



Revealing the extent of process heterogeneity in choice analysis: An empirical assessment

Sean M. Puckett, David A. Hensher *

Institute of Transport and Logistics Studies (ITLS), Faculty of Economics and Business, The University of Sydney, NSW 2006, Australia

ARTICLE INFO

Article history:

Received 3 October 2007

Accepted 24 July 2008

Keywords:

Process heterogeneity

Event description

Stated choice

Design complexity

Relevance

Subadditivity

Packaging

ABSTRACT

Choice analysts increasingly use a mix of revealed preference and stated choice data paradigms to identify preferences of samples of individuals that are used to infer behavioural response and willingness to pay for specific attributes. These data are in a sense artificial constructs that are developed to approximate real choice settings of the way that individuals process relevant information in making choices. As such, all data designs formalized through a survey instrument seek information through questions that become *descriptions* of events and as such the probabilities of choice that are of interest are strictly probabilities attached to event descriptions and not choice probabilities of events per se. The recognition of this distinction, initially noted by Kahneman et al. [Kahneman, D., Slovic, P., Tversky, A., 1982. *Judgement under uncertainty: Heuristics and biases*. Cambridge University Press, New York], can be captured, at least in part, through the idea of process heterogeneity, as a way of recognizing and accounting for the many ways in which individuals process information, and in part is influenced by the way the analyst describes the context in which preference data is sought. Building on previous contributions on attribute processing, this paper draws on recent empirical evidence to further reinforce the importance of joint modelling of process and outcome in choice analysis. This study adds to the evidence of a trend emerging on the upward bias of mean estimates of marginal willingness to pay when ignoring process heterogeneity.

© 2008 Published by Elsevier Ltd.

1. Introduction

Choice modelling has advanced significantly in recent years in terms of the econometric representation of preference heterogeneity within a sample (see Train 2003) and in terms of ways in which data are obtained, notably the design of choice experiments (see Rose and Bliemer 2007) and the tailored mechanisms available to capture behaviourally rich choice responses through internet and computer aided personal survey instruments (CAPI) (see Hensher et al. 2007).

What has been given less attention, although by no means ignored (see Bonini et al. 2004 and Hensher 2008), has been the questioning of the appropriateness of the frameworks used to obtain preference information from a sample of respondents (Rabin 1998). Specifically, the nature of the stated or revealed choice setting may not be sufficiently 'realistic' to provide the necessary information required to obtain choice probabilities and estimates of willingness to pay for specific attributes in the context of other attributes (and their levels). The recognition of process heterogeneity (Hensher 2008) suggests that more effort should be invested in understanding the relationship between what is offered up to each respondent and how they process such information, including the extent to which the experimental or revealed setting (neither of which is a true

* Corresponding author. Tel.: +61 2 93510071; fax: +61 2 93510088.

E-mail addresses: Seanp@itls.usyd.edu.au (S.M. Puckett), Davidh@itls.usyd.edu.au (D.A. Hensher).

definition of an actual event) increases the ‘gap’ between the choice probability of interest and the ‘choice’ probability that is obtained. This ‘gap’ or bias is in part a reflection of a potential over-simplification of the way that the heuristics that individuals adopt in choice making in real-markets are clouded by the use of *descriptions* of events instead of events per se.

The ability to portray events in the exact way that individuals perceive real world events is virtually impossible in a deterministic sense, and the best the analyst can do is to try and represent the events as realistically as possible, including the possibility of event variation over time (e.g., through new products and/or levels of attributes of existing products). What the analyst is unable to do adequately through a focus on choice outcome is to identify the influence that specific process strategies have on outcome. Specifically, individuals bring a whole raft of processing (including judgmental) capability and rules to the table when faced with scenarios designed with great statistical ingenuity by analysts, and it is the judgments of evidence strength or support that underlies numerical judgments of preference (be it through a first preference, a rank or a rate) and subsequent influence on choice probability.

Building on the contribution by Hensher and colleagues on attribute processing and the important distinction between complexity and relevancy, this paper draws on recent empirical evidence to further reinforce the importance of joint modelling of process and outcome in choice analysis. This research considers the effects of processing heterogeneity utilised by respondents for *every alternative in every choice set* faced, acknowledging that varying process rules may be enacted not only across decision makers, but across choice tasks faced by a given decision maker.

2. Empirical context to assess the presence of process heterogeneity

The data utilised in the empirical discussion within this section are from a 2005 study of road freight transport providers and their customers in Sydney, Australia. The study was designed to elicit preferences under a hypothetical road user charging system. The first step in the process involved administering the experiment to representatives of freight transport firms. Centred on a CAPI survey with a d-optimal experimental design (discussed in Puckett and Hensher (2008)), the stated choice experiment involves three distinct procedures: (1) non-stated-choice questions intended to capture the relevant deliberation attributes and other contextual effects; (2) choice menus corresponding to a freight-contract-based setting (see Fig. 1 and Puckett and Hensher, 2008); and (3) questions regarding the attribute processing strategies enacted by respondents within each choice set. The resulting estimation sample, after controlling for outliers and problematic respondent data¹, includes 108 transporters, yielding 432 choice sets. The response rate was 45%.

In all cases except for the variable charges, the attribute levels for each of the stated choice alternatives are pivoted off of the levels of the reference alternative, as detailed below. The levels are expressed as deviations from the reference level, which is the exact value specified in the corresponding non-stated choice questions, unless noted:

- (1) Free-flow time: –50%, –25%, 0, +25%, +50%.
- (2) Congested time: –50%, –25%, 0, +25%, +50%.
- (3) Waiting time at destination: –50%, –25%, 0, +25%, +50%.
- (4) Probability of on-time arrival: –50%, –25%, 0, +25%, +50%, with the resulting value rounded to the nearest 5% (e.g., a reference value of 75% reduced by 50% would yield a raw figure of 37.5%, which would be rounded to 40%). If the resulting value is 100%, the value is expressed as 99%. If the reference level is greater than 92%, the pivot base is set to 92%. If the pivot base is greater than 66% (i.e., if 1.5 times the base would be greater than 100%) let the pivot base equal X , and let the difference between 99% and X equal Y . The range of attribute levels for on-time arrival when $X > 66\%$ are (in percentage terms): $X - Y$, $X - 0.5 * Y$, X , $X + 0.5 * Y$, $X + Y$. This yields five equally-spaced attribute levels between $X - Y$ and 99%.
- (5) Fuel cost: –50%, –25%, 0, +25%, +50% (representing changes in fuel taxes of –100%, –50%, 0, +50%, +100%).
- (6) Distance-based charges: Pivot base equals $0.5 * (\text{reference fuel cost})$, to reflect the amount of fuel taxes paid in the reference alternative. Variations around the pivot base are: –50%, –25%, 0, +25%, +50%.

The extant literature, with rare exception (see Hensher et al. 2006; Hensher 2008) either ignores process heterogeneity altogether, utilises the information structure within choice sets as an indicator of cognitive burden, or uses global (i.e., across all choice sets faced) indicators of attribute exclusion and agglomeration. This research considers the effects of attribute processing strategies utilised by respondents for *every alternative in every choice set* faced, acknowledging that varying process

¹ Preliminary analysis revealed that the degree of heterogeneity in reference trips was sufficiently high that some outliers obscured the inferential power of the data. After careful consideration, the following observations were removed from the final sample: (a) trips based on a fuel efficiency over 101 l per 100 km (or approximately twice the average fuel consumption for the larger trucks in the sample); (b) trips based on a probability of on-time arrival less than 33%; (c) round trips (or tours) of less than 50 km; and (d) round trips of more than 600 km. The trips eliminated, based on low fuel efficiency, may have obscured the results due to significantly prohibitive values for fuel cost and variable charges, reflecting reference trips that are too atypical to be pooled with other trips. An alternative source of obscuring effects via low fuel efficiency may be that the implied values of fuel efficiency were inaccurate, and hence either made the trade-offs implausible to respondents or reflect an inability of the respondent to offer meaningful information on which to base the alternatives. The trips eliminated, based on low probability of on-time arrival, are likely to have obscured the results because the trips involved travel quality significantly worse than the remainder of the sample, making the pooling of these trips into the sample problematic. Similarly, extremely short or long trips may have involved trade-offs that are significantly different to the trade-offs made by respondents in the sample at large.

Download English Version:

<https://daneshyari.com/en/article/312527>

Download Persian Version:

<https://daneshyari.com/article/312527>

[Daneshyari.com](https://daneshyari.com)