

Original Research Article

Explaining Heterogeneity in Risk Preferences Using a Finite Mixture Model

Narges Hajimoladarvish*

Received: 14 Jun 2020

Approved: 27 Feb 2022

This paper studies the effect of the space (distance) between lotteries' outcomes on risk-taking behavior and the shape of estimated utility and probability weighting functions. Previously investigated experimental data shows a significant space effect in the gain domain. As compared to low spaced lotteries, high spaced lotteries are associated with higher risk aversion for high probabilities of gain and higher risk-seeking for low probabilities of gain. Hence, the investigation is carried under cumulative prospect theory that respects framing effect and characterizes risk attitudes with respect to probabilities and outcomes. The observed certainty equivalents of lotteries are assumed to be driven by cumulative prospect theory. To estimate the parameters of cumulative prospect theory with maximum likelihood, the usual error term is added. The cumulative prospect theory is incapable of explaining the space effect as its parameters cannot explain the average behavior. Taking account of heterogeneity, a two-component mixture model shows that behavioral parameters of around 25% of the sample can explain the observed differences in relative risk aversions. The results confirm the previous findings of aggregation bias associated with representative-agent models. Furthermore, the results have implications for experimental designs as high space between lotteries' outcomes is required to guarantee the curvature of utility functions.

Keywords: Space Effect, Cumulative Prospect Theory, Decision Making Under Risk, Finite Mixture Model.

JEL Classification: C91, D81.

1 Introduction

There has been growing interest in examining different sources of systematic variations in individuals' behaviour under risk (Etchart-Vincent 2004, 2009; Feltovich et al. 2012; Einav et al. 2012; Fehr-Duda et al. 2010). In economics, risky situations are modeled by lotteries; hence it is reasonable to study the characteristic of lotteries and their effect on risk-taking behaviour. This paper

* Department of Economics, Faculty of Social Sciences and Economics, Alzahra University, Tehran, Iran; n.moladarvish@alzahra.ac.ir

aims to examine the effect of the space between lotteries' outcomes on individuals' risk attitudes and structural estimation of preferences by incorporating the heterogeneity within and among individuals.

The space effect is a relevant area of research as, for many people, financial investments are usually triggered or characterized by a large difference between potential outcomes. Etchart-Vincent's (2009) study suggests that as compared to low-spaced lotteries, high-spaced lotteries tend to increase probability overweighting. Etchart-Vincent (2009) showed that the impact of space is both opposite to and stronger than the impact of level in the loss domain. The effect of the level of outcomes on risk attitudes has been debated for decades. Hence, it is interesting to examine the space effect and search for its plausible explanation.

Furthermore, when we want to elicit utility and probability weighting functions from lottery choices, there is no consensus on the choice of payoffs distribution, as the theory makes no distinction. For making comparisons, outcomes are usually transformed to be between 0 and 1 in the gain domain and -1 and 0 in the loss domain. This transformation is needed to reduce the effect of a payoff distribution. This study examines the effect of the space between the outcomes on risky choices after transforming the data and including appropriate error specifications. It is especially important for experimental designs where changes in payoffs' distribution might confound conclusions based on hypothesis testing.

Psychologists have demonstrated the sensitivity of individuals' risk-taking behaviour to their daily experience (Ungemach et al. 2011), the level of outcomes (Stewart et al. 2015), the size of the sample (Hilbig and Glockner 2011), and the information sources (Hertwig et al. 2004). Stewart et al. (2006) proposed a new theory called the decision by sampling to take account of the distribution of payoffs. Since the essence of the decision by sampling is based on outcomes' rank in the sample and subjects' memory recall, it cannot accumulate the space effect either. However, the important question is whether the shape of utility and probability weighting functions or their degree of concavity/convexity changes according to subjects' level of risk aversion determined by their choices. Only then, we can dissociate the variability of choices due to the characteristic of lotteries from different sources of heterogeneity.

Experimental data are subject to heterogeneity within and among individuals and differentiating the choices' noise from true variations can depend on the error specification (Loomes 2005). Variability within individuals' choices, which concerns about the mood of subjects, their

characteristics, their mistakes, and inattentiveness, is modelled by adding the appropriate stochastic component to deterministic choice models. The heterogeneity among individuals due to the existence of distinct behavioural types can be modelled by means of finite mixture models. In other words, the substantive heterogeneity in individual risk-taking behaviour makes a single representative-agent model inadequate to describe behaviour.

I use previously investigated data by Bruhin et al. (2010) that elicit certainty equivalents of two outcome lotteries. I divide the data into two treatments. High spaced treatment in which the space between the outcomes of the lotteries is higher than 30, which is the median of outcomes, and low spaced treatment where the space between the outcomes is less or equal to 30. Then, the behaviour of subjects are classified according to their relative risk premia. The space effect shows an increase in relative risk aversion for high probabilities of gain and a decrease in relative risk aversion for low probabilities of gain. In the loss domain, there is no coherent picture. Thus, a single representative model that does not take account of framing effect and characterizes risk attitudes irrespective of probabilities cannot accumulate all the observed differences. Hence, I carry the investigation under cumulative prospect theory with different error specifications. The results suggest that the space effect cannot be explained by changes in the parameters of cumulative prospect theory due to the heterogeneity among individuals.

As recent studies suggest that behaviour might be better characterised with more than one decision making processes, I fit the data from each treatment to a two-component mixture model. It is as if we assume there are two types of individuals. Considering heterogeneity through the mixture model, I find that the minority group's risk-taking behavior can explain the space effect. This finding shows that sometimes it is worth introducing extra parameters to gain explanatory power and tackle the problem associated with a single representative model that cannot explain the variation in the data.

The schematic outline of the paper is the following. Section 2 presents a literature review. Section 3 explains the experiment and the data. Section 4 outlines the econometric model used to measure risk attitudes. Section 5 provides the results. Section 6 discusses the results and concludes.

2 Literature Review

Economic models strive to take account of the heterogeneity of the population by different means. When observed outcomes are drawn from a finite mixture of distributions, finite mixture models are used to model population heterogeneity while estimating structural parameters. McLachlan and Peel

(2004) provide an excellent explanation of finite mixture models and how to estimate them. Here, I present the basic definition of finite mixture models from Cameron and Trivedi (2005). If the sample is a probabilistic mixture from two sub-population with *pdf* $f_1(y|\mu_1(x))$ and $f_2(y|\mu_2(x))$, then $\pi f_1(.) + (1 - \pi)f_2(.)$, where $0 \leq \pi \leq 1$, defines a two-component finite mixture model. That is, observations are drawn from f_1 and f_2 , with probabilities π and $1 - \pi$ respectively. The parameters to be estimated are (π, μ_1, μ_2) . The parameter π may be treated as constant or maybe further parameterized. Thus, we think of types of individuals, those that come from $f_1(.)$, and those that come from $f_2(.)$. The interpretation is that a linear combination of densities makes a good approximation to the observed distribution of y . In many applications such as this study, μ_1 and μ_2 are further parameterized. Although generalization to additive mixtures with three or more components is in principle straightforward, it is subject to potential problems of identifiability of the components. Therefore, it is very helpful in an empirical application if the components have a natural interpretation. We can simply think of each sub-population as a type or as a representation of population heterogeneity.

The estimation of the finite mixture model may be carried out under the assumption of either known or unknown number of components, which will be denoted by $c = 1, \dots, C$. If the fraction π_c s are known, maximum likelihood estimates of the component distributions can be carried out. Usually, the proportions π_c are unknown and the estimation involves both the π_c and the parameters of each component. The latter is by iterative expectation-maximization (EM) algorithm (Dempster et al., 1977) that is a general method of finding the maximum likelihood estimate from a given data set when the data is incomplete.

Mixture models are widely used throughout applied statistics, including labour economics, industrial organization, and computer science (Adams, 2016). Finite mixture models are featured in many areas of econometrics and several strategies for the identification of finite mixtures have been developed (Henry et al. 2014). For example, Eckstein and Wolpin (1990) and Keane and Wolpin (1997) model individual heterogeneity in labor markets by a finite number of types. Finite mixture models are used in experimental economics as well. In public goods experiments, Bardsley and Moffatt (2007) allow four types of individuals; reciprocators, strategists, altruists, and free-riders. They estimate a finite mixture 2-limit tobit with tremble and show while most subjects act selfishly, a substantial proportion are reciprocal with altruism playing only a marginal role.

Similarly, some researchers have examined the possibility of having different subjects following different decision theories for choices under risk (Harrison and Rutstrom, 2009; Conte et al., 2011; Bruhin et al., 2010). By allowing the plausible competing theories to coexist in the sample, mixture models provide a reconciliation of the debate over dominant theories of choice under risk (Harrison and Rutstrom, 2009). Mixture models take account of the fact that different individuals may have different preference functions. According to Kasahara and Shimotsu (2009), the attractive feature of the finite mixture model is that it provides flexible ways to account for unobserved heterogeneity. Thus, if we assume that individuals' behavior can be explained by a finite (C) number of data generating processes, mixture models assume that each individual behavior can be regarded as a draw from one of these C latent decision-making processes. Hence, it enables us to estimate the magnitude of these latent decision-making processes, as well as their parameters of interest.

For modeling choices under risk, most studies consider outright choices between lotteries and assume that data is generated by the expected utility and one or more additional competing decision-making processes. For example, Harrison and Rutstrom (2009) assumed that data is generated by expected utility and prospect theory. They found almost equal support for each theory and examined the role of demographics on theories' performances. Harrison and Rutstrom (2009) consider the finite mixture model as a statistical device for estimating the parameters of the model at the same time as one estimates the probability that each model applies to the sample. This approach indicates that one should not assume any model as the correct model, but as a model for a distinct type of individuals.

With similar outright choice data, Conte et al. (2011) considered expected utility and rank dependent utility with different specifications and concluded that a representative agent model gives a distorted view. Their findings suggest that 20% of the population are expected utility maximizers and 80% are rank dependent utility maximizers. The work of Bruhin et al. (2010) is interesting in a sense that it uses certainty equivalents of lotteries rather than outright choices between them and investigates the number of mixture components as well. Furthermore, their approach allows for more flexibility as the expected utility types are not imposed *a priori* in the estimation procedure. When assuming two components, their results suggest the same mixing proportion between the expected value and cumulative prospect theory from three different data sets. Nearly 20% of the subjects are classified as

expected value maximizers and 80% are classified as cumulative prospect theory maximizers.

3 The Experiment and the Data

The data used in this study is previously collected by Bruhin et al. (2010). The experiment was conducted in Zurich in 2006 when 1 Swiss franc equaled about 0.84 U.S. dollars. 118 subjects were drawn from the subject pool of the Institute for Empirical Research in Economics, which consists of students of all fields of the University of Zurich and the Swiss Federal Institute of Technology Zurich. The certainty equivalent of 40 lotteries were elicited for each subject.¹ Half of the lotteries were framed as gains and the other half were framed as losses.

Table 1
Gain lotteries ($x_1, p; x_2, 1 - p$)

p	x_1	x_2	$ x_1 - x_2 $	p	x_1	x_2	$ x_1 - x_2 $
0.9	10	20	10	0.50	20	50	30
0.5	0	10	10	0.25	10	40	30
0.5	10	20	10	0.25	20	50	30
0.1	10	20	10	0.05	10	40	30
0.5	0	20	20	0.05	20	50	30
0.95	10	40	30	0.95	0	40	40
0.95	20	50	30	0.75	0	40	40
0.75	10	40	30	0.05	0	50	50
0.75	20	50	30	0.95	50	150	100
0.5	10	40	30	0.90	0	150	150

For each loss situation, subjects were endowed with a guaranteed amount of money to compensate for any potential losses. The probabilities and outcomes defining the 40 lotteries in the gain domain are listed in Table 1. Accordingly, Table 2 shows loss lotteries. To prevent any order effect, the ordering of the lotteries were randomized. The random lottery incentive system was applied where at the end of the final session, one of the subject's choices was selected at random and played for real. Besides, each subject received 10 Swiss francs as a show-up fee.

¹ The elicitation procedure can be found in Bruhin et al. (2010, p.1379).

Table 2
Loss lotteries ($x_1, p; x_2, 1 - p$)

p	x_1	x_2	$ x_1 - x_2 $	p	x_1	x_2	$ x_1 - x_2 $
0.1	-10	-20	10	0.50	-20	-50	30
0.5	0	-10	10	0.75	-10	-40	30
0.5	-10	-20	10	0.75	-20	-50	30
0.9	-10	-20	10	0.95	-10	-40	30
0.5	0	-20	20	0.95	-20	-50	30
0.05	-10	-40	30	0.05	0	-40	40
0.05	-20	-50	30	0.25	0	-40	40
0.25	-10	-40	30	0.95	0	-50	50
0.25	-20	-50	30	0.05	-50	-150	100
0.5	-10	-40	30	0.10	0	-150	150

4 Econometric model

Using lotteries' certainty equivalents, I estimate the parameters of cumulative prospect theory (CPT) for each treatment (Tversky and Kahneman, 1992). I assume CPT as the underlying theory of decisions under risk as it captures framing effect and non-linear probability weighting. Moreover, CPT nests expected utility as a special case. Since there are no mixed lotteries in the data, I assume a zero reference point and use the term utility function instead of the value function. Furthermore, as a representative-agent model might not take account of the heterogeneity among individuals, a two-component mixture model is used. The maximum likelihood procedure used in the estimation of CPT is a special case of the mixture model when the number of components is one and hence, not reported. Conte et al. (2011) briefly discussed the problems with fitting data subject by subject and aggregate data fitting. The individual by individual analysis has the problem of heterogeneity within the subjects, as demonstrated by Hey (2001). Hence, many studies focused to incorporate different sources of heteroskedastic errors. While aggregate data analysis saves on degrees of freedom, it cannot take account of heterogeneity among subjects, as individuals are different in terms of their preference functions and their parameters.

Recently, the possibility of having more than one decision-making process for risky choices have been examined. Mixture models take account of the fact that different individuals may have different preference functions. According to Cameron and Trivedi (2005), the attractive feature of the finite mixture model is that it leads to a flexible parametric distribution. Moreover, it is a natural and simple way to treat population heterogeneity. Thus, if we assume that individuals' behaviour can be explained by a finite (C) number of data generating processes, mixture models assume that each individual behaviour

can be regarded as a draw from one of these C latent decision-making processes. Hence, it enables us to estimate the magnitude of these latent decision-making processes, as well as their parameters of interest.

Bruhin et al. (2010) used mixture models to estimate the parameters and the proportion of decision-making processes from certainty equivalents of simple lotteries. CPT as a deterministic choice model dictates the amounts of certainty equivalents. Hence, in order to estimate the parameters of the model, we have to add an error term, ε_{ig} . While there might be different sources of error, I follow Bruhin et al. (2010) and allow for three different sources of heteroskedasticity in the error variance. First, the error variance is proportional to the lottery's range. This is because, in the elicitation procedure, subjects had to choose between a lottery and 20 certain amounts, which are equally spaced throughout the lottery's range. Secondly, as subjects are indeed different, I allow the error variance to differ by individuals. Lastly, because of an asymmetry between gains and losses, I allow for domain-specific errors.

According to CPT, the value of any binary lottery $G_g = (x_g; p_g; y_g; 1 - y_g)$, with $y_g > y_g > 0$ is given by:

$$u(g) = w(1 - p_g)u(y_g) + (1 - w(1 - p_g))u(x_g) \quad (1)$$

Hence, the lottery's certainty equivalent $\widehat{ce}_g$ can be written as:

$$\widehat{ce}_g = u^{-1} \left(w(1 - p_g)u(y_g) + (1 - w(1 - p_g))u(x_g) \right) \quad (2)$$

We have to assume specific functional forms for the utility function $u(x)$ and the probability weighting function $w(p)$. As the point of this study is to examine the space effect and not finding the best fit, I focus only on one specification for the utility and probability weighting functions. A natural candidate for $u(x)$ is a power function:

$$u(x) = \begin{cases} x^\alpha & \text{if } x \geq 0 \\ -(-x)^\beta & \text{if } x < 0 \end{cases}$$

With $\alpha > 0$ and $\beta > 0$. This is proven to be the best fit for experimental data (Stott 2006). The absence of the loss aversion parameter in this specification is because of a lack of mixed lotteries when the loss aversion parameter is not identifiable. For the probability weighting function, I use the two-parameter specification by Goldstein and Einhorn (1987) abbreviated by GE. The GE specification is given by:

$$w(p) = \frac{\delta p^\gamma}{\delta p^\gamma + (1-p)^\gamma}, \delta > 0, \gamma > 0$$

Following Bruhin et al. (2010), I favour this specification because it can be conveniently interpreted. While γ captures the curvature, δ governs elevation and determines the intersection point. Note that the information about subjects is captured by their certainty equivalents. By incorporating three sources of heteroskedasticity, the standard deviation of the error term can be written as $\sigma_{ig} = \xi_i (x_g - y_g)$ where ξ_i denotes an individual domain-specific parameter and y_g , and x_g refer to the lottery's range. Thus, the observed certainty equivalent ce_{ig} can then be written as $ce_{ig} = \widehat{ce}_{ig} + \varepsilon_{ig}$, where the subscript i refers to individuals and g refers to lotteries.¹

The two-component mixture model assumes that individuals' choices are characterized by one of the two types of behaviour, each characterized by a distinct vector of parameters $\theta_c = (\alpha_c, \beta_c, \hat{\gamma}_c, \hat{\delta}_c)^2$. Note that the distinction between the types is by restrictions on parameters and there might be for example no expected utility types in the sample if the parameters of the probability weighting function for both types are significantly different from one, as for $\gamma = 1$ and $\delta = 1$ the GE function becomes linear. Since, $ce_{ig} \sim N(\widehat{ce}_{ig}, \sigma_{ig})$ the density of type C for the i th individual can be expressed as:

$$f(ce_{ig}; \theta_c, \xi_i) = \prod_{g=1}^G \frac{1}{\sigma_{ig}\sqrt{2\pi}} e^{-\frac{1}{2} \left(\frac{ce_{ig} - \widehat{ce}_{ig}}{\sigma_{ig}} \right)^2}$$

Pooling over all individuals and summing over all C components give us the model's likelihood L.

$$f(ce, \psi) = \sum_{c=1}^C \pi_c \prod_{i=1}^N f(ce_{ig}; \theta_c, \xi_i)$$

The log-likelihood of the finite mixture model is then given by:

$$\ln L(\psi|ce) = \sum_{i=1}^N \ln \sum_{c=1}^C \pi_c f(ce_{ig}; \theta_c, \xi_i)$$

where the vector $\psi = (\theta_1, \dots, \theta_c, \pi_1, \dots, \pi_{c-1}, \xi_1, \dots, \xi_N)$ summarises all the parameters of interest. The log-likelihood is maximized by the iterative expectation maximization (EM) algorithm (Dempster et al. 1977), which also

¹ $\xi_i = \xi$ is rejected as likelihood ratio test has a zero p-value

² $\hat{\gamma}_c$ and $\hat{\delta}_c$ are vectors containing the domain-specific parameters of the probability weighting functions.

provides individuals' posterior probability of belonging to group *C* by Bayesian updating.

5 Result

5.1 Raw Data Analysis

After clustering the data into high spaced treatment and low spaced treatment, we are left with 3502 observations for the low spaced treatment and 1167 observations for the high spaced treatment. Using relative risk premia = $\frac{EV-CE}{EV}$, where *EV* denotes expected value and *CE* stands for certainty equivalent, subjects' behaviour is classified to risk-averse ($RRP > 0$), risk-loving ($RRP < 0$) and risk-neutral ($RRP = 0$). In the high spaced treatment, on average 38% of choices exhibit risk-averse behaviour. However, 58% of choices exhibit risk aversion in the low spaced treatment. Hence, on average subjects made more risky choices in the high spaced treatment as compared to low spaced ones. Moreover, there was not a single risk-neutral observation.

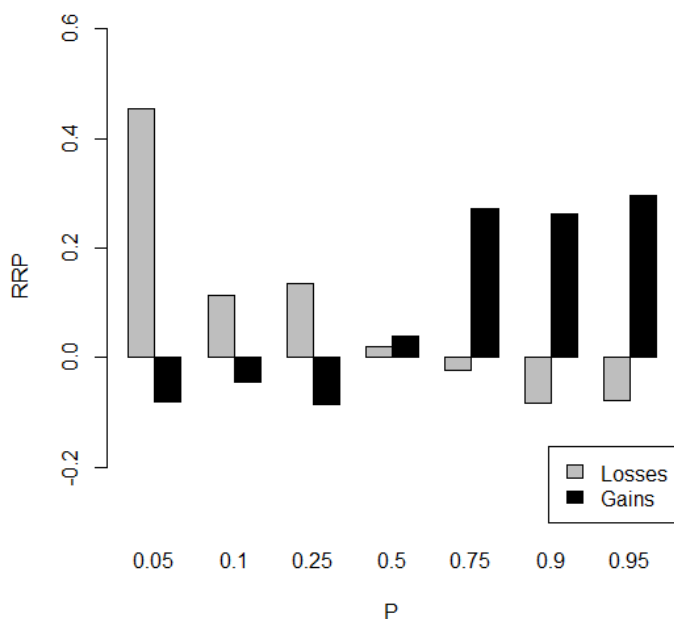


Figure 1. Median RRP.

Source: Research Findings

The data is consistent with the fourfold pattern of risk attitudes. It predicts risk-seeking behaviour over low probability gains and high probability losses. Furthermore, the pattern predicts risk aversion for low probability losses and high probability gains. Figure 1 shows the median RRP sorted by p , the probability of the more extreme outcome. Since there is a systematic relationship between RRP and p , average behaviour is better described by a model such as CPT.

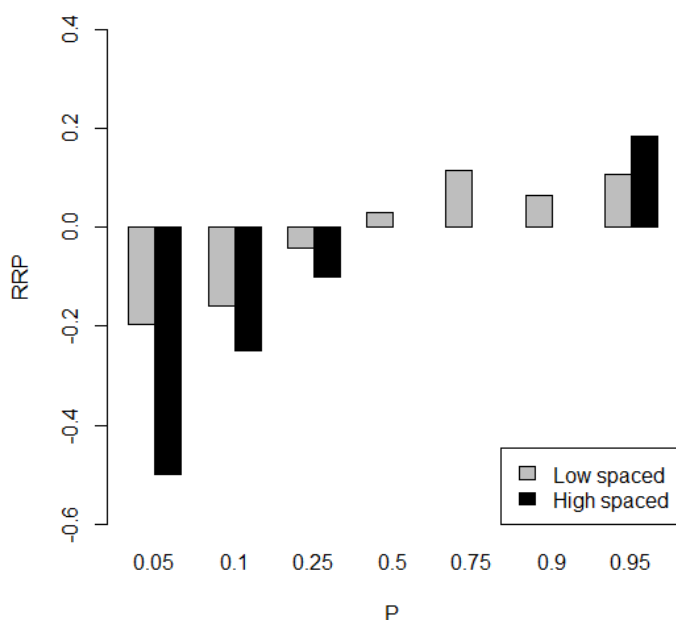


Figure 2. Gain Domain.

Source: Research Findings

In the gain domain, Figure 2 depicts the median RRP for both treatments under consideration. The black bars correspond to high spaced lotteries, and the grey bars represent low spaced lotteries. As it appears, the magnitude of RRP increases as the space between lotteries increases. It can be interpreted as subjects' more dramatic reaction when dealing with high spaced lotteries. Furthermore, choices in the high spaced treatment reveal a higher degree of risk aversion for the only high probability under consideration and a substantially higher degree of risk tolerance for low probabilities. It is contrary

to the stake effect found by Fehr-Duda et al. (2010), in which high stakes induce a higher degree of risk aversion overall probabilities.

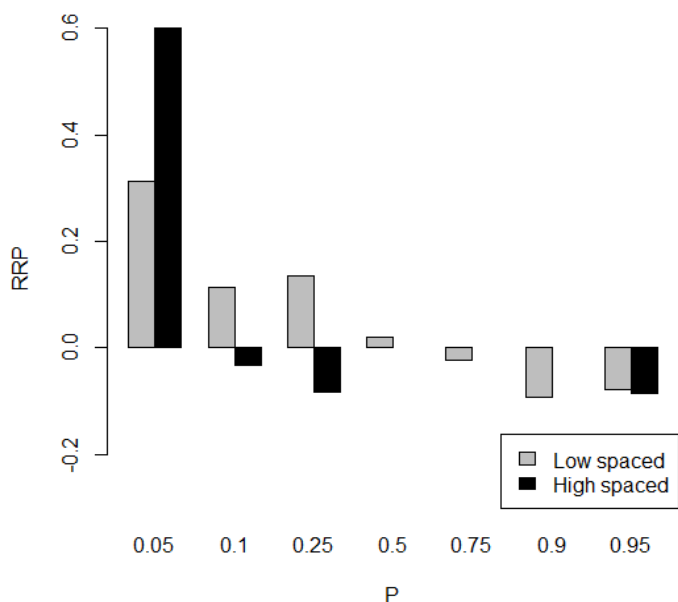


Figure 3. Loss Domain.

Source: Research Findings

The inspection of Figure 2 implies a significant difference between RRP from high and low spaced lotteries. Wilcoxon rank-sum tests confirm the space effect as the absolute values of RRP from the low spaced treatment except for $p = 0.25$ are significantly smaller than the high spaced ones.

Figure 3 shows the median RRP for both treatments in the loss domain. Choices in the high spaced treatment are not consistent with the fourfold pattern of risk attitudes as for two low probabilities they exhibit risk-seeking whereas they should indicate risk aversion. While there is a substantially higher risk aversion for the lowest probability under consideration (0.05) in the high spaced treatment, we observe risk-seeking behaviour for the other two low probabilities. Moreover, there is not a significant difference between risk-seeking behaviour for high probabilities of loss between the two treatments. High spaced RRP are insignificantly different from low-spaced ones at the probability level of 0.05 and 0.25. Therefore, analogous to the size effect, there is no coherent change in relative risk aversion in the loss domain.

In the next section, I examine whether the observed changes in relative risk aversion can be attributed to a specific component of lotteries evaluation. The stake effect implies lower degree of risk-seeking for low probabilities of gain when dealing with high stakes.

5.2 Single Model

When assuming CPT preferences, Table 3 reports the parameter estimates. The first interesting result is that the utility functions in both domains become almost linear over the low spaced lotteries, which is consistent with the previous findings that low outcomes induce linear utility functions. However, in the low spaced treatment we have outcomes as high as 50. Therefore, in order to guarantee the curvature of the utility function, we are required to take account of both the size of the outcomes and the space between them.

Table 3
CPT parameters for low and high spaced lotteries

Parameters	$\xi_i = \xi$					
	High	Low	Pooled	High	Low	Pooled
Gains						
α	0.9556 (0.0027)	1.007 (0.0012)	0.9164 (0.0005)	0.8445 (0.0023)	1.0567 (0.0044)	0.8259 (0.0004)
γ	0.6381 (0.0003)	0.4947 (0.0001)	0.5189 (0.0001)	0.4859 (0.0004)	0.4605 (0.0001)	0.4844 (0.0001)
δ	0.8017 (0.0053)	0.8517 (0.0005)	0.8859 (0.0004)	0.8085 (0.0061)	0.765 (0.0011)	0.8489 (0.0007)
Losses						
β	0.8972 (0.0134)	1.0079 (0.0015)	1.0934 (0.0011)	2.1083 (0.2212)	0.9591 (0.0048)	1.3809 (0.0041)
γ	0.5825 (0.0012)	0.5674 (0.0003)	0.5785 (0.0002)	0.6916 (0.0071)	0.4808 (0.0001)	0.548 (0.0001)
δ	1.4963 (0.0749)	1.0602 (0.001)	1.0387 (0.0009)	0.4388 (0.0373)	1.1800 (0.0028)	0.9487 (0.0027)
Ln L	1968	5129	10670	1442	4186	8264
Parameters	242	242	242	7	7	7
Individuals	118	118	118	118	118	118
Observations	3502	1167	4669	3502	1167	4669

Source: Research Findings

As risk attitudes under CPT are decomposed to attitudes towards probabilities and attitudes towards outcomes (payoffs), the observed difference in RRP in the data is not explained by a specific component of lotteries evaluation. The reason might be that when we jointly estimate the

utility and probability weighting functions, we allow for some flexibility between the two functions. Recall the fourfold pattern of risk attitudes and the difference in RRP's for each probability. As compared to low spaced gain lotteries, high spaced gain lotteries increase risk aversion over high probabilities and increase risk-seeking over low probabilities. Thus, if the observed difference were driven by a change in the probability weighting function, a more pronounced probability weighting function was expected in the high spaced treatment. However, δ which controls the elevation is higher in the low spaced treatment and γ , which controls the curvature of the probability weighting function is lower in the low spaced treatment. Furthermore, the curvature of the utility functions cannot explain the average behaviour either.

Prelec (2000) demonstrated the interesting feature of using power function in CPT model, the unique measure of risk attitudes in terms of structural parameters. Here, I extend Prelec example to two outcome simple lotteries. In a choice between a risky lottery $L = (x; p; ax; 1 - p)$, with $x > ax > 0$ and the sure receipt of its expectation $b = (px + ax - pax)$, the lottery will be preferred if:

$$\begin{aligned} w(p)x^\alpha + (1 - w(p))a^\alpha x^\alpha &> (px + ax - pax)^\alpha \\ w(p)(x^\alpha - a^\alpha x^\alpha) &> b^\alpha - a^\alpha x^\alpha \\ w(p) &> \frac{b^\alpha - a^\alpha x^\alpha}{x^\alpha - a^\alpha x^\alpha} \end{aligned}$$

If we denote $\frac{b^\alpha - a^\alpha x^\alpha}{x^\alpha - a^\alpha x^\alpha} = d$, then the condition is:

$$w(p) > d$$

As it appears from the unique measure of risk attitudes; for a given two outcome lottery, risk-seeking behaviour depends on both outcomes, their probability of occurrence, the space between them and the subjective weight of outcomes and probabilities.¹ Therefore, for $w(p) - d > 0$ we observe risk-seeking behaviour and for $w(p) - d < 0$ we observe risk aversion. Like RRP we use the magnitude of $|w(p) - d|$ to differentiate between the two treatments. In both treatments, on average, subjects exhibit risk-seeking behaviour over $p = (0.05, 0.10, 0.25)$. For these probabilities, $|w(p) -$

¹ In case of $a = 0$, we get risk seeking behaviour when probability weighting function overrides the curvature of the utility function ($w(p) > p^\alpha$).

d is higher in the low spaced treatment as compared to the high spaced treatment. Hence, the unique measure of risk attitudes cannot explain the observed difference in $RRPs$.¹

In order to rectify the pattern of risk attitudes observed in the row data, I imposed typical utility functions and then estimated the probability weights. With the imposed parameters of utility functions ($\alpha = \beta = 0.80$), the result is the following: It appears from Table 4 that the elevation parameter of the probability weighting function changes dramatically. In the gain domain, δ becomes higher in the high spaced treatment, which can explain the difference in $RRPs$ for each probability between the two treatments. With the imposed utility functions Figure 4 depicts the probability weighting functions in each domain that clearly represent the observed risk attitudes sorted by probabilities.

Table 4
Estimation with the imposed utility functions

Parameters	High	Low	Pooled
Gains			
γ	0.6190 (0.0003)	0.4907 (0.0001)	0.5057 (0.0001)
δ	1.0281 (0.0015)	0.9874 (0.0002)	0.9810 (0.0002)
Losses			
γ	0.5619 (0.0006)	0.5606 (0.0002)	0.5787 (0.0002)
δ	1.7517 (0.0038)	1.2332 (0.0004)	1.3174 (0.0004)
Ln L	1961	5085	10600
Parameters	240	240	240
Individuals	118	118	118
Observations	3502	1167	4669

Source: Research Findings

¹ For the only common high probability, 0.95, $w(p) < d$ in the both treatments with $|w(p) - d|$ being higher in the low spaced treatment.

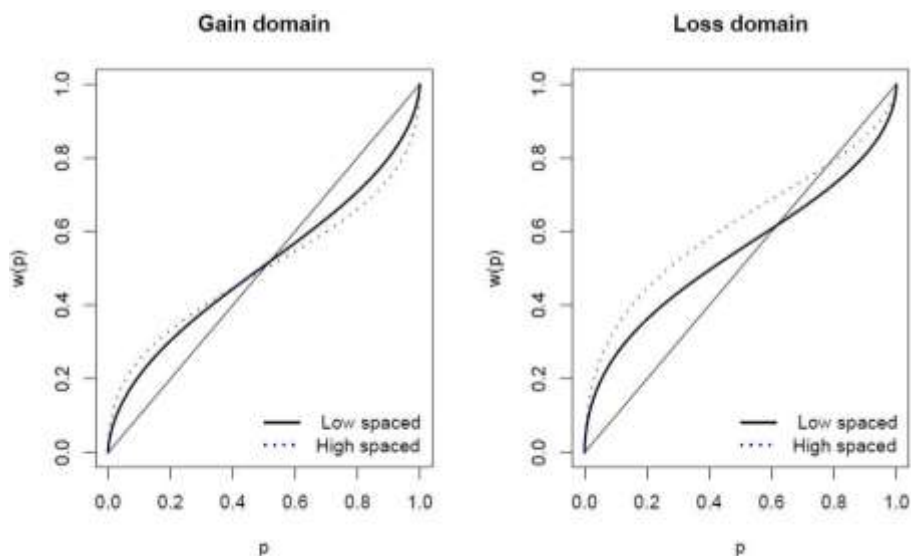


Figure 4. Probability Weighting Functions by the Imposed Utility Functions.

Source: Research Findings

5.3 Mixture Model

As a representative-agent model cannot explain the observed difference between the high and low spaced treatments, a two-component mixture model is used for each treatment. Table 5 reports the mixing proportions and the behavioural parameters for each type which are called CPT type I and CPT type II. The interesting feature of the mixture model is the segregation of behavioural types. Since the estimation procedure provides us with each individual posterior probability of group membership, we can test the classification power of the mixture model. As such, Figure 5 shows the distribution of the individuals' posterior probabilities of group membership for each treatment. In this Figure c_1 represents the posterior probability of belonging to CPT type I. As the Figure shows, the c_1 's are either close to 1 or close to 0 for most individuals, indicating a clear classification of types. Interestingly, the mixing proportion is almost the same under the two treatments, which also indicates a robust classification of types.

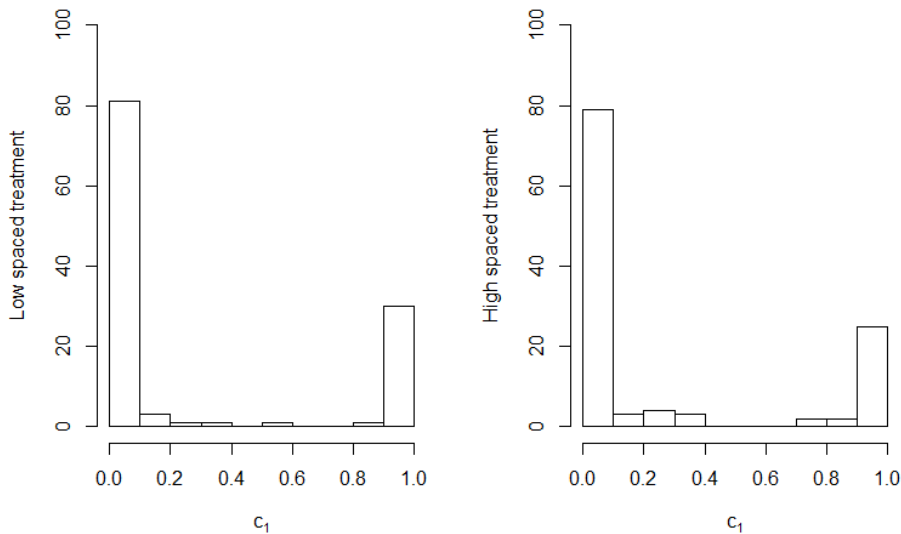


Figure 5. Distribution of Individuals' Posterior Probability of CPT Type I.

Source: Research Findings

While in the low spaced treatment, the utility and probability weighting parameters are close to 1 at least for CPT type I, in the high spaced treatment they are not. However, because the standard errors are very small, the 95%-confidence interval of every single estimate does not contain unity. Thus, it seems there is no expected utility type in the data set. In the low spaced treatment, the estimated parameters α , β and δ display high degree of conformity in both domains, whereas in the high spaced treatment only in the gain domain the estimated parameters α and δ are close. Note that the major difference between the types is the difference in γ that controls the curvature of probability weighting function. This in turn, confirms the conjecture of Tversky and Kahneman (1992) on the curvature dependency of probability weighting functions on the space between the outcomes.

I use the parameter estimates of the utility functions and the probability weighting functions to characterize risk-taking behaviour. Because of the dual structure of risk preferences under CPT, I first analyze the risk attitudes in terms of probabilities. Since there is only a coherent space effect in the gain domain, I focus only on the gain domain. Figure 6 illustrates the probability weighting functions. For CPT type I individuals, both treatments induce probability underweighting over the whole probability interval. Moreover, as it appears from Figure 6, the high spaced treatment induces more pronounced

probability weights. It is analogous to the size effect appeared in Fehr-Duda et al. (2010) in a sense that higher stakes induce more degree of high probabilities underweighting and less degree of low probabilities overweighting. Hence, in the minority group, more underweighting of high probabilities can capture a higher degree of risk aversion in the high spaced treatment.

With regard to the majority (CPT type II) types who have more pronounced probability weights, the space effect induces less probability weighting in general. Low probabilities are over-weighted more in the low spaced treatment and high probabilities are underweighted more under the low spaced treatment, which does not capture the space effect emerged from the median *RRPs*.

As the probability weighting functions cannot explain higher degree of risk-seeking over low probabilities of gain appeared in the high spaced treatment, I turn to the curvature of utility functions. While the characteristic of utility functions can be used to explain changes in relative risk aversion, it will only give a measure over the whole distribution of outcomes irrespective of their possibility of occurrence. In the minority group, changes in risk tolerance can be attributed to the way subjects weight outcomes. The coefficients of relative risk aversion are negative which indicates risk-seeking behaviour. As α is higher in the high spaced treatment, it can take account of a higher degree of risk-seeking behaviour as compared to the low spaced treatment. In the majority group, however, the utility function from the low spaced treatment exhibits risk-seeking behaviour while the high spaced treatment shows risk aversion.

Table 5
Classification of behaviour

Parameters	CPT type I			CPT type II		
	High	Low	Pooled	High	low	Pooled
π	0.2474 (0.0019)	0.2394 (0.0012)	0.2237 (0.0015)	0.7526	0.7606	0.7763
Gains						
α	1.1159 (0.0077)	1.0101 (0.0004)	0.9884 (0.0003)	0.9092 (0.0036)	1.0301 (0.0016)	0.9015 (0.0006)
γ	0.9387 (0.0014)	0.9482 (0.0001)	0.9448 (0.0001)	0.5061 (0.0004)	0.3948 (0.0002)	0.4247 (0.0001)
δ	0.6814 (0.0103)	0.8988 (0.0003)	0.9087 (0.0003)	0.7692 (0.0066)	0.8077 (0.0005)	0.8620 (0.0005)
Losses						
β	1.8317 (0.0121)	0.9896 (0.0006)	1.0136 (0.0006)	0.8997 (0.0171)	1.0136 (0.0024)	1.1214 (0.0016)
γ	1.3046 (0.0021)	0.9459 (0.0005)	0.9527 (0.0001)	0.4732 (0.0007)	0.4197 (0.0008)	0.4515 (0.0001)
δ	0.3974 (0.0027)	1.0704 (0.0008)	1.0490 (0.0006)	1.5579 (0.0950)	1.0942 (0.0009)	1.0591 (0.0013)
Ln L	2119	5625	11336			
Parameters	249	249	249			
Individuals	118	118	118			
Observations	3502	1167	4669			

Source: Research Findings

As demonstrated in section 4.2, in the gain domain risk-seeking behaviour holds if $w(p) > d$. Thus, while in the minority group all probabilities are underweighted which mistakenly might be considered as risk aversion, the unique measure shows risk-seeking behaviour over $p = (0.05, 0.1)$ under both treatments. As compared to the low spaced treatment, $|w(p) - d|$ derived from the high spaced treatment is higher for both probabilities. Hence, in accordance with the observed difference in RRP, the minority types are more risk-seeking when confronted with high spaced lotteries over low probabilities of gains. Moreover, the minority types exhibit a higher degree of risk aversion for $p = 0.95$ under high spaced treatment as well.

The majority types exhibit risk-seeking behaviour over $p = (0.05, 0.1, 0.25)$ under both treatments. As compared to the low spaced treatment, the $|w(p) - d|$ derived from the high spaced treatment is lower for all probabilities. Hence, subjects exhibit a higher degree of risk-seeking in the low spaced treatment.

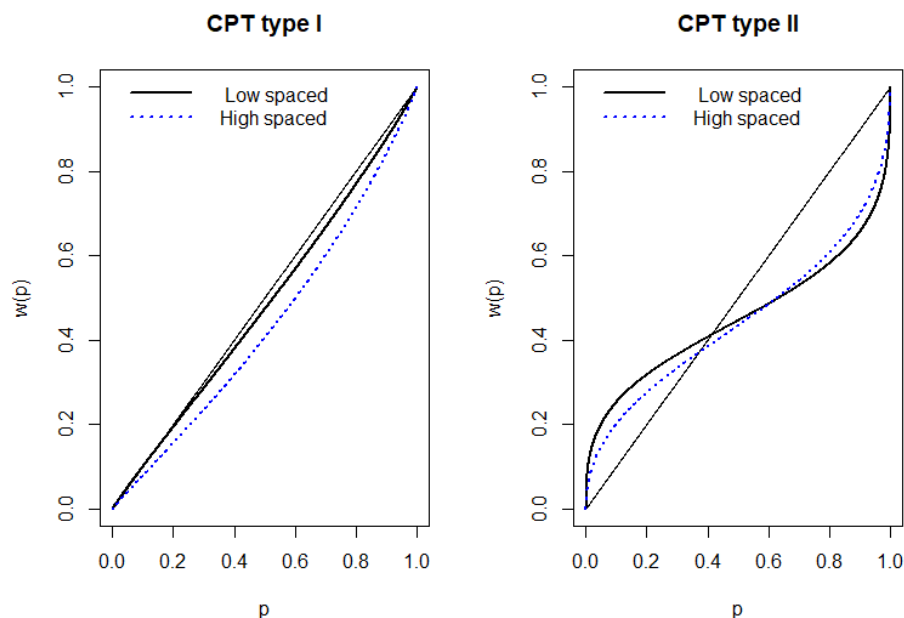


Figure 6. Probability Weighting Functions by Types in the Gain Domain.

Source: Research Findings

6 Discussion and conclusion

This paper has three objectives. Firstly, it studies the effect of an increase in the space between lotteries' outcomes on risk attitudes. Interestingly, we do find the space effect from data that was not collected to test it. The median of *RRPs* sorted by the probability of the highest absolute outcome demonstrates a significant difference between the high and low spaced treatments. While there is no coherent picture in the loss domain, in the gain domain, the space effect implies greater reaction of subjects in terms of *RRPs* magnitudes. The space effect induces higher degree of risk-seeking for low probabilities of gain and higher degree of risk aversion for high probabilities of gain.

Concerning the second objective, when assuming CPT preferences, the observed difference in *RRPs* cannot be attributed to any of the components of lotteries evaluation. In other words, changes in the parameter estimates are not meaningful. This misrepresentation of risk preferences might be due to interaction between probability weights and outcomes evaluation or the interaction between the stake and the space effect or all of them together. As many financial decisions involve subsequent follow-up decisions, carefully

designed experiments are called for, disentangling the size and space effect on risk attitudes and structural estimation of behavioural parameters.

With regard to the heterogeneity among individuals, which is the third objective, this study demonstrates that heterogeneity can be pivotal in explaining the observed difference in *RRPs*. The two-component mixture model shows that a small percentage of the sample can be responsible for the observed difference in the row data, which is not recovered by the average behaviour determined by a representative-agent model.

Furthermore, one might look at the reference point for an explanation of increasing *RRP* as subjects might frame a positive gain based on the characteristic of the alternatives; the size and the space between them. However, as stated by Bruhin et al. (2010), simultaneous estimation of model parameters with the reference point is questionable when there are no mixed lotteries from the onset.

References

- Adams, C. P. (2016). Finite mixture models with one exclusion restriction. *The Econometrics Journal*, 19(2), 150-165.
- Bardsley, N., & Moffatt, P. G. (2007). The Experimentics of Public Goods: Inferring Motivations from Contributions. *Theory and Decision*, 62(2), 161-193.
- Bruhin, A., Fehr-Duda, H., Epper, Th. (2010). Risk and Rationality: Uncovering Heterogeneity in Probability Distortion. *Econometrica*, 78(4), 1375-1412.
- Cameron, A. C. & Trivedi, P. K. (2005). *Microeconometrics: Methods and Applications*. Cambridge University Press.
- Conte, A., Hey, J. D., Moffatt, P. G. (2011). Mixture Models of Choice under Risk. *Journal of Econometrics*, 162(1), 79-88.
- Dempster, A., Laird, N., Rubin, D. (1977). Maximum Likelihood from Incomplete Data via the EM Algorithm. *Journal of the Royal Statistical Society, Series B*, 39(1), 1-38.
- Harrison, G. W. & Rutstrom, E. E. (2009). Representative Agents in Lottery Choice Experiments: One Wedding and a Decent Funeral. *Experimental Economics*, 12(2), 133-158.
- Eckstein, Z., & Wolpin, K. I. (1990). Estimating a Market Equilibrium Search Model from Panel Data on Individuals. *Econometrica: Journal of the Econometric Society*, 58(4), 783-808.
- Einav, L., Finkelstein, A., Pascu, I., Cullen, M.R. (2012). How general are Risk Preferences? Choices under Uncertainty in Different Domains. *American Economic Review*, 102(6), 2606-2638.
- Etchart-Vincent, N. (2004). Is Probability Weighting Sensitive to the Magnitude of Consequences? An Experimental Investigation on Losses. *Journal of Risk and Uncertainty*, 28(3), 217-235.
- Etchart-Vincent, N. (2009). Probability Weighting and the Level and Spacing of Outcomes: An Experimental Study over Losses. *Journal of Risk and Uncertainty*, 39(1), 45-63.

- Fehr-Duda, H., Bruhin, A., Epper, Th., Schubert, R. (2010). Rationality on the Rise: Why Relative Risk Aversion Increases with Stake Size. *Journal of Risk and Uncertainty*, 40(2), 147-180.
- Feltovich, N., Iwasaki, A., Oda, S.H. (2012). Payoff Levels, Loss Avoidance, and Equilibrium Selection in Games with Multiple Equilibria: An Experimental Study. *Economic Inquiry*, 50(4), 932-952.
- Goldstein, W. M. & Einhorn, H. J. (1987). Expression Theory and the Preference Reversal Phenomenon. *Psychological Review*, 94(2), 236-254.
- Harrison, G. W. & Rutstrom, E. E. (2009). Representative Agents in Lottery Choice Experiments: One Wedding and a Decent Funeral. *Experimental Economics*, 12(2), 133-158.
- Henry, M., Kitamura, Y., & Salanié, B. (2014). Partial Identification of Finite Mixtures in Econometric Models. *Quantitative Economics*, 5(1), 123-144.
- Hertwig, R., Barron, G., Weber, E. U., Erev, I. (2004). Decisions from Experience and the Effect of Rare Events in Risky Choice. *Psychological Science*, 15(8), 534-539.
- Hey, J.D. (2001). Does Repetition Improve Consistency? *Experimental Economics*, 4(1), 5-54.
- Hilbig, B. E., & Glockner, A. (2011). Yes, they can! Appropriate Weighting of Small Probabilities as a Function of Information Acquisition. *Acta Psychologica*, 138(3), 390-396.
- Kasahara, H. and K. Shimotsu 2009, Nonparametric Identification and Estimation of Finite Mixture Models of Dynamic Discrete Choice. *Econometrica*, 77(1), 135--175.
- Keane, M. P., & Wolpin, K. I. (1997). The Career Decisions of Young Men. *Journal of Political Economy*, 105(3), 473-522.
- Loomes, G. (2005). Modeling the Stochastic Component of Behaviour in Experiments: Some Issues for the Interpretation of Data. *Experimental Economics*, 8(4), 301-323.
- McLachlan, G. J., & Peel, D. (2004). *Finite mixture models*. John Wiley & Sons.
- Prelec, D. (2000). Compound Invariant Weighting Functions in Prospect Theory. In: Kahneman, D., Tversky, A. *Choices, Values and Frames* (pp. 67-92). Cambridge University Press: New York.
- Stewart, N., Reimers, S., & Harris, A. J. L. (2015). On the origin of Utility, Weighting, and Discounting Functions: How they get their Shapes and How to Change their Shapes. *Management Science*, 61(3). 687-705.
- Stewart, N., Chater, N., & Brown, G. D. (2006). Decision by Sampling. *Cognitive Psychology*, 53(1), 1-26.
- Stott, H. P. (2006). Cumulative Prospect Theory Functional Menagerie. *Journal of Risk and Uncertainty*, 32(2), 101-130.
- Tversky, A., & Kahneman, D. (1992). Advances in Prospect Theory: Cumulative Representation of Uncertainty. *Journal of Risk and uncertainty*, 5(4), 297-323.
- Ungemach, C., Stewart, N., Reimers, S. (2011). How Incidental Values from our Environment Affect Decisions about Money, Risk, and Delay. *Psychological Science*, 22(2), 253-260.