

STOCKHOLM SCHOOL OF ECONOMICS
Department of Economics
5350 – Thesis in Economics
Fall 2023

Evaluating probability weighting among lottery bond investors: an observational study on Cumulative Prospect Theory

Anton Eriksson Nyberg (42413) & Olof de Blanche (24693)

Abstract: Cumulative Prospect Theory is a leading descriptive theory of decisions under uncertainty, backed up by a plethora of experimental evidence. This paper is one of the first to apply it to observational data. Risk attitudes in the bond market have not been studied in a CPT framework. We fill this research gap. By using maximum likelihood estimation, CPT parameters are estimated for the Swedish lottery bond market. We find drastic probability weighting that opposes linearity in probabilities espoused by Expected Utility Theory.

Keywords: Cumulative Prospect Theory, CPT, observational study, field study, lottery bonds

JEL: D81, C93, G40, C51

Supervisor: Karl Wärneryd
Date Submitted: December 4th
Date Examined: December 13th
Discussants: Qingrong Wu & Markus Sallkvist
Examiner: Kelly Ragan

Acknowledgments

We extend thanks to our supervisor, Karl Wärneryd, for his helpful input. Further, we are grateful to Kristian Rydqvist for providing extensive data on Swedish lottery bonds. Last, we want to give our warmest gratitude to Agnes Durbeej Hjalt and Adriana Eriksson for their moral support and suggestions.

Table of Contents

1) Introduction	4
1a) Overview.....	4
1b) Background.....	4
1c) Relation to previous work.....	6
2) Theory	8
2a) Overview.....	8
2b) Further discussion about PDWs	9
3) Method and Data.....	12
3a) Lottery bonds	12
3b) Construction of variables.....	14
3c) Estimation of population parameters	16
3d) Inferences.....	17
4) Results	19
5) Discussion.....	22
5a) Representative agent assumption.....	22
5b) Independence assumption.....	22
5c) Validity of analysis	22
5d) Generalization.....	23
5e) Further research	24
6) Conclusion.....	25
7) References	26

1) Introduction

1a) Overview

The treatment of decisions under uncertainty has a rich literature behind it. This comes as no surprise; most decisions in the real world are made under uncertainty. The more economics can explain and make predictions about such decisions, the more the field gains applicability. We aim to contribute with an observational study on the probability weights that agents assign to lottery outcomes. Weights are estimated in a Cumulative Prospect Theory (CPT) model. The lotteries we deal with are drawings of Swedish lottery bonds. We consider two conflicting theories' views on probability weights. In Expected Utility Theory (EUT), weights assigned to events are the probabilities. In CPT, weights are allowed to depend on probabilities nonlinearly. Our goal in this paper is to determine whether participants in the Swedish lottery bond market evaluate probabilities consistent with EUT or CPT. We find the latter to be true. The implications of probability weighting are broad. With a more comprehensive understanding of probability weighting dynamics, government-funded lotteries can be structured in a more socially optimal way. Insurance policies can be more efficiently constructed to increase profits.

1b) Background

In the framework of EUT, formalized by von Neumann and Morgenstern (1945), the weights assigned to the events of an uncertain prospect¹ are their respective probabilities². That is, the expected value function is linear in probabilities³. Under EUT, preferences (or risk aversion) are determined solely by the shape of the utility function, which can be, and generally is, described by one parameter, e.g., the exponent in power utility. The results from Allais's (1953) choice problem⁴ challenge EUT, revealing that agents deviate from linearity in probabilities when making decisions under uncertainty. Allais paints a scenario⁵ in which you can win a substantial prize with a probability of 0.99 and win 0 with a probability of 0.01. Subjects are willing to pay a lot to make the big prize certain (by increasing the probability of winning by 0.01). Allais presents another scenario in which the probability of winning the same substantial prize is only 0.10. Here, people are not willing to pay as much as in the previous setting to increase the probability by 0.01. Hence, the probability 0.01 carries a different weight depending on the starting point, and subjects in Allais's experiment exhibit nonlinear probability weighting of outcomes.

Tversky and Kahneman (1992) present a prominent critique and generalization of EUT titled "Cumulative Prospect Theory", reconciling theory with these findings. CPT's primary distinction from EUT is weighting events by probability decision weights (PDW or π) instead of probabilities. Each PDW considers the entire cumulative distribution of the prospect and assigns

¹ Basically the same as a gamble or a lottery.

² In this paper, we also use the term "objective probability". This refers to the *true (objective)* probability of an event occurring. With "probability", we refer to a subjective or objective probability, used when a distinction is not needed or should not be made.

³ A gamble in which you gain x_1 with probability p_1 and x_2 with probability $1 - p_1$ is combined linearly in its expected value as $EV(p_1) = p_1 u(x_1) + (1 - p_1) u(x_2)$.

⁴ In 1952, Allais presents a set of gambles to an audience at an economics conference in Paris. He clarifies the prize winnings and probabilities of the gambles and makes the participants pick the most preferable.

⁵ We have changed the scenario in some ways while keeping it in the spirit of the original.

a corresponding weight between 0 and 1. CPT adds one⁶ additional parameter, connected to the PDWs, over EUT to characterize preferences. Now, two parameters jointly determine risk aversion. Under CPT, the utility function measures changes in utility from a reference point. It is important to stress that Neumann and Morgenstern's axioms do not preclude such a utility function. In the present study, the utility function measures changes in utility, both in the context of CPT and EUT.

Tversky and Kahneman's (1992) paper serves not only as an axiomatic but also as a behavioral treatment. While PDWs and the utility function can technically take on any arbitrary parametric forms, Tversky and Kahneman argue that they should adhere to certain empirical findings, Allais's being among them. Specifically, they posit that each PDW is a nonlinear function that favors overweighting of low probabilities and shows increasing sensitivity around 0 and 1. With this in mind, CPT is consistent with Allais's findings. As with probabilities, the sum of the PDWs will always equal 1 for positive and negative prospects⁷, but how this mass of 1 is distributed is still up for question. Concerning the utility function, Tversky and Kahneman argue that it is characterized by diminishing marginal sensitivity and is steeper for losses than for gains. They back up their claims about PDWs and the utility function with an analysis of experimental data.

A range of studies have examined CPT in experiments and largely confirmed Tversky and Kahneman's findings (Wu and Gonzalez 1996; Abdellaoui 2000; Camerer and Ho 1994; Bruhin, Fehr-Duda, and Epper 2010). In the wake of CPT, EUT's status as the go-to framework for decisions under uncertainty is fading. However, most research on CPT has been based on laboratory experiments to gauge preferences. While experiments can be instructional, the true test for theory is its application to the real world. Our goal is to survey CPT in precisely this.

Is probability weighting in bond markets consistent with EUT; does our data support linearity in probabilities?

Rejection of our hypothesis, under the CPT paradigm, will lend credence to nonlinear probability weighting in the real world. Non-rejection will make the dismissal of EUT more questionable. A plethora of research has established a fourfold pattern in risk attitudes. For gains, there is risk seeking towards outcomes of low probability and risk aversion towards outcomes of high probability. Conversely, for losses, there is risk seeking towards outcomes of high probability and risk aversion towards outcomes of low probability. Such risk profiles can explain why an agent would be open to buying both lottery tickets and insurance, as noted by Friedman and Savage (1948). They argue that EUT can sufficiently account for the fourfold pattern through the utility function. Tversky and Kahneman (1992) address the pattern with both the diminishing sensitivity of the utility function *and* nonlinear probability weighting. While the topic of this paper revolves around the latter, we cannot disregard the former. We must estimate a complete CPT model, including the parameter determining marginal utility, to assess PDWs. Only then can we thoroughly scrutinize the model's properties, with one caveat. The lotteries we deal with only permit us to consider gains. The loss aversion aspect of CPT is a salient feature of the model that we miss out on.

To test our hypothesis, agents need to reveal their preference relationships between different prospects. We focus on binary choices constructed from a pool of lotteries. For any given

⁶ π can take any functional form incorporating arbitrarily many parameters. However, the convention is to describe it with one parameter.

⁷ For mixed prospects—involving both gains and losses—PDWs can sum to less or greater than 1.

parameter values, the respective Cumulative Prospect Value⁸ (CPV) of two lotteries implies a theoretical preference relationship, with the prospect commanding the highest CPV deemed preferable. Theoretical preference relationships are of no value without data on actual choices. In the lab, subjects are generally prompted to decide between two or more prospects, and thus, true preferences are revealed. We are not afforded such straightforwardness in using real-world data.

Nonetheless, from daily prices of Swedish lottery bonds, we can deduce preference relationships between lotteries. We proceed to estimate parameters α and γ , respectively connected to the utility function and the PDWs (see Eq. (3) and Eq. (4) in Chapter 2). In this, we follow the approach of Camerer and Ho (1994), using stochastic utility and maximum likelihood estimation. α and γ are chosen to maximize the likelihood that, under certain assumptions about stochastic terms, theoretical preference relationships match observed choices. To make inferences, we simulate 1 000 samples by repeated drawings with replacement from our full sample of pairwise lottery comparisons (bootstrapping). Each simulation gives new point estimates of parameters α and γ . Together, these make out the estimated sampling distributions from which we derive confidence intervals for both parameters to make inferences and test our hypothesis.

1c) Relation to previous work

CPT is an extension of Kahneman and Tversky's seminal paper "Prospect Theory" (PT) (Kahneman and Tversky 1979). In PT, each PDW is a function of an individual probability. While conceptionally simpler, such a construction can give rise to violations of first-order stochastic dominance⁹ (Bernheim and Sprenger 2020). With broad unanimity on the subject of dominance, PT is significantly constrained. As a remedy, CPT incorporates the weighting of cumulative—instead of individual—probabilities, introduced by Quiggin (1982). With this, CPT accommodates the same behavioral findings as PT and is more readily extended to prospects with more than two outcomes. As it stands, CPT has overtaken PT as the leading model of decisions under uncertainty.

On the topic of CPT, relatively few studies have focused on estimating model parameters, turning instead toward the qualitative properties of the theory. This may be a consequence of Kahneman and Tversky's heeding in their 1992 paper:

“The estimation of a complex choice model, such as cumulative prospect theory, is problematic. If the functions associated with the theory are not constrained, the number of estimated parameters for each subject is too large. To reduce this number, it is common to assume a parametric form (e.g., a power utility function), but this approach confounds the general test of the theory with that of the specific parametric form.”

However, it seems that Tversky and Kahneman are aware of the merits of model estimation because, in that very same paper, they present parametric forms for both the utility function and the PDWs. They go on to estimate the parameters and discuss their implications. We

⁸ The value attained from a prospect in which you gain x_1 with probability p_1 and x_2 with probability $1 - p_1$ is denominated by its Cumulative Prospect Value as $CPV(p_1) = \pi_1(p_1)u(x_1) + \pi_2(1 - p_1)u(x_2)$.

⁹ Consider a choice between two lotteries that are identical in every way except that in one lottery all outcomes have one additional dollar over the outcomes of the other lottery. Preferring the lottery with a lower payout at each state is one example of a violation of stochastic dominance.

see a distinction between investigating if the general behavioral properties posed by Kahneman and Tversky reflect reality and investigating if a specific CPT model exhibits these properties when applied to real-world data. There is value in both alternatives; however, we are focused on the latter. Bernheim and Sprenger (2020) present a powerful critique of CPT. Still, we are not explicitly testing if CPT is true; instead, we assume it is and use the generality it provides to test it against a specific case—EUT.

While maximum likelihood estimation paired with a stochastic element has not been applied to market data before the present study, it is a prominent approach in estimating the parameters of CPT in laboratory and survey settings (Camerer and Ho 1994; Stott 2006; Harrison and Rutström 2009; Zeisberger, Vrecko, and Langer 2012; Jou and Chen 2013; Bocquého, Jacquet, and Reynaud 2014). The leading approach is to estimate individual CPT parameters for each subject and then pool the estimates to describe the population. We go against the grain, following Camerer and Ho (1994), by pooling all subjects before the estimation phase, implicitly assuming that all have identical preferences. In the aforementioned studies, “subjects” exclusively refer to individual human beings. Our subjects are market behaviors. In this, we use the convention of modeling a single representative agent exhibiting these market behaviors. We follow Zeisberger, Vrecko, and Langer (2012) and Rieger, Wang, and Hens (2017) in using bootstrapping for inferences.

Other papers modify the maximum likelihood methodology by constraining individuals to have parameters distributed according to a group-level distribution (Nilsson, Rieskamp, and Wagenmakers 2011; Murphy and ten Brincke 2018). With an appropriate a priori hierarchical structure, they should retrieve more reliable estimates of individual subjects’ parameters. In contrast to using maximum likelihood estimation, Rieger, Wang, and Hens (2017) follow the approach of Tversky and Kahneman (1992) to estimate their model, minimizing the squared error between elicited certainty equivalents and theoretical certainty equivalents¹⁰ to estimate parameters.

Gurevich, Kliger, and Levy (2009) present the first field study on CPT. They infer risk preferences from pricing data of stock options (for the top 30 companies on S&P100). Theoretical state prices are inferred from option prices, and probabilities are derived from historical stock prices. Estimates of α and γ are retrieved by multiple regression analysis. Unlike their approach, we do not rely on theoretical relationships to the same extent; objective probabilities are given exogenously by the Swedish Treasury. However, they study CPT for both losses and gains. They find nonlinear weighting of probabilities in line with CPT but less drastic than in Tversky and Kahneman’s (1992) paper. We are inspired by Gurevich, Kliger, and Levy’s work and want to extend it. Our study considers a different market in a different country during the same period—the 1990s. This paper is the first observational study on CPT involving lottery bonds and maximum likelihood estimation. We introduce a novel way to use a cash-equivalent function to map market prices to theoretical preferences.

The paper is organized as follows. Chapter 2 outlines theoretical concepts. Chapter 3 introduces lottery bonds and our dataset, maps theory to observed variables, goes deeper into the estimation method, and presents confidence intervals retrieved by bootstrapping. Chapter 4 reveals key findings. Chapter 5 discusses underlying assumptions, internal and external validity, and ideas for further research.

¹⁰ Theoretical certainty equivalents involve inverting CPV to get a corresponding cash value.

2) Theory

2a) Overview

In examining lottery preferences over market data, we derive aggregate functions and thus model the decisions of the representative lottery bond investor. We assume that CPT is the primary driver of investor preferences between different lotteries, but we also allow for a stochastic element. In fact, we need an additional component in the decision process to explain any non-unanimous preferences over identical lotteries—of which there are many.

We represent lottery drawings as a sequence $f = \{(p_1, x_1), (p_2, x_2), \dots, (p_n, x_n)\}$ of pairs of objective probabilities p_i and prize money outcomes x_i where $x_i > x_j$ for $i > j$. The static CPV of lottery a with n different outcomes is defined as $v(f_a) = \sum_{i=1}^n \pi_i u(x_i)$, where π_i is the PDW assigned to outcome i and $u(x_i)$ is the utility measure of x_i amount of prize money. In EUT, π_i is the probability of outcome i . CPT replaces the probabilities with the more general concept $\pi_i = w(\sum_{j=i}^n p_j) - w(\sum_{j=i+1}^n p_j)$ and $\pi_n = w(p_n)$, where $w: \sum_{j=i}^n p_j \in [0,1] \rightarrow [0,1]$ with $w(0) = 0$ and $w(1) = 1$. Thus, for each PDW, the entire probability distribution is transformed according to some weighting function $w(\cdot)$.

We introduce a probabilistic element by defining stochastic CPV as $V(f_a) = v(f_a) + r_a$, where r_a is the stochastic term for lottery a . Now, a lottery is evaluated not only by its static CPV but also by random trembles in preferences. The representative agent makes a binary choice between lotteries a and b by evaluating their respective stochastic CPV. Specifically, the agent follows the decision criteria “choose f_a over f_b if $y_{ab} = 1$ ”, and we define $y_{ab} = 1[y_{ab}^* > 0]$, where $y_{ab}^* = V(f_a) - V(f_b) = v(f_a) - v(f_b) + z_{ab}$ and $z_{ab} = r_a - r_b$.

Further, we assume the relationship:

$$s_a = C(V(f_a)), \quad \text{Eq. (1)}$$

where $C(\cdot)$ is a strictly increasing cash-equivalent function and s_a is the price of lottery a . Thus, we have the equivalence:

$$s_a - s_b > 0 \Leftrightarrow V(f_a) > V(f_b). \quad \text{Eq. (2)}$$

We proceed by deriving the estimation of our parameters of interest from preferences over lotteries a and b . We let z_{ab} have the logistic distribution with cumulative distribution $G(x) = \frac{e^x}{1+e^x}$ (as in Camerer and Ho's (1994) paper), and we let $\Delta_{ab} = v(f_a) - v(f_b)$. To construct the log likelihood function, expressions for $P(\text{choose } f_a)$ and $P(\text{choose } f_b)$ must be derived.

$$\begin{aligned}
P(\text{choose } f_a) &= P(y_{ab} = 1) \\
&= P(y_{ab}^* > 0) \\
&= P(V(f_a) > V(f_b)) \\
&= P(-z_{ab} < v(f_a) - v(f_b)) \\
&= P(-z_{ab} < \Delta_{ab}) \\
&= G(\Delta_{ab}).
\end{aligned}$$

Conversely,

$$\begin{aligned}
P(\text{choose } f_b) &= P(y_{ab} = 0) \\
&= P(y_{ab}^* < 0) \\
&= P(V(f_b) > V(f_a)) \\
&= P(z_{ab} < v(f_b) - v(f_a)) \\
&= P(z_{ab} < -\Delta_{ab}) \\
&= 1 - G(\Delta_{ab}).
\end{aligned}$$

We let $\ell_{ab} = y_{ab} \text{Log}(G(\Delta_{ab})) + (1 - y_{ab}) \text{Log}((1 - G(\Delta_{ab})))$ be the log likelihood of the observed choice between lotteries a and b and $\mathbb{F} = \{(f_a, f_b), (f_c, f_d), \dots\}$ be the set of all pairs of lotteries to be compared. We use the functional forms specified by Kahneman and Tversky:

$$u(x_i) = x_i^\alpha \quad \text{Eq. (3)}$$

$$w(p_i) = \frac{p_i^\gamma}{(p_i^\gamma + (1 - p_i)^\gamma)^{1/\gamma}}. \quad \text{Eq. (4)}$$

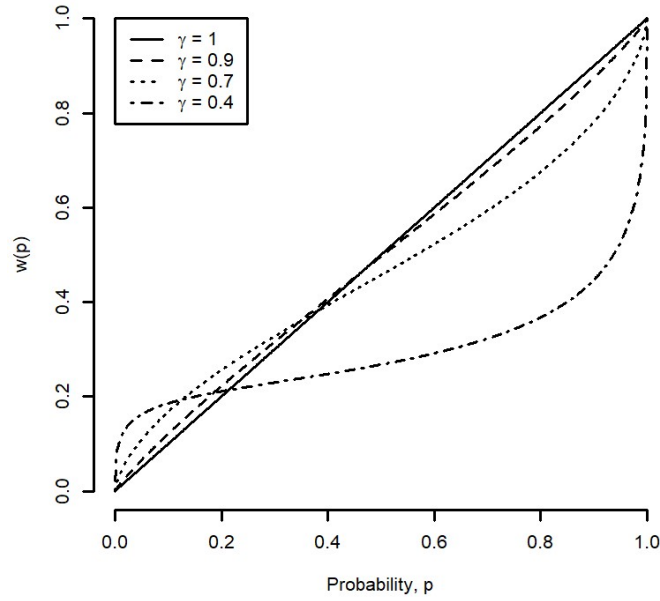
We define the log likelihood function $\mathcal{L}_{\mathbb{F}} = \ell_{ab} + \ell_{cd} + \dots$ and estimate α and γ as $\arg \max_{\alpha, \gamma} \mathcal{L}_{\mathbb{F}}$. There are two main underlying assumptions. First is the assumption that lottery preferences are independent (independence assumption); $y_{ab} \perp y_{cd}$ for all a, b, c, d . Second is the assumption that the representative agent does not change across lottery comparisons (representative agent assumption); preferences can be represented by a single static CPV function $v(\cdot)$ with an identically distributed stochastic term added. The two assumptions together yield $V(f_i) = v(f_i) + r_i$, with r_i i.i.d. The hypothesis test becomes $H_0: \gamma = 1$, $H_A: \gamma < 1$, with $\gamma = 1$ corresponding to $w(p_i) = p_i \Rightarrow \pi_i = p_i$ and $\gamma < 1$ corresponding to $w(p_i) \neq p_i \Rightarrow \pi_i \neq p_i$. Thus $\gamma = 1$ implies EUT, and $\gamma < 1$ implies nonlinear probability weighting in line¹¹ with CPT.

2b) Further discussion about PDWs

To understand how PDWs are determined under functional forms (Eq. (3) and Eq. (4)), it is important to note that the PDWs depend both on the weighting function $w(\cdot)$ and the order of the cash prizes to which the probabilities are assigned. The PDWs can be interpreted as the difference in weight (determined by $w(\cdot)$) put on the events: “the outcome is at least as good as x_i ” and “the outcome is strictly better than x_i ”. The following section will try to shed some light on the dynamics of the weighting function and PDWs and their rank-dependent nature with some examples.

¹¹ $\gamma > 1$ gives nonlinear probability weighing opposite to that proposed by Tversky and Kahneman.

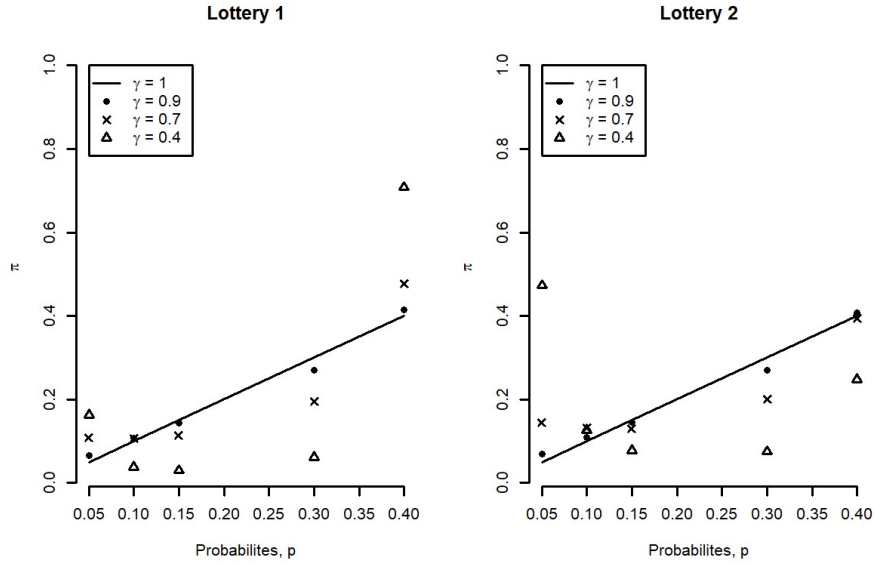
The weighting function (Figure 1)



A plot of the weighting function (Eq. (4)) for different parameter values. The solid line ($\gamma = 1$) represents EUT/linearity in probabilities.

As the value of γ changes, so do the essential characteristics of the weighting function (Figure 1). First, a lower γ increases the steepness of the function around the endpoints of the interval. Second, a lower γ decreases the steepness of the function on a large part of the interior of the interval. This implies high sensitivity to changes in low and high probabilities and low sensitivity to changes in moderate probabilities. Third, a lower γ decreases the fixed point (where $w(p) = p$) of the function. If this weighting function alone determined the PDWs, every probability lower than the fixed point would be overweighted, and all other probabilities would be underweighted. But now, the most we can say is that very low probabilities are *probable* to be overweighted. Further, the overweighting will be more severe with a lower γ . Still, this depends on the ordering of the probabilities. We introduce two example lotteries to demonstrate (Figure 2). Let Lottery 1 be structured as $(p_1, p_2, p_3, p_4, p_5) = (0.4, 0.3, 0.15, 0.1, 0.05)$, with p_1 assigned to the lowest prize, p_2 to the second lowest, etc. We call the same lottery but with the opposite order of probabilities Lottery 2.

PDWs (Figure 2)



Two graphs of two different lotteries showing how PDWs depend on both probabilities and the order of prizes. In the left graph, the prizes are decreasing in probabilities (0.05 is the probability of the largest prize; 0.40 is the probability of the smallest). In the right graph, the reverse is true.

The smallest γ (0.4) gives PDWs that differ the most from probabilities. Over the probabilities (0.1, 0.15, 0.3), which can be seen as moderate, we see lower sensitivity to changes in γ . Interestingly, in the case of Lottery 1, the highest probability is overweighted the most (in absolute terms). However, the weighting of this probability changes drastically when the order of probabilities changes; in Lottery 2, it is instead underweighted. This is due to the fact that $\pi_1^1 = w(\sum_{j=1}^5 p_j) - w(\sum_{j=2}^5 p_j) = w(1) - w(0.6)$, while $\pi_1^2 = w(0.4)$, where π_1^1 is the weight put on the outcome with probability 0.4 in Lottery 1, and π_1^2 is the weight put on the outcome with the same probability in Lottery 2.

Last, the functional forms of CPT are merely an attempt to test the qualitative properties of the theory against observed behavior in a relatively comprehensive and easy way. Our estimates of γ and α should be seen as a “summary” of market behavior put into two numbers.

3) Method and Data

3a) Lottery bonds

The dataset used for our empirical tests is constructed from daily price data over Swedish lottery bonds issued between 1918 and 1994. Lottery bonds are bonds issued by the Swedish Treasury that were introduced to make saving more attractive to the average Swede. The term of Swedish lottery bonds ranges between 5 to 10 years. They differ from regular bonds by the fact that the size of each coupon payment is determined via lottery instead of being certain. The distribution of coupon payments is announced by the Treasury together with each issue and can vary over time. Every issue has at least two lottery dates per year. For example, the prize distribution of the second lottery of the first bond issued in 1994 is presented by the Treasury as follows.

Prize distribution of a lottery bond (Figure 3)		
Prize	Number of Prizes	Number of bonds
500	52 000	2 600 000
1 000	940	2 600 000
5 000	74	2 600 000
10 000	35	2 600 000
20 000	17	2 600 000
100 000	17	2 600 000
500 000	1	2 600 000
1 000 000	1	2 600 000

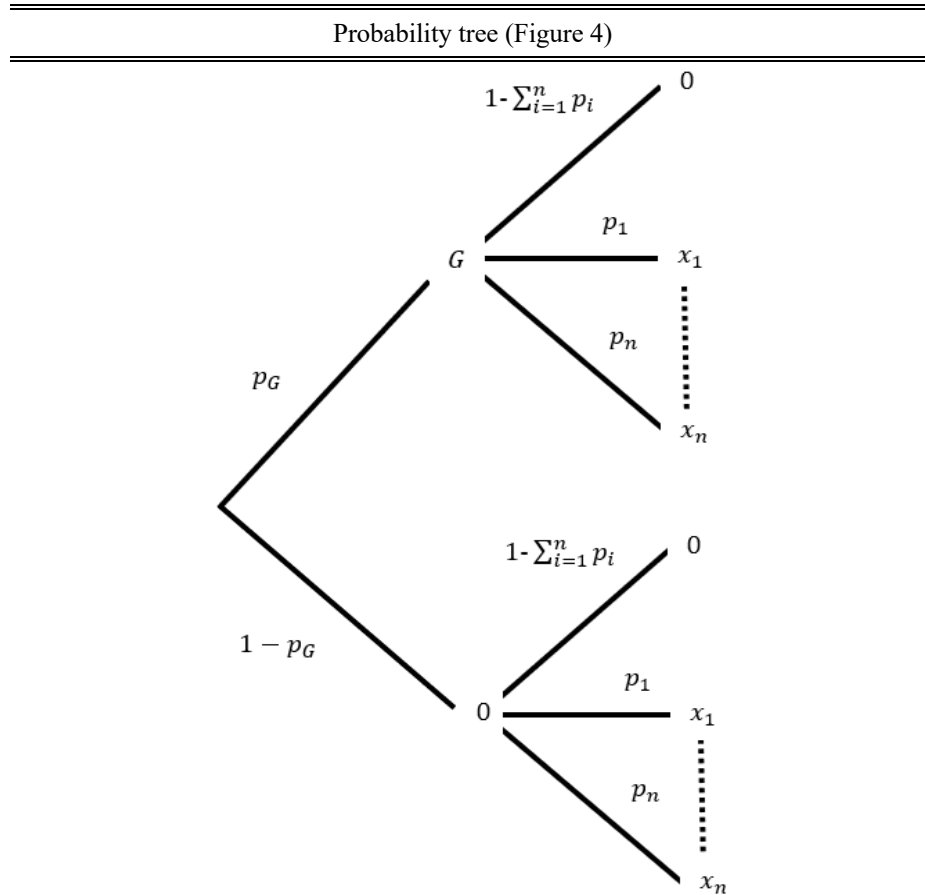
A figure displaying Lottery 94:12 as the Swedish Treasury presented it to potential buyers of the bond issue 94:1. The prize of 500 is the partial guarantee. Prizes are denoted in SEK.

Over the period of our data, the Treasury often issues more than one lottery bond per year. The first bond issued in 1994 is denoted 94:1. The next bond issued in 1994 is denoted 94:2, etc. During the period of our data, most lottery bonds were held in accounts at Swedish banks. To allow time for banks to determine which bonds held in their accounts had winning numbers, the trading of lottery bonds was suspended six business days before a lottery and six business days after, a total of 12 business days.

An important concept is that of the “partial guarantee”. Every lottery bond has a series number and an order number. For instance, 95:1 has 540 series with bonds of order number 1-1 000 in each, making a total of 540 000 bonds. Each lottery is structured such that the winner of all prizes, except the smallest, is determined by drawing a series number and order number corresponding to one particular bond. In contrast, the winners of the smallest prize are determined by drawing only an order number corresponding to the bond with that order number in all series. Further, the order numbers are divided into intervals (or sequences), which are presented together with each issue. For the smallest payments, one order number is drawn from each of these intervals. By owning bonds with order numbers corresponding to a complete interval, one is guaranteed to win at least the smallest payment at every lottery: the partial guarantee. For issue 95:1, one order number is drawn from every interval of 50 bonds, making a total of 200 intervals. For instance, by holding bonds of order number 251 to 300, one is guaranteed at least 500 SEK

every lottery. Because of the partial guarantee, lottery bonds are traded in bundles of sequenced order numbers called *sequenced bundles*. Bonds are also traded in *mixed bundles*, which are bundles of bonds that are randomly put together. Green and Rydqvist (1997) show that prices of sequenced bundles and mixed bundles systematically differ.

As a result of the partial guarantee, the lotteries are of a two-tier structure. First, the winners of the partial guarantee are drawn, and the winning bonds are then “replaced”. Second, the winners of the larger prizes are drawn. From the view of a holder of a single bond, you can win nothing or the partial guarantee *and* nothing or one of the larger prizes. Hence, Figure 3 does not paint the entire picture; the probability tree (Figure 4) fills in the gaps.



A probability tree explaining the two-stage nature of lotteries tied to Swedish lottery bonds. G is the partial guarantee, and x_i is prize level i (1000 to 1 million SEK in Figure 3). p_G is the probability of winning the partial guarantee in stage one. p_i is the probability of winning prize level x_i in stage 2.

The second stage of the lottery adds seven prize levels not depicted in Figure 3: the prizes 1 000 to 1 million SEK, each with the partial guarantee of 500 SEK added. The probability measures in Figure 3 must be slightly adjusted to account for the additional outcomes. The correct objective probabilities can be derived from the information in Figure 3 with the help of the probability tree (Figure 4). The objective probability of winning exactly x_i is the probability of winning nothing in the first stage multiplied by the probability of winning x_i in the second stage, i.e., $(1 - p_G)p_i$. Figure 10 in Chapter 4 showcases the complete two-stage version of Lottery 94:12.

Most prior research on Swedish lottery bonds has centered around the ex-day behavior of bond prices. “Ex-day” refers to the day after the right to a dividend (the lottery in our case) of a

security has been decided. An investor that buys a security (and does not sell) before its ex-day will receive the dividend corresponding to this ex-day. An investor buying the security on the ex-day or later will not be eligible to receive the dividend. Green and Rydqvist (1999) show that bond price behavior around ex-days is mainly driven by the tax rate because of the unique tax treatment of coupon payments (they are tax-exempt) and capital losses on lottery bonds. Before 1980, capital losses on lottery bonds could be used to offset income from any other source. An investor could, by buying a bond cum-lottery and selling post-lottery, use the capital loss to offset income from other sources while gaining a tax-exempt coupon. This kind of tax-arbitrage behavior was popularized by Roger Akelius (1974). Bond turnover around lottery drawings increased dramatically during this period. In 1980, the tax treatment was changed. Now, capital losses on bonds issued 1980 or earlier could be used to offset gains only on equities or other lottery bonds, and capital losses on bonds issued 1981 or later could be used to offset gains only from other lottery bonds. In 1990, the tax treatment of all living bonds was changed again, but the details are not important for our study. Green and Rydqvist (1999) demonstrate that under all tax regimes, the average price depreciation of bonds during lotteries is larger than the expected coupon.

3b) Construction of variables

To bridge theory and reality, we expand on f_a from Chapter 2. Let each bond m together with a time t correspond to a lottery $\delta \in a, b, c, \dots$. We create the variable $\hat{s}_\delta = k_{m,t}^{cum} - k_{m,t}^{post}$, where $k_{m,t}^{cum}$ is the cum lottery price of bond m at time t and $k_{m,t}^{post}$ is the post lottery price of the same bond. The goal of this section is to recreate Eq. (2) with observed data.

The representative agent who bought the bond for k_m^* and sells the bond cum lottery gains $k_{m,t}^{cum} - \tau(k_{m,t}^{cum} - k_m^*)$, where τ is the marginal tax rate and $\tau(k_{m,t}^{cum} - k_m^*)$ is the tax liability caused by the capital gain/loss. Selling post lottery, the agent will gain, $k_{m,t}^{post} + C(V(f_\delta)) - \tau(k_{m,t}^{post} - k_m^*)$, where $C(V(f_\delta))$ is her monetary valuation of lottery δ . Under a no-arbitrage condition¹², the agent must be indifferent between selling cum and post lottery.

Setting these two expressions equal and solving for $C(V(f_\delta))$ yields:

$$\begin{aligned} C(V(f_\delta)) &= k_{m,t}^{cum} - k_{m,t}^{post} + \tau(k_{m,t}^{post} - k_m^*) - \tau(k_{m,t}^{cum} - k_m^*) \\ &= k_{m,t}^{cum} - k_{m,t}^{post} - \tau(k_{m,t}^{cum} - k_{m,t}^{post}) \\ &= (k_{m,t}^{cum} - k_{m,t}^{post})(1 - \tau) \\ &= \hat{s}_\delta(1 - \tau) \end{aligned}$$

Combining with Eq. (1) and incorporating the tax effect $(1 - \tau)$ in an error term, we use this equality to estimate the price of a lottery as

$$s_\delta = \hat{s}_\delta + \epsilon_\delta. \quad \text{Eq. (5)}$$

Eq. (5) asserts that the price of the lottery is the price drop of the lottery bond with some error added. Substituting Eq. (5) into the left side of Eq. (2), the difference in the price of two lotteries a, b can be estimated as

$$s_a - s_b = \hat{s}_a - \hat{s}_b + \epsilon_{ab},$$

where $\epsilon_{ab} = \epsilon_a - \epsilon_b$ is the combined error.

¹² In deriving the no-arbitrage condition, we follow Green and Rydqvist (1999) but incorporate our novel $C(V(f_\delta))$ term.

To determine the representative agent's preferences (Eq. (2)) using observed price drops, the property $E[\epsilon_{ab}] = 0$ is not enough; we need $\epsilon_{ab} = 0$. First, the tax effect must be accounted for. The ex-day behavior of lottery bonds implies that $E[\epsilon_a] < E[\epsilon_b]$ if lottery a is tied to a bond subject to a more liberal tax policy than that which b is tied to because investors pay for both the lottery and the tax benefits. Therefore, we compare only lotteries taking place under the same tax policy. With this restriction, we have $E[\epsilon_a] = E[\epsilon_b]$ but not $\epsilon_a = \epsilon_b$. For the latter to be true, we compare only lotteries taking place on the same date and tied to bonds in mixed bundles. By including only lotteries tied to mixed bonds, the systematic price differences between sequenced and mixed bundles of bonds are controlled for. By comparing exclusively lotteries taking place on the same date, we fully account for time fixed effects. With the above conditions in place, we argue that we are left with

$$s_a - s_b = \hat{s}_a - \hat{s}_b. \quad \text{Eq. (6)}$$

Combining Eq. (6) and Eq. (2), the goal of this section is reached: $\hat{s}_a - \hat{s}_b > 0 \Leftrightarrow V(f_a) > V(f_b)$. Theoretical preference relationships (y_{ab}) can be inferred from observed variables ($\hat{s}_a$ and $\hat{s}_b$). To summarize, we argue that the difference in the price drop over two lotteries, tied to mixed bonds and taking place on the same date, is caused only by differences in lottery preferences. With these constraints in place, the available dataset of bonds shrinks from 1918-1994 to 1987-1994.

Last, we want the time span of our sample to be as recent and short as possible while still retaining an appropriate number of observations. Preferences are more likely to be stable over a shorter time span, increasing the validity of the representative agent assumption. By using the latest data available, our results are more applicable to the current society. We are left with the pairs of bonds (89:1,89:2), (91:1,91:2), (92:1,92:2), (93:1,93:2), and (94:1,94:2), giving a sample size of 116 lottery comparisons over a period of circa 11 years¹³.

Complete sample of lottery comparisons (Figure 5)

Pair of bonds		Number of lottery comparisons	Number of $y_{ab} = 1$	Date of first lottery comparison	Date of last lottery comparison	Number of distributions		Average expected coupon	
89:1	89:2	28	13	1990-01-18	1999-01-14	2	1	21.02	22.50
91:1	91:2	24	12	1991-11-14	1999-06-23	9	8	19.12	18.88
92:1	92:2	22	9	1993-01-07	2000-01-13	7	5	19.42	19.09
93:1	93:2	22	9	1994-02-03	2001-02-01	4	4	17.29	14.04
94:1	94:2	20	14	1995-03-16	2001-06-29	2	1	11.50	12.00

A table giving an overview of our sample of 116 lottery comparisons. "Number of distributions" and "Average expected coupon" correspond to the bonds in the "Pair of bonds" column. For example, 91:1 has nine different distributions; 91:2 has eight. "Number of $y_{ab} = 1$ " indicates the number of times a lottery of the left bond (from the "Pair of bonds" column) is valued higher than that of the right. "Average expected coupon" is denoted in SEK. No issue from 1990 is included because there is no other issue with coinciding lotteries.

Note that most of the issues have variability in their lottery distribution. For instance, 91:1 has nine different lotteries denoted 91:11, 91:12 etc. The first lottery of issue 91:1 has the distribution 91:11, the second 91:12, etc. From the ninth lottery onwards, every lottery has the distribution 91:19. This process follows for all bonds.

¹³ While we use bonds issued between 1989 and 1994, they come to term several years later. For instance, bonds issued in 1994 have lottery drawings until 2001.

The turnover of lottery bonds is small compared to other securities. On some days, there is no turnover for some bundles of bonds. This was common in the earlier parts of the 20th century but much less so during the 1980s and 1990s. When compiling our dataset, we include observations in which bonds were last traded at a maximum of three days before the suspension of trading tied to a lottery and traded first at a maximum of three days after the suspension. This is a potential source of bias in our data. If one bond that is traded directly after the suspension is compared with one that is first traded three days after, we could be confounding lottery preferences with other external factors occurring during those days. However, for the period of our sample, this potential issue only surfaces four times, merely in the case of a single day with no trading before or after the suspension. It is hard to argue that four infractions would significantly bias our estimates, and the exclusion of these observations could very well be worse than including them.

All lotteries in our dataset follow a similar structure to that of lottery 94:12 (Figure 3). There are seven prize levels, ranging from 1 000 to 1 million SEK, not counting the partial guarantee. The partial guarantee ranges from 500 to 750 SEK. The probability of winning nothing dominates (around 98%), and most of the residual probability mass is assigned to the lowest non-zero prize¹⁴. The probability distributions we deal with are quite different from those of laboratory experiments in which the probabilities can be more evenly distributed. This has implications for the conclusions we can draw based on our results.

3c) Estimation of population parameters

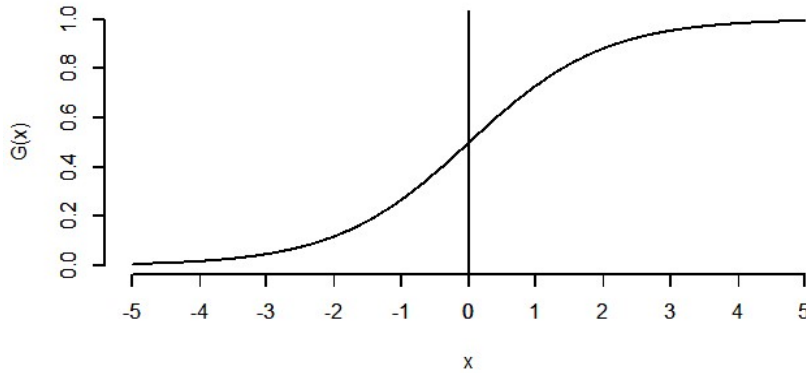
The decision of how to group subjects in the estimation phase is an important one. Gurevich, Kliger, and Levy (2009) are, like us, faced with market decisions instead of individual decisions. They estimate parameters separately for each stock in addition to providing a pooled estimate. With grouping, they can survey heterogeneity in risk attitudes between stocks. We are constrained by the number of observations. While bond pairs (Figure 5) seem a fitting stratification, no group would reach 30 observations¹⁵. We proceed to pool all observations to estimate α and γ .

To apply maximum likelihood estimation, the joint likelihood of observing our sample (116 binary choices over 232 lotteries) must be a function of CPT parameters α and γ . From the independence assumption, we can write the joint likelihood as the sum of log individual probabilities. From the representative agent assumption, market choices are modeled by a single stochastic CPV function. Through this function, we compare stochastic CPVs of lotteries, in which the difference (z_{ab}) in the stochastic terms (r_a and r_b) is cumulatively distributed according to $G(x) = \frac{e^x}{1+e^x}$.

¹⁴ The probability of winning the smallest non-zero prize conditional on winning any non-zero prize is also around 98%.

¹⁵ Evincen by the “Number of lottery comparisons” column.

Logistic cumulative probability function (Figure 6)



A graph showing the function we use to model the probability of choosing a given lottery over another. For instance, a difference of 1.5 in static CPV between two lotteries gives a 0.8 probability of picking the one with the higher static CPV.

Our G-function (Figure 6) has a range between 0 and 1 and is commonly used as the cumulative probability of the realization of observed variables (e.g., logit model). The properties of the G-function limit us to consider only binary decisions over lotteries. In Chapter 2, we show that introducing a stochastic term implies that the probability of choosing lottery a over lottery b is solely determined by the difference in their respective static CPV, i.e., Δ_{ab} . The cumulative probability of choosing a over b is then $G(\Delta_{ab})$. Intuitively, the more utility lottery a brings the agent, the more likely she is to prefer it to lottery b .

We see the application of stochastic utility as uniquely fitting in describing the representative agent on the bond market. The markets for bonds and stocks are innately characterized by random market noise: various forms of volatility, such as price corrections and price fluctuations. We posit that this volatility arises from the probabilistic nature of human decisionmaking (Rieskamp 2008). A great deal of contradictions are amended by modeling decisionmaking as a probabilistic process instead of a deterministic one. In our sample of 116 decisions, several involve choices over identical pairs, i.e., $f_a = f_c$ and $f_b = f_d$. Letting $\Delta_{ab} > 0$, identical lottery pairs, $\Delta_{ab} = \Delta_{cd}$, implies $\hat{s}_a > \hat{s}_b$ and $\hat{s}_c > \hat{s}_d$ in a deterministic world. But in our sample, we are confronted with instances of $\hat{s}_a > \hat{s}_b$ and $\hat{s}_c < \hat{s}_d$ and vice versa. A more realistic point of view is that of the probabilistic: identical lottery pairs increase the likelihood of the inequalities pointing in the same direction but do not preclude the possibility of market noise, such as intraday volatility, distorting the stochastic CPV sufficiently to switch the preference relationship. Our G-function and observed y_{ab} -variables enable us to estimate γ and α with a maximum likelihood function that performs a grid search.

3d) Inferences

To test our hypothesis, we need to estimate confidence intervals for $\hat{\gamma}$. Confidence intervals are provided for $\hat{\alpha}$ as well. To this end, we use a bootstrapping process, the foundation of which rests on repeated resampling. From our realized sample, observations are drawn with replacement until there are again 116 observations. We use the estimation method from 3c to estimate γ_j^b and α_j^b , which are the $\hat{\gamma}$ and $\hat{\alpha}$ of this particular bootstrap sample. The process is repeated 1 000 times, and every pair of estimates (γ_j^b and α_j^b) is stored. Under certain conditions—which we assume fulfilled—the characteristics of the distributions of γ_j^b and α_j^b

will reflect those of the sampling distributions that are important when performing statistical inference (Horowitz 2001). The bias in our point estimates is estimated by $Bias(\hat{\theta}) = E(\hat{\theta}) - \theta \approx \overline{\theta^b} - \hat{\theta}$, where θ is the parameter of interest. Below are three ubiquitous confidence intervals presented by Hesterberg (2011) and Johnson (2001).

The simplest is the *percentile interval*:

$$(\hat{F}^{-1}(\beta/2), \hat{F}^{-1}(1 - \beta/2)),$$

where $\hat{F}$ is the bootstrap approximated sampling distribution and β is the confidence level.

If we want to take bias into account, we use the *pivot confidence interval*. It is derived from the approximated distribution of the variable $\theta - \hat{\theta}$.

$$\begin{aligned} P(\hat{\theta} - \hat{F}^{-1}(1 - \beta/2) < \hat{\theta} - \theta^b < \hat{\theta} - \hat{F}^{-1}(\beta/2)) &\approx \\ \approx P(\hat{\theta} - \hat{F}^{-1}(1 - \beta/2) < \theta - \hat{\theta} < \hat{\theta} - \hat{F}^{-1}(\beta/2)) &= \\ = P(2\hat{\theta} - \hat{F}^{-1}(1 - \beta/2) < \theta < 2\hat{\theta} - \hat{F}^{-1}(\beta/2)), \end{aligned}$$

which yields the confidence interval:

$$(2\hat{\theta} - \hat{F}^{-1}(1 - \beta/2), 2\hat{\theta} - \hat{F}^{-1}(\beta/2)).$$

The last confidence interval is the *bias-corrected and accelerated interval* (BCa). This interval considers not only the bias but also the “acceleration”. With acceleration, the standard error of the point estimate depends on the point estimate itself and does not reflect the true standard deviation of the sampling distribution. The BCa confidence interval is defined as:

$$(\hat{F}^{-1}(p(\beta/2)), \hat{F}^{-1}(p(1 - \beta/2))),$$

where $p(x) = \Phi\left(z_0 + \frac{z_0 + \Phi^{-1}(x)}{1 - a(z_0 + \Phi^{-1}(x))}\right)$, and $\Phi(\cdot)$ is the cumulative normal density function. The bias estimate, $z_0 = \Phi^{-1}(\#(\theta^b < \hat{\theta})/B)$, is estimated using the number of bootstrap estimates that fall below the point estimate, where $\#(x)$ is a function counting the number of times the statement x is true, and B is the number of bootstrap estimates. The acceleration parameter is defined as $a = \frac{\sum_{i=1}^n (\hat{\theta}_i - \bar{\theta}_i)^3}{6(\sum_{i=1}^n (\hat{\theta}_i - \bar{\theta}_i)^2)^{3/2}}$, where $\hat{\theta}_i$ is the point estimate calculated by excluding observation i and $\bar{\theta}_i$ is the average of those values.

Confidence intervals are calculated for both α and γ using an appropriate bootstrapping package in R. We follow convention by estimating 1 000 bootstraps to create our confidence intervals. To test our hypothesis, we survey the three confidence intervals presented. If no interval for $\hat{\gamma}$ intersects 1, there is a rejection of our null hypothesis.

4) Results

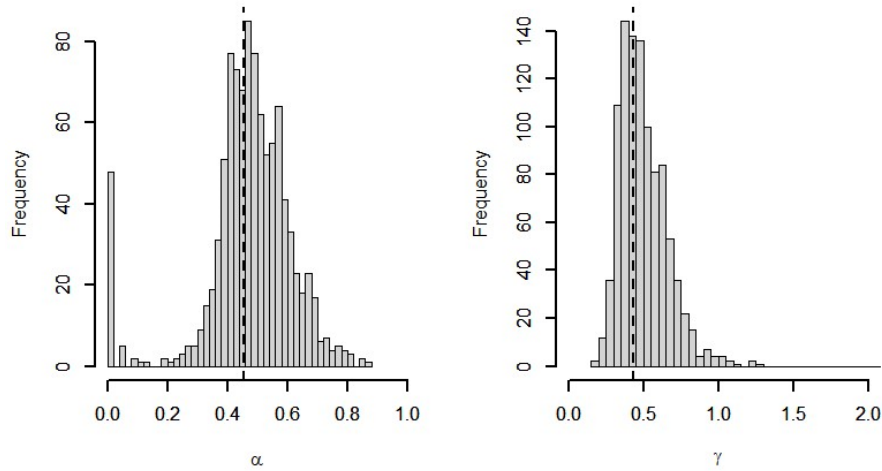
Main results (Figure 7)

	α	γ
Point estimate	0.46	0.43
Bias	0.02	0.13
Percentile interval	(0.00 , 0.73)	(0.28 , 0.92)
Pivot confidence interval	(0.18 , 0.91)	(-0.05 , 0.59)
BCa interval	(0.00 , 0.65)	(0.23 , 0.74)

A table presenting the CPT parameters of our study. Confidence intervals and bias are presented as well. The different confidence intervals are described in section 3d. Every confidence interval is a 95% confidence interval. The leftmost endpoint of two α -confidence intervals is rounded to zero (the bootstrap estimates are of order 10^{-6}).

For both point estimates, none of the confidence intervals include 1. We reject our null hypothesis at the 0.025 confidence level¹⁶; the representative agent displays nonlinear weighting of objective probabilities (see Figure 9 for a graphical depiction). The estimated bias in $\hat{\alpha}$ is almost negligible, whereas there is a relatively large estimated upward bias in $\hat{\gamma}$.

Histograms of bootstraps (Figure 8)

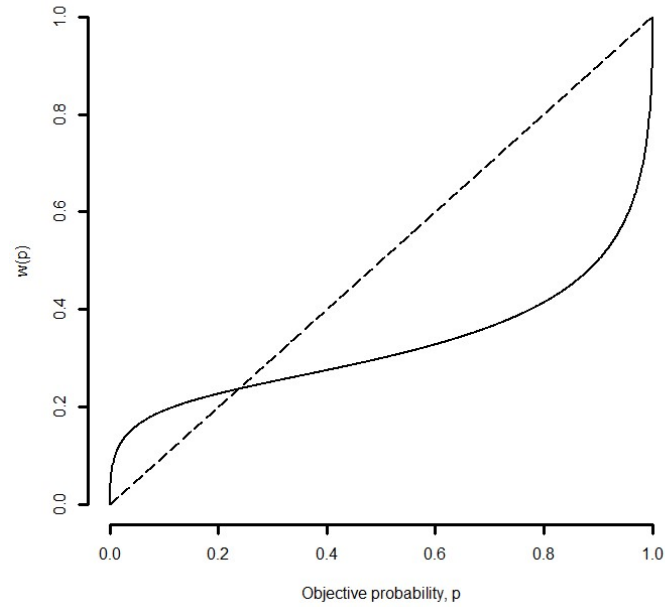


Two graphs showing the distribution of bootstrap estimates for our parameters. The dashed lines mark our point estimates. Seven outliers (ranging from 9 to 11) are excluded from the right graph to illustrate the properties of the distribution of γ -estimates in more detail.

Both distributions in Figure 8 are far from normally distributed, and there is a significant bunching of α -estimates at 0.

¹⁶ Because of our one-sided hypothesis test.

The weighting function of the representative agent (Figure 9)



A graph depicting the weighting function implied by our estimate $\gamma = 0.43$. The dashed line represents EUT ($\gamma = 1$), i.e., linearity in probabilities.

Figure 9 gives a clear account of the significant nonlinear weighting of objective probabilities exhibited by the representative agent. The function has a fixed point at roughly 0.23. Moreover, the function is very steep around the endpoints of the interval but relatively flat on most of the interior.

PDWs contrasted with objective probabilities (Figure 10)			
Prize	Obj. Probability ($\gamma = 1$)	PDW ($\gamma = 0.43$)	$\frac{PDW}{Obj. Prob}$
0	0.979591	0.87	0.89
500	0.020	0.096	4.78
1 000	0.00035	0.017	48.25
1 500	$7.2 \cdot 10^{-6}$	0.00073	100.20
5 000	$2.8 \cdot 10^{-5}$	0.0035	125.53
5 500	$5.7 \cdot 10^{-7}$	$9.0 \cdot 10^{-5}$	157.51
10 000	$1.3 \cdot 10^{-5}$	0.0025	189.35
10 500	$2.7 \cdot 10^{-7}$	$6.3 \cdot 10^{-5}$	234.30
20 000	$6.4 \cdot 10^{-6}$	0.0018	278.46
20 500	$1.3 \cdot 10^{-7}$	$4.4 \cdot 10^{-5}$	339.51
100 000	$6.4 \cdot 10^{-6}$	0.0035	539.78
100 500	$1.3 \cdot 10^{-7}$	0.00016	1185.96
500 000	$3.8 \cdot 10^{-7}$	0.00056	1481.11
500 500	$7.7 \cdot 10^{-9}$	$1.4 \cdot 10^{-5}$	1832.28
1 000 000	$3.8 \cdot 10^{-7}$	0.0013	3550.55
1 000 500	$7.7 \cdot 10^{-9}$	0.00030	39151.56

A table showing the difference between our estimated weights (PDW) and those predicted by EUT (Obj. Probability) over outcomes of the complete two-stage version of lottery 94:12 (Figure 3).

In our CPT model, the zero-prize outcome is underweighted compared to its objective probability ($0.87 < 0.98$). The other prizes are considerably overweighted; the weight mass assigned to all non-zero prizes under EUT is 0.02 (summing objective probabilities) compared to 0.13 under CPT (summing PDWs). The relative difference between the PDWs and their corresponding objective probabilities widens as the prize level increases. The weight assigned to the event of winning at least 1 million SEK in prizes is overweighted by a factor of 4300¹⁷. This pattern follows for every lottery of our sample.

¹⁷ Calculated by dividing the sum of the PDWs assigned to prize levels equal to or greater than 1 million SEK by the sum of the objective probabilities assigned to the same prize levels.

5) Discussion

5a) Representative agent assumption

The representative agent assumption is crucial for us to facilitate estimation. Camerer and Ho (1994) employ it as a remedy for a scarcity of observations on individual subjects. We face the same predicament and argue that individual purchasing decisions over lottery bonds can be aggregated into a single parameter estimate for a single representative agent. Two factors strengthen its use. First, with the narrow timespan of our sample, the composition of the lottery bond market is likely to be stable, making the representative agent assumption more fitting. Second, market data is especially geared towards the assumption. Markets are an aggregation of thousands of people exhibiting the modal characteristics of its participants. Still, we are surely missing important dynamics between market participants by describing the market's risk attitudes with a single parameter estimate. Bruhin, Fehr-Duda, and Epper (2010) survey the heterogeneity of risk preferences within and across countries. They discover a persistent share (around 20 %) of people across countries to be EUT maximizers. The remaining share is split between agents exhibiting moderate probability weighting and more radical probability weighting. Even if the lottery bond market is characterized by a similar split of risk attitudes, the nature of our dataset makes differentiating between them impossible.

5b) Independence assumption

The independence assumption is an implication of the construction of our log likelihood function. A main concern is self-perpetuating valuations. To demonstrate, let two lotteries be compared on multiple occasions. Assume that the second lottery is preferred at the first decision. If, in the following instances, the market reminisces about and confirms the previous relative valuation, independence is broken. With self-reinforcing decisionmaking stemming from the first decision, stochastic terms are serially correlated. Still, we believe the concern is unlikely. Self-perpetuating valuations of this kind stand in conflict with efficient markets.

5c) Validity of analysis

While our point estimates are reasonable and in line with other empirical findings (Figure 11), some of the bootstrap estimates are not. A few bootstrapped γ 's are extremely large. While a γ greater than 1 is not necessarily unreasonable, it goes against the empirical findings that motivated CPT. Specifically, it indicates that lower probabilities are prone to be underweighted or even neglected, while higher ones are prone to be overweighted. For γ 's greater than 7, which are encountered in our sample, probabilities as high as 0.3 could be neglected; this, we find unreasonable. Further, 6% of estimated α 's are approximately zero. In these cases, the representative agent exhibits no marginal utility in money. The fact that our sample of lottery comparisons and/or our estimation process can imply something that is so incoherent with economic theory suggests that some degree of caution is needed.

A likely cause of these breaches of logic and the significant bias in γ is limitations in functional forms (Eq. (3) and Eq. (4)). We are faced with more complex lotteries than other studies. Further, there are repeated close calls of stochastic dominance in which the almost dominated lottery is preferred. In some cases, it is a question of a basis point difference in one or two probabilities that makes a lottery not dominated. With just one parameter to determine the form of our weighting function (Figure 9), these slight differences in probabilities cannot be effectively modeled. Instead, our maximum likelihood process is forced to lower the marginal

utility of money. As this reaches 0, the stochastic dominance factor is attenuated. Now, say that such observations are drawn adversely many times in a bootstrap sample. Then, the parameter values that maximize the likelihood of this observed sample will likely be unreasonable. One approach to these kinds of issues, used by Nilsson, Rieskamp, and Wagenmakers (2011), is to set the domain of the parameters to a priori specified interval. If applied in our case, it would decrease the standard errors of the point estimates but not solve the root cause of our unreasonable values. Still, our point estimates are reasonable when put in the context of other CPT model estimation papers (Figure 11). That we get there with fewer restrictions is a sign of the strength of our estimation process. Ultimately, the existence of unreasonable values could be an indication that choices under uncertainty in this market are better described by other functional forms than the ones suggested by Tversky and Kahneman (1992).

An important factor in both lab and field experiments is the framing of probabilities. That is, how should the probabilities of a prospect be depicted for an experiment to reflect a decision made in real life? In real life, you rarely have access to stated objective probabilities, and if you do, the framing of these might affect your decisionmaking. To demonstrate, a probability of 0.3 can be stated as

$$30\% , \frac{3}{10} , 0.3 , 3/10 , \text{“3 out of 10”} , 3 \cdot 10^{-1} , \text{●} , 6/20 , 0.03 \cdot 10.$$

Although some of these depictions are quite unorthodox, they emphasize that the subjective probability of an event can very well depend on how its objective probability is framed. This becomes even more apparent when smaller and more complex probabilities are introduced. Most lab experiments use the standard decimal form (Tversky and Kahneman 1992) or percentages (Zeisberger, Vrecko, and Langer 2012; Stott 2006) to express objective probabilities, but some use alternatives such as pie charts (Harrison and Rutström 2009). In our case, a 17 in 2 600 000 chance to win 100 000 SEK (Figure 3) could be interpreted more favorably than a probability of 0.00000653846 to win 100 000 SEK. Concurrently, the two-tier structure (Figure 4) of the lotteries and the extremely small probabilities (Figure 3) may make it difficult for the representative agent to weight outcomes linearly in objective probabilities even if she wants to. Ultimately, we cannot rule out the possibility that the framing process and complexity of the lotteries are the main drivers of observed behavior. We can only say that the representative agent exhibits nonlinear weighting of *objective* probabilities.

5d) Generalization

By focusing on the estimation of CPT, we inherently constrain our study’s implications. The estimated parameters of our model reflect not only the chosen parametric form but also a specific population at a specific point in time, making decisions under a particular context. For studies set in the laboratory, the limitations are even more stark. It is in no way obvious that decisionmaking in fictional scenarios is in any way generalizable. By carrying out a field study, we are at the least making conclusions about the specific “field” of lottery bonds in Sweden, which we find valuable in itself. Moreover, the specific population under consideration is considerably larger than most experimental studies to date. But exactly who makes up this population and, in turn, whom our representative agent embodies are difficult questions. Lottery bonds in Sweden were launched to entice the average Swede to start saving. By the 1990s, the wild turnover (a result of traders chasing arbitrage opportunities) around lottery drawings had subsided with changed tax laws. Taken together, we see the representative agent to reasonably approximate the average retail bond investor in Sweden during the 1990s and 2000s. Still, to draw wide generalizations from our results is problematic. First, during the period of our sample,

Sweden was dealing with a severe financial crisis and, later, a boom. It is conceivable that people's risk attitudes during such periods do not reflect their behavior in more stable times. Second, Swedes might exhibit different weighting of probabilities than other nationalities; Rieger, Wang, and Hens (2017) show wide differences in CPT parameters between countries. Third, decisionmaking under uncertainty is likely to be context-dependent (Figure 11). Considering the significant difference between the estimates of Gurevich, Kliger, and Levy (2009) and ours (we find more nonlinear functions), risk behavior in options markets seemingly varies from risk behavior in lottery bond markets. Options markets are likely to be characterized by a different kind of investor than bond markets.

CPT parameters for different studies (Figure 11)

	α	γ
Present paper	0.46	0.43
Tversky and Kahneman 1992	0.88	0.61
Jou and Chen 2013	0.92	0.80
Rieger, Wang, and Hens 2017	0.46	0.50
Bocquého, Jacquet, and Reynaud 2014	0.28	0.66
Gurevich, Kliger, and Levy 2009	0.98	0.84

A table presenting CPT estimates for six different studies (ours included) set in different contexts. In Kahneman and Tversky's study, graduate students at Berkeley and Stanford choose between different money gambles. Jou and Chen's study deals with the risk attitudes of Taiwanese highway drivers. Rieger, Wang, and Hens survey the risk preferences of a large number of undergraduate students in 53 countries. Bocquého, Jacquet, and Reynaud research the risk aversion of French farmers to risky gambles. In Gurevich, Kliger, and Levy's paper, the preferences of investors in US options markets are explored.

Further, the lotteries in our sample are of a very specific character; only a very small portion of the probability mass is allocated to winning a non-zero prize. Our study exclusively focuses on the shape of the weighting function around the endpoints of the unit interval, and interpolation of these results to the whole interior should be done with caution. For our results to be more comparable to other studies, we need greater variance in the probability density of our lotteries.

5e) Further research

With the above discussion in mind, there are multiple ways our sample could be examined further. Parameters could be estimated using other functional forms for $u(\cdot)$ and $w(\cdot)$. With more advanced forms, the underlying decision processes may be modeled more accurately. Other estimation methods than maximum likelihood can be used to capitalize on the fact that we can infer cash equivalences of the lotteries up to some, mostly tax-dependent, error. In the present paper, these cash equivalences are only used in an ordinal manner. Further, lotteries in our sample frequently share prize levels and probabilities. This might provide ample testing ground for the concept of *Cancellation* presented by Kahneman and Tversky (1979). Another angle worth exploring is how, from the perspective of evolutionary economics, the behavior observed in this paper can survive (Alchian 1950). Intuitively, the optimal behavior in a financial market to ensure survival (i.e., positive profits) should be to maximize expected returns.

6) Conclusion

We establish Swedish lottery bond investors to significantly overweight low objective probabilities of winning large prizes. Specifically, events of low probability are overweighted by a factor of 4300. We conclude that investors do not adhere to linearity in objective probabilities. Our estimates for α and γ imply more diminishing sensitivity in the representative agent's utility function and more drastic probability weighting than those of Tversky and Kahneman (1992) imply. Estimates from Gurevich, Kliger, and Levy's (2009) field study are even more linear. The difference may indicate that decisionmaking under uncertainty is highly dependent on the context. Further, nonsensical values from the bootstrapping process suggest limitations in our functional forms. The former poses problems in generalizing our findings, and the latter underscores the necessity for further investigation.

7) References

- Abdellaoui, Mohammed. 2000. "Parameter-Free Elicitation of Utility and Probability Weighting Functions." *Management Science* 46 (11): 1497–1512. <https://doi.org/10.1287/mnsc.46.11.1497.12080>.
- Akelius, Roger. 1974. *Allt om premieobligationer*. Jönköping: Akelius' databöcker.
- Alchian, Armen A. 1950. "Uncertainty, Evolution, and Economic Theory." *Journal of Political Economy* 58 (3): 211–21. <https://doi.org/10.1086/256940>.
- Allais, M. 1953. "Le Comportement de l'Homme Rationnel Devant Le Risque: Critique Des Postulats et Axiomes de l'Ecole Americaine." *Econometrica* 21 (4): 503. <https://doi.org/10.2307/1907921>.
- Bernheim, B. Douglas, and Charles Sprenger. 2020. "On the Empirical Validity of Cumulative Prospect Theory: Experimental Evidence of Rank-Independent Probability Weighting." *Econometrica* 88 (4): 1363–1409. <https://doi.org/10.3982/ECTA16646>.
- Bocquého, Géraldine, Florence Jacquet, and Arnaud Reynaud. 2014. "Expected Utility or Prospect Theory Maximisers? Assessing Farmers' Risk Behaviour from Field-Experiment Data." *European Review of Agricultural Economics* 41 (1): 135–72. <https://doi.org/10.1093/erae/jbt006>.
- Bruhin, Adrian, Helga Fehr-Duda, and Thomas Epper. 2010. "Risk and Rationality: Uncovering Heterogeneity in Probability Distortion." *Econometrica* 78 (4): 1375–1412.
- Camerer, Colin F., and Teck-Hua Ho. 1994. "Violations of the Betweenness Axiom and Nonlinearity in Probability." *Journal of Risk and Uncertainty* 8 (2): 167–96. <https://doi.org/10.1007/BF01065371>.
- Friedman, Milton, and L. J. Savage. 1948. "The Utility Analysis of Choices Involving Risk." *Journal of Political Economy* 56 (4): 279–304. <https://doi.org/10.1086/256692>.
- Green, R. 1999. "Ex-Day Behavior with Dividend Preference and Limitations to Short-Term Arbitrage: The Case of Swedish Lottery Bonds." *Journal of Financial Economics* 53 (2): 145–87. [https://doi.org/10.1016/S0304-405X\(99\)00019-7](https://doi.org/10.1016/S0304-405X(99)00019-7).
- Green, Richard C., and Kristian Rydqvist. 1997. "The Valuation of Nonsystematic Risks and the Pricing of Swedish Lottery Bonds." *Review of Financial Studies* 10 (2): 447–80. <https://doi.org/10.1093/rfs/10.2.447>.
- Gurevich, Gregory, Doron Kliger, and Ori Levy. 2009. "Decision-Making under Uncertainty – A Field Study of Cumulative Prospect Theory." *Journal of Banking & Finance* 33 (7): 1221–29. <https://doi.org/10.1016/j.jbankfin.2008.12.017>.
- Harrison, Glenn W., and E. Elisabet Rutström. 2009. "Expected Utility Theory and Prospect Theory: One Wedding and a Decent Funeral." *Experimental Economics* 12 (2): 133–58. <https://doi.org/10.1007/s10683-008-9203-7>.
- Hesterberg, Tim. 2011. "Bootstrap." *WIREs Computational Statistics* 3 (6): 497–526. <https://doi.org/10.1002/wics.182>.
- Horowitz, Joel L. 2001. "The Bootstrap." In *Handbook of Econometrics*, 5:3159–3228. Elsevier. [https://doi.org/10.1016/S1573-4412\(01\)05005-X](https://doi.org/10.1016/S1573-4412(01)05005-X).
- Johnson, Roger W. 2001. "An Introduction to the Bootstrap." *Teaching Statistics* 23 (2): 49–54. <https://doi.org/10.1111/1467-9639.00050>.
- Jou, Rong-Chang, and Ke-Hong Chen. 2013. "An Application of Cumulative Prospect Theory to Freeway Drivers' Route Choice Behaviours." *Transportation Research Part A: Policy and Practice* 49 (March): 123–31. <https://doi.org/10.1016/j.tra.2013.01.011>.
- Kahneman, Daniel, and Amos Tversky. 1979. "Prospect Theory: An Analysis of Decision under Risk." *Econometrica* 47 (2): 263. <https://doi.org/10.2307/1914185>.
- Kattsoff, L. O. 1945. "Theory of Games and Economic Behavior. By John von Neumann and Oskar Morgenstern. Princeton, N. J.: Princeton University Press, 1944. 625 Pp. \$10.00." *Social Forces* 24 (2): 245–46. <https://doi.org/10.2307/2572550>.

- Murphy, Ryan O., and Robert H. W. ten Brincke. 2018. "Hierarchical Maximum Likelihood Parameter Estimation for Cumulative Prospect Theory: Improving the Reliability of Individual Risk Parameter Estimates." *Management Science* 64 (1): 308–26. <https://doi.org/10.1287/mnsc.2016.2591>.
- Nilsson, Håkan, Jörg Rieskamp, and Eric-Jan Wagenmakers. 2011. "Hierarchical Bayesian Parameter Estimation for Cumulative Prospect Theory." *Journal of Mathematical Psychology* 55 (1): 84–93. <https://doi.org/10.1016/j.jmp.2010.08.006>.
- Quiggin, John. 1982. "A Theory of Anticipated Utility." *Journal of Economic Behavior & Organization* 3 (4): 323–43. [https://doi.org/10.1016/0167-2681\(82\)90008-7](https://doi.org/10.1016/0167-2681(82)90008-7).
- Rieger, Marc Oliver, Mei Wang, and Thorsten Hens. 2017. "Estimating Cumulative Prospect Theory Parameters from an International Survey." *Theory and Decision* 82 (4): 567–96. <https://doi.org/10.1007/s11238-016-9582-8>.
- Rieskamp, Jörg. 2008. "The Probabilistic Nature of Preferential Choice." *Journal of Experimental Psychology: Learning, Memory, and Cognition* 34 (6): 1446–65. <https://doi.org/10.1037/a0013646>.
- Stott, Henry P. 2006. "Cumulative Prospect Theory's Functional Menagerie." *Journal of Risk and Uncertainty* 32 (2): 101–30. <https://doi.org/10.1007/s11166-006-8289-6>.
- Tversky, Amos, and Daniel Kahneman. 1992. "Advances in Prospect Theory: Cumulative Representation of Uncertainty." *Journal of Risk and Uncertainty* 5 (4): 297–323. <https://doi.org/10.1007/BF00122574>.
- Wu, George, and Richard Gonzalez. 1996. "Curvature of the Probability Weighting Function." *Management Science* 42 (12): 1676–90. <https://doi.org/10.1287/mnsc.42.12.1676>.
- Zeisberger, Stefan, Dennis Vrecko, and Thomas Langer. 2012. "Measuring the Time Stability of Prospect Theory Preferences." *Theory and Decision* 72 (3): 359–86. <https://doi.org/10.1007/s11238-010-9234-3>.