



JENA ECONOMIC RESEARCH PAPERS



2010 – 042

The econometric modeling of social Preferences

by

**Anna Conte
Peter G. Moffatt**

www.jenecon.de

ISSN 1864-7057

The JENA ECONOMIC RESEARCH PAPERS is a joint publication of the Friedrich Schiller University and the Max Planck Institute of Economics, Jena, Germany. For editorial correspondence please contact markus.pasche@uni-jena.de.

Impressum:

Friedrich Schiller University Jena
Carl-Zeiss-Str. 3
D-07743 Jena
www.uni-jena.de

Max Planck Institute of Economics
Kahlaische Str. 10
D-07745 Jena
www.econ.mpg.de

© by the author.

THE ECONOMETRIC MODELLING OF SOCIAL PREFERENCES

Anna Conte

Strategic Interaction Group, Max-Planck-Institut für Okonomik, Jena, Germany

aconte@econ.mpg.de

Peter G. Moffatt

School of Economics, University of East Anglia, Norwich, UK

p.moffatt@uea.ac.uk

ABSTRACT

Experimental data on social preferences present a number of features that need to be incorporated in econometric modelling. We explore a variety of econometric modelling approaches to the analysis of such data. The approaches under consideration are: the random utility approach (in which it is assumed that each possible action yields a utility with a deterministic and a stochastic component, and that the individual selects the action yielding the highest utility); the random behavioural approach (which assumes that the individual computes the maximum of a deterministic utility function, and that computational error causes their observed behaviour to depart stochastically from this optimum); and the random preference approach (in which all variation in behaviour is attributed to stochastic variation in the parameters of the deterministic component of utility). These approaches are applied in various ways to an experiment on fairness conducted by Cappelen et al. (2007). At least two of the models that we estimate succeed in capturing the key features of the data set.

Keywords: Econometric modelling and estimation; model evaluation; individual behaviour; fairness

JEL codes: C51; C52; C91; D63

Data

The data set used in this paper was downloaded from the AER website:

http://www.aeaweb.org/articles/issue_detail_datasets.php?journal=AER&volume=97&issue=3&issue_date=June%202007

Acknowledgements

We thank A. Cappelen, A. Hole, E. Sorensen and B. Tungodden for detailed comments on an earlier version of this paper. That feedback led to significant improvements in the paper.

1. Introduction

Clearly a variety of possible approaches are possible in the econometric modelling of experimental data. In deciding what sort of approach is most suitable for a particular application, the most obvious criteria are: consistency with the economic theoretic model that is being posited; and the ability to explain all features of the data. Often these two criteria are in conflict: a model which is fully consistent with the theory often fails to account for data features; a model which is developed with the objective of explaining all data features is unlikely to be consistent with theory.

This conflict is particularly marked in the context of experiments in which social preferences are somehow being elicited. Examples are public goods games, in which subjects decide how much to contribute to a social fund, and bargaining games in which subjects decide how much of an endowment to keep for themselves, or equivalently how much to donate to an opponent. The amount contributed or claimed by the subject is the dependent variable under analysis. A theory, or a combination of different theories, based on received behavioural norms, may then be applied to such data. The problem is that some features of the data cannot be explained by such theories alone. One straightforward example is the tendency in public goods games for an individual to contribute exactly half of their endowment to the public fund. According to Bardsley and Moffatt (2002, p.179), “It is hard to envisage a model of rational decision making which would accommodate [over-representation of 50% contributions]”.

Rather than attempting to extend theories in ways that seem unnatural, it is preferable for this type of feature of the data to be incorporated somehow into the model’s stochastic specification. We are proposing that theory and the stochastic specification are combined so as to “meet half-way”, with the theory being called upon to explain the broad patterns of the data (and if it fails to do even this, we would surely discard it and search for a new theory), and the finer intricacies of the data being left to the stochastic component.

In attempting to elicit the social preferences of experimental subjects, two types of experimental design are possible. The first is one in which the subject is presented with a finite array of possible allocations from which they are invited to select one. This type of design has been used in the context of the ultimatum game by Bellemare et al. (2008), and in the context of a public goods game by Conte and Levati (2010). The second approach to design is one in which subjects are asked to state their preferred allocation in an open-ended protocol. This sort

of design has been used by Bardsley (2000) in the context of a public goods game, and by Cappelen et al. (2007) in the context of a Fairness experiment.

The distinction between these two types of experimental design is important because they lend themselves to different types of econometric model. In the first case, subjects have been asked to choose from a finite set of alternatives each with its own “attributes” (i.e. the allocation), and it is natural to apply a discrete choice model, on the assumption that the utility derived from each alternative contains a random component (the random utility assumption). For example, Bellemare et al. (2008) applied an econometric model based on the mixed logit approach to data on the proposer’s decision in their ultimatum game; Costa-Gomes and Weizsäcker (2008) applied a similar approach to players’ actions in normal-form games.

In the second case, in which an open-ended response has been elicited, it is natural to assume that the decision variable is continuous, and to proceed accordingly. For example Bardsley and Moffatt’s (2007) model of voluntary contributions is built around a linear regression equation with subject’s contribution as the dependent variable. In this second case, there are two possible assumptions concerning the source of the continuous variation: a continuously-distributed deviation from optimising behaviour (the random behavioural assumption); and variation in the optimal decision itself (the random preference assumption).

The principal objective of this paper is to demonstrate these various approaches to estimation in the context of the fairness experiment conducted by Cappelen, Hole, Sorensen and Tungodden (2007), henceforth CHST. The objectives of this experiment were to determine what sorts of fairness ideals are brought into play by individuals in deciding on income allocations between themselves and others, and to determine the importance that individuals attach to such fairness ideals in relation to the motivation of self-interest. We use many of the features of the CHST approach (e.g. assuming a discrete mixture of “types”, and assuming heterogeneity in certain parameters) because we view these features as innovative and in step with recent developments in the literature (Bardsley and Moffatt, 2007; Conte et al., 2009; Botti et al., 2008; Harrison and Rutstrom, 2009).

Because CHST use an open-ended protocol in their elicitation of fairness ideals, and also because their theoretical model results in a continuous decision variable, it is our opinion that either the random behavioural or the random preference approach should be followed, and not the random utility approach. In fact, CHST themselves used the random utility approach. A

secondary objective of this paper is therefore to settle the issue of which of the approaches is best able to explain the data in this particular case.

A plan of the remainder of the paper is as follows. In Section 2 we summarise the theoretical model of fairness developed by CHST. In Section 3 we provide a conceptual outline of the three possible econometric approaches: Random Utility, Random Behavioural and Random Preference. In Section 4, we draw attention to the particular features of the CHST data set that the stochastic component of each model will be called upon to explain. In Section 5, we develop five different econometric models, each of which follows one of the three approaches outlined in Section 3. We also explain how each model is estimated and comment on the estimates results. We also investigate the predictive performance of each model, in terms of their ability to explain the key features of the data identified in Section 4. In Section 6, we discuss mixing proportions and derive posterior probabilities of types. Section 7 concludes.

2. A Theoretical Model of Fairness

The model summarised here is the same as that of CHST. Here we draw out only those features of the model that are essential for understanding the econometric models constructed in later sections of the paper.

Consider a game involving 2 players ($i = 1, 2$). The game consists of two phases: the production phase and the distribution phase. In the first (production) phase, subject i ($i = 1, 2$) decides how much of her initial endowment will be “invested” (q_i); this investment is then multiplied by the individual’s exogenously assigned “rate of return” (a_i) in order to determine the income that the individual generates ($a_i q_i$). The incomes generated by the two individuals are added together to give total income:

$$X(\mathbf{a}, \mathbf{q}) = a_1 q_1 + a_2 q_2. \quad (1)$$

In the second (distribution) phase, each player takes the role of Proposer in a dictator game: each decides, on the basis of the phase-1 investments and incomes of themselves and of the other individual, how much of the total income they would like to allocate to themselves (y), and how much they would like to leave for their opponent ($X-y$). It is then randomly determined which of the two proposed allocations is implemented. Under reasonable assumptions, we may safely assume that each player has an incentive to reveal their utility maximising allocation of income.

What is an individual's utility maximising allocation? Following CHST, it is assumed that individuals are motivated by income, and also that they are motivated by fairness. Each individual has a "fairness ideal", $m(\mathbf{a}, \mathbf{q})$, defined as the income they would allocate to themselves that they consider to be perfectly "fair". It is then assumed that utility of individual i is given by:

$$V_i(y; \mathbf{a}, \mathbf{q}) = \gamma y - \beta \frac{[y - m(\mathbf{a}, \mathbf{q})]^2}{2X(\mathbf{a}, \mathbf{q})}, \quad (2)$$

where the parameters $\gamma > 0$ and $\beta \geq 0$ respectively represent the importance the individual places on income and fairness considerations.

An individual maximising (2) will choose to allocate the following income to herself:

$$y^* = m(\mathbf{a}, \mathbf{q}) + \frac{\gamma}{\beta} X(\mathbf{a}, \mathbf{q}) = m(\mathbf{a}, \mathbf{q}) + \theta X(\mathbf{a}, \mathbf{q}). \quad (3)$$

Note that, in (3), we have introduced a new parameter θ as the ratio of the two parameters of the utility function (2). We shall refer to the term $\theta X(\mathbf{a}, \mathbf{q})$ as the "selfishness-premium" of an individual, since it is the amount over and above the fairness ideal that the individual chooses to keep for herself.

The next question arising is what is the "fairness ideal". Here it is assumed that there are three different "types" of individual in the population, each with a different rule for computing their fairness ideal. The first type is the *Strict Egalitarian*, who considers the fairest allocation to be one in which the total is divided equally between the two individuals, regardless of how the total has been determined. This type is defined as follows:

$$\text{Strict Egalitarian (SE; Type 1): } m_1(\mathbf{a}, \mathbf{q}) = \frac{X(\mathbf{a}, \mathbf{q})}{2}.$$

The second type is the *Liberal Egalitarian*, whose ideal is for each individual to receive an amount proportional to the amount they themselves have invested. This type is defined as:

$$\text{Liberal Egalitarian (LE; Type 2): } m_2(\mathbf{a}, \mathbf{q}) = \frac{q_1}{q_1 + q_2} X(\mathbf{a}, \mathbf{q}).$$

The third type is the *Libertarian*, whose ideal is for each individual to receive an amount equal to their own contribution. This type is defined as:

$$\text{Libertarian (L; Type 3): } m_3(\mathbf{a}, \mathbf{q}) = a_1 q_1.$$

3. Econometric Modelling Approaches

In this Section, we draw an important distinction between three econometric approaches: the *Random Utility* (RU) model; the *Random Behavioural* (RB) model; and the *Random Preference* (RP) model.

The *random utility* approach consists of estimating the parameters of the utility function (2), using a choice model, in which it is assumed that a subject selects their allocation from a discrete set of alternatives. It is assumed that the utility of each alternative contains a random error term, and that the subject chooses the one with the highest utility. The essence of the approach can be seen in Figure 1, in which the “self-allocation” is measured on the horizontal axis, and utility on the vertical. The curve represents the deterministic component of utility (2). This curve is inverted U-shaped as a consequence of the assumption of the existence of a fairness ideal: utility rises until a certain allocation (y^*) is reached, but then falls as a consequence of the individual’s fairness principles outweighing their desire for higher income. Note that y^* will, in accordance with (3), be somewhat to the right of the individual’s “fairness ideal”, this distance representing the “selfishness-premium”.

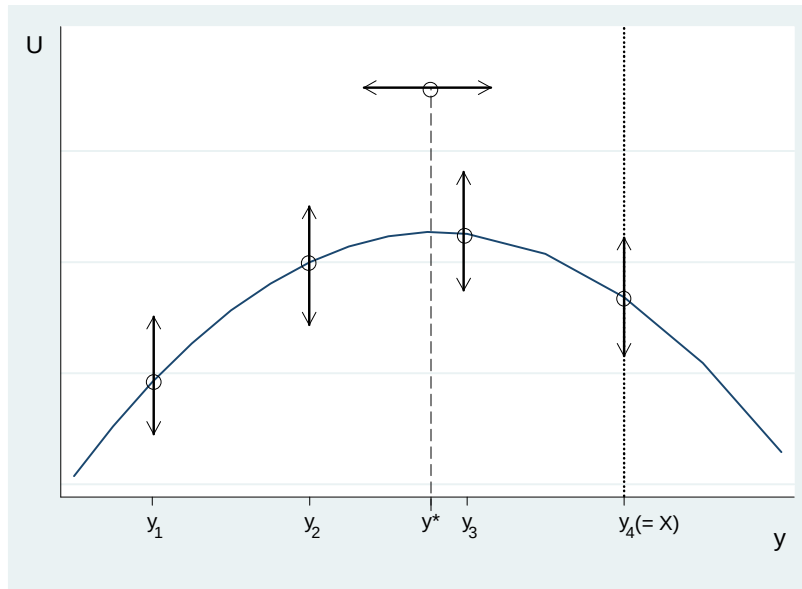


Figure 1: A utility function over income allocations. The utility-maximising allocation is y^* . Total income is X .

In the example presented in Figure 1, there are four permitted allocations, $y_1 < y_2 < y_3 < y_4$, the largest of which, y_4 , is equal to total income (X). It is assumed that the utility at each of these allocations is the sum of the deterministic component indicated by the

vertical position on the curve, and an i.i.d. random term. The random term is represented loosely in Figure 1 by the vertical arrows. Finally, it is assumed that the allocation with the highest utility is chosen. In the context of the example presented in Figure 1, it is clear that, although allocation y_3 is the most likely to be chosen, being the one that is closest to y^* , any of the four allocations could actually be chosen, since the choice depends also on the realisations of the four random components.

It should be noted that the RU approach is a version of what is known as the Fechner (1860/1966) model, in the sense that stochastic terms are being applied additively to the utilities on whose comparison the individual's decision is based.

The *random behavioural* approach consists of modelling the behavioural equation (3) directly, instead of the utility function (2). Equation (3) indicates the position of the optimal allocation y^* in Figure 1. It is assumed that the actual allocation is this optimal allocation plus a random error term with a continuous distribution. This random term is represented loosely by the *horizontal* arrow in Figure 1. An implication of this assumption is that the chosen allocation can, in theory, be *any real number* between zero and X ; it is not restricted to being one of a discrete set of values.

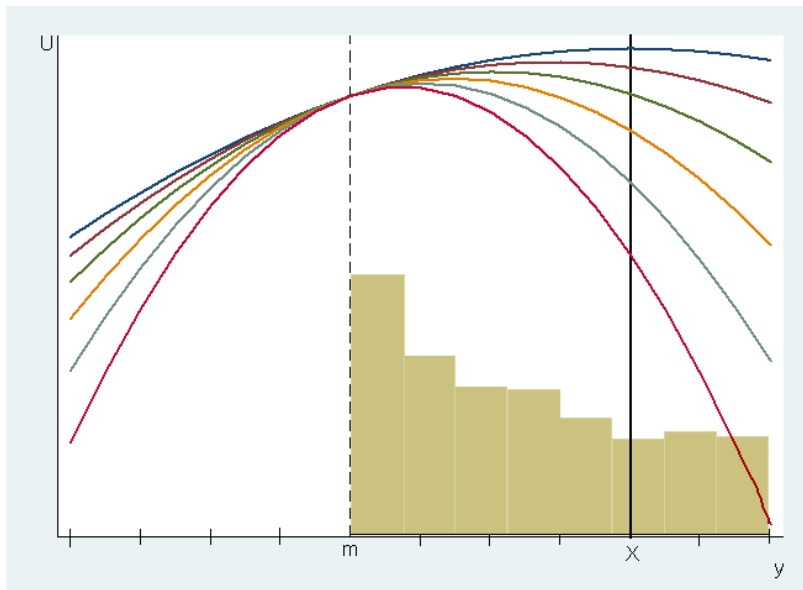


Figure 2: Utility functions under the Random Preference (RP) assumption. m is the fairness ideal; X is total income. The curves represent utility functions (2) with various values of β (hence of θ). The distribution of (latent) utility-maximising allocations y^* resulting from a log-normal distribution for β (hence of θ) is superimposed.

The *random preference* approach is built on the premise that all variation in behaviour is explained in terms of stochastic variation in the parameters of (2) or (3). Figure 2 illustrates the situation in which it is assumed that the parameter β in (2), or equivalently θ in (3), varies randomly. Note that the position of the utility function, and hence the position of the utility-maximising allocation (y^*), are varying, with y^* moving to the right as β falls, although note that y^* is always to the right of the fairness ideal (m). The superimposed histogram shows the distribution of the (latent) utility-maximising allocation resulting from a simulation in which β follows a lognormal distribution. This allocation is “latent” in the sense that it sometimes exceeds the maximum permitted allocation (X).

Although the RP approach has the important advantages of theoretical consistency and intuitive elegance, there are situations in which it breaks down, in the sense of being unable to explain particular data patterns. For example, if a subject is observed claiming an allocation that is lower than the lowest of all of their possible fairness ideals (e.g. to the left of m in Figure 2), the RP model cannot account for this. This sort of problem has been encountered in other applications of the RP model: Loomes et al. (2002), in their econometric analysis of risky choices, note that the RP model breaks down whenever a subject chooses a stochastically dominated alternative. These considerations usually lead to the recommendation that the RP assumption is used in conjunction with some other, more ad-hoc, stochastic component.

4. Issues with the CHST data set

The experiment conducted by CHST is an exact representation of the 2-player game described in Section 2. There were 96 subjects, of whom nearly all played the game twice, with different opponents. The amount jointly earned in phase 1 varies between NOK400 and NOK1600. The amount of this that is claimed by each player is the focus of the analysis.

Here, we identify the distinctive features of CHST’s data. These features are of great importance in guiding the choice of econometric specification.

Figure 3 shows the distribution, over all 190 observations, of the proportion of total income that the individual chooses to keep for themselves. There is undoubtedly a strong element of discreteness in the data. There are prominent modes at 50% of the endowment, and at 100% of the endowment. More than half (58%) of the data points are at one of these two “focal points”. However, the economic model proposed by CHST, outlined in Section 2 of this paper,

dictates that the dependent variable (the amount of the endowment claimed by the individual) has a continuous distribution, depending partly on a continuously distributed variable which we label the “selfishness premium”. It is therefore important to consider why such a high proportion of the observations draw together at these two points, and then to incorporate the likely explanations into the econometric modelling strategy.

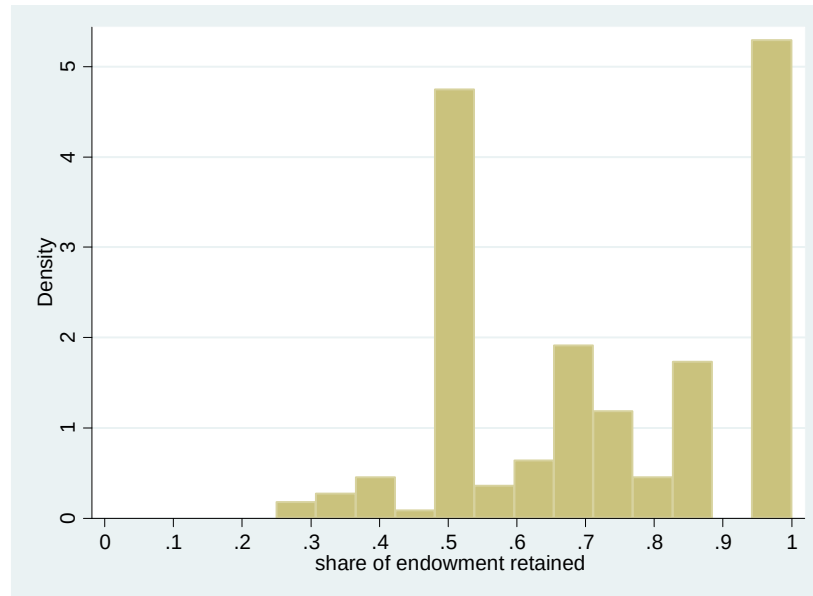


Figure 3: Distribution (over all 190 tasks) of amount retained as proportion of endowment.

The most likely reason for the probability mass at 50% of the endowment is that this proportion of the allocation corresponds exactly to at least one of the “fairness ideals”. Given this, if we assume that individuals might behave exactly in accordance with their fairness ideal (i.e. with a selfishness premium of zero), we are able to explain this probability mass.

The probability mass at 100% of the endowment is *not* attributable to a fairness ideal. This is because no reasonable notion of “fairness” could dictate that an individual takes all of the available endowment for themselves. These observations are the result of selfish behaviour, and the reason for the probability mass is simply that 100% of the endowment is the upper limit to the allocation. Therefore we attribute this probability mass to upper censoring of the allocation variable. Upper censoring is therefore another necessary feature of any econometric model purporting to explain this data.

5. Econometric Specifications

In this section, we present five econometric specifications, each following one of the three modelling approaches explained in section 3. In conformity with CHST, the econometric models are all based on the assumption that individuals are of different types, but that they cannot switch type between tasks. An individual's type determines uniquely that individual's fairness ideal for a given task. However, the actual allocation decision made in any task may or may not coincide with the fairness ideal. This is one of the aspects that differs between the models.

The available data is an unbalanced panel, since most (but not all) subjects engage in more than one task¹. We therefore define $\mathbf{a}_{it}$ and $\mathbf{q}_{it}$ respectively to be the vectors $\mathbf{a}$ and $\mathbf{q}$ that apply to subject i in task t , $t=1, \dots, T_i$; $i=1, \dots, n$. We further define $m_{k,it} = m_k(\mathbf{a}_{it}, \mathbf{q}_{it})$; $X_{it} = X(\mathbf{a}_{it}, \mathbf{q}_{it})$. For sake of comparison, we present the estimation results from all the models in Table 1. In each column the parameter estimates of one single model are reported.²

In order to assess the relative performance of each model, we investigate how accurately each model is able to predict the two key features of the data that were identified in Figure 3.³ These features were: the high proportion of observations (27.37%) at exactly half of the endowment; and the similarly high proportion of observations (30.53%) at the whole of the endowment. To this end, we use the estimated model to simulate a large number (10,000) of data sets with the same sample dimensions, and the same explanatory variable values, as the actual data. We then examine the distribution of the relevant proportions over the simulated samples. If the actual value of the relevant proportion lies somewhere in the main part of this simulated distribution, we may infer that the model predicts the proportion well; if the actual value is in the tails, we infer a poor predictive performance.

¹ Of the 96 subjects, 94 performed two tasks, and two performed only one task.

² All models, except the RU model (whose estimation results are taken directly from the CHST paper), are estimated in STATA version 11.0. The programs are available from the authors on request.

³ It should be made clear that the structure of the models estimated in the paper is such that the maximised log-likelihoods are not comparable between most pairs of models. This is one of our reasons for focusing instead on the models' predictive performance.

With this in mind, the discussion of each model will be followed by a reference to Figure 4, which consists of five rows and three columns: each row corresponds to one of the models we estimate; the first two columns report, respectively, the histograms of the proportion of observation at 50% of the endowment and at 100% of the endowment in the simulated samples; the vertical lines represent the actual proportion observed in the real data. The third column, perhaps more informatively, provides a means of jointly assessing the prediction of both features. Each point in these scatter plots represents a simulated data set, with the proportion of observations at 50% of the endowment measured on the horizontal axis, and the proportion at 100% on the vertical. The true proportions are represented by the two straight lines in each plot. A good predictive performance would be indicated by a scatter which is roughly centred on the intersection of the two lines. 2-tailed p -values of the univariate test and the joint test are reported.⁴

In the remainder of this section, we will discuss five econometric models, presenting and also discussing the estimates' results and the tests we use to assess their relative performance.

5.1 *The Random Utility model*

This model is centred around the RU approach discussed Section 3. The assumptions of the model are as follows. Each subject i , $i = 1, \dots, n$, draws a value β_i for β in (2), from a log-normal distribution, and this value applies to all tasks, $t = 1, \dots, T_i$. This determines the deterministic component of utility. An error term, independent between alternatives and between tasks, is added to the utility of each alternative. The alternative with the highest resulting utility is chosen.

These assumptions in combination with the utility function defined in eq. 2 give rise to the model:

$$U_i(y_{jit}; \mathbf{a}_{it}, \mathbf{q}_{it}) = V_i(y_{jit}; \mathbf{a}_{it}, \mathbf{q}_{it}) + \varepsilon_{iy} = \gamma y_{jit} - \beta_i \frac{[y_{jit} - m(\mathbf{a}_{it}, \mathbf{q}_{it})]^2}{2X(\mathbf{a}_{it}, \mathbf{q}_{it})} + \varepsilon_{iy} \quad (4)$$

$$\log(\beta_i) \sim N(\zeta, \sigma^2)$$

⁴ The p -value for the joint test is constructed as follows. The two straight lines divide the simulated data into four quadrants. We consider the proportion of the data falling into each quadrant, and ask which of these four proportions is smallest. We then multiply this proportion by 4 to obtain the p -value for the test. This p -value may be interpreted in the usual way, with a value less than 0.05 representing evidence against the specification under test.

In task t , subject i is confronted with the $s_{it} + 1$ alternatives $j_{it} \in \{0, 1, \dots, s_{it}\}$, with $s_{it} = X(\mathbf{a}_{it}, \mathbf{q}_{it})/50NOK$; choosing alternative j_{it} results in a self-allocation of $y_{j_{it}} = j_{it} \times 50NOK$.⁵ The i.i.d. error term ε_{iy} is taken to follow a Type I extreme value distribution, with the consequence that across the alternatives the difference between any two ε_{iy} is distributed logistic. Subject i in task t chooses the alternative that maximises (4).

The likelihood contribution of subject i choosing alternative j_{it} is:

$$L_{i_RU}(\lambda_1, \lambda_2, \lambda_3, \zeta, \sigma) = \sum_{k=1}^3 \lambda_k \int_0^\infty \left\{ \prod_{t=1}^{T_i} \frac{\exp[U_i(y_{j_{it}}; \mathbf{a}_{it}, \mathbf{q}_{it}, m_{k,it})]}{\sum_{j_{it} \in \{0,1,\dots,s_{it}\}} \exp[U_i(y_{j_{it}}; \mathbf{a}_{it}, \mathbf{q}_{it}, m_k)]} \right\} f(\beta; \zeta, \sigma) d\beta, \quad (5)$$

where $f(\beta; \zeta, \sigma)$ is the density function of the lognormal distribution evaluated at β , with ζ and σ being the parameters of the underlying normal distribution. Moreover, $m_{k,it}$, with $k = 1, 2, 3$, are the *fairness ideals* and the three parameters λ_k , with $k = 1, 2, 3$, are the *mixing proportions*, representing the proportions of the population who are respectively Strict Egalitarian, Liberal Egalitarian, and Libertarian, as defined at the end of Section 2.

(5) is the model used by CHST. In column 1 of Table 1, we report the estimation results for this model presented in their paper. The estimated mixing proportions show that the Strict Egalitarian type is the most common in the population (0.435), followed by the Liberal Egalitarian type (0.381) and finally the Libertarian type (0.184).

To assess the performance of this model at reproducing the peculiarities of the data set under analysis, in the first row of Figure 4 the simulated distribution for the RU model are reported. The first histogram shows the distribution of the proportion at one half of the endowment; the second shows that of the proportion at the whole endowment. We see that this model successfully predicts the proportion at half, but appears to under-predict the proportion at the full endowment. In fact, 99.58% of the simulated samples have a proportion at the full endowment which is less than the actual proportion of 30.53%. This translates into a 2-tailed p -value of 0.0084, representing strong evidence of prediction bias.

These results seem to militate towards the incorporation of censoring in the discrete choice model, but unfortunately that is not an easy task in such a setting. As we shall

⁵ The number of permitted allocations depends on the amount of total income, X . In CHST, the number of permitted allocations varies between 9 and 33.

demonstrate in what follows, dealing with censoring is much more natural in models based both on the RB and on the RP approaches.

	Specification				
	RU	RB	Mod-RB	RP	RP-Poisson
λ_1 (proportion Strict Egalitarian)	0.435 (0.090)	0.516 (0.106)	0.508 (0.084)	0.441 (0.106)	0.478 (0.085)
λ_2 (proportion Liberal Egalitarian)	0.381 (0.088)	0.460 (0.125)	0.285 (0.079)	0.293 (0.104)	0.181 (0.067)
λ_3 (proportion Libertarian)	0.184 (0.066)	0.024 (0.092)	0.207 (0.065)	0.266 (0.090)	0.341 (0.075)
γ	28.359 (3.589)	- -	- -	- -	- -
μ (mean of $\log(\theta_i)$), ζ (mean of $\log(\beta_i)$)	5.385 (0.349)	-1.755 (0.129)	-0.793 (0.075)	-1.022 (0.111)	- -
η (s.d. of $\log(\theta_i)$), σ (s.d. of $\log(\beta_i)$)	3.371 (0.530)	2.036 (0.336)	1.116 (0.146)	0.900 (0.055)	-
α	- -	- -	- -	- -	8.286 (0.156)
ρ	- -	- -	0.392 (0.036)	0.239 (0.039)	0.390 (0.032)
σ_v (s.d. of v_{it})	- -	70.490 (6.336)	86.096 (12.980)	- -	- -
number of subjects (n)	96	96	96	94	94
number of observations $\left(\sum_{i=1}^n T_i\right)$	190	190	190	186	186
Log likelihood	-337.584	-874.803	-581.986	-145.026	-442.454

Table 1: Maximum likelihood estimates of parameters of all five models defined in Section 5.⁶ All models are estimated in STATA version 11.0. The RB model (9) and the modified RB model (11) are maximised using 20-point Gauss-Hermite quadrature.

5.2 The Random Behavioural Model

The RB model, as discussed in Section 3, is defined as follows. The *desired* allocation by subject i (of type k) in task t is:

⁶ It is worth noting that the RU apparently succeeds in estimating both γ and β_i . However, what it is in fact estimating is the ratio of each of these parameters to another unknown parameter representing the standard deviation of the error term in their model. Hence, effectively, it is only the ratio of γ to β_i that it is succeeding in estimating. For further discussion of this sort of identification problem see Train (2003, p. 45).

$$\begin{aligned}
\tilde{y}_{it} &= m_k(\mathbf{a}_{it}, \mathbf{q}_{it}) + \theta_i X(\mathbf{a}_{it}, \mathbf{q}_{it}) + v_{it} \\
\theta_i &\sim \text{Lognormal}(\mu, \eta^2) \\
v_{it} &\sim N(0, \sigma_v^2) \quad t=1, \dots, T_i \quad i=1, \dots, n
\end{aligned} \tag{6}$$

We still assume that each subject draws a “selfishness” parameter, θ_i , from a log-normal distribution and that this value applies to all tasks. Variation in behaviour between tasks is then explained using the two-sided error, v_{it} , for which a new value is (independently) drawn for each task.

A feature of the data identified in Figure 3 is that for 58 of the 190 observations, the chosen allocation equals the total income. As already remarked, this is a clear manifestation of upper censoring, and any model should take this into account. We do so as follows. Since the maximum possible allocation is the total income $X(\mathbf{a}_{it}, \mathbf{q}_{it})$, the *observed* allocation (y_{it}) is obtained from the *desired* allocation ($\tilde{y}_{it}$ defined in (6)), via the following censoring rule:

$$\begin{aligned}
y_{it} &= \tilde{y}_{it} \quad \text{if } \tilde{y}_{it} < X(\mathbf{a}_{it}, \mathbf{q}_{it}) \\
y_{it} &= X(\mathbf{a}_{it}, \mathbf{q}_{it}) \quad \text{if } \tilde{y}_{it} \geq X(\mathbf{a}_{it}, \mathbf{q}_{it})
\end{aligned} \tag{7}$$

We further define a censoring indicator d_{it} , to be:

$$\begin{aligned}
d_{it} &= 0 \quad \text{if } y_{it} < X(\mathbf{a}_{it}, \mathbf{q}_{it}) \\
d_{it} &= 1 \quad \text{if } y_{it} = X(\mathbf{a}_{it}, \mathbf{q}_{it})
\end{aligned} \tag{8}$$

The individual likelihood function for RB Model is then:

$$\begin{aligned}
L_{i_RB}(\lambda_1, \lambda_2, \lambda_3, \mu, \eta, \sigma_v) \\
&= \sum_{k=1}^3 \lambda_k \int_0^\infty \prod_{t=1}^{T_i} \left[(1 - d_{it}) \frac{1}{\sigma_v} \phi\left(\frac{y_{it} - m_{k,it} - \theta X_{it}}{\sigma_v}\right) \right. \\
&\quad \left. + d_{it}(1 - p) \Phi\left(\frac{m_{k,it} + \theta X_{it} - X_{it}}{\sigma_v}\right) \right] f(\theta; \mu, \eta) d\theta
\end{aligned} \tag{9}$$

where $\phi(\cdot)$ and $\Phi(\cdot)$ are respectively the standard normal density function and the standard normal cumulative distribution function, and $f(\theta; \mu, \eta)$ is the (lognormal) density function evaluated at θ .

The estimation results for this model are reported in column 2 of Table 1. It reaches different conclusions from the previous model: although it agrees that “Strict Egalitarians” are the dominant type, it finds that “Libertarians” barely exist.

The second row of Figure 3 shows the simulated distribution for the RB. Contrarily to the RU model, this model seriously under-predicts the proportion at one half (p -value=0.0000), but successfully predicts the proportion at the full endowment (p -value = 0.7096). The p -value of 0.0000 in the joint test depicted in the third column represents strong overall evidence of prediction bias. The problem with this model is that apparently the assumed distribution of the selfishness premium does not incorporate the behaviour of those whose allocation corresponds exactly to their fairness ideal (which, as previously noted, is often one half).

Anyhow, in contrast to the RU model, the RB model is malleable enough to allow us to mould the distribution of the selfishness premium so to explain this probability mass. The new distributional assumptions give rise to the Modified-RB model that we present in the next subsection.

5.2.1 *The Modified-RB Model*

As remarked earlier, the data resulting from the experiment has the feature that a high concentration of the amount claimed is at exactly 50% of the endowment. Since 50% of the endowment always corresponds exactly to one or more of the fairness ideals, it is reasonable to infer that individuals claiming 50% of the endowment are behaving *exactly* in accordance with their fairness ideal, that is, that their “selfishness premium” is zero. This reasoning leads us to modify our assumption concerning the distribution of the selfishness premium. We assume that the allocation claimed by an individual in any task is either exactly equal to the fairness ideal, or it exceeds it. By how much it exceeds the fairness ideal depends on the degree of selfishness that the individual is experiencing at the time the decision is made. As with the RB model, we assume that an individual draws a selfishness premium from a log-normal distribution, and this value applies to all tasks. However, we further assume that for any task, with probability p , the selfishness premium is exactly zero, and there is no two-sided error. That is, with probability p , the individual behaves exactly in accordance with their fairness ideal. With probability $1-p$, their behaviour is in accordance with their positive selfishness-premium, with a two-sided error applied. The two-sided error is clearly necessary to allow variation in behaviour between tasks, when behaviour is away from the Fairness ideal. It can also explain those cases in which the amount subjects claim for themselves is below the fairness ideal. We shall refer to this model as the “Modified-Random-Behavioural model”.

The intuitive rationale of the Modified-RB model is as follows. Economic theoretic reasoning leads to the fairness ideal as the starting point in the construction of an econometric model. We then look at the data. One striking feature is that *no* observations are to the left of all of the fairness ideals. This amounts to straightforward non-parametric evidence of selfishness relative to fairness ideals. However, another important feature is that much of the data appears to be exactly on the fairness ideal. It is therefore not appropriate to assume that the amount by which a claim exceeds the fairness ideal (i.e. the selfishness premium) comes from a distribution with strictly positive support. This distribution must include a positive probability of zero.

These considerations lead us to modify the RB model, based on the behavioural equation (2), in the following way. Upper censoring is incorporated as in the RB model (eq. (7) and (8)).

$$\begin{aligned}\tilde{y}_{it} &= m(\mathbf{a}_{it}, \mathbf{q}_{it}) + \theta_i X(\mathbf{a}_{it}, \mathbf{q}_{it}) + v_{it} \\ v_{it} &\sim N(0, \sigma_v^2) \quad t = 1, \dots, T_i \quad i = 1, \dots, n \\ \theta_i &\sim \text{Lognormal}(\mu, \eta^2) \text{ with probability } 1 - p \\ \theta_i &= 0 \text{ with probability } p.\end{aligned}\tag{10}$$

The individual likelihood for this model is:

$$\begin{aligned}L_{i_M-RB}(\lambda_1, \lambda_2, \lambda_3, \mu, \eta, \sigma_v, p) \\ = \sum_{k=1}^3 \lambda_k \int_0^\infty \prod_{t=1}^{T_i} \left\{ (1 - d_{it}) \left[p I(y_{it} = m_{k,it}) \right. \right. \\ \left. \left. + (1 - p) I(y_{it} \neq m_{k,it}) \frac{1}{\sigma_v} \phi\left(\frac{y_{it} - m_{k,it} - \theta X_{it}}{\sigma_v}\right) \right] \right. \\ \left. + d_{it} (1 - p) \Phi\left(\frac{m_{k,it} + \theta X_{it} - X_{it}}{\sigma_v}\right) \right\} f(\theta; \mu, \eta) d\theta.\end{aligned}\tag{11}$$

Estimation results from this model are presented in column three of Table 1. A key feature of the modified-RB model is that it contains a parameter p representing the probability that the selfishness premium is exactly zero, leading to behaviour that is in exact accordance to the individual's fairness ideal. We see that the estimate of this probability is 0.392 with standard error 0.036. The magnitude and significance of this estimate provide clear evidence of the importance of this component of the model.

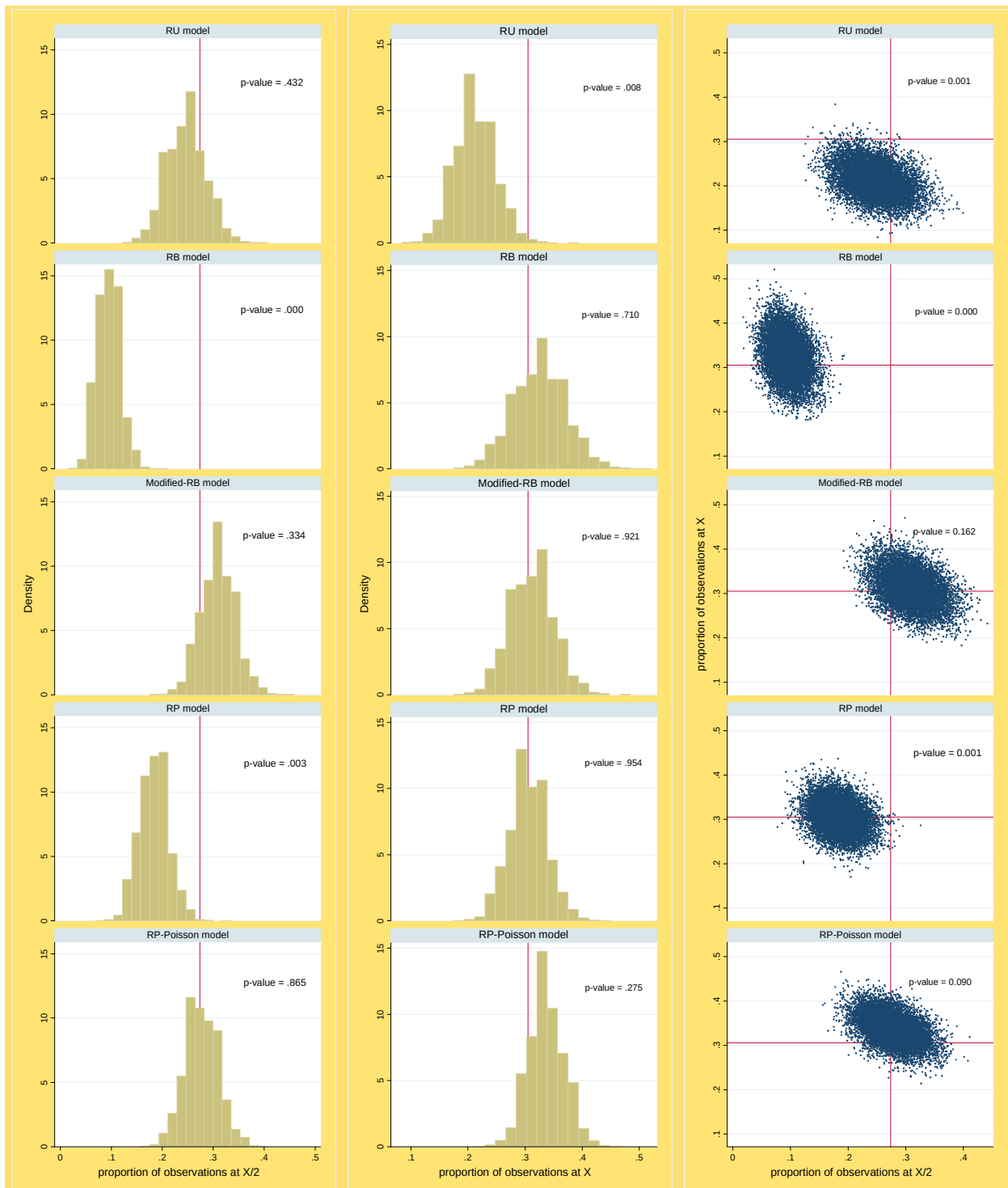


Figure 4: Simulated distributions of proportions of observations at: 50% of endowment (left pane); 100% of endowment (central pane) and joint testing of ability of models to predict proportion of observations at $X/2$ (horizontal axis) and proportion at X (vertical axis) (right pane). 10,000 replications. The red lines represent the true proportions. Top pane: RU model; second pane: RB model; third pane: Modified-RB model; fourth pane: RP model; bottom pane: RP-Poisson model. Univariate (left and central pane) and joint (right pane) tests' p -values are superimposed.

The third row of Figure 4 shows predictions from the Modified-RB model. We see that this model predicts both proportions accurately (p -value = 0.3344; p -value = 0.9214), and also when considered jointly (p -value = 0.162). The clear outcome from the prediction tests that have been considered so far is that the Modified-RB model is the first model that successfully explains both of the prominent features of the data simultaneously. For this reason we focus on the results of this model in interpretation.

The principal way in which the conclusions of the modified RB-model differ from those of the RU model and the RB model concerns the mixing proportions of the two types LE and L. Specifically, we again estimate the proportion of strict egalitarians (λ_1) to be around 0.5, but here we estimate the proportion of liberal egalitarians (λ_2) to be lower (our 0.285 compared to the 0.381 of the RU model and 0.460 of the RB model), and the proportion of Libertarians to be correspondingly higher.

5.3 The Random Preference model

The fourth model we arrive at is based on the Random Preference approach explained in Section 3. In this model, the allocation claimed by an individual in any task is either exactly equal to the fairness ideal, or it exceeds it, but it cannot fall below it.⁷ By how much it exceeds the fairness ideal depends on the degree of selfishness that the individual is experiencing at the time the decision is made. Each time a decision is made, the individual draws a selfishness parameter (θ_{it}) from a distribution which may be seen as a mixture of a lognormal distribution and zero, with mixing proportions $1 - p$ and p respectively. Since variation in behaviour is being explained in terms of variation in one of the model's parameters, the model is classified as a Random Preference model.⁸

Applying the random preference assumption to equation (3), we obtain the *desired* allocation by subject i (of type k) in task t :

$$\tilde{y}_{it} = m_k(\mathbf{a}_{it}, \mathbf{q}_{it}) + \theta_{it}X(\mathbf{a}_{it}, \mathbf{q}_{it}) \quad t = 1, \dots, T_i \quad i = 1, \dots, n \quad (12)$$

$\theta_{it} \sim \text{Lognormal}(\mu, \eta^2)$ with probability $1 - p$

$\theta_{it} = 0$ with probability p .

⁷ In the Modified RB model, allocations falling below the fairness ideals are explained by the two-sided error term.

⁸ This approach is inspired by Loomes and Sugden (1995)'s Random Preference model derived in the context of choice under risk.

This approach, by explaining variation in behaviour in terms of variation in the selfishness parameter, avoids the need for the additive error terms that characterise the first three approaches. Upper censoring is incorporated as in the RB and Modified-RB models. The sample likelihood function for subject i is then constructed as:

$$\begin{aligned}
 L_{i,RP}(\lambda_1, \lambda_2, \lambda_3, \mu, \eta, p) &= \sum_{k=1}^3 \lambda_k \prod_{t=1}^{T_i} I(y_{it} \geq m_{k,it}) \\
 &\times \left\{ (1 - d_{it}) \left[p I(y_{it} = m_{k,it}) + (1 - p) I(y_{it} > m_{k,it}) f\left(\frac{y_{it} - m_{k,it}}{X_{it}}; \mu, \eta\right) \right] \right. \\
 &\left. + d_{it} (1 - p) \Phi\left(\frac{\mu - \ln(X_{it} - m_{k,it})}{\eta}\right) \right\}. \tag{13}
 \end{aligned}$$

The likelihood in (13) is convenient in the sense that it is quite parsimonious and easy to estimate because it does not require any numerical integration.

As noted at the end of Section 3, logical problems are sometimes encountered with the RP model. There are some subjects whose behaviour cannot be explained for the following reason. For one of their decisions, the amount claimed indicates that they can *only* be of one “type”; for the other decision, the amount claimed indicates that they are of a *different* type. That is, there is no k for which $y_{it} > m_{k,it}$ for all t . This is inconsistent with the model, since it must be assumed that individuals cannot change type between tasks. For this reason two subjects are excluded from the estimation of this model. Obviously, if for a certain fairness type in one or both occasions the amount claimed is below the fairness ideal, the likelihood contribution for that type is zero.

The estimation results for this model are reported in the fourth column of Table 1. The mixing proportions are quite similar to those estimated under the Modified-RB model, but the estimated probability of being exactly at the fairness ideal is much lower, being 0.239 instead of 0.392.

From the fourth row of Figure 4, we see that the RP model appears to under-predict the proportion at one half (p -value = 0.0032) but to predict accurately the proportion at the full endowment (p -value = 0.9582). The joint test leads us to the conclusion that the RP model in this form is incapable of reproducing both the relevant features of the data (p -value = 0.001).

5.3.1. The Random Preference-Poisson Model

The Random Preference model is nevertheless useful in the modelling of another element of discreteness in the data that has so far been neglected: subjects always appear to claim for themselves amounts that are multiples of 50NOK, even though they are free to claim *any* amount between 0 and X . Such a feature of the data might lead to the mistaken belief that subjects are choosing between a discrete set of alternatives. Instead, it simply reflects the level of refinement subjects achieve when revealing (i.e. providing a measure of) their preferences in monetary terms. It may be useful to imagine that, when asked to report a measure of their preferences, subjects use a tape measure calibrated in units of 50NOK.⁹

The Random Preference approach is suitable for the modelling of this particular characteristic of the data. Let us rewrite eq. (3) in the following way:

$$\frac{\tilde{y}_{it}}{50NOK} = \frac{m_{k,it} + (\tilde{y}_{it} - m_{k,it})}{50NOK} = \frac{m_{k,it}}{50NOK} + n_{k,it} \quad t = 1, \dots, T_i \quad i = 1, \dots, n \quad (14)$$

$$\frac{(\tilde{y}_{it} - m_{k,it})}{50NOK} = n_{k,it} \sim \text{Poisson}(\delta_{it}) \text{ with probability } 1 - p$$

$$n_{k,it} = 0 \text{ with probability } p.$$

In order to capture the rationale behind such a re-formulation of the RP model and the new distributional assumptions, let us picture an individual, endowed with a certain number ($X_{it}/50NOK$) of 50NOK notes, in the act of splitting such an endowment between herself and the other individual. We expect that she starts counting the notes in her hands in order to separate the notes that she assigns to herself from those intended for the other individual. Of the number of notes she keeps for herself ($\tilde{y}_{it}/50NOK$), a certain number ($m_{k,it}/50NOK$) constitute her fairness ideal, while the remainder ($n_{k,it}$) reflect her selfishness. In this sense, this model can be set up as a count data model with the selfishness premium assumed to follow a zero-inflated Poisson distribution. We refer to this model as the *Random Preference-Poisson model*.

The censoring rule and the censoring indicator defined in eqq. (7) and (8) still apply. Here, we only need to notice that in CHST data set, for type Liberal Egalitarian, the fairness ideal is sometimes not a whole number, implying that the selfishness premium ($n_{k,it}$) is not a whole number either. In these cases, we assume that the number of notes retained is either the

⁹ Here, the fact that 50NOK is the smallest note in the Norwegian currency is likely to be relevant. On the assumption that subjects perceive the endowment as consisting of a bundle of banknotes, see Anderson et al. (1998), and in particular, footnote 7.

integer immediately below $n_{k,it}$ ($\lfloor n_{k,it} \rfloor$), or the integer immediately above $n_{k,it}$ ($\lceil n_{k,it} \rceil$), and the likelihood contribution is the sum of the poisson probabilities for these two adjacent integers.¹⁰

The following indicator is required for such cases:

$$\begin{aligned} H_{it} &= 1 \text{ if } (n_{k,it}) \neq \lfloor n_{k,it} \rfloor \\ H_{it} &= 0 \text{ otherwise.} \end{aligned} \quad (15)$$

The likelihood contribution of subject i is then constructed as:

$$\begin{aligned} L_{i,RP-P}(\lambda_1, \lambda_2, \lambda_3, \alpha, p) &= \sum_{k=1}^3 \lambda_k \prod_{t=1}^{T_i} I(y_{it} \geq m_{k,it}) \\ &\times \left\{ (1 - d_{it}) \right. \\ &\times \left[(1 - H_{it}) [pI(y_{it} = m_{k,it}) + (1 - p)I(y_{it} > m_{k,it})g(n_{k,it}; \alpha)] \right. \\ &\left. + H_{it} [pI(y_{it} = m_{k,it}) + (1 - p)I(y_{it} > m_{k,it}) \times [g(\lfloor n_{k,it} \rfloor; \alpha) + g(\lceil n_{k,it} \rceil; \alpha)]] \right] \\ &\left. + d_{it}(1 - p)[1 - G(n_{k,it}; \alpha)] \right\}, \end{aligned} \quad (16)$$

where $g(\cdot)$ and $G(\cdot)$ are respectively the Poisson probability distribution function and the Poisson cumulative probability distribution function evaluated at α .

In column 5 of Table 1, the estimation results for this model are listed. We can see that again the SE is the leading type with a proportion close to 0.5, but this time the Libertarian type are more represented than the Liberal Egalitarian type, with mixing proportions of respectively 0.34 and 0.18. The estimated value of p , that is the probability of being exactly at the fairness, is 0.39, almost exactly the same as for the Modified-RB model.

The last row of Figure 4 displays the simulated distribution for the RP-Poisson model. We see that this model successfully predicts the proportion at half (p -value = 0.865) and also the proportion at the full endowment (p -value = 0.275). In the joint test, the p -value is 0.090, representing prediction unbiasedness.

In summary, the RP-Poisson model provides an explanation of why discrete amounts are claimed by individuals. This is an alternative modelling strategy to the assumption of discrete choices underlying the RU approach. This model captures both the mass at 50% and at 100% of the endowment, both singularly and jointly, suggesting that this is a superior approach for dealing with discreteness.

¹⁰ For an introduction to the Interval Poisson model, see Moffatt (1995).

6. The mixing proportions and the posterior probabilities of types

So far, we have focused on the ability of the various fairness models presented to explain the distinctive features of the CHST dataset. All of the models share the characteristic of being mixture models. With each model, the proportions of the population who are of each type are estimated. A further criterion of model success is how well they assign subjects to types. In this sense, while the mixing proportions convey an idea of the proportions of the population who are of each type, we can go further than this by using Bayes' Rule to compute the posterior probability of each subject in the sample being of each type. This technique has previously been used by Conte et al. (2009). The posterior probabilities for our models can be computed as follows:

$$\Pr[\text{type } k | \text{obs}_i] = \frac{\Pr[\text{type } k] \times \Pr[\text{obs}_i | \text{type } k]}{\Pr[\text{obs}_i]} = \frac{\lambda_k \times \Pr[\text{obs}_i | \text{type } k]}{\Pr[\text{obs}_i]} = \frac{\lambda_k \times L_{i,h}^k}{L_{i,h}}, \quad (17)$$

where $k \in \{SE, LE, L\}$, $h \in \{RU, RB, M-RB, RP, RP-P\}$ and $L_{i,h}^k$ represent the component of the likelihood function corresponding to type k .

In this section, we make use of graphical displays of posterior probabilities following model estimation. These graphs convey useful information that the estimates themselves do not. In particular, they indicate how well the model discriminates between types for each individual subject.

Figure 5 shows the posterior probabilities obtained using (17) from each estimated model. Each of the subjects is represented by a single point in the graph. The RU model does not appear to be very successful at assigning individuals to types; in fact, all subjects are far from the vertices and concentrated in restricted area in the centre of the triangle. In the figure representing the posterior probabilities for the RB model, the fact that the vast majority of subjects are close to the outer edge of the triangle is consistent with the very low estimate of the proportion of Libertarians (0.024) that we obtain with the RB model. In the other three figures, we see a more even spread of the subjects. The Figures for the Modified-RB and the RP-Poisson models, show a greater tendency for subjects to be close to the corners of the triangle, suggesting that the Modified-RB model and the RP-Poisson model are the best able to assign subjects to types.

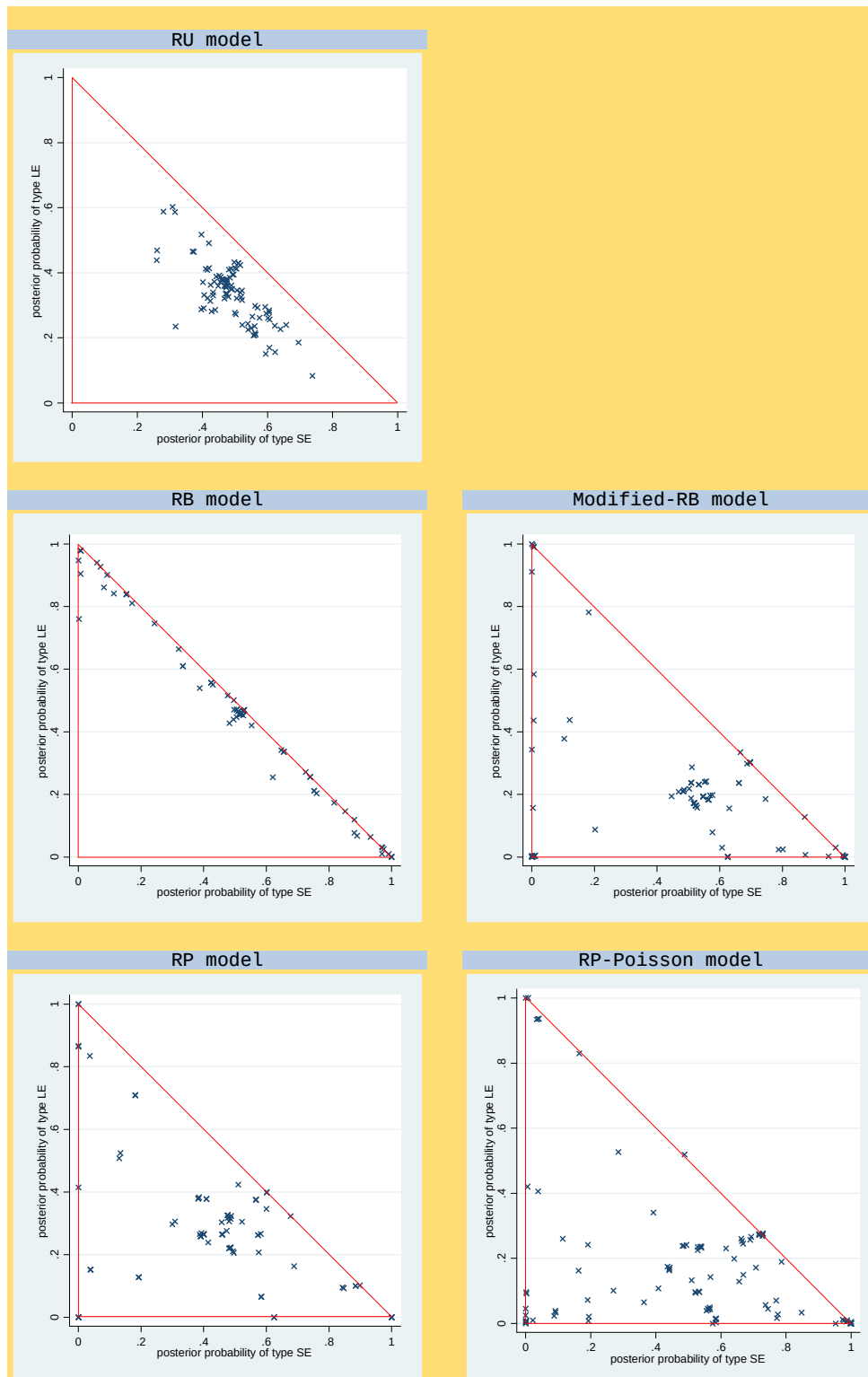


Figure 5: Posterior probabilities distribution of the three types from the five models estimated in Table 1.

In this regard, some of the graphical analyses may be seen as disappointing, since in our preferred models, the posterior probabilities indicate that we are unable to allocate many subjects to a single type with confidence. This may be a consequence of the low number of repeated observations per subject, and if the number of tasks were higher, we might find subjects gravitating more clearly to types. In any case, we must not lose sight of what we are ultimately interested in: the population. The estimated mixing proportions are consistent estimates of the proportions of the population who are of each type, regardless of how clearly the model is able to allocate *individuals* to types.

7. Conclusion

The importance of the stochastic component in modelling is now widely accepted among experimental economists. This paper has been concerned with the important question of exactly how the stochastic component should be introduced, with particular attention being paid to the particular problems that arise in the modelling of social preferences. A variety of approaches have been considered and evaluated, with the objective of fine-tuning the stochastic specification in such a way as to strike the ideal balance between theory and data.

The data set from the fairness experiment of Cappelen et al. (2007) has been very useful in the demonstration of these various approaches. This data set has features that can only be described as inconvenient from the standpoint of any theory of decision making. The features we have focused upon have been the large concentrations of allocations at 50% and 100% of the endowment. The first of these has been dealt with by tailoring the specification in such a way that there is a positive probability of behaviour coinciding *exactly* with the individual's "fairness ideal". The parameter (p) representing this probability is estimated to be between 0.24 and 0.40, depending on which specification is assumed, and therefore is clearly an important, if not essential, component of the models.

The large concentration of observations at 100% has been dealt with by assuming upper censoring. However, dealing with this data feature is another example of the delicate interplay between theory and data. It would be possible to rely on the theory, by assuming the presence of a fourth type: the "totally selfish" type. The peculiar characteristics of such a type would be that of keeping everything for themselves regardless of the dimension of the pie and of any other considerations about productivity, contribution to the total income, and so on. In Section 4,

concerns have already been expressed about the idea of introducing a notion of fairness that contradicts any principle of fairness. Also, by allowing individuals to have a positive selfishness premium, the possibility of extreme selfish behaviour is already taken into account. Finally, a totally selfish type would not solve the censoring problem at least in a within-subject framework. The problem of censored observations would anyway arise for those who show themselves as a fair type in some cases and behave totally selfishly in others. Hence the introduction of a totally selfish type would not itself be enough to explain all of the observations at 100%.

Consideration also needs to be given to observations which are neither at 50% nor at 100% of the endowment. According to the theory, the extent of selfishness has a continuous distribution, so the allocation also has a continuous distribution. Once again the data challenges the theory: all observed allocations are seen to be exact multiples of NOK50. In recommending ways for dealing with this type of discreteness, we have attached importance to the question of *why* it arises. If it were a consequence of the experimental design, with subjects induced to choose between a discrete array of alternatives, then we might favour the discrete choice modelling approach¹¹. However, this is not the case in this experiment. As discussed in Section 5.3.1, the discreteness is apparently the result of subjects' perception of their decision as one of choosing the number of banknotes to keep for themselves. One possible approach would be simply to follow the example of mainstream econometric modellers by treating the discreteness as the result of rounding of the decision variable, and therefore as a form of minor measurement error which is of little consequence in estimation. Another approach would be to be faithful to the idea of the subject "counting" from a bundle of banknotes, by assuming a discrete distribution for the selfishness premium that is suitable for count data. This approach was followed in Section 5.3.1 in the guise of the RP-Poisson model.

Whatever our prior beliefs about the manner in which individuals make decisions, and whatever our personal preferences between modelling approaches, we eventually need to consider which model is best able to explain the data. To this end, we have, in figure 4, tested

¹¹ In the experimental literature it is quite common to analyse open-ended questions with a continuous logit-type approach. See, for an example in the context of public good games, Anderson et al. (1998). This is essentially a version of the RU model described here with an infinite number of alternatives as a support. A further reason why we have not considered such an approach in this paper is because, as a discrete choice model, it is not malleable enough when modelling censored data.

the predictive accuracy of each model using simulation methods. Of the five models that have been considered, the most successful in this regard are the Modified-RB and RP-Poisson models, both of which give unbiased predictions of both of the features of the data that we have scrutinised. Another criterion of model performance is the ability to assign each subject to a type with high posterior probability. This was investigated in section 6, where we found that the Modified-RB and RP-Poisson models again performed best. On statistical grounds, we are therefore robustly led to favour these two models.

It is very important to identify the “best” model, in the way we have attempted to do in the previous paragraph. This is because the estimates obtained, in particular for the mixing proportions, are highly sensitive to the choice of econometric specification. This means that the conclusions we reach concerning fairness-related behaviour depend crucially on this choice.

Finally, we like to add that, in this paper, we have tried to make a contribution towards the building of a bridge that connects the theory to the data, passing through the inner details of the econometric modelling. We do not wish to suggest that our analysis is in any way exhaustive of the multifaceted characteristics of experimental data. However, we at least hope that this paper will have the effect of stimulating further research in this direction.

References

- Anderson, Simon P., Goeree, Jacob K., and Charles A. Holt. 1998. "A theoretical analysis of altruism and decision error in public goods games." *Journal of Public Economics*, 70, 297-323.
- Bardsley, Nicholas. 2000. "Control without deception: individual behaviour in free-riding experiments revisited." *Experimental Economics*, 3, 215-241.
- Bardsley, Nicholas and Peter G. Moffatt. 2007. "The experimentics of public goods: inferring motivations from contributions." *Theory and Decision*, 62, 161-193.
- Bellemare, Charles, Kröger, Sabine, and Arthur van Soest. 2008. "Measuring inequity aversion in a heterogeneous population using experimental decisions and subjective probabilities," *Econometrica*, 76(4), 815-839.
- Botti, Fabrizio, Anna Conte, Daniela T. Di Cagno, and Carlo D'Ippoliti. 2008. "Risk Attitude in Real Decision Problems." *The B.E. Journal of Economic Analysis & Policy – Advances*, Vol. 8 (1).
- Cappelen, Alexander, W., Astri D. Hole, Erik Ø. Sorensen, and Bertil Tungodden. 2007. "The pluralism of fairness ideals: an experimental approach." *American Economic Review*, 97(3), 818-827.
- Conte, Anna, John D. Hey, and Peter G. Moffatt. 2009. "Mixture models of choice under risk." *Journal of Econometrics*, in press.
- Conte, Anna, and Vittoria Levati. 2010. "Contribution plans in a public good game", *mimeo*.
- Costa-Gomes and Weizsäcker. 2008. "Stated beliefs and play in normal-form games." *The Review of Economic Studies*, 75, 729-762.
- Fechner, Gustav. 1860/1966. *Elements of Psychophysics*, Vol. 1. New York: Holt, Rinehart and Winston.
- Harrison, Glenn, and Elisabet Rutström. 2009. "Expected Utility And Prospect Theory: One Wedding and Decent Funeral." *Experimental Economics*, 12(2), 133-158.
- Loomes G., P.G. Moffatt, and R. Sugden. 2002. "A microeconomic test of alternative stochastic theories of risky choice." *Journal of Risk and Uncertainty*, 24, 103-130.
- Loomes G., and R. Sugden. 1995. "Incorporating a stochastic element into decision theories." *European Economic Review*, 39, 641-8.
- Moffatt P.G.. 1995. "Grouped Poisson regression models: theory and an application to public house visit frequency." *Communications in Statistics-Part A: Theory and Methods*, 24(11), 2779-2796.
- Train, Kenneth. 2003. *Discrete choice methods with simulation*. Cambridge: Cambridge University Press.