction that there is $x \in \mathbf{R}^S$ for which $\pi \cdot x < u(x)$. Observe that π supports $\{z \in \mathbf{R}^S : u(z) \geq \pi \cdot x\}$ at the act y , which returns $\pi \cdot x$ in each state. Observe that $u(z) > \pi \cdot x$ implies $\pi \cdot z > \pi \cdot x$, by continuity of u and definition of the supporting hyperplane, that is, $\{z \in \mathbf{R}^S : u(z) \geq \pi \cdot x\} \subseteq \{z \in \mathbf{R}^S : \pi \cdot z \geq \pi \cdot x\}$ implies $\{z \in \mathbf{R}^S : u(z) > \pi \cdot x\} \subseteq \{z \in \mathbf{R}^S : \pi \cdot z > \pi \cdot x\}$ because the latter sets are the interiors of the former. Therefore, if $u(x) > \pi \cdot x$, we conclude $\pi \cdot x > \pi \cdot x$, a contradiction.

$$\frac{\pi_s \exp(-\alpha x_s^k)}{p_s^k} = \frac{\pi_t \exp(-\alpha x_t^k)}{p_t^k}.$$

Conclude that $\frac{\pi_s^k \pi_t}{p_t^k \pi_s} = \exp(-\alpha(x_s^k - x_t^k))$. By taking logs, the system becomes

$$\log(\pi_s) - \log(\pi_t) + \alpha(x_t^k - x_s^k) = \log(p_s^k) - \log(p_t^k). \quad [5]$$

In the case of CRRA, the existence of a rationalizing π and parameter α imply a first-order condition of the form

$$\log(\pi_s) - \log(\pi_t) + \alpha \log(x_t^k / x_s^k) = \log(p_s^k) - \log(p_t^k). \quad [6]$$

We can denote $\log(\pi_s)$ by z_s in Eqs. 5 and 6. Thus, we obtain that D is rationalizable if and only if there exist $z_s \in \mathbf{R}$ and $\alpha > 0$ such that the following equation is solved for all s, t, k with $s \neq t$:

$$z_s - z_t + \alpha(y_t^k - y_s^k) = \log(p_s^k) - \log(p_t^k),$$

where $y_t^k = x_t^k$ for CARA rationalizability, and $y_t^k = \log x_t^k$ for CRRA rationalizability.

Now the necessity of the axioms is obvious. Let $k \neq k'$, then

$$\alpha(y_t^k - y_s^k) - \log(p_s^k / p_t^k) = z_s - z_t = \alpha(y_t^{k'} - y_s^{k'}) - \log(p_s^{k'} / p_t^{k'})$$

for any s and t . Thus

$$\alpha(y_t^k - y_s^k - y_t^{k'} + y_s^{k'}) = \log\left(\frac{p_s^k p_t^{k'}}{p_t^k p_s^{k'}}\right).$$

So, i is satisfied for the case of CARA rationalizability, and ii is satisfied for the case of CRRA rationalizability.

To prove sufficiency, let

$$\begin{aligned} d^p(s, t, k) &= \log(p_s^k / p_t^k) \\ d^x(s, t, k) &= y_s^k - y_t^k. \end{aligned}$$

Let α^* be such that for all k, k', s, s' and t ,

$$\alpha^*(y_t^k - y_s^k - y_t^{k'} + y_s^{k'}) = \log\left(\frac{p_s^k p_t^{k'}}{p_t^k p_s^{k'}}\right).$$

Then in particular, for all k, k', s, s' and t ,

$$d^p(s, t, k) + \alpha^* d^x(s, t, k) + d^p(t, s, k') + \alpha^* d^x(t, s, k') = 0. \quad [7]$$

Note also that

$$\begin{aligned} d^p(s, t, k) + d^p(t, s', k) + d^p(s', s, k) \\ + \alpha^*(d^x(s, t, k) + d^x(t, s', k) + d^x(s', s, k)) \end{aligned} \quad [8]$$

Fix $s_0 \in S$ and let $z_{s_0} \in \mathbf{R}$ be arbitrary. For any $s \in S$, define z_s by

$$z_s = z_{s_0} + \alpha^* d^x(s_0, s, k) + d^p(s, s_0, k),$$

for some k . In fact, by Eq. 7, this definition is independent of k because $d^p(s, s_0, k) + \alpha^* d^x(s, s_0, k) = d^p(s, s_0, k') + \alpha^* d^x(s, s_0, k')$.

Given this definition, note that

$$\begin{aligned} z_s - z_t &= \alpha^*(d^x(s_0, s, k) - d^x(s_0, t, k)) + d^p(s, s_0, k) - d^p(t, s_0, k) \\ &= \alpha^*(d^x(s_0, s, k) - d^x(s_0, t, k)) + d^p(s, s_0, k) - d^p(t, s_0, k) \\ &\quad + d^p(s, t, k) + d^p(t, s_0, k) + d^p(s_0, s, k) + \alpha^*(d^x(s, t, k) \\ &\quad + d^x(t, s_0, k) + d^x(s_0, s, k)) = d^p(s, t, k) + \alpha^* d^x(s, t, k), \end{aligned}$$

where the second equality uses Eq. 8.

Hence, with the constructed $(z_t)_{t \in S}$ we have

$$z_s - z_t + \alpha^*(y_t^k - y_s^k) = \log(p_s^k / p_t^k),$$

for all s, t , and k . The first-order conditions for rationalizability are therefore satisfied.

ACKNOWLEDGMENTS. We thank Simon Grant and Massimo Marinacci, who posed questions to us that led to some of the results in this paper. We also thank John Hey and Noemi Pace for allowing us to use

their data and Taisuke Imai for carrying out our test. This work was supported by National Science Foundation Grant SES-1426867 (to C.P.C.).

1. Ellsberg D (1961) Risk, ambiguity, and the Savage axioms. *Q J Econ* 75(4):643–669.
2. Hey J, Pace N (2014) The explanatory and predictive power of non two-stage-probability theories of decision making under ambiguity. *J Risk Uncertain* 49(1):1–29.
3. Maccheroni F, Marinacci M, Rustichini A (2006) Ambiguity aversion, robustness, and the variational representation of preferences. *Econometrica* 74(6):1447–1498.
4. Gilboa I, Schmeidler D (1989) Maxmin expected utility with non-unique prior. *J Math Econ* 18(2):141–153.
5. Cerreia-Vioglio S, Maccheroni F, Marinacci M, Montrucchio L (2013) Classical subjective expected utility. *Proc Natl Acad Sci USA* 110(17):6754–6759.
6. Cerreia-Vioglio S, Maccheroni F, Marinacci M, Montrucchio L (2011) Uncertainty averse preferences. *J Econ Theory* 146(4):1275–1330.
7. Hansen LP, Sargent TJ (2008) *Robustness* (Princeton Univ Press, Princeton).
8. Green RC, Srivastava S (1986) Expected utility maximization and demand behavior. *J Econ Theory* 38(2):313–323.
9. Varian HR (1988) Estimating risk aversion from Arrow-Debreu portfolio choice. *Econometrica* 43(4):973–979.
10. Kubler F, Selden L, Wei X (2014) Asset demand based tests of expected utility maximization. *Am Econ Rev* 104(11):3459–3480.
11. Echenique F, Saito K (2015) Savage in the market. *Econometrica* 83(4):1467–1495.
12. Bayer R, Bose S, Polisson M, Renou L (2012) Ambiguity revealed (Univ of Essex, Essex, UK).
13. Polisson M, Quah J (2013) Revealed preference tests under risk and uncertainty. Working Paper 13/24 (Univ of Leicester, Leicester, UK).
14. Chambers CP, Echenique F, Shmaya E (2014) The axiomatic structure of empirical content. *Am Econ Rev* 104(8):2303–2319.
15. Ahn D, Choi S, Gale D, Kariv S (2014) Estimating ambiguity aversion in a portfolio choice experiment. *Quant Econom* 5(2):195–223.
16. Strzalecki T (2011) Axiomatic foundations of multiplier preferences. *Econometrica* 79(1):47–73.
17. Brown DJ, Calsamiglia C (2007) The nonparametric approach to applied welfare analysis. *Econ Theory* 31(1):183–188.
18. Rockafellar RT (1997) *Convex Analysis* (Princeton Univ Press, Princeton).
19. Schmeidler D (1989) Subjective probability and expected utility without additivity. *Econometrica* 57(3):571–587.
20. Chambers CP, Echenique F, Saito K (2015) Testable implications of translation invariance and homotheticity: Variational, maxmin, cara and CRRRA preferences. Working paper (California Institute of Technology, Pasadena, CA). Available at www.hss.caltech.edu/content/testable-implications-translation-invariance-and-homotheticity-variational-maxmin-cara-and.