

Estimating Risk Preferences in the Field*

Levon Barseghyan
Cornell University

Francesca Molinari
Cornell University

Ted O'Donoghue
Cornell University

Joshua C. Teitelbaum
Georgetown University

February 28, 2015

Prepared for the *Journal of Economic Literature*

COMMENTS WELCOME

Abstract

We survey the literature on estimating risk preferences using field data—i.e., data on individuals' real-world economic behavior. We mostly limit our attention to studies in which risk preferences are the focal object and estimating their structure is the core enterprise. We motivate shining a spotlight on this literature with an example that highlights why knowledge about the structure of risk preference matters for applied work in economics. To further set the stage, we provide a detailed review of a number of models of risk preferences—including both expected utility (EU) theory and non-EU models—that have been estimated using field data, and we discuss issues related to identification and estimation of such models using field data. We then survey the literature, giving separate treatment to research that uses individual-level data (e.g., property insurance data) and research that uses aggregate data (e.g., betting market data). We conclude by discussing work at the frontier of the literature and directions for future research.

*We thank Steven Durlauf and two anonymous referees for comments, and Heidi Verheggen and Lin Xu for excellent research assistance. Financial support from the National Science Foundation (SES-1031136) is gratefully acknowledged. In addition, Molinari acknowledges financial support from National Science Foundation grant SES-0922330.

Contents

1	Introduction	1
2	Motivating Example	3
2.1	Out-of-Sample Predictions	4
2.2	Welfare Analysis	6
2.3	Identification and Estimation Challenges	7
3	Models of Risk Preferences	8
3.1	Expected Utility (EU)	9
3.2	Rank-Dependent Expected Utility (RDEU)	12
3.3	Cumulative Prospect Theory (CPT)	16
3.4	Koszegi-Rabin Loss Aversion	18
3.5	Disappointment Aversion	20
3.6	Model Predictions and Identification	22
3.6.1	Expected Utility	23
3.6.2	Rank Dependent Expected Utility	24
3.6.3	EU and RDEU with Heterogeneity in Preferences	26
3.6.4	Distinguishing More Complex Models	27
4	Estimation with Individual-Level Data	30
4.1	The General Approach	30
4.1.1	Translation into Lotteries	30
4.1.2	Introducing “Noise”	32
4.1.3	Moral Hazard and Adverse Selection	34
4.2	Property Insurance Data	36
4.3	Game Shows	44
4.4	Health Insurance Data	46
5	Estimation with Aggregate Data	52
5.1	The General Approach	52
5.2	Betting Markets	52
5.2.1	Representative Agent Framework	53
5.2.2	Heterogeneity in Beliefs and Preferences	60
5.3	Consumption, Asset Returns, and Labor Supply Data	64
6	Directions for Future Research	70
6.1	Consistency across contexts	71
6.2	Combining Experimental and Field Data	79
6.3	Direct Measurement of Beliefs	80
6.4	Assumptions about “Mental Accounting”	81

1 Introduction

Risk preferences are integral to modern economics. They are the primary focus of the literature on decision making under uncertainty. They play a central role in insurance and financial economics. The topics of risk sharing and insurance are prominent in development, health, labor, and public economics, particularly in the study of incentives and social insurance programs. And risk preferences are a major driver in models of consumption, investment, and asset pricing in macroeconomics.

The vast majority of the literature uses expected utility (EU) theory to model risk preferences, and much of the literature is theoretical in nature, deriving qualitative predictions in different environments. At the same time, there is a large empirical literature that estimates risk preferences, both EU and non-EU models, using data from a variety of settings. Early studies in the empirical literature on risk preferences focus on the EU model and rely on data from laboratory experiments (e.g., Preston and Baratta 1948; Yaari 1965); for early reviews, see Camerer (1995) and Starmer (2000). Laboratory experiments generated many insights about risk preferences, and most notably demonstrated both substantial heterogeneity in risk preferences and substantial deviations from EU theory. However, the limitations commonly associated with the laboratory setting motivated economists to look for suitable data from field settings. By field setting, we mean an environment in which people’s real-world economic behavior is observable.¹

As a result, there is a relatively small (but growing) literature that takes on the difficult task of estimating risk preferences using field data. Our goal in this paper is to review and assess this literature. We mostly limit our attention to studies in which risk preferences are the focal object and estimating their structure is the core enterprise. In particular, we generally exclude papers that estimate a structural model of risk preferences as an incidental matter and treat the risk preference parameters as nuisance parameters and not as the parameters of main interest. Although there are many excellent papers in this category which make important contributions to numerous fields of economics, they are beyond the scope of this review.²

We begin in Section 2 with a motivating example designed to address the question of why economists should care about the structure of risk preferences. More and more, economists are engaging in analyses that investigate the quantitative impact of a change in the underly-

¹By comparison, laboratory experiments have two major advantages. First, the researcher has precise control over the choice sets that subjects face, and thus can design those choice sets to be maximally informative. Second, the researcher can elicit multiple choices from each subject, and thus in principle could estimate risk preferences at the individual level.

²As we explain below, however, we discuss a handful of papers that, although they fall into this category, make valuable contributions to the methodology of estimating risk preference using field data.

ing economic environment (e.g., a legal reform). In such analyses, risk preferences are often a required input, even if only as a nuisance parameter in a broader model. Our example highlights two reasons why the specification of risk preferences matters. First, many quantitative analyses attempt to make out-of-sample predictions for behavior based on the broader model. We demonstrate in our example how different assumptions about risk preferences can lead to very different out-of-sample predictions for behavior. Second, many quantitative analyses attempt to reach welfare conclusions. Again, we demonstrate in our example how different assumptions about risk preferences can lead to very different welfare conclusions.

In Section 3, we provide a detailed review of several models of risk preferences that have been estimated or otherwise studied using field data. We begin with EU theory, and proceed to describe several non-EU models that were originally motivated by experimental evidence but which have subsequently been studied using field data, including rank-dependent expected utility (RDEU) and cumulative prospect theory (CPT). We then provide a discussion of identification, and in particular describe what types of data are needed to estimate and distinguish the various models.

In Section 4, we discuss research that estimates risk preferences, and sometimes heterogeneity in risk preferences, using individual-level data. We begin with an overview of the general approach used throughout this literature. Next, we describe in detail research that estimates risk preferences using data on property insurance choices. We then briefly discuss studies that use data from television game shows. However, we limit the depth of our coverage, because, although these studies focus on estimating risk preferences, television game shows fall outside our definition of a field setting.³ Lastly, we review a handful of recent papers that analyze data on health insurance choices. We again limit the depth of our coverage, but for a different reason. Health insurance is an important field context, but the papers that use health insurance data do not focus on estimating risk preferences. We believe this is because estimating risk preferences using health insurance data is especially challenging. Nevertheless, we highlight a few recent papers that address some of these challenges and whose contributions could facilitate future work that focuses on estimating risk preferences.

In Section 5, we turn to research that estimates risk preferences, and sometimes heterogeneity in risk preferences, using market-level, or aggregate data. Once again, we begin with an overview of the general approach of the literature, highlighting how the use of aggregate data naturally requires a stronger set of assumptions in order to identify risk preferences. Next, we describe in detail research that estimates risk preferences using data on betting

³This is because factors other than their risk preferences and beliefs influence agents' decisions in a game show. For example, the entertainment value, the length of the game, the TV studio environment, the interaction with the game host, and the influence of the audience.

markets, specifically data on betting in pari-mutuel horse races. We then discuss a select assortment of papers that use macroeconomic data to estimate risk preferences, including data on consumption and investment (asset returns) and on labor supply, or that combine aggregate and individual consumption data to study risk sharing behavior.

Finally, in Section 6 we discuss a number of directions for future research. An under-researched issue is the extent to which risk preferences are stable across contexts. We review the few studies that use field data to investigate this issue, and we highlight the questions left open by these studies. Relatedly, we also discuss the possibility of combining data from laboratory and field contexts in order to paint a more complete picture of risk preferences—and also to gain insight on the question of whether experimental results can be directly applied to make field predictions. Next, we describe the recent literature on using surveys to directly measure risk perceptions, and we discuss the extent to which survey data might be usefully combined with field data to identify and estimate risk preferences under weaker assumptions. Finally, we discuss the importance of various "mental accounting," by which we mean assumptions about how agents translate a complex field context into a set of concrete lotteries to be evaluated. We encourage future research to pay more careful attention to such assumptions.

2 Motivating Example

In this section, we present a stylized example designed in part to motivate why economists should care about the structure of risk preferences and in part to illustrate some of the challenges that economists face in identifying and estimating risk preferences using field data. The setting of our example is a hypothetical insurance market. Initially, we make a number of strong assumptions—about the setting and the data—that in turn make identification and estimation straightforward. We then revisit those strong assumptions to highlight some of the identification and estimation challenges that economists face in more realistic field settings.

Imagine that there is a continuum of households of measure one who each face the possibility of a loss L that occurs with probability μ . Both L and μ are the same across households, and their values are known. To fix ideas, let $L = 10,000$ and $\mu = 0.05$. There is insurance available to the households—full insurance at a price p . Moreover, there is sufficient exogenous price variation (e.g., over time or across various identical subsets of the households, see Einav, Finkelstein, and Cullen (2010) for an example) to non-parametrically identify the market demand function for full insurance, $Q^F(p)$, which returns the fraction of households willing to purchase full insurance at price p . Panel (a) of Figure 1 depicts one

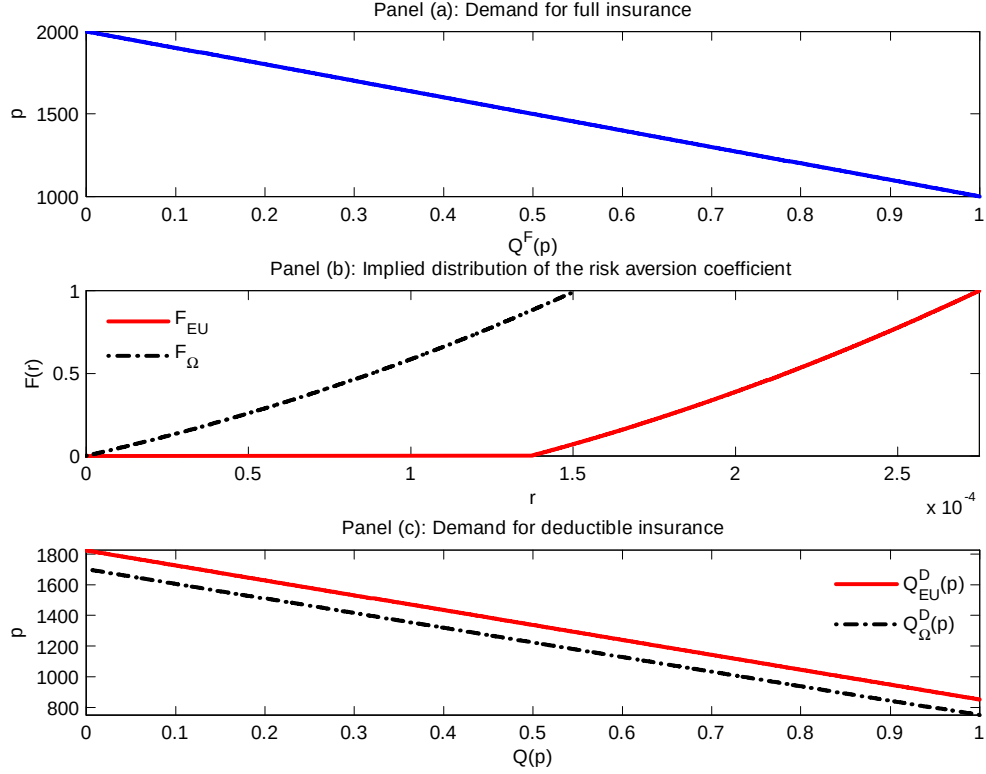


Figure 1: Demand for insurance and underlying risk preferences

such demand function, namely $Q^F(p) = 2 - 0.001p$.⁴ It is a typical demand function—as the insurance gets cheaper, the fraction of households willing to purchase it gets larger. It also reflects aversion to risk—households demand insurance at actuarially unfair prices.

2.1 Out-of-Sample Predictions

We first illustrate how the underlying structure of risk preferences can matter for making out-of-sample predictions. Consider a regulatory proposal to require all insurance policies to carry a deductible $D < L$. In order to assess this proposal, we need to know how the demand for insurance would respond to the introduction of the deductible D . The demand function for full insurance $Q^F(p)$ —which we observe—provides, by itself, limited information about the market demand function for deductible insurance, $Q^D(p)$. However, if we know the underlying model that generates $Q^F(p)$, we can use that model to construct $Q^D(p)$.

⁴More precisely, $Q^F(p) = \begin{cases} 1 & \text{if } p < 1000 \\ 2 - 0.001p & \text{if } 1000 \leq p \leq 2000 \\ 0 & \text{if } p > 2000 \end{cases}$.

Assume for the moment that the underlying model is EU. In addition, assume that (i) the utility function exhibits constant absolute risk aversion (CARA), specifically $u(y) = -\exp(-ry)/r$, where r is the coefficient of absolute risk aversion, and (ii) the slope of the demand function $Q^D(p)$ arises from heterogeneity in r . Given these assumptions, we can recover the population distribution of r , denoted F_{EU} , from the market demand function $Q^F(p)$. Observe that a household's willingness to pay for full insurance is the z such that⁵

$$\exp(rz) = \mu \exp(rL) + (1 - \mu). \quad (1)$$

Equation (1) defines $r^F(z)$ —the coefficient of absolute risk aversion of a household with willingness to pay z for full insurance. A household purchases full insurance when $r^F(z) > r^F(p)$, and hence the demand for full insurance satisfies $Q^F(p) = 1 - F_{EU}(r^F(p))$. It follows that, given $Q^F(p)$, we can recover F_{EU} . Panel (b) of Figure 1 displays the F_{EU} that corresponds to the $Q^F(p)$ depicted in panel (a).

Given F_{EU} , it is straightforward to construct the demand for deductible insurance $Q^D(p)$. A household's willingness to pay for deductible insurance is the z such that

$$\mu \exp(r(z + D)) + (1 - \mu) \exp(rz) = \mu \exp(rL) + (1 - \mu). \quad (2)$$

Equation (2) defines $r^D(z)$ —the coefficient of absolute risk aversion of a household with willingness to pay z for deductible insurance. A household purchases deductible insurance when $r^D(z) > r^D(p)$, and hence the demand for deductible insurance is $Q_{EU}^D(p) = 1 - F_{EU}(r^D(p))$. Panel (a) of Figure 1 depicts the $Q_{EU}^D(p)$ that corresponds to the $Q^F(p)$ depicted in panel (a), assuming $D = 2,500$. Because deductible insurance provides less coverage than full insurance, naturally $Q_{EU}^D(p) < Q^F(p)$.

Making a different assumption about the underlying model, however, can lead to different predictions about the level of demand for deductible insurance. Suppose, for example, that the underlying model is the probability distortion (PD) model featured in Barseghyan, Molinari, O'Donoghue, and Teitelbaum (2013b). The PD model posits that households, instead of weighting outcomes by their objective probabilities, weight outcomes using distorted probabilities.⁶ Under the PD model (and maintaining the additional assumptions specified above), a household's willingness to pay for full insurance is the z such that

$$\exp(rz) = \Omega(\mu) \exp(rL) + (1 - \Omega(\mu)), \quad (3)$$

⁵For an arbitrary utility function, z is defined implicitly by $u(w - z) = \mu u(w - L) + (1 - \mu)u(w)$, where w is the household's status quo wealth.

⁶In the setting of our example, such probability distortion can emerge from several prominent alternative models of choice under risk. See Section 3 for further discussion.

where $\Omega(\mu) \neq \mu$ is the weight on the loss outcome. Suppose that $\Omega(\mu) = \bar{\Omega} > \mu$ is the same across households, and that $\bar{\Omega}$ is known. Given $\bar{\Omega}$, we can proceed as before to use the known demand for full insurance $Q^F(p)$ to construct the counterfactual demand for deductible insurance $Q_\Omega^D(p)$.

Let F_Ω denote the distribution of r given the PD model with loss weight $\bar{\Omega}$. We can recover F_Ω from the demand for full insurance, $Q^F(p) = 1 - F_\Omega(r_\Omega^F(p))$, where $r_\Omega^F(z)$ is defined by equation (3) with $\Omega(\mu) = \bar{\Omega}$. Panel (b) of Figure 1 displays the F_Ω that corresponds to the $Q^F(p)$ depicted in panel (a), assuming $\bar{\Omega} = 0.10$. Given F_Ω , we can construct the demand for deductible insurance, $Q_\Omega^D(p) = 1 - F_\Omega(r_\Omega^D(p))$, where $r_\Omega^D(z)$ is defined by

$$\bar{\Omega} \exp(r(z + D)) + (1 - \bar{\Omega}) \exp(rz) = \bar{\Omega} \exp(rL) + (1 - \bar{\Omega}),$$

the equation that implicitly defines a household's willingness to pay z for deductible insurance. Panel (c) of Figure 1 depicts the $Q_\Omega^D(p)$ that corresponds to the $Q^F(p)$ depicted in panel (a), assuming $D = 2,500$ and $\bar{\Omega} = 0.10$. Observe that $Q_\Omega^D(p) < Q_{EU}^D(p)$.

In short, we see that the two models generate different predictions for the level of demand for deductible insurance. In particular, the EU model predicts a higher level of demand than the PD model. The intuition for this difference follows the nature of concave utility. Under both models, a concave utility function implies that the concern for reducing risk is stronger the more risk one bears.⁷ Moreover, this effect becomes stronger as the concavity of the utility function increases (i.e., as r gets larger). For a given (observed) demand for actuarially unfair full insurance, the concavity of the utility function is greater under the EU model than under the PD model with $\Omega(\mu) > \mu$,⁸ thus the implied demand for deductible insurance is greater under EU.

2.2 Welfare Analysis

The underlying structure of risk preferences can matter for welfare analysis, even when the analysis does not entail making out-of-sample predictions. In an important paper, Einav, Finkelstein, and Cullen (2010) propose an approach to empirical welfare analysis in insurance markets that does not require specifying or estimating the underlying structure of households' risk preferences. Instead, their approach—which takes the characteristics of insurance

⁷Take our example: although the deductible insurance provides 75 percent of the coverage of full insurance, under both modes a household's willingness to pay for the deductible insurance is greater than 75 percent of the willingness to pay for full insurance (see 1).

⁸Intuitively, this is because under the EU model a household's aversion to risk (which generates its insurance demand) is driven solely by the concavity of its utility function, whereas under the PD model a household's aversion to risk is driven also by the overweighting of its distortion function.

contracts as given but allows their prices to be determined endogenously—imposes certain "reduced-form" restrictions on households, firms, and the marketplace. It then relies on exogenous price variation to estimate the market demand and cost functions and uses the estimates to conduct welfare analysis of regulatory interventions that change the prices of existing contracts. (The authors acknowledge that their approach cannot be used to analyze interventions that introduce contracts not observed in the data.)

This type of welfare analysis is valid provided that the demand function is a sufficient statistic for consumer welfare. Yet, this is not the case for all possible underlying risk preference structures. Perhaps the cleanest example of when the demand function is not a sufficient statistic for consumer welfare, is a PD model in which $\Omega(\mu)$ represents risk misperceptions (i.e., incorrect subjective beliefs) rather than rank-dependent probability weighting or some other deviation from EU.⁹ In this case, the demand for insurance is generated by the EU model with the misperceived probability $\Omega(\mu)$; however, consumer welfare arguably should be evaluated using the EU model with the true probability μ .

2.3 Identification and Estimation Challenges

We conclude this section by highlighting several challenges that economists confront when estimating risk preferences in the field. We list them below and return to them from time to time in the rest of the paper.

Identification Our stylized example assumes that the weight $\bar{\Omega}$ is known. In the field this would not be the case; instead, it would be necessary to estimate $\bar{\Omega}$. The challenge would be to collect data that would allow for separate identification of $\bar{\Omega}$ and the distribution of r . For instance, the "data" in our stylized example is not sufficient to separately identify $\bar{\Omega}$ and the distribution of r . After all, Figure 1 illustrates two combinations of a value of $\bar{\Omega}$ and a distribution of r that are equally consistent with the observed demand for full insurance depicted in Figure 1 (namely, (i) $\bar{\Omega} = 0.05$ and F_{EU} and (ii) $\bar{\Omega} = 0.10$ and F_{Ω}). Indeed, there are an infinite number of combinations that are consistent with the data in our example. Thus, a major challenge that economists face when estimating risk preferences in the field is collecting data that allow for joint identification under plausible restrictions on the model of the various underlying sources of aversion to risk. We discuss this problem in Section 3.6.

Risk heterogeneity Our stylized example assumes that every household faces the same magnitude and probability of loss. In field contexts, however, this would not be the case.

⁹See Section 3.

Indeed, the large literature on adverse selection is entirely premised on the idea that different households (even with identical observable characteristics) face different risks. Accounting for such risk heterogeneity creates a major challenge to estimating risk preferences. In some instances, access to observable household characteristics permits one to estimate how households' risks depend on those observables. In addition, more sophisticated analyses also permit (and estimate the degree of) unobserved heterogeneity in the risks that different households face — we return to this idea in Section 4.1.2.

Exogenous price variation Our stylized example assumes that the data contain sufficient exogenous price variation to non-parametrically identify the market demand function for full insurance. However, a recurring theme throughout the literature is the difficulty of collecting data with sufficient price variation that is plausibly exogenous to households' risk preferences. Indeed, as we evidence in this review, the literature on estimating risk preferences using field data got off to a slow start due in part to a paucity of "good" data. This literature has been picking up steam in recent years due in part to economists collecting better data.

Preference heterogeneity This is an important issue that is intertwined with identification. In reality, no field data are "ideal," implying that most models are identified up to one or more strong assumptions about the structure of unobserved heterogeneity in risk preferences. These assumptions may constrain the dimensionality of unobserved heterogeneity (e.g., allowing for unobserved heterogeneity either in the curvature of the utility function or in the shape of the probability distortion function, but not in both), or may impose parametric structure on its distribution. We discuss these matters in Section 3.6.3.

Risk preferences versus subjective beliefs Our stylized example also illustrates the importance of being able to distinguish probability distortions that arise because of subjective beliefs (i.e., risk misperceptions) from those that arise because of the underlying structure of risk preferences (e.g., rank-dependent probability weighting). We discuss this problem in Section 3.6.2. More broadly, the nature of subjective beliefs and their relation to the objective probabilities is the subject of an expanding literature that focuses on elicitation of probabilistic expectations. We discuss this literature in Section 6.3.

3 Models of Risk Preferences

To estimate models of risk preferences, one must understand those models well, both the general structure of a particular model and the possible parametric restrictions that one

might use for that model (along with the limitations of each parametric restriction). Hence, in this section we describe in detail several models of risk preferences. We begin by reviewing the standard EU model, and then we move on to introduce several alternative models. Our goal is not to provide an exhaustive list but rather to focus on models of risk preferences that have been prominent in the literature that uses field data to estimate risk preferences.¹⁰

We start by introducing notation that we use throughout this section.

Definition 1. Let $X \equiv (x_1, \mu_1; x_2, \mu_2; \dots; x_N, \mu_N)$ denote a lottery which yields outcome x_n with probability μ_n , where $\sum_{n=1}^N \mu_n = 1$.

Models of risk preferences describe how a household chooses among lotteries of this form, where we often use $\mathcal{X}$ to denote a choice set.¹¹ Throughout, we express lottery outcomes in terms of increments added to (or subtracted from) the household's prior wealth w . In other words, if outcome x_n is realized, then the household will have final wealth $w + x_n$. The probabilities should be taken to be a household's subjective beliefs. In particular, the models below describe how a household's subjective beliefs impact its choices. The models are silent on the source of those subjective beliefs — we return to this issue in Section 4.1.1.

We illustrate the different models by describing each theory's predictions in two examples:

Example 1. What is a household's willingness to pay for full insurance against the possibility of losing L with probability μ ? In other words, what is the z that makes the household indifferent between the lottery $(-z, 1)$ and the lottery $(-L, \mu; 0, 1 - \mu)$?

Example 2. What is a household's willingness to pay for an asset that pays out x_1 with probability μ_1 , x_2 with probability μ_2 , and x_3 with probability μ_3 ? In other words, what is the z that makes the household indifferent between the lottery $(0, 1)$ and the lottery $(-z + x_1, \mu_1; -z + x_2, \mu_2; -z + x_3, \mu_3)$?

3.1 Expected Utility (EU)

Under expected utility (EU), given a choice set $\mathcal{X}$, a household will choose the option $X \in \mathcal{X}$ that maximizes

$$EU(X) \equiv \sum_{n=1}^N \mu_n u(w + x_n),$$

¹⁰Experimental analyses of risk preferences often consider a broader set of models. Indeed, the experimental approach permits one to identify more nuanced features of risk preferences because one is able to design the data in the needed way. One might worry, however, that the more nuanced the behavior one seeks to identify in an experiment, the less valid it is to extrapolate from the artificial experimental environment.

¹¹Note that different lotteries could of course have a different number of outcomes — that is, N could differ across lotteries.

Panel 1: Utility functions	
CARA	$u(y) = \begin{cases} -\frac{1}{r} \exp(-ry) & \text{for any } r \neq 0, \\ y & \text{for } r = 0. \end{cases}$
CRRA	$u(y) = \begin{cases} \frac{1}{1-\rho} y^{1-\rho} & \text{for any } \rho \neq 1, \\ \ln y & \text{for } \rho = 1. \end{cases}$
HARA	$u(y) = \begin{cases} \frac{\gamma}{1-\gamma} \left(\eta + \frac{y}{\gamma} \right)^{1-\gamma} & \text{for any } \gamma \neq 1, \\ \gamma \ln \left(\eta + \frac{y}{\gamma} \right) & \text{for } \gamma = 1. \end{cases}$
NTD	$\tilde{u}(y) = y - \frac{r}{2} y^2.$
Panel 2: Probability weighting functions	
Karmarkar (1978)	$\pi(\mu) = \frac{\mu^\gamma}{\mu^\gamma + (1-\mu)^\gamma}.$
Tversky and Kahneman (1992)	$\pi(\mu) = \frac{\mu^\gamma}{[\mu^\gamma + (1-\mu)^\gamma]^{1/\gamma}}.$
Lattimore, Baker, and Witte (1992)	$\pi(\mu) = \frac{\delta \mu^\gamma}{\delta \mu^\gamma + (1-\mu)^\gamma}.$
Prelec (1998)	$\pi(\mu) = \exp(-(-\ln p)^\alpha).$
Panel 3: Value function	
Tversky and Kahneman (1992)	$v(y) = \begin{cases} y^\alpha & \text{for } y \geq 0, \alpha \in (0, 1), \\ -\lambda (-y)^\beta & \text{for } y < 0, \beta \in (0, 1), \lambda > 1. \end{cases}$

Table 1: Functional forms used in this article

where u is a utility function that maps final wealth onto the real line. Hence, in our two examples:

Example 1 (Under EU). *A household's willingness to pay for full insurance against the possibility of losing L with probability μ is the z such that*

$$u(w - z) = \mu u(w - L) + (1 - \mu)u(w).$$

Example 2 (Under EU). *A household's willingness to pay for an asset that pays out x_1 with probability μ_1 , x_2 with probability μ_2 , and x_3 with probability μ_3 is the z such that*

$$u(w) = \mu_1 u(w - z + x_1) + \mu_2 u(w - z + x_2) + \mu_3 u(w - z + x_3).$$

Under EU, a household's attitude towards risk is fully captured by its utility function u (and its prior wealth). In broad terms, a household will be risk averse if u is concave, risk loving if u is convex, and risk neutral if u is linear. More narrowly, we can derive a local measure of absolute or relative risk aversion (or risk lovingness) that characterizes how a household will react locally to choices between lotteries — e.g., in Example 1, a household

with higher risk aversion will have a higher willingness to pay for full insurance.

Hence, when one estimates an EU model, the main object to estimate is the utility function u . As we shall see, occasionally researchers have taken a non-parametric approach to estimating u , but most often they assume a specific parametric functional form for u . Perhaps the most common functional forms are the Constant Absolute Risk Aversion (CARA), the Constant Relative Risk Aversion (CRRA), and the Hyperbolic Absolute Risk Aversion (HARA) families, reported in Panel 1 of Table 1.

When one uses the CARA family, one estimates the parameter r , which is the coefficient of absolute risk aversion (higher r means more risk averse). The CARA family has the advantage that it implies a household's prior wealth w , which frequently is unobserved, is irrelevant to the household's decisions. However, the CARA family has a drawback. Economists typically believe that households exhibit decreasing absolute risk aversion — that is, as a household becomes wealthier, it becomes less averse to risk.

When one uses the CRRA family, one estimates the parameter ρ , which is the coefficient of relative risk aversion (higher ρ means more risk averse). The CRRA family has the advantage of implying decreasing absolute risk aversion (among those who are risk averse). However, the CRRA family has the major drawback that it requires prior wealth w as an input. Hence, when researchers use the CRRA family and do not observe prior wealth, they typically either posit some reasonable value for prior wealth (and check robustness for other values), or proxy for wealth using some aspect of the data (e.g., home value).

Finally, when one uses the HARA family, one estimates the parameters η and γ , which together determine the degree of absolute risk aversion $r(y) = (\eta + y/\gamma)^{-1}$. The HARA family has the property that it nests the CARA and CRRA families as special cases, respectively $\gamma \rightarrow +\infty$ yielding CARA, and $\eta = 0$ yielding CRRA.¹²

A third technique is to use an approximation approach (see Cohen and Einav (2007); Barseghyan, Prince, and Teitelbaum (2011); Barseghyan, Molinari, O'Donoghue, and Teitelbaum (2013b)). Specifically, if one takes a second-order Taylor approximation of the utility function around prior wealth w and then normalizes by marginal utility evaluated at prior wealth w , one gets

$$\tilde{u}(\Delta) \equiv \frac{u(w + \Delta)}{u'(w)} - \frac{u(w)}{u'(w)} \cong \Delta - \frac{r}{2}\Delta^2, \quad (4)$$

where $r \equiv -u''(w)/u'(w)$ is local absolute risk aversion. This approximation is accurate when the third- and higher-order derivatives of the utility function u are negligible, at least relative to the increments to wealth that are relevant in a particular application. As such,

¹²Some researchers (e.g., Cicchetti and Dubin (1994) and Jullien and Salanié (2000)) assume a simpler HARA specification $u(y) = (\eta + y)^\gamma$. This simplification necessitates limiting γ to lie in the interval $(0, 1]$.

we label this approach the Negligible Third Derivative (NTD) approach.

The NTD family is convenient to work with because it does not require prior wealth as an input. However, one must be careful to assess whether the approximation method is appropriate for the particular application under consideration. This will depend on the magnitude of the increments to wealth relative to the estimated degree of risk aversion.¹³

3.2 Rank-Dependent Expected Utility (RDEU)

Rank-dependent expected utility (RDEU) emerged from a tradition in psychology of relaxing the feature of EU that outcomes are weighted by their probabilities. In other words, we replace the expected utility equation with

$$V(X) \equiv \sum_{n=1}^N \omega_n u(w + x_n)$$

where ω_n is a decision weight associated with outcome x_n and may not be equal to a person's belief μ_n . The original idea was proposed by Edwards (1955, 1962) and popularized in Kahneman and Tversky's 1979 prospect theory, which assumes $\omega_n = \pi(\mu_n)$. In other words, there is an increasing function π — often labelled a probability weighting function — that transforms each probability into a decision weight (still normalizing that $\pi(0) = 0$ and $\pi(1) = 1$). With this formulation, however, for any $\pi(\mu) \neq \mu$ it is possible to construct examples in which the theory predicts violations of stochastic dominance — i.e., that people would choose a lottery over another that stochastically dominates it. The source of such predictions is that, unlike under EU, when evaluating lotteries, the weights need not sum to one.¹⁴

Quiggin (1982) proposed a rank-dependent model to solve this problem. Under the rank-dependent approach, when evaluating a lottery $X \equiv (x_1, \mu_1; x_2, \mu_2; \dots; x_N, \mu_N)$, a household first ranks the outcomes from best to worst. Specifically, if the outcomes are ordered such that $x_1 < x_2 < \dots < x_N$, then the weight on outcome n is

$$\omega_n = \begin{cases} \pi(\mu_1) & \text{for } n = 1, \\ \pi\left(\sum_{j=1}^n \mu_j\right) - \pi\left(\sum_{j=1}^{n-1} \mu_j\right) & \text{for } n \in \{2, \dots, N-1\}, \\ 1 - \pi\left(\sum_{j=1}^{n-1} \mu_j\right) & \text{for } n = N, \end{cases}$$

¹³One obvious concern is that utility must be increasing, which for risk-averse individuals (with $r > 0$) holds only for $\Delta < 1/r$.

¹⁴For instance, if $\pi(1/3) > 1/3$, then there exists $\bar{y} > 0$ such that the model predicts a person would choose the lottery $(x, 1/3; x - y, 1/3; x - 2y, 1/3)$ over the lottery $(x, 1)$ for all $y \in (0, \bar{y})$.

where again π is a probability weighting function. With this approach, when evaluating a lottery, the weights sum to one by construction, and there are no violations of stochastic dominance.

For our two examples, RDEU generates the following equations:

Example 1 (Under RDEU). *A household's willingness to pay for full insurance against the possibility of losing L with probability μ is the z such that:*

$$u(w - z) = \pi(\mu)u(w - L) + (1 - \pi(\mu))u(w).$$

Example 2 (Under RDEU). *A household's the willingness to pay for an asset that pays out x_1 with probability μ_1 , x_2 with probability μ_2 , and x_3 with probability μ_3 is the z such that*

$$\begin{aligned} u(w) = & \pi(\mu_1)u(w - z + x_1) + (\pi(\mu_2 + \mu_1) - \pi(\mu_1))u(w - z + x_2) \\ & + (1 - \pi(\mu_2 + \mu_1))u(w - z + x_3) \end{aligned}$$

The implications of RDEU of course depend on the specific probability weighting function that is used. The literature — in large part based on experimental results — has emphasized an inverse-S-shaped probability weighting function: For small μ , $\pi(\mu)$ is concave and has $\pi(\mu) > \mu$, while for large μ , $\pi(\mu)$ is convex and has $\pi(\mu) < \mu$. Figure 2 depicts a stylized version of such a function, where this particular version has an inflection point at $\mu = 1/2$ and is symmetric around $\mu = 1/2$.¹⁵

To understand the implications of RDEU, consider the implications of the probability weighting function in Figure 2 for Examples 1 and 2. For binary lotteries, as in Example 1, there will be overweighting of low probability events (events with $\mu < 1/2$) and underweighting of high probability events (events with $\mu > 1/2$). Hence, in Example 1, if the probability of a loss $\mu < 1/2$, then the weight $\pi(\mu)$ on the loss will be greater than the probability, and thus probability weighting generates a source of risk aversion. In contrast, if the probability of a loss $\mu > 1/2$, then the weight $\pi(\mu)$ on the loss will be less than the probability, and thus probability weighting generates a source of risk seeking.

For lotteries with more than two outcomes, such as Example 2, an inverse-S-shaped probability weighting function instead generates overweighting of extreme outcomes and underweighting of intermediate outcomes — and importantly two outcomes that are equally likely need not have the same decision weights. For instance, consider Example 2 when $\mu_1 = \mu_2 = \mu_3 = 1/3$. Given the probability weighting function in Figure 2, the extreme outcomes of x_1 and x_3 will both be overweighted (i.e., $\pi(1/3) > 1/3$ and $1 - \pi(2/3) > 1/3$),

¹⁵Figure 2 is consistent with the Karmarkar (1978) functional form reported in Table 1, with $\gamma = 1/2$.

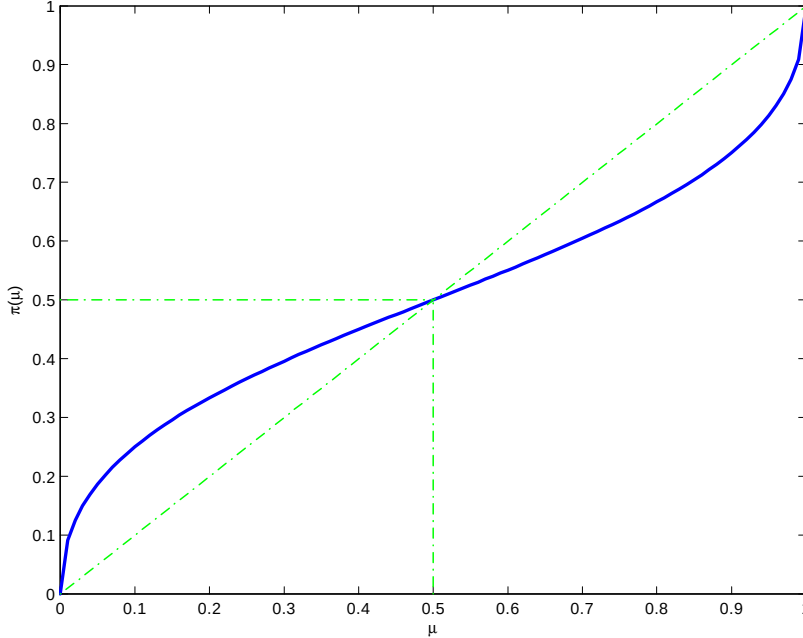


Figure 2: Symmetric probability weighting function.

while the intermediate outcome x_2 is underweighted (i.e., $\pi(2/3) - \pi(1/3) < 1/3$).

Beyond the general inverse-S shape, a number of parametrized functional forms have been proposed in the literature on probability weighting. Some prominent functional forms are reported in Table 1-Panel 2.¹⁶

Whereas Figure 2 presents the Karmarkar form for $\gamma = 1/2$, Figure 3 presents estimated versions of the other three forms.¹⁷ Note two features in Figure 3. First, unlike in Figure 2, the functions are not symmetric around $\mu = 1/2$, but rather they typically cross the 45-degree line at $\mu < 1/2$. We return to the importance of this feature in a moment. Second, the functions exhibit excess steepness near $\mu = 0$ and $\mu = 1$ — in the sense of $\pi'(\mu) \gg 1$. In fact, in their original discussion of probability weighting, Kahneman and Tversky (1979) instead suggested that probability weighting is discontinuous at the endpoints, reflecting a notion that as the probability of an event gets small enough, people ignore that possible event. The

¹⁶To be precise, Lattimore, Baker, and Witte (1992) propose the functional form

$$\pi(\mu_i|\mu_{-i}) = \frac{\delta\mu_i^\gamma}{\delta\mu_i^\gamma + \sum_{j \neq i} \mu_j^\gamma}, \quad i = 1, \dots, N,$$

with μ_{-i} denoting the entries of the probability vector μ other than μ_i . For $N = 2$, the above expression coincides with what appears in Table 1, and this functional form (used also in the case of $N > 2$) is commonly referred to in the literature as the Lattimore, Baker, and Witte's form.

¹⁷Figure 3 closely parallels Figure 1 from Prelec (1998). As in that figure, we use $\alpha = .65$ for the Prelec function, and we use $\gamma = .61$ for the Tversky and Kahneman (1992) function. For the Lattimore, Baker, and Witte (1992), we use $\delta = .77$ and $\gamma = .44$, which are estimates from Gonzalez and Wu (1999).

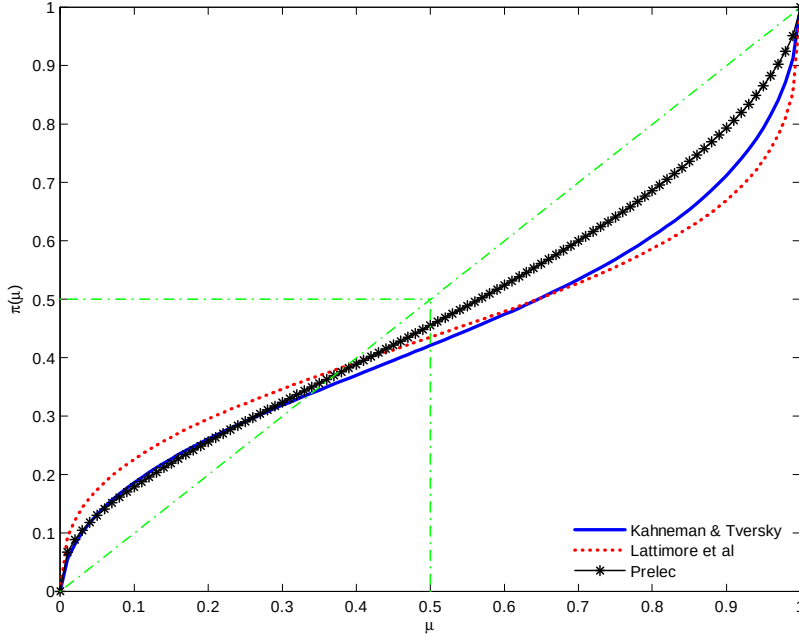


Figure 3: Assymmetric probability weighting functions.

subsequent literature seems to have introduced the excess steepness near $\mu = 0$ and $\mu = 1$ to eliminate this discontinuity. However, it is unclear how much evidence there is for this excess steepness. As we shall see, in field applications, it is important to assess whether and how low-probability events are incorporated into household's decision calculations.

In the RDEU model outlined above, the probability-weighting function is first applied to the worst outcome, and then successively applied to better and better outcomes — which we refer to as RDEU from the bottom. Some analyses do the reverse, so the probability weighting function is first applied to the best outcome, and then successively applied to worse and worse outcomes — which we refer to as RDEU from the top. Formally, if the outcomes are again ordered such that $x_1 < x_2 < \dots < x_N$, then, under RDEU from the top, the weight on outcome n is

$$\omega_n = \begin{cases} \pi(\mu_N) & \text{for } n = N, \\ \pi\left(\sum_{j=n}^N \mu_j\right) - \pi\left(\sum_{j=n+1}^N \mu_j\right) & \text{for } n \in \{2, \dots, N-1\}, \\ 1 - \pi\left(\sum_{j=n+1}^N \mu_j\right) & \text{for } n = 1. \end{cases}$$

In principle, whether one applies RDEU from the bottom versus from the top does not necessarily add a dimension along which the model can be misspecified. In particular, if we use π^T in RDEU from the top and π^B in RDEU from the bottom, the two forms yield identical predictions as long as these functions are symmetric, in the sense that $\pi^B(\mu) = 1 - \pi^T(1 - \mu)$

for all μ . Similarly, if one estimates an RDEU model with a non-parametric approach to π , it is not important which version one uses. However, if one uses a parametric functional form that does not satisfy the symmetry property, the direction along which RDEU is applied yields an additional dimension along which the model can be misspecified. For instance, suppose we use the Prelec (1998) functional form in Example 1 with $\mu = .4$. Under RDEU from the bottom, because $\pi(.4) < 0.4$ we would underweight the loss event, whereas under RDEU from the top, because $1 - \pi(.6) > 0.4$ we would overweight the loss event.

In field applications, both variants have been used. Typically, the version used depends on which outcome is more “focal” in a particular application. For insurance applications, where the loss event arguably is the focal event, researchers most often use RDEU from the bottom. For gambling applications, where the win event arguably is the focal event, researchers most often use RDEU from the top.

3.3 Cumulative Prospect Theory (CPT)

Kahneman and Tversky (1979) “prospect theory” has two key features: probability weighting and loss aversion. As discussed above, probability weighting derived from an older tradition in psychology, and is fully incorporated into RDEU. Loss aversion represents a second departure from the standard expected utility model: Instead of there being a utility function u defined over final wealth, there is a value function v defined over gains and losses relative to some reference point.

Tversky and Kahneman (1992) propose an improved version of their theory, labelled “cumulative prospect theory” (CPT). CPT requires as an input a reference outcome s , and each outcome is coded as a gain or loss relative to this reference outcome.¹⁸ Consider a lottery $X \equiv (x_1, \mu_1; \dots; x_N, \mu_N)$ and a reference point s , and suppose $x_1 < \dots < x_{\bar{n}-1} \leq s < x_{\bar{n}} < \dots < x_N$. Under CPT, this lottery is evaluated as

$$V(X; r) \equiv \sum_{n=1}^N \omega_n v(x_n - s)$$

¹⁸The discussion in the text will focus on a reference outcome expressed in increments to wealth, and thus the comparison is x to s . One could equivalently use a reference outcome expressed in final wealth, in which case the comparison would be $(w + x)$ to s — i.e., the value function $v(x - s)$ would be replaced with $v((w + x) - s)$.

where the weight on outcome x_n is

$$\omega_n = \begin{cases} \pi^-(\mu_1) & \text{for } n = 1 \\ \pi^-\left(\sum_{j=1}^n \mu_j\right) - \pi^-\left(\sum_{j=1}^{n-1} \mu_j\right) & \text{for } n \in \{2, \dots, \bar{n} - 1\} \\ \pi^+\left(\sum_{j=n}^N \mu_j\right) - \pi^+\left(\sum_{j=n+1}^N \mu_j\right) & \text{for } n \in \{\bar{n}, \dots, N - 1\}. \\ \pi^+(\mu_N) & \text{for } n = N. \end{cases}$$

In this formulation, π^- and π^+ are probability-weighting functions applied to the loss and gain events, respectively. Thus, the theory permits differential weighting for gains and losses.¹⁹ Of course, this differential weighting creates the potential for violations of stochastic dominance.

The value function v is assumed to have three key properties: (i) it has $v(0) = 0$ and assigns positive value to gains and negative value to losses, (ii) it is concave over gains and convex over losses (often labelled “diminishing sensitivity”), and (iii) it is steeper in the loss domain than in the gain domain (often labelled “loss aversion”).

For our two examples, CPT with a reference point $s = 0$ generates the following equations:

Example 1 (Under CPT). *A household’s willingness to pay for full insurance against the possibility of losing L with probability μ is the z such that:*

$$v(-z) = \pi^-(\mu)v(-L).$$

Example 2 (Under CPT). *A household’s willingness to pay for an asset that pays out x_1 with probability μ_1 , x_2 with probability μ_2 , and x_3 with probability μ_3 , where $x_1 < s < x_2 < x_3$ is the z such that²⁰*

$$\begin{aligned} v(0) &= \pi^-(\mu_1)v(-z + x_1) + (\pi^+(\mu_3 + \mu_2) - \pi^+(\mu_3))u(w - z + x_2) \\ &\quad + \pi^+(\mu_3)u(w - z + x_3). \end{aligned}$$

In order to estimate a CPT model, one often needs functional form assumptions (although occasionally researchers have attempted more non-parametric approaches). In terms of the probability-weighting functions π^- and π^+ , the literature which estimates CPT has used the same functional forms as the RDEU literature — indeed, the Tversky and Kahneman (1992) form reported above was suggested as part of CPT. The value function proposed by Tversky and Kahneman (1992) is reported in Table 1, Panel 3. In that specification, $\alpha \in (0, 1)$ and $\beta \in (0, 1)$ generate diminishing sensitivity in both the gain and loss domains, respectively.

¹⁹If $\pi^+(\mu) = 1 - \pi^-(1 - \mu)$, then the distinction between π^- and π^+ is irrelevant.

²⁰Here we use π^+ to weigh the second event because we are assuming that z is such that $x_2 - z \geq 0$.

The parameter $\lambda > 1$ reflects loss aversion, as it implies the negative value generated by a loss tends to be larger in magnitude than the positive value generated by an equal sized gain. Based on their experimental data, Tversky and Kahneman (1992) suggest that $\lambda = 2.25$, $\alpha = \beta = 0.88$, and for probability weighting as reported in Table 1, Panel 2, $\gamma^- = 0.69$ and $\gamma^+ = 0.61$.

When applying CPT, researchers must specify a reference point, and typically this is done using some external intuitive argument. For instance, in experiments it is typically argued that the reference point should be zero or experimentally endowed wealth.²¹ In field applications, researchers argue for a natural reference point given that application (e.g., in his recent analysis of tax evasion Rees-Jones (2014) argues that a zero balance due is a natural reference point). This extra “degree of freedom” in CPT is often seen as a limitation by some and it has led to attempts to tie down the reference point, which we review in the following sections.

3.4 Kőszegi-Rabin Loss Aversion

Kőszegi and Rabin (2006, 2007) proposed a form of “rational expectations” loss aversion wherein gains and losses are defined relative to expectations about outcomes — that is, the reference point is taken to be one’s expectations about outcomes. Moreover, because such expectations could involve uncertainty about future outcomes, they extend the model of loss aversion to use a reference lottery instead of a reference outcome.

Specifically, under Kőszegi-Rabin (KR) loss aversion, the utility from choosing lottery $X \equiv (x_n, \mu_n)_{n=1}^N$ given a reference lottery $\tilde{X} \equiv (\tilde{x}_m, \tilde{\mu}_m)_{m=1}^M$ is

$$V(X|\tilde{X}) \equiv \sum_{n=1}^N \sum_{m=1}^M \mu_n \tilde{\mu}_m [u(x_n) + v(x_n|\tilde{x}_m)].$$

The function u represents standard “intrinsic” utility defined over final wealth, just as in EU. The function v represents “gain-loss” utility that results from experiencing gains or losses relative to the reference lottery. Gain-loss utility depends on how a realized outcome x_n is compared to all possible outcomes that could have occurred in the reference lottery. For the value function, KR use

$$v(y|\tilde{y}) = \begin{cases} \eta [u(y) - u(\tilde{y})] & \text{if } u(y) > u(\tilde{y}) \\ \eta\lambda [u(y) - u(\tilde{y})] & \text{if } u(y) \leq u(\tilde{y}) \end{cases}.$$

²¹For the latter, if subjects are initially given $\$Z$ before making choices, it is sometimes argued that the reference point should be $\$Z$.

In this formulation, the magnitude of gain-loss utility is determined by the intrinsic utility gain or utility loss relative to consuming the reference point. Moreover, gain-loss utility takes a two-part linear form, where $\eta \geq 0$ captures the importance of gain-loss utility relative to intrinsic utility and $\lambda \geq 1$ captures loss aversion. The model reduces to expected utility when $\eta = 0$ or $\lambda = 1$.

KR propose that the reference lottery equals recent expectations about outcomes—i.e., if a household expects to face lottery $\tilde{X}$, then its reference lottery becomes $\tilde{X}$. However, because situations vary in terms of when a household deliberates about its choices and when it commits to its choices, KR offer multiple solution concepts for the determination of the reference lottery. Here, we focus on two solution concepts that are perhaps most relevant for field data:

Definition 2 (KR-PPE). *Given a choice set $\mathcal{X}$, a lottery $X \in \mathcal{X}$ is a personal equilibrium if for all $X' \in \mathcal{X}$, $V(X|X) \geq V(X'|X)$, and it is a preferred personal equilibrium if there does not exist another $X' \in \mathcal{X}$ such that X' is a personal equilibrium and $V(X'|X') > V(X|X)$.*

Definition 3 (KR-CPE). *Given a choice set $\mathcal{X}$, a lottery $X \in \mathcal{X}$ is a choice-acclimating personal equilibrium if for all $X' \in \mathcal{X}$, $V(X|X) \geq V(X'|X')$.*

KR suggest that PPE is appropriate when, faced with a choice set $\mathcal{X}$, a person thinks about the choice situation, decides on a planned choice $X \in \mathcal{X}$, and then makes that choice shortly before the uncertainty is resolved. An option X is a personal equilibrium if, when a person plans on that option and thus that option determines her reference lottery, it is indeed optimal to make that choice. Among the set of personal equilibria, the PPE is the personal equilibrium that yields the highest “utility”. In terms of field contexts, then, PPE is an appropriate solution concept when a person is able to think about a choice situation for some duration and then make a choice shortly before the uncertainty is revealed. Among those that we discuss in Sections 4 and 5, the field context which perhaps fits this scenario is betting on horse races.

The idea behind CPE is that, when faced with a choice set $\mathcal{X}$, a person commits to a choice well in advance of the resolution of uncertainty. By the time the uncertainty is resolved, the person will have become accustomed to its choice and hence expect the lottery induced by its choice. Hence, the person chooses the lottery that yields the largest utility conditional on that lottery being the reference lottery. Field contexts that have the feature of committing to a choice well in advance of the resolution of uncertainty are property and health insurance.

To illustrate these solution concepts, consider the implications in Example 1:²²

²²These equations are derived in Appendix 1. Because the equations get quite complex, we do not present

Example 1 (Under KR loss aversion with PPE). *A household's willingness to pay for full insurance against the possibility of losing L with probability μ is the z such that:*

$$u(w - z) = \left(\frac{\mu(1 + \eta\lambda)}{1 + \mu\eta\lambda + (1 - \mu)\eta} \right) u(w - L) + \left(1 - \frac{\mu(1 + \eta\lambda)}{1 + \mu\eta\lambda + (1 - \mu)\eta} \right) u(w).$$

Example 1 (Under KR loss aversion with CPE). *A household's willingness to pay for full insurance against the possibility of losing L with probability μ is the z such that:*

$$u(w - z) = \mu [1 + \eta(\lambda - 1)(1 - \mu)] u(w - L) + [1 - \mu [1 + \eta(\lambda - 1)(1 - \mu)]] u(w).$$

When estimating the KR model, one needs to estimate the parameters η and λ along with the utility function $u(y)$. Because the latter is meant to be standard utility over final wealth as used in EU, any of the functional forms for $u(y)$ in Table 1 might be used.

Finally, there are two further points with regard to CPE. First, note that in Example 1 the parameters η and λ always appear as a product $\eta(\lambda - 1)$. In fact, under CPE this always holds, and thus the model really has only one parameter $\Lambda = \eta(\lambda - 1)$. Second, CPE can sometimes yields strange predictions. For instance, one might naturally think that, ceteris paribus, if the probability μ of a loss increases, the willingness to pay z for full insurance should also increase. However, if $\Lambda > 1$, one can show that for μ close enough to one, the willingness to pay declines with μ . The intuition for this result is that, under CPE, a person has a strong aversion to risk. Hence, because intermediate probabilities (close to $1/2$) involve more risk than probabilities close to one, the person is willing to pay more for full insurance (even though the expected loss is smaller).²³

3.5 Disappointment Aversion

Models of “disappointment aversion” also assume that people’s choices are influenced by their expectations. The concept of disappointment aversion was proposed by Bell (1985) and further developed by Loomes and Sugden (1986) and Gul (1991). The basic idea is that a person is disappointed (or elated) if the realized outcome of a lottery is worse (or better) than “expected”.

the equations for Example 2, although we do present those equations in Appendix 1.

²³In fact, one can also show that, for any $\Lambda > 0$, under CPE a person can choose a dominated lottery. The intuition is much the same: the aversion to risk can be so strong that the person would rather choose a certain outcome over a risky lottery that dominates that certain outcome.

Bell (1985) proposes a variant of disappointment aversion in which disappointment is determined from a comparison of one's realized utility to one's expected utility, and the person accounts for expected disappointment when making a choice. Formally, a lottery $X \equiv (x_n, \mu_n)_{n=1}^N$ is evaluated as

$$V(X) = \sum_{n=1}^N \mu_n u(x_n) - \beta \sum_{n=1}^N \mu_n \left[I \left(u(x_n) < \sum_{n'=1}^N \mu_{n'} u(x_{n'}) \right) \left(\sum_{n'=1}^N \mu_{n'} u(x_{n'}) - u(x_n) \right) \right]$$

where I is an indicator function. In this formulation, the first term is the standard expected utility of lottery X . The second term reflects the expected disutility from disappointment that arises when the realized utility from an outcome is less than the standard expected utility of the lottery. The parameter β captures the magnitude of disappointment aversion, where the model reduces the expected utility for $\beta = 0$.²⁴

Gul (1991) proposes another variant of disappointment aversion in which disappointment is determined from a comparison of one's realized outcome to one's certainty equivalent for the lottery. Formally, a lottery $X \equiv (x_n, \mu_n)_{n=1}^N$ is evaluated as $V(x) = u(z)$ such that

$$u(z) = \sum_{n=1}^N \mu_n u(x_n) - \beta \sum_{n=1}^N \mu_n [I(u(x_n) < u(z)) (u(z) - u(x_n))]$$

where again I is an indicator function, and z represents one's certainty equivalent for lottery X .

To illustrate these solution concepts, consider the implications in Example 1.²⁵

Example 1 (Under Bell disappointment aversion). *A household's willingness to pay for full insurance against the possibility of losing L with probability μ is the z such that:*

$$\begin{aligned} u(w - z) &= \mu [1 + \beta(1 - \mu)] u(w - L) \\ &\quad + [1 - \mu [1 + \beta(1 - \mu)]] u(w). \end{aligned}$$

Example 1 (Under Gul disappointment aversion). *A household's willingness to pay for full*

²⁴Bell (1985) further assumes that (i) $u(x) = x$ and (ii) a person might also experience utility from elation when the realized outcome is larger than the expected utility. Even with the latter, however, his model reduces to the model in the text where β represents the difference between the marginal disutility from disappointment and the marginal utility from elation. Loomes and Sugden (1986) also use this formulation, except they study non-linear disappointment.

²⁵These equations are derived in Appendix 1. Because the equations get quite complex, we do not present the equations for Example 2, although we do present those equations in Appendix 1.

Model	WTP
EU	$\mu u(w - d - z) + (1 - \mu)u(w - z) = \mu u(w - L) + (1 - \mu)u(w)$
RDEU	$\pi(\mu)u(w - d - z) + (1 - \pi(\mu))u(w - z) = \pi(\mu)u(w - L) + (1 - \pi(\mu))u(w)$
CPT	$\pi^-(\mu)v(-d - z) + \pi^-(1 - \mu)v(-z) = \pi^-(\mu)v(-L)$
KR (CPE)	$\{\mu u(w - d - z) + (1 - \mu)u(w - z) - \Lambda(1 - \mu)\mu[u(w - z) - u(w - d - z)]\} = \{\mu u(w - L) + (1 - \mu)u(w) - \Lambda(1 - \mu)\mu[u(w) - u(w - L)]\}$
Bell	$\{\mu [1 + \beta(1 - \mu)] u(w - d - z) + [1 - \mu [1 + \beta(1 - \mu)]] u(w - z)\} = \{\mu [1 + \beta(1 - \mu)] u(w - L) + [1 - \mu [1 + \beta(1 - \mu)]] u(w)\}$
Gul	$\left\{ \left(\frac{(1+\beta)\mu}{1+\beta\mu} \right) u(w - d - z) + \left(1 - \frac{(1+\beta)\mu}{1+\beta\mu} \right) u(w - z) \right\} = \left\{ \left(\frac{(1+\beta)\mu}{1+\beta\mu} \right) u(w - L) + \left(1 - \frac{(1+\beta)\mu}{1+\beta\mu} \right) u(w) \right\}$

Table 2: Willingness to pay (z) for insurance with deductible d , against the possibility of losing L with probability μ .

insurance against the possibility of losing L with probability μ is the z such that:

$$u(w - z) = \left(\frac{(1 + \beta)\mu}{1 + \beta\mu} \right) u(w - L) + \left(1 - \frac{(1 + \beta)\mu}{1 + \beta\mu} \right) u(w).$$

When Bell disappointment aversion is applied to binary lotteries, the model is equivalent to the KR-CPE model — e.g., in Example 1, the equation that determines z is identical, where β replaces $\eta(\lambda - 1)$. Gul disappointment aversion yields a slightly different equation, though the structure is still quite similar. In contrast, for lotteries with more than two outcomes, while Bell disappointment aversion and Gul disappointment aversion are still structured similarly, their structure is qualitatively different from the KR model. For details, see Appendix 1.

When estimating models of disappointment aversion, one needs to estimate the parameter β along with the utility function $u(y)$. Because the latter is again meant to be standard utility over final wealth as used in EU, any of the functional forms for $u(y)$ in Table 1 might be used.

3.6 Model Predictions and Identification

Our goal in this section is to develop intuition for the type of data that may yield point identification of the model's parameters. We consider an idealized scenario, in which a researcher has data that allow her to learn households' willingness to pay for different lotteries,

and show what one can infer about risk preferences in this scenario. To aid our discussion, we use the framework in Example 1, in which households incur a loss L with probability μ . In Table 2 we report the households' willingness to pay for an insurance policy with a deductible $d \geq 0$, implied by each of the models presented in the previous sections; that is, we report the equation that one needs to solve to obtain the value z such that

$$(w - z, 1 - \mu ; w - z - d, \mu) \sim (w, 1 - \mu ; w - L, \mu).$$

Of course, z is a function of d and μ .

With the value $z(d, \mu)$ in hand, we examine whether point identification of the parameter vector characterizing the model can be achieved. Point identification requires that only one value of that vector is concordant with the joint distribution of the observable variables. When we consider models that encompass several features leading to aversion to risk (e.g., curvature of the utility function, and probability distortion function), a minimum requirement towards identification of the parameters characterizing each of these features, is that when these parameters take distinct values, they lead to distinct model predictions. In our setting, perhaps the most transparent way to illustrate how the model's predictions change with the parameters' values is to show how the implied willingness to pay for insurance (or for financial assets) changes.

3.6.1 Expected Utility

Table 2 reports an equation that implicitly defines households' willingness to pay for deductible insurance against a loss L that may occur with probability μ . Suppose we observe one specific deductible-probability combination denoted (d_0, μ_0) , and the associated willingness to pay, which in turn solves

$$\frac{u(w - z(d_0, \mu_0) - d_0) - u(w - L)}{u(w) - u(w - z(d_0, \mu_0))} = \frac{1 - \mu_0}{\mu_0}.$$

In order to learn risk preferences, much of the literature assumes a parametric functional form for $u(w)$, e.g., CARA, CRRA, or NTD, with a single parameter capturing the magnitude of risk aversion. These functional forms all fall in a class of utility functions that satisfy Assumption 1 below.

Denote by $u(x; \phi)$ the parametric utility function, where x is a final wealth state and ϕ is a taste parameter. Assume that u is continuous in both $x > 0$ and $\phi \in \mathbb{R}$, and that $\phi = 0$ if and only if $u(x; \phi) = x$.

Assumption 1. (i) $u(x; \phi)$ is increasing in x , and for any $x_0 > x_1 > x_2$, the ratio $R \equiv [u(x_1; \phi) - u(x_2; \phi)] / [u(x_0; \phi) - u(x_1; \phi)]$ is strictly increasing in ϕ .
(ii) $\lim_{\phi \rightarrow \infty} R = \infty$ and $\lim_{\phi \rightarrow -\infty} R = 0$.²⁶

Assumption 1 naturally associates ϕ with the magnitude of an individual's risk aversion. In particular, Assumption 1 holds if and only if for any $x_0 > x_1 > x_2$ and $\mu \in (0, 1)$, there exists a $\bar{\phi}$ such that $(x_0, 1 - \mu; x_2, \mu) \succ (x_1, 1)$ for $\phi \in [0, \bar{\phi})$, $(x_0, 1 - \mu; x_2, \mu) \sim (x_1, 1)$ for $\phi = \bar{\phi}$, and $(x_0, 1 - \mu; x_2, \mu) \prec (x_1, 1)$ for $\phi > \bar{\phi}$. In words, whenever a person compares a binary risky lottery to a certain amount in the support of the risky lottery, the person chooses the riskier lottery if her risk aversion (ϕ) is small enough (risk loving enough), and she chooses the less risky lottery if her risk aversion is high enough.

Under EU, for any $u(x; \cdot)$ satisfying Assumption 1, each ϕ implies a unique $z(d_0, \mu_0)$, and in particular, the larger is ϕ (the more risk averse the agent is) the larger is $z(d_0, \mu_0)$ (the more the agent is willing to pay for insurance). Hence, under EU with a parametric functional form for u , observation of individuals' willingness to pay associated with a single combination (d_0, μ_0) yields point identification of ϕ .

The literature most often assumes a parametric functional form for u , not only when estimating EU but also when estimating alternative models (as we discuss below). As we have seen, this assumption dramatically simplifies identification, but it is a strong restriction. It would be desirable to be able to trace out the utility function nonparametrically over the relevant support. However, this would also require strong assumptions about the nature of heterogeneity across people and information about their wealth (we return to this discussion in Section 5.2.1).

3.6.2 Rank Dependent Expected Utility

We next consider model predictions under RDEU. From the second line in Table 2, we obtain that under RDEU,

$$\frac{u(w - z(d, \mu) - d) - u(w - L)}{u(w) - u(w - z(d, \mu))} = \frac{1 - \pi(\mu)}{\pi(\mu)}.$$

Suppose we observe the households' willingness to pay associated with a specific combination (d_0, μ_0) . In this case, model predictions depend on both the utility function u and the decision weight $\pi(\mu_0) \equiv \pi_0$.²⁷ As a result, observing willingness to pay for insurance with

²⁶The limit assumption is made merely to guarantee interior solutions in any formal results below. In practice, this assumption is unlikely to be important. NTD does not satisfy this assumptions, but Result 1 goes through for NTD as well.

²⁷The discussion here assumes a fixed μ_0 , and thus focuses on identifying π evaluated at μ_0 . In order to identify the entire function $\pi(\cdot)$ on a given range of values for μ , one needs to observe willingness to pay for insurance at all probabilities of loss over the range of μ .

deductible d_0 does not suffice to point identify both π_0 and u , even when the utility function is parametrically specified. In particular, even in the parametric case, there is a set of (ϕ, π_0) pairs consistent with willingness to pay $z(d_0, \mu_0)$. Indeed, if we define $\bar{\pi}_0(\phi|z(d_0, \mu_0))$ to be the required π_0 as a function of ϕ that generates willingness to pay $z(d_0, \mu_0)$, we can rearrange the equality above to derive

$$\bar{\pi}_0(\phi|z(d_0, \mu_0)) = \frac{u(w; \phi) - u(w - z(d_0, \mu_0); \phi)}{[u(w; \phi) - u(w - z(d_0, \mu_0); \phi)] + [u(w - z(d_0, \mu_0) - d_0; \phi) - u(w - L; \phi)]}.$$

For any u that satisfies Assumption 1, $\bar{\pi}_0(\phi|z(d_0, \mu_0))$ is decreasing in ϕ (as depicted in Figure 4). Intuitively, both an increased risk aversion and an increased decision weight on the loss state, imply an increased willingness to pay for insurance. Hence, for a fixed willingness to pay, as risk aversion increases, the decision weight on the loss state must decline in order to keep the willingness to pay unchanged.

In order to separately point identify ϕ and π , we therefore need to observe richer data. In the insurance example, for instance, it would suffice to also observe the willingness to pay for an insurance with a different deductible d_1 , denoted $z(d_1, \mu_0)$. This willingness to pay yields another set of (ϕ, π_0) pairs, which can be represented by the curve $\bar{\pi}_0(\phi|z(d_1, \mu_0))$. As we establish in Result 1 below, these two curves cross at only one point, yielding a unique (ϕ, π_0) pair consistent with individuals' willingness to pay for each insurance plan.

Result 1. *If $u(x; \phi)$ satisfies Assumption 1, then for any $0 \leq d_0 < d_1 < L$ there exists a unique $\bar{\phi}$ such that*

- (i) $\bar{\pi}_0(\bar{\phi}|z(d_0, \mu_0)) = \bar{\pi}_0(\bar{\phi}|z(d_1, \mu_0))$,
- (ii) $\bar{\pi}_0(\phi|z(d_0, \mu_0)) < \bar{\pi}_0(\phi|z(d_1, \mu_0))$ for all $\phi < \bar{\phi}$, and $\bar{\pi}_0(\phi|z(d_0, \mu_0)) > \bar{\pi}_0(\phi|z(d_1, \mu_0))$ for all $\phi > \bar{\phi}$.

Proof: First note that, because u is increasing in x , $z(d_0, \mu_0) > z(d_1, \mu_0)$ while $z(d_0, \mu_0) + d_0 < z(d_1, \mu_0) + d_1$ (otherwise the individual would violate dominance). Define $A(\phi) \equiv u(w; \phi) - u(w - z(d_0, \mu_0); \phi)$, $B(\phi) \equiv u(w - z(d_0, \mu_0) - d; \phi) - u(w - L; \phi)$, $A'(\phi) \equiv u(w; \phi) - u(w - z(d_1, \mu_0); \phi)$, and $B'(\phi) \equiv u(w - z(d_1, \mu_0) - d'; \phi) - u(w - L; \phi)$, in which case $\bar{\pi}_0(\phi|z(d_0, \mu_0)) = A(\phi)/[A(\phi) + B(\phi)]$ and $\bar{\pi}_0(\phi|z(d_1, \mu_0)) = A'(\phi)/[A'(\phi) + B'(\phi)]$. Hence

$$\begin{aligned} \bar{\pi}_0(\phi|z(d_0, \mu_0)) \geq \bar{\pi}_0(\phi|z(d_1, \mu_0)) &\iff \frac{A(\phi)}{A(\phi)+B(\phi)} \geq \frac{A'(\phi)}{A'(\phi)+B'(\phi)} \\ &\iff \frac{B'(\phi)}{B(\phi)} \geq \frac{A'(\phi)}{A(\phi)}. \end{aligned}$$

Assumption 1 yields that $A'(\phi)/A(\phi)$ is a strictly decreasing function of ϕ , where $\lim_{\phi \rightarrow \infty} \frac{A'(\phi)}{A(\phi)} = 0$ and $\lim_{\phi \rightarrow -\infty} \frac{A'(\phi)}{A(\phi)} = 1$. Analogously, Assumption 1 yields that $B'(\phi)/B(\phi)$ is a strictly

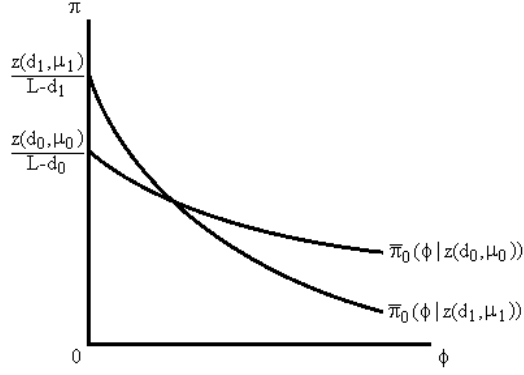


Figure 4: Identification in the RDEU model

increasing function of ϕ , where $\lim_{\phi \rightarrow \infty} \frac{B'(\phi)}{B(\phi)} = 1$ and $\lim_{\phi \rightarrow -\infty} \frac{B'(\phi)}{B(\phi)} = 0$. The result follows. ■

Figure 4 illustrates this result. The key intuition behind this result is that, for standard risk aversion (ϕ), the concern for reducing risk is stronger the more risk one bears, and so, for instance, the willingness to pay for full insurance is less than double the willingness to pay to eliminate half the risk (in our example, for $d = L/2$). Moreover, the larger one's standard risk aversion is, the stronger this asymmetry will be. Hence, the ratio of the willingness to pay to reduce some of the risk relative to the willingness to pay to reduce more of the risk serves to identify how much standard risk aversion is present. This intuition can be applied in other contexts as well.

We reiterate, however, that Result 1 relies on having a parametric specification for the utility function. Similar considerations as in the previous section apply a fortiori here, if one endeavors to trace out the utility function nonparametrically.

3.6.3 EU and RDEU with Heterogeneity in Preferences

The discussion above abstracts from unobserved heterogeneity in risk preferences. The presence of such heterogeneity implies that observationally equivalent people facing the same choice set, may make different choices (a feature that occurs in virtually any dataset). How can one point identify risk preferences in this case? Intuitively, in the presence of unobserved heterogeneity, the researcher needs to be able to observe how households in the population choose different options when facing a fixed menu. Variation in the choice menu can then be used to identify risk preferences and unobserved heterogeneity.

To illustrate, consider an EU model, with a parametric utility function $u(x; \phi)$ that satisfies Assumption 1. Suppose that in our insurance example the choice set consists of three options for one's deductible, $d_A > d_B > d_C$, with respective premiums $p_A < p_B < p_C$. Suppose that this menu partitions the support of ϕ , so that there exist ϕ_{AB} and ϕ_{BC} with $\phi_{BC} > \phi_{AB}$, such that individuals with $\phi < \phi_{AB}$ choose deductible d_A , individuals with $\phi \in (\phi_{AB}, \phi_{BC})$ choose deductible d_B , and individuals with $\phi > \phi_{BC}$ choose deductible d_C . If we then let F denote the cumulative distribution function of ϕ , the model predicts that a fraction $F(\phi_{AB})$ chooses deductible d_A , a fraction $F(\phi_{BC}) - F(\phi_{AB})$ chooses deductible d_B , and a fraction $1 - F(\phi_{BC})$ chooses deductible d_C . In other words, a menu composed of three choices identifies the value of $F(\cdot)$ at most at two points. If there is variation in choice sets conditional on observed characteristics, F can be uncovered from the data. We return to discussing this strategy in Section 5.2.2.

If instead we consider an RDEU model, there are potentially two sources of unobserved heterogeneity (in standard risk aversion, and in the decision weight π). In principle, a similar approach to what was delineated for the EU model might be applied. However, to date, point identification of multidimensional heterogeneity in risk preferences has relied upon parametric assumptions about their joint distribution. It remains a question for future research, to find a field setting and the proper set of assumptions to obtain nonparametric identification. We return to this discussion in Section 6.1.

3.6.4 Distinguishing More Complex Models

Finally, we discuss identification of more complex models, that include several features generating aversion to risk. For example, this could be a model that includes Koszegi-Rabin loss aversion (KR) or disappointment aversion (DA), as well as probability weighting and/or concave utility function. We organize our discussion in two parts: first we consider domains in which households or individuals only make choices over binary lotteries. Then we consider domains in which households or individuals make choices over lotteries with more than two outcomes.

If the data contains only households choices over binary lotteries $X \equiv (x_1, \mu; x_2, 1 - \mu)$, with $x_1 < x_2$, then all models considered here (except for KR-PPE) can be reduced to one in which households choose the lottery that maximizes

$$U(X) = \Omega(\mu)u(w + x_1) + (1 - \Omega(\mu))u(w + x_2), \quad (5)$$

where $\Omega(\mu)$ is a generic probability distortion function, as labeled in Barseghyan, Molinari, O'Donoghue, and Teitelbaum (2013b). As reflected in our presentation of Example 1 for the

various models that we consider in the previous sections, the function $\Omega(\mu)$ may take one of the following forms:

$$\begin{aligned} \text{Under RDEU:} \quad & \Omega(\mu) = \pi(\mu), \\ \text{Under KR-CPE:} \quad & \Omega(\mu) = \mu(1 + \Lambda(1 - \mu)), \\ \text{Under Bell-DA:} \quad & \Omega(\mu) = \mu(1 + \beta(1 - \mu)), \\ \text{Under Gul-DA:} \quad & \Omega(\mu) = (1 + \beta)\mu/(1 + \beta\mu). \end{aligned}$$

Hence, even if one has data that allows for nonparametric point identification of $\Omega(\mu)$, the underlying models can be distinguished only to the extent that they impose different restrictions on $\Omega(\mu)$. Clearly, KR-CPE and Bell-DA cannot be distinguished from each other. Each of KR-CPE/Bell-DA and Gul-DA impose strong parametric assumptions, and so does RDEU if a parametric functional form is assumed for $\pi(\mu)$ (e.g., any of those in Table 1). With these parametric assumptions, the models can be tested to find which best fits the data. However, a more flexible approach would allow for nonparametric probability weighting. If either loss aversion or disappointment aversion are also allowed for, identification of the resulting $\Omega(\mu)$ function would not allow one to disentangle these sources of aversion to risk. For example, KR-CPE and RDEU together yield $\Omega(\mu) = \pi(\mu)(1 + \Lambda(1 - \pi(\mu)))$, and similarly for RDEU together with Bell-DA or with Gul-DA, so that $\pi(\mu)$ and Λ (or β) cannot be separately identified.

Hence, with binary lotteries, without relying strongly on functional form assumptions, the best one can do is to focus on identification and estimation of $\Omega(\mu)$. Potential exceptions are represented by models that feature CPT or KR-PPE. Under CPT, $U(X)$ does not take the form in equation (5), because there is an exogenous reference point and utility varies with whether outcomes are above or below that reference point. Hence, one can potentially separately identify the CPT component of the model by studying how behavior changes as outcomes move above and below the reference point (or, alternatively, by studying how behavior changes as the reference point changes). Under KR-PPE, one cannot define a $U(X)$ independently of the other lotteries in the choice set. Hence, one can potentially separately identify the KR-PPE component of the model by studying how behavior changes as the choice set varies. Of course, one needs data with the right type of variation to pursue either of these approaches. At this time, we are not aware of any empirical work that attempts to distinguish CPT or KR-PPE from RDEU.

A key maintained assumption in the literature is that households subjective beliefs μ coincide with objective expectations (e.g., "objective" claim probabilities), which in turn the econometrician can estimate. However this assumption may fail in a given application. In that case, when μ is assumed to equal objective expectations, the estimated $\Omega(\mu)$ function

captures a mapping Ψ from the estimated objective probabilities to subjective beliefs, thereby yielding another possible source of probability distortions. In turn, this implies that under the assumption that households' behavior is governed by RDEU, for example, the estimated $\Omega(\mu)$ function would not necessarily correspond to $\pi(\mu)$, but to $(\pi \circ \Psi)(\mu)$.

When data contains lotteries with more than two outcomes, more refined inference is possible. For example, as pointed out in Barseghyan, Molinari, O'Donoghue, and Teitelbaum (2013a), for lotteries with more than two outcomes, a model with RDEU only, and a model with subjective beliefs only, as formalized by the mapping Ψ , generate different predictions. Intuitively, under the mapping Ψ alone, the weight on a particular event is independent of the magnitude of the outcome associated with that event. In contrast, under RDEU alone, the magnitude of the outcome associated with an event impacts the rank-ordering of outcomes, and thereby can influence the weight. In addition, KR-CPE and RDEU may generate predictions that are different predictions from those under either Bell-DA or Gul-DA. Under either KR-CPE or RDEU, the decision weight assigned to an event depends only on the rank-order of the outcome associated with that event. In contrast, under either model of DA, the decision weight assigned to an event depends on whether the outcome is above or below the relevant benchmark that determines disappointment—i.e., the expected utility of the lottery under Bell-DA, or the certainty equivalent under Gul-DA. Hence, variation in outcomes that does not change the rank-order but does change the magnitudes of intermediate outcomes can be used to distinguish these classes of models.

However, KR-CPE and RDEU cannot be separately identified even with choice data on lotteries with more than two outcomes. Indeed, one can show that, for lotteries with any number of outcomes, the combination of KR-CPE and RDEU reduces to an equivalent RDEU model using effective probability weighting $\Omega(\mu) = \pi(\mu)(1 + \lambda(1 - \pi(\mu)))$ applied from the same direction as in the original RDEU model. Hence, it is never possible to separately identify KR-CPE and RDEU.²⁸

We conclude with two important and related implications of this discussion. First, for models that purport to estimate an RDEU model, the estimated probability weighting functions might in fact reflect a combination of probability weighting and some other phenomenon—e.g., KR-CPE, Bell-DA, or Gul-DA. Second, given this fact, it would seem valuable to take a nonparametric approach to estimating RDEU models, as opposed to restricting attention to the functional forms in Table 1. As we summarize in Section 4.2, the more recent literature has taken a first step in this direction.

²⁸Masatlioglu and Raymond (2014) make a similar point using a decision theoretic approach.

4 Estimation with Individual-Level Data

In this section, we describe research that estimates risk preferences using individual-level data. We begin by discussing the general approach that is broadly used in all of these papers. We then review how that approach has been applied in several different domains: (i) property insurance, (ii) health insurance, and (iii) game shows.

4.1 The General Approach

Estimating risk preferences from individual data typically requires three main steps:

- Step 1: Identify a field context in which economic agents make choices between options that involve risk and for which the researcher can obtain data on both the agents' choice sets and the agents' choices.
- Step 2: Translate the (typically) rich field choice environment into a choice between a well defined set of lotteries (as formalized in Definition 1) — so that each model of risk preferences defined in Section 3 makes a prediction for which option should be chosen as a function of the taste parameters.
- Step 3: Enrich the basic models of risk preferences with some form of “noise,” because in practice observationally equivalent agents facing identical choice sets are observed to make different choices.

The following subsections describe Steps 2 and 3 in some detail, along with issues that arise in each. It is within these steps that the methodological toolkits of economics need to be applied. Step 1 requires the researcher to (i) identify field contexts that are likely to best reveal risk preferences, and (ii) be able to obtain the corresponding data. As we review below, the existing work using individual-level data has focused on some of the field contexts traditionally emphasized by the theoretical literature (e.g., insurance), for which researchers have been able to obtain sufficiently rich datasets. As more data sets become available, and as economists develop a better understanding of how to approach estimation of risk preferences, we expect the general approach described in this section to be applied more broadly.

4.1.1 Translation into Lotteries

Virtually any field choice environment involves a rich context and does not present itself in the form of lotteries to which one can immediately apply a model of risk preferences. Hence,

to begin the analysis, one needs to first make a series of assumptions that permits translating the rich field context into a choice between a well defined set of lotteries.

Consideration Sets. To begin this translation, one needs to make an assumption about what exactly is the set of options under consideration — i.e., the consideration set. In some cases, this appears straightforward. For instance, in papers which study people’s deductible choices for property insurance, a natural assumption is that the consideration set includes each of the possible deductibles. Even here, however, assumptions are being made — e.g., in some cases, there are available deductibles that are very rarely chosen, and so one might wonder whether people are actually even considering these possibilities, and one might assume these deductibles are not part of the choice set. Also, most analyses get data from one company, and thus one cannot observe options that a household might have considered purchasing from another insurance company. We return in Section 6.4 to discuss when assumptions about the consideration set are likely to be important, and when they are not.

Outcomes. A closely related assumption the researcher needs to make, concerns the set of possible outcomes that might occur. In most field contexts, each option is associated with many possible outcomes. Returning to the example of deductible choices, during the policy period an individual might incur zero losses, one loss, two losses, three losses, and so forth. However, the researcher often restricts the set of possible outcomes. For example, some analyses reviewed below in the context of property insurance assume that households focus only on the possibilities of either zero or one loss, and they ignore the (small probability event) of multiple losses during a single policy period. We return in Section 6.4 to discuss when assumptions about the set of possible outcomes are likely to be important, and when they are not.

Subjective Beliefs. A particularly important step in the translation from the field context to a set of lotteries, concerns what one assumes about agents’ subjective beliefs about the likelihood of the possible outcomes. In some field contexts, there exist clear objective probabilities for outcomes — e.g., in games of chance such as state-run lotteries or casino roulette. In most field contexts, however, objective probabilities either do not exist or are very hard to assess. For such situations, an ideal approach would be to simultaneously estimate both the household’s beliefs and preferences. As we shall see in Section 6.3, however, this presents a fundamental identification problem. Hence, the most common approach to date has been to assume “rational expectations,” in the sense that agents’ subjective beliefs correspond to objective probabilities (often, but not always, as reflected in past outcomes). The researcher then estimates probabilities over outcomes based on realized outcomes, and then takes these estimated probabilities as “data,” in the sense that they are treated as an

observed input when estimating preferences.²⁹

Moral Hazard. A final important assumption concerns whether beliefs over outcomes depend on one’s choices — e.g., whether there is moral hazard in an insurance context. In fact, most analyses that estimate risk preferences assume there is no moral hazard, although a few directly study it. We discuss moral hazard more in Section 4.1.3.

Before moving on to Step 3, it is worth comparing the issues raised in this section, to experimental analyses. A major advantage of the experimental context is that it offers significantly more control over the translation from the observed choice context to choices between a well defined set of lotteries. This is because the experimenter fully defines the choice set, the set of possible outcomes associated with each choice set, and the probabilities over outcomes.

4.1.2 Introducing “Noise”

Once Step 2 is carried out, for each agent the researcher observes a choice set $\mathbf{X}$ and the agent’s optimal choice $X^* \in \mathbf{X}$. In addition, upon making some assumptions on beliefs, the researcher often proceeds as if she has data that reveal individuals’ subjective beliefs, here denoted μ . Each model of risk preferences from Section 3 generates a utility function $V(X, \mu, \theta)$, where θ denotes a set of taste parameters (this set differs, depending on the model). And each model of risk preferences generically predicts a unique optimal choice for any specific beliefs μ and taste parameters θ . In virtually any data set, however, observationally equivalent agents facing identical choice sets are observed to make different choices.³⁰ Hence, one must enrich the basic models of risk preferences with some form of “noise”. There are three main approaches used in the literature (and some analyses incorporate two or even all three) to accomplish this.

First, and perhaps most common, one can apply *random utility* (McFadden (1974)). Here, one assumes that agents choose the option $X \in \mathbf{X}$ that maximizes “total utility” $W(X, \mu, \theta) \equiv V(X, \mu, \theta) + \varepsilon(X)$, where $\varepsilon(X)$ is an idiosyncratic noise term that makes the agent more or less likely to choose option X relative to what the agent’s risk preferences $V(X, \mu, \theta)$ would predict. Then one specifies a joint distribution for the $\varepsilon(X)$ ’s for all $X \in \mathbf{X}$, and the random utility model predicts a probability distribution over the set of possible options. Typically, the $\varepsilon(X)$ ’s are assumed to be iid with a type-1 extreme value distribution with scale parameter σ ,³¹ in which case the predicted probability of choosing option $X^* \in \mathbf{X}$

²⁹In most cases, estimation error in this stage is not accounted for when reporting standard errors for the estimates of preferences.

³⁰Indeed, in experiments it is not uncommon for individual subjects to make different choices when presented the same choice situation more than once.

³¹The scale parameter σ is a monotone transformation of the variance of $\varepsilon(X)$, and thus a larger σ means

is

$$\Pr(X^*|\theta, \sigma) = \frac{\exp(V(X^*, \mu, \theta)/\sigma)}{\sum_{X \in \mathbf{X}} \exp(V(X, \mu, \theta)/\sigma)}.$$

Second, one can assume unobserved heterogeneity in preferences. Here, one posits that agents choose the option $X \in \mathbf{X}$ that maximizes their risk preferences as reflected by $V(X, \mu, \theta)$. However, one further assumes that there is unobserved (to the researcher) heterogeneity in preferences, and thus the researcher will observe different choices even among observationally equivalent agents facing identical choice sets. One then specifies a distribution for θ , and the model predicts a probability distribution over the set of possible outcomes. In particular, if $\Theta(X^*)$ is the set of θ such that X^* is the optimal choice, and if $F(\theta)$ reflects the distribution of θ , then the predicted probability of choosing option $X^* \in \mathbf{X}$ is

$$\Pr(X^*|\theta, F) = \int_{\theta \in \Theta(\mathbf{X}^*)} dF(\theta).$$

Third, one can assume unobserved heterogeneity in beliefs. Here, one again posits that agents choose the option $X \in \mathbf{X}$ that maximizes their risk preferences as reflected by $V(X, \mu, \theta)$. However, one relaxes the assumption that subjective beliefs μ are “data”, and instead assumes that there is unobserved (to the researcher) heterogeneity in beliefs. One then specifies a distribution for μ — typically where average beliefs correspond to objective probabilities, but there is a distribution around that average — and the model predicts a probability distribution over the set of possible outcomes. In particular, if $\mu(\mathbf{X}^*)$ is the set of μ such that X^* is the optimal choice, and if $G(\mu)$ reflects the distribution of μ , then the predicted probability of choosing option $X^* \in \mathbf{X}$ is

$$\Pr(X^*|\theta, G) = \int_{\mu \in \mu(\mathbf{X}^*)} dG(\mu).$$

Once one enriches the model by incorporating one or more of these approaches, one then estimates the parameters — both the taste parameters and the noise parameters — by making the predicted distribution of choices “match” the empirically observed distribution of choices (using MLE, MCMC, or some other econometric technique).

Two final comments about the three approaches above. First, in principle, the noise could reflect either permanent differences across economic agents or random variation within each agent across different choice situations. In cross-sectional data, one could never distinguish one from the other, but if one observes multiple choices from each agent, one can start to

larger variance. In general, one must either normalize the scale of the utility function $u(w)$ or the scale of the choice noise σ . The literature typically uses parametrizations of u that normalize its scale (see the three examples in Section 3), and thus σ is typically estimated.

tease these apart. In Section 6.1, we discuss in detail the extent to which economic agents exhibit stable risk preferences vs. risk preferences that change from choice to choice. Second, returning again to the comparison with the experimental setting, we note that while the latter permits great control over the issues raised in Step 2, experimental analyses must deal with “noise” — especially random utility or unobserved heterogeneity in risk preferences — just as much as field analyses do.

4.1.3 Moral Hazard and Adverse Selection

Much of the literature estimating risk preferences using individual-level data focuses on the insurance context. In this domain, economists have long discussed the problems of moral hazard and adverse selection, and thus we comment on these issues here.

In the context of insurance, moral hazard refers to the idea that individuals who have more insurance coverage will have less incentive to take care, and thus are more likely to incur a loss — in other words, people’s risk is endogenous to their choices. Such endogenous risk — which in principle could arise in other environments as well — can create problems in two ways for estimation of risk preferences. First, because it is hard to account for it, many analyses assume that it does not exist. For such analyses, if in fact moral hazard does exist, then estimates of risk preferences can be distorted. Second, if moral hazard exists, then there must be a reason why behavior is changing — e.g., in the insurance context, the fact that people take less care when they have more coverage presumably reflects that they get some form of utility from taking less care. If so, then this feature of preferences should be incorporated into the analysis.

Concerning adverse selection, in the insurance context this refers to the idea that individuals who bear more risk (which is not observable or cannot be priced) are more likely to purchase higher insurance coverage. As such, it is worth highlighting that adverse selection is a problem for the insurance company. However, it is not of concern for a researcher trying to estimate risk preferences insofar as unobserved heterogeneity in risk is taken into account, even if this unobserved heterogeneity is correlated with unobserved heterogeneity in risk preferences.³²

For many of the property insurance domains, moral hazard appears to play a small role (for a recent review of the literature, see Cohen and Siegelman (2010)). Most studies that test for the presence of asymmetric information in auto insurance markets using cross-sectional data do not find evidence of a positive correlation between risk and coverage (Chiappori 1999; Chiappori and Salanié 2000; Chiappori, Jullien, Salanié, and Salanié

³²For a leading example, see Cohen and Einav (2007), which is discussed in detail in the next Section.

2006; Dionne, Gouriéroux, and Vanasse 1999, 2001).³³ These studies have been interpreted as casting doubt on the presence of moral hazard (Cohen 2005), at least in auto insurance. Recently, a handful of papers separately test for moral hazard in longitudinal auto insurance data using a dynamic approach pioneered by Abbring, Chiappori, and Pinquet (2003) and Abbring, Chiappori, Heckman, and Pinquet (2003). Abbring, Chiappori, and Pinquet (2003) find no evidence of moral hazard in French auto insurance data. Israel (2004) reports a small, but statistically significant moral hazard effect for drivers in Illinois. Ceccarini (2007), Dionne, Michaud, and Dahchour (2007), Abbring, Chiappori, and Zavadil (2008), and Pinquet, Dionne, Vanasse, and Mathieu (2008) present stronger evidence of moral hazard using auto insurance data from Italy, France, the Netherlands, and Quebec, respectively. Each of the foregoing papers, however, identifies a moral hazard effect with respect to either liability coverage or a composite coverage that confounds liability coverage with other coverages. None of them identifies or quantifies the separate moral hazard effect (if any) directly attributable to collision and comprehensive auto coverages.³⁴

There is a significant body of work assessing the presence of asymmetric information in other markets, in particular the health and life insurance markets. Cardon and Hendel (2001) estimate a structural model of health insurance and health care choices using individual-level data. They find no evidence of informational asymmetries. Cawley and Philipson (1999) look for the presence of asymmetric information in term life insurance markets, using measures for both actual and self reported risk. They too find no evidence of asymmetric information. More recent work has found increasing evidence for informational asymmetries. As in the prior two works, Finkelstein and Poterba (2004) find no informational asymmetries in the U.K. annuity market based only on annuity amounts or “insurance payout”. However, they find informational asymmetries for other annuity characteristics, e.g. payout timing and whether an estate is guaranteed payment. Further, in long-term care insurance, Finkelstein and McGarry (2006) find evidence for multiple dimensions of private information, which can separately lead to moral hazard, adverse or advantageous selection. Thus, focusing on the failure to reject a positive correlation between insurance coverage and risk occurrence ignore the fact that these selection forces may “cancel” each other out. The findings of Fang, Keane, and Silverman (2008) support the notion of advantageous selection in Medigap

³³We are aware of two exceptions: Puelz and Snow (1994) and Cohen (2005). Recent work criticizes the methodology of Puelz and Snow (1994) and suggests that their results are spurious (Chiappori 1999; Dionne, Gouriéroux, and Vanasse 1999). Cohen (2005) presents mixed evidence of asymmetric information. She finds that a positive coverage-accident correlation exists for new customers with more than three years of driving experience, but not for new customers with fewer than three years of driving experience. To our knowledge, there are no studies that test for asymmetric information in home insurance.

³⁴These are the auto coverages that have been used to date, to estimate risk preferences with property insurance data

insurance. Bajari, Dalton, Hong, and Khwaja (2014), using semi-parametric methods, report significant moral hazard/adverse selection in a recent (2002 - 2004) data from a large self-insured U.S. employer. He (2009) finds evidence of asymmetric information in life insurance markets in a sample of potential new buyers. Bundorf, Levin, and Mahoney (2012) calculate that asymmetric information between consumers and the risk-adjustment software used by health insurance providers, lowers annual welfare by a \$35–\$100 per enrollee in data from small employers. Einav, Finkelstein, Ryan, Schrimpf, and Cullen (2013) find evidence of moral hazard in individuals' choices of health insurance plans.

We conclude with raising one more issue closely related to unobserved heterogeneity. If there is unobserved heterogeneity in risk preferences, then the researcher needs to be confident that the choice sets faced by the agents are independent of those risk preferences. In particular, risk preferences are estimated by investigating how agents react to changes in choice sets, and thus our estimates would be biased if a change in choice sets were correlated with a change in risk preferences.³⁵ In principle, there is a natural reason to be very worried about this issue. In the insurance context, for instance, if an insurance company can easily get a sense of a person's risk preferences, the company might be able to increase profits by adjusting the insurance pricing in reaction to those preferences. In practice, most insurance pricing is not done in this way — often due to heavy regulation — and thus this issue is perhaps less of a problem.

4.2 Property Insurance Data

Insurance choices are a natural domain in which to estimate risk preferences. Surprisingly, due to difficulties in obtaining data prior to the 2000s (when firms were not equally willing as they are today to share their data) there are relatively few papers which use individual-level insurance data to carry out this task.

The first paper to use individual-level data on insurance choices to estimate risk preferences is Cicchetti and Dubin (1994). They analyze data from Mountain Bell on roughly 10,000 residential telephone customers in Colorado in 1990. The choice of interest is whether customers purchased inside-wire insurance. This insurance cost roughly \$0.45 per month, and protected against telephone-wire problems inside one's residence. Without the insurance, in the event of a problem, the household would need to pay the service cost to fix the problem, which averaged about \$55. The probability of experiencing a problem was roughly 0.5% (see below), and thus the expected benefit of the insurance was roughly \$0.275. Hence,

³⁵One may think of this problem as "adverse selection for the researcher." We avoid this terminology as it may cause confusion with respect to the typical use of the terms adverse selection in economics.

purchasing this insurance for \$0.45 is a clear sign of aversion to risk. Cicchetti and Dubin set out to get a more precise understanding of this aversion to risk.

In order to translate the field context into the domain of preferences (a choice between lotteries, as discussed in Section 4.1), Cicchetti and Dubin assume that households effectively choose between the following two (incremental) lotteries:

$$(-p, 1) \quad \text{vs.} \quad (-R, \mu; 0, 1 - \mu)$$

The first lottery is that associated with the choice to purchase the insurance, where p is the premium charged for the insurance. The second lottery is that associated with the choice not to insure, where μ is the probability of experiencing a problem within any given month, and R is the expected service cost in the event of a problem.³⁶ The authors estimate μ using data on historical (1982-1986) trouble calls, where they divide the data into nine zones, and then they take the observed ratio of trouble calls to customers in each zone to be the probability of a problem for customers in that zone. Then, in their empirical analysis, they treat p , R , and μ as data.

Cicchetti and Dubin then estimate an EU model with a HARA utility specification. For prior wealth — a required input in the HARA specification — they use a measure of monthly income generated from census data. They also allow the curvature of the utility function to depend on a household’s average monthly bill.³⁷ In addition to estimating an EU model, they also estimate an RDEU model with a parametrized form of overweighting of probabilities (the one in Lattimore, Baker, and Witte (1992)).

They estimate these models using MLE, where they use a McFadden (1974) random utility specification to account for observationally equivalent households making heterogeneous choices. From their estimates, Cicchetti and Dubin conclude that “the overall pattern of results is remarkably consistent with expected utility theory” (p. 183). In particular, they find virtually no evidence of overweighting of probabilities — i.e., RDEU does no better than EU — and moreover 78 percent of households have an estimated utility function consistent with EU. Finally, they note that, for the average household, the estimated degree of absolute risk aversion is relatively small and yields a willingness to pay for the insurance virtually identical to the expected benefit from the insurance.

The Cicchetti and Dubin paper has a number of limitations. Perhaps most important

³⁶This formulation implicitly assumes that households expect at most one problem in any given month. As we mentioned in Section 4.1, this type of assumption is common in the literature.

³⁷They motivate this assumption based on a reduced-form finding that household’s with larger average monthly bills are more likely to purchase inside-wire insurance. While this assumption will capture that feature in their structural estimation of risk preferences, it is not clear that there is a good primitive justification for it.

is a data limitation: their data contain very little variation in p , R , and μ .³⁸ Hence, their estimation primarily identifies the impact of monthly income and average monthly bill on the (local) degree of absolute risk aversion. Moreover, as we discussed in Section 3.6, the limited variation in p , R , and μ in their data creates a major impediment to separately identifying both utility curvature and probability weighting, and thus their RDEU estimates depend heavily on their functional-form assumptions for utility and probability weighting. As such, one should be cautious in interpreting their estimates for the degree of probability overweighting. Finally, upon closer look, it is not so clear that EU fares so well. For the 22 percent of households with an estimated utility function inconsistent with EU, the inconsistency is in the form of a decreasing utility function (i.e., less is better), which is clearly counterfactual. In addition, among the 78 percent of households with an estimated utility function consistent with EU, nearly half of them (37 percent of all households) are estimated to be risk loving. Finally, in addition to permitting choices to depend on risk preferences, Cicchetti and Dubin also permit an idiosyncratic direct preference for the insurance (perhaps due to a belief that households with insurance will get priority service relative to households without insurance). In the estimated model, this direct preference is estimated to be quite large and appears to be the main factor explaining the choice to purchase insurance for many households. Despite these limitations, the Cicchetti and Dubin paper was the first of its kind in the domain of insurance, and inspired later papers that were able to work with better data and more sophisticated models.

A key contribution in this direction is given by Cohen and Einav (2007), which provides a much more sophisticated analysis, although the authors limit themselves to the EU framework. They use individual-level data from an insurance provider in Israel, and analyze deductible choices among households who purchased one particular form of auto insurance (similar to comprehensive automobile insurance in the United States). In their data, all households have purchased the insurance, and the choice of interest is which of four deductible options they chose. The data contain the full menu of premium-deductible combinations offered to each household, along with that household's chosen deductible. In addition, they observe actual claims made by these households during the policy year.

Since virtually all households (98.9 percent) chose one of the two lowest deductibles, Cohen and Einav limit attention to the choice between those two.³⁹ Hence, they translate the field context into the domain of preferences — i.e., into a choice between lotteries — by

³⁸The paper seems to suggest that, in the estimation, a single value of p and R is used for all households. While there is some variation in μ across the nine zones, its range of [0.32%, 0.74%] is very tight.

³⁹Cohen and Einav do not drop households which chose one of the two high deductibles, but rather they proceed as if those households had chosen the highest remaining deductible. For virtually all of these households, this assumption is consistent with the estimated model.

assuming that households were effectively choosing between the following two (incremental) lotteries:

$$(-p_l - d_l, \mu; -p_l, 1 - \mu) \quad \text{vs.} \quad (-p_h - d_h, \mu; -p_h, 1 - \mu)$$

The first lottery is that associated with the choice of the lower deductible, and the second lottery is that associated with the choice of the higher deductible. For each choice $k \in \{l, h\}$, p_k is the premium charged for the insurance, d_k is the deductible that must be paid by the household in the event of a loss, and μ is the probability of experiencing a loss during the policy period. The insurance is typically priced such that lower deductibles are actuarially unfair — i.e., a risk neutral household would choose the higher deductible d_h , and d_l would be optimal if the household is risk averse enough.⁴⁰

This formulation reflects three simplifying assumptions. First, it assumes that a household's claim probability μ is independent of its deductible choice — i.e., that there is no moral hazard with respect to the deductible choice. Second, it assumes that every possible loss is larger than the larger deductible d_h . Cohen and Einav explicitly discuss these assumptions, and argue that they are supported in their data. The third simplifying assumption is that households ignore the possibility of incurring more than one claim during the policy period. This assumption is more implicit and follows from their explicit assumption that households make decisions while considering a policy period that is infinitesimally small. But since multiple claims are rare (in their data, only 2.70 percent of households experience more than one claim over the one-year policy period), this assumption is probably not restrictive (as we discussed in Section 4.1, similar assumptions are made in subsequent analyses).⁴¹

Cohen and Einav then estimate an EU model with an NTD specification for the utility function.⁴² Importantly, they permit both observed heterogeneity (i.e., depending on household observables) and unobserved heterogeneity in both the degree of (absolute) risk aversion r and the likelihood of a claim μ . Indeed, the main goal of their analysis is to assess (i) the extent of such heterogeneity, (ii) the relative importance of heterogeneity in risk vs. in risk preferences, and (iii) the correlation between the unobserved elements of the heterogeneity. Formally, they assume that claims are determined by a Poisson process with Poisson claim

⁴⁰Under their benchmark estimates of claim rates, this is true for 98.7 percent of households (p. 752)).

⁴¹Although Cohen and Einav's approach effectively assumes that households ignore the possibility of multiple claims during the policy period, their estimation still uses the full distribution of number of claims over the policy period, because this distribution is needed to identify the variance and correlations of unobserved heterogeneity in risk (as we discuss below).

⁴²In fact, their utility equation is slightly different from the one in Section 3 because they consider a policy period that is infinitesimally small.

rate λ , and they assume that r and λ have a bivariate lognormal distribution with

$$\begin{pmatrix} \ln \lambda_i \\ \ln r_i \end{pmatrix} \sim^{iid} N \left(\begin{bmatrix} x'_i \beta_\lambda \\ x'_i \beta_r \end{bmatrix}, \begin{bmatrix} \sigma_\lambda^2 & \rho \sigma_\lambda \sigma_r \\ \rho \sigma_\lambda \sigma_r & \sigma_r^2 \end{bmatrix} \right).$$

In the above expression, σ_λ^2 is the variance of λ , σ_r^2 is the variance of r , and ρ is their correlation.

Because they permit such unobserved heterogeneity, Cohen and Einav do not introduce any other source of heterogeneity in choices — i.e., unobserved heterogeneity in risk preferences and in risk account for all heterogeneity in choices among observationally equivalent households.

Cohen and Einav estimate this model using an MCMC approach. The data contain several key features that permit them to parametrically identify the unobserved heterogeneity in risk preferences and in risk. First, they observe the full distribution of (the number of) household claims over the course of the one-year policy period. This distribution permits them to directly estimate the mean and variance of risk (for a fixed set of observables) without any reference to premium-deductible menus or to choices.⁴³ Second, different households in the data are offered different menus of (p_l, d_l) vs. (p_h, d_h) .⁴⁴ Seeing how households respond to different menus (while knowing the mean and variance of risk) permits identification of the mean and variance of risk aversion. Finally, seeing how these responses differ between households with different claims experiences (i.e., households with a different number of actual claims experienced during the one-year policy period) permits identification of the correlation between the unobserved components of risk and risk preferences.

Cohen and Einav’s estimation yields several interesting conclusions. First, in terms of the degree of risk aversion, mean risk aversion is quite large ($r = 0.0019$), but in fact the distribution is quite skewed and the median is more reasonable ($r = 0.0000073$). In other words, in the estimated model, the roughly 18 percent of households that choose the low deductible are primarily explained by very high risk aversion. Second, they find more unobserved heterogeneity in risk aversion than in risk, and moreover, given their estimates, the unobserved heterogeneity in risk preferences is more important for profits and pricing. Finally, they find a strong positive correlation between unobserved heterogeneity in risk preferences and unobserved heterogeneity in risk. This is driven by the fact that as the observed number

⁴³Note that if they only observed whether or not households made any claims, they could only estimate the mean claim rate. The variance of claim rates is identified by comparing the likelihood of at least one claim to the likelihoods of at least two or more claims.

⁴⁴This variation is not purely idiosyncratic. For instance, for two-thirds of households, the menu satisfies $d_h = 0.5p_h$, $d_l = 0.3p_h$, and $p_l = 1.06p_h$, where the variation across households comes entirely from variation in p_h . But there is some additional variation, which Cohen and Einav discuss in some detail.

of claims increases, the proportion of households that choose the low deductible increases rapidly. They are cautious in concluding too much from this positive correlation, because the unobservables which drive a correlation between risk aversion and risk might be very context specific.

Sydnor (2010) uses similar data to also study the implications of EU for insurance deductible choices. However, he does not pursue an estimation of preferences, but rather a calibration approach in the spirit of Rabin (2000). Nonetheless, this paper is instructive for those who want to estimate risk preferences using insurance data. Sydnor uses individual-level data from a large home insurance company, from which he obtained a random sample of 50,000 home insurance policyholders in a single state in a single year. At this company, each policyholder must choose one of four deductibles — \$1000, \$500, \$250, or \$100 — where a lower deductible implies a higher premium. Much like Cohen and Einav (2007), Sydnor observes the full menu of premium-deductible combinations offered to each household, along with that household’s chosen deductible. In addition, he observes actual claims made by these households during the policy year, from which he can derive claim rates.

Sydnor translates this field context into the domain of preferences exactly as in Cohen and Einav (2007), by assuming that households are effectively choosing between four lotteries of the form:

$$(-p_d - d, \mu; -p_d, 1 - \mu)$$

where p_d is the premium charged for the insurance, d is the deductible that must be paid by the household in the event of a loss, and μ is the probability of experiencing a loss during the policy period. This formulation reflects the same three simplifying assumptions as in Cohen and Einav (2007).

In his main analysis, Sydnor considers an EU model with CRRA utility, and he calibrates the degree of relative risk aversion (ρ) that households would need to have to explain their choices. More precisely, for the 6268 customers who were new at the sample firm in the sample year — and who thus are more likely to have actively chosen their deductible — he focuses on the 3791 customers who chose either the \$500 or the \$250 deductible.⁴⁵ Choosing a deductible smaller than the maximum (\$1000) reveals an aversion to risk, and Sydnor calibrates for each household a lower bound on how risk averse that household must be. In his baseline calibration, he assigns to each household a claim rate equal to the average claim rate among those who chose that deductible, and he assumes a prior wealth of \$1 million. He demonstrates that, for the vast majority of the households that chose deductibles smaller than \$1000, this specification implies implausibly large risk aversion. He further

⁴⁵Because only 3 of the 6268 new customers choose the \$100 deductible, he does not analyze that group.

demonstrates that this finding is robust to assuming a variety of values for prior wealth and to assigning to each household a fitted claim rate (in much the same way as described below for Barseghyan, Molinari, O’Donoghue, and Teitelbaum (2013b)). Based on these calibrations, Sydnor concludes that EU is not a good explanation of people’s deductible choices, and he then discusses several potential alternative explanations, including risk misperceptions, probability weighting, and Koszegi-Rabin loss aversion.

Barseghyan, Molinari, O’Donoghue, and Teitelbaum (2013b, BMOT henceforth) use similar data to estimate a probability distortion model (as described in Section 3). In other words, they consider a model that, in the insurance-deductible context, nests EU, RDEU, KR loss aversion, and risk misperceptions, and they investigate which combination best explains the variation in the data.⁴⁶ Their data come from a large insurance company that sells multiple lines of coverage. The full dataset comprises yearly information on more than 400,000 households who held auto or home policies between 1998 and 2006. For their main analysis, BMOT use a core dataset of 4170 households who were new customers at the firm in 2005 or 2006 and who purchased both home and auto insurance. The authors focus on households’ initial choices. In this group, every household made a deductible choice for three coverages: home all perils insurance, auto collision insurance, and auto comprehensive insurance.⁴⁷ For each choice, the data contain the full menu of premium-deductible combinations offered to each household, along with that household’s chosen deductible.⁴⁸ In addition, the data contain the history of claims for the full dataset.

The translation of the field domain into the domain of lotteries is exactly as in Cohen and Einav (2007) and Sydnor (2010), except that BMOT make one further assumption: they assume that households bracket the three choices narrowly in the sense that, while each household has fixed risk preferences that apply to all three choices, the household treats these choices as three independent decisions of the form studied by Cohen and Einav (2007) and

⁴⁶As discussed in Section 3.6, BMOT can only identify an overall probability distortion, and cannot break it down into these possible underlying forces. Nonetheless, this approach identifies the extent to which these underlying forces together might help better explain deductible choices (relative to EU).

⁴⁷Auto collision coverage pays for damage to the insured vehicle caused by a collision with another vehicle or object, without regard to fault. Auto comprehensive coverage pays for damage to the insured vehicle from all other causes (e.g., theft, fire, flood, windstorm, glass breakage, vandalism, or hitting or being hit by an animal or by falling or flying objects), without regard to fault. Home all perils coverage pays for damage to the insured home from all causes (e.g., fire, windstorm, hail, tornado, vandalism, or smoke damage), except those that are specifically excluded (e.g., flood, earthquake, or war).

⁴⁸The available deductible options were the same for all households. For home, the options were 100, 250, 500, 1000, 2500, and 5000; for collision, the options were 100, 200, 250, 500, and 1000; and for comprehensive, the options were 50, 100, 200, 250, 500, and 1000. In order to make it plausible that every potential claim is larger than the largest deductible, BMOT exclude the two highest deductibles in home (chosen by only 1.6 percent of households); households that chose such deductibles are kept in the analysis and treated as if they chose a 1000 deductible (this approach is analogous to that in Cohen and Einav (2007)).

Sydnor (2010). It is worth highlighting that the assumption of narrow bracketing is in fact implicit in Cohen and Einav (2007) and in Sydnor (2010), in the sense that they assume that people have risk preferences that apply to their single observed deductible choice without reference to all the other risk choices that households are making.⁴⁹

BMOT estimate a probability distortion model. Applied to the context of choosing deductibles, the utility from choosing deductible d is

$$\Omega(\mu)u(w - p_d - d) + (1 - \Omega(\mu))u(w - p_d),$$

where μ is the (objective) probability of a claim, u is a utility function, and Ω is a probability distortion function. As a preliminary step, BMOT use the full sample and the full history of claims to estimate, for each coverage, a Poisson panel regression with random effects. They then use the output from these claim-rate regressions to assign to each household a fitted probability of incurring a loss (μ) on each coverage. In the estimation of preferences, these claim probabilities are treated as data.⁵⁰

In their benchmark analysis, BMOT estimate a model of homogeneous preferences, i.e., where all households have the same utility function u and the same probability distortion function Ω . They assume a NTD specification for u (see Table 1), but they take a nonparametric approach to Ω . Specifically, they consider three nonparametric approaches: (i) a Chebyshev polynomial expansion of $\ln \Omega(\mu)$, (ii) a Chebyshev polynomial expansion of $\Omega(\mu)$, and (iii) a cubic spline. All three approaches yield the same conclusion: While there is statistically significant curvature in u and statistically significant probability distortions, economically the lion’s share of households’ observed aversion to risk is attributed to probability distortions. This result represents a more convincing demonstration than Sydnor (2010) of the limitations of EU in explaining deductible choices. Whereas Sydnor (2010) merely provides calibration arguments against EU based on the required curvature of the utility function being “too large”, BMOT find that a model with large probability distortions better explains the data without imposing any restrictions on the magnitude of curvature in the utility function. BMOT compare special cases of models yielding probability distortions, including KR loss aversion and Gul disappointment aversion (see Section 3 for details on these models), with the nonparametric estimate of $\Omega(\mu)$, to assess whether these models are

⁴⁹In their benchmark estimation, the only advantage (relative to Cohen and Einav (2007)) of having three choices per household is that it yields increased variation in claim probabilities and prices represented in the data. When BMOT incorporate observed and unobserved heterogeneity into the analysis, this feature will then add a restriction that each household must have the same risk preferences across all three coverages.

⁵⁰They demonstrate that the main results are robust to instead assigning to each household a fitted distribution of claim probabilities for each coverage. As such, the approach in BMOT is roughly equivalent to that in Cohen and Einav (2007), except that they do not permit the unobserved heterogeneity in risk to be correlated with the unobserved heterogeneity in risk preferences.

consistent with the empirical evidence. The results show that neither KR loss aversion alone nor Gul disappointment aversion alone can explain the estimated probability distortions, thereby suggesting a key role for probability weighting in individuals' deductible choices.

BMOT next expand their analysis to incorporate both observed and unobserved heterogeneity in both the curvature of the utility function and in the magnitude of the probability distortions. When allowing for unobserved heterogeneity, BMOT's econometric model takes the form of a mixed logit (with parametric assumptions on the distribution of the unobserved heterogeneity terms), which they estimate via MCMC. The mean estimated probability distortions in a model that allows for observed heterogeneity only, for unobserved heterogeneity only, or for both, are nearly identical to each other and to the estimated probability distortions in a model with homogeneous preferences. Hence, whether BMOT assume preferences are homogeneous or allow for observed or unobserved heterogeneity, their main message is the same: large probability distortions characterized by substantial overweighting of claim probabilities and only mild insensitivity to probability changes. By contrast, the estimated degree of standard risk aversion is somewhat sensitive to the modeling approach. With observed heterogeneity only, the mean fitted value of r is 0.00073, slightly higher than in their model with homogeneous preferences (in which $r = 0.00064$). Allowing for unobserved heterogeneity leads to an additional increase to $r = 0.00156$, and allowing for both observed and unobserved heterogeneity further increases the estimate to $r = 0.00147$. The variance estimates for the unobserved heterogeneity terms suggest a substantial presence of unobserved heterogeneity (though smaller than in Cohen and Einav (2007)); unobserved heterogeneity in r and in probability distortions are estimated to be negatively correlated.

4.3 Game Shows

Several papers estimate risk preferences using data from television game shows. Game shows provide an attractive setting for estimating both EU and non-EU models. This is because contestants "are presented with well-defined choices where the stakes are real and sizeable, and the tasks are repeated in the same manner from contestant to contestant" (Andersen, Harrison, Lau, and Rutström 2008, p. 361). Field settings rarely have such desirable characteristics, and in fact a game show does not meet our definition of a field setting (a setting in which we observe individuals' real-world economic behavior). Rather, a game show is really a "controlled natural experiment" (*ibid.*) conducted in a peculiar environment that raises concerns about ecological and external validity, observer bias, and selection bias.⁵¹ For this reason, we provide below only a brief discussion of the pioneers of

⁵¹Indeed, concerns about observer and selection bias would seem to be greater in game shows than in controlled laboratory experiments.

the literature and references to a selection of subsequent studies.

The pioneers of the game show literature include Gertner (1993), who uses data from "Card Sharks," Metrick (1995), who uses data from "Jeopardy!," and Beetsma and Schotman (2001), who use data from the Dutch show "Lingo."

Gertner (1993) assumes that Card Sharks contestants are expected utility maximizers with CARA utility and pursues two methods for estimating the lower bound on the coefficients of absolute risk aversion r . First, Gertner looks solely at the contestants' bets in the final round of the game and obtains a lower bound on each contestant's r by assuming that a contestant who wagers her entire stake is risk neutral and a contestant who wagers the minimum amount (half of her stake) wishes to bet exactly that. He reports an average lower bound of 0.000310. Second, Gertner looks at the contestants' bets in all rounds and compares the sample distribution of outcomes with the distribution of outcomes if a contestant plays the optimal strategy for a risk-neutral contestant. By revealed preference, an "average" contestant prefers the former distribution to the latter, and so the degree of absolute risk aversion that would make a contestant indifferent between the two distributions provides an estimate of the average lower bound on r . Under this method, Gertner reports an average lower bound of 0.0000711.⁵²

Metrick (1995) also models Jeopardy! contestants as expected utility maximizers with CARA utility. In the pertinent part of the paper, Metrick looks only at the bets of the leaders in the final round of "runaway" games, i.e., games in which the leader is so far ahead entering the final round that she can guarantee a win by betting a sufficiently small amount. From the first-order necessary condition of the leader's utility maximization problem, Metrick derives an expression for the leader's subjective probability of answering correctly given her bet and coefficient of absolute risk aversion r . This expression is exactly the form of a logit regression. He then obtains an ML estimate of r for the "representative" contestant (i.e., he obtains the r that maximizes the likelihood of the observed sample of bets and correct/incorrect answers). He reports a statistically insignificant point estimate of 0.000066.⁵³

Beetsma and Schotman (2001) consider expected utility maximization with both CRRA and CARA utility functions. They analyze the decision at the start of a round in the

⁵²In the second part of the paper, Gertner presents evidence that contestants' bets are inconsistent with expected utility maximization. More specifically, he finds that contestants' bets exhibit sensitivity to accumulated winnings in the current round (stakes) but not to accumulated winnings in previous rounds (wealth).

⁵³In another part of the paper, Metrick looks at the bets of the first and second place contestants entering the final round of games in which these contestants can mostly ignore the actions of the third contestant. In this part, however, Metrick does not focus on estimating the contestants' risk preferences. Rather, he focuses on testing whether the contestants play "empirical-best-responses," i.e., best-responses to the empirical frequency of strategies played by their opponents in his sample of similar games. He finds that while first-place contestants generally play empirical-best-responses, second-place contestants do not.

Lingo finals to stop or continue play of the game. They first show that, under fairly weak conditions, the stop/play decision amounts to a choice between receiving the current stake x with certainty and a lottery in which they receive $2x$ with probability p and a with probability $1 - p$, where a represents the option value of coming back in the next show.⁵⁴ They then specify a probit model for the stop/play decision and estimate the degree of risk aversion. For the CRRA specification, they estimate coefficients of relative risk aversion ranging from 0.42 (assuming zero wealth) to 6.99 (assuming wealth of 50,000 Dutch guilders) to 13.08 (assuming wealth of 100,000 Dutch guilders). For the CARA specification, they estimate a coefficient of absolute risk aversion of 0.12 (which, assuming wealth of 50,000 Dutch guilders, implies a degree of relative risk aversion of 6.0).

Most of the subsequent papers in the game show literature use data from the U.S. and international versions of "Deal or No Deal." For surveys, see Andersen, Harrison, Lau, and Rutström (2008), Post, van den Assem, Baltussen, and Thaler (2008), and Hartley, Lanot, and Walker (2014). Two exceptions are Fullenkamp, Tenorio, and Battalio (2003), who use data from "Hoosier Millionaire," and Hartley, Lanot, and Walker (2014), who use data from "Who Wants to be a Millionaire?" Like the pioneering studies, many of the subsequent papers work with EU models with CRRA and/or CARA utility (e.g., Fullenkamp, Tenorio, and Battalio 2003; Andersen, Harrison, Lau, and Rutström 2008; Deck, Lee, and Reyes 2008; Conte, Moffatt, Botti, Di Cagno, and D'Ippoliti 2012; Hartley, Lanot, and Walker 2014). Others go beyond EU and consider RDEU, prospect theory, or other non-EU models (e.g., Botti and Conte 2008; Post, van den Assem, Baltussen, and Thaler 2008; Mulino, Scheelings, Brooks, and Faff 2009; De Roos and Sarafidis 2010; Bombardini and Trebbi 2012).

4.4 Health Insurance Data

Recently, a growing empirical literature has made use of structural models of decision making under uncertainty to answer a variety of important questions about health insurance markets. These include, *inter alia*, the welfare cost of inertia and imperfect information in health insurance and the implications of adverse selection and moral hazard in these markets.

However, the health insurance field context has not been used for estimation of risk preferences as the object of fundamental interest. We suspect the reason is that estimation of risk preferences when using health insurance choices is more challenging than it is when using property insurance choices for at least three reasons. First, the set of outcomes associated with each lottery is significantly more complex; for example, health expenses have (essentially) a continuous distribution. Second, individuals may care about more than mere

⁵⁴Under the rules of the game, if contestants lose they can come back in the next show unless the current show is their third.

monetary costs of care; for example, the quality of their life is differentially impacted by different health outcomes. The researcher then either needs to model the utility from each possible health status directly, or to monetize health status. Third, moral hazard is likely to be a larger concern in this context. In particular, if individuals' choices of health insurance plans are subject to selection on moral hazard, this selection cannot be ignored in the estimation of preferences.

Even when the ultimate target of the analysis is not estimation of risk preferences, the literature studying health insurance choices needs to contend with the difficulties listed above. One approach has been to not model directly individuals' risk aversion and simplify the problem by posing a reduced form equation for individuals' valuation of health insurance plans; see for example Starc (2014). Another approach has been to exogenously impose risk preferences, as opposed to estimating risk preferences within a larger model. See for example Einav, Finkelstein, and Schrimpf (2010), who use external data to calibrate the values for risk aversion (as well as other parameters that would only be identified via functional form assumptions in their model).

Because of this, we substantially limit the depth of our coverage of this literature. Nevertheless, we highlight a few recent papers that address some of these challenges and estimate a fully specified expected utility model with CARA preferences. In certain cases, these papers also allow for preference heterogeneity across individuals. We expect that these contributions could facilitate future work that focuses on estimating risk preferences.

Handel (2013) leverages a major change to insurance provision that occurred at a large firm to quantify the welfare cost of consumer inertia in health insurance markets,⁵⁵ and to study policies that could mitigate this inertia. In order to make progress on this important question, Handel assumes away two of the challenges listed above: he assumes that there is no moral hazard and that individuals base their choice of insurance plan only on the monetary costs of each option. However, Handel uses a very careful approach to account for complex lotteries over outcomes.

Handel observes households choosing among three "preferred provider organization" (PPO) plans, denoted PPO₂₅₀, PPO₅₀₀ and PPO₁₂₀₀. These PPO plans differ among each other in their premiums and cost sharing characteristics (e.g., deductible, coinsurance, and out-of-pocket maximums). These characteristics in turn determine the mapping from total medical expenditures to employee out-of-pocket expenditures.

The households in the sample are observed making insurance plan choices at multiple points in time. Specifically, at time $t = 0$ they are observed making an "active" choice, because the company changed the plans offered (their premiums and cost sharing character-

⁵⁵Within the welfare analysis, the author also considers the effects of adverse selection.

istics) and the individuals had to choose a plan among the new offerings. At times $t = 1$ and $t = 2$ individuals are observed making (potentially) "passive" choices, in the sense that they may simply continue with the plan chosen at time $t = 0$ without reassessing their options. Of notice is the fact that over time, the PPO₂₅₀ (the plan which yields the more comprehensive coverage) becomes substantially more expensive, while the PPO₅₀₀ (a plan which yields a lower coverage) becomes substantially less expensive; Figure 2 in Handel (2013) illustrates this fact. As such, choices that were optimal at time $t = 0$ need not be optimal at times $t = 1, 2$.

Handel's interest is in quantifying the effect of the *inertia* displayed by individuals who do not change their plan over time. To achieve this goal, he sets up a random utility model of insurance choice, similar to the one in BMOT described above. One important difference is that, in the case of health insurance, the lottery is defined over out of pocket expenditures, rather than over binary outcomes. The distribution of out of pocket expenditures is assumed known by the households. In practice, it is estimated using medical predictive software, household characteristics, and ex-post claim realizations. Specifically, for each individual and open enrollment period, the author uses the past year of diagnoses, drugs, and expenses, along with age and gender, to predict mean total medical expenditures for the upcoming year using the Johns Hopkins ACG Case-Mix software package. He then incorporates medically relevant metrics such as type and duration of specific conditions, as well as comorbidities. This is done for four distinct types of expenditures: (i) pharmacy; (ii) mental health; (iii) physician office visit; and (iv) hospital, outpatient, and all other. Individuals are then grouped into cells determined by the mean predicted future utilization. For each expenditure type and risk cell, Handel estimates a spending distribution for the upcoming year based on ex-post observed cost realizations, combining the marginal distributions across expenditure categories into joint distributions using empirical correlations and copula methods. He then maps individual total expense projections into the family out-of-pocket expense projections, taking into account family-level plan characteristics. This yields the distribution of out of pocket expenditures to be used in the choice model.

The effect of inertia is modelled as an additive "consumption" term which is positive for the chosen option in the previous year, yielding

$$\begin{aligned} c_0 &= w_0 - pp_0 - oe_0, \\ c_1 &= w_1 - pp_1 - oe_1 + \eta 1(IC_1 = IC_0), \end{aligned}$$

where c_t denotes consumption, w_t denotes wealth, pp_t denotes price of insurance, oe_t denotes out-of-pocket expenditure, each at time $t \in \{0, 1\}$, IC_t denotes insurance plan chosen at

time $t \in \{0, 1\}$, and $1(\cdot)$ is the indicator function of the event in parenthesis. In addition, for households with a very high medical expenditure risk, the consumption equation includes an additional dummy variable that puts additional weight on the most comprehensive choice, P_{250} , to capture the idea that essentially all of these households have chosen P_{250} .

The so specified consumption enters a CARA utility function, which in turn households use to evaluate the three PPO options (by integrating it against the distribution of out of pocket expenditure). Then households choose the one that delivers the highest expected utility.

The coefficient of risk aversion r is modeled as a random coefficient (following a Normal distribution with mean specified as a linear function of observable characteristics), while the coefficient of inertia η is modeled as a linear function of family status and demographics. The resulting econometrics model is a mixed logit and is estimated via MCMC. Unobserved heterogeneity in risk, however, is not allowed for.

The estimates of the coefficient of absolute risk aversion are similar in magnitude to the ones in BMOT, and they lie in $[1.9, 3.25] \cdot 10^{-4}$. The intercept in the inertia function η , is estimated to be large in magnitude with values of \$1,729 for single employees and of \$2,480 for employees who cover at least one dependent. An employee who enrolls in a flexible spending account (FSA) is estimated to forgo \$551 less due to inertia than one who does not. Because inertia is linked to multiple dimensions of observable heterogeneity, Handel also reports the estimated mean and variance of inertia: The mean total money left on the table per employee due to inertia is \$2,032 with a standard deviation of \$446.

A natural extension of Handel’s model is to study additional frictions that may affect an individual’s choice of insurance plans. For example, Handel and Kolstad (2015) study the effect of information frictions and hassle costs, measured via a survey administered to individuals for whom health insurance information is available. The survey elicits individuals’ information about available medical providers/treatments and perceived time and hassle costs to learn the characteristics of a high deductible health insurance plan (which is cheaper than the other options). The additional friction measures appear to be important predictors of choices and to significantly decrease risk preference estimates.

Einav, Finkelstein, Ryan, Schrimpf, and Cullen (2013) address the difficulties associated with evaluation of outcomes in the health insurance context, as well as the presence of moral hazard. In particular, moral hazard and its impact on selection, is the key object of interest in that paper.

The authors begin their analysis with the key observation that households’ utilization rate of medical services may depend significantly on the characteristics of their health insurance coverage. They propose the following simple model, which builds on the original model in

Cardon and Hendel (2001). Each period is divided into two sub-periods. In the second sub-period, households take insurance coverage as given, and their utility function is assumed separable in health and money

$$u(m; \lambda, \omega) = h(m - \lambda; \omega) + y(m),$$

where m is the monetary value of the health care utilization chosen, λ is the monetary value of the health shock, $y(m)$ is the residual income, which is decreasing in m at a rate that depends on the health insurance coverage, and ω is a parameter that captures households' responsiveness to the price of medical utilization.

The authors assume that $h(m - \lambda; \omega) = m - \lambda - \frac{1}{2\omega}(m - \lambda)^2$ and that, under insurance contract j , the marginal cost of health care is c_j (i.e., $y'(m) = -c_j$). It follows that the optimal amount of utilization is

$$m_j^*(\lambda) = \lambda + \omega(1 - c_j). \quad (6)$$

While these assumptions are quite strong, they play a key role in keeping the model tractable. In particular, through equation (6) the authors assure that the health shock contributes to utilization additively separably from the contribution of moral hazard.

Denoting by $F_\lambda(\cdot)$ the distribution of the health shock, the authors aim at identifying $F_\lambda(\cdot)$ and ω . To argue that these functionals are point identified, the authors rely heavily on (i) the fact that incremental utilization, which is observed, does not depend on risk aversion, but on ω ; and (ii) that the data contains a plausibly exogenous change in the entire menu of health coverage options. First, the authors consider a counter-factual (but ideal) data setting in which one observes each household's entire distribution of medical expenditure for two different coverages (denoted I and II), call them $G_i^I(m)$ and $G_i^{II}(m)$. Assuming that the distribution of observed health shocks $F_\lambda(\cdot)$ is the same under both coverages, using the fact that expenditure should equal $\lambda + \omega(1 - c_j)$, and denoting by $E_{G_i^k}(m)$ the expected value of utilization under distribution $G_i^k(m)$, $k \in \{I, II\}$, it immediately follows that the difference $\left(E_{G_i^I}(m) - E_{G_i^{II}}(m)\right) / (c_I - c_{II})$ uncovers the parameter ω . The distribution $F_{i,\lambda}(\cdot)$ for each individual i can also be learned. This is because the authors observe a panel data. If individuals within this panel are observed for a sufficiently long period under the different coverages, and if $F_{i,\lambda}(\cdot)$ is time invariant, then $F_{i,\lambda}(\cdot)$ can be learned recalling that $\lambda_{it} = m_{it} - \omega_i(1 - c_t)$, with t denoting time period.

Risk aversion is introduced in the analysis by monetizing the second period utility $u(m_j^*(\lambda); \lambda, \omega)$ and using this object as the argument of a CARA utility function with coef-

ficient of risk aversion r . Hence, the expected utility from coverage j in this model is

$$EU_j^\omega = -\frac{1}{r} \int \exp(-u(m_j^*(\lambda); \lambda, \omega)) dF_\lambda(\lambda).$$

An individual specific coefficient of standard risk aversion r_i can then be identified if individuals face a continuous option set.

The authors argue that the data they have comes sufficiently close to this ideal scenario. Specifically, their data hails from Alcoa, Inc., a large multinational producer of aluminum and related products, and it comprises health insurance choices and medical care utilization of their United States-based workers (and their dependents). The analysis in this paper focuses mostly on data from the years 2003 and 2004. The data is very rich, documenting health insurance options and choices, claim information, demographic characteristics of households, etc., as well as a summary proxy of individuals health care utilization based on predictive medical software.

For this paper, the distinctive feature of the data is that in 2004, Alcoa introduced a new set of health insurance PPO options, which were phased in gradually to different employees based on their union affiliation, since new benefits could only be introduced when an existing union contract expired. The staggered timing in the transition from one set of insurance options to another provides a plausibly exogenous source of variation in coverage (mimicking the possibility of observing distributions $G_i^I(m)$ and $G_i^{II}(m)$ from the ideal scenario) which is, as discussed above, the key to identification. Prior to 2004 there were three PPO options under the old benefits and five entirely different PPO options under the new benefits. Hence workers were forced to make an active choice.

In practice, to estimate the model the authors impose a parametric structure on the distribution function $F_\lambda(\cdot)$ and on the form of unobserved heterogeneity. Because of the model's fairly complex structure, estimation is carried out via MCMC. The estimation exercise yields an average coefficient of absolute risk aversion of $1.9 \cdot 10^{-3}$, with a large standard deviation of $2 \cdot 10^{-3}$. The average value for ω is about \$1,300 which corresponds to one third of the average health risk (about \$4,340) per employee-year. In other words, on average, moving from no insurance to full insurance will cause an increase in medical utilization of \$1,300.

The authors interpret the parameter ω as (a measure of) moral hazard. In essence, ω maps the change in health insurance coverage into a change in utilization rate. And just like in the textbook version of adverse selection, ceteris paribus individuals with higher moral hazard (higher ω) will have higher willingness to pay for insurance and will be more costly to insure.

5 Estimation with Aggregate Data

In this section, we describe research that estimates risk preferences using aggregate data. We begin by discussing the general approach that is broadly used in all of these papers. We then review how that approach has been applied in several different domains: (i) betting markets, and (ii) consumption, asset returns and labor supply data.

5.1 The General Approach

Estimating risk preferences from aggregate data requires steps that are similar to the ones faced when working with individual-level data, as well as some steps that are different.

- Step 1: Identify a field context in which economic agents make choices between options that involve risk, and for which the researcher can obtain data on both the choice set faced *in aggregate* by the agents, and the agents *aggregate choices*.
- Step 2: Translate the (typically) rich field choice environment into a choice between a well defined set of lotteries (as formalized in Definition 1). Several of the considerations presented in Section 4.1.1 continue to apply. In particular, while the set of possible outcomes might be clear in some contexts, e.g., betting markets, assumptions are going to be imposed in others, e.g. when the analysis is based on consumption, asset returns and labor supply data. In addition, subjective beliefs (of the representative agent) play the same role and pose the same challenges when working with aggregate data, as they do in the case of individual-level data. On the other hand, moral hazard is unlikely to be a concern in these contexts.
- Step 3: Determine if and to what extent the *representative agent* assumption which is appropriate in the specific context, can be weakened. The main departures entail allowing for either heterogeneity in beliefs or for heterogeneity in risk preferences. Due to the limitations of aggregate data, it seems unlikely that both dimensions of heterogeneity can be allowed for, see Section 5.2.2 below.

5.2 Betting Markets

At the center of the literature that estimates risk preferences with aggregate data are studies that use data on betting markets, and in particular those that use data on betting in pari-mutuel horse races. In a pari-mutuel horse race, the total amount wagered, the "betting pool," net of a house take, is distributed among the winning bettors in proportion to their individual bets. Specifically, suppose that there are n horses, $i = 1, \dots, n$, in each race,

and denote by τ the house take. Then the net return on a one dollar bet on horse i , or equivalently the "odds" on horse i , is given by

$$R_i = \frac{1 - \tau}{s_i} - 1, \quad (7)$$

where s_i is the "share" of the betting pool wagered on horse i . The house take is computed using the fact that the betting shares must sum to one:

$$1 - \tau = \left(\sum_{i=1}^n \frac{1}{1 + R_i} \right)^{-1}. \quad (8)$$

The betting shares are then computed using equation 7:

$$s_i = \frac{(1 - \tau)}{(1 + R_i)}. \quad (9)$$

Typically, researchers working in this field context have access to data that collect, for each race, an observed vector of odds $\mathbf{R} = (R_1, \dots, R_n)$ and the identity of the winning horse.

The conceptual framework behind the estimation of preferences from such data is as follows. The underlying primitives of a race are the house take τ , and a vector of probabilities $\mathbf{p} \equiv (p_1, \dots, p_n) \gg 0$, where p_i the probability that horse i wins the race. Bettors' preferences determine which horse they bet on, where the equilibrium market shares $\mathbf{s} = (s_1, \dots, s_n)$ determine the equilibrium odds $\mathbf{R}$ according to equation (7), and the market shares must be consistent with bettors maximizing their preferences given the probabilities $\mathbf{p}$ and the odds $\mathbf{R}$. Importantly, given observed odds $\mathbf{R}$, one can immediately derive the market shares $\mathbf{s}$ and the house take τ using equations (9) and (8) respectively. Hence, the literature typically takes $\mathbf{R}$, $\mathbf{s}$, and τ to be observed, whereas $\mathbf{p}$ and preferences are unobserved. The papers that we review in this section differ in terms of the estimation approach adopted and the nature of preferences considered.

5.2.1 Representative Agent Framework

The bulk of the studies that use pari-mutuel horse race data, the first wave of which considers only EU preferences, adopt a representative agent framework. The key implication of there being a representative agent is that, for each race, the equilibrium market shares $\mathbf{s}$ and odds $\mathbf{R}$ must be such that the representative agent is indifferent between betting on all horses.

In an early paper, Griffith (1949) uses data on a sample of 1,386 U.S. horse races to study the risk perceptions of the racetrack crowd. Griffith groups horses by odds and effectively compares, at each odds group, (i) the reciprocal odds after correction for the house take (i.e.,

$s/(1-s)$), which he takes as expressing the representative bettor's subjective win probability, with (ii) the empirical frequency of winners, which he takes as the objective win probability. He finds that the corrected reciprocal odds exceed the frequency of winners at short odds (less than 5.1-to-1), while the reverse holds at long odds. From this he concludes that there is a systematic undervaluation of the chances of (and underbetting on) short-odded horses and overvaluation of the chances of (and overbetting on) long-odded horses. This phenomenon has come to be known as the favorite-longshot bias. Griffith (1961) reports similar findings in a follow-up study in which he turns attention from win bets to "show" bets.

McGlothlin (1956) repeats and extends Griffith's study using data from 9,248 California horse races. Like Griffith, McGlothlin groups horses by odds. He then calculates the expected value of a one dollar win bet at each odds group and examines how it varies across odds groups. In calculating the expected value of the bet, McGlothlin (like Griffith) treats the empirical frequency of winners as the objective win probability. He finds that the bet's expected value is positive at short odds and negative at long odds, with the crossing point located between 3.5-to-1 and 5.5-to-1, in agreement with Griffith's prior findings. What is more, McGlothlin finds the same patterns when he breaks down the data by the position of the race in the daily program (first race, second race, etc.). At the same time, he finds that both the average win bet and the proportion of win bets (relative to higher probability place and show bets) increase as the racing day proceeds. From this he concludes that the data are inconsistent with the EU model. In particular, he argues that if the subjective evaluation of low probability-high stakes bets increases over a series of races, EU would predict a concomitant intensification of the pattern of positive EV at short odds and negative EV at long odds, *contra* to the data which displays a stable pattern over the racing day.

Weitzman (1965) analyzes data from over 12,000 New York horse races to estimate the utility function of the representative bettor (the "average man at the race track," or "Avmart"). Like Griffith and McGlothlin, Weitzman first groups horses by odds and treats the empirical frequency of winners $\hat{p}_j$ at each odds group j as the objective win probability, yielding 257 data points $\{\hat{p}_j, R_j\}_{j=1}^{257}$. He then uses these data points to estimate by weighted least squares a relationship between $\hat{p}$ and R . In the estimation, he tests several function forms and concludes that the most appropriate function is the "corrected" hyperbolic form, $\hat{p} = A/R + B \log(1 + R)/R$, which yields a remarkable fit ($R^2 = 0.9852$, although he does not account for the fact that $\hat{p}$ was estimated in the first stage). Weitzman argues that this curve represents Avmart's indifference map between lotteries (probability-returns pairs), as Avmart is assumed to place his bets to maximize his expected utility. Weitzman then uses this indifference map to recover the shape of Avmart's utility function, which appears convex (increasing marginal utility) in the range of his data, suggesting a region of local risk loving

similar to that proposed by Markowitz (1952).

The first wave of papers culminates with Ali (1977). Ali uses data on more than 20,000 New York harness races to revisit the relationship between subjective and objective win probabilities. Unlike his predecessors, Ali does not group horses according to their odds. Instead, he groups horses according to the order statistics of their odds. That is, horses are grouped according to "favorites"—the horse with the lowest odds in a race is the first favorite, the horse with the next lowest odds is the second favorite, and so forth. As the vast majority of races in his sample have at most eight horses, Ali carries out his analysis under the assumption that there are no more than eight horses in a race.

Let the h -th favorite in a race be called horse h , $h = 1, \dots, 8$, and let there be N races indexed by $k = 1, \dots, N$. Similar to his predecessors, Ali estimates the subjective win probability for horse h by $\bar{s}_h = \sum_{k=1}^N s_{hk}/N$, where s_{hk} is the (observed) share of win bets placed on horse h in race k . He estimates the objective win probability by the empirical frequency of wins by horse h over all races, $\hat{p}_h = \sum_{k=1}^N y_{hk}/N$, where $y_{hk} = 1$ if horse h wins race k equals zero otherwise. Ali then uses these estimates to document the existence of the favorite-longshot bias. He finds that every subjective win probability, except for that of the first favorite, exceeds the corresponding objective win probability, and moreover that the difference increases (when expressed as a ratio $(\bar{s}_h - \hat{p}_h)/SE(\hat{p}_h)$, with $SE(\cdot)$ denoting the standard error of the estimator in parenthesis) as the horse becomes less favorite. Ali also uses these estimates to recover the utility function of the representative bettor. It is remarkably similar to the one estimated by Weitzman.

In addition to reporting these empirical findings, Ali makes two important theoretical points. First, he observes that the utility function of the representative bettor is inconsistent with the EU hypothesis when the betting opportunity set is expanded to include parlay, martingale and other compound bets (a point that is leveraged by Snowberg and Wolfers (2010); see below). Second, he proves that the favorite-longshot bias can be explained with bettors that are risk neutral EU maximizers but have heterogeneous risk perceptions (a point that is developed by Gandhi and Serrano-Padial (2014); see below).

The second wave of papers that use pari-mutuel horse race data consider EU and non-EU preferences. Jullien and Salanié (2000) use data from more than 34,000 horse races in Britain ("in which bets are placed with a bookmaker who offers odds on all horses," p. 506), to estimate EU, RDEU, and CPT models. Critically, Jullien and Salanié assume that bettors do not spread their bets among horses, that every horse is bet on by at least one bettor, and that all bettors bet the same amount a in all circumstances. They posit a representative bettor whose behavior corresponds to the aggregate behavior of the crowd at the track and assume the bettor's subjective win probabilities correspond to the objective win probabilities

$(p_1, \dots, p_n)$. After placing a one dollar bet on horse i , the bettor receives R dollars if the horse wins, which happens with probability p_i , and loses one dollar if the horse does not win, which happens with probability $1 - p_i$.

The representative bettor's risk preferences are parametrized through a vector θ which depends on the model considered (e.g., EU, RDEU, etc.), and for each model Jullien and Salanié derive its implication for what the equilibrium probability p_i of winning the race for each horse $i = 1, \dots, n$ should be as a function of R and θ . Recognizing this relationship between objective win probabilities, on the one hand, and returns and model parameters, on the other, is a key contribution of this paper. The parameters are then estimated using a standard maximum likelihood procedure for multinomial models based on which horse won the race.

Jullien and Salanié begin by estimating EU models. Because they assume that bettors do not spread their bets among horses and that every horse is bet on by at least one bettor, the equilibrium value of betting a on each horse must be the same.⁵⁶ Hence,

$$p_i u(M + aR_i, \theta) + (1 - p_i) u(M - a, \theta) = W \quad \text{for all } i = 1, \dots, n,$$

where W is a placeholder for the value of betting a on horse i , and M is the bettor's initial wealth. Solving this equation for p_i and then using the fact that $\sum_{i=1}^n p_i = 1$ to learn W , they obtain

$$p_i = \frac{[u(M + aR_i, \theta) - u(M - a, \theta)]^{-1}}{\sum_{j=1}^n [u(M + aR_j, \theta) - u(M - a, \theta)]^{-1}}.$$

Assume that horses are numbered across races, so that horse 1 always wins. Then indexing all races by $k = 1, \dots, N$, Jullien and Salanié obtain the log-likelihood for the sample as

$$\begin{aligned} L^N(\theta) &= \sum_{k=1}^N \log p_1(R^N, \theta) \\ &\stackrel{EU}{=} \sum_{k=1}^N \log \frac{[u(M + aR_{1k}, \theta) - u(M - a, \theta)]^{-1}}{\sum_{j=1}^n [u(M + aR_j, \theta) - u(M - a, \theta)]^{-1}}, \end{aligned}$$

where R_{1k} are the odds of the horse that won race k , and p_1 is the objective win probability for the horse implied by the specific EU model. Intuitively, this estimation approach relies on the fact that, in equilibrium, there is a one-to-one mapping from risk preference parameters

⁵⁶Jullien and Salanié formalize risk preferences through a functional $V(F)$, such that if F_1, F_2 are two risky distributions, bettors will prefer F_2 to F_1 if and only if $V(F_2) \geq V(F_1)$. They show that if $V(\cdot)$ is a continuous functional in F and satisfies first order stochastic dominance, the set of equalities among the values of betting on each horse uniquely define p_i as a function of all R_i 's and θ .

to the win probability of a horse, given the odds on the horse. Hence, maximum likelihood estimation will find the value of θ that makes the model implied win probabilities as close as possible to the empirical win frequencies observed in the data.

Estimation of the EU model with a CARA utility function (for which clearly only θa is identified) yields an estimate of $\hat{\theta}a = -0.055$, thereby indicating that bettors are risk loving (a result which was not "obvious" ex ante, because the authors do not model the decision to bet). Similar results are obtained with a HARA utility function.

Jullien and Salanié next estimate various RDEU models. With RDEU, the value of betting a on horse i is

$$\pi(p_i, \theta)u(M + aR_i, \theta) + (1 - \pi(p_i, \theta))u(M - a, \theta),$$

and must be equal across all horses in equilibrium. Hence one can solve for $\pi(p_i, \theta)$, and using again the adding up constraint $\sum_{i=1}^n p_i = 1$, obtain an implicit formula for p_i as a function of R and θ based on the inverse function of $\pi(\cdot, \theta)$. Specific functional forms for $\pi(\cdot, \theta)$ considered in the paper include those proposed by Lattimore, Baker, and Witte (1992) and Prelec (1998). The estimation results indicate that the Lattimore, Baker, and Witte (1992) model do not fit the data better (in the sense of value of the likelihood function) than the EU model, that the estimate $\hat{\theta}a$ is essentially unaffected, and that the estimated coefficient for the power function is very close to one. Estimation of the Prelec (1998) model yields a better fit to the data and suggests rejection of EU in favor of RDEU. However, the estimated weighting function is quite close to the diagonal.

Finally, Jullien and Salanié estimate CPT models. With CPT, using M as a reference point and assuming that $u(0) = 0$ with $u(\cdot)$ a continuous increasing function, the value of betting a on horse i is

$$\pi^+(p_i, \theta)u(aR_i, \theta) + \pi^-(1 - p_i, \theta)u(-a, \theta),$$

and again must be equal across all horses in equilibrium. The authors then solve numerically for p_i , $i = 1, \dots, n$, using the same logic as in the previous cases. In the estimation of the model, the π^+ and π^- functions take the parametric forms proposed by Lattimore, Baker, and Witte (1992). Power functions are considered as well. In each case, the results show that the estimate $\hat{\theta}a$ is not far from what is obtained for the EU model, that the probability weighting function for gains is slightly convex but not significantly so, but that the probability weighting function for losses is highly and significantly concave, leading to a clear rejection of EU.

Snowberg and Wolfers (2010) revisit the favorite-long shot bias to investigate whether it is driven by risk love (increasing marginal utility) or by risk misperceptions (probability distortions). The fundamental insight used for identification is that while both an EU model and an RDEU model can be parametrized to explain bettors' choices in one part of their choice set (betting on which horse wins the race), each of the models so parametrized will have a different implication for the bettors' decisions over a wider choice set, including compound bets in the exacta, quinella, or trifecta pools.⁵⁷ Snowberg and Wolfers' data is ideally suited to pursue this identification strategy, because their data include information on every horse race run in North America from 1992 to 2001, including for each race the finishing position of each horse that started in the race, the win odds on each horse, and the payoffs for win bets and compound bets (exactas, quinellas, and trifectas). Because only data on equilibrium prices (odds, denoted O_i) is available, the authors resort to a representative agent model.

The analysis proceeds in three steps. First, Snowberg and Wolfers' use the race results to estimate the objective win probability p_i for each horse, and then treat these estimates as data, along with the win odds on each horse which are observed. Although the authors do not discuss this step detail, it appears that they follow the approach of Griffith, McGlothlin, and Weitzman and group horses according to their win odds.

Next, Snowberg and Wolfers separately fit the EU model and the risk misperceptions model to the win data via nonparametric smoothing,⁵⁸ leveraging the fact that in equilibrium bettors must be indifferent between betting on each horse to win. For the EU model, the authors assume that bettors have perfect knowledge of the win odds on each horse, and that each bettor bets on one horse only. They normalize utility to zero if the bet is lost and to one if the bettor chooses not to bet. Then in equilibrium

$$p_i u(O_i) = 1 \Rightarrow p_i = \frac{1}{u(O_i)}, \quad i = 1, \dots, n. \quad (10)$$

The results indicate that "a risk-loving utility function is required to rationalize the bettor accepting lower average returns on long shots, even as they are riskier bets" (p. 729). The results also indicate that a CARA utility function fits the data reasonably well. For the risk misperception model, the authors assume that bettors have a linear utility function, and

⁵⁷An exacta is a bet on both which horse will come first and which will come second. A quinella is a bet on two horses to come first and second in either order. A trifecta is a bet on which horse will come in first, which second, and which third.

⁵⁸In the EU model, the authors can nonparametrically estimate the utility function because they assume that all bettors have the same wealth, in addition to the other assumptions specified below.

they normalize utility to zero if the bettor chooses not to bet. Then in equilibrium

$$\pi(p_i)O_i + (1 - \pi(p_i))(-1) = 0 \Rightarrow \pi(p_i) = \frac{1}{O_i + 1}. \quad (11)$$

The results indicate a risk misperception function π that maps tiny win probabilities into moderate win probabilities, "similarly to some of the features of the decision weights in prospect theory (Kahneman and Tversky (1979))" (p. 731). The results also indicate that a one-parameter probability weighting function as in Prelec (1998) fits the data reasonably well.

Without restrictive parametric assumptions, the EU model and the risk misperceptions model are each just-identified, and so each yields identical predictions for the pricing of win bets. In the third step, therefore, Snowberg and Wolfers turn to compound bets in order to compare the EU and risk misperceptions models. For an exacta bet, in which the bettor wagers that horse A will come first and horse B will come second, the EU model (via equation 10) yields

$$p_A p_{B|A} u(E_{AB}) = 1 \Rightarrow E_{AB} = u^{-1}(u(O_A)u(O_{B|A})),$$

where E_{AB} denotes the odds of horse A winning and horse B being second, and $p_{B|A}$ denotes the probability that B is second given that A wins. On the other hand, the risk misperceptions model (via equation 11) yields

$$\pi(p_A)\pi(p_{B|A})(E_{AB} + 1) = 1 \Rightarrow E_{AB} = (O_A + 1)(O_{B|A} + 1) - 1.$$

The authors use data on odds O_A and (estimated) odds $O_{B|A}$, to infer implied odds E_{AB} for each model (in the EU model, the utility function needs to be estimated from data on win bets). The odds $O_{B|A}$ are not directly observable, and therefore they are inferred via $p_{B|A}$, assuming that bettors believe in conditional independence (i.e., the race for second is considered as a "race within the race"), using the formulae:⁵⁹

$$\left\{ \begin{array}{ll} p_{B|A} = \frac{p_B}{1-p_A} & \Rightarrow O_{B|A} = u^{-1}\left(\frac{1-p_A}{p_B}\right) \quad \text{for the EU model;} \\ \pi(p_{B|A}) = \frac{\pi(p_B)}{1-\pi(p_A)} & \Rightarrow O_{B|A} = \frac{1-\pi(p_A)}{\pi(p_B)} - 1 \quad \text{for the risk misperceptions model.} \end{array} \right.$$

They then compare models, by computing the mean absolute prediction error between the

⁵⁹Snowberg and Wolfers discuss the possibility that instead of assessing the likelihood first that horse A wins, and then that horse B comes in second, bettors might directly assess the likelihood of the first and second combination, hence considering $\pi(p_A p_{B|A})$ and therefore reducing compound lotteries into their equivalent simple lotteries form. However, this would yield a pricing rule for the misperception class equivalent to that of the EU class.

prices of exacta bets observed in the data, and the prices implied by the values of E_{AB} predicted by each model (via the formula: price = $1/(\text{odds} + 1)$). They find that the risk misperceptions model is six percentage points closer to the actual prices of exactas. While this is not a formal statistical test, the authors conclude that their "results are more consistent with the favorite-long shot bias being driven by misperceptions rather than by risk love" (p. 744). Robustness of the results to the assumption of conditional independence is checked by using nonparametric estimates of $p_{A|B}$, obtained using data on races in which a horse with odds O_B came in second, when a horse with odds O_A won the race.

5.2.2 Heterogeneity in Beliefs and Preferences

Two recent papers (Gandhi and Serrano-Padial (2014) and Chiappori, Salanié, Salanié, and Gandhi (2012)) investigate the extent to which one can use aggregate data on horse races to estimate heterogeneity in beliefs and heterogeneity in preferences. The key idea in both is that, instead of assuming a representative bettor who is indifferent between betting on all horses in a race, it is assumed that heterogeneous bettors partition themselves across the horses in a race based on either their beliefs or their preferences. If one can observe how bettors partition themselves across races with different characteristics, one can draw inference on the distribution of heterogeneity.

Gandhi and Serrano-Padial (2014) are motivated by the favorite long-shot bias, and the idea that belief heterogeneity could drive it (as suggested by (Ali 1977, Theorem 2)). In the first part of their paper, they build a formal model of risk-neutral agents with heterogeneous beliefs—specifically, beliefs that are on average correct but with heterogeneity around that average—and prove that such a model would generate a favorite long-shot bias. Here, we focus on the latter part of the paper, where they estimate the extent of belief heterogeneity.

The basic structure of their model is as follows. As above, a race consists of a house take τ and a vector of probabilities $\mathbf{p} \equiv (p_1, \dots, p_n) \gg 0$, where p_i is the probability that horse i wins the race. Bettor t holds beliefs $\boldsymbol{\pi}_t \equiv (\pi_{1t}, \dots, \pi_{nt})$, where beliefs are heterogeneous and thus some bettors have $\boldsymbol{\pi}_t \neq \mathbf{p}$. To put some structure on the nature of belief heterogeneity, Gandhi and Serrano-Padial (2014) define $\nu_{it} = \log(p_1/\pi_{1t}) - \log(p_t/\pi_{it})$, and focus on the population distribution of $\boldsymbol{\nu}_t = (\nu_{1t}, \dots, \nu_{nt})$, which they denote by $P(\boldsymbol{\nu}_t; \theta)$ with θ a vector of parameters to be estimated. They assume that $P(\boldsymbol{\nu}_t; \theta)$ is continuous, and in the estimation they further assume that belief heterogeneity is “idiosyncratic” in the sense that $E(\boldsymbol{\nu}_t) = 0$ and $E(\nu_{it}|\nu_{-it}) = 0$ for any ν_{-it} .

Given this structure, bettors partition themselves across horses to generate equilibrium market shares $\mathbf{s} = (s_1, \dots, s_n)$ and resulting odds $\mathbf{R} = (R_1, \dots, R_n)$. Assuming that bettors

are risk-neutral, in equilibrium each bettor t will bet on the horse with the largest $\pi_{it}R_i$, and the result of all bettors behaving in this way must yield the equilibrium shares $\mathbf{s}$.⁶⁰

Gandhi and Serrano-Padial (2014) have the key insight that, again with risk-neutral bettors, this market is isomorphic to a discrete-choice horizontally differentiated products market of the form studied by Berry, Levinsohn, and Pakes (1995) and Berry, Gandhi, and Haile (2013) (see the paper for details).⁶¹ Applying results from that literature, for any continuous $P(\boldsymbol{\nu}_t; \theta)$, one can translate any observed odds $\mathbf{R}$ into a unique vector of model-implied underlying probabilities $\mathbf{p}^*(\theta)$. One can then set up a likelihood function exactly as in Jullien and Salanié (2000) — numbering horses such that the winning horse is horse 1, and denoting races by $k = 1, \dots, N$, the log-likelihood for the sample is

$$L^N(\theta) = \sum_{k=1}^N \log p_1^*(\theta).$$

The authors estimate this model using a sample of more than 176,000 pari-mutuel races that were collected from North-American tracks over 2003-2006, consisting of more than 1,400,000 horse starts.⁶² With risk-neutral bettors, the only object to estimate is the distribution of belief heterogeneity $P(\boldsymbol{\nu}_t; \theta)$. They assume that $\boldsymbol{\nu}_t$ is distributed according to a variance mixture of logit errors. Specifically, they assume

$$P(\nu_{1t}, \dots, \nu_{nt}; \theta) = \sum_{m=1}^M \left(\prod_{i=2}^n F(\nu_{it} | \sigma_m) \right) g(\sigma_m)$$

where F is a standard logistic distribution, and $g(\sigma_m)$ is the mixing probability. The authors report results for $M = 1, 2$, and 3, where they interpret M to be the number of types in the population and σ_m to be a measure of belief heterogeneity for type m . A likelihood ratio test rejects the one-type model in favor of the two-type model, while adding a third type does not significantly improve the log-likelihood. The authors therefore focus on the two-type model. In the two-type model with risk-neutral bettors, they estimate that 70% of the population have a small $\sigma = .028$, while 30% have a large $\sigma = .503$. The authors interpret the former group as having roughly correct beliefs (informed traders), and they interpret the latter group as having dispersed beliefs (noise traders).⁶³

⁶⁰The authors effectively assume that all bettors bet the same amount, because while they permit heterogeneity in the amount bet, they assume this heterogeneity to be independent of the belief heterogeneity, and thus the average amount bet by agents of each belief-heterogeneity type is equal to the average amount bet by the population.

⁶¹In an appendix, they describe how this approach can be extended to permit risk-averse bettors.

⁶²As a preliminary step, the authors regress ex-post return on (log) market shares, and indeed find evidence of the favorite-long shot bias in their data.

⁶³They also estimate a model with risk-averse bettors who all have the same CARA utility function, and

Finally, the authors do some additional analysis to further buttress their argument that belief heterogeneity is playing an important role in this domain. First, they use their data to estimate a representative agent model using the preferred specification of Jullien and Salanié (2000), and they conclude using a Vuong test that their belief heterogeneity model is statistically better than the representative agent model. Second, they separately estimate their model on the subsamples of maiden and non-maiden races. Maiden races are races in which participating horses have yet to win a single race, and thus there is less handicapping information about horses in maiden races than in non-maiden races. Their estimates suggest stable proportions of informed vs. noise traders across the two types of races, but lower variance estimates for each type of trader in the non-maiden races. This result is exactly what one might expect if belief heterogeneity were playing an important role in this environment.

In an alternative approach, Chiappori, Salanié, Salanié, and Gandhi (2012) investigate how one can use aggregate data on horse races to estimate a model with heterogeneity in preferences. In particular, they demonstrate that preference heterogeneity can be identified provided some key restrictions are satisfied: (i) heterogeneity is one dimensional and satisfies the single-crossing property described below; (ii) utility from betting on one horse is independent from the utility from betting on a different one (hence not allowing for regret theory); (iii) agents make their decisions based on the correct win probabilities $\mathbf{p}$ (hence not allowing for subjective beliefs).

Chiappori, Salanié, Salanié, and Gandhi (2012) assume that there is a continuum of bettors with utility function $V(R, p, \theta)$, where θ is a heterogeneity parameter normalized to be uniformly distributed on $[0, 1]$. The authors prove that, under a set of regularity conditions (among which a single crossing property explained below plays a particularly important role), the utility function $V(R, p, \theta)$ is identified. Estimation and inference, however, are conducted based on a feature of V which they label the “normalized fear of ruin”:

$$NF(R, p, \theta) = \frac{p}{R+1} \frac{V_p}{V_R}(R, p, \theta).$$

Because different models of risk preferences have different implications for $NF(R, p, \theta)$, the authors are able to use the estimated NF to test which, among different models of risk preferences, better fits the data.

The utility function $V(R, p, \theta)$ is assumed to be continuously differentiable almost everywhere, increasing in R and p , and such that any higher return can be compensated by a lower probability of winning and vice versa. The crucial assumption is that V satisfies a single-crossing property: For any two gambles (R, p) and (R', p') with $p' < p$, if, for some

reach much the same conclusion about the nature of belief heterogeneity.

θ , $V(R, p, \theta) \leq V(R', p', \theta)$, then for all $\theta' > \theta$, $V(R, p, \theta') < V(R', p', \theta')$. Hence, θ can be interpreted as a taste for risk, where the larger θ is, the more a bettor is prone to prefer horses with a lower probability but higher expected return. This single-crossing assumption is satisfied for an EU model with CARA utility or with CRRA utility. For an RDEU model, it can be satisfied, though additional restrictions are required. A particularly important restriction under RDEU is that if there is heterogeneity in both the utility function and the probability-weighting function, the two types of heterogeneity must have a structure that can be reduced to a single dimension of heterogeneity that satisfies the single-crossing condition.

To characterize the equilibrium given these preferences, Chiappori, Salanié, Salanié, and Gandhi (2012) label horses in each race such that $p_1 > p_2 > \dots > p_n$. If preferences satisfy the single-crossing condition, then for any vector of odds $\mathbf{R}$ satisfying $R_1 < R_2 < \dots < R_n$ we can define, for each $i \in \{1, \dots, n-1\}$, $\theta_i(\mathbf{R})$ to be the unique θ' such that $V(R_i, p_i, \theta') = V(R_{i+1}, p_{i+1}, \theta')$. For any race $(\mathbf{p}, \tau)$, the market shares $\mathbf{s}$ and resulting odds $\mathbf{R}$ will adjust, such that

$$\begin{aligned} \theta_i(\mathbf{R}) &< \theta_{i+1}(\mathbf{R}) \text{ for all } i \in \{1, \dots, n-1\} \\ s_1 &= \theta_1(\mathbf{R}) \\ \text{and } s_i &= \theta_{i+1}(\mathbf{R}) - \theta_i(\mathbf{R}) \text{ for all } i \in \{2, \dots, n\}. \end{aligned}$$

In other words, bettors partition themselves according to their θ 's, where the bettors with the lowest θ 's, who are most averse to risk, bet on the horse with the largest probability of winning, and bettors with larger and larger θ 's bet on horses with smaller and smaller probabilities of winning. The authors prove that for any race $(\mathbf{p}, \tau)$ there exists a unique equilibrium $(\mathbf{s}, \mathbf{R})$, and moreover for any observed equilibrium $(\mathbf{s}, \mathbf{R})$ there is a unique race $(\mathbf{p}, \tau)$ that can generate those odds. Hence, one can define a function $\mathbf{p}(\mathbf{R}) \equiv (p_1(\mathbf{R}), \dots, p_n(\mathbf{R}))$ that reflects the underlying vector of probabilities associated with any observed odds $\mathbf{R}$.

The intuition behind the authors' approach to identification is that if one observes, for each θ , enough indifference conditions of the form $(R, p) \sim (R', p')$, then one can make inferences about $V(R, p, \theta)$ —in particular, about $NF(R, p, \theta)$. Specifically, they define $G(R, p, R', \theta)$ to be the p' such that a bettor of type θ has $(R, p) \sim (R', p')$. For any observed race, the equilibrium conditions imply

$$\forall i < n, \quad p_{i+1}(\mathbf{R}) = G(R_i, p_i(\mathbf{R}), R_{i+1}, \theta_i(\mathbf{R})).$$

With this structure in place, the authors show that knowledge of $G(R_i, p_i(\mathbf{R}), R_{i+1}, \theta_i(\mathbf{R}))$ (which can be acquired from a regression of $p_{i+1}(\mathbf{R})$ on $R_i, p_i(\mathbf{R}), R_{i+1}, \theta_i(\mathbf{R})$), yields knowl-

edge of

$$NF(R, p, \theta) = \frac{p}{R+1} \frac{G_p}{G_R}(R, p, R', \theta).$$

The authors estimate their model—that is, the function G —using a dataset of more than 53,000 thoroughbred races in the United States from 2001 to 2004.⁶⁴ For each race, the authors observe the odds $\mathbf{R}$, from which they can immediately derive the market shares $\mathbf{s}$ and thus the cutoff values $\theta_i(\mathbf{R})$, $i = 1, \dots, n$. To estimate G , they follow a two-step procedure. First, they use empirically observed outcomes to estimate the function $p_i(\mathbf{R})$.⁶⁵ In this step, they use a semi-nonparametric sieve method with orthogonal polynomials, where the number of terms to be included is selected using the Akaike Information Criterion and the Bayesian Information Criterion. Second, with the function $p_i(\mathbf{R})$ in hand (in this step, they use the BIC selected model), they estimate the function $G(R_i, p_i(\mathbf{R}), R_{i+1}, \theta_i(\mathbf{R}))$ via nonparametric regression, performed using Generalized Additive Models.

Finally, the authors investigate the implications of the estimated $\hat{G}$ function. In particular, they use $\hat{G}$ to compute the NF index at each observation in the data. They then test how well different models of risk preferences can explain the variation in this estimated NF index. Of the models considered, a heterogenous non-expected utility model with non-additive probability weights, in which the estimated probability weighting function is concave and then convex, performs best.

5.3 Consumption, Asset Returns, and Labor Supply Data

In addition to the studies that use betting data, there are a number of studies that use macroeconomic data to estimate risk preferences.

The bulk of these studies use data on consumption and investment (asset returns). A seminal paper is Hansen and Singleton (1983). They study a single-good economy of identical, infinitely-lived agents with time-additive EU preferences and CRRA utility. The representative agent in this economy chooses a stochastic consumption plan $\{c_t\}_{t=0}^{\infty}$ to maximize

$$E_0 \left[\sum_{t=0}^{\infty} \beta^t u(c_t) \right], \quad (12)$$

where E_0 is the expectation conditional on information available in period zero, $\beta > 0$ is a

⁶⁴Observable heterogeneity across races is dealt with by conditioning on observable covariates X . In the analysis, these are discrete and incorporated via cell mean estimation. In particular, the results reported are for the subsample of weekday races run on urban tracks.

⁶⁵This step is analogous to the approach in Weitzman (1965) and Snowberg and Wolfers (2010), except that they estimate how the probability of horse i winning depends on the full vector of odds $\mathbf{R}$ as opposed to merely depending on R_i .

discount factor, and $u(c_t) = (c_t^{1-\rho} - 1)/(1 - \rho)$, $\rho > 0$, subject to the sequence of budget constraints

$$c_t + \mathbf{p}'_t \mathbf{x}_{t+1} \leq (\mathbf{p}_t + \mathbf{d}_t)' \mathbf{x}_t, \quad (13)$$

where $\mathbf{x}_t$ is a vector of the agent's holdings of n assets in period t , $\mathbf{p}_t$ is the vector of prices of the n assets in $\mathbf{x}_t$ net of any distributions, and $\mathbf{d}_t$ is the vector of distributions in period t .⁶⁶

The first-order necessary conditions for the maximization of (12) subject to (13) are given by the Euler equations,

$$E_t \left[\beta \left(\frac{c_{t+1}}{c_t} \right)^{-\rho} R_{it+1} \right] = 1, \quad i = 1, \dots, n, \quad (14)$$

where E_t is the expectation conditional on information available in period t and $R_{it+1} = (p_{it+1} + d_{it+1})/p_{it}$ is the one-period return on asset i .

By assuming, *inter alia*, that the joint distribution of consumption and asset returns is lognormal, Hansen and Singleton are able to obtain maximum likelihood (ML) estimates of ρ and β using monthly U.S. data for the period February 1959 through December 1978. In general, their estimates of the coefficient of relative risk aversion ρ range between zero and two, and their estimates of the discount factor β are less than but close to unity. Perhaps more importantly, however, they perform various chi-square and likelihood ratio tests that provide substantial evidence against the model. Of course, these tests are joint tests of the model's several restrictions, including the preference assumptions (identical agents, EU preferences, time additivity, CRRA utility, exponential discounting, etc.) and the distributional assumption (jointly lognormally distributed consumption and asset returns), making it impossible to say which restrictions are being rejected. For this reason, the authors lay out a dual plan of pursuing both "models . . . with more general specifications of preferences and distribution-free methods of estimation and inference" (Hansen and Singleton 1983, p. 264).

In Hansen and Singleton (1982), the authors progress their plan to pursue distribution-free methods of estimation and inference.⁶⁷ They develop a generalized instrumental variables procedure for estimating the parameters of model (12)-(13) and implement their procedure using the same monthly data on consumption and asset returns used in Hansen and Singleton (1983). In brief, they first use the Euler equations (14) to generate a set of

⁶⁶In Hansen and Singleton (1983), the right-hand side of the budget constraint also includes a term explicitly measuring the agent's labor income in period t . This term can be suppressed, however, without loss of generality (Epstein and Zin 1989, 1991, p. 267).

⁶⁷Hansen and Singleton (1982) was published before but apparently written after Hansen and Singleton (1983).

population orthogonality conditions,

$$E \left[\begin{pmatrix} \beta \left(\frac{c_{t+1}}{c_t} \right)^{-\rho} R_{1t+1} - 1 \\ \vdots \\ \beta \left(\frac{c_{t+1}}{c_t} \right)^{-\rho} R_{nt+1} - 1 \end{pmatrix} \otimes \mathbf{z}_t \right] = \mathbf{0}, \quad (15)$$

where E is unconditional expectation and the vector of instruments $\mathbf{z}_t$ comprises lagged values of $(R_{1t+1}, \dots, R_{nt+1}, c_{t+1}/c_t)'$. They then use the sample analog of the orthogonality conditions to construct a generalized method of moments (GMM) estimator of the model parameters. The GMM estimates of ρ and β are similar to the ML estimates reported in Hansen and Singleton (1983). What is more, chi-square tests again provide evidence against the model, which here may be interpreted as direct evidence against the model's preference assumptions.

The rejection by Hansen and Singleton (1982, 1983) of the standard preference assumptions is echoed in the subsequent literature on the equity premium puzzle (Mehra and Prescott 1985).⁶⁸ In their famous paper, Mehra and Prescott consider a variation of the standard model studied by Hansen and Singleton in which there are two assets: a risky equity security and a risk-free debt security. Unlike Hansen and Singleton, however, Mehra and Prescott do not estimate the model's preference parameters. Instead, they calibrate the model's non-preference parameters so that the population mean, variance, and autocorrelation of consumption growth match their corresponding sample values for the U.S. economy between 1889 and 1978. They then derive analytical expressions for the expected returns on the equity and risk-free securities, R^e and R^f , in terms of the model's preference parameters, ρ and β . Restricting $\rho \leq 10$, $\beta \leq 1$, and $R^f \leq 0.04$ (more than four times the average return on the three month Treasury bill during the sample period), they find that the maximum equity premium ($R^e - R^f$) that is consistent with the model is less than one percent, whereas the historical equity premium (the difference between the average annual returns on the S&P 500 and the three month Treasury bill during the sample period) at the time was more than six percent.

Motivated in part by the poor empirical performance of the standard model in Hansen and Singleton (1983) and Mehra and Prescott (1985), and in part by concurrent developments in the microeconomics literature on non-EU preferences, Epstein and Zin (1989, 1991) take the path of pursuing a model with a more general specification of preferences.⁶⁹ The key

⁶⁸For a survey of the literature on the equity premium puzzle, see Kocherlakota (1996). See also Siegel and Thaler (1997). Campbell (2003) summarizes the larger literature on asset price puzzles in macroeconomics.

⁶⁹Weil (1989, 1990) takes a parallel path.

innovation of their model is that it disentangles the representative agent's degree of risk aversion from his elasticity of intertemporal substitution. In the standard model, the coefficient of relative risk aversion is constrained to equal the reciprocal of the elasticity of intertemporal substitution. As a result, an agent with standard preferences that is very averse to risk must also be very averse to consumption growth. However, this is what makes the equity premium a puzzle: the high degree of risk aversion that is required to explain the historical equity premium implies an implausibly low elasticity of intertemporal substitution given the historical rate of consumption growth.⁷⁰ Accordingly, Epstein and Zin investigate whether a model that delinks preferences over risk and intertemporal substitution can outperform the standard model.

In Epstein and Zin's model, the representative agent's utility in period t is defined recursively as

$$U_t = \left[(1 - \beta) c_t^\varsigma + \beta (E_t [U_{t+1}]^{1-\rho})^{\frac{\varsigma}{1-\rho}} \right]^{\frac{1}{\varsigma}}, \quad (16)$$

where $\varsigma = 1 - (1/\sigma)$ and σ is the agent's elasticity of intertemporal substitution, and his intertemporal budget constraint is given by (13). The Euler equations that characterize the first-order necessary conditions of the optimal consumption path are

$$E_t \left[\beta \left(\frac{c_{t+1}}{c_t} \right)^{\theta(\varsigma-1)} (m_{t+1})^{\theta-1} R_{it+1} \right] = 1, \quad i = 1, \dots, n, \quad (17)$$

where $\theta = (1 - \rho)/\varsigma$ and $m_{t+1} = [(\mathbf{p}_{t+1} + \mathbf{d}_{t+1})' \mathbf{x}_{t+1}] / \mathbf{p}_t' \mathbf{x}_{t+1}$ is the one-period return on the agent's asset holdings in period t . Observe that when $\rho = 1/\sigma$, the Euler equations (17) reduce to (14).

Following the GMM procedure of Hansen and Singleton (1982), Epstein and Zin estimate the model using monthly U.S. data for the period February 1959 through December 1986. Although Wald and likelihood ratio tests of the restriction $\rho = 1/\sigma$ generally reject the standard model, Epstein and Zin's model does not resolve the equity premium puzzle. In general, they find that the coefficient of relative risk aversion ρ is close to one and the elasticity of intertemporal substitution σ is less than one. Moreover, they find that the discount factor β is often greater than one.

Another generalization of the standard EU model which features prominently in the macro-finance literature is the model with habit formation (e.g., Abel 1990, 1999; Constantinides 1990; Campbell and Cochrane 1999). Habit formation models posit that marginal

⁷⁰ Stated differently, the equity premium puzzle may be viewed as a risk-free rate puzzle (Weil 1989): the high degree of risk aversion that is required to explain the historical equity premium implies a low elasticity of intertemporal substitution, but a low elasticity of intertemporal substitution necessitates a counterfactually high risk-free rate to induce savings sufficient to explain the historical rate of consumption growth.

utility of consumption in period t is an increasing function of consumption in period $t - 1$. Consider, as in Campbell and Cochrane (1999), the following modification to the standard maximization problem in (12):

$$E_0 \sum_{t=0}^{\infty} \beta^t \frac{(c_t - x_t)^{1-\rho} - 1}{1 - \rho},$$

where x_t is the level of habit. Defining $S_t \equiv (c_t - x_t)/c_t$ as the surplus consumption ratio, it is immediate to see that positive habit, x_t , breaks the one-to-one link between the degree of relative risk aversion and the parameter ρ :

$$-\frac{c_t u_{cc}(c_t, x_t)}{u_c(c_t, x_t)} = \frac{\rho}{S_t}.$$

In particular, for a given ρ , the lower is the surplus consumption ratio (i.e., the lower is consumption c_t relative to habit x_t), the higher is the agent's relative risk aversion. Moreover, depending on the specification of x_t the coefficient of relative risk aversion will be time varying. The implications for asset pricing can be seen from the corresponding inter-temporal Euler equation:⁷¹

$$E_t \left[\beta \left(\frac{S_{t+1}}{S_t} \right)^{-\rho} \left(\frac{c_{t+1}}{c_t} \right)^{-\rho} R_{it+1} \right] = 1.$$

Without habit (i.e., $x_t = 0$), this collapses to the standard Euler equation in (14). However, with positive habit, asset returns become untangled from consumption growth. Naturally, the success of models with habit formation rests entirely on a convincing specification of a process for S_t . Campbell and Cochrane postulate the following heteroskedastic AR(1) process for the (log) surplus consumption ratio:

$$\log S_{t+1} = (1 - \phi)\bar{s} + \phi \log S_t + \lambda(\log S_t)(\log c_{t+1} - \log c_t - g),$$

where ϕ and $\bar{s}$ are parameters, g is the average consumption growth, and $\lambda(\cdot)$ is a decreasing convex function over a finite support chosen such that the risk free rate is constant. They show that with a realistically low parameter ρ (and other calibrated features of the model),

⁷¹There are two ways to specify habit formation: as either internal or external habit formation. With internal habit formation, habit depends on an agent's own consumption and the agent takes into account the fact that increasing consumption today will increase the habit level, and thus the marginal utility of consumption, in future periods. With external habit formation, habit depends on aggregate consumption (which is unaffected by any one agent's behavior), and can be viewed as formalizing the notion of "keeping up with Jones." The Euler equation here is the one derived assuming external habit formation.

the data generated by the model exhibits short- and long-run equity returns that come close to their empirical counterparts.

Another type of macroeconomic data that has been used to estimate risk preferences is data on labor supply (Chetty 2006). Chetty's basic insight is that the wage elasticity of labor supply, which has been estimated in numerous studies in labor economics, provides information about the curvature of the marginal utility of consumption. To see the intuition behind Chetty's insight, note that in standard models with endogenous labor supply and competitive labor markets, the first order condition with respect to labor is given by

$$U_l(c, l) = -wU_c(c, l),$$

where c is consumption, l is labor supplied, and w is the wage. For illustrative purposes, assume that the environment is static and that agents' preferences are separable in consumption and labor:

$$U(c, l) = \frac{c^{1-\rho}}{1-\rho} - \Lambda \frac{l^{1+\psi}}{1+\psi},$$

where Λ and ψ are positive constants that govern the disutility from labor. Finally, assume that agents do not have any non-labor income: $c = wl$. It follows that the first-order condition with respect to labor becomes

$$wc^{-\rho} = \Lambda l^{\psi} \Leftrightarrow w^{1-\rho} = \Lambda l^{\psi+\rho} \Leftrightarrow (1-\rho) \log w = \log \Lambda + (\psi + \rho) \log l.$$

It then follows that

$$\frac{d \log l}{d \log w} = \frac{1-\rho}{\psi + \rho} \Rightarrow \text{if } \frac{d \log l}{d \log w} > 0 \text{ then } \rho < 1.$$

In words, if the wage elasticity of labor supply is positive, then the curvature of the marginal utility of consumption cannot be too high. Indeed, *ceteris paribus*, a high curvature implies that the marginal utility of consumption diminishes quickly, and hence faced with increased wages, agents would be willing to forgo consumption to have more leisure—that is, they will work less.

Chetty develops this insight in a dynamic life cycle model with a generic utility function. He uses the estimates of the degree of complementarity between consumption and labor, as well as the estimates of wage elasticity of labor supply, obtained in various other studies, to conclude, through a calibration exercise, that the curvature of the marginal utility of consumption has to be modest, corresponding to a value of the coefficient of relative risk aversion around one.

Finally, aggregate consumption data has been coupled with individual consumption data and the theory of optimal risk sharing to shed light on a number of interesting implications of the EU model. A notable example of this type of research is Chiappori, Samphatharak, Schulhofer-Wohl, and Townsend (2014), who study risk sharing in village economies in Thailand. The goal of the paper is to assess to which degree (if any) households in rural areas share risk and to assess the degree of heterogeneity in their risk preferences. Based on a test that relies on measuring (conditional) correlation between households' consumption and their respective incomes, the authors cannot reject the null hypothesis of full risk sharing. Full risk sharing allows for measuring preference heterogeneity via consumption correlation patterns across households. Consider, for example, a village where all households have identical risk preferences. Then, each household's share of aggregate village consumption is constant over time. However, if some households are more risk averse than others, optimal risk sharing would imply that their consumption is more volatile and co-moves more closely with aggregate village consumption. In fact, with full insurance, two households will have co-moving consumption only when their consumption also co-moves with aggregate consumption. Similarly, if the consumption levels of a pair of households are not correlated with each other, one of these households must be more risk averse, and its consumption should be less correlated with the aggregate consumption level of the village. The paper formalizes this intuition into a mechanism that identifies relatively more or less risk averse households, by looking at the pairwise correlations in their consumption.⁷² The authors find substantial heterogeneity in risk preferences. However, as they make clear, the analysis relying on the second moment properties of consumption does not speak directly to the scale of the (average) risk aversion across the population in question.⁷³

6 Directions for Future Research

Our analysis of identification and estimation of risk preferences highlights specific assumptions that the literature has imposed to date. These can be divided into three broad categories: (1) assumptions about the possible sources of aversion to risk; (2) assumptions about subjective beliefs; and (3) assumptions about how agents themselves transform—and

⁷²A companion paper, Chiappori, Samphatharak, Schulhofer-Wohl, and Townsend (2013), shows that the measures of risk aversion found in this study correlate with those that can be inferred from households' portfolio choices.

⁷³The inability to pin down the exact scale of risk preferences can also be found in Chiappori and Paiella (2011), who study how the riskiness of households' financial portfolios change with changes in their total financial wealth. The authors find a negligible response of households' portfolio composition to wealth shocks, which points to constant relative risk aversion. They also find that there is a negative correlation between wealth and risk, though the latter can be measured only up to scale.

simplify—a choice situation in their own minds before making a decision. In this section we discuss various ways in which, as the literature moves forward, these assumptions can be refined or weakened.

Within the first category of assumptions, we include not merely functional form restrictions, but more substantially the choice of whether a model of decision making under uncertainty should include curvature of the utility function, and/or elements of probability weighting, cumulative prospect theory, loss aversion, and disappointment aversion. In making this choice, researchers often take lessons from the related literature. A natural question is then to what extent these lessons are "portable." In Section 6.1 we review a very small literature that has studied the extent of stability of risk preferences across contexts. We believe that more research is needed to answer this important question. And to the extent that stability of risk preferences is assumed, we propose that field data in which agents make more than one choice can be leveraged to draw inference on the extent to which different models of aversion to risk are useful in explaining agents' choices. We then suggest in Section 6.2 an additional avenue for future research, which also builds on the assumption of stability: combining data from the laboratory with field data to obtain a better understanding of risk preferences.

Within the second category of assumptions, we refer to the identification problem associated with disentangling risk preferences from beliefs. In Section 6.3 we propose that future research augments choice data with survey data that measure probabilistic expectations directly.

Finally, within the third category of assumptions, we refer to how the typically complex field context is translated into concrete choice data that can be used to estimate risk preferences (see Section 4.1.1). Typically, the resulting assumptions are motivated by the need of the researcher to obtain a tractable model, rather than by their psychological realism. An important agenda for future research is to pay more careful attention to these assumptions, and to investigate directly the impact of such assumptions on estimates of risk preferences.

6.1 Consistency across contexts

A common assumption in economics is that risk preferences are stable across decision contexts (or domains), which implies that multiple risky choices by the same economic unit should reflect the same degree of risk aversion (or risk loving), even if the contexts of the decisions are different. If this assumption is correct, then the estimates of risk preferences derived from choices in one context can be used to understand and make predictions about the behavior of the households in other contexts. Assessing the empirical validity of this assumption is a

difficult task. Moreover, it is quite possible that risk preferences are stable across a certain set of contexts, but not others.

There has been a handful of papers on this issue, which we survey below. More work is needed to fully assess whether and to what extent risk preferences are context-specific versus domain-general.

Wolf and Pohlman (1983) compare the risk preferences of a single dealer in U.S. government securities in two domains, a hypothetical gambling context and an actual investment context. Throughout their study, Wolf and Pohlman assume that the dealer is an expected utility maximizer with the following hyperbolic absolute risk aversion (HARA) utility function,

$$u(w) = \frac{(1-g)}{g} \left[\frac{w}{1-g} + n \right]^g, \quad (18)$$

where $[w/(1-g)] + n > 0$.

In the gambling context, Wolf and Pohlman recover the utility parameters g and n from the dealer's direct assessments of six hypothetical lotteries. In each assessment, the dealer selected a level of wealth that made him indifferent between the lottery $\{w_i, 0.5; w_k, 0.5\}$ and the status quo wealth w_j , where $w_i < w_j < w_k$. Specifically, for each $w_k \in \{1.500, 1.667, 2.000\}$, the dealer selected (i) w_i assuming $w_j = 1$ and (ii) w_j assuming $w_i = 1$. From each pair (w_i, w_j) selected by the dealer, the authors recover the implied values of g and n and then compute the implied coefficients of absolute and relative risk aversion. They report that, in each case, the dealer's coefficient of relative risk aversion was equal to one, and that his coefficient of absolute risk aversion likewise was close to one.

In the investment context, Wolf and Pohlman estimate g and n from the dealer's bid decisions over a series of 28 Treasury bill auctions conducted during a 20 week interval in 1976. According to the authors, the dealer's bid decision can be separated into two decisions. The dealer first determines the optimal bid composition, which determines the random net return per dollar invested, r . He then chooses the optimal bid size, λ , to maximize

$$E[u(w_0(1 + \lambda r))]$$

subject to $\lambda \geq 0$, where w_0 is the initial capital assigned to the auction. With $u(w)$ given by (18), the first-order necessary condition is

$$E \left[w_0 r \left(\frac{w_0(1 + \lambda r)}{1-g} + n \right) \right]^{g-1} = 0. \quad (19)$$

For each auction in the sample, the dealer reported the composition and size of his bid. He

also reported his subjective forecasts (at the time he submitted his bid) of the prices at which Treasury bills could be bought and sold, expressed as discrete probability distributions over a set of stop-out prices of the current auction and the following week's auction.⁷⁴ Using the two forecasts and the bid composition, the authors construct the distribution of r . In addition, they fix $w_0 = 1$ and set λ equal to the bid size per dollar of assigned capital. The authors obtain NLS estimates of g and n by finding the values that minimize, across the 28 auctions in the sample, the sum of the squared residuals between λ and $\hat{\lambda}$, where $\hat{\lambda}$ is defined implicitly by (19). They then compute the implied coefficients of absolute and relative risk aversion. The authors find that both coefficients are four times larger than the dealer's directly assessed coefficients, leading them to conclude that people's "degree of risk aversion may depend on the specific context in which their choices are made" (Wolf and Pohlman 1983, p. 849).

Barseghyan, Prince, and Teitelbaum (2011) investigate the stability hypothesis by examining households' insurance choices. The authors assemble a dataset that records the insurance choices of 702 households in three lines of personal property coverage: auto collision, auto comprehensive, and home all perils.⁷⁵ They use the data to test whether the households' deductible choices across coverage lines reflect the same degree of absolute risk aversion. The test relies on a model of deductible choice that assumes, *inter alia*, that households are expected utility maximizers and that households know their coverage-specific claim rates (i.e., the average rate at which they will experience claims in each line of coverage).

Under the assumptions of the model, Barseghyan, Prince, and Teitelbaum derive an approximation of the coefficient of absolute risk aversion at which a household is indifferent between any two deductible options, $d_l < d_h$:

$$\tilde{r}_{l,h} \approx \frac{\frac{p_l - p_h}{\lambda(d_h - d_l)} - 1}{\frac{1}{2}(d_h - d_l)},$$

where p_l and p_h are the premiums for coverage with deductibles d_l and d_h , respectively, and λ is the household's claim rate. Although the deductible options and associated premiums are observable variables, the household's claim rate is a latent variable. The authors assume that a household's claims are generated by a Poisson process having rate λ , and they use data on claim realizations and demographics to estimate λ .

The authors' test leverages the idea that each deductible choice by a household implies that its coefficient of absolute risk aversion lies between two indifference points. For exam-

⁷⁴The stop-out price of a Treasury bill auction action is the lowest price that the Treasury will accept for bills in the auction.

⁷⁵These are the same contexts as analyzed in BMOT, although for a different dataset. See footnote 47.

ple, if the household chooses a deductible of \$250 from a menu of \$100, \$250, and \$500, then $\tilde{r}_{250,500}$ provides a lower bound on its coefficient of absolute risk aversion and $\tilde{r}_{100,250}$ provides an upper bound. In this fashion, a household's deductible choice in each line of coverage identifies an interval that must contain its coefficient of absolute risk aversion. Because a household makes deductible choices in three lines of coverage—auto collision, auto comprehensive, and home all perils—its choices imply three intervals. The test simply asks whether the intervals intersect. If the intervals intersect, then any coefficient of absolute risk aversion contained in the intersection can rationalize the household's choices. If the intervals do not intersect, however, then the household's choices cannot be rationalized by the same coefficient of absolute risk aversion.

Barseghyan, Prince, and Teitelbaum find that the hypothesis of stable risk preferences is rejected by the data. More specifically, they find that only 23 percent of households "pass" the test—i.e., the three intervals intersect for only 23 percent of households. The authors argue that this is rather low considering that the pass rate would be 14 percent if households were randomly assigned their deductible choices.⁷⁶ As for the other 77 percent, the authors find that these households typically exhibit greater risk aversion in their home deductible choices than they do in their auto deductible choices. In the home domain, the average household would pay \$45 to avoid facing a gamble offering an equal chance of winning and losing \$100. In the auto domain, by comparison, the average household would pay only \$30 to avoid facing the same gamble.

To further illustrate their results, the authors compare the joint distribution of auto collision and home deductibles in the data with the joint distribution generated by the model under the null hypothesis of stable risk preferences—see Figure 5. Conditional on their auto collision deductibles, the model predicts that households would choose "high" home deductibles (\$500 or higher) with significantly greater frequency than is observed in the data. Indeed, the authors report that a Wald test rejects at the 1 percent level the equality of the marginal distribution of the actual home deductibles and the marginal distribution of the model-generated home deductibles. Because lower deductibles correspond to more insurance, the figure illustrates not only the conclusion that the hypothesis of stable risk preferences is rejected by the data, but also the finding that households exhibit greater risk aversion in the home domain than they do in the auto domain.

In their concluding remarks, Barseghyan, Prince, and Teitelbaum discuss several potential

⁷⁶That said, even if every household had stable risk preferences, one should not expect a pass rate of 100 percent. This is because the households' true claim rates are bound to differ from their predicted claim rates due to unobserved heterogeneity. Therefore, as a relevant point of comparison, the authors simulate the expected pass rate under the null hypothesis of stable risk preferences. They find that the expected pass rate is 50 percent, more than twice the actual percentage.

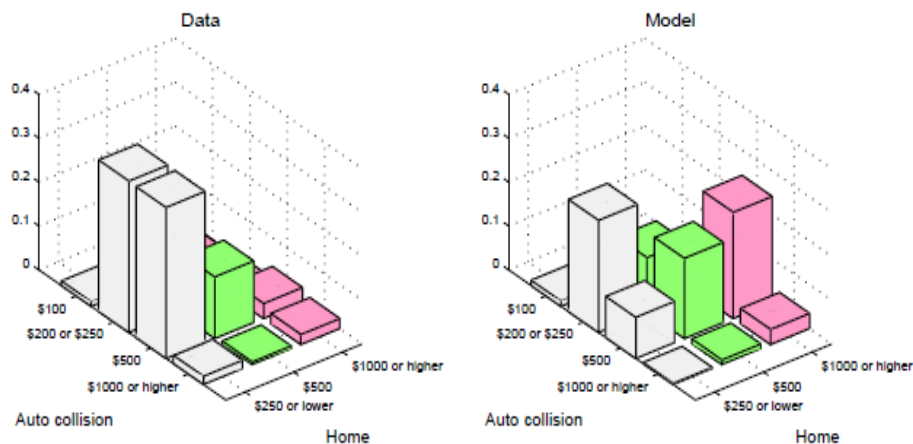


Figure 5: Joint distribution of deductibles – data vs. model

explanations for their results, including the possibility that people have systematic risk misperceptions. They also acknowledge that rejecting the hypothesis of stable risk preference is not equivalent to rejecting the hypothesis that there is no domain-general component to risk preferences. After all, the authors’ test is a joint test of the stability hypothesis and the expected utility hypothesis, and the latter may be what is being rejected by their test.

Einav, Finkelstein, Pascu, and Cullen (2012) analyze whether and to what extent people’s employee benefits choices exhibit systematic patterns, as would be implied by a domain-general component of risk preferences. More specifically, the authors examine the benefits choices of 12,752 Alcoa employees in six domains: health insurance, drug insurance, dental insurance, short-term disability insurance, long-term disability insurance, and 401(k) investments. Using these data, they investigate the stability in ranking across domains of an employee’s willingness to bear risk relative to his or her peers. In contrast to Barseghyan, Prince, and Teitelbaum (2011), who focus on testing the hypothesis of stable risk preferences, Einav, Finkelstein, Pascu, and Cullen focus on testing the hypothesis of nonstable risk preferences and on quantifying the empirical importance of any domain-general component of risk preferences.

Einav, Finkelstein, Pascu, and Cullen take two complementary approaches in their analysis. First, they take a model-free approach in which they rank by risk the options within each domain and compute the pairwise rank correlations in the employees’ choices across domains. Their Table 3A-Panel A, reproduced in Figure 6 for the reader’s convenience, reports their results. They find that an employee’s choice in every domain is positively correlated to some extent with his or her choice in every other domain. Thus, they conclude that they can reject the null hypothesis of zero correlation and hence the null hypothesis of no domain-general

	Health	Drug	Dental	STD	LTD
Drug	0.400				
Dental	0.242	0.275			
STD	0.226	0.210	0.179		
LTD	0.180	0.199	0.173	0.593	
401(k)	0.057	0.061	0.036	0.029	0.028

Note: All correlations are statistically different from zero at the 1 percent level.

Figure 6: Table 3A-Panel A from Einav, Finkelstein, Pascu, and Cullen (2012)

component of risk preferences.

Second, Einav, Finkelstein, Pascu, and Cullen take a model-based approach that is conceptually similar to the test proposed by Barseghyan, Prince, and Teitelbaum (2011). They first specify a model of coverage choice that assumes that employees are expected utility maximizers and that an employee’s utility function has one free parameter that confounds the degree of risk aversion and other domain-specific effects (e.g., beliefs). The parameter is allowed to vary across domains (but not across employees within a given domain). Given this model—and a specific parametric utility function such as CRRA or CARA—each coverage choice of an employee identifies an interval that must contain its utility parameter. Because an employee makes coverage choices in six domains, his or her choices imply six intervals. If the intervals overlap, then the employee’s choices can be rationalized by a single coefficient of risk aversion, subject to other domain-specific effects (that do not vary across employees). In other words, if the intervals overlap, then the employee’s choices reflect a stable ranking (relative to his or her peers) of risk aversion across domains.⁷⁷ The authors then search for the set of utility parameters (one for each domain) that maximize the fraction of employees whose intervals overlap. They find that this maximum fraction is 30 percent.

In summary, on the one hand, Einav, Finkelstein, Pascu, and Cullen find evidence that people’s risky choices are rank correlated across contexts, leading them to reject the null hypothesis of nonstable risk preferences and to conclude that risk preferences have an important domain-general component. On the other hand, like Barseghyan, Prince, and Teitelbaum (2011), they also find evidence that most people do not exhibit stable risk preferences under the assumption of expected utility maximization, suggesting that while people’s risk preferences may have a domain-general component, they may not be well represented by the expected utility model.

⁷⁷Contrast this with the test in Barseghyan, Prince, and Teitelbaum (2011). Under that test, if a household’s intervals intersect, then the household’s choices reflect a stable level (in absolute terms) of risk aversion across domains.

Rather than treat stability of risk preferences as a testable hypothesis, Barseghyan, Molinari, and Teitelbaum (2014) treat stability as a partial identification restriction and explore what they can learn about the structure of households' risk preferences from their risky choices across multiple domains. Working with the same insurance data (featuring two auto coverages and one home coverage) and probability distortion model studied by Barseghyan, Molinari, O'Donoghue, and Teitelbaum (2013b), Barseghyan, Molinari, and Teitelbaum show that a household's choice of deductible d^* in a given line of coverage implies lower and upper bounds on its distorted probability of experiencing a claim in that coverage:

$$LB \leq \Omega(\mu) \leq UB,$$

where

$$LB \equiv \max \left\{ 0, \max_{d > d^*} \Delta \right\} \quad \text{and} \quad UB \equiv \min \left\{ 1, \min_{d < d^*} \Delta \right\}$$

and

$$\Delta \equiv \frac{u(w - p(d)) - u(w - p(d^*))}{\begin{cases} [u(w - p(d)) - u(w - p(d) - d)] \\ -[u(w - p(d^*)) - u(w - p(d^*) - d^*)] \end{cases}},$$

and where $u(\cdot)$ and $\Omega(\cdot)$ are the household's utility and probability distortion functions, μ and w are the household's claim probability and wealth, and $p(d)$ is the household's premium for coverage with deductible d . Because the authors observe three choices per household (one choice per coverage) and a household's μ and $p(\cdot)$ differ across coverage, they obtain three pairs of bounds—or intervals—per household (one interval per coverage). Under the assumption of stability across contexts—i.e., assuming $u(\cdot)$ and $\Omega(\cdot)$ are domain-general and thus do not differ across coverages—the authors show to partially identify $u(\cdot)$ and $\Omega(\cdot)$ through shape restrictions.

To begin, they assume that $u(\cdot)$ belongs to the CARA family—the implication of which is that $u(\cdot)$ is fully characterized by the coefficient of absolute risk aversion r (which is assumed to be household-specific)—and they exclude households that makes choices which cannot be rationalized by any $r \in [0, 0.01]$. In addition, they consider five shape restrictions on $\Omega(\cdot)$: monotonicity, quadraticity, linearity, unit slope, and zero intercept.⁷⁸ They first use the intervals to recover the distribution of the lower bound on r under each shape restriction on $\Omega(\cdot)$. They find that the distribution is skewed to the right under each shape restriction, and that the median is zero under each non-degenerate shape restriction (i.e.,

⁷⁸The shape restrictions are cumulative. That is, quadraticity requires that $\Omega(\cdot)$ is both monotone and quadratic, i.e., $\Omega(\mu) = a + b\mu + c\mu^2$, $b \geq 0$ and $c \geq -b/2$; linearity requires $\Omega(\mu) = a + b\mu$, $b \geq 0$; unit slope requires $\Omega(\mu) = a + \mu$; and zero intercept requires $\Omega(\mu) = \mu$.

each shape restriction other than zero intercept). Indeed, they find that the vast majority of households—between 70 and 80 percent—can be rationalized by a model with linear utility given monotone, quadratic, or even linear probability distortions. By contrast, they find that fewer than 40 percent can be rationalized by a model with concave utility but no probability distortions. Next, they use the intervals to perform kernel regressions of the lower and upper bounds on $\Omega(\mu)$ as a function of μ . Under each non-degenerate shape restriction, the results evince a distortion function that substantially overweights small probabilities. Under monotonicity, for instance, the midpoints of the estimated bounds imply $\Omega(0.02) = 0.11$, $\Omega(0.05) = 0.17$, and $\Omega(0.10) = 0.25$, and even the estimated lower bounds imply $\Omega(0.02) = 0.07$, $\Omega(0.05) = 0.12$, and $\Omega(0.10) = 0.15$.⁷⁹

In the remainder of the paper, Barseghyan, Molinari, and Teitelbaum show how one can utilize the intervals to (i) classify households into preference types (i.e., special cases of the model) and (ii) point estimate $\Omega(\cdot)$ for the representative household. They also demonstrate a close connection between rank correlation of choices and stability of risk preferences under the probability distortion model. More specifically, they document that households' deductible choices are rank correlated across lines of coverage, echoing the finding by Einav, Finkelstein, Pascu, and Cullen (2012), and they show that it is "monotone" households who are driving these rank correlations.

In particular, the classification exercise, in addition to measuring the fraction of households whose behavior is consistent with various special cases of the model (e.g., EU vs KR-CPE), allows the researcher to measure the marginal contribution of each feature of the general model in rationalizing households' choices. For example, the analysis recovers that allowing for KR-CPE loss aversion *in addition to* standard curvature in the utility function increases the fraction of rationalizable households by a very small amount. In contrast, allowing for a monotone probability distortion function more than doubles the fraction of rationalizable households.

Barseghyan, Molinari, and Teitelbaum construct the point estimate of $\Omega(\cdot)$ without relying on parametric assumptions about the nature of unobserved heterogeneity. Rather, they built an estimator that comes closest to rationalizing the behavior of the average household. The estimated $\Omega(\cdot)$ can rationalize all three choices of nearly one in five households whose behavior is consistent with monotone probability distortions. The residual deviation between the households' intervals and the estimated $\Omega(\cdot)$ allows one to uncover the lower bound on the degree of heterogeneity in probability distortions among households. An interesting next

⁷⁹The results are very similar under quadraticity and linearity. Under quadraticity, for example, the midpoints imply $\Omega(0.02) = 0.11$, $\Omega(0.05) = 0.16$, and $\Omega(0.10) = 0.21$, and the lower bounds imply $\Omega(0.02) = 0.07$, $\Omega(0.05) = 0.12$, and $\Omega(0.10) = 0.15$.

step for future research is to generalize this approach to broader models of decision making under uncertainty.

6.2 Combining Experimental and Field Data

A special case of the question of consistency across contexts is the question of consistency between behavior in laboratory experiments and behavior in the field. Of course, if there is such consistency, one might be able to combine experimental data from the laboratory and field data to get a more complete picture of risk preferences. More importantly, however, answering this question allows one to establish the extent of portability of risk preference estimates obtained from the laboratory to real world applications.

In fact, there is some research that starts to address this question. One approach is to gather experimental and field data *for the same agents*, although typically covering very different contexts. For instance, some research elicits subjects' risk preferences using a standard experimental paradigm, while simultaneously collecting data on some aspect of those same subjects' field behaviors related to risk. The research question is then, whether subjects who are measured experimentally to have a stronger aversion to risk are also less likely to engage in risky behaviors in the field. Examples of such work are Liu (2013) and Liu and Huang (2013), which study how risk preferences elicited in a survey correlate, respectively, with the adoption of new crops and with the usage of pesticide by Chinese farmers. In a similar spirit, Meier and Sprenger (2010) study the link between time preference (patience) and credit card borrowing by the U.S. consumers.

An alternative approach is to gather experimental and field data on arguably the same or very similar choices—perhaps even with different agents—and then investigate the extent to which risk preferences estimated on experimental data correspond to risk preferences estimated on field data. One paper that follows this second approach is Barseghyan, O'Donoghue, and Xu (2015). In an online survey, subjects are presented with three deductible choices for property insurance coverages. The choice menus are constructed to match closely the deductible options and prices associated with specific households in the field data used by Barseghyan, Molinari, O'Donoghue, and Teitelbaum (2013b). The key difference is that agents in the laboratory are provided with their underlying (objective) claim probabilities, while Barseghyan, Molinari, O'Donoghue, and Teitelbaum (2013b) assume households beliefs correspond to observed claim probabilities.

Qualitatively, the findings of Barseghyan, O'Donoghue, and Xu (2015) confirm the patterns documented with field data (Barseghyan, Molinari, O'Donoghue, and Teitelbaum (2013b), Barseghyan, Molinari, and Teitelbaum (2014)): probability distortions can ratio-

nalize the behavior of the vast majority of households, while the curvature of the utility function alone cannot do so. The estimated (average) probability distortion function is increasing in the relevant range and exhibits significant over-weighting. Quantitatively, there are interesting differences between laboratory and field findings. There is more overweighting in the laboratory than in the field, as well as more "choice noise" and heterogeneity in agents' behavior.

Going forward, we envision more work emerging in this area, as more field data become available, and researchers are granted ways to design surveys that can reach (a subset of) subjects in the data (e.g., Handel and Kolstad (2015)).

6.3 Direct Measurement of Beliefs

When estimating risk preferences from choice data, a researcher typically faces a fundamental identification problem: Observed choices are consistent with many combinations of decision makers' risk preferences and their subjective expectations about various outcome probabilities. In the literature reviewed in the previous sections (insurance and betting), this fundamental identification problem is solved in one of two ways. One approach is to assume individuals hold objective expectations. In the insurance context, this translates into assuming that individuals know the objective claim probabilities or the objective health shocks distribution. And in the betting context this translates into assuming that individuals know horses' objective odds of winning races. The analysis then focuses on identification and estimation of risk preferences. A second approach is to assume that individuals are risk neutral, and the analysis focuses on identification and estimation of subjective beliefs.

A different and promising approach for solving this identification problem, is based on using survey data to measure probabilistic expectations directly, as advocated in Manski (2004). For example, one may elicit subjective beliefs on likelihood of a claim with questions such as: "What do you think is the percent chance that you will experience an auto collision claim within the next twelve months?". An important aspect of this approach is that expectations are not elicited in verbal form (e.g., by asking individuals how likely to occur an event is, with modifiers "very", "fairly", "not too" or "not at all" likely), but in numerical form (e.g., on a scale out of 100).

Extensive work in this literature has been devoted to show that respondents are willing and able to provide this information in probabilistic format. Manski (2004) and Hurd (2009) survey these findings in connection with data elicited in developed countries, while Delavande (2014) does the same in connection with data elicited in developing countries. The findings include, inter alia, that there is substantial heterogeneity in beliefs (at least in the contexts

analyzed so far), and that elicited beliefs correlate with individuals’ observable characteristics similarly to how actual outcomes do; see, e.g., Dominitz (1998) and Hurd and McGarry (2002). While it is not possible to evaluate directly whether the reported expectations are in fact those that respondents truly hold (because there cannot be validation data for this information), numerous studies in different contexts have shown that respondents give internally consistent answers that are sensible, when asked about questions that are relevant to their lives. Examples include Dominitz and Manski (1997), who study individuals’ income expectations, and Manski and Straub (2000), who analyze individuals’ expectations of their job security, in both cases analyzing data from the Survey of Economic Expectations. Manski (2004, Sections 5 and 6) summarizes their findings and the findings of many other studies. Probabilistic expectations data have been used to enrich econometric analysis of field data in several contexts already, including retirement behavior (e.g., Hurd, Smith, and Zissimopoulos (2004) and van der Klaauw and Wolpin (2008)), criminal behavior (e.g., Lochner (2007)), contraceptive choices and updating of beliefs on contraceptive effectiveness (e.g., Delavande (2008b) and Delavande (2008a)), and schooling choices (e.g., Giustinelli (2015) and Wiswall and Zafar (2015)). In the experimental literature, probabilistic expectations data have been used, for example, by Nyarko and Schotter (2002) and Dominitz and Hung (2009).

Within the analysis of risk preferences and risk perceptions, expectation data can be used to avoid the rational expectations and risk neutrality assumptions. We envision at least three ways in which these richer data can supplement choice data. First, it would be of interest to compare subjective beliefs with objective probabilities. Second, subjective beliefs could be compared with estimates of the probability distortion function estimated as in Barseghyan, Molinari, O’Donoghue, and Teitelbaum (2013b). Third, estimation of a rich model encompassing standard risk aversion and a probability distortion function could be conducted, using the reported subjective beliefs in place of the objective probabilities. Recovering no further probability distortions would indicate that in fact subjective beliefs drive individuals’ decision. Recovering a probability distortion function that does not come close to the identity function may further corroborate the finding that probability weighting plays a key role in individuals’ decisions.

6.4 Assumptions about “Mental Accounting”

As we discuss in Section 4.1.1, in order to estimate risk preferences in a field context, one must make assumptions about how the typically complex field context is translated into concrete choice data that can be used to estimate risk preferences. One justification for such assumptions is that the agents themselves transform—and simplify—a choice situation

in their own minds before making a decision. This type of mental operation on the part of agents is often labelled “mental accounting”, and thus we refer to these assumptions as mental accounting assumptions. However, these assumptions are often not motivated by psychological realism, but rather as “simplifying assumptions” that help the researcher make progress—e.g., to overcome limitations of the data or to simplify computations. An important agenda for future research is to pay more careful attention to these assumptions, and to investigate directly the impact of such assumptions on estimates of risk preferences.

An important dimension on which one must make a mental accounting assumption is how broadly vs. narrowly households bracket their decisions. On one extreme, households could bracket all their decisions together into one grand “life” decision—indeed, theoretical economic models are often written in this way. On the other extreme, households could bracket very narrowly and evaluate each decision in isolation from all others. There are also many possibilities in between.

In fact, virtually all papers that estimate risk preferences implicitly –and occasionally explicitly– assume very narrow bracketing. They estimate risk preferences reflected in one particular choice in isolation from how that choice might interact with the many other choices that households make. In most cases, we suspect that this narrow bracketing is assumed merely to help the researcher. If one has data on only one decision per household, it is hard to assume anything other than narrow bracketing. Even when one has data on multiple decisions per household, it can be computationally burdensome to collect them all together into one grand decision.⁸⁰ There may, however, be some psychological realism to the assumption of narrow bracketing. Indeed, there is a literature which suggests that when people make multiple choices, they frequently do not assess the consequences in an integrated way, but rather tend to make each choice in isolation (e.g., Read, Loewenstein, and Rabin 1999).

In future research, it is worth investigating more carefully how broadly households bracket their decisions. In simple terms, we need to understand whether incorrect assumptions about bracketing bias estimates of risk preferences. Even beyond this, if households do in fact bracket multiple decisions together, then they are choosing between more complex lotteries, and, as highlighted in Section 3.6.4, this might permit one to separately identify multiple sources of aversion to risk. Indeed, Barseghyan, Molinari, O’Donoghue, and Teitelbaum (2013a) propose exactly this approach.

It is also worth noting a more reduced-form approach to account for broader bracketing,

⁸⁰For instance, if one assumed daily bracketing at a horse track, and if on any day the track holds 10 races with 8 horses in each, then there are 8^{10} possible lotteries that one could choose (and even this is restricting attention to win bets and ignoring any dynamics associated with basing later bets on earlier results).

by incorporating “background risk” into one’s analysis. In other words, instead of directly modeling all the other decisions that a household faces, one can reduce all of those decisions into the resulting background risk that a person already faces when making the present decision. Operationally, this means replacing the assumption that a household has a certain prior wealth with an assumption that prior wealth is a lottery itself.

A second dimension on which one must make a mental accounting assumption is what options enter a household’s consideration set—that is, what options does a household take to be in its choice set. Because estimates of risk preferences can depend on what options are considered but not chosen, assumptions about the consideration set can alter estimates. To illustrate, consider a stylized example: Suppose we observe an individual in a casino who chooses to bet \$10 on BLACK in roulette (on a typical roulette wheel that has both 0 and 00). This choice yields a risky lottery $(+\$10, \frac{18}{38}; -\$10, \frac{20}{38})$ that has an expected value of $-\$0.526$. This person could have also bet \$10 on #1, which instead would yield a riskier lottery $(+\$360, \frac{1}{38}; -\$10, \frac{37}{38})$ that has a larger expected value of $-\$0.263$. If we estimated, for example, an EU model focusing on the fact that the person chose to bet on BLACK rather than bet on #1, we would conclude that the person is risk averse, and we could infer a lower bound on the magnitude of this risk aversion. However, if instead we estimated an EU model focusing on the fact that the person chose to bet on BLACK rather than to not bet at all, we would instead conclude that the person is risk loving, and we could infer a lower bound on the magnitude of this risk lovingness.

The intuition of this example extends almost immediately to research that estimates risk preferences using data from horse races. In such analyses, one must make an assumption about whether the option not to bet is included in the consideration set. Because the expected return on most horses is negative, including the option not to bet in the consideration set will yield estimates of risk preferences that are more risk loving. But this issue also applies to property insurance. When estimating preferences from deductible choices, one might wonder whether households consider the possibility of not insuring at all, and also whether they turn down any insurance options from other firms that are not in the dataset. In future research, it is worth investigating more carefully the determinants and importance of consideration sets.

A third dimension on which one must make a mental accounting assumption is what are the possible outcomes that households account for. To illustrate, consider property insurance. When researchers estimate risk preferences using data on property insurance, they typically assume that households only consider the possibility of incurring no loss or a single loss during the policy period. However, in principle, one might incur two, three, or even more losses during a policy period. If so, then the set of lotteries from which households are choosing are

different. The assumption of zero or one loss is often made for simplicity, but again it could reflect a psychological realism, as people seem to have a hard time imagining and accounting for all the possibilities that could occur in life. This is especially true in more complex domains, such as health insurance, where it seems quite likely that households approach decisions with a simplified conceptualization of all the possible outcomes that might occur.

Moving forward, we think it important that the literature considers more carefully and more directly these and other mental accounting assumptions when estimating risk preferences. Such assumptions can matter under EU, and they become even more important under RDEU and other more complex models of risk preferences. Hence, these assumptions need to be taken more seriously.

References

- ABBRING, J. H., P.-A. CHIAPPORI, J. HECKMAN, AND J. PINQUET (2003): “Adverse Selection and Moral Hazard in Insurance: Can Dynamic Data Help to Distinguish?,” *Journal of the European Economic Association*, 1(2-3), 512–521.
- ABBRING, J. H., P.-A. CHIAPPORI, AND J. PINQUET (2003): “Moral Hazard and Dynamic Insurance Data,” *Journal of the European Economic Association*, 1(4), 767–820.
- ABBRING, J. H., P.-A. CHIAPPORI, AND T. ZAVADIL (2008): “Better Safe than Sorry? Ex Ante and Ex Post Moral Hazard in Dynamic Insurance Data,” CentER Discussion Paper 2008-77, Tilburg University.
- ABEL, A. B. (1990): “Asset Prices Under Habit Formation and Catching Up With the Joneses,” *American Economic Review: Papers and Proceedings*, 80(2), 38–42.
- (1999): “Risk Premia and Term Premia in General Equilibrium,” *Journal of Monetary Economics*, 43(1), 3–33.
- ALI, M. M. (1977): “Probability and Utility Estimates for Racetrack Bettors,” *Journal of Political Economy*, 85, 803–815.
- ANDERSEN, S., G. W. HARRISON, M. I. LAU, AND E. E. RUTSTRÖM (2008): “Risk Aversion in Game Shows,” in *Risk Aversion in Experiments*, ed. by J. C. Cox, and G. W. Harrison, vol. 12 of *Research in Experimental Economics*, p. 3610406. JAI Press, Bingley.
- BAJARI, P., C. DALTON, H. HONG, AND A. KHWAJA (2014): “Moral hazard, adverse selection, and health expenditures: A semiparametric analysis,” *The RAND Journal of Economics*, 45(4), 747–763.
- BARSEGHYAN, L., F. MOLINARI, T. O’DONOGHUE, AND J. C. TEITELBAUM (2013a): “Distinguishing Probability Weighting from Risk Misperceptions in Field Data,” *American Economic Review: Papers and Proceedings*, 103(3), 580–585.
- (2013b): “The Nature of Risk Preferences: Evidence from Insurance Data,” *American Economic Review*, 103, 2499–2529.
- BARSEGHYAN, L., F. MOLINARI, AND J. C. TEITELBAUM (2014): “Inference Under Stability of Risk Preferences,” Working paper.
- BARSEGHYAN, L., T. O’DONOGHUE, AND L. XU (2015): “How different are insurance purchases in experiments and the real world?,” in preparation.

- BARSEGHYAN, L., J. PRINCE, AND J. C. TEITELBAUM (2011): “Are Risk Preferences Stable Across Contexts? Evidence from Insurance Data,” *American Economic Review*, 101(2), 591–631.
- BEETSMA, R. M. W. J., AND P. C. SCHOTMAN (2001): “Measuring Risk Attitudes in a Natural Experiment: Data from the Television Game Show Lingo,” *Economic Journal*, 111(474), 821–848.
- BELL, D. E. (1985): “Disappointment in Decision Making under Uncertainty,” *Operations Research*, 33(1), 1–27.
- BERRY, S., A. GANDHI, AND P. HAILE (2013): “Connected Substitutes and Invertibility of Demand,” *Econometrica*, 81, 2087–2111.
- BERRY, S., J. LEVINSOHN, AND A. PAKES (1995): “Automobile Prices in Market Equilibrium,” *Econometrica*, 63, 841–890.
- BOMBARDINI, M., AND F. TREBBI (2012): “Risk Aversion and Expected Utility Theory: An Experiment with Large and Small Stakes,” *Journal of the European Economic Association*, 10(6), 1348–1399.
- BOTTI, F., AND A. CONTE (2008): “Risk Attitude in Real Decision Problems,” *B.E. Journal of Economic Analysis & Policy: Advances*, 8(1), Article 6.
- BUNDORF, M. K., J. LEVIN, AND N. MAHONEY (2012): “Pricing and Welfare in Health Plan Choice,” *American Economic Review*, 102(7), 3214–48.
- CAMERER, C. (1995): “Individual Decision Making,” in *Handbook of Experimental Economics*, ed. by J. H. Hagel, and A. E. Roth, chap. 8, pp. 587–703. Princeton University Press, Princeton.
- CAMPBELL, J. Y. (2003): “Consumption-Based Asset Pricing,” in *Handbook of the Economics of Finance*, ed. by G. M. Constantinides, M. Harris, and R. M. Stulz, vol. 1B, chap. 13, pp. 803–887. Elsevier, Amsterdam.
- CAMPBELL, J. Y., AND J. H. COCHRANE (1999): “By Force of Habit: A Consumption-Based Explanation of Aggregate Stock Market Behavior,” *Journal of Political Economy*, 107(2), 205–251.
- CARDON, J. H., AND I. HENDEL (2001): “Asymmetric Information in Health Insurance: Evidence from the National Medical Expenditure Survey,” *RAND Journal of Economics*, 32, 408–442.

- CRAWLEY, J., AND T. PHILIPSON (1999): “An Empirical Examination of Information Barriers to Trade in Insurance,” *American Economic Review*, 89(4), 827–846.
- CECCARINI, O. (2007): “Does Experience Rating Matter in Reducing Accident Probabilities? A Test for Moral Hazard,” Discussion paper, <http://www.econ.upenn.edu/~oliviace/jobmkt.pdf>.
- CHETTY, R. (2006): “A New Method of Estimating Risk Aversion,” *American Economic Review*, 96(5), 1821–1834.
- CHIAPPORI, P.-A. (1999): “Asymmetric Information In Automobile Insurance: An Overview,” in *Automobile Insurance: Road Safety, New Drivers, Risks, Insurance Fraud and Regulation*, ed. by G. Dionne, and C. Laberge-Nadeau, pp. 1–11. Kluwer Academic Publishers, Boston.
- CHIAPPORI, P.-A., B. JULLIEN, B. SALANIÉ, AND F. SALANIÉ (2006): “Asymmetric Information in Insurance: General Testable Implications,” *RAND Journal of Economics*, 37(4), 783–798.
- CHIAPPORI, P.-A., AND M. PAIELLA (2011): “Relative Risk Aversion Is Constant: Evidence from Panel Data,” *Journal of the European Economic Association*, 9(6), 1021–1052.
- CHIAPPORI, P. A., AND B. SALANIÉ (2000): “Testing for Asymmetric Information in Insurance Markets,” *Journal of Political Economy*, 108(1), 56–78.
- CHIAPPORI, P. A., B. SALANIÉ, F. SALANIÉ, AND A. GANDHI (2012): “From Aggregate Betting Data to Individual Risk Preferences,” Working paper.
- CHIAPPORI, P. A., K. SAMPHATHARAK, S. SCHULHOFER-WOHL, AND R. M. TOWNSEND (2013): “Portfolio Choices and Risk Preferences in Village Economies,” Federal Reserve Bank of Minneapolis Working Paper 706.
- (2014): “Heterogeneity and Risk-Sharing in Village Economies,” *Quantitative Economics*, forthcoming.
- CICCHETTI, C. J., AND J. A. DUBIN (1994): “A Microeconometric Analysis of Risk Aversion and the Decision to Self-Insure,” *Journal of Political Economy*, 102, 169–186.
- COHEN, A. (2005): “Asymmetric Information and Learning: Evidence from the Automobile Insurance Market,” *Review of Economics and Statistics*, 87(2), 197–207.

- COHEN, A., AND L. EINAV (2007): “Estimating Risk Preferences from Deductible Choice,” *American Economic Review*, 97(3), 745–788.
- COHEN, A., AND P. SIEGELMAN (2010): “Testing for Adverse Selection in Insurance Markets,” *Journal of Risk and Insurance*, 77(1), 39–84.
- CONSTANTINIDES, G. M. (1990): “Habit Formation: A Resolution of the Equity Premium Puzzle,” *Journal of Political Economy*, 98(3), 519–543.
- CONTE, A., P. G. MOFFATT, F. BOTTI, D. T. DI CAGNO, AND C. D’IPPOLITI (2012): “A Test of the Rational Expectations Hypothesis Using Data from a Natural Experiment,” *Applied Economics*, 44, 4661–4678.
- DE ROOS, N., AND Y. SARAFIDIS (2010): “Decision Making Under Risk in *Deal or No Deal*,” *Journal of Applied Econometrics*, 25, 987–1027.
- DECK, C., J. LEE, AND J. REYES (2008): “*Risk attitudes in Large Stake Gambles: Evidence from a Game Show*,” *Applied Economics*, 40(1), 41–52.
- DELAVANDE, A. (2008a): “*Measuring Revisions to Subjective Expectations*,” *Journal of Risk and Uncertainty*, 36(1), 43–82.
- (2008b): “*Pill, Patch or Shot? Subjective Expectation and Birth Control Choice*,” *International Economic Review*, 49(3), 999–1042.
- (2014): “*Probabilistic Expectations in Developing Countries*,” *Annual Review of Economics*, 6, 1–20.
- DIONNE, G., C. GOURIÉROUX, AND C. VANASSE (1999): “*Evidence of Adverse Selection in Automobile Insurance Markets*,” in *Automobile Insurance: Road Safety, New Drivers, Risks, Insurance Fraud and Regulation*, ed. by G. Dionne, and C. Laberge-Nadeau, pp. 13–46. Kluwer Academic Publishers, Boston.
- (2001): “*Testing for Evidence of Adverse Selection in the Automobile Insurance Market: A Comment*,” *Journal of Political Economy*, 109(2), 444–451.
- DIONNE, G., P.-C. MICHAUD, AND M. DAHCHOUR (2007): “*Separating Moral Hazard from Adverse Selection and Learning in Automobile Insurance: Longitudinal Evidence from France*,” *Discussion paper*, <http://www.rmi.gsu.edu/Research/Dionne/Separating-JEEA.pdf>.

- DOMINITZ, J. (1998): “*Earnings Expectations, Revisions, and Realizations*,” Review of Economics and Statistics, 80(3), 374–388.
- DOMINITZ, J., AND A. A. HUNG (2009): “*Empirical models of discrete choice and belief updating in observational learning experiments*,” Journal of Economic Behavior and Organization, 69(2), 94–109.
- DOMINITZ, J., AND C. F. MANSKI (1997): “*Using Expectations Data to Study Subjective Income Expectations*,” Journal of the American Statistical Association, 92(439), 855–867.
- EDWARDS, W. (1955): “*The Prediction of Decisions among Bets*,” Journal of Experimental Psychology, 50, 201–214.
- (1962): “*Subjective Probabilities Inferred from Decisions*,” Psychological Review, 69, 109–135.
- EINAV, L., A. FINKELSTEIN, AND M. CULLEN (2010): “*Estimating Welfare in Insurance Markets using Variation in Prices*,” Quarterly Journal of Economics, 125(3), 877–921.
- EINAV, L., A. FINKELSTEIN, I. PASCU, AND M. R. CULLEN (2012): “*How General Are Risk Preferences? Choices under Uncertainty in Different Domains*,” American Economic Review, 102(6), 2606–2638.
- EINAV, L., A. FINKELSTEIN, S. P. RYAN, P. SCHRIMPF, AND M. R. CULLEN (2013): “*Selection on moral hazard in health insurance*,” American Economic Review, 103(1), 178–219.
- EINAV, L., A. FINKELSTEIN, AND P. SCHRIMPF (2010): “*Optimal Mandates and the Welfare Cost of Asymmetric Information: Evidence From the U.K. Annuity Market*,” Econometrica, 78(3), 1031–1092.
- EPSTEIN, L. G., AND S. E. ZIN (1989): “*Substitution, Risk Aversion, and the Temporal Behavior of Consumption and Asset Returns: A Theoretical Framework*,” Econometrica, 57(4), 937–969.
- (1991): “*Substitution, Risk Aversion, and the Temporal Behavior of Consumption and Asset Returns: An Empirical Analysis*,” Journal of Political Economy, 99(2), 263–286.
- FANG, H., M. P. KEANE, AND D. SILVERMAN (2008): “*Sources of Advantageous Selection: Evidence from the Medigap Insurance Market*,” Journal of Political Economy, 116(2), 303–350.

- FINKELSTEIN, A., AND K. MCGARRY (2006): “Multiple Dimensions of Private Information: Evidence from the Long-Term Care Insurance Market,” *American Economic Review*, 96(4), 938–958.
- FINKELSTEIN, A., AND J. POTERBA (2004): “Adverse selection in insurance markets: Policyholder evidence from the U.K. Annuity Market,” *Journal of Political Economy*, vol. 112(1), 183–208.
- FULLENKAMP, C., R. TENORIO, AND R. BATTALIO (2003): “Assessing Individual Risk Attitudes Using Field Data from Lottery Gameshow,” *Review of Economics and Statistics*, 85(1), 218–225.
- GANDHI, A., AND R. SERRANO-PADIAL (2014): “Does Belief Heterogeneity Explain Asset Prices: The Case of the Longshot Bias,” *Review of Economic Studies*, forthcoming, Working paper.
- GERTNER, R. (1993): “Game Shows and Economic Behavior: Risk-Taking on “Card Sharks,”” *Quarterly Journal of Economics*, 108(2), 507–521.
- GIUSTINELLI, P. (2015): “Group Decision Making with Uncertain Outcomes: Unpacking Child-Parent Choice of the High School Track,” *International Economic Review*, forthcoming.
- GONZALEZ, R., AND G. WU (1999): “On the Shape of the Probability Weighting Function,” *Cognitive Psychology*, 38, 129–166.
- GRIFFITH, R. M. (1949): “Odds Adjustments by American Horse-Race Bettors,” *The American Journal of Psychology*, 62, 290–294.
- (1961): “A Footnote on Horse Race Betting,” *Transactions, Kentucky Academy of Science*, 22, 78–81.
- GUL, F. (1991): “A Theory of Disappointment Aversion,” *Econometrica*, 59(3), 667–686.
- HANDEL, B. R. (2013): “Adverse Selection and Inertia in Health Insurance Markets: When Nudging Hurts,” *American Economic Review*, 103(7), 2643–2682.
- HANDEL, B. R., AND J. T. KOLSTAD (2015): “Health Insurance for “Humans”: Information Frictions, Plan Choice, and Consumer Welfare,” *American Economic Review*, forthcoming.

- HANSEN, L. P., AND K. T. SINGLETON (1982): “Generalized Instrumental Variables Estimation of Nonlinear Rational Expectations Models,” *Econometrica*, 50(5), 1269–1286.
- (1983): “Stochastic Consumption, Risk Aversion, and the Temporal Behavior of Asset Returns,” *Journal of Political Economy*, 91, 249–265.
- HARTLEY, R., G. LANOT, AND I. WALKER (2014): “Who Really Wants to Be a Millionaire? Estimates of Risk Aversion from Gameshow Data,” *Journal of Applied Econometrics*, 29, 861–879.
- HE, D. (2009): “The life insurance market: Asymmetric information revisited,” *Journal of Public Economics*, 93(9–10), 1090–1097.
- HURD, M. D. (2009): “Subjective Probabilities in Household Surveys,” *Annual Review of Economics*, 1, 543–562.
- HURD, M. D., AND K. MCGARRY (2002): “The Predictive Validity of Subjective Probabilities of Survival,” *The Economic Journal*, 112(482), 966–985.
- HURD, M. D., J. P. SMITH, AND J. M. ZISSIMOPOULOS (2004): “The effects of subjective survival on retirement and Social Security claiming,” *Journal of Applied Econometrics*, 19(6), 761–775.
- ISRAEL, M. (2004): “Do We Drive More Safely When Accidents are More Expensive? Identifying Moral Hazard from Experience Rating Schemes,” *Csio working paper 0043*, Northwestern University.
- JULLIEN, B., AND B. SALANIÉ (2000): “Estimating Preferences Under Risk: The Case of Racetrack Bettors,” *Journal of Political Economy*, 108(3), 503–530.
- KAHNEMAN, D., AND A. TVERSKY (1979): “Prospect Theory: An Analysis of Decision under Risk,” *Econometrica*, 47, 263–291.
- KARMAKAR, U. S. (1978): “Subjectively Weighted Utility: A descriptive extension of the Expected Utility model,” *Organizational Behavior and Human Performance*, 21, 61–72.
- KÖSZEGI, B., AND M. RABIN (2006): “A Model of Reference-Dependent Preferences,” *Quarterly Journal of Economics*, 121(4), 1133–1166.
- (2007): “Reference-Dependent Risk Attitudes,” *American Economic Review*, 97(4), 1047–1073.

- KOCHERLAKOTA, N. R. (1996): “*The Equity Premium: It’s Still a Puzzle*,” *Journal of Economic Literature*, 34, 42–71.
- LATTIMORE, P. K., J. R. BAKER, AND A. D. WITTE (1992): “*The influence of probability on risky choice : A parametric examination*,” *Journal of Economic Behavior and Organization*, 17, 377–400.
- LIU, E. M. (2013): “*Time to Change What to Sow: Risk Preferences and Technology Adoption Decisions of Cotton Farmers in China*,” *Review of Economics and Statistics*, forthcoming.
- LIU, E. M., AND J. HUANG (2013): “*Risk Preference and Pesticide Use by Cotton Farmers in China*,” *Journal of Development Economics*, 103, 202–215.
- LOCHNER, L. (2007): “*Individual Perceptions of the Criminal Justice System*,” *American Economic Review*, 97(1), 444–460.
- LOOMES, G., AND R. SUGDEN (1986): “*Disappointment and Dynamic Consistency in Choice under Uncertainty*,” *Review of Economic Studies*, 53(2), 271–282.
- MANSKI, C. F. (2004): “*Measuring Expectations*,” *Econometrica*, 72(5), 1329–1376.
- MANSKI, C. F., AND J. D. STRAUB (2000): “*Worker Perceptions of Job Insecurity in the Mid-1990s: Evidence from the Survey of Economic Expectations*,” *The Journal of Human Resources*, 35(3), 447–479.
- MARKOWITZ, H. (1952): “*The Utility of Wealth*,” *Journal of Political Economy*, 60(2), 151–158.
- MASATLIOGLU, Y., AND C. RAYMOND (2014): “*A Behavioral Analysis of Stochastic Reference Dependence*,” *Mimeo, University of Michigan*.
- McFADDEN, D. L. (1974): “*Conditional Logit Analysis of Qualitative Choice Behavior*,” in *Frontiers in Econometrics*, ed. by P. Zarembka, pp. 105–142. Academic Press, New York.
- MCGLOTHLIN, W. H. (1956): “*Stability of Choices among Uncertain Alternatives*,” *The American Journal of Psychology*, 69, 604–615.
- MEHRA, R., AND E. C. PRESCOTT (1985): “*The Equity Premium: A Puzzle*,” *Journal of Monetary Economics*, 15, 145–161.
- MEIER, S., AND C. SPRENGER (2010): “*Present-Biased Preferences and Credit Card Borrowing*,” *American Economic Journal: Applied Economics*, 2, 193–210.

- METRICK, A. (1995): "A Natural Experiment in "Jeopardy!"," American Economic Review, 85(1), 240–253.
- MULINO, D., R. SCHEELINGS, R. BROOKS, AND R. FAFF (2009): "Does Risk Aversion Vary with Decision-Frame? An Empirical Test Using Recent Game Show Data," Review of Behavioral Finance, 1, 44–61.
- NYARKO, Y., AND A. SCHOTTER (2002): "An Experimental Study of Belief Learning Using Elicited Beliefs," Econometrica, 70(3), 971–1005.
- PINQUET, J., G. DIONNE, C. VANASSE, AND M. MATHIEU (2008): "Point-Record Inventives, Asymmetric Information and Dynamic Data," Discussion paper, <http://www.rmi.gsu.edu/Research/Dionne/RES-11853-08-07-22.pdf>.
- POST, T., M. J. VAN DEN ASSEM, G. BALTUSSEN, AND R. H. THALER (2008): "Deal or No Deal? Decision Making under Risk in a Large-Payoff Game Show," American Economic Review, 98(1), 38–71.
- PRELEC, D. (1998): "The Probability Weighting Function," Econometrica, 66, 497–527.
- PRESTON, M. G., AND P. BARATTA (1948): "An Experimental Study of the Auction-Value of an Uncertain Outcome," The American Journal of Psychology, 61, 183–193.
- PUELZ, R., AND A. SNOW (1994): "Evidence on Adverse Selection: Equilibrium Signaling and Cross-Subsidization in the Insurance Market," Journal of Political Economy, 102(2), 236–257.
- QUIGGIN, J. (1982): "A Theory of Anticipated Utility," Journal of Economic Behavior and Organization, 3(4), 323–343.
- RABIN, M. (2000): "Risk Aversion and Expected-Utility Theory: A Calibration Theorem," Econometrica, 68(5), 1281–1292.
- READ, D., G. LOEWENSTEIN, AND M. RABIN (1999): "Choice Bracketing," Journal of Risk and Uncertainty, 19(1-3), 171–197.
- REES-JONES, A. R. (2014): "Loss Aversion Motivates Tax Sheltering: Evidence from U.S. Tax Returns," Working Paper, SSRN.
- SIEGEL, J. J., AND R. H. THALER (1997): "The Equity Premium Puzzle," Journal of Economic Perspectives, 11(1), 191–200.

- SNOWBERG, E., AND J. WOLFERS (2010): “*Explaining the Favorite-Long Shot Bias: Is it Risk-Love or Misperceptions?*,” *Journal of Political Economy*, 118(4), 723–746.
- STARC, A. (2014): “*Insurer pricing and consumer welfare: evidence from Medigap*,” *The RAND Journal of Economics*, 45, 198–220.
- STARMER, C. (2000): “*Developments in Non-expected Utility Theory: The Hunt for a Descriptive Theory of Choice under Risk*,” *Journal of Economic Literature*, 38(2), 332–382.
- SYDNOR, J. (2010): “*(Over)Insuring Modest Risks*,” *American Economic Journal: Applied Economics*, 2(4), 177–199.
- TVERSKY, A., AND D. KAHNEMAN (1992): “*Advances in Prospect Theory: Cumulative Representation of Uncertainty*,” *Journal of Risk and Uncertainty*, 5(4), 297–323.
- VAN DER KLAUW, W., AND K. I. WOLPIN (2008): “*Social security and the retirement and savings behavior of low-income households*,” *Journal of Econometrics*, 145(1-2), 21–42.
- WEIL, P. (1989): “*The Equity Premium Puzzle and the Risk-Free Rate Puzzle*,” *Journal of Monetary Economics*, 24, 401–421.
- (1990): “*Nonexpected Utility in Macroeconomics*,” *Quarterly Journal of Economics*, 105(1), 29–42.
- WEITZMAN, M. (1965): “*Utility Analysis and Group Behavior: An Empirical Study*,” *Journal of Poli*, 73, 18–26.
- WISWALL, M. J., AND B. ZAFAR (2015): “*Determinants of College Major Choices: Identification from an Information Experiment*,” *Review of Economic Studies*, *forthcoming*.
- WOLF, C., AND L. POHLMAN (1983): “*The Recovery of Risk Preferences from Actual Choices*,” *Econometrica*, 51(3), 843–850.
- YAARI, M. E. (1965): “*Convexity in the Theory of Choice under Risk*,” *Quarterly Journal of Economics*, 79, 278–90.