

Choices under Risk in Rural Peru^{*}

Francisco B. Galarza[†]
University of Wisconsin–Madison

August, 2009

Abstract

This paper estimates the risk preferences of cotton farmers in Southern Peru, using the results from a multiple-price-list lottery game. Assuming that preferences conform to two of the leading models of decision under risk—Expected Utility Theory (EUT) and Cumulative Prospect Theory (CPT)—we find strong evidence of moderate risk aversion. Once we include individual characteristics in the estimation of risk parameters, we observe that farmers use subjective nonlinear probability weighting, a behavior consistent with CPT. Interestingly, when we allow for preference heterogeneity via the estimation of mixture models—where the proportion of subjects who behave according to EUT or to CPT is endogenously determined—we report that the majority of farmers’ choices are best explained by CPT. We further hypothesize that the multiple switching behavior observed in our sample can be explained by nonlinear probability weighting made in a context of large random calculation mistakes; the evidence found on this regard is mixed. Finally, we find that attaining higher education is the single most important individual characteristic correlated with risk preferences, a result that suggests a connection between cognitive abilities and behavior towards risk.

Keywords: risk aversion, probability weighting, mixture models, experimental economics, Peru.

JEL classification numbers: D81, Q12.

^{*}I am indebted to my adviser, Michael R. Carter, for his continuous encouragement and insightful suggestions. I also thank Brad Barham, Jed Frees, Paul Mitchell, and Laura Schechter for helpful comments, and seminar participants at the University of Wisconsin and the 2009 Pacific Conference for Development Economics for their suggestions. Steve Boucher, Carlos de los Rios, Conner Mullally, and Carolina Trivelli provided helpful feedback on the experimental design. Ramón Díaz, Oscar Madalengoitia, Roberto Piselli, Chris Rue, Raphael Saldaña, Jessica Varney, Josh Weinberg, and especially Johanna Yancari, provided a valuable assistance in the field. Financial support from the USAID Cooperative Agreement No. EDH-A-00-06-0003-00 through the Assets and Market Access Collaborative Research Support Program is gratefully acknowledged. Usual disclaimer applies.

[†]Department of Agricultural and Applied Economics. E-mail: fbgalarzaare@wisc.edu.

1 Introduction

The study of decisions under risk has been one of the main subjects in economics since at least Arrow (1971). From a development economics perspective, risk has been considered as a factor that can slow down the adoption of financial or production innovations (Feder, 1980; Feder et al., 1985). It is argued that risk considerations may prevent subjects from undertaking potentially profitable investments. Similar concerns may also encourage subjects to continue using traditional farming technologies instead of new, more profitable technologies because of the uncertainty that may be involved in their adoption process. Experimental evidence from China for the case of *Bt* cotton (Liu, 2008) and India (Binswanger et al., 1980) supports this claim. As a consequence, production and financial decisions may result in an underoptimal accumulation of assets, with a deterioration of one's ability to cope with large shocks in the long run.

Recent evidence from the laboratory (e.g., Hey and Orme, 1994; Holt and Laury, 2002) and the field (e.g., Harrison et al., 2007, 2009; Schechter, 2007) suggests that subjects are, in general, risk averse over the gains domain.¹ On the experimental ground, several methods of measuring risk preferences exist (for a review of these methods, see Cox and Harrison, 2008), but the multiple price listing (MPL) is one of the most commonly used.² Under the MPL format, subjects are given a menu of choices with ordered prices, one per row. The task is to indicate “yes” or “no” for each price, and at the end the experimenter implements one of the rows at random, and subjects get the choice made in that row.

In order to measure risk using the MPL format, choices are binary lotteries and the prices are given for the probability structure associated with each lottery; thus, in each row, subjects should decide whether to take one lottery or the other, given the probability that each prize has within a given lottery. The main attractive feature of the MPL is its simplicity to explain, implement, and elicit true valuations, a trait that is especially important in a farming context where a large proportion of the individuals have typically low levels of schooling. This method has also been found to yield less noise than alternative elicitation methods that attempt to elicit certainty equivalents, such as: the measurement of willingness to pay in an auction, in which subjects report a maximum buying price for a lottery; the measurement of willingness to accept in an auction, in which subjects report a minimum selling price for a lottery; or the Becker et al. (1964) method, in which subjects report a certainty equivalent for a lottery. (See Hey et al. [2007] for details). One of the disadvantages of this method, however, is that it allows for multiple switching between lotteries, a behavior that is not expected from fully rational players. Multiple switching behavior (MSB) has been reported by previous studies (e.g., Andersen et al., 2006; Bruner et al., 2008; Eckel

¹One of the exceptions include Henrich and McElreath (2002), who find that while Chilean peasants are risk loving, Tanzanian peasants seem risk averse. They used three binary lotteries, with a safe choice and a risky bet, with the idea being to use the indifference point between the safe choice and the lottery as a measure of certainty equivalence value.

²Other methods using nonexperimental data include contingent valuation methods and structural estimation.

and Wilson, 2004; Holt and Laury, 2002; Jacobson and Petrie, 2009),³ and it has been attributed to indifference between the two lotteries (e.g., Andersen et al., 2006) or to the lack of salience in the lottery prizes (e.g., Bruner, 2007).⁴ However, MSB could also be due to confusion, incomplete understanding of the game rules, or to the lack of attention due to boredom or hurry.

The typical solution for the multiple switching in MPL designs has been to discard these observations under the implicit assumption that there is nothing to learn from those seemingly irrational choices or *mistakes*. This solution can be understandable when the proportion of such inconsistent choices is small (as found by most of the previously mentioned studies), but when such proportion is large (as we found in this paper), then we believe that a more constructive approach would be to examine whether those inconsistent choices can be rationalized under a particular framework.

Our estimation of the risk parameters will consider the Expected Utility Theory (EUT) as the work horse. We will then generalize EUT to account for the possibility that, instead of objective probabilities, individuals make subjective distortions to those probabilities in their decisions under risk. These distortions lead people to underweight or overweight probabilities, a phenomenon examined by Tversky and Kahneman’s (1992) Cumulative Prospect Theory (CPT), and supported by substantial experimental evidence (Starmer, 2000).⁵ An additional appealing feature that motivates us to consider CPT is that, as we state later in this paper, it can help explain the MSB observed in our data.

This paper uses data gathered in rural Peru, where we created an experimental economics laboratory to elicit measures of risk aversion. The experimental sessions were conducted with small-scale cotton farmers over six weeks in a southern coastal valley, using the MPL format popularized by Holt and Laury (2002).⁶ Experimental subjects were asked to choose between a relatively safe and a relatively risky lottery along ten decision rows (indifference between lotteries was not an option). Given that a large proportion of our subjects (52 percent) switched back and forth from one lottery to the other at least twice,⁷ rather than discarding the observations where MSB was observed, we will adopt a more constructive approach and will keep them in our sample,

³Andersen et al. (2006) find that 5.8 percent of their subjects switch multiple times when allowing for indifference (taken by 24.3 percent of subjects). In Bruner et al. (2008) such proportion is 20 percent; in Eckel and Wilson (2004), 12.9 percent; in Holt and Laury (2002), 13.2 percent in hypothetical choices; and in Jacobson and Petrie (2009), 55 percent.

⁴In a setting where game instructions are written and displayed on computers, Bruner (2007) finds that reinforcing verbally that earnings for the risk game would be determined *only* [added emphasis] for one decision row, significantly reduces the proportion of multiple switches.

⁵Another feature of prospect theory is the concept of *loss aversion* that states that individuals are more sensitive to losses than gains of similar magnitude. We do not explore loss aversion in this paper, given that our game is defined over gains, not losses.

⁶This research project was carried out in partnership with an insurance company and a vendor of insurance contracts bundled with loans. At all times, we emphasized our participation as researchers was simply to inform farmers about the main features of this new financial product and to examine their willingness to buy it. We also stressed the fact that participating in these sessions should not make them feel obliged to buy insurance.

⁷The magnitude of multiple switching is much less in most of the existing experimental studies (5-13 percent in Holt and Laury, 2002; 5.8 percent in Andersen et al., 2006, when the *indifferent* choice is allowed, and taken by 24.3 percent of subjects). The exception is Jacobson and Petrie (2009), with over 50 percent of Rwandan subjects making at least one inconsistent choice.

with the tenet that those choices reflect people’s preferences and could potentially be rationalized by probability weighting made in the presence of large random mistakes. Our sample includes 378 experimental subjects, who made a total of 3,780 choices over monetary gains, and were surveyed to get information about personal and background characteristics that are used as controls in the empirical analysis. We report the following results.

First, on average subjects exhibit a moderate degree of relative risk aversion, which is consonant with most of the existing field experiments eliciting risk preferences. Risk estimates are somewhat similar regardless of the functional form governing the preferences (we assumed Constant Relative Risk Aversion—CRRA preferences under EUT and CPT). Moreover, the main individual characteristic that predicts risk preferences is higher education (more educated people are less risk averse); while age and gender do not play a statistically significant role in predicting risk aversion. Second, when we allow the data to be explained by EUT *and* CPT, where the proportion of subjects behaving according to each model is endogenously estimated, we find that 76 percent of the choices or observations are explained by CPT and 24 percent by EUT, a result that suggests that most of our subjects do actually make subjective distortions of the underlying probabilities in their decisions under risk. Again, both groups of subjects are risk averse under this mixture model specification.

Third, since our sample is composed of individuals with typically low levels of schooling, we could expect a large proportion of them to make mistakes in their lottery choices. Indeed, we find evidence of large random errors when subjects calculate the values of the lotteries. These random errors or mistakes can be attributed to several factors, including the lack of attention to or understanding of the rules of the game. More interestingly, we explore whether the nonlinear probability weighting that characterizes CPT provides insights to explain how those mistakes can be consistent with the MSB we observe in our sample. Contrary to what we expected, we find that CPT does not explain MSB.

Lastly, we find that our risk aversion estimates are significantly negatively correlated with higher education. Moreover, risk aversion is positively correlated with financial literacy for our entire sample, age (positively correlated), gender (female are more risk averse), and wealth (positively); however, in the last three cases, the coefficients are not statistically significant at conventional levels.

This paper contributes to the discussion about decision making under risk in developing countries. By examining two of the leading theories of decision under risk (EUT and CPT), we provide novel evidence of various risk preference estimates that will be used later (in Chapter 3) to predict financial decisions made in an experimental context.

The remainder of this paper is structured as follows. Section 2 discusses several representative related works. Section 3 reviews the experimental procedures followed in the field and the data used in the empirical analysis. Section 4 explains the analytical framework used to estimate the relevant parameters under EUT and CPT and mixing both models. Section 5 presents the estimation results for the entire sample and examines whether nonlinear probability weighting, in the presence of large

random errors, can help explain MSB. Section 6 restricts the analysis to the subset of subjects who correctly chose the higher prize lottery in the last decision row. Section 7 concludes.

2 Related Games

In recent years, field experiments in the economics of development have analyzed a wide array of topics, applying and adapting laboratory tools to investigate development problems. In a survey of the literature about experiments conducted in developing countries, Cardenas and Carpenter (2008) report that the measurement of trust, cooperation, and risk, have been three of the main topics examined by recent field experiments.

Pioneered by Binswanger’s (1980, 1981) studies in India, experimental economics methods have since been used to elicit risk preferences in several regions of the developing world. Some of the most recent experimental works implemented in developing countries include Engle-Warnick et al. (2007) in Peru, Liu (2008) in China, Tanaka et al. (2009) in Vietnam, and Jacobson and Petrie (2009) in Rwanda. We will also review the study by Holt and Laury (2002), originally run with American students, since their method has been widely used in experiments implemented in developing economies.

Risk experiments are typically run using one of the following two approaches: several choices are all given at once, from which subjects should choose *only* one; or the choices are presented using a multiple price list (MPL) design, in which subjects must choose between two prospects along several decision rows. In both of these designs, risky prospects can vary in terms of either prizes (keeping probabilities constant) or probabilities (keeping prizes constant). Now, to get risk estimates, one can simply compare the expected gains in two consecutive choices—the bounds of the risk intervals would be given by the values that yield the same value or utility in those two choices under the desired specification—or one can use econometric techniques to fit the choices data (techniques available include Ordinary Least Squares, interval regressions, and maximum likelihood methods).

Rather than describe in detail the methods used by the previous studies, we will report some of their main results. Roughly speaking, the methods used can be divided into two groups: those in which it is the probability of getting the prizes that varies, and those in which the prizes vary along several decision rows.

In Binswanger’s (1980, 1981) studies, Indian subjects were given eight alternatives (each with a bad luck outcome and a good luck outcome, both equally likely) where a higher expected value can only be obtained at the cost of having a larger standard deviation. Real and large payoffs were used,⁸ and partial risk aversion (denoted as s) is then measured by computing the indifference points between any two “efficient” alternatives.⁹ Binswanger finds that most farmers exhibit moderate

⁸The highest expected payoff in a single decision was higher than an average monthly wage for an unskilled worker.

⁹The “inefficient” alternatives had the same expected return than any two other efficient “alternatives”, but had bigger variances, which should make them unattractive to any risk-averse person. The partial risk-aversion coefficient (s) was computed using the following utility function, defined over game monetary prizes, M : $U(M) = (1 - s)M^{1-s}$.

partial risk aversion ($0.51 < s < 1.19$); and that risk aversion tends to increase with higher payoffs.

Variations of Binswanger’s procedure include comparing risky prospects to a “safe” choice, like the design used by Eckel and Grossman (2008) with American students. These authors use a five-alternative Binswanger-like design (with one choice being a sure thing), where the expected payoffs increase linearly with their standard deviation (risk). They assume Constant Relative Risk Aversion (CRRA) preferences and get risk estimate intervals by comparing two adjacent choices. These authors also provide evidence supporting moderate risk aversion. A similar version of the game was used by Engle-Warnick et al. (2007), with five risky choices (no sure thing), with Peruvian farmers. Because this experiment was run with subjects with low levels of schooling, Engle-Warnick et al. presented each of the five lotteries in circles split by half (to denote the 50/50 odds) and indicating the prizes in each half of the circle. Interestingly, the authors exploit this design to measure for *ambiguity* aversion, in which five pairs of lotteries with two prizes each are presented, one with 50/50 odds and the other with the same prizes but with unknown probabilities. Subjects were asked whether they would choose the uncertain or the risky lottery and how much they would be willing to pay to avoid facing the uncertain lotteries. This indicator was used to measure ambiguity aversion, which, unlike their risk aversion indicator (number of risky choices made), was correlated with a higher likelihood of diversifying crops.

One of the most widely cited papers on measuring risk preferences in the experimental literature is by Holt and Laury (2002) (henceforth HL). Played with American students, their ten-round lottery game allows a finer elicitation of risk preferences than in the aforementioned works. In this design, subjects are given a binary lottery (one relatively risky and the other relatively safe) in each round. In this *multiple pricing* list design, the expected value of the safe lottery is greater, but this relationship reverses as the rounds progress. Thus, even the most risk-averse player should switch before the tenth round. This switching point provides an estimate of the subject’s degree of risk aversion (the farther they switch in the rounds, the more risk averse they are), which, however, becomes very noisy when multiple switches are observed. Since there are only two choices available, only ranges of relative risk aversion can be obtained. Using CRRA preferences to compute those intervals, Holt and Laury find that even for (real) low-stake payoffs (i.e., around \$4), players are risk averse; and most importantly, that the level of risk aversion increases as the (real) stakes increase (prizes are scaled up by factors of 20, 50, and 90). Even though Harrison et al. (2005b) show that the risk aversion estimates are upward biased due to order effects (i.e., there is an increase in risk aversion when payoffs are higher that is confounded by order effects: players were “experienced” when they played the high-stake tasks), that last result still holds after controlling for those game effects.¹⁰

The previous studies consider that risk preferences are fully determined by a single parameter. However, recent experimental evidence suggests that such a picture of reality is incomplete, because psychological factors also seem to play a role in explaining preferences over the gains domain.

¹⁰Subsequent response to Harrison et al.’s critique by Holt and Laury (2005) confirms this finding.

Unlike Expected Utility Theory (EUT), where risk preferences are completely determined by the curvature of the utility function, in Cumulative Prospect Theory (CPT) such preferences are also determined by the curvature of the probability weighting function, a function that captures such probability distortions. In particular, there is substantial experimental evidence that subjects tend to overweight small probabilities and underweight large probabilities, a fact that would imply that the probability weighting function is concave for small probabilities and convex for large ones, over the gains domain (Camerer and Ho, 1994; Gonzalez and Wu, 1999; Tversky and Kahneman, 1992).

Tanaka et al. (2009) is perhaps the first study to design an experiment with gains and losses in a developing country (Vietnam in this case) to measure risk preferences. In their design, the authors enforce monotonicity in preferences by asking subjects the decision row at which they would shift from one lottery to the other (no switching back and forth was allowed). The risk game design included 3 series of decision rows (the task is to choose between two binary lotteries, A and B) with a total of 35 rows. The first set is composed of 14 rows of a HL-like design, in which only the higher prize of lottery B increases as one moves down the table and everything else (probabilities and the other prizes) is held constant; thus, while in the first row option A has higher expected value, after row 7 such a relationship reverses. In the second set of 14 rows, the design is the same as in the first set, but option A has higher expected value in every row. Finally, in the third set, each lottery (with 50/50 odds) involves gains and losses. Thus, the first two sets of rows are used to obtain intervals of risk aversion and the curvature of the probability weighting function, while the last set of rows is used to obtain the intervals of the loss aversion parameter. These authors find that their Vietnamese subjects are risk averse (and risk aversion is negatively correlated with education) and that they overweight small probabilities and underweight larger probabilities. In a similar study, Liu (2008) replicates Tanaka et al.'s method and provides evidence confirming risk aversion and probability weighting under PT in her sample of Chinese cotton farmers. Liu further finds that risk averse subjects are less prone to have adopted *Bt* cotton.

An alternative experimental method to measure risk aversion uses bidding mechanisms. Becker-DeGroot-Marschak's (Becker et al., 1964) (henceforth BDM) procedure allows elicitation of certainty equivalent values of lotteries (CE) in two stages. In the first stage, a subject is endowed with a lottery gain prospect that offers a prize G with probability q . In the second stage, the subject is asked to set a minimum selling price, wtp_i , for that lottery. Then, a buying price (p^b) for that lottery is randomly drawn. If the buying price exceeds the selling price ($p^b > wtp$), the subject gets p^b ; otherwise, he has to play the lottery. Risk aversion coefficients can then be measured for each response comparing the CE s and the expected values of the lotteries (prizes).¹¹ Varying the lottery prizes, the sensitivity of risk aversion to scale effects can further be examined. While BDM may be used to elicit true valuation provided that the independence axiom of EUT holds, the risk aversion estimated under this procedure has been reported to be very sensitive to the experimental conditions (Harrison and Rutström, 2008). Kachelmeir and Shehata's (1992) paper using the BDM

¹¹More clearly, risk aversion coefficients under CRRA preferences would be determined by solving $1 - [\frac{\ln q}{\ln CE - \ln G}]$.

procedure in their study for China illustrate some of the difficulties faced using this method to measure risk aversion.¹²

While a substantial effort has been spent to use different techniques to estimate risk (and time) preferences, little is still known about the relationship between economic preferences and cognitive skills. Despite this, a few remarkable results seem to consistently emerge: in general, such preferences are significantly correlated with cognitive abilities, in the sense that subjects with better cognitive skills tend to be more likely to take risks (Benjamin et al., 2006; Burks et al., 2009; Dohmen et al., 2007; Frederick, 2005). In the same vein, little research has been done to analyze the correlation between elicited preferences and economic decisions, an issue that is crucial in development economics. One of these works, focused on a developing country (Rwanda), is by Jacobson and Petrie (2009), who examine whether risk preferences are correlated with financial decisions made in real life, and the role that randomness plays in risky decisions. In their 5-row lottery game, the authors find that risk averse subjects are more likely to taking out a formal loan. However, once they control for the probability of making mistakes in the lottery game, risk aversion also turns significant as a predictor of being member of a savings group (positive sign) and taking out informal loans (negative sign). Further, for those who are more likely to make mistakes, such relationship is reversed. This result, the authors maintain, implies that those who are more likely to make mistakes also tend not to choose optimally or are excluded from savings groups. In another work, Guiso and Paiella (2007) find that their (survey) measure of risk aversion is correlated with a lower probability to move from one region to another, to change a job, and even to have a chronic disease.

In this paper, we adopt the approach of fitting several functional forms to our risk data with the expectation that once we unravel true preferences, they will be correlated with real life economic decisions. The analysis of these issues is deferred to Chapter 3, which examines the link between risk preferences and financial decisions in Peru.

3 Experimental Procedures

This section describes the sampling design, reports the main characteristics of our experiment participants, and discusses the experimental procedures followed in our *artefactual* field risk experiment (using the terminology coined by Harrison and List, 2004); that is, we designed a conventional laboratory experiment, which was conducted with farmers in Peru, a nonstandard subject pool. Chapter 3 will discuss the results of a different experiment, a *framed* field experiment, conducted with the same subjects.

¹²Other indicators of experimental risk aversion may be the amount bet in a gamble with varying amounts of bets and expected returns, one choice being not to gamble at all. Schechter’s (2007) risk game in Paraguay uses this method and finds an average lower bound of 0.8 when defining utility only over the risk game winnings, and an average lower bound of about 2 when utility is defined over income and winnings from the risk game.

3.1 Recruitment of Subjects

The experimental sessions were held in several zones of the Province of Pisco, a valley located in the Department of Ica, in coastal Peru.¹³ A map of the Pisco valley is shown on the left panel of Appendix A, while a map of Peru is displayed on the right panel. With a total of 3,600 producers, owning 24,000 hectares, the agricultural production in Pisco is heavily concentrated on a single crop: cotton. Cotton production concentrates about 45 percent of the total sown area in Pisco (about 11,000 hectares), and about 77 percent of all Pisco agricultural producers grow cotton. The average size of a cotton parcel is 3.8 hectares, and the typical farm comprises 6.6 hectares in total (figures as of 2007-2008). This small-scale production is mostly the result of the Peruvian Agrarian Reform carried out in the nineteen seventies, where a military government expropriated the land from the landlords and redistributed it to farmers. Agricultural production is greatly concentrated in three Districts—Independencia, Humay, and San Clemente—as shown in Appendix A, where the dark area on the left panel depicts the agriculture parcels.

We conducted 24 experimental sessions in 12 different locations across the Pisco valley. These sessions were held in locations where electricity was available and enough room to host 25-30 persons was ensured. In all of the cases, farmers were familiar with the locations chosen—public schools, private houses, a church eatery, and municipal auditoriums. Whenever possible, we found locations that were focal points to the majority of the selected farmers. However, we had to hold several sessions in locations with very scattered households, and with very little or no access to public transportation, where the two conditions mentioned earlier—electricity and space—were met.¹⁴

Given that surface water for irrigation is a crucial factor in Pisco, the agricultural area is divided into 42 Irrigation Blocks, which are overseen by 20 Water Irrigation Commissions (WICs). Every WIC controls irrigation water administration and use, which is closely supervised by the *Junta de Usuarios*, an entity that operates as the superintendent of water administration in the valley. We relied on information from the Junta de Usuarios of Pisco to implement our sampling procedure. We selected our subjects using two-stage stratified sampling. In the first stage, on the basis of 40 Irrigation Blocks (two were dropped because of their small acreage), we constructed 23 conglomerates with the following criteria: to have a minimum of 600 hectares and a maximum of 1,500 hectares of cotton in total, and to be geographically adjacent. We then randomly chose 13 conglomerates, having a total of 1,604 farmers. In the second stage, to carry out the individual-level randomization, nearest-neighbor farmers with respect to farm size were broken down into pairs within each of those 13 conglomerates, and one member of each pair was randomly chosen to participate in our experiment. The final sample size consisted of 804 farmers, spread along 12 WICs.

Invitations to attend the experimental sessions were sent out for all of those farmers, 745 of

¹³The political division of Peru includes: regions, departments (akin to a U.S. state), provinces, and districts. The Peruvian territory comprises 24 departments and a Constitutional Province, Callao (our main port).

¹⁴In one (extreme) case, a farmer told us he walked for 45 minutes (one way) to get to the session. The only means to get to his house was by riding a horse or a motorcycle.

whom were reached by our messengers.¹⁵ In order to deliver the invitations, we hired the *sectoristas de riego*, persons in charge of the water administration, in each WIC because of their familiarity with farmers (where they have parcels and live). Invitations were personalized and included information about the experimental sessions (date, location, time), the promise of 7 Soles for just showing up,¹⁶ and an estimate of their winnings for participating in the sessions (between 10 and 30 Soles, or between \$3.6 and \$10.7). Farmers were also told to bring their invitations to the session as a way of checking their identity. We had 410 experimental subjects come to the sessions, 399 of whom stayed until the end of the risk games.¹⁷ The core of our analysis will be based on 378 subjects for whom we have information on most of the variables examined.

The entire experimental sessions were conducted in Spanish. Each experimental session was composed of six parts, conducted in the following order: entry survey, farming games (the design and results from these games are discussed in Galarza [2009]), *risk game*, exit survey, a short video produced by the insurance company advertising the new insurance product, and a brief lecture on valley-wide insurance (*charla*). During the *charla*, questions about any unclear issues during the sessions were answered.

3.2 Characteristics of Participants

Participants in our experimental sessions are on average older than 50, mostly male, and have typically completed only elementary school (which takes six years in Peru), as shown in Table B.1 in the Appendix. In terms of their farming activity, almost all of them own the parcels they work on, and the typical farmer reports owning almost 6 hectares in total, a figure that is slightly smaller than the average size of the farm reported at the valley level (6.6 hectares). The size of the cultivated land is 5 hectares. Moreover, 82 percent of experimental subjects planted cotton in the last farming season (2007-2008), obtaining an average yield of 47 quintals (or 2,162 Kilograms) per hectare.¹⁸ Subjects also report extensive experience managing their own agricultural parcels (an average of 24 years). Compared to subjects who did *not* attend but were invited to the sessions, participants do not have statistically significant differences in terms of the total owned area, total sown area, or area sown with cotton,¹⁹ a result that provides evidence that no sample selection

¹⁵The remaining 59 subjects were not at home when our messengers repeatedly visited their households, because they were working far from home and could not be reached, or they had simply moved out of town.

¹⁶In most of the cases, this fee was sufficient to pay for round-trip travel from subjects' houses to the locations where sessions were held. In some instances, however, no public transportation was available, and producers could only ride their own motorcycles, horses, or walk.

¹⁷Explanations for this seemingly low participation rate include: scattered households, poor explanation of the incentives for participating, higher opportunity cost (probably non-pecuniary), unclear explanation or comprehension of the benefits from attending the sessions. Farmers in this valley have low participation rates for any meeting organized, even by the same WIC, excluding the ones related to irrigation water use.

¹⁸According to official statistics, the valley-wide average yield in 2005 and 2006 was 47 and 42 quintals per hectare, respectively. No official statistics for 2007 were reported at the time of writing this paper. One quintal is equivalent to 46 Kilograms.

¹⁹These are the only variables available to look for selection in the sample who showed up for the experiments. This comparison was done using information from the Junta de Usuarios de Riego de Pisco for the 2007-2008 farming

would seem to exist.

Our sample is mainly composed of poor farmers, judging by the self-reported value of their assets (20,000 Soles, or roughly US\$7,000, including house and land, but only 6,900 Soles considering land alone²⁰). With 80 percent of the subjects being owners of the houses they live in, the average value of a house is 16,300 Soles. In addition to their lack of access to formal insurance mechanisms, individuals in our sample report being widely exposed to external shocks that affect their agricultural production, as well as to individual shocks (mainly injuries or illnesses), and robbery, that affected their farming activities. Furthermore, the majority of subjects (61 percent) report having access to credit markets. This access to credit is concentrated in the formal sector (39 percent), followed by the informal sector (34 percent), and the cotton gins (27 percent).

In terms of their participation in local organizations, 29 percent of the respondents belong to a farmer's association and 39 percent have ever played a role as some type of local authority (this includes any participation in the Water Irrigation Commission). However, more than reflecting a selection of more active people in our sample, this relatively great presence of local leaders could simply reflect the high degree of involvement of farmers in their communities, since this variable considers any type of involvement in local organizations.

We further obtained information about how informed farmers were about any type of insurance before starting the experiments and how much they learned during the experiments. The typical participant in our experiments has only a basic knowledge of any type of insurance. We will get back to this subject in Chapter 3, when we will examine information contained in the bottom panel of Table B.1.

3.3 Experimental Sessions

In all of our 24 conducted sessions, participants were assigned to numbered seats at random upon arrival, and received a binder containing the experiment worksheets and a pencil to record their choices. We divided the participants into a maximum of four "valleys" with a minimum of 3 members in each one. Splitting the experimental subjects into several valleys allowed us to have closer monitoring and to accelerate the tasks. Two persons were in charge of each valley. A senior assistant, well versed in the game rules and procedures, recorded the choices of players, and a helper assisted with the implementation of the randomizing device used to pick the decision for play (dice rolling) in each valley. The experiment instructions were read aloud to all participants as a group. To ensure that farmers understood the mechanics and rules of the games, we allowed them to ask questions during the presentation of the instructions. Game winnings and attendance fees were paid at the end of the entire experimental session, which lasted on average five hours. The total game winnings from both types of experiments ranged between 11 and 30 Soles, with an average of 20 Soles. These winnings compare well with a daily unskilled wage (*jornal*) of 15 to 20 Soles at

season.

²⁰Farmers were asked to self-report their rental value of land, which can be considered a lower bound of the land value.

the time of running the games. Winnings from the risk game, which lasted an average of half hour, averaged 3 Soles, with a minimum of 0.15 Soles and a maximum of 5.75 Soles.

3.4 The Risk Game

We used a relatively simple game to measure risk aversion introduced by Holt and Laury—HL (2002), in which players chose between a relatively safe lottery (which we call option *Sol*) and a relatively risky lottery (which we call option *Luna*) along ten decision rows. In this design, lotteries' characteristics in row t are as follows: *Sol*: $(t/10, 1800; 1400)$, and *Luna*: $(t/10, 3500; 90)$. That is, as shown in Table 1 below, in the first row (i.e., for $t = 1$), subjects choosing lottery *Sol* have a 10 percent chance ($prob = 1/10$) of getting 1,800 Soles and a 90 percent ($[1 - prob] = 9/10$) chance of getting 1,400 Soles. Similarly, if they choose lottery *Luna*, there is a 10 percent chance of getting 3,500 Soles, and a 90 percent chance of getting 90 Soles. In the second row, there is a 20 percent chance of getting the higher prize in each lottery, and so on.

Table 1: Matrix Payoff in the Risk Experiment

Row	Sol				Luna				$EV^S - EV^L$	CRRA interval if switches to L	Risk Preference Classification
	p	Prize	1-p	Prize	p	Prize	1-p	Prize			
1	0.1	1800	0.9	1400	0.1	3500	0.9	90	1009	$-\infty, -1.61$	Extremely RL ¹
2	0.2	1800	0.8	1400	0.2	3500	0.8	90	708	-1.61, -0.88	Highly RL
3	0.3	1800	0.7	1400	0.3	3500	0.7	90	407	-0.88, -0.43	Very RL
4	0.4	1800	0.6	1400	0.4	3500	0.6	90	106	-0.43, -0.10	RL
5	0.5	1800	0.5	1400	0.5	3500	0.5	90	-195	-0.10, 0.18	RN ²
6	0.6	1800	0.4	1400	0.6	3500	0.4	90	-496	0.18, 0.44	Slightly RA ³
7	0.7	1800	0.3	1400	0.7	3500	0.3	90	-797	0.44, 0.69	RA
8	0.8	1800	0.2	1400	0.8	3500	0.2	90	-1098	0.69, 0.98	Very RA
9	0.9	1800	0.1	1400	0.9	3500	0.1	90	-1399	0.98, 1.38	Highly RA
10	1.0	1800	0.0	1400	1.0	3500	0.0	90	-1700	1.38, $+\infty$	Extremely RA

Note: The last four columns were not shown to experimental subjects.

¹ RL: risk loving; ² RN: risk neutral; ³ RA: risk averse. This classification is based on the mid-CRRA intervals.

Note that in this design prizes are held constant across the decision rows and we vary only the probabilities of the higher and smaller prizes in each row. Also, the probabilities of each prize in a given row are the same for both lotteries, so that subjects would focus on the changes in probabilities (the higher prize in each lottery increased its probability) across rows. As a result, the difference in the expected values of the lotteries decreases as subjects move down in the decision rows, as shown in column 10 of Table 1, where we also see the risk intervals associated with switches made at every decision row (columns 11-12), as well as the risk categories associated with each interval (column 13), which is based on the mid-point interval in each row. Thus, only risk loving subjects would

choose the lottery Luna in the first decision rows, while only risk averse subjects would choose the lottery Sol in the last decision rows. In turn, risk neutral subjects would switch from choosing lottery Sol to lottery Luna in row 5, when the expected values of both lotteries are about the same. This switching point from the safe to the risky lottery provides an estimate of subjects' degree of risk aversion (the farther they switch in the rounds, the higher the risk aversion, as shown above). In this design, subjects *could* start by choosing the lottery Luna in the first row (if they were highly risk seeking) and stick to it until row ten (in which case they would be infinitely risk seeking), or switch to lottery Sol before that. Note that this design does not prevents switching back and forth (from Sol to Luna or viceversa).²¹

The risk game was implemented as follows (the experiment instructions are provided in Appendix C): we first showed subjects the prizes associated with each lottery, the task involved (i.e., choose one of those lotteries along ten decision rows), and the way they could win those prizes. We next showed the prizes in decision row 2, putting emphasis on the probabilities associated with each of the prizes for *both* lotteries. Then, we showed the prizes in decision row 8 and proceeded likewise. We then displayed rows 2 and 8 together in order to show the symmetry in probabilities of the bigger and smaller prizes: while in row 2, there is a 20 percent chance of getting the higher prize in each lottery, in row 8 such odds are 80 percent. Figure 1 shows the slide shared with our experimental subjects at this point.

Figure 1: Risk Experiment: Characteristics of Rows 2 and 8

	SOL 	LUNA 																																																												
2	 <table><tr><td>1</td><td>2</td></tr><tr><td colspan="2">S/.</td></tr><tr><td colspan="2">1800</td></tr></table><table><tr><td>3</td><td>4</td><td>5</td><td>6</td><td>7</td><td>8</td><td>9</td><td>10</td></tr><tr><td colspan="8">S/.</td></tr><tr><td colspan="8">1400</td></tr></table>	1	2	S/.		1800		3	4	5	6	7	8	9	10	S/.								1400								 <table><tr><td>1</td><td>2</td></tr><tr><td colspan="2">S/.</td></tr><tr><td colspan="2">3500</td></tr></table><table><tr><td>3</td><td>4</td><td>5</td><td>6</td><td>7</td><td>8</td><td>9</td><td>10</td></tr><tr><td colspan="8">S/.</td></tr><tr><td colspan="8">90</td></tr></table>	1	2	S/.		3500		3	4	5	6	7	8	9	10	S/.								90							
1	2																																																													
S/.																																																														
1800																																																														
3	4	5	6	7	8	9	10																																																							
S/.																																																														
1400																																																														
1	2																																																													
S/.																																																														
3500																																																														
3	4	5	6	7	8	9	10																																																							
S/.																																																														
90																																																														
8	 <table><tr><td>1</td><td>2</td><td>3</td><td>4</td><td>5</td><td>6</td><td>7</td><td>8</td></tr><tr><td colspan="8">S/.</td></tr><tr><td colspan="8">1800</td></tr></table><table><tr><td>9</td><td>10</td></tr><tr><td colspan="2">S/.</td></tr><tr><td colspan="2">1400</td></tr></table>	1	2	3	4	5	6	7	8	S/.								1800								9	10	S/.		1400		 <table><tr><td>1</td><td>2</td><td>3</td><td>4</td><td>5</td><td>6</td><td>7</td><td>8</td></tr><tr><td colspan="8">S/.</td></tr><tr><td colspan="8">3500</td></tr></table><table><tr><td>9</td><td>10</td></tr><tr><td colspan="2">S/.</td></tr><tr><td colspan="2">90</td></tr></table>	1	2	3	4	5	6	7	8	S/.								3500								9	10	S/.		90	
1	2	3	4	5	6	7	8																																																							
S/.																																																														
1800																																																														
9	10																																																													
S/.																																																														
1400																																																														
1	2	3	4	5	6	7	8																																																							
S/.																																																														
3500																																																														
9	10																																																													
S/.																																																														
90																																																														

We next projected a slide containing all the decision rows, and showed subjects the pattern of increasing probability of getting the higher payoffs and the resulting decreasing probability of getting the lower prizes in both lotteries as one goes down in the table. The explanation of the game ended with a mention of the last row, when the monitor said that subjects will get the higher prize for certain in each lottery, so that the choice would be between 1,800 if they choose Sol and

²¹Two other critiques to this procedure follow (Harrison et al., 2005a): it only elicits intervals of risk aversion, and it can be vulnerable to framing issues, since subjects may be drawn to some *focal* choice picking (e.g., switching at the middle of the table). Although refinements have been suggested to overcome those potential obstacles—offering more choices to subjects within each interval, and randomizing the order of the lottery choices—such methodological improvements may have come at a high cost: most likely, the resulting higher complexity would have overwhelmed our subjects.

3,500 if they choose Luna. In Appendix D we provide a sample worksheet used in the experiment.

To sum up, the main factors we asked farmers to consider in making decisions were the minimum and maximum payoffs within each lottery, and the likelihood of those payoffs in each decision row. Also, by showing them all the decision rows rather than one by one sequentially, we wanted them to see the decreasing probability pattern of the lower payoff in each lottery. Despite the fact that we explained this, a large proportion of subjects (105 out of 378) chose the safe lottery in the last round, which clearly denotes lack of attention or understanding of the game. These mistakes could in turn be due to fatigue, as the risk game was played after an intensive section of farming experiments.

After the explanation of the game procedures and rules, subjects played a practice game, in which they selected their preferred lotteries along ten decision rows, and learned how to calculate their winnings in real Soles, but did not earn money in cash. They thus learned that their winnings would be determined by only one decision row, chosen at random in each valley by rolling a ten-sided die (row for play), and that their specific winnings in game Soles would correspond to the prize associated in the option chosen in the row for play with the number resulting in a second, individual die roll. To make it clear, if, for example, the first die roll for the valley of certain subject landed on the number 5 (i.e., the fifth row will be played), and if this subject's second die roll landed on the number 6, she would win 1,400 game Soles if she chose lottery Sol in row 5 and 90 game Soles if she selected lottery Luna. Immediately after practicing the game, subjects played the game *for real*, following the same rules, and having an exchange rate of 1 Sol in cash for every 600 game Soles.

4 Structural Estimation of Risk Parameters

In this section, we will describe the estimation procedures used to measure risk preferences, first assuming that the data are entirely generated by one model (either Expected Utility Theory [EUT] or Cumulative Prospect Theory [CPT]), and then using a mixture model that allows to estimate simultaneously the risk parameters under each model, in addition to the proportion of subjects who are best described by EUT and by CPT. The estimation is done using the maximum likelihood method,²² and in every case we will correct the standard errors for clustering at the individual level, in order to account for the possibility that choices made by the same individual are correlated across decision rows. We will analyze 378 subjects' decisions made over monetary gains along 10 decision rows, which makes a total of 3,780 choices.²³

²²The estimation was done in STATA. The algorithm used was the BFGS (Broyden, Fletcher, Goldfarb, Shanno). The codes were graciously shared by Glenn Harrison from the University of Central Florida (UCF).

²³However, when we include individual characteristics as covariates, the sample size shrinks to 365 subjects with 3,650 choices, due to missing information.

4.1 Assuming Expected Utility Theory

As is typical under EUT, we will assume that the utility of income from outcome $j \in \{1, 2\}$ in lottery $k \in \{Sol(S), Luna(L)\}$ that individual $i \in \{1, \dots, N\}$ gets, denoted by $M_i^{k,j}$, is defined by the following CRRA preferences:

$$U(M_i^{k,j}) = \frac{(M_i^{k,j})^{1-r_i^{EU}}}{1-r_i^{EU}}, \quad r_i^{EU} \neq 1, \quad (1)$$

Note that since the prizes are constant across rows in every lottery (see Table 1 above), no row index is needed for M_i . In this specification, risk aversion is completely determined by the curvature parameter, r_i^{EU} , with $r_i^{EU} = 0$ denoting risk neutrality; $r_i^{EU} > 0$, risk aversion; and $r_i^{EU} < 0$, risk seeking behavior. Recall that in our risk game, each lottery k in row m has two possible outcomes, each with probability p_m^j . Then, when confronted with a binary lottery, farmers are assumed to make the following expected utility (EU) calculation at every decision row, m :

$$EU_{i,m}^k = \sum_{j=1}^2 p_m^j * U(M_i^{k,j}), \quad (2)$$

and to choose, either lottery *Sol* or *Luna*, according to the value of the following latent index or choice rule:

$$\Delta EU_{i,m} = EU_{i,m}^S - EU_{i,m}^L + \mu_i, \quad (3)$$

where μ_i denotes the errors made by subject i (as a result of carelessness, hurry, or insufficient motivation) in the process of calculating the expected utilities,²⁴ which will be assumed to be white noise (i.e., with mean zero²⁵ and constant variance). This additive error was first proposed by Fechner (1860/1966), and we will refer to them as “Fechner errors.”²⁶ The prior function can be interpreted as the *perceived* advantage of lottery *S* over lottery *L*, while such function excluding the error term represents the *true* advantage of lottery *S* over lottery *L*. Appealing to the Central Limit Theorem, we can assume that μ_i is also normally distributed:

$$\mu_i \sim N(0, \sigma_{\mu_i}^2). \quad (4)$$

We will further assume that errors are uncorrelated across decision rows. In this Fechner error story, a careful individual i would have a relatively small error or noise,²⁷ represented by a small standard deviation (σ_{μ_i}), in her decisions. On the other hand, when σ_{μ_i} becomes large, her decision

²⁴For a discussion about the different stages at which randomness can play a role in the decision making in lotteries, see Loomes et al. (2002).

²⁵This means that respondents are assumed not to have left or right bias in their answers.

²⁶This error type was also used by Hey and Orme (1994). A popular alternative error type, due to Luce (1959), is succinctly examined in Appendix F. Preliminary results suggest that the main qualitative results hold assuming either Fechner errors or Luce errors.

²⁷The terms *noise*, *error*, and *mistake* are used interchangeably throughout the text.

would respond less to the differences in subjective values and more to randomness.

For estimation purposes, under the assumption made in eqn.[4], we will use a probit linking function to transform the latent index given by eqn.[3] (which has a value between $\pm\infty$) into a binary variable that denotes the observed choices. Thus, the probability that S is chosen over L will be:

$$\begin{aligned} \Pr(EU_{i,m}^S - EU_{i,m}^L + \mu_i > 0) &= \Pr\left(\frac{\mu_i}{\sigma_{\mu_i}} > -\frac{EU_{i,m}^S - EU_{i,m}^L}{\sigma_{\mu_i}}\right) \\ &= 1 - \Phi\left(-\frac{EU_{i,m}^S - EU_{i,m}^L}{\sigma_{\mu_i}}\right) = \Phi\left(\frac{EU_{i,m}^S - EU_{i,m}^L}{\sigma_{\mu_i}}\right), \end{aligned} \quad (5)$$

where the expression in the last parenthesis has a standard normal distribution, and will be referred to as the stochastic expected utility indicator:

$$\Delta SEU_{i,m} = \frac{EU_{i,m}^S - EU_{i,m}^L}{\sigma_{\mu_i}}. \quad (6)$$

Thus, $\Phi(\Delta SEU_{i,m})$ denotes the probability of choosing lottery S , and $[1 - \Phi(\Delta SEU_{i,m})]$ represents the probability of choosing lottery L . By the symmetry of the normal distribution, the last expression is equivalent to $\Phi(-\Delta SEU_{i,m})$. To be clear, lottery S would be chosen when $\Delta SEU_{i,m} > 0$ or $\Phi(\Delta SEU_{i,m}) > 0.5$; otherwise, lottery L would be selected. We will use the prior eqn. in the optimization routine to estimate the risk parameter, r_i , and the standard deviation of the noise, σ_{μ_i} . Then, individual i 's contribution to the model's likelihood can be written as:

$$L_i^{\text{EU}}(r_i^{\text{EU}}, \sigma_{\mu_i}; I_i^m, X_i) = \prod_{m=1}^{10} [\Phi(\Delta SEU_{i,m})^{I_{i,m}}] * [\Phi(-\Delta SEU_{i,m})^{1-I_{i,m}}], \quad (7)$$

where $I_{i,m}$ is an indicator variable of the choice made by individual i in row $m \in \{1, \dots, 10\}$, which takes the value of 1 when lottery S is chosen in row m , and 0 otherwise. X_i is a vector of individual i 's characteristics. The model's log-likelihood would then be the log of the product of L_i^{EU} over all the individuals i :

$$l^{\text{EU}} = \sum_{i=1}^N \ln L_i^{\text{EU}}(r_i^{\text{EU}}, \sigma_{\mu_i}; I_i^m, X_i) = \sum_{i=1}^N \sum_{m=1}^{10} [\ln(\Phi(\Delta SEU_{i,m})^{I_{i,m}}) + \ln(\Phi(-\Delta SEU_{i,m})^{1-I_{i,m}})]. \quad (8)$$

Note that the inclusion of a vector of covariates (X_i) in the likelihood functions above allows the estimation of subject specific parameters (including the standard deviation of the error, σ_{μ_i}), where we consider the parameter to be a linear function of the covariates. Doing so permits to unveil the existence of heterogeneity at the parameter level, in contrast to other sources of heterogeneity, such as heterogeneity at the preference functional level (which is analyzed in section 4.3, by estimating finite mixture models). Obviously, if we do not include such covariates in the parameter function,

we will simply estimate an aggregate parameter for the entire sample.

In the case of the risk parameter, we will include covariates to estimate subject-specific parameters whenever it is possible in this and the other specifications considered later on (prospect theory and mixture models). In particular, r_i will be estimated as a linear function of dummy variables of gender, age, education, and geographic location (we label this as the *heterogeneous agents* case): $\hat{r}_i = \hat{r}(X_i)$:

$$\hat{r}_i = \hat{r}_0 + \hat{r}_{age} * Age_i + \hat{r}_{gender} * Gender_i + \hat{r}_{edu} * Education_i + \hat{r}_{geog} * Geog.Location_i, \quad (9)$$

which will allow us to examine how idiosyncratic risk preferences are. Clearly, not including individual characteristics would yield estimates only for a representative subject, which we label as the *homogeneous agent* case: $\hat{r}_i = \hat{r}_0$. Similarly, we will estimate heteroskedastic random errors below, but with a restricted set of individual characteristics. Similarly, when no covariates are included in the regression, we would be estimating an aggregate standard deviation of the errors for the entire sample: $\hat{\sigma}_{\mu_i} = \hat{\sigma}_0$.

4.2 Assuming Cumulative Prospect Theory

As mentioned in the Introduction of this paper, unlike EUT that can be entirely defined by the curvature parameter, Tversky and Kahneman (TK)'s (1992) Cumulative Prospect Theory (CPT) adds two psychological features: the notion that losses are more heavily felt than gains of similar magnitude (loss aversion); and the notion that subjects make subjective assessments that distort the probabilities of lotteries in evaluating prospects (nonlinear probability weighting). Since our risk game considers only gains (not losses), we can only test the existence of risk aversion *and* whether subjects underweight and overweight probabilities in making risky choices. Further note that the utility function, renamed by TK as *value function*, is defined over *gains* and *losses* from the lottery game; not over terminal wealth. Gains and losses are, in turn, defined with respect to a *reference point*, which is usually assumed to be the status quo, or the current level of wealth. We will adopt this approach, thus defining gains as a situation better than the *status quo* and losses, as a situation that is less favorable than the *status quo*. Further, we will continue to assume CRRA preferences, with r^{PT} denoting the risk parameter, for ease of comparison with the EUT risk parameter, r^{EU} :

$$U(M_i^{k,j}) = \frac{(M_i^{k,j})^{1-r_i^{PT}}}{1-r_i^{PT}}, \quad r_i^{PT} \neq 1. \quad (10)$$

In CPT, instead of probabilities, we have *decision weights*, $\mathbf{dw}^j(p)$, associated with each of the two outcomes in lottery k . Defined over the cumulative probabilities,²⁸ these decision weights

²⁸Defining the decision weights over the cumulative density instead of the probability outcomes helped CPT avoid the stochastic dominance problem that affected the original Kahneman and Tversky's (KT) (1979) Prospect Theory (PT). Under the original PT, subjects could choose a stochastically dominated lottery.

reflect the subjective distortion of probabilities that has been found by several previous studies (e.g., Camerer and Ho, 1994; Gonzalez and Wu, 1999), and which can explain why the same subjects can be risk averse over some prospects (i.e., buying insurance), while being risk loving over some others (e.g., gambling). Thus, when given a binary lottery k to choose, subjects are assumed to make the following *Cumulative Prospect Utility* (CPU) calculation:

$$CPU_{i,m}^k = \sum_{j=1}^2 \mathbf{dw}_m^j(p) * U(M_i^{k,j}) = \mathbf{dw}_m^1(p)U(M_i^{k,1}) + \mathbf{dw}_m^2(p)U(M_i^{k,2}), \quad (11)$$

where $\mathbf{dw}_m^j(p)$ at every decision row is given by:

$$\mathbf{dw}_m^j(p) = \begin{cases} 1 - \mathbf{w}(p_m^2), & \text{for } j = 1 \\ \mathbf{w}(p_m^2), & \text{for } j = 2 \end{cases}, \quad (12)$$

and the weighting function, $\mathbf{w}(p_m^j)$, will be represented by:

$$\mathbf{w}(p_m^j) = \frac{(p_m^j)^\gamma}{\left[(p_m^j)^\gamma + (1 - p_m^j)^\gamma\right]^{1/\gamma}}, \quad \gamma > 0 \quad (j = 1, 2), \quad (13)$$

where $\mathbf{w}(0) = 0$ and $\mathbf{w}(1) = 1$. This one-parameter function was proposed by Quiggin (1982) and popularized by TK (1992). Looking at eqns.[11] and [12] we see that CPU^k is the sum of two rank-dependent outcomes, where different weights are given to different utilities of outcomes. Note also that the decision weights in eqn.[12] add up to one, since this function is defined over the *cumulative* density function. This specification overcomes the potential problem that nonlinear weighting can cause violations of stochastic dominance (Fox and Poldrack, 2009). Further note that under Cumulative Prospect Theory, there are *two* sources of risk aversion: the curvature of the utility function, defined by r^{PT} , and the parameter of the probability weighting function, γ (Tversky and Kahneman, 1992). Thus, for a given curvature of the value function, risk aversion is reinforced by underweighting of middle to large probabilities and offset by overweighting of small probabilities (Fox and Poldrack, 2009). We will explain what we mean by overweighting and underweighting next, when we examine the different shapes that TK's weighting function can yield, depending on the value of the curvature and elevation parameter, γ :²⁹

- (i) If $0 < \gamma < 1$, then $\mathbf{w}(p)$ would have an inverse S-shape, implying that subjects overweight small probabilities (concave section, where $\mathbf{w}(p) > p$) and underweight large probabilities (convex section, where $\mathbf{w}(p) < p$).³⁰ Examples of this case are depicted by the solid and dot-dashed

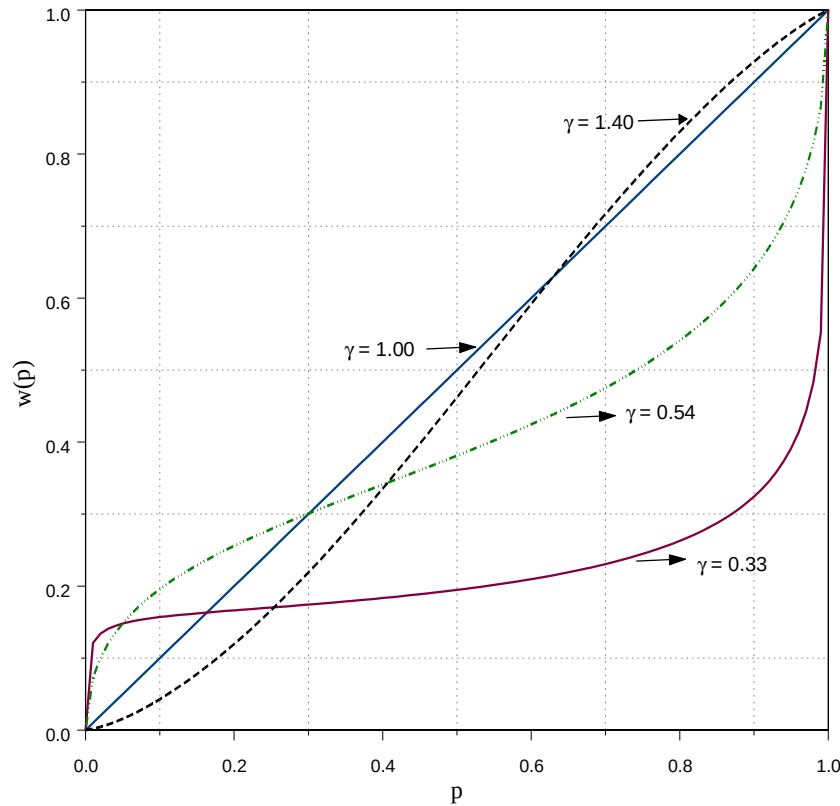
²⁹The notion of elevation in the weighting functions becomes relevant with the estimation of two-parameter functions. In such case, elevation refers to the attractiveness to gambling (or a higher weight assigned to the larger prize in our context). For more details, see Gonzalez and Wu (1999). Popular weighting functions include Rieger and Wang's (2006), Prelec's (1998), and Lattimore et al.'s (1992). We will introduce Prelec's function in Appendix F.

³⁰Underweighting (overweighting) happens when subjects behave as if the chances of occurring a given outcome

curves in Figure 2.

- (ii) If $\gamma = 1$, then $\mathbf{w}(p) = p$, and we would be back to the EU framework with linear probabilities; this case is represented by the 45° degree line in Figure 2.
- (iii) If $\gamma > 1$, then $\mathbf{w}(p)$ will have a S-shape, with convexity for small and moderate probabilities and concavity for larger probabilities. An example of this case is depicted by the dashed curve in Figure 2.

Figure 2: Tversky and Kahneman's (1992) Subjective Weighting Function



While TK's (1992) weighting function has shown to fit well the data by several empirical studies, it is not without drawbacks. In particular, it is not increasing in p for small values of γ ,³¹ and it does not have axiomatic foundations. While the former limitation does not represent a practical problem for us (since we generally find estimates of γ greater than 0.3, as we will see in the next section), the latter implies that TK's function, as well as other *ad hoc* functions (such as Lattimore

are lower (greater) than they actually are.

³¹Simulations show that such a function is partially decreasing for $\gamma \leq 0.278$, although for greater values of γ , such a problem does not seem to exist (see Rieger and Wang, 2006).

et al.'s [1992]), could not fit data from subjects who reduce simple compound lotteries (Luce, 2000).³²

Similar to the EUT case, in order to get estimates of the risk and weighting parameters under CPT, we will maximize the following log-likelihood function:

$$l^{\text{PT}} = \sum_{i=1}^N \ln L_i^{\text{PT}}(r_i^{\text{PT}}, \gamma_i, \sigma_{\mu_i}; I_i^m, X_i) = \sum_{i=1}^N \underbrace{\left\{ \sum_{m=1}^{10} [\ln(\Phi(\Delta SPU_{i,m})^{I_i^m}) + \ln(\Phi(-\Delta SPU_{i,m})^{1-I_i^m})] \right\}}_{\ln L_i^{\text{PT}}(\cdot)}, \quad (14)$$

where the new choice rule, $\Delta CPU_{i,m}$, is given by the sign of the following expression:

$$\Delta CPU_{i,m} = CPU_{i,m}^S - CPU_{i,m}^L + \mu_i, \quad (15)$$

and the stochastic prospect utility indicator, SPU , used in the maximization algorithm, can be written as:

$$\Delta SPU_{i,m} = \frac{CPU_{i,m}^S - CPU_{i,m}^L}{\sigma_{\mu_i}}. \quad (16)$$

To account for individual heterogeneity, the parameters r_i^{PT} and γ_i will be estimated, whenever is possible, as linear functions of the individual characteristics, X_i , using the specification shown in eqn.[9]. We will proceed similarly with the estimation of σ_{μ_i} when we consider the heteroskedastic case.

4.3 Mixture Model Specification

In this section, we will move one step forward towards a more accurate depiction of the true but unknown underlying risk preferences, by allowing for heterogeneity at the functional form level. So far, we have assumed that a single decision model, either EUT or CPT, explains the preferences of *all* subjects. In other words, we assume that all the observations are explained by the same underlying mechanism or decision rule. However, it is plausible to think that different groups of subjects may exhibit distinct risk preferences (i.e., they may be explained by different underlying mechanisms). Whether we can distinguish heterogeneity in risk preferences in our sample is an empirical question that can be addressed via the estimation of *finite mixture models*.³⁴ While we could endogenously estimate the number of groups that behaves according to different models,³⁵ we will consider here only two latent potential types: those whose choices are consistent with

³²Reduction of compound lotteries states that an individual should be indifferent between two lotteries with the same probability of winning and the same prize for winning. In notational terms, $((x, p; y), q; y) \sim (x, pq; y)$, where x and y are prizes and p and q are probabilities.

³³Note that this stochastic variable, μ , has the properties indicated in eqn.[4].

³⁴Harrison and Rutström (2009), Conte et al. (2008), and Bruhin et al. (2007), provide evidence supporting heterogeneity in preferences using these mixture models.

³⁵There is little to gain from doing so, given that our sample size is relative small. The reader interested in learning this method may consult McLachlan and Peel (2000).

utility maximization with linear probabilities (labeled as EUT-type), and those whose choices are consistent with nonlinear probability weighting (labeled as CPT-type).

Thus, in order to estimate a two-component mixture model, we sum the likelihood of each group, denoted by $L_i^{\text{EU}}(\cdot)$ and $L_i^{\text{PT}}(\cdot)$ and shown in eqns.[7] and [14]), multiplied by its respective *mixing* proportion, θ^{EU} and $\theta^{\text{PT}} = 1 - \theta^{\text{EU}}$. These proportions denote the probability that the EUT (CPT) model is the correct specification for a given observation. We could think of θ^{EU} as the unobservable proportion of subjects who do not distort probabilities when making risky choices, while θ^{PT} captures the unobservable proportion of subjects distorting probabilities in a nonlinear way. And what the mixture model estimation does is to cluster the observations into two groups within which the behavior is homogeneous. Then, as a result of maximizing the following weighted log-likelihood function,

$$l^{\text{MIXED}} = \sum_{i=1}^N \ln L_i^{\text{MIXED}}(r_i^{\text{EU}}, r_i^{\text{PT}}, \gamma_i, \theta^{\text{EU}}, \sigma_{\mu_i}^{\text{EU}}, \sigma_{\mu_i}^{\text{PT}}; I_i^m, X_i) = \sum_{i=1}^N \ln [(\theta^{\text{EU}} * L_i^{\text{EU}}(\cdot)) + ((1 - \theta^{\text{EU}}) * L_i^{\text{PT}}(\cdot))], \quad (17)$$

the mixing proportions are estimated together with the risk and weighting parameters and the noise corresponding to each model.³⁶ The econometrics behind the mixture models is similar to that of the unobserved exogenous regime switching models, with the (unobserved) probability of belonging to either one of those regimes being exogenous to the calculation errors. Note that in our case one regime (the EUT model) is nested in the other (the CPT model), but clearly the way that probabilities are *perceived* and *assessed* under those models or regimes is markedly different. It should be therefore clear that this estimation method is *not* equivalent to assume that CPT explains all the data and then just estimate the weighting function parameter, γ_i , for every individual i and subsequently test whether such estimate is equal to 1 or not: while in the mixture model specification we assume that the population is composed of two homogeneous subpopulations (EUT-type and CPT-type subjects), here we would be assuming that the population is composed only of CPT-type subjects.

Further note that in the same spirit as we did for the heteroskedastic errors case estimated earlier (i.e., when $\sigma_{\mu_i} = \sigma_{\mu}[X_i]$), we could also estimate the unobservable mixing proportions as a linear function of individual characteristics (i.e., $\theta_i = \theta[X_i]$). It is in this case where the similarity of mixture models with the switching regression models becomes more transparent.³⁷

³⁶It is worth mentioning that we are estimating an *aggregate* noise here for each model. We could also work on a heteroskedastic specification, but given our sample size it is likely that the estimation will not converge. Estimating a common noise for both models yields similar risk estimates, but with a lower weighting function parameter estimate.

³⁷Note the similarity to express the problem above in the exogenous switching regression context (Maddala, 1983):
 Regime 1: $y_i = X_{1i}\beta_1 + \mu_{1i}$ with probability θ_i
 Regime 2: $y_i = X_{1i}\beta_2 + \mu_{2i}$ with probability $(1 - \theta_i)$,
 with $\mu_{1i} \sim N[0, (\sigma_{\mu_i}^{\text{EUT}})^2]$ and $\mu_{2i} \sim N[0, (\sigma_{\mu_i}^{\text{CPT}})^2]$, where y_i is a binary variable, and β_1 (β_2) denotes the parameter estimate under EUT (CPT).

5 Risk Estimation Results

In this section, we report the maximum likelihood estimates assuming Normal Fechner errors with a standard deviation (σ_μ) that is also estimated (sections 2.5.1 and 2.5.2). We further examine whether nonlinear probability weighting in the presence of large random mistakes or errors can explain the multiple switching behavior observed in our data (section 5.3).

5.1 EUT and CPT Estimates

Table 2 reports the relative risk estimates assuming that choices are *entirely* explained by EUT (from eqn.[8]), and making such risk estimate be a linear function of selected individual characteristics. Three main results can be drawn from the Table. First, we find a moderate degree of risk aversion (average $\hat{r}^{\text{EUT}} = 0.45$). Second, there is evidence of a large degree of randomness in choices, as indicated by the big standard deviation of the random mistakes, $\hat{\sigma}_\mu = 2.79$, a value that is 5 percent lower than the expected utility obtained by subjects who chose lottery *Luna* (average across all decision rows), and 12 percent lower than the expected utility obtained by those choosing lottery *Sol* (average across all rows). Third, subjects with higher education are less risk averse, as shown by the decreasingly negative coefficients of the variables *some secondary* and *skilled*, though only in the latter case the effect is statistically significant ($p\text{-value} < 0.01$). The indicators of age and gender do not appear to predict risk preferences. While the variables included in the regression are jointly significant ($p\text{-value} < 0.001$), the fact that only one variable included results significantly correlated with risk preferences might suggest that our estimates would seem unreliable.

Table 2: Expected Utility Estimates with Fechner Normal Errors
Heterogeneous Agent Case

Coefficient	Variable	Estimate	Std.Error	$p\text{-value}$	95% Conf. Interval
r_i^{EUT}	Intercept	0.44	0.16	0.00	0.13 0.75
	Female	0.02	0.12	0.87	-0.22 0.26
	Young (Age < 40)	-0.18	0.18	0.31	-0.53 0.17
	Middle (Age: [50-60])	0.05	0.16	0.77	-0.26 0.36
	Old (Age > 60)	0.11	0.19	0.56	-0.26 0.48
	Illiterate	-0.33	0.28	0.24	-0.87 0.21
	Some secondary	-0.23	0.15	0.13	-0.54 0.07
	Skilled (> sec. educ.)	-0.53	0.20	0.00	-0.92 -0.14
	Low Pisco	0.19	0.13	0.14	-0.06 0.44
	High Pisco	0.45	0.42	0.28	-0.38 1.28
<i>Predicted r^{EUT} at average values</i>		<i>0.45</i>			
σ_u	Intercept	2.79	0.25	0.00	2.31 3.27
N		3,650			

Notes: S.E. clustered at the individual level. The omitted category for *age* is for those aged between 40-50. The omitted category for *education* is for those with some primary education.

We also estimated risk preferences using only subject-specific dummy variables in a separate regression (unreported in this paper), where we confirm the finding of risk averse preferences, with an average risk estimate of 1.56. These risk estimates have expectedly more variability (the standard deviation is 4.12) than the estimates reported above, and show a low correlation (of around 0.16) with those. Moreover, those risk estimates are less correlated with education and age than the estimates reported in Table 2 (-0.11 versus -0.65 and 0.08 versus 0.63).

Turning to the maximum likelihood estimates under CPT from eqn.[14]), we observe a relatively high average risk estimate ($\hat{r}^{\text{PT}} = 0.74$, as seen in Table 3).³⁸ Second, we find a significant subjective probability weighting, given that the parameter estimate, $\hat{\gamma} = 0.54 < 1$ is significantly different from 0 or 1 (p -values in both cases < 0.001). This estimate of γ is pictured by the dot-dashed curve in Figure 2 above, where we see that subjects overweight probabilities (i.e., $w(p) > p$) until $p = 0.3$, and thereafter they underweight (i.e., $w(p) < p$) middle and large probabilities. Third, the mistakes in choices made, assumed to be mean-zero normal random variables, have a large standard deviation, $\hat{\sigma}_{\mu} = 1.38$, which represents 29 percent of the expected utility obtained by subjects who chose lottery *Luna* (average across all decision rows), and 27.6 percent of the expected utility obtained by those choosing lottery *Sol* (average across all rows).

Fourth, we continue to find that subjects with higher education are less risk averse, a finding that can have important implications in terms of subjects' willingness to undertake risky but potentially profitable investments. In particular, we could expect those persons to be more able to assess risk and have a better understanding of the salient features of any new technology. Furthermore, we find that our risk estimates under EUT and CPT are significantly correlated with financial literacy³⁹ (the coefficients of correlation are -0.26 and -0.35, respectively): more financially literate subjects, who also happen to be better educated, are less risk averse. Other individual characteristics, such as age and gender, do not enter with statistically significant coefficients, though women (variable *female*) appear to be slightly more risk averse than men and older persons seem more risk averse than younger ones.

In order to have an idea about the magnitude of the effect of higher education on the CRRA coefficient estimates, in an auxiliary regression, not included in this paper, we estimated the risk and weighting function parameters using as covariates *female*, *age* (in years), and *education* (in years), and found that having 10 more years of education would decrease the risk estimate by between 0.3 (under EUT) and 0.5 (under CPT); the associated coefficients resulted barely significant in both cases (p -values are 0.081 and 0.001). The effect of age on risk aversion thus estimated is negligible, while being female implies a higher risk aversion by 0.07 (EUT) or 0.22 (CPT), with only the latter being barely statistically significant (p -value is 0.105). Overall, all the variables included in the regressions are jointly statistically significant under both models.

³⁸The correlation coefficient between the risk estimates under EUT and CPT (heterogeneous cases) is 0.82.

³⁹This indicator measures the level of comprehension of the rules of the insurance game played by our subjects. It includes a variable of self-reported comprehension of the experiment, and objective measures of how well they knew the consequences of loan default, and the features of the indemnity function.

Table 3: Cumulative Prospect Theory Estimates with Fechner Normal Errors
Heterogeneous Agent Case

Coefficient	Variable	Estimate	Std.Error	<i>p-value</i>	95% Conf.	Interval
r_i^{CPT}	Intercept	0.85	0.18	0.00	0.50	1.19
	Female	0.15	0.15	0.31	-0.14	0.43
	Young (Age < 40)	-0.07	0.12	0.58	-0.30	0.17
	Middle (Age: [50-60])	0.15	0.13	0.27	-0.11	0.41
	Old (Age > 60)	0.11	0.22	0.63	-0.33	0.54
	Illiterate	-0.28	0.42	0.50	-1.10	0.54
	Some secondary educ.	-0.49	0.16	0.00	-0.80	-0.19
	Skilled (> sec. educ.)	-0.71	0.19	0.00	-1.08	-0.34
	Low Pisco	0.07	0.13	0.57	-0.17	0.32
	High Pisco	0.03	0.22	0.88	-0.40	0.47
<i>Predicted r^{CPT} at average values</i>		<i>0.74</i>				
γ_i	Intercept	0.44	0.08	0.00	0.29	0.60
	Female	-0.09	0.06	0.17	-0.21	0.04
	Young (Age < 40)	0.09	0.17	0.59	-0.24	0.43
	Middle (Age: [50-60])	-0.09	0.07	0.15	-0.22	0.03
	Old (Age > 60)	-0.07	0.09	0.47	-0.25	0.12
	Illiterate	0.03	0.13	0.82	-0.22	0.28
	Some secondary educ.	0.22	0.08	0.00	0.06	0.38
	Skilled (> sec. educ.)	0.53	0.29	0.06	-0.03	1.09
	Low Pisco	0.04	0.07	0.50	-0.08	0.17
	High Pisco	0.16	0.15	0.29	-0.14	0.46
<i>Predicted γ at average values</i>		<i>0.54</i>				
σ_μ	Intercept	1.38	0.18	0.00	1.02	1.73
N		3,650				

Notes: S.E. clustered at the individual level. The omitted category for *age* is for those aged between 40-50. The omitted category for *education* is for those with some primary education.

Before turning to examine the individual characteristics correlated with the weighting function parameter estimate ($\hat{\gamma}$), we should note that, as long as $\hat{\gamma} < 1$, greater values of $\hat{\gamma}$ imply a lower sensitivity to probability changes, which is reflected by a flatter curve (with respect to the 45 degree line) both in the overweighting and underweighting regions of such function, and also that the intersection with the 45 degree line happens at larger values of probabilities. From Table 3, we see that only education is significantly correlated with the shape of the weighting function.⁴⁰ Further, we see that women display a more curved weighting function, which implies that they underweight medium and large probabilities more strongly than males, but the coefficient on the

⁴⁰Using the results of our auxiliary regression, we find that 5 additional years of education would imply an increase in $\hat{\gamma}$ by 0.24 (e.g., from 0.54 to 0.78 using the average values of the other covariates); which will cause a substantial reduction in the curvature of the weighting function, especially in the convex, underweighting region.

gender indicator is not statistically significant ($p\text{-value} = 0.17$).⁴¹ While examining this gender difference goes beyond the scope of our analysis, recent evidence suggests that women’s probability weighting is sensitive to mood states, while men’s is not (Fehr-Duda et al., 2006b).⁴²

So far, we have been estimating the random calculation errors as being homoskedastic, with a constant standard deviation of the errors across subjects. However, it seems plausible to hypothesize that the magnitude of such standard deviation may be correlated with some observable characteristics. In particular, we will test the significance of age and education (both expressed in years)⁴³ in the following errors equation:

$$\hat{\sigma}_{\mu_i} = \hat{\sigma}_0 + \hat{\sigma}_a \text{Age}_i + \hat{\sigma}_e \text{Education}_i. \quad (18)$$

In the case of EUT, the average standard deviation of the mistake thus estimated is large, 3.30, while the resulting risk aversion evaluated at the mean of the regressors is 0.57 (see Table 4).

Table 4: EUT Estimates Assuming Heterogeneous Subjects
With Heteroskedastic Fechner Normal Errors

Coefficient	Variable	Estimate	Std.Error	<i>p-value</i>	95% Conf. Interval	
r_i^{EUT}	Intercept	0.45	0.19	0.02	0.08	0.83
	Female	0.10	0.10	0.34	-0.10	0.30
	Young (Age < 40)	-0.0003	0.12	1.00	-0.24	0.24
	Middle (Age: [50-60])	0.04	0.14	0.76	-0.24	0.33
	Old (Age > 60)	0.23	0.27	0.40	-0.30	0.75
	Illiterate	-0.61	0.41	0.14	-1.42	0.20
	Some secondary educ.	-0.22	0.19	0.24	-0.60	0.15
	Skilled (> sec. educ.)	-0.35	0.20	0.08	-0.75	0.05
	Low Pisco	0.15	0.11	0.16	-0.06	0.36
	High Pisco	0.77	0.42	0.07	-0.06	1.60
<i>Predicted r^{EUT} at average values</i>		<i>0.57</i>				
σ_{u_i}	Intercept	2.30	0.68	0.00	0.95	3.64
	Age (years)	0.04	0.01	0.00	0.02	0.06
	Education (years)	-0.20	0.05	0.00	-0.29	-0.11
<i>Predicted σ_{μ} at average values</i>		<i>3.30</i>				
N		3,650				

Notes: S.E. clustered at the individual level. The omitted category for *age* is for those aged between 40-50. The omitted category for *education* is for those with some primary education.

Given the nature of our risk game, we would expect farmers with less education to display greater calculation mistakes in their choices than those with higher education. Indeed, we find that

⁴¹The effect of gender on the shape of weighting function was also addressed by Fehr-Duda et al. (2006a), who find a similar result to ours.

⁴²The authors find that women in a good mood tend to underestimate probabilities of gains more heavily than do women in a bad mood. A better mood than usual is also correlated with weighting probabilities more “optimistically”.

⁴³Expressing age and education as indicator variables in the model did not yield convergence.

mistakes are positively correlated with age ($\hat{\sigma}_a = 0.04$) and negatively correlated with education ($\hat{\sigma}_e = -0.20$), with their respective coefficients being statistically significant at one percent. The regressors on the risk parameter equation are also jointly significant at 5 percent.

For the CPT case, the average risk aversion estimate is 0.67, the average weighting function parameter is 0.86, and the average standard deviation of the noise, 2.11 (Table 5). In the errors equation, *education* enters with an insignificant coefficient, but the coefficient on age, $\hat{\sigma}_e = 0.03$, is significant at 1 percent, and indicates that older subjects are more prone to make mistakes.

Table 5: CPT Estimates Assuming Heterogeneous Subjects
With Heteroskedastic Fechner Normal Errors

Coefficient	Variable	Estimate	Std.Error	<i>p-value</i>	95% Conf. Interval	
r_i^{CPT}	Intercept	0.68	0.22	0.00	0.26	1.10
	Female	0.16	0.13	0.23	-0.10	0.42
	Young (Age < 40)	0.10	0.12	0.43	-0.15	0.34
	Middle (Age: [50-60])	0.09	0.13	0.51	-0.17	0.35
	Old (Age > 60)	0.08	0.23	0.74	-0.38	0.53
	Illiterate	-0.03	0.54	0.95	-1.09	1.03
	Some secondary	-0.35	0.19	0.06	-0.72	0.02
	Skilled (> sec. educ.)	-0.54	0.23	0.02	-1.00	-0.08
	Low Pisco	0.05	0.10	0.62	-0.15	0.26
	High Pisco	0.30	0.30	0.32	-0.29	0.89
<i>Predicted r^{CPT} at average values</i>		<i>0.67</i>				
γ_i	Intercept	0.55	0.17	0.00	0.23	0.88
	Female	-0.15	0.13	0.25	-0.40	0.10
	Young (Age < 40)	-0.15	0.12	0.19	-0.39	0.08
	Middle (Age: [50-60])	-0.11	0.12	0.38	-0.34	0.13
	Old (Age > 60)	-0.01	0.17	0.97	-0.33	0.32
	Illiterate	3.58	3.56	0.31	-3.39	10.55
	Some secondary	0.19	0.12	0.12	-0.05	0.42
	Skilled (> sec. educ.)	0.39	0.30	0.19	-0.19	0.98
	Low Pisco	0.11	0.13	0.38	-0.14	0.36
	High Pisco	0.47	0.19	0.01	0.10	0.84
<i>Predicted γ at average values</i>		<i>0.86</i>				
σ_{μ_i}	Intercept	0.73	0.76	0.34	-0.75	2.22
	Age (years)	0.03	0.01	0.00	0.01	0.06
	Education (years)	-0.08	0.06	0.16	-0.20	0.03
<i>Predicted σ_{μ} at average values</i>		<i>2.11</i>				
N		3,640				

Notes: S.E. clustered at the individual level. The omitted category for *age* is for those aged 40-50. The omitted category for *education* is for those with some primary education.

5.2 Mixing EUT and CPT: Estimation Results

In this section, we will relax the assumption made so far that the data are *entirely* explained by a single decision model. Instead, we will now allow the data to be explained by two decision models, EUT *and* CPT, in order to account for heterogeneity in preferences (we will estimate the model given by eqn.[17]). In this two-component *mixture* model specification, in addition to the risk and weighting parameters under each model, the proportion of the sample that behaves according to either EUT or CPT is estimated. Maximum likelihood estimates of those parameters, as well as the standard deviation of the Fechner errors, are reported in Table 6.

We can see that 76 percent of subjects behave as prospect utility maximizers (i.e., $\hat{\theta}^{\text{PT}} = 0.76$), overweighting small probabilities and underweighting medium and large probabilities, while the remaining 24 percent behave as expected utility maximizers, treating probabilities linearly. These mixing proportions are statistically different from 0.5 or 0 (both *p-values* < 0.05). Moreover, we continue to find evidence of risk aversion: the risk estimate for the EUT-type subjects is 0.20, while that for the CPT-type subjects is 1.21. Furthermore, the estimated weighting parameter, $\hat{\gamma} = 0.33$, implies a substantial curvature in the weighting function, especially in the convex region (see the solid inverse-S shaped curve depicted in Figure 2 above). Lastly, we report that the estimated standard deviations of the Fechner errors are still large, in particular the one related with the CPT model, $\hat{\sigma}_{\mu}^{\text{PT}} = 1.4$. We further observe that several

Table 6: Mixture Model Estimates with Homogeneous Mixing Proportions
Assuming Homogeneous Subjects and Normal Fechner Errors

Variable		Expected Utility	Prospect Theory
r	Coefficient	0.20***	1.21*
	Standard Error	0.07	0.65
	95% C. Interval	[0.07, 0.34]	[-0.08, 2.49]
γ	Coefficient	n.a.	0.33*
	Standard Error		0.19
	95% C. Interval		[-0.04, 0.71]
θ	Coefficient	0.24**	0.76***
	Standard Error	0.10	0.10
	95% C. Interval	[0.06, 0.43]	[0.57, 0.94]
σ_{μ}	Coefficient	0.51*	1.40
	Standard Error	0.26	0.88
	95% C. Interval	[-0.01, 1.03]	[-0.33, 3.13]
N		3,780	3,780

n.a.: not applicable.

*** (**) [*] denotes significance at 1% (5%) [10%] level. S.E. clustered at the individual level.

In sum, from this mixture model we find that the EU maximizers are significantly less risk

averse than the PU maximizers, and this model allows us to capture the existence of nonlinear probability weighting in the choices made. A clear implication of the statistical significance of the mixing proportions is that aggregation of preferences, namely assuming that only one model governs risk preferences, becomes problematic.

While it could be interesting to estimate the heterogeneous agent case (i.e., to include covariates as linear functions of the risk and weighting function parameters) under this mixture model, which it would provide a clearer picture of the heterogeneity *within* each type of decision maker, doing so did not yield convergence. Fortunately, we still can examine the distinctive characteristics that define the regime (EUT-type or CPT-type) subjects belong. We proceed in this direction in the next subsection.

5.2.1 Heterogeneous Mixing Weights: Explaining the Regime Switching Behavior

In order to test for the existence of heterogeneity *within* the subset of EUT-type or CPT-type subjects, in terms of the characteristics associated with their belonging to one regime or another, we next estimate the mixing proportion parameter as a linear function of gender (*female*) and farming experience expressed in years:⁴⁴

$$\hat{\theta}_i^{\text{EUT}} = \hat{\theta}_0 + \hat{\theta}_{fe} \text{Female}_i + \hat{\theta}_{fa} \text{Farm_exp}_i \quad (19)$$

in the likelihood function shown in eqn.[17]. Recall that $\hat{\theta}_0$ is the parameter we have been estimating thus far. The maximum likelihood estimates shown in Table 7 indicate that gender and farming experience are statistically significant (they are further jointly but barely significantly: *p-value* = 0.107). While these results should be taken with caution, since the regression does not capture the heterogeneity at the parameter level (i.e., no covariates are included for the risk parameters),⁴⁵ they suggest that subjects with more farming experience behave according to expected utility theory, while less experienced farmers's behavior is explained by prospect theory, a result also found by List (2004) for a different set of subjects. Interestingly, these EUT-type of subjects also attained higher education levels than the CPT-type of subjects (12 versus 5 years), and are substantially younger (40 versus 58 years old). In both cases, the means *T*-tests reject the null hypothesis that the means are equal at 1 percent of significance. Evaluated at the average values of those covariates, the estimated value of $\hat{\theta}^{\text{EUT}}$ is 0.18, thus implying that most of the observations (82 percent) are explained by CPT ($\hat{\theta}^{\text{CPT}}$ is 0.82).

5.3 Switching Behavior and Probability Weighting with Large Mistakes

As mentioned in the introduction, we observe a large proportion (to be precise, 196 or 52 percent) of subjects switching back and forth from one lottery to the other in our risk game.⁴⁶ The current

⁴⁴Including education in the equation results in globally insignificant regression coefficients.

⁴⁵The regression including covariates in the risk and weighting function parameters did not converge.

⁴⁶121 subjects switched only once.

Table 7: Mixture Model Estimates with Heterogeneous Mixing Proportions
Assuming Homogeneous Agents & Normal Fechner Errors

Variable		Estimate	S.E.	95% Conf.	Interval
r^{EUT}	Intercept	0.20**	0.10	0.01	0.40
r^{CPT}	Intercept	0.71***	0.25	0.21	1.21
γ	Intercept	0.54***	0.17	0.21	0.88
θ_i^{EUT}	Intercept	0.72***	0.14	0.46	0.99
	Female	0.24*	0.15	-0.04	0.53
	Farming Experience-Yrs.	0.47***	0.01	0.44	0.50
σ_μ^{EUT}	Intercept	0.62**	0.25	0.13	1.13
σ_μ^{CPT}	Intercept	2.19***	0.64	0.94	3.43
N		3,680			

*** (**) [*] denotes significance at 1% (5%) [10%] level.

S.E. clustered at the individual level.

literature falls short to explain this multiple switching behavior (MSB) and, to the best of our knowledge, only provides two potential explanations: lack of salience in the monetary incentives (Bruner, 2008), and a revealed preference for indifference between the two lotteries (Harrison and Rutström, 2008).

In this section, we examine an alternative explanation to MSB: a combination of large random calculation errors made by our subjects *and* the existence of subjective distortions of underlying probabilities found in our sample. To analyze this matter, we will consider the risk estimate of an average subject in the heterogeneous agent case under CPT (i.e., we will use $\hat{r} = 0.74$) (see Table 3). We will then picture the decision rules functions in the absence of random mistakes under EUT and CPT (heterogeneous agent case). These decision rules are given by the difference in expected utilities (or prospective utilities, depending on the model considered) between lottery *Sol* (*S*) and lottery *Luna* (*L*), and are a function of the estimated risk parameter, $\hat{r}$, conditional on the probabilities or decision weights used:⁴⁷

$$\Delta EU_m(\hat{r} | p) = EU_m^S(\hat{r}) - EU_m^L(\hat{r}) \quad (20)$$

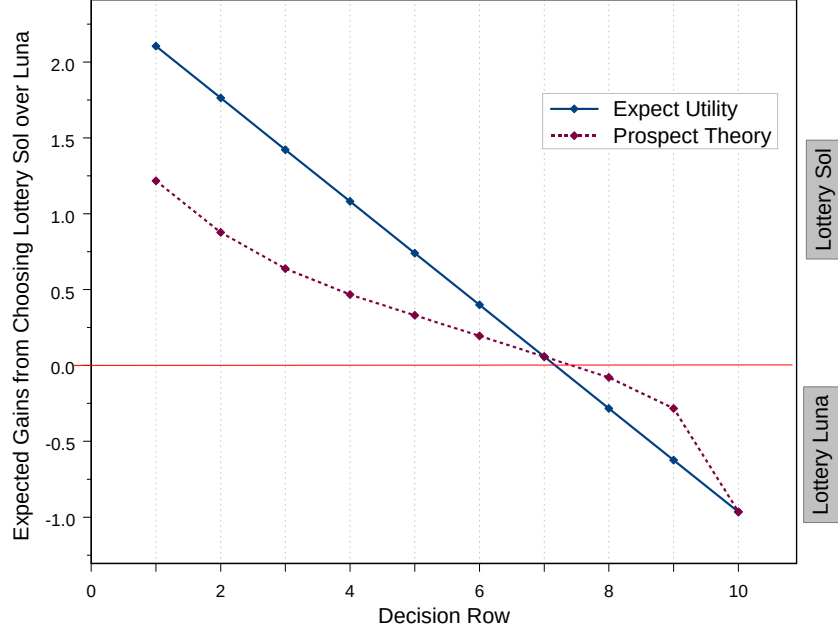
$$\Delta PU_m(\hat{r}, \hat{\gamma} | \widehat{\mathbf{dw}}[p]) = PU_m^S(\hat{r}) - PU_m^L(\hat{r}). \quad (21)$$

Pictured in Figure 3, we see that those functions are monotonically decreasing over the decision rows, m , displayed in the horizontal axis. A positive value of $\Delta EU_m(\cdot)$ or $\Delta PU_m(\cdot)$ would obviously mean that lottery *S* should be selected; while a negative value would imply that lottery *L* should be chosen. The solid line depicts the *true* expected gains from choosing lottery *S* over lottery *L* under EUT, while the dashed line pictures the same function for CPT. In both cases, those functions attain a positive value until decision row 7, and thereafter their values become negative, meaning

⁴⁷Note that these equations are similar to the ones used above (eqns.[15] and [3], respectively), but without the noise parameter, and with the value of the risk estimate under each model is already plugged in the equations.

that lottery L should be chosen, in the absence of calculation mistakes. We can see that for a given risk parameter (equal to 0.74), while linear probability weighting results in a straight line with negative slope (the solid line, EUT), non-linear probability weighting implied by the weighting function parameter $\hat{\gamma}$ of 0.54 (the dotted curve, CPT) flattens the expected gains function for the first nine decision rows.

Figure 3: Expected Gains under EUT and CPT



How can we then rationalize MSB in this context? Figure 3 can be misleading in answering such a question, given that it shows expected gains under different utility functions, which are ordinal measures and have different underlying distributions of calculation mistakes. Thus, one way to examine MSB is to estimate the probability of making mistakes under the error distributions estimated above ($\hat{\mu} \mid \text{EUT} \sim N[0, 2.79]$ and $\hat{\mu} \mid \text{CPT} \sim N[0, 1.38]$). Thus, the probability of *not* making mistakes (i.e., of choosing lottery S when it has higher expected gains; and selecting lottery L , otherwise) under EUT is given by eqn.[5]:

$$\Pr(\text{no mistakes}) = \Pr(EU_m^S - EU_m^L + \mu > 0) = \Phi\left(\frac{EU_m^S - EU_m^L}{\hat{\sigma}_\mu}\right),$$

where Φ denotes the c.d.f. of the standard normal distribution. A similar expression holds for CPT. The probability of making mistakes ($\Pr[\text{mistakes}]$) is then $[1 - \Pr(\text{no mistakes})]$. We then plugged the estimated values of the true expected gains from choosing lottery S into the prior expression and computed such probabilities for each decision row m . We report these calculations in Table

8, which also presents the decision weights ($\mathbf{dw}[p]$), expressed by eqn.[12], used to produce Figure 3. Defined over the cumulative probability distribution, these decision weights are implied by the weighting functions reported in eqn.[13] for the above-indicated value of $\hat{r} = 0.74$. We see in the table that in the CPT heterogeneous case, subjects overweight (i.e., $\mathbf{dw}[p] > p$) small and medium probabilities (up to row 6), and thereafter they underweight (large) probabilities, a result that is reflected by the shape of the expected gains function for CPT pictured above.

So, can the nonlinear probability weighting observed under CPT help explain the MSB? The results are not as strong as we expected. To address the question posed, note that as long as the probability of making mistakes is *greater* under CPT than EUT, the former model would explain better the MSB. And, the larger the gap between the probabilities of making mistakes, the greater the “explanatory effect” of the model showing the higher probability. Thus, as seen in the table, CPT explains slightly better than EUT the MSB in rows 3, 4, 5, 8, and 9. In turn, EUT does better in rows 1, 7, and 10. (In rows 2 and 6, both models explain MSB equally well.) Further note that the probability of making mistakes under both models gets large from decision rows 5 through 9 (40 to almost 50 percent), a result that reflects the large magnitude of the calculation errors made by our subjects. A more accurate estimation of subject-specific calculation mistakes, which cannot be done in this paper because of our limited sample size, would likely help to better examine the MSB.

Table 8: Expected Gains from Lottery Sol and Decision Weights under EUT and CPT

Variables	Decision Row									
	1	2	3	4	5	6	7	8	9	10
<i>EUT Heterogeneous agent case, with $\hat{\sigma}_\mu = 2.79^a$</i>										
Probability, p^b	0.10	0.20	0.30	0.40	0.50	0.60	0.70	0.80	0.90	1.00
Expected gain from choosing lottery S over L	2.10	1.76	1.42	1.08	0.74	0.40	0.06	-0.28	-0.62	-0.97
Probability of mistake ^c	0.23	0.26	0.31	0.35	0.40	0.44	0.49	0.46	0.41	0.36
<i>CPT Heterogeneous agent case, with $\hat{\gamma} = 0.54$ & $\hat{\sigma}_\mu = 1.38^d$</i>										
Decision weight, $\mathbf{dw}(p)^b$	0.36	0.46	0.53	0.58	0.62	0.66	0.70	0.74	0.80	1.00
Expected gain from choosing lottery S over L	1.22	0.88	0.64	0.47	0.33	0.19	0.06	-0.08	-0.28	-0.97
Probability of mistake ^c	0.19	0.26	0.32	0.37	0.41	0.44	0.48	0.48	0.42	0.24

^a Estimates from Table 2. ^b These are the probabilities, weights, and decision weights, of the higher prize under each lottery. E.g., p_1 (higher prize)=0.1, p_2 (lower prize)=0.9, $w(p_2)$ =0.64, $\mathbf{dw}(p_1)$ =1- $w(p_2)$ =0.36.

^c This is the probability of choosing the lottery with lower expected gains.

^d Estimates from Table 3.

In terms of the experimental design, the multiple switching behavior could be reduced in two

ways: by doing a better job explaining the mechanics of the game.⁴⁸ Some authors suggest that MSB can be due to indifference (Harrison and Rutström, 2008); if this is the case, introducing a “I am indifferent between lotteries” choice could help mitigate the problem. However, when such behavior is more a result of confusion, such modification would likely be insufficient. A second way to address the problem could be by providing bigger (probably monetary) incentives to pay attention. This solution goes in line with what Bruner (2007) proposed. In terms of what we find in this section, any factor that may improve the subjects’ understanding of the mechanics of the game would likely reduce the prevalence of the MSB.

6 Estimation Results Excluding Irrational Subjects

As mentioned in the introduction, a large number of subjects (105 to be precise) mistakenly chose the lottery *Sol* in the 10th decision row, thus preferring to get 1,800 Soles for sure, instead of the 3,500 Soles for sure they could have gotten by choosing lottery *Luna*. We call these subjects, *irrational*. As one may expect, those subjects attained lower levels of schooling (two years less of education), a much lower financial literacy indicator (0.48 versus 0.56), and are significantly older (in more than 4 years),⁴⁹ than subjects who did not make such mistake. Dropping those 105 individuals from our original sample of 378 subjects (3,780 observations), results in a restricted sample of 273 subjects (2,730 observations). From that subsample, we have individual-level information for 265 individuals (2,650 observations).

Summarizing the regression results, reported in Appendix E, we find that the expected utility maximizers in this subsample are much *less* risk averse than in our full sample: $\hat{r}^{EU} = 0.11$ (see column 3 of Table E.1), and although the estimated magnitude of the errors in the choices made is still substantial, it is expectedly lower than in the full sample, $\hat{\sigma}_\mu = 2.0$ versus $\hat{\sigma}_\mu = 2.8$. Further, unlike the full sample, age indicators become significant: *young* (subjects who are 40 years old or less) and *middle* (subjects who are between 50-60 years old) are less risk averse with respect to subjects aged between 40-50 (the excluded age category), while *illiterate* subjects are less risk averse than those who have some primary education (the excluded education category).

On the other hand, under CPT, we continue to see strong evidence of subjective probability weighting, $\hat{\gamma} = 0.46$ (see column 3 of Table E.2), and we further find that more educated subjects are significantly more likely to be less risk averse than those with only primary education, a result also found in the full sample. Two predictors of the shape of the weighting function are higher education and younger age: more educated and younger persons are less sensitive to probability changes. Lastly, the standard deviation of the mistakes is relatively small in the restricted sample; this result is because a large proportion of the noise has been removed with the irrational subjects.

Further exploring the data, we estimated a *mixture* model for the restricted sample. The results,

⁴⁸Also, playing a practice round for real money would likely help reduce the magnitude of the random component in choices.

⁴⁹All the mean *T*-tests performed in those cases have a *p-value* < 0.001.

reported in Table E.3, show evidence of the higher proportion of observations best fitted by CPT than EUT ($\hat{\theta}^{\text{PT}} = 0.71$), and of moderate (EUT-type, $\hat{r}^{\text{EU}} = 0.42$) to high (CPT-type, $\hat{r}^{\text{PT}} = 1.14$) risk aversion levels. We should mention, however, that the weighting function parameter thus estimated, $\hat{\gamma} = 0.27$, implies a non-monotonic curve, very close to a step function.

To sum up then, even for the restricted sample, we find that a large proportion of subjects makes risky choices using nonlinear probability weighting, a result recurrently found by the recent literature.

7 Conclusion

In this paper, we examine the results of an artefactual field experiment designed to measure risk preferences in a southern Peruvian valley. We fit the experimental data to two of the leading models of decision under risk, Expected Utility Theory (EUT) and Cumulative Prospect Theory (CPT), and in both cases find evidence of risk averse preferences. This qualitative result remains unaltered when a subset (a proportion that is also estimated) of observations is allowed to be explained by different decision-making processes or models (i.e., either EUT *or* CPT), although the degree of risk aversion thus estimated varies substantially.

In terms of the individual characteristics that predict risk preferences, in general, only higher education appears to be positively correlated with a greater propensity to take risks, a result that not only suggests a linkage between cognitive abilities and risk preferences, but can also have interesting implications for the diffusion of new technologies involving risks. We could, for instance, hypothesize that if the elicited preferences indeed reflect preferences held in real life, the “Schumpeterean” farmers (innovators) would likely come from the higher end of the education sample distribution. This conjecture implicitly assumes that preferences are somewhat stable, a result that needs to be properly tested. Examining preferences stability would certainly require a more complicated experimental design than ours. The result that only higher education appears significant in the regressions may suggest that our estimates would seem to be unreliable, which advises us to take these results with caution.

On the other hand, when we assume that only *one* model explains the entire data, the risk estimates are higher under CPT than EUT, and nonlinear probability weighting often plays a role in explaining risky choices. If we instead let the observations be endogenously classified as EUT-type or CPT-type by estimating mixture models, we find that CPT explains typically about 70 percent or a higher proportion of the observations; the remaining percent of observations is explained by EUT, where subjects do not distort objective probabilities. This result shows that considering a single decision model as representative of the preferences of an entire set of people could lead to biased results. Statistical power issues prevent us from conducting further analysis with this mixture model specification, in particular that related to the heterogeneity within each type of subjects (EUT-type and CPT-type).

In addition to the evidence of preference heterogeneity at the functional form level, we offer some partial evidence that underweighting of small and medium probabilities and overweighting of large probabilities, in the presence of large random calculation errors, may explain better than EUT the multiple switching across decision rows observed in our data; however. However, nonlinear probability weighting under CPT does not explain multiple switching as much as we expected.

While in principle, those errors are capturing the effects of several factors, including confusion and indifference between lotteries, we find that they are correlated with education and age: less educated and older subjects are more likely to display larger noises in their choices.

This research has several limitations, which suggest directions for future research. First, including a loss dimension in the lotteries would have allowed us to test loss aversion in preferences, a subject that may be important if we would like to examine the predictive power of our estimates in terms of decisions involving potential losses. Ignoring loss preferences may thus be biasing our results. Second, we require an *ad hoc* experimental design in order to account for other variables that may affect risk preferences. For instance, some recent works suggest that psychological factors, such as emotions (e.g., mood, anger, fear), can also be correlated with risk preferences. Third, a more sophisticated modeling of the random choices could allow us to tease out its nonrandom component. Fourth, in a world where rationally bounded people are exposed to an experimental environment where decisions need to be taken in a short time, the assumption that subjects *can* make complex calculations with little difficulty is clearly unrealistic. In that context, it could be interesting to unveil some heuristics (if any) used by subjects to make risky choices. As we move towards the use of more comprehensive approaches to explain people’s behavior, this promises to be an active area of research. Finally, in principle, our results may be sensitive to the errors type assumed (Fechner’s) or the specification of the weighting function; therefore exploring alternative error stories and weighting functions may prove to be helpful to examine the robustness of our findings. In an effort to advance in this direction, Appendix F provides details about some of those proposed extensions to our analysis. As in that Appendix, the main qualitative results continue to hold under an alternative error type and weighting function.

References

- [1] Andersen, Steffen, Glenn Harrison, Morten Igel Lau, and E. Elisabet Rutström (2006). "Elicitation using multiple price list formats," *Experimental Economics*, 9: 383-405.
- [2] Arrow, Kenneth (1971). *Essays in the Theory of Risk-Bearing*. Amsterdam: North-Holland Pub.Co.
- [3] Becker, Gordon M., Morris H. DeGroot and Jacob Marschak (1964). "Measuring Utility by a Single-Response Sequential Method," *Behavioral Science*, 9: 226-232.
- [4] Benjamin, Daniel, Sebastian A. Brown and Jesse M. Shapiro (2006). "Who is 'Behavioral'? Cognitive Ability and Anomalous Preferences." Working Paper, University of Chicago.
- [5] Binswanger, Hans (1980). "Attitudes Toward Risk: Experimental Measurement in Rural India," *American Journal of Agricultural Economics*, 62(3): 395-407.
- [6] Binswanger, Hans (1981). "Attitudes Toward Risk: Theoretical Implications of an Experiment in Rural India," *Economic Journal*, 91(364): 867-890.
- [7] Binswanger, Hans, Dayanatha Jha, T. Balaramaiah and Donald Sillers (1980). "The impacts of risk aversion on agricultural decisions in semi-arid India," International Crops Research Institute for the Semi-Arid Tropics, Hyderabad, India. Mimeo.
- [8] Bruhin, Adrian, Helga Fehr-Duda and Thomas F. Epper (2007). "Risk and Rationality: Uncovering Heterogeneity in Probability Distortion," *SOI Working Paper* 0705, Socioeconomic Institute, University of Zurich, July.
- [9] Bruner, David (2007). "Multiple Switching Behavior in Multiple Price Lists," Mimeo, University of Calgary, Department of Economics.
- [10] Bruner, David, Michael McKee and Rudy Santore (2008). "Hand in the Cookie Jar: An Experimental Investigation of Equity-based Compensation and Managerial Fraud," *Southern Economic Journal*, 75(1): 261-78.
- [11] Burks, Stephen, Jeffrey Carpenter, Lorenz Goette and Aldo Rusticini (2009). "Cognitive Skills Explain Economic Preferences, Strategic Behavior and Job Attachment," Working Paper.
- [12] Camerer, Colin and Teck-Hua Ho (1994). "Violations of the betweenness axiom and nonlinearity in probability," *Journal of Risk and Uncertainty*, 8(2): 167-196.
- [13] Cardenas, Juan Camilo and Jeffrey Carpenter (2008). "Behavioural Development Economics: Lessons from Field Labs in the Developing world," *Journal of Development Studies*, 44 (3): 337-364.
- [14] Conte, Anna, Peter Moffat and John Hey (2008). "Mixture models of choice under risk", *Journal of Econometrics*, forthcoming.
- [15] Cox, James C. and Glenn Harrison (eds.) (2008). *Risk Aversion in Experiments*. Bingley, U: Emerald, Research in Experimental Economics, Volume 12.

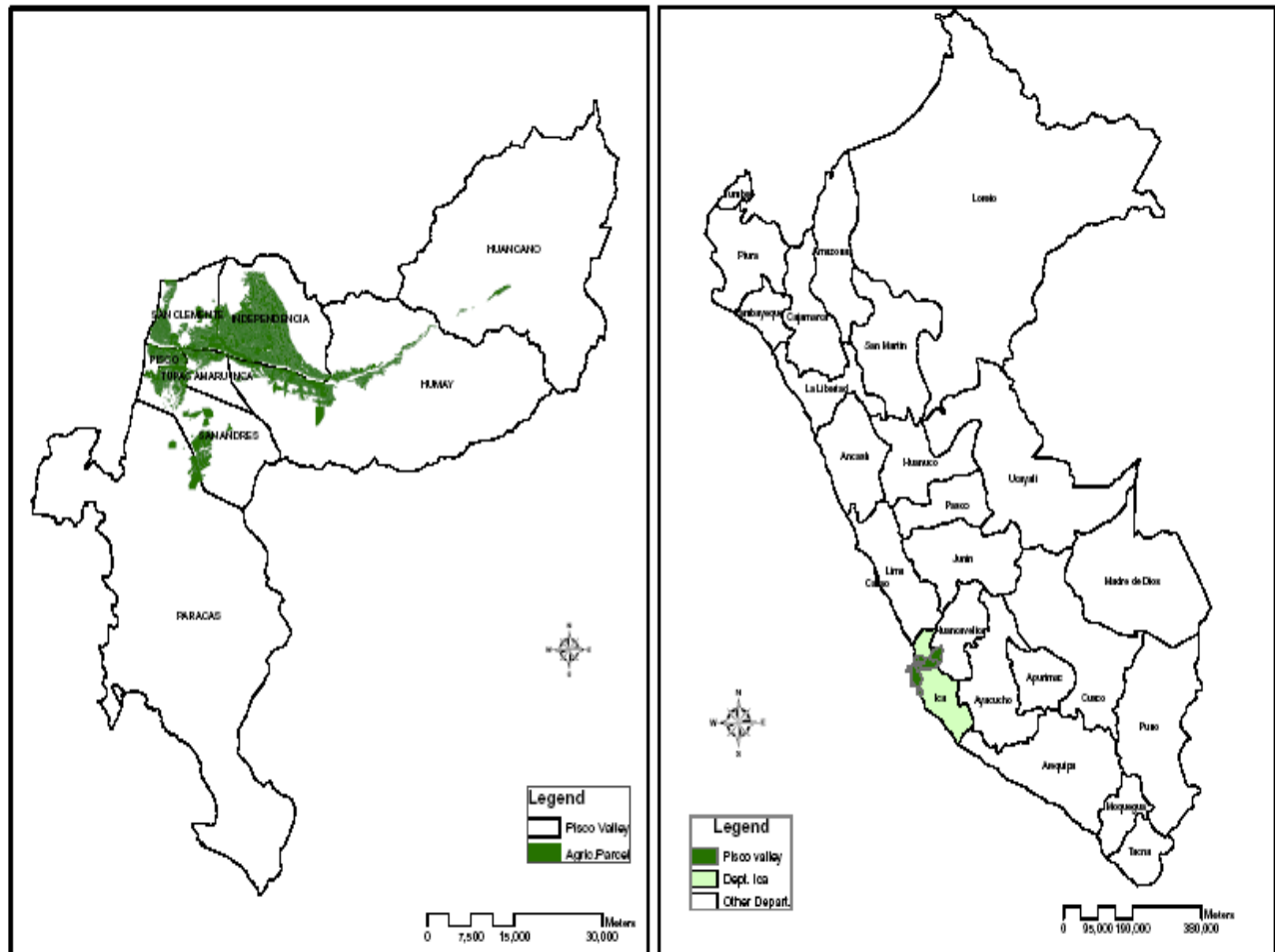
- [16] Dohmen, Thomas, Armin Falk, David Huffman and Uwe Sunde (2007). "Are Risk Aversion and Impatience Related to Cognitive Ability?," *IZA Discussion Paper* No. 2735. Bonn, Germany: Institute for the Study of Labor (IZA).
- [17] Eckel, Catherine and Philip Grossman (2008). "Forecasting Risk Attitudes: An Experimental Study Using Actual and Forecast Gamble Choices," *Journal of Economic Behavior and Organization*, 68(1): 1-17.
- [18] Eckel, Catherine and Rick Wilson (2004). "Is Trust a Risky Decision?," *Journal of Economic Behavior and Organization*, 55(4): 447-465.
- [19] Engle-Warnick, Jim, Javier Escobal and Sonia Laszlo (2007). "Ambiguity Aversion as a Predictor of Technology Choice: Experimental Evidence from Peru," *Scientific Series* 2007s-01. Montreal, CIRANO.
- [20] Fechner, Gustav T. (1860/1966). *Elements of Psychophysics*. New York: Holt Rinehart and Winston.
- [21] Feder, Gershon (1980). "Farm Size, Risk Aversion and the Adoption of New Technology under Uncertainty," *Oxford Economic Papers*, 32(2): 263-283.
- [22] Feder, Gershon, Richard Just and David Zilberman (1985). "Adoption of Agricultural Innovations in Developing Countries: A Survey," *Economic Development and Cultural Change*, 33: 255-297
- [23] Fehr-Duda, Helga, Manuele de Gennaro and Renate Schubert (2006a). "Gender, Financial Risk, and Probability Weights," *Theory and Decision*, 60: 283-313.
- [24] Fehr-Duda, Helga, Marc Schürer and Renate Schubert (2006b). "What Determines the Shape of the Probability Weighting Function," Economics Working Papers 06/54, Center for Economic Research at ETH Zurich.
- [25] Fox, Craig R. and Russel A. Poldrack (2009). "Prospect Theory and the Brain," In: Glimcher, Paul W., Colin F. Camerer, Ernst Fehr and Russell A. Poldrack (eds.), *Neuroeconomics. Decision Making and the Brain* (pp.145-173). London; San Diego, CA: Academic Press.
- [26] Frederick, Shane (2005). "Cognitive Reflection and Decision Making," *Journal of Economic Perspectives*, 19(4): 25-42.
- [27] Galarza, Francisco (2009). "Risk, Credit, and Insurance in Peru: Field Experimental Evidence," Working Paper, Department of Agricultural and Applied Economics, University of Wisconsin-Madison.
- [28] Gonzalez, Richard and George Wu (1999). "On the Shape of the Probability Weighting Function," *Cognitive Psychology*, 38: 129-166.
- [29] Guiso, Luigi and Monica Paiella (2007). "Risk Aversion, Wealth, and Background Risk," *EUI Working Paper* No. 47, Department of Economics, European University Institute (EUI), Italy.
- [30] Harrison, Glenn (2007). "Maximum Likelihood Estimation of Utility Functions Using *Stata*," Working Paper # 06-12, Department of Economics, College of Business Administration, University of Central Florida, April.

- [31] Harrison, Glenn, Morten I. Lau and Elisabet Rutström (2007). "Estimating Risk Attitudes in Denmark: A Field Experiment," *Scandinavian Journal of Economics*, 109(2): 341-368.
- [32] Harrison, Glenn and John List (2004). "Field Experiments," *Journal of Economic Literature*, 42 (4): 1013-1059.
- [33] Harrison, Glenn and Elisabet Rutström (2008). "Risk Aversion in the Laboratory." In: Cox, James and Glenn Harrison (eds.), *Risk Aversion in Experiments*. Bingley, U: Emerald, Research in Experimental Economics, Volume 12, pp. 41-196.
- [34] Harrison, Glenn and Elisabet Rutström (2009). "Expected Utility Theory and Prospect Theory: One Wedding and A Decent Funeral," *Experimental Economics*, 12. Forthcoming.
- [35] Harrison, Glenn, Morten Lau, Elisabet Rutström and Melanie Williams (2005a). "Eliciting Risk and Time Preferences Using Field Experiments: Some Methodological Issues," in: Carpenter, Jeffrey, Glenn Harrison and John List (eds.), *Field Experiments in Economics*. Greenwich, CT: JAI Press, Research in Experimental Economics, Volume 10.
- [36] Harrison, Glenn, Eric Johnson, Melayne McInnes and Elisabet Rutström (2005b). "Risk Aversion and Incentive Effects: Comment," *American Economic Review*, 95(3): 900-904.
- [37] Harrison, Glenn, Steven J. Humphrey and Arjan Verschoor (2009). "Choice Under Uncertainty: Evidence from Ethiopia, India and Uganda," *The Economic Journal*, 119, Forthcoming.
- [38] Henrich, Joseph and Richard McElreath (2002). "Are Peasants Risk-Averse Decision Makers?," *Current Anthropology*, 43(1): 172-181.
- [39] Hey, John, Andrea Morone and Ulrich Schmidt (2007). "Noise and Bias in Eliciting Preferences," *Discussion Paper in Economics* No. 2007/04, Department of Economics and Related Studies, University of York, UK.
- [40] Hey, John D. and Chris Orme (1994). "Investigating Generalizations of Expected Utility Theory Using Experimental Data," *Econometrica*, 62(6): 1291-1326.
- [41] Holt, Charles and Susan Laury (2002). "Risk Aversion and Incentive Effects," *American Economic Review*, 92(5): 1644-1655.
- [42] Holt, Charles and Susan Laury (2005). "Risk Aversion and Incentive effects: New Data without Order Effects," *American Economic Review*, 95(3): 902-904.
- [43] Jacobson, Sarah and Ragan Petrie (2009). "Learning from Mistakes: What Do Inconsistent Choices Over Risk Tell Us?," *Journal of Risk and Uncertainty*, 38(2): 143-158.
- [44] Kahneman, Daniel and Amos Tversky (1979). "Prospect Theory: An Analysis of Decision Under Risk," *Econometrica*, 47(2): 263-292.
- [45] Kachelmeir, Steven J. and Mohamed Shehata (1992). "Examining Risk Preferences Under High Monetary Incentives: Experimental Evidence from the People's Republic of China," *American Economic Review*, 82(5): 1120-1141.
- [46] Lattimore, Pamela K., Joanna R. Baker and Ann D. Witte (1992). "The Influence of Probability on Risky Choice," *Journal of Economic Behavior and Organization*, 17: 377-400.

- [47] Liu, Elaine (2008). "Time to Change What to Sow: Risk Preferences and Technology Adoption Decisions of Cotton Farmers in China," Job Market Paper, Princeton University, Department of Economics.
- [48] List, John (2004). "Neoclassical Theory versus Prospect Theory: Evidence from the Marketplace," *Econometrica*, 72(2): 615-625.
- [49] Loomes, Graham, Peter G. Moffat and Robert Sugden (2002). "A Microeconomic Test of Alternative Stochastic Theories of Risky Choice," *Journal of Risk and Uncertainty*, 24(2): 103-130.
- [50] Luce, R. Duncan (1959). *Individual Choice Behavior; A Theoretical Analysis*. New York: Wiley.
- [51] Luce, R. Duncan (2000). *Utility of Gains and Losses: Measurement-Theoretical and Experimental Approaches*. Mahwah, N.J.; London: Lawrence Erlbaum Associates (Scientific Psychology Series).
- [52] Maddala, G.S. (1983). *Limited-dependent and qualitative variables in econometrics*. Cambridge, New York: Cambridge University Press.
- [53] McLachlan, Geoffrey and David Peel (2000). *Finite Mixture Models*, New York: Wiley.
- [54] Prelec, Drazen (1998). "The Probability Weighting Function," *Econometrica*, 66: 497-527.
- [55] Quiggin, John (1982). "A Theory of Anticipated Utility," *Journal of Economic Behavior and Organization*, 3(4): 323-343.
- [56] Rieger, Marc Oliver and Mei Wang (2006). "Cumulative Prospect Theory and the St. Petersburg Paradox," *Economic Theory*, 28: 665-679.
- [57] Schechter, Laura (2007). "Risk Aversion and Expected Utility Theory: A Calibration Exercise," *Journal of Risk and Uncertainty*, 35(1): 67-76.
- [58] Starmer, Chris (2000). "Developments in Non-Expected Utility Theory: The Hunt for a Descriptive Theory of Choice Under Risk," *Journal of Economic Literature*, 38: 332-382.
- [59] Tanaka, Tomomi, Colin Camerer and Quang Nguyen (2009): "Risk and Time Preferences: Linking Experimental and Household Survey Data from Vietnam," *American Economic Review*, forthcoming.
- [60] Train, Kenneth (2003). *Discrete Choice Methods with Simulation*. Cambridge University Press.
- [61] Tversky, Amos and Daniel Kahneman (1992). "Advances in Prospect Theory: Cumulative Representations of Uncertainty," *Journal of Risk and Uncertainty*, 5(4): 297-323.
- [62] Wilcox, Nathaniel (2008). "Stochastic Models for Binary Discrete Choice Under Risk: A Critical Primer and Econometric Comparison," In: Cox, James and Glenn Harrison (eds.), *Risk Aversion in Experiments*. Bingley, U: Emerald, Research in Experimental Economics, Volume 12 (pp. 197-292).
- [63] Wu, George and Richard Gonzalez (1996). "Curvature of the probability weighting function," *Management Science*, 42: 1676-1690.

Appendix A. Our Research Area

Figure A.1: Maps of Peru (Right) and Pisco Valley (Left)



Appendix B. Experimental Subjects: Basic Statistics

Table B.1 Summary Statistics

Variable	Mean	Std. Dev.	N
<i>Dependent variable</i>			
Insured loan take-up rate (high stakes)	0.57	0.49	378
<i>Demographics and Education</i>			
Age (years)	54.9	13.3	367
Aged less than 40	0.14	0.35	367
Aged between than 40 and 50	0.19	0.39	367
Aged between than 50 and 60	0.33	0.47	367
Aged over 60	0.33	0.47	367
Female (Yes=1)	0.27	0.44	367
Education (years)	6.33	4.11	365
Illiterate	0.05	0.23	365
Some year of primary school	0.51	0.50	365
Some year of secondary school	0.34	0.47	365
Completed more than secondary school	0.09	0.29	365
Financial literacy indicator ¹	0.54	0.20	378
<i>Agriculture and Assets</i>			
Farming experience (years)	23.9	12.7	368
Size of owned agricultural plot (hectares)	6.03	5.57	367
Size of cultivated land (hectares) ²	5.01	4.13	365
Planted cotton (Yes=1) ²	0.83	0.39	368
Cotton yields (quintals per hectare) ²	46.8	14.8	293
Self-reported value of owned ag plot (000 Soles) ³	7.43	8.78	307
Self-reported value of house (000 Soles) ⁴	15.92	21.0	321
Self-reported value of assets (000 Soles) ⁵	20.42	21.8	362
<i>Networks</i>			
Talked to somebody in her “valley” about farming(Yes=1)	0.67	0.47	378
Number of “valley” members in her agricultural network	1.73	1.61	378
Has ever been a local authority (Yes=1)	0.39	0.49	365
Belongs to a farmer association (Yes=1)	0.29	0.46	364
<i>Credit and Insurance</i>			
Got credit for farming activities (Yes=1) ²	0.61	0.49	378
Got formal credit (Yes=1)	0.39	0.49	232
Got credit from cotton mills (Yes=1)	0.27	0.45	232
Has any other type of insurance (Yes=1)	0.44	0.50	378

Table B.1 Summary Statistics (continued)

Variable	Mean	Std.Dev.	N
<i>Experimental Variables</i>			
Risk rationed (Baseline Game) (Yes=1)	0.24	0.43	378
Risk parameter estimate, EUT ⁶	0.45	0.29	365
Risk parameter estimate, CPT ⁶	0.74	0.32	365
Probability weighting parameter estimate, CPT ⁶	0.54	0.21	365
Overweighting (Yes=1) ⁷	0.20	0.40	365
Faced a bad year, low stakes Insurance Game(Yes=1) ⁸	0.29	0.45	378
Drew a black chip in last low-stake round, Insurance Game	0.08	0.28	378
Winnings from all the farming games (Soles)	16.8	2.6	378
Winnings from low stakes rounds, Insurance Game (Soles)	3.04	0.85	378

¹ Indicator calculated using knowledge of insurance and loan project, as well as a self-reported degree of . comprehension. ² It refers to the 2007-2008 farming season. ³ The question was “how much do you think you’d have to pay to rent a hectare of land with similar characteristics to your main parcel. ⁴ The question was “how much do you think you’d have to pay to buy a house with similar characteristics to your own”.

⁵ *Wealth* includes the values of land and house. ⁶ EUT (CPT): Risk estimate assuming Expected Utility Theory (Cumulative Prospect Theory). ⁷ Overweighting means that the weighting parameter is greater than 0.7. ⁸ It means that subject chose the *uninsured* loan project the season a black chip was drawn in her valley (hence loan default).

Appendix C. Experiment Instructions for the Risk Game

What follows are the instructions given to subjects. Recall that the monitor used slides that were projected on a wall. The parts in [square brackets] are directions the monitor should follow, and the parts in {curly brackets} are notes to clarify information provided to subjects or procedures followed.

The objective of this activity is to win money. There are 4 possible prizes: 90, 1400, 3500, and 1800 soles. How are you going to win these prizes?

[Show slide 1: Lottery prizes]

To win these prizes, first you will have to choose between two options, the option Sol and the option Luna. If you choose the option Sol, you can win a maximum prize of 1,800 soles and a minimum prize of 1,400 soles. And if you choose the option Luna, you can win a maximum prize of 3,500 soles and a minimum prize of 90 soles. Note that with the option Sol the difference between the maximum and minimum prize is small, while it is large in the case of the option Luna.

In addition, in the option Sol the maximum prize of 1,800 soles is smaller than the maximum prize of 3,500 soles in the option Luna, and the minimum prize of 1,400 soles in the option Sol is greater than the minimum prize of 90 soles in the option Luna.

Thus, you will pick between Sol and Luna in 10 rows, one after another. Once you have picked between Sol and Luna, the prize you will receive in each row will depend on the number that you obtain by rolling a 10-sided die like the one your assistants are now showing you. This die has 10 sides numbered from 1 to 10; that is, it is equally likely to get any of the 10 numbers.

We will now see an example of the prizes you may earn. Please look at the second row on page 7 in your binders {page 7 displayed the 10 decision rows}.

[Show slide 2: Lottery prizes in row 2]

[Pick a volunteer] Let's see... [Say name], in the second row, which do you prefer: Sol or Luna? {s/he will say Sol/Luna}. Now throw the die. The die is showing the number [1,2,3,...,or 10]. Since [say name] chose [Sol/Luna], we see that in the second row the prize that corresponds to the number [1,2,3,...,10] is [say amount]. Now, if [say name] would have chosen [Luna/Sol], the prize that corresponds to [1,2,3,...,10] is [say amount].

Let's do another example.

[Show slide 3: Lottery prizes in row 8]

[Pick another volunteer]. Let's see... [say name], in the eighth row, which do you prefer: Sol or Luna? [s/he will say Sol/Luna]. Now throw the die. The die is showing the number [1,2,3,...,10]. Since [say name] chose [Sol/Luna], we see that in the eighth row the prize that corresponds to the number [1,2,3,...,10] is [say amount]. Now, if [say name] would have chosen [Luna/Sol], the prize that corresponds to [1,2,3,...,10] is [say amount].

Now we will compare the prizes that could be obtained from the two rows we have used in our examples.

[Show slide 4: Lottery prizes in rows 2 and 8]

If throwing the die results in the number 4, those who chose Sol in row 2 would win 1,400 soles and those that chose Luna in row 2 would win 90 soles. In row 8, if throwing the die results in number 4, those who chose Sol would win 1,800 soles and those who chose Luna would win 3,500 soles.

Note that in row 8 there are more chances of winning the maximum prize than there are in row 2, in both Sol and Luna. This is because in row 8 there are 8 chances to win the maximum prize and only 2 chances to win the minimum prize, while in row 2, there are only 2 chances to win the maximum prize and 8 chances to win the minimum prize.

Now look at page 7 [Show slide 5: Lottery prizes in the 10 rows]... in which you can see that as one goes down the table, the chances to win the larger prize are greater; and the chances of winning the smaller prize are fewer.

Notice that in the last row, row 10, regardless of which number appears on the die, the prize you will receive will be 1,800 soles if you choose Sol and 3,500 soles if you choose Luna.

Practice Round

We are now going to do a practice round. Please look at your worksheets on page 7 of your binders. In this sheet you have to choose, for each of the 10 rows, between the option Sol and the option Luna, marking with an “X” on the drawing of the sun or the moon.

Please mark on your practice sheet for each of the 10 rows the option that you prefer in each row.

After that, we will choose one row to determine the prize you would win. Assistants, please begin the practice round in your valleys.

{Assistants advised the monitor when all choices in their valleys were made.}

Now, we will see what you would have won. To determine the row you will play, we will throw the 10-sided die in each valley. Assistants, throw the die one time in each of your valleys...

{Assistants advised the monitor when dice were rolled one time in their valleys.}

Now that we know which row was chosen in each valley, each of you will throw the 10-sided die to determine your prize. Assistants, have each farmer in your valley throw the die.

{Assistants advised the monitor when all choices in their valleys were made.}

{The monitor then chose two valleys to illustrate the procedure to determine winnings for this game.}

Let's see... [say name in valley 1], we are going to see what would have been your prize. First, which row was selected in your valley?... And in that row, did you pick Sol or Luna? ... And what number did you get when you threw the die?.....Then [say name] in row [say resulting row], with the option [Sol/Luna] and the number [say number of die roll], would win [say amount]. If [say name] would have chosen the other option [Luna/Sol], he would win [say amount].

Let's see... [say name in valley 2], we are going to see what would have been your prize. First, which row was selected in your valley?... And in that row, did you pick Sol or Luna?... And what

number did you get when you threw the die?.....Then [say name] in row [say resulting row], with the option [Sol/Luna] and the number [say number of die roll], would win [say amount]. If [say name] would have chosen the other option [Luna/Sol], he would win [say amount].

This was a practice round. Now we are going to do it for real money.

High Payout Round

How will we determine how much money you will win for participating in this activity? This will depend on the prizes obtained in the row that is chosen at random, in the same way that your earnings were selected in the earlier activities. To choose the row, you will throw a 10-sided die one time for each valley, just as you did in the practice round. We will pay you one sol in cash for each 600 soles of prize winnings. That is, the minimum amount that you could win is $90/600 = 0.15$ soles, and the maximum amount is $3,500/600 =$ almost 6 soles.

Please look at your sheet for this activity, on page 8 of your binders.

[show slide 6: lottery prizes in the 10 rows]

In this activity you have to choose, just like in the practice round, between the options Sol and Luna, marking with an “X” on the drawing of the sun or the moon in each row, from 1 to 10. Before we start, are there any questions?

{Pause for questions.}

Now, please, mark on your sheets for this activity in each one of the 10 rows the option that you prefer.

Appendix D. Risk Game Worksheet Sample

		SOL 	LUNA 																																																													
1		<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 90%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">1800</td> <td style="text-align: center;">1400</td> </tr> </table>	1		S/.	S/.	1800	1400	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 90%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">3500</td> <td style="text-align: center;">90</td> </tr> </table>	1		S/.	S/.	3500	90																																																	
	1																																																															
S/.	S/.																																																															
1800	1400																																																															
1																																																																
S/.	S/.																																																															
3500	90																																																															
	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 80%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">1800</td> <td style="text-align: center;">1400</td> <td style="text-align: center;">90</td> </tr> </table>	1	2		S/.	S/.	S/.	1800	1400	90	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 80%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">3500</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2		S/.	S/.	S/.	3500	90	90																																												
1	2																																																															
S/.	S/.	S/.																																																														
1800	1400	90																																																														
1	2																																																															
S/.	S/.	S/.																																																														
3500	90	90																																																														
3		<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 70%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">1800</td> <td style="text-align: center;">1400</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3		S/.	S/.	S/.	S/.	1800	1400	90	90	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 70%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">3500</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3		S/.	S/.	S/.	S/.	3500	90	90	90																																					
	1	2	3																																																													
S/.	S/.	S/.	S/.																																																													
1800	1400	90	90																																																													
1	2	3																																																														
S/.	S/.	S/.	S/.																																																													
3500	90	90	90																																																													
	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 10%; text-align: center;">4</td> <td style="width: 60%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">1800</td> <td style="text-align: center;">1400</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3	4		S/.	S/.	S/.	S/.	S/.	1800	1400	90	90	90	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 10%; text-align: center;">4</td> <td style="width: 60%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">3500</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3	4		S/.	S/.	S/.	S/.	S/.	3500	90	90	90	90																																
1	2	3	4																																																													
S/.	S/.	S/.	S/.	S/.																																																												
1800	1400	90	90	90																																																												
1	2	3	4																																																													
S/.	S/.	S/.	S/.	S/.																																																												
3500	90	90	90	90																																																												
5		<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 10%; text-align: center;">4</td> <td style="width: 10%; text-align: center;">5</td> <td style="width: 40%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">1800</td> <td style="text-align: center;">1400</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3	4	5		S/.	S/.	S/.	S/.	S/.	S/.	1800	1400	90	90	90	90	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 10%; text-align: center;">4</td> <td style="width: 10%; text-align: center;">5</td> <td style="width: 40%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">3500</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3	4	5		S/.	S/.	S/.	S/.	S/.	S/.	3500	90	90	90	90	90																									
	1	2	3	4	5																																																											
S/.	S/.	S/.	S/.	S/.	S/.																																																											
1800	1400	90	90	90	90																																																											
1	2	3	4	5																																																												
S/.	S/.	S/.	S/.	S/.	S/.																																																											
3500	90	90	90	90	90																																																											
	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 10%; text-align: center;">4</td> <td style="width: 10%; text-align: center;">5</td> <td style="width: 10%; text-align: center;">6</td> <td style="width: 30%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">1800</td> <td style="text-align: center;">1400</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3	4	5	6		S/.	S/.	S/.	S/.	S/.	S/.	S/.	1800	1400	90	90	90	90	90	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 10%; text-align: center;">4</td> <td style="width: 10%; text-align: center;">5</td> <td style="width: 10%; text-align: center;">6</td> <td style="width: 30%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">3500</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3	4	5	6		S/.	S/.	S/.	S/.	S/.	S/.	S/.	3500	90	90	90	90	90	90																				
1	2	3	4	5	6																																																											
S/.	S/.	S/.	S/.	S/.	S/.	S/.																																																										
1800	1400	90	90	90	90	90																																																										
1	2	3	4	5	6																																																											
S/.	S/.	S/.	S/.	S/.	S/.	S/.																																																										
3500	90	90	90	90	90	90																																																										
7		<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 10%; text-align: center;">4</td> <td style="width: 10%; text-align: center;">5</td> <td style="width: 10%; text-align: center;">6</td> <td style="width: 10%; text-align: center;">7</td> <td style="width: 20%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">1800</td> <td style="text-align: center;">1400</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3	4	5	6	7		S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	1800	1400	90	90	90	90	90	90	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 10%; text-align: center;">4</td> <td style="width: 10%; text-align: center;">5</td> <td style="width: 10%; text-align: center;">6</td> <td style="width: 10%; text-align: center;">7</td> <td style="width: 20%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">3500</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3	4	5	6	7		S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	3500	90	90	90	90	90	90	90													
	1	2	3	4	5	6	7																																																									
S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.																																																									
1800	1400	90	90	90	90	90	90																																																									
1	2	3	4	5	6	7																																																										
S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.																																																									
3500	90	90	90	90	90	90	90																																																									
	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 10%; text-align: center;">4</td> <td style="width: 10%; text-align: center;">5</td> <td style="width: 10%; text-align: center;">6</td> <td style="width: 10%; text-align: center;">7</td> <td style="width: 10%; text-align: center;">8</td> <td style="width: 10%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">1800</td> <td style="text-align: center;">1400</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3	4	5	6	7	8		S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	1800	1400	90	90	90	90	90	90	90	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 10%; text-align: center;">4</td> <td style="width: 10%; text-align: center;">5</td> <td style="width: 10%; text-align: center;">6</td> <td style="width: 10%; text-align: center;">7</td> <td style="width: 10%; text-align: center;">8</td> <td style="width: 10%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">3500</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3	4	5	6	7	8		S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	3500	90	90	90	90	90	90	90	90								
1	2	3	4	5	6	7	8																																																									
S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.																																																								
1800	1400	90	90	90	90	90	90	90																																																								
1	2	3	4	5	6	7	8																																																									
S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.																																																								
3500	90	90	90	90	90	90	90	90																																																								
9		<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 10%; text-align: center;">4</td> <td style="width: 10%; text-align: center;">5</td> <td style="width: 10%; text-align: center;">6</td> <td style="width: 10%; text-align: center;">7</td> <td style="width: 10%; text-align: center;">8</td> <td style="width: 10%; text-align: center;">9</td> <td style="width: 10%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">1800</td> <td style="text-align: center;">1400</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3	4	5	6	7	8	9		S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	1800	1400	90	90	90	90	90	90	90	90	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 10%; text-align: center;">4</td> <td style="width: 10%; text-align: center;">5</td> <td style="width: 10%; text-align: center;">6</td> <td style="width: 10%; text-align: center;">7</td> <td style="width: 10%; text-align: center;">8</td> <td style="width: 10%; text-align: center;">9</td> <td style="width: 10%;"></td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">3500</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3	4	5	6	7	8	9		S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	3500	90	90	90	90	90	90	90	90	90	
	1	2	3	4	5	6	7	8	9																																																							
S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.																																																							
1800	1400	90	90	90	90	90	90	90	90																																																							
1	2	3	4	5	6	7	8	9																																																								
S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.																																																							
3500	90	90	90	90	90	90	90	90	90																																																							
	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 10%; text-align: center;">4</td> <td style="width: 10%; text-align: center;">5</td> <td style="width: 10%; text-align: center;">6</td> <td style="width: 10%; text-align: center;">7</td> <td style="width: 10%; text-align: center;">8</td> <td style="width: 10%; text-align: center;">9</td> <td style="width: 10%; text-align: center;">10</td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">1800</td> <td style="text-align: center;">1400</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3	4	5	6	7	8	9	10	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	1800	1400	90	90	90	90	90	90	90	90	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 10%; text-align: center;">1</td> <td style="width: 10%; text-align: center;">2</td> <td style="width: 10%; text-align: center;">3</td> <td style="width: 10%; text-align: center;">4</td> <td style="width: 10%; text-align: center;">5</td> <td style="width: 10%; text-align: center;">6</td> <td style="width: 10%; text-align: center;">7</td> <td style="width: 10%; text-align: center;">8</td> <td style="width: 10%; text-align: center;">9</td> <td style="width: 10%; text-align: center;">10</td> </tr> <tr> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> <td style="text-align: center;">S/.</td> </tr> <tr> <td style="text-align: center;">3500</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> <td style="text-align: center;">90</td> </tr> </table>	1	2	3	4	5	6	7	8	9	10	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	3500	90	90	90	90	90	90	90	90	90		
1	2	3	4	5	6	7	8	9	10																																																							
S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.																																																							
1800	1400	90	90	90	90	90	90	90	90																																																							
1	2	3	4	5	6	7	8	9	10																																																							
S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.	S/.																																																							
3500	90	90	90	90	90	90	90	90	90																																																							

Appendix E. Regression Results for the Restricted Sample

Table E.1: EUT Estimates with Heterogeneous Subjects and Normal Fechner Errors
(*Restricted Sample: $N = 2,650$*)

Coefficient	Variable	Estimate	Std.Error	<i>p-value</i>	95% Conf. Interval
r_i^{EUT}	Intercept	0.18	0.09	0.04	0.00 0.35
	Female	0.04	0.07	0.58	-0.10 0.18
	Young (Age < 40)	-0.18	0.10	0.06	-0.38 0.01
	Middle (Age: [50-60])	-0.15	0.09	0.09	-0.32 0.03
	Old (Age > 60)	-0.02	0.09	0.80	-0.20 0.16
	Illiterate	-0.28	0.14	0.04	-0.56 0.00
	Some secondary	-0.07	0.07	0.36	-0.22 0.08
	Skilled (> sec. educ.)	-0.10	0.12	0.41	-0.32 0.13
	Low Pisco	0.12	0.07	0.09	-0.02 0.27
	High Pisco	0.10	0.11	0.33	-0.10 0.31
<i>Predicted r^{EUT} at average values</i>		<i>0.11</i>			
σ_u	Intercept	2.00	0.16	0.00	1.70 2.32

Notes: S.E. clustered at the individual level. The omitted category for *age* is for those aged 40-50. The omitted category for *education* is for those with some primary education.

Table E.2: CPT Estimates Assuming Heterogeneous Subjects and Fechner Normal Errors
(Restricted Sample: $N = 2,650$)

Coefficient	Variable	Estimate	Std.Error	p -value	95% Conf. Interval
r_i^{CPT}	Intercept	0.70	0.11	0.00	0.50 0.92
	Female	0.06	0.06	0.32	-0.06 0.19
	Young (Age < 40)	-0.21	0.09	0.03	-0.39 -0.03
	Middle (Age: [50-60])	-0.09	0.08	0.28	-0.26 0.07
	Old (Age > 60)	-0.03	0.08	0.72	-0.19 0.13
	Illiterate	-0.12	0.12	0.34	-0.36 0.12
	Some secondary educ.	-0.15	0.07	0.03	-0.29 -0.01
	Skilled (> sec. educ.)	-0.23	0.11	0.03	-0.43 -0.02
	Low Pisco	0.07	0.06	0.29	-0.06 0.19
	High Pisco	-0.06	0.09	0.49	-0.23 0.11
<i>Predicted r^{CPT} at average values</i>		<i>0.59</i>			
γ_i	Intercept	0.39	0.05	0.00	0.30 0.48
	Female	-0.03	0.03	0.38	-0.09 0.03
	Young (Age < 40)	0.13	0.07	0.05	-0.0001 0.26
	Middle (Age: [50-60])	0.01	0.04	0.76	-0.07 0.10
	Old (Age > 60)	0.01	0.04	0.80	-0.07 0.09
	Illiterate	-0.02	0.05	0.67	-0.12 0.08
	Some secondary educ.	0.08	0.04	0.05	-0.0003 0.15
	Skilled (> sec. educ.)	0.15	0.07	0.04	0.01 0.28
	Low Pisco	0.01	0.03	0.86	-0.06 0.07
	High Pisco	0.10	0.06	0.09	-0.02 0.22
<i>Predicted γ at average values</i>		<i>0.46</i>			
σ_μ	Intercept	0.66	0.09	0.00	0.50 0.83

Notes: S.E. clustered at the individual level. The omitted category for *age* is for those aged 40-50. The omitted category for *education* is for those with some primary education.

Table E.3: Mixture Model Estimates: EUT and CPT with Fechner Normal Errors
(*Restricted Sample: $N = 2,730$*)

Variable		Expected Utility	Prospect Theory
r	Intercept	0.42***	1.14***
	Standard Error	0.08	0.03
	95% C. Interval	[0.26, 0.58]	[0.85, 1.42]
γ	Intercept	n.a.	0.27***
	Standard Error		0.03
	95% C. Interval		[0.22, 0.33]
θ	Intercept	0.29***	0.71***
	Standard Error	0.04	0.04
	95% C. Interval	[0.21, 0.36]	[0.64, 0.79]
σ_μ	Intercept		0.31***
	Standard Error		0.07
	95% C. Interval		[0.18, 0.44]

n.a.: not applicable.

*** denotes significance at 1 % level. S.E. clustered at the individual level.

Appendix F. Extensions: Alternative Error Type and Weighting Function

F.1 Luce's (1959) Errors

Previous studies have found that risk estimates are sensitive to the errors' specification (e.g., Wilcox, 2008). In this Appendix, we will examine whether this is the case in our sample by estimating an alternative error type, due to Luce (1959). Luce errors' standard deviation will be denoted as σ_{μ}^{Luce} . In this case, the stochastic choice rule under EUT is given by the following expression:

$$\Delta SEU_i^{Luce} = \frac{(EU_i^S)^{1/\sigma_{\mu_i}^{Luce}}}{(EU_i^S)^{1/\sigma_{\mu_i}^{Luce}} + (EU_i^L)^{1/\sigma_{\mu_i}^{Luce}}}, \quad (22)$$

which is bounded between 0 and 1, and can thus be used to predict probabilities in its ratio form: $\Delta SEU_i^{Luce} = \Pr(\text{choosing lottery } S)$. The parameter $\sigma_{\mu_i}^{Luce} > 0$ represents the noise that can be understood as the deviation from the correct choices. In the limit, when $\sigma_{\mu_i}^{Luce} \rightarrow 0$, the lottery choices would reflect exactly the differences in expected utilities of lotteries, and when $\sigma_{\mu_i}^{Luce} \rightarrow \infty$, choices would rather be random (and $\Delta SEU_i^{Luce} \rightarrow 0.5$). Obviously, when $\sigma_{\mu_i}^{Luce} \rightarrow 1$, we would be back to the ratio form:

$$\Delta EU_i^{Luce} = \frac{EU_i^S}{EU_i^S + EU_i^L}, \quad (23)$$

which is the decision rule under the assumption of Luce standard normal error (i.e., Luce error with unit variance). The corresponding decision rule under CPT would have the same ratio form, but with the cumulative prospect utility (*CPU*) replacing the expected utility (*EU*) shown in the prior eqn.

The conditional likelihood function under EUT assuming that Luce errors are *normally* distributed with variance, $(\sigma_{\mu_i}^2)^{Luce}$, will be then:

$$l^{\text{EU}} = \sum_{i=1}^N \ln L_i^{\text{EU}}(r_i^{\text{EU}}, \sigma_{\mu_i}^{Luce}, I_i^m, X_i) = \sum_{i=1}^N \sum_{m=1}^{10} [\ln(\Delta SEU_{i,m}^{Luce})^{I_i^m} + \ln(1 - \Delta SEU_{i,m}^{Luce})^{1-I_i^m}], \quad (24)$$

where I_i^m is equal to 1 when *S* is chosen in row *m*; and 0 otherwise. Note that since the choice rule is already in a cumulative probability form, we do not need to assume a particular density function to transform the values from the choice rule into a [0,1] value, and we therefore simply take logs to ΔSEU_i^{Luce} to construct the likelihood function. As we did earlier with the Fechner error, we will estimate risk preferences under the heterogeneous agent case (i.e., with $r_i = r_0 + f[X_i]$), where X_i denotes the individual characteristics.

Furthermore, we will use the corresponding cumulative prospect utility (CPU_i) from eqn.[11]) instead of the EU_i in eqn.[22] in order to estimate the parameters under CPT.

Tables F.1 and F.2 below show the maximum likelihood estimates under EUT and CPT for

the homoskedastic Luce errors. In sum, subjects appear moderately risk averse ($\hat{r}^{\text{EUT}} = 0.56$ and $\hat{r}^{\text{CPT}} = 0.72$), and higher education is negatively correlated with risk aversion in both models. Moreover, the estimated S-shaped weighting function implies a substantial overweighting of low probabilities and underweighting of large probabilities. These risk parameter estimates are highly correlated with the ones obtained under the Fechner errors assumption (the correlation coefficients are 0.74 under EUT, and 0.94 under CPT). Lastly, the magnitude of the estimated standard deviation ($\hat{\sigma}_{\mu}^{\text{Luce}}$) is in accordance with prior studies that use the same errors structure (e.g., Holt and Laury, 2002).

Table F.1: EUT Estimates Assuming Heterogeneous Subjects & Luce Errors

Coefficient	Variable	Estimate	Std.Error	<i>p-value</i>	95% Conf. Interval	
r_i^{EU}	Intercept	0.75	0.12	0.00	0.52	0.98
	Female	0.06	0.08	0.77	-0.09	0.21
	Young (Age < 40)	-0.15	0.19	0.41	-0.52	0.21
	Middle (Age: [50-60])	0.08	0.12	0.47	-0.15	0.32
	Old (Age > 60)	-0.002	0.13	0.99	-0.25	0.25
	Illiterate	0.10	0.12	0.44	-0.15	0.34
	Some secondary	-0.46	0.21	0.03	-0.87	-0.05
	Skilled (> sec. educ.)	-0.70	0.19	0.00	-1.08	-0.32
<i>Predicted r at average values</i>		<i>0.56</i>				
$\sigma_{\mu}^{\text{Luce}}$	Intercept	0.49	0.07	0.00	0.36	0.62
N		3,650				

Table F.2: CPT Estimates Assuming Heterogeneous Subjects & Luce Errors

Coefficient	Variable	Estimate	Std.Error	<i>p-value</i>	95% Conf. Interval
r_i^{CPT}	Intercept	0.73	0.12	0.00	0.50 0.96
	Female	0.05	0.05	0.25	-0.04 0.15
	Young (Age < 40)	-0.11	0.10	0.24	-0.30 0.08
	Middle (Age: [50-60])	0.09	0.07	0.20	-0.05 0.23
	Old (Age > 60)	0.15	0.11	0.19	-0.07 0.37
	Illiterate	-0.18	0.10	0.07	-0.38 0.02
	Some secondary educ.	-0.16	0.13	0.20	-0.42 0.09
	Skilled (> sec. educ.)	-0.41	0.14	0.00	-0.67 -0.14
	Low Pisco	0.06	0.07	0.34	-0.07 0.19
	High Pisco	-0.06	0.07	0.39	-0.18 0.07
<i>Predicted r^{CPT} at average values</i>		<i>0.72</i>			
γ_i	Intercept	0.45	0.05	0.00	0.35 0.55
	Female	-0.001	0.04	0.97	-0.08 0.08
	Young (Age < 40)	0.08	0.05	0.11	-0.02 0.19
	Middle (Age: [50-60])	-0.05	0.05	0.29	-0.14 0.04
	Old (Age > 60)	0.07	0.13	0.61	-0.19 0.32
	Illiterate	-0.17	0.16	0.28	-0.48 0.14
	Some secondary educ.	0.01	0.06	0.85	-0.10 0.12
	Skilled (> sec. educ.)	0.14	0.09	0.13	-0.04 0.32
	Low Pisco	0.11	0.09	0.23	-0.07 0.29
	High Pisco	0.21	0.09	0.02	0.03 0.39
<i>Predicted γ at average values</i>		<i>0.54</i>			
σ_μ^{Luce}	Intercept	0.17	0.06	0.00	0.04 0.29
N		3,650			

F.2 Prelec’s (1998) Weighting Function

Our limited sample size prevents the estimation of some of the most popular two-parameter weighting functions (Rieger and Wang, 2006; Lattimore et al., 1992; and Prelec, 1998). In addition to the curvature parameter, these functions allow for the estimation of the “elevation” parameter. While the curvature alludes to the sensitivity to probability changes, the elevation is referred to the *overall* risk aversion, or to how attractive subjects find gambling (Gonzalez and Wu, 1999).

While we could obtain the parameter estimates without including covariates in the parameter equations (i.e., estimating an overall parameter for the entire sample), including covariates usually results in lack of convergence or in large estimated standard errors. For that reason, we will only estimate a simplified version of Prelec’s (1998) weighting function:

$$\mathbf{w}_i(p) = \exp\{-\eta(-\ln p)^{\phi_i}\}, \quad \phi_i, \eta > 0, \quad (25)$$

that imposes $\eta = 1$ in the prior eqn. To be clear, the parameter η measures the elevation, and ϕ , the curvature. For a given elevation parameter, a lower $\phi_i < 1$ Prelec’s function becomes more curved,

which implies that the function exhibits more rapidly diminishing sensitivity to probabilities near 0 or 1. Further note that as $\phi_i \rightarrow 1$, then $\mathbf{w}_i(p) \rightarrow p$, and thus prospect theory will collapse to the EUT framework.

Table F.3 below reports the maximum likelihood estimates using Prelec weighting function and the original version of prospect theory (i.e., not defined over cumulative probabilities). In addition to evidence of moderate risk aversion (evaluated at the covariates mean, $\hat{r}_i = 0.44$), we find evidence of an inverse S-shape weighting function. In the equation for the weighting function parameter, ϕ_i , we included education expressed in years in the curvature parameter equation, because only doing so let this variable enter with a significant sign in the regression; otherwise (i.e., expressing education levels with dummy variables), in addition to insignificant coefficients, the standard errors estimated result very large. These results, which validate the existence of an inverse S-shaped weighting function, should be seen as preliminary, since the parameter estimates sharply change when we adopt alternative model specifications.

Table F.3: CPT Estimates Assuming Heterogeneous Subjects & Fechner Errors
With Prelec Weighting Function

Coefficient	Variable	Estimate	Std.Error	<i>p-value</i>	95% Conf. Interval
r_i^{PT}	Intercept	0.41	0.14	0.00	0.14 0.68
	Female	0.04	0.11	0.72	-0.18 0.26
	Young (Age < 40)	-0.14	0.15	0.36	-0.44 0.16
	Middle (Age: [50-60])	0.04	0.14	0.74	-0.22 0.31
	Old (Age > 60)	0.09	0.17	0.57	-0.23 0.42
	Illiterate	-0.32	0.28	0.25	-0.87 0.22
	Some secondary educ.	-0.16	0.13	0.23	-0.41 0.10
	Skilled (> sec. educ.)	-0.41	0.17	0.02	-0.75 -0.07
	Low Pisco	0.17	0.11	0.13	-0.05 0.39
	High Pisco	0.36	0.29	0.21	-0.21 0.93
<i>Predicted r^{PT} at average values</i>		<i>0.44</i>			
$\phi_i^\dagger$	Intercept	0.18	0.12	0.12	-0.05 0.41
	Female	0.52	0.24	0.03	0.04 0.99
	Education in years	0.59	0.02	0.00	0.55 0.63
<i>Predicted $\phi^\dagger$ at average values</i>		<i>0.71</i>			
σ_μ	Intercept	2.15	0.18	0.00	1.79 2.51
N		3,650			

Note: This weighting function was estimated using the original (not cumulative) prospect theory.

† The coefficients on female and education were computed using the Delta method. The original estimate was ψ , with $\phi = 1/(1 + \exp(\psi))$. Similarly, the predicted ϕ was obtained by plugging the predicted ψ (computed at the average of female and education) into the prior formula.

Furthermore, we picture the weighting function associated to the value of $\phi = 0.71$ in the Figure 4, where we also depict the TK,s (1992) weighting function for the CPT case with Fechner errors

(reported in Table 3). In choosing one weighting function over another, it would be interesting to test which one performs better. Such analysis is deferred to future work.

Figure 4: Tversky and Kahneman and Prelec Weighting Functions
With Fechner Errors and Heterogeneous Subjects

