

Birchby, Jeff; Gigliotti, Gary; Sopher, Barry

Working Paper

Consistency and aggregation in individual choice under uncertainty

Working Paper, No. 2013-01

Provided in Cooperation with:

Department of Economics, Rutgers University

Suggested Citation: Birchby, Jeff; Gigliotti, Gary; Sopher, Barry (2011) : Consistency and aggregation in individual choice under uncertainty, Working Paper, No. 2013-01, Rutgers University, Department of Economics, New Brunswick, NJ

This Version is available at:

<https://hdl.handle.net/10419/94232>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.

Consistency and Aggregation in Individual Choice Under Uncertainty

Jeff Birchby
Ph.D. Student in Economics

Gary Gigliotti
Professor of Economics

Barry Sopher
Professor of Economics

Department of Economics
Rutgers University
New Brunswick, NJ 08901
USA

June 2011

1. Introduction

It is common in studies of individual choice behavior to report averages of the behavior under consideration. In the social sciences the mean is, indeed, often the quantity of interest, but at times focusing on the mean can be misleading. For example, it is well known in labor economics that failure to account for individual differences may lead to incorrect inference about the nature of hazard functions for unemployment duration (e.g., Burdett, Keifer, Mortensen and Neumann (1984) and Keifer (1988)). If all workers have constant hazard functions independent of duration, simple aggregation will nonetheless lead to the inference that the hazard function is state-dependent, with the hazard of leaving unemployment declining with duration of unemployment. Similarly, recent studies in psychology (Gottlieb (2008), Papachristos and Gallistel (2006)) have shown that the “learning curve,” a monotonically increasing function of response to a stimuli, is better understood as an average representation of individual response functions that are, in fact, more step-function-like. As such, the learning curve as commonly understood is a misleading representation of the behavior of any one individual.

These observations motivate us to consider the question of possible aggregation bias in the realm of choice under uncertainty. In particular, Cumulative Prospect Theory (Tversky and Kahneman (1992)) posits a weighting function through which probabilities are transformed into decision weights. An inverted S-shaped weighting function is commonly taken to be “the” appropriate weighting function, based on quite a number of experimental studies. This particular version of the weighting function implies, in simple two outcome lotteries, that an individual will tend to overweight small (near 0) probabilities and to underweight large (near 1) probabilities. A natural question to ask, suggested by both the

hazard function and the learning curve examples, is whether this weighting function is not, similarly, an artifact of aggregation. Of course, no one believes that every individual's behavior can be accounted for by a single weighting function. Studies have shown that there can be considerable variation in estimated weighting functions across individuals. But no one, to our knowledge, has systematically addressed the question of whether, in fact, one can meaningfully use a single weighting function, even as a rhetorical device, to accurately discuss individual choice behavior. If most individuals indeed do have an inverted S-shaped weighting function, then this representation of choice behavior is not misleading, provided it is clear that one is discussing the behavior of "most," not all, individuals.

We focus on the reliability of estimated weighting functions. We study the problem of determining the parameters of the cumulative prospect theory function. Using responses to paired sets of choice questions, it is possible to derive estimates for a two-parameter version of the Cumulative Prospect Theory choice function, using a power function for the value function and a one-parameter version of the weighting function (Prelec (1998)). By analyzing multiple such pairs of choice questions, we are able to also investigate the consistency of these estimates. Our main finding is that there is, in general, considerable variation at the individual level in the choice parameters implied by the responses to the different pairs of choice questions. The modal choice pattern observed is one consistent with expected value maximization, and there is considerably less variation (again, at the individual level) in the parameters implied by those who appear to be maximizing expected value on one pair of choice questions than for those who never choose in this way. But these individuals account for only about one-fifth to one-sixth of subjects. For the rest of the subjects, it is rare that any two pairs of estimates are the same, and often the implied parameters are quite

different. For example, it is not usual for one pair of estimates to imply an inverted S-shaped weighting function, and for another to imply an S-shaped function, for the same individual. We conclude that there is a sizable stochastic element in the typical individual's choice behavior, and that a theory that really describes behavior is going to have to try to explicitly account for this element.

One interesting regularity, which we do not believe has been observed or noted before, is that, in repeated observations on essentially the same choice pairs, the distribution of responses typically does not change, but the composition of the distribution does change. That is, the group of decision makers exhibits a great deal of consistency in repeated observations, but many individuals are not consistent in repeated observations. So, on average (and this is consistent with the history of empirical testing that led to the adoption by many of CPT as the appropriate descriptive model for choice under uncertainty) in repeated sampling on (essentially) the same choice question, a group of subjects will perform the same way—that is, we will arrive at similar estimates of the parameters of the CPT functions. But there is a startling lack of consistency by the typical individual in these repeated samplings. Moreover (and in this we depart from the historical trend mentioned above) we find that parameter estimates at the group level are quite unstable when the underlying choice questions are perturbed only mildly. In particular, modifying lottery choice questions in a way that makes the alternatives less extreme, in probability terms, leads to very different estimates for the probability weighting function.

First, for the aggregate issue (that the CPT parameter estimates are unstable across different types of choice questions), we suggest a non-obvious analogy to the history of the Phillips Curve (Phillips (1958), Blanchard (2000)). Phillips discovered the negative

relationship between the current rate of inflation and the current unemployment rate. At the time, it was thought that governments could direct fiscal and monetary policy so that one could either lower the rate of inflation and pay for this by an increase in the unemployment rate, or, lower the unemployment rate and pay for this by an increase in the inflation rate. Empirical studies indicated, though, that the Phillips relationship was not stationary, but moved in response to the effects of changes in fiscal and monetary policy. The Long Run Phillips curve analysis was an argument that the short run relationship between the rate of inflation and the rate of unemployment was dependent upon the inflationary expectations of the populace. The critical relationship was the ‘natural’ rate of unemployment, dictated by the institutional and technological structure of the economy. When unemployment reached this level, any attempt to lower it would generate additional inflation, and, a readjustment of inflationary expectations, so that the entire short run Phillips relationship would shift vertically, worsening the trade-off between the current rate of inflation and the current rate of unemployment.

We speculate that something similar is occurring when individual weighting functions are estimated. We postulate that each individual has a ‘long-run’ or ‘structural’ attitude towards probabilities. Provocatively, we postulate that this structural attitude of each individual is ‘linear’, i.e., the individual’s structural weighting function is linear. But, in any ‘short-run’ choice scenario, such as those faced by a subject in a choice experiment, the short run weighting function used by the individual is influenced by particular circumstances of the choice problem.

This implies quite a bit about the results of choice experiments with uncertainty. First, on any particular choice any individual could have an s-shaped, an inverted s-shaped or a linear weighting function. Until we know what determines the short run

attitude, the analogue of the shift in inflationary expectations in the long run Phillips curve model, we cannot predict the particular short run weighting function any individual will use. Our empirical study shows that individuals, even in a very simple choice problem, with a one-parameter weighting function (a la Prelec) and a power function as the value function, can exhibit many different weighting functions in a varying group of choice problems. Averaging these functions for individual, or worse, averaging weighting functions across individuals, can only give misleading results about just what is going on.

To understand the problem of choice, we must first admit that expected utility theory is a normative theory for a good reason: it makes sense to anyone with the quantitative skill to analyze the choice problem. But, other factors influence just what an individual does in any particular choice situation. With the Phillips curve, it was the inflationary expectations of agents in the economy that changed to change the position of the short run Phillips curve. In this instance, we postulate that individual expectations about probabilities influence the weighting function, i.e., the parameters of the weighting function are not stable across choice problems.

We only suggest an alternative account of decision making that seems to have some promise in accounting for these observations. The account, rooted in the computational theory of mind, models individuals as choosing according to an ordered hierarchy of violable lexicographic constraints. The basic idea of the CTM is that the mind operates in a distributed fashion (using “modules”) to solve problems. That is, there is not necessarily a single overarching conscious executive decision making authority, as is implicit in expected utility theory or its variants. Instead, the different agents/modules (whichever one is at the forefront at the time the decision must be made) operate on the

question at hand. But is it possible, within this framework, that a person can act as an executive, or chairman, and try to aggregate the different agents preferences into something sensible and coherent that is best for the person (person==society). This is the big question. We plan to explore this idea in a new experiment, which we mention briefly at the end of the paper.

The rest of the paper is organized as follows. Section 2 contains details of the questionnaire and the empirical strategy for estimating the CPT and other parameters. Section 3 contains empirical analysis of the data. Section 4 concludes.

2. Questionnaire and Empirical Strategy

A common strategy in choice experiments is to collect data in order to estimate a choice functional, either parametrically or, with enough data, non-parametrically. Our strategy here is to impose a lot of structure, which reduces the amount of data needed to generate parameter estimates dramatically (2 choice questions is sufficient to get estimates of the parameters of interest.) This highly structured approach seems justified, since those who defend CPT tell us that an inverted S – shaped weighting function is the correct form. But, given the parsimonious structure, it seems wise to take things a step further and get multiple readings on the parameters of the CPT representation for each individual (rather than using lots of data to get one set of parameter estimates). This allows us to judge the stability of the estimates, both within-subjects and between-subjects and to draw conclusions on how consistent individuals are, conditional on the assumption that the general form of preferences is represented by CPT. We also considered (though we do not report those results here) two alternatives that could be thought of as special cases of CPT, expected utility theory and the

dual theory of choice (Yaari (1987), as we can easily generate estimates for these functionals as well.

d

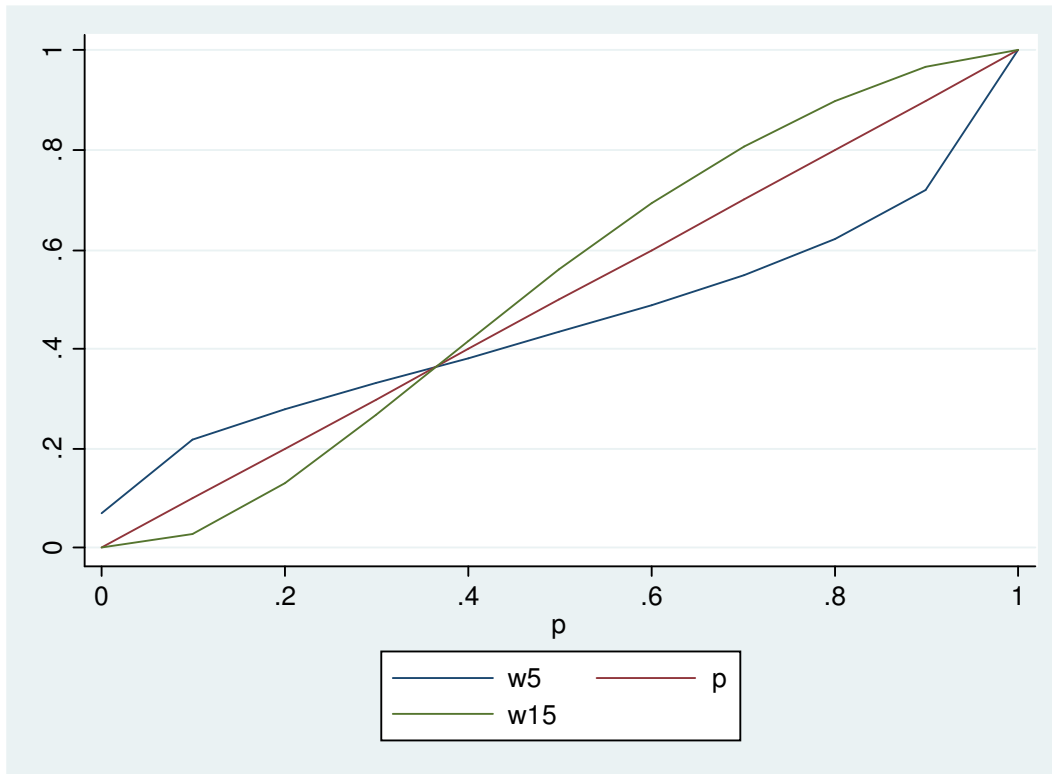


Figure 1: Probability Weighting Functions

The CPT Probability weighting function is given by $w(p) = \exp[-[-\ln(p)]^\alpha]$. This functional form, proposed by Prelec (1998), has the property of lying between zero and 1 and crossing the 45 degree line at $p=.4$. Values of α less than 1 give an inverted S-shaped function (above the 45 degree line for $p<.4$ and below the 45 degree line for $p>.4$), while values of $\alpha > 1$ imply an S-shaped function. A value of $\alpha = 1$ gives a linear weighting function, i.e., expected utility

(EU). Figure 1 illustrates an inverted S-shaped weighting function for $\alpha = .5$, an S-shaped function for $\alpha = 1.5$, and the linear function for $\alpha = 1$.

We specify the CPT Value function by $v(x) = x^\sigma$, a simple power function. Note that a value of $\sigma = 1$ gives a linear value function which, combined with the CPT weighting function, can be interpreted as Yaari's Dual Theory (DT). The CPT representation of utility of lottery $(x, p; y, q)$, $0 < x < y$, is: $U(x, p; y, q) = w(p+q)v(x) + w(q)[v(y) - v(x)]$

Under these assumptions on the function forms of the probability weighting function and the value function, it is possible to derive estimates for the two parameters of interest from only two choice questions. See Tanaka, Camerer and Nguyen (2006) for details. An example will suffice for us. Consider the following question, illustrated in Table 1. A subject is asked to consider the lotteries in the left column (option A) and to compare them to the lotteries in the right column (option B). Option A is unchanging, while Option B improves as one moves down the list (the large prize becomes larger). Subjects are asked to indicate the point at which they would switch from Option A to Option B. The “top” question and “bottom” questions are so-named because they were displayed on a single sheet of paper to the subjects in this fashion, one at the top and one at the bottom. By considering the switching point for each question and equating the value of the CPT functional for Option A and Option B at that point, one can derive values of the parameters α and σ that satisfy both equations.

TABLE 1: Baseline Questions for the Study
“Top” Question

Option A	Option B
3/10 chance at \$4, 7/10 chance at \$1	1/10 chance at \$6.80 9/10 chance at \$.50
3/10 chance at \$4, 7/10 chance at \$1	1/10 chance at \$7.50 9/10 chance at \$.50
3/10 chance at \$4, 7/10 chance at \$1	1/10 chance at \$8.30 9/10 chance at \$.50
3/10 chance at \$4, 7/10 chance at \$1	1/10 chance at \$9.30 9/10 chance at \$.50
3/10 chance at \$4, 7/10 chance at \$1	1/10 chance at \$10.60 9/10 chance at \$.50
3/10 chance at \$4, 7/10 chance at \$1	1/10 chance at \$12.50 9/10 chance at \$.50
3/10 chance at \$4, 7/10 chance at \$1	1/10 chance at \$15.00 9/10 chance at \$.50
3/10 chance at \$4, 7/10 chance at \$1	1/10 chance at \$18.50 9/10 chance at \$.50
3/10 chance at \$4, 7/10 chance at \$1	1/10 chance at \$22.00 9/10 chance at \$.50
3/10 chance at \$4, 7/10 chance at \$1	1/10 chance at \$30.00 9/10 chance at \$.50
3/10 chance at \$4, 7/10 chance at \$1	1/10 chance at \$40.00 9/10 chance at \$.50

“Bottom” Question

9/10 chance at \$4, 1/10 chance at \$3	7/10 chance at \$5.40 3/10 chance at \$.50
9/10 chance at \$4, 1/10 chance at \$3	7/10 chance at \$5.60 3/10 chance at \$.50
9/10 chance at \$4, 1/10 chance at \$3	7/10 chance at \$5.80 3/10 chance at \$.50
9/10 chance at \$4, 1/10 chance at \$3	7/10 chance at \$6.00 3/10 chance at \$.50
9/10 chance at \$4, 1/10 chance at \$3	7/10 chance at \$6.20 3/10 chance at \$.50
9/10 chance at \$4, 1/10 chance at \$3	7/10 chance at \$6.50 3/10 chance at \$.50
9/10 chance at \$4, 1/10 chance at \$3	7/10 chance at \$6.80 3/10 chance at \$.50
9/10 chance at \$4, 1/10 chance at \$3	7/10 chance at \$7.20 3/10 chance at \$.50
9/10 chance at \$4, 1/10 chance at \$3	7/10 chance at \$7.70 3/10 chance at \$.50
9/10 chance at \$4, 1/10 chance at \$3	7/10 chance at \$8.30 3/10 chance at \$.50
9/10 chance at \$4, 1/10 chance at \$3	7/10 chance at \$9.00 3/10 chance at \$.50

The “baseline” pair of questions shown in Table 1 are the basis for the full questionnaire, which consisted of six similar pairs of questions in all. Two more pairs are generated from the baseline simply by doubling, in one case, and halving, in the other case, all of the monetary payoffs in the baseline. Three more pairs of questions are generated from these first three pairs by shifting 2/10 probability weight from the payoffs with the higher probability

weight in each option to the payoff with the smaller probability weight. For example, Option A, “top” question in Table 1, so transformed has 5/10 weight on \$4 and \$1, and Option B, in the first line, has 3/10 weight on \$6.80 and 7/10 weight on \$0.50. The other questions are transformed in a similar fashion.

With these six pairs of questions we are thus able to generate multiple readings on the two parameters of interest for each individual. Each of the six “top” questions can be combined with each of the six “bottom” questions to generate new estimates of the parameters α and σ , which generates 36 (α, σ) combinations for each individual subject. This allows us to do econometric analysis on the variability of the parameters, rather than using all of the observations to generate a single estimate of each parameter for each individual. This schema is illustrated in Table 2. The lower right corner represents the typical approach of reporting a single set of parameter estimates for a group of subjects. The large central cell and the cells immediately to the right and below will be the focus of our analysis.

	TABLE 2: Dimensions of the Parameters Estimated	
	Choice Pair 1 Choice Pair K	
Observations on Subject 1 to Subject N	α_{11}, σ_{11} α_{1k}, σ_{1k} α_{21}, σ_{21} α_{2k}, σ_{2k} α_{n1}, σ_{n1} α_{nk}, σ_{nk} (Individual parameter estimates)	Average parameters over K Choice pairs by subject
	Average parameters over N subjects by choice pair	Average over choice pairs and subjects

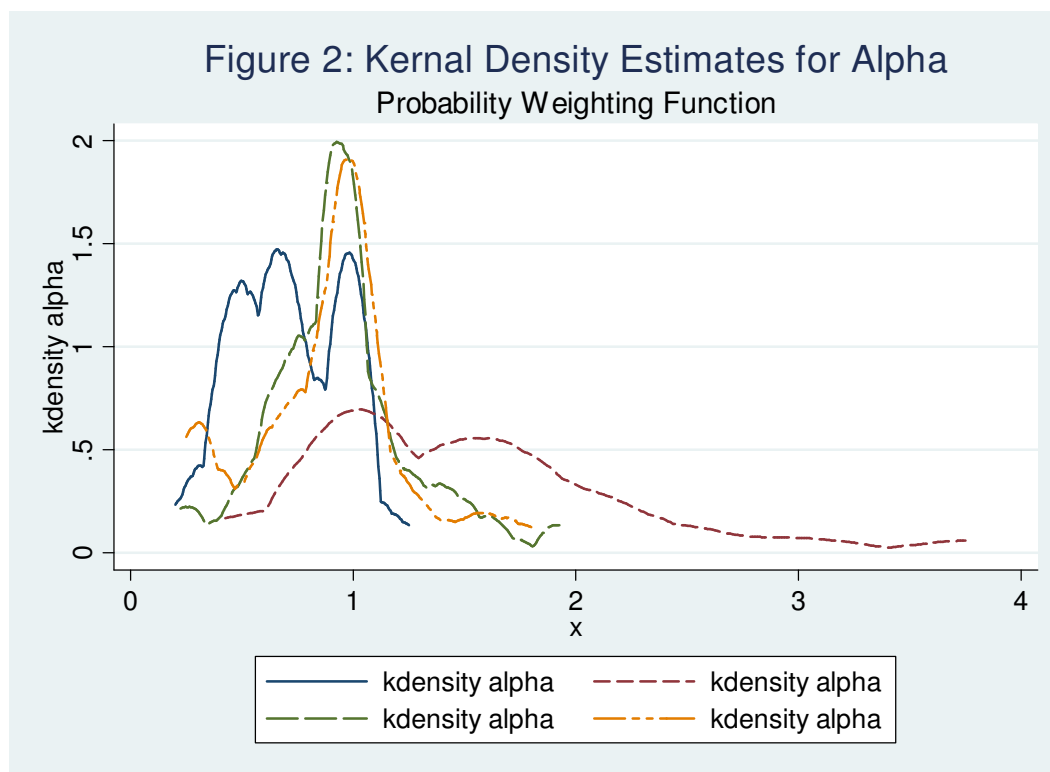
3. Empirical Analysis

Table 3 contains the main findings, based on administering the questionnaire to 121 student subjects at Rutgers University in the Spring of 2006. Subjects answered all 12 choice questions and then one of the lotteries that they chose was randomly played out, and their earnings were based on the outcome of the lottery, plus a \$5 show up fee. Earnings ranged from \$5 to \$65, with an average of about \$11 for a 30 minutes session.

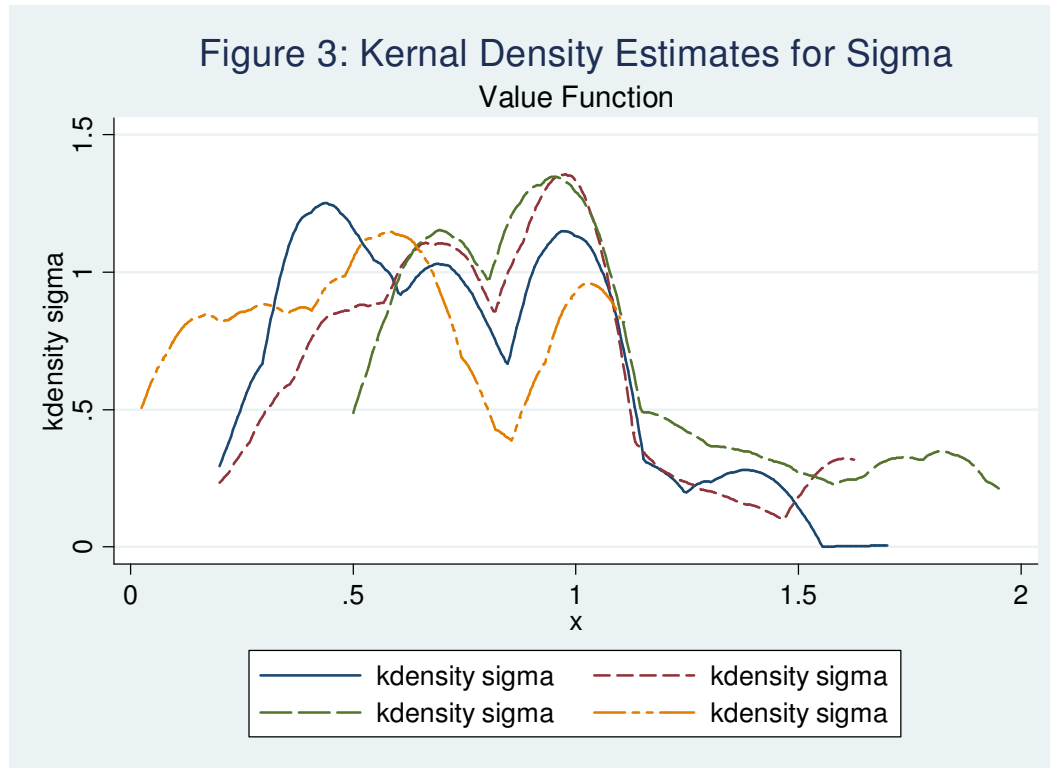
Overall, using all 36 possible choice pairs, the average (over all 121 subjects in our experiment), shown in row (i), for α is 1.01 and for σ is .80. That is, the average implies essentially a linear probability weighting function (no distortion at all) and a degree of curvature in the value function comparable to that found in other studies. Line (ii) shows estimates based only on the subset of choice pairs that correspond to those used in Tanaka, Camerer and Nguyen (2006), and our findings ($\alpha = .71$ and $\sigma = .73$) are comparable to their findings of an average α of about .60 and an average σ of about .70. The other two rows show (iii) estimates derived from “top” questions (see Table 1) comparable to those used by Tanaka, et al. and transformed “bottom” questions and (iv) the reverse.

TABLE 3: Estimates of α and σ		
Questions used for estimates	α (std. dev.)	σ (std. dev.)
(i) All choice pairs (36)	1.01 (.54)	.80 (.38)
(ii) Choice pairs based on Tanaka, Camerer and Nguyen only (9 Pairs)	.71 (.25)	.73 (.31)
(iii) Transformed choice pairs only (9 pairs)	1.51 (.72)	.81 (.35)
(iv) Mixtures of “top” questions from (ii) and “bottom” questions from (iii) (9 pairs)	1.21 (.62)	.93 (.38)
(v) Mixtures of “top” questions from (iii) and “bottom” questions from (ii) (9 pairs)	.88 (.35)	.57 (.33)

What is going on here? Clearly the shift to less extreme lottery choice questions is the source of at least part of the instability in the parameters, particularly of the parameter on the probability weighting function. Figures 2 and 3 provide more information in displaying kernel density estimates of the underlying distributions from which the parameters are drawn.



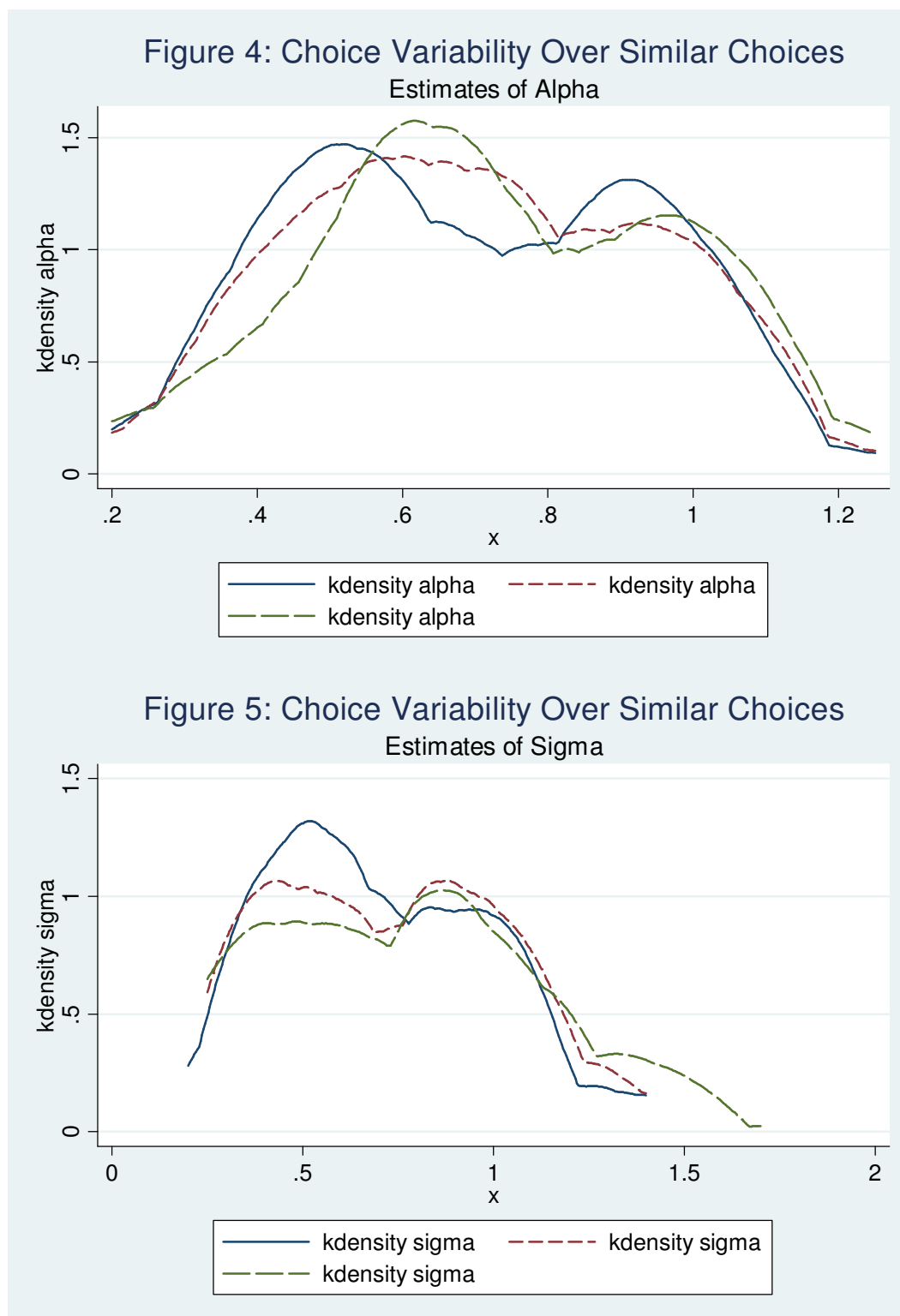
The solid, short dashed, long dashed, and long-short dashed lines in Figures 2 and 3 correspond to the sets of questions listed in (i) – (iv) (**or is it (ii) to (v)?**), respectively, of Table 3. Note that in Figure 2 all of the estimated density functions for α have a mode at 1, but the density for the (ii) curve, corresponding to the transformed choice questions, has a second mode to the right and a very long right tail. In Figure 3 we see that the estimated densities of the value function parameter σ are more similar to one another, but all have multiple modes. There is always a mode at a value of 1, and a mode at a lower value as well.



Given the instability of the parameter estimates for the probability weighting function, it is clearly inappropriate to speak of “the” value of α that determines the function without reference to a particular space of lotteries. It is interesting that previous studies have so consistently come up with a parameter value less than 1, implying an inverted S-shaped weighting function, when it was so easy for us to find dramatically different estimates for a simple shift in the probabilities in the choice questions to less extreme values. One possibility, consistent with our findings here, is that the weighting function is more of a concave shape above the 45 degree line. Non-parametric estimates of the weighting function in Sopher, Gigliotti and Klein (2002) tend to imply such a shape. If this turns out to be true, then the wide variance in the estimates found here would be due to the extreme structure imposed by the single-parameter weighting function assumed. In particular, to the extent that probabilities in the neighborhood of .5 are not underweighted, the parameter estimate is going to tend to be

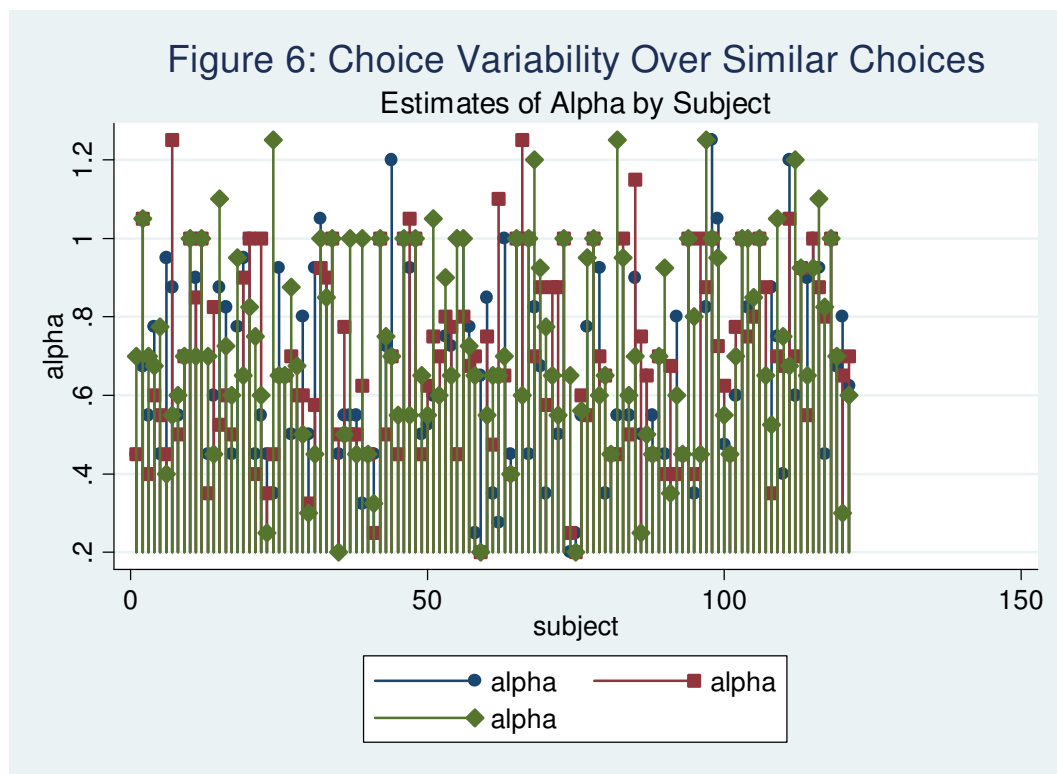
larger than 1, as this yields an S-shaped function that is above the 45 degree line in the neighborhood of $p=.5$.

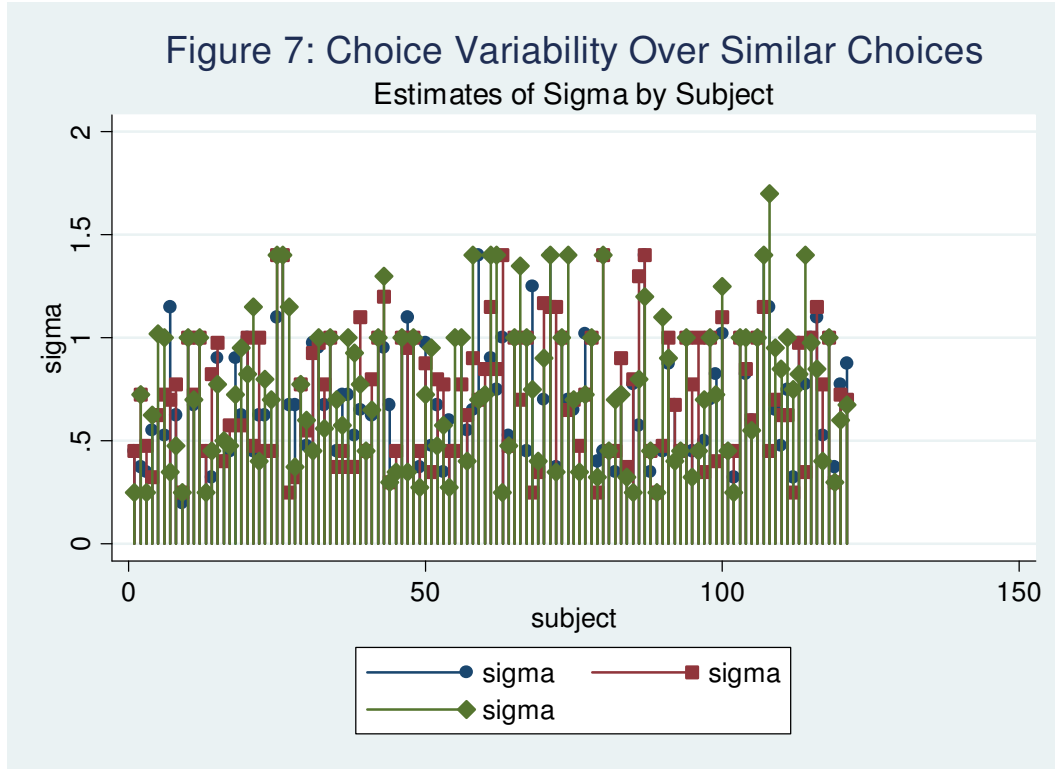
We now briefly consider individual heterogeneity in the estimates. Figures 4 and 5 are illustrative of a general finding in our data, that individual estimated parameters are more variable than one might expect on repetitions of (essentially) the same choice questions.



The figures show kernel density estimates for parameters derived from rescaled versions of a given choice pair—thus, the choices are essentially the same question, only slightly disguised.

(Solid line is original question, short dashed is with prizes doubled, long dashed is with prizes halved.) Although statistical tests may not reject the hypothesis that the distributions are the same overall, there is a surprising amount of variability at the individual level, given the simplicity of the questions asked. Figures 6 and 7 show the same data as in Figures 4 and 5, but organized by subject, so that one can see the variability for individuals (blue, red and green are for original, doubled, and halved prizes). The vertical length of the lines shows the extent of the variability. We plan to address the sources of this heterogeneity in a new experiment motivated by the discussion of the Computational Theory of Mind alluded to in the Introduction.





4. Conclusions

We have conducted a simple experiment that has a rather simple but surprising result. Without any maneuvering or scheming, we found it very easy, making assumptions about the structure of the CPT choice functional that are widely regarded as acceptable among a large set of researchers who subscribe to CPT as a descriptive theory of choice under uncertainty, to demonstrate a large instability in the estimated value of the parameter that determines the shape of the probability weighting α function. By simply shifting probability weights so as to make the lotteries in the choice questions less extreme, in probability terms, we have found that the average estimated value of α is significantly greater than 1, implying an S-shaped weighting function rather than the inverted S-shape normally assumed to be the case.

The average value of α over all choice pairs is essentially 1, implying a perfectly linear probability weighting function, as in expected utility theory. We believe, however, that

reporting averages alone is misleading, and suggest that further analysis of the root causes of this instability, both at the group level and particularly at the individual level is warranted.

5. References

Blanchard, Olivier (2000). *Macroeconomics*, Second edition, Prentice Hall.

Burdett, K., N.M. Kiefer, D.T. Mortensen, and G.R. Neumann (1984). "Earnings, Unemployment and the Allocation of Time over Time." *Review of Economic Studies*, 51(4), 559-578.

Gigerenzer, G. and D. Goldstein (1996). "Reasoning the Fast and Frugal Way: Models of Bounded Rationality." *Psychological Review*, 103(4), 650-669.

Gottlieb, D. (2008). "Is the number of trials a primary determinant of conditioned responding?" *Journal of Experimental Psychology: Animal Behavior Processes*. 34(2), 185-201.

Kiefer, N.M. (1988). "Economic Duration Data and Hazard Functions." *Review of Economic Literature*, 26(2), 646-679.

Papachristos, E.B. and C.R. Gallistel (2006). "Autoshaped head poking in the mouse: a quantitative analysis of the learning curve." *Journal of Experimental Analysis of Behavior*, 85(3), 293-308.

Phillips, A W (1958). "The Relationship between Unemployment and the Rate of Change of Money Wages in the United Kingdom 1861-1957". *Economica* **25** (100): 283-299.

Prelec, D. (1998). "The Probability Weighting Function." *Econometrica*, 66, 497-527.

Quiqin, J.(1982). "Theory of Anticipated Utility," *Journal of Economic Behavior and Organization*, 3, 323-343.

Sopher, B., G. Gigliotti, and R. Klein (2002) "Testing for Stability of the Transformation Function in Rank-Dependent Expected Utility Theory," mimeo, Rutgers University.

Tanaka, T., C. Camerer and Q. Nguyen (2006). "Poverty, Politics and Preferences: Field Experiments and Survey Data from Vietnam." Mimeo, California Institute of Technology.

Tversky, Amos; Daniel Kahneman (1992). "[Advances in prospect theory: Cumulative representation of uncertainty](#)". *Journal of Risk and Uncertainty* **5**, 297-323.

Wu, G. and R. Gonzalez (1995). "Curvature of the Probability Weighting Function," *Management Science*.

Yaari, M.E. (1987). "The Dual Theory of Choice Under Risk." *Econometrica*, 55(1), 95-1115.