

Institute for the Study of Social Change

# ISSC DISCUSSION PAPER SERIES

## ANCHORING BIAS AND COVARIATE NONRESPONSE.

*Dr. Rosalia Vazquez-Alvarez*

*Swiss Institute for International Economics and Applied Economic Research (SIAW),  
University of St. Gallen.*

*This paper is produced as part of the Policy Evaluation Programme at ISSC; however the  
views expressed here do not necessarily reflect those of ISSC. All errors and omissions  
remain those of the author.*

# Anchoring Bias and Covariate Nonresponse

Rosalia Vazquez-Alvarez

*Swiss Institute for International Economics and Applied Economic Research (SIAW), University of  
St.Gallen, Switzerland.*

This version: May 2003<sup>\*</sup>

## Abstract

Non-random item nonresponse makes identification of parameters problematic. Such nonresponse can occur with respect to both dependent and conditioning variables. A method often used to reduce nonresponse is that of adding unfolding brackets as follow up to open-ended questions. With these, initial non-respondents can provide additional (incomplete) information on the missing value. However, recent studies suggest that responses to unfolding brackets can lead to a type of bias as a result of ‘the anchoring effect’. In this paper, bounding intervals of the type as presented in Horowitz and Manski (1998) are extended to incorporate information provided by bracket respondents while allowing for different types of anchoring, and, therefore, accounting for significant nonresponse in the conditioning set. The theoretical framework is illustrated with empirical evidence based on the 1996 wave of the Health and Retirement Study.

**JEL Classification:** C13, C14, C42, C81, D31.

**Key Words:** Unfolding brackets design, anchoring effects, survey nonresponse, bounding intervals.

---

<sup>\*</sup> I am grateful to Arthur vanSoest and Bertrand Melenberg for advice and past co-operation on related work, and to Charles F. Manski for suggesting this topic as an extension to Chapter 5 in my doctoral thesis. This paper has benefited from comments at ISSC’s AEW Seminar Series (Ireland) and local seminar participants in HSG (CH). The author is also affiliated to ISSC (UCD-Ireland) and the Health Policy Reset Group at the ESRI in Dublin, Ireland. Address for correspondence: Rosalia Vazquez-Alvarez, SIAW, Dufourstrasse, 48, University of St.Gallen, 9000-SG, Switzerland. Phone: +41 71 224 23 50, Fax: +41 71 224 22 98, Email: [rosalia.vazquez-alvarez@unisg.ch](mailto:rosalia.vazquez-alvarez@unisg.ch).

# 1. Introduction

Lets assume that you are interested on testing if there is a significant difference between the smoking habits of males and females, relative to their position in the income distribution. That is, your interest is to test if the absolute value between  $P(\text{smoking} | \text{income} \in B, \text{male})$  and  $P(\text{smoking} | \text{income} \in B, \text{female})$  is significantly different than zero in the target population. One problem we might face when performing such test is that of regressor (nonrandom) nonresponse. In the above example, this implies income nonresponse, a problem often encounter in household surveys. Since the seminal work by Heckman on sample selection issues (for example, see Heckman, 1979), it is well known that ignoring non-respondents (i.e., assuming exogeneity), is often unrealistic and can lead to severe selection bias on the estimated measures. A possible solution is to impose a particular distributional assumption on the missing values, for example, the use of selection models to jointly model the response behaviour and the variable of interest, conditional on a set of covariates (see Vella (1998) for a survey on this topic). Both parametric and semi-parametric selection models avoid the assumption that item nonresponse is random conditional on a set of variables, but require alternative assumptions such as a single index assumption or independence between covariates and error terms.

Since the early 1990s a new approach has emerge which aims at estimating the parameter of interest without imposing any assumptions or very weak data assumptions. Manski (1989, 1990, 1994, 1995) has shown how bounds on either the distribution function or conditional quantile of the distribution can be derived while imposing no assumptions and allowing for any type of nonrandom nonresponse. Manski's framework is intuitively appealing, empirically easy to apply and very flexible. The result of this approach is to estimate an upper and lower bound on the measure of interest, where such bounds allow for both sampling error as well as error due to nonresponse. The implication is that the flexibility of the method is at the expense of increased uncertainty.

Before the estimation process begins, a solution often used in household surveys to reduce the problem of missing data is to give initial non-respondents the choice to provide partial information classifying the missing value into a particular category. When such questioning strategy is available, it is often the case that a significant percentage of initial non-respondents will provide some information, even if incomplete. Juster and Smith (1997) suggest cognitive factors (e.g., confidentiality issues, or an initial lack of accurate information and/or the realisation

from respondents that the interviewer does not require precise information) to explain why people might prefer to provide information in the form of a category. One of the possible types of categorical questions often used in household surveys is that of an unfolding bracket design. The reason why data collectors might prefer such type is because it is easier to administer over the telephone, thus is less costly than other types of categorical questioning (e.g., range cards). In an unfolding bracket design initial non-respondents to an open-ended question are routed to an ordered sequence of bids the result of which is a set of categories. For example, if respondents who are asked to declare their annual income (an amount) answer 'don't know' or 'refuse', they can be routed to a second question where they are given a bid, say B1, and asked if their income exceeds such bid. According to their answer ('yes', 'no' or 'don't know'), they might be further routed to a second bid B2k, for  $k=0,1$  (with  $B21 > B1$  if 'yes' to B1, or  $B20 < B1$  if 'no' to B1). Faced with this second bid they might be asked to declare, again, if income is greater than B2k. The number of bids in a sequence is defined by the data collectors, and although it can vary according to the variables in question, it is often no greater than 2 or 3 bids. The number of categories formed by the unfolding bracket design is endogenously defined by respondent's answers.<sup>1</sup>

One problem with an unfolding bracket design is that partial information provided by bracket responses is subject to 'the anchoring effect', a phenomenon well documented in the psychological literature. The idea is to think that the bid creates a fictitious belief in the respondent's mind: faced with a question related to an unknown quantity, the respondent treats the question as a problem solving situation, and the given bid becomes an anchor with respondents using it as a cue to solve the problem. This can result in responses that are influenced by the design of the unfolding sequence. A leading example of the phenomena is found in Jacowitz and Kahneman (1995). In their study they use experimental data to show that uninformative anchors (with respect to the true answer) given arbitrarily can have large significant effect on subject responses. Another example is that of Hurd et al. (1998). In their study the use of an experimental module with randomised bids shows that the distribution of partial responses is biased towards the categories closed to the initial bid. Other parametric models of the anchoring effect are introduced by Cameron and Quiggin (1994) and Herriges and

---

<sup>1</sup> As opposed to a situation where individuals are faced with a range card. A range card shows each respondent a set of possible categories in which to classify the missing value. The categories are exogenous defined.

Shogren (1996).

In general, the message sent by studies on the anchoring effect from both the psychological and economic literature tell us that answers given in a sequence of unfolding brackets might be wrong, that is, they provide incorrect information with respect to the true distribution of information from partial respondents. Thus, whereas bracket respondents might reduce nonresponse considerably, estimates based on such data (in the presence of both nonresponse and unfolding bracket response) need to account for the possible bias created by the anchoring effect.

This paper extends the paper by Horowitz and Manski (1998) in that it further allows for a sub-population of bracket respondents (rather than full and non-respondents only) to derive bounds in the presence of covariate nonresponse.<sup>2</sup> Because information from bracket responses can be subject to the anchoring effect, the extension draws from Vazquez-Alvarez, R., B. Melenberg and A. vanSoest (2000) to incorporate the possibility of bias due to the anchoring effect. Such derivations draw from three competing models of anchoring, namely, Jacowitz and Kahneman (1995), Heringes and Shogren (1996) and Hurt et al. (1998).<sup>3</sup> The theoretical framework is applied to the 1996 wave of the Health and Retirement Study, to test for any significant difference in the smoking behaviour of males versus females at different levels of the income distribution, where income is defined as annual labour income. We compare bounds which allow with those which do not allow for bias due to anchoring. When ignoring anchoring effects but accounting for bracket response information, the data suggests a significant difference in the smoking habits between genders, with the probability of smoking been significantly higher for males than for females at all intervals of income. Once the possible existence of bias due to anchoring is incorporated, the horizontal distance for each pair of estimated bounds in each sub-sample widens, thus increasing the overlap of the identification regions for the unknown probabilities of smoking between sub-populations. This results in a statistical test of no difference between gender's smoking probabilities that is weaker in power than the test

---

2 In Horowitz and Manski (1998), as well as exploring the case of regressor nonresponse, they also derive bounds in the presence of joint nonresponse (of dependent and conditional variables) and the mixed case where nonresponse affects some but not all the variables in the estimation process. Appendix A summarises similar arguments as in Section 2 for the cases of either joint or mixed nonresponse.

3 Sections 2 and 4 use the interpretation of anchoring as in Jacowitz and Kahneman (1995). Appendix B draws from Vazquez-Alvarez, R., B. Melenberg and A van Soest (2000) to show how bounds can also be derived with alternative models of anchoring.

performed assuming no anchoring effects, since, for all income intervals, a wider overlapping region implies that the null of equality between gender's smoking probability cannot be rejected. In fact, once anchoring is allowed for, and relative to estimates of worst case bounds without bracket information, the result of adding partial information from bracket respondents does not help to improved the identifying power of the bounds.

The reminder of the paper is organised as follows. Section 2 sets forth the theoretical framework deriving bounds under unfolding brackets and three sub-populations of individuals, namely full respondents, bracket respondents and full non-respondents. Section 3 describes the Health and Retirement Study data used in the empirical illustration. Section 4 explains the estimation procedure and presents the empirical results. Section 5 concludes.

## 2. Theoretical framework

### 2.1 Regressor censoring without bracket respondents

First we consider identification with regressor censoring without bracket respondents. Suppose that we have a representative sample of the target population, and we want to make inferences about the distribution of an outcome  $Y = y \in R$ , conditional on  $X = x \in R$ , where the conditional variable suffers from non-negligible item nonresponse<sup>4</sup>, while the outcome variable  $Y$  is observed for the full sample of size  $n$ . Let  $FR$  indicate that  $x$  is observed while  $NR$  indicates that  $x$  is missing, such that  $P(FR) + P(NR) = 1$ . Drawing from Horowitz and Manski (1998), we can apply Bayes Theorem to partition the measure of interest  $E[g(y) | x \in A]$ :

$$E[g(y) | A] = E[g(y) | A, FR] \times \frac{P(A | FR)P(FR)}{P(A | FR)P(FR) + P(A | NR)P(NR)} + E[g(y) | A, NR] \times \frac{P(A | NR)P(NR)}{P(A | FR)P(FR) + P(A | NR)P(NR)} \quad (1)$$

where the left hand side of (1) is not identified since neither  $P(A | NR)$  or  $E[g(y) | A, NR]$  are identified by the sampling process. With respect to  $p = P(A | NR)$ , the only knowledge we have

---

<sup>4</sup> For simplicity of exposition it is assumed there is only one variable in the conditioning set. Similar arguments apply to a situation with more than one variable in the conditioning set. Appendix A further looks at the situation where nonresponse affects both the dependent and the conditioning set.

is that  $p \in [0,1]$ . On the other hand, although covariate nonresponse precludes identification of  $E[g(y) | A, NR]$ , the sampling process allows for identification of  $E[g(y) | NR]$ . Drawing from Horowitz and Manski (1998) and applying Proposition 1 in Horowitz and Manski (1995), expression  $E[g(y) | NR]$  can be partition as:

$$E[g(y) | NR] = E[g(y) | NR, A]p + E[g(y) | NR, \bar{A}](1 - p) \quad (2)$$

where  $\bar{A}$  is the complement space of  $A$ . Assuming that  $p$  is known, a sharp restriction on  $E[g(y) | A, NR]$  is given by

$$E[g(y) | A, NR] \in [g_0(p), g_1(p)], \quad (3)$$

with  $g_0(p) = \inf[h : h \in G(p)]$ ,  $g_1(p) = \sup[h : h \in G(p)]$  and  $G(p) = \{ \int g(y) d\psi, \psi \in \Psi(p) \}$  where  $\Psi(p)$  denotes the set of all distributions of  $Y$ , for a given  $p \in [0,1]$ . For example, if  $g(y) = I[y = 1]$ ,  $\Psi(p) = [0,1]$  and  $E[g(y) | A, NR] = P(y | A, NR)$ , so that using (2) expression (3) becomes:<sup>5</sup>

$$P(y | A, NR) \in \left[ \max \left( 0, \frac{P(y | NR) - (1 - p)}{p} \right), \min \left( 1, \frac{P(y | NR)}{p} \right) \right] \quad (4)$$

whereas for  $g(y) = I[y \leq t]$ ,  $E[g(y) | A, NR] = P(Y \leq y | A, NR)$ , such that:

$$P(Y \leq y | A, NR) \in \Psi(p) \cap \left[ \frac{P(Y \leq y | NR) - (1 - p)\psi}{p}, \psi = P(Y \leq y | \bar{A}, NR) \in \Psi(p) \right] \quad (5)$$

However, the sampling process does not identify the measure  $P(Y \leq y | \bar{A}, NR)$ . The only thing we know is that for any  $y \in Y$  the measure falls in the  $[0,1]$  interval. Therefore, (5) can also be

---

<sup>5</sup> This example anticipates the empirical illustration in Section 4, where we treat the case where  $g(y)=I[y \leq t]$  and  $g(y)=I[y=1]$ . Expression (4) provides also the interpretation for the first case where the aim is to bound the conditional distribution function of a continuous variable  $y$ . This is because in the continuous case the numerators inside the min and max expression in (4) would be substituted by  $[P(y \leq t|A)-p]$

expressed as:

$$P(Y \leq y | A, NR) \in \left[ \max \left( 0, \frac{P(Y \leq y | NR) - (1-p)}{p} \right), \min \left( 1, \frac{P(Y \leq y | NR)}{p} \right) \right] \quad (6).$$

Using the generic form in (3), bounds on (1) are given by:

$$\begin{aligned} E[g(y) | A, FR] \times \frac{P(A | FR)P(FR)}{P(A | FR)P(FR) + P(NR)p} + g_0(p) \frac{P(NR)p}{P(A | FR)P(FR) + P(NR)p} \\ \leq E[g(y) | A] \leq \\ E[g(y) | A, FR] \times \frac{P(A | FR)P(FR)}{P(A | FR)P(FR) + P(NR)p} + g_1(p) \frac{P(NR)p}{P(A | FR)P(FR) + P(NR)p} \end{aligned} \quad (7)$$

But (7) assumes that  $p = P(A | NR)$  is known. For unknown  $p = P(A | NR)$  a sharp bound on  $E[g(y) | A]$  is obtained from (7) by minimising and maximising the lower and upper bound, respectively, with respect to  $p = P(A | NR) \in [0,1]$ . Therefore, a computable bound on  $E[g(y) | A]$  given non-negligible non-random nonresponse on  $x \in A$ , is given by:

$$\begin{aligned} \inf_p \left\{ E[g(y) | A, FR] \times \frac{P(A | FR)P(FR)}{P(A | FR)P(FR) + P(NR)p} + g_0(p) \frac{P(NR)p}{P(A | FR)P(FR) + P(NR)p} \right\} \\ \leq E[g(y) | A] \leq \\ \sup_p \left\{ E[g(y) | A, FR] \times \frac{P(A | FR)P(FR)}{P(A | FR)P(FR) + P(NR)p} + g_1(p) \frac{P(NR)p}{P(A | FR)P(FR) + P(NR)p} \right\} \end{aligned} \quad (8)$$

## 2.2 Regressor censoring with bracket respondents

Assume the variable in question refers to a continuous variable, for example, income. Surveys are often designed so that initial non-respondents can provide partial information by classifying the missing amount into a category within a range of categories. If so, a sample of  $n$  respondents can be partitioned into three sub-categories, namely full respondents (FR), full non-respondents (NR) and bracket respondents (BR). With this, expression (1) is modified such that:

$$\begin{aligned}
E[g(y) | A] = & E[g(y) | A, FR] \times \frac{P(A | FR)P(FR)}{P(A | FR)P(FR) + P(A | NR)P(NR) + P(A | BR)P(BR)} \\
& + E[g(y) | A, NR] \times \frac{P(A | NR)P(NR)}{P(A | FR)P(FR) + P(A | NR)P(NR) + P(A | BR)P(BR)} \\
& + E[g(y) | A, BR] \times \frac{P(A | BR)P(BR)}{P(A | FR)P(FR) + P(A | NR)P(NR) + P(A | BR)P(BR)}
\end{aligned} \tag{9}$$

As was the case with (1), expression (9) is not identified by the data since the sampling process provides no information on either  $E[g(y) | A, NR]$  or  $P(NR | A)$ . Moreover, expression (9) implies that part of those who were previously classified as full non-respondents are not bracket respondent. From this latter group we can only attain partial information on both,  $E[g(y) | A, BR]$  and  $P(BR | A)$ . Both the data collection method and the interpretation of such partial information determines the derivation of upper and lower bounds. We assume from the start that bracket respondents provide partial information using a follow-up unfolding bracket design. Initially, all surveyed individuals are given an open ended question. An unfolding bracket design consists on routing all who are initial non-respondents to the open-ended question towards a particular category using sequential bids. Let  $B1$  be the initial bid, and assume that all initial non-respondents are given the same sequence of bids.<sup>6</sup> The first bracket question is given by

$$\text{"Is the ammount \$B1 or more?"} \tag{10}$$

to which individuals answer “yes”, “no” or “don’t know”.<sup>7</sup> Individuals who answer “yes” receive the same question with a new bid  $B21$ , ( $\infty > B21 > B1$ ), whereas those who answer “no” are faced with a new bid  $B20$ , ( $0 < B20 < B1$ ). Although bracket respondents can face more than two bids, this is rare in practice, therefore, we confine our analysis to a two-bid unfolding bracket design.<sup>8</sup> For example, if the covariate in question is income, and  $B1 = \$25,000$ ,  $B20 = \$5,000$

---

6 In many studies the introduction of randomised bids implies different starting values  $B1$  for different initial non-respondents. Although the design can be such that all respondents end up classified within similar categories, the advantage is that randomising the starting bid can lead to data able to test for general starting up bias. See Hurd et al (2000) for an example using the 1996 module from the Health and Retirement Study.

7 Here the assumption is that there is no distinction between the “don’t know” and “Refuse”, and both are treated as “don’t know” answers.

8 With it, the theoretical framework anticipates the empirical illustration, although we can generalise this

and  $B_{21} = \$50,000$ , such design would define the following sequence of categories:

**Table 1: Example of derived categories according to information provided by bracket respondents.**

| Group                                           | Anchor 1:<br>B1 | Answer to<br>Anchor 1 | Anchor 2:<br>B20/B21 | Answer to<br>Anchor 2 | Resulting Categories |
|-------------------------------------------------|-----------------|-----------------------|----------------------|-----------------------|----------------------|
| <b>Complete Bracket<br/>respondent (CBR)</b>    | >\$25,000       |                       |                      | Yes                   | [\$50,000-infinity)  |
|                                                 |                 | Yes                   | >\$50,000            | No                    | [\$25,000-\$50,000)  |
|                                                 |                 | No                    | >\$5,000             | Yes                   | [\$5,000-\$25,000)   |
|                                                 |                 |                       |                      | No                    | [\$0-\$5,000)        |
| <b>Incomplete Bracket<br/>Respondents (ICB)</b> | >\$25,000       | Yes                   | >\$50,000            | DK/RF                 | [\$25,000-infinity)  |
|                                                 |                 | No                    | >\$5,000             | DK/RF                 | [\$0-\$25,000)       |

The distinction between “complete” and “incomplete” bracket respondents is necessary because some bracket respondents might not complete the sequence. In the case of two bids, this implies that some individuals respond either “don’t know” or “refuse” when faced with the second bid. However, the distinction between Complete and Incomplete bracket respondents has implication only with reference to empirical applications, so without any loose, deriving bounding intervals with bracket information can be done assuming that all bracket respondents belong to the CBR sub-category.<sup>9</sup>

## Case 1: Not allowing for the Anchoring Effect

With reference to the first anchor, let  $Q_1 = 1$  if the answer to (10) is ‘yes’, and  $Q_1 = 0$  if the answer is ‘no’. With this, bracket respondents identify  $P(Q_1 = 1 | A, BR)$  such that,

---

presentation to any  $K$  number of bids (bracket categories).

<sup>9</sup> With Complete and Incomplete bracket respondents, the measure of interest  $E[g(y) | A, BR]$  can be expressed as  $E[g(y) | A, BR] = E[g(y) | A, CBR, BR]P(CBR | A, BR) + E[g(y) | A, IBR, BR]P(IBR | A, BR)$ . Those who do not complete the sequence (IBR sub-group) provide partial information as if faced with an unfolding sequence design with one anchor  $B_1$ , so the bounding interval is defined over two regions in the distribution of  $A$ , namely  $A < B_1$  and  $A \geq B_1$ . The difference with the sub-group of CBR is that, in the case of two anchors, the bounds are defined over 4 partitions on the distribution of  $A$ , i.e.,  $[0, B_{20})$ ,  $[B_{20}, B_1)$ ,  $[B_1, B_{21})$  and  $[B_{21}, \max)$ . Therefore the distinction between IBR and CBR is only empirically relevant. For more detail derivation of bounds with such distinction the reader is referred to Vazquez-Alvarez, R., B. Melemborg and A. vanSoest, (2000).

$$\begin{aligned}
P(Q1 = 1 | BR) &= P(Q1 = 1 | BR, A < B1)P(A < B1 | BR) \\
&+ P(Q1 = 1 | BR, A \geq B1)P(A \geq B1 | BR)
\end{aligned} \tag{11}$$

If there is no anchoring, all bracket respondents answer correctly to question (10). This implies that  $P(Q1 = 1 | A \leq B1, BR) = 0$ ,  $P(Q1 = 1 | A > B1, BR) = 1$ ,  $P(Q1 = 1 | BR) = 1 - P(A \leq B1 | BR)$  and therefore  $P(A > B1 | BR)$  is identified by the data on bracket respondents. This leads to the following bounds on  $P(A | BR)$ :

$$\begin{aligned}
\text{for } A \leq B1, \quad & 0 \leq P(A | BR) \leq P(Q1 = 0 | BR) \\
\text{for } A > B1, \quad & P(Q1 = 0 | BR) \leq P(A | BR) \leq 1
\end{aligned} \tag{12}$$

A similar argument applies to the case where we have an unfolding bracket design with two bids, rather than just one. Define dummy variables  $Q20$  and  $Q21$  for those who answer the second bracket question on  $B20$  and  $B21$  with  $Q1=1$  and  $Q1=0$ , respectively. For example,  $Q20 = 1$  if  $Q1=0$  and the individual respondent declares the amount to be greater than  $B20$ . With the introduction of  $Q20$  and  $Q21$ , two further probabilities are not identified by the data, namely,  $P(Q20 = 1 | BR, Q1 = 0)$  and  $P(Q21 = 1 | BR, Q1 = 1)$ . Again, these probabilities can be expressed as partitions with respect to the distribution of  $A$  such that:

$$\begin{aligned}
P(Q2k = 1 | BR, Q1 = k) &= P(Q2k = 1 | BR, A \leq B2k, Q1 = k)P(A \leq B2k | BR, Q1 = k) \\
&+ P(Q2k = 1 | BR, A \geq B2k, Q1 = k)P(A \geq B2k | BR, Q1 = k)
\end{aligned} \tag{13}$$

for  $k = 0, 1$ . Under the assumption of no anchoring effect, the implication is that respondents answer correctly to the question with the second bid. Thus, under no anchoring effect  $P(Q2k = 1 | A \leq B2k, BR, Q1 = k) = 0$ , such that

$$P(Q2k = 1 | BR, Q1 = k) = P(A \geq B2k | BR, Q1 = k) \tag{14}$$

for both  $k = 0, 1$ . Expression (14), together with (12) define a bounding interval for different regions of the distribution according to the range of values for the conditional variable  $A$ . In

general, for the case where we assume no anchoring effect, the implication is that those who answer the bracket question do so correctly, so that for each bracket respondent it is known whether  $A$  is in  $[0, B_{20}), [B_{20}, B_1), [B_1, B_{21}),$  or  $[B_{21}, \infty)$ . The information is identical to a situation where initial non-respondents are given a range card question and each has to choose one of the four (simultaneously given) categories to classify the missing value. Denoting the category containing  $A$  by  $[B_j, B_{j+1}]$ <sup>10</sup>, we have that

$$L_1(B_j) \leq P(A | BR) \leq U_1(B_{j+1}) \quad (15)$$

We now turn our attention to the measure  $E[g(y) | A, BR]$  in expression (9). With partial information provided by bracket respondents the latter can now be identified (at least) up to an interval according to the categories defined by the unfolding bracket design. That is, if conditioning on  $A$  implies a range such that our interest is  $E[g(y) | A \in [B_j, B_{j+1}), BR]$  where both  $B_j$  and  $B_{j+1}$  are identical to the given anchors, then  $E[g(y) | A \in [B_j, B_{j+1}), BR]$  is fully identified by the data. Likewise, if our interest is on  $E[g(y) | A \leq B_j, BR]$  where, again,  $B_j$  is an bid within the sequence of bids, then partial information is sufficiently informative to identify  $E[g(y) | A \leq B_j, BR]$ . However, conditioning on a particular value of the distribution (i.e.,  $A = t$ ) or conditioning within a range  $A \in [C_j, C_{j+1})$  with such range  $[C_j, C_{j+1})$  not given by the bids, implies that additional information provided by bracket respondents does not help to improve the informative power of the bounds. In this case, bound on  $E[g(y) | A, BR]$  are attained in a similar manner as bounds on  $E[g(y) | A, NR]$ , thus,

---

<sup>10</sup> Is straight forwards to show that for  $A \in [0, B_{20})$ ,  $L(a) = 0$  and  $U(a) = P(Q_{20} = 0 | BR, Q_1 = 0)$ , for  $A \in [B_{20}, B_1)$ ,  $L(a) = P(Q_{20} = 0 | BR, Q_1 = 0)$  and  $U(a) = P(Q_1 = 0 | BR)$ , for  $A \in [B_1, B_{21})$ ,  $L(a) = P(Q_1 = 0 | BR)$  and  $U(a) = P(Q_{21} = 0 | BR, Q_1 = 1)$  and finally, for  $A \in [B_{21}, \infty)$ ,  $L(a) = P(Q_{21} = 0 | BR, Q_1 = 1)$  and  $U(a) = 1$ .

for  $A = t$  and  $t \in \mathbb{R}$ , or,  $A \leq t$ ,  $t \in \mathbb{R}$ ,  $t \neq B_1$  and  $t \neq B_{2k}$ , for  $k = 0, 1$ ,

$$E[g(y) | A, BR] \in \Psi(p_b) \cap \left[ \frac{E[g(y) | BR] - (1 - p_b)\psi}{p_b}, \psi \in \Psi(p_b) \right]$$

where  $p_b = P(A | BR) \in (0, 1)$

(16)

Otherwise, for  $A \in [B_j, B_{j+1}]$  or  $A \leq B_j$  where  $B_j$  and  $B_{j+1}$  are bids,

$E[g(y) | A, BR]$  is identified.

In this case,  $p_b = P(A | BR) \in [L(B_j), U(B_{j+1})]$ .

As in section 2.1, the interpretation of  $g(y)$  determines the shape of the bounds in (16), so that if  $g(y) = I[y = t]$ ,  $t \in \mathbb{R}$ , bounds on  $E[g(y) | A, BR]$  are analogous to (4) – in case that the measure is not identified – and if  $g(y) = I[y \leq t]$ ,  $t \in \mathbb{R}$ , then the appropriate interpretation is that given in (6). Lets assume a generic form for (16) such that  $E[g(y) | A, BR] \in [g_{ob}(p_b), g_{ib}(p_b)]$ , where  $g_{ob}(p_b) = g_{ib}(p_b)$  if  $A$  is fully defined by the sequence of anchors. With this, and following a similar argument for  $E[g(y) | A, NR]$  as in section 2.1, a set of computable bounds on  $E[g(y) | A]$  is given such that,

for any partition  $A \in [C_j, C_{j+1}]$ ,  $C_j \geq 0$ ,  $C_{j+1} < \infty$

$$\begin{aligned}
& \inf_{p, p_b} \left\{ \begin{aligned} & E[g(y) | A, FR] \times \frac{P(A | FR)P(FR)}{P(A | FR)P(FR) + P(BR)p_b + P(NR)p} \\ & + g_{0b}(p_b) \times \frac{P(A | BR)P(BR)}{P(A | FR)P(FR) + P(BR)p_b + P(NR)p} \\ & + g_0(p) \times \frac{P(A | NR)P(NR)}{P(A | FR)P(FR) + P(BR)p_b + P(NR)p} \end{aligned} \right\} \\
& \leq E[g(y) | A] \leq \\
& \sup_{p, p_b} \left\{ \begin{aligned} & E[g(y) | A, FR] \times \frac{P(A | FR)P(FR)}{P(A | FR)P(FR) + P(BR)p_b + P(NR)p} \\ & + g_{1b}(p_b) \times \frac{P(A | BR)P(BR)}{P(A | FR)P(FR) + P(BR)p_b + P(NR)p} \\ & + g_1(p) \times \frac{P(A | NR)P(NR)}{P(A | FR)P(FR) + P(BR)p_b + P(NR)p} \end{aligned} \right\} \tag{17}
\end{aligned}$$

If  $[C_j, C_{j+1})$  is a range identified by the sequence of anchors so that  $[C_j, C_{j+1}) = [B_j, B_{j+1})$ , then  $g_{0b}(p_b) = g_{1b}(p_b)$  and  $p_b \in [L_1(B_j), U_1(B_{j+1})]$ . Otherwise,  $g_{0b}(p_b) \leq g_{1b}(p_b)$  and  $p_b \in (0, 1)$ .

## Case 2: Allowing for the Anchoring Effect

In the case where we assume the existence of an anchoring effect on partial information provided by bracket respondents, assumptions in (12) and (14) are no longer valid for the cases of one and two bid unfolding bracket design, respectively. This is because in the presence of anchoring the measures  $P(Q1 = 1 | BR, Y \leq B1)$  and  $P(Q2k = 1 | BR, Y \leq B2k)$ , for  $k = 0, 1$ , can be nonzero, so that under the assumption of anchoring<sup>11</sup>,

---

<sup>11</sup> In fact, we need to assume the possible existence of anchoring since testing for anchoring is only possible if we had randomised the starting bid among all partial respondents. In this case, some respondents would receive different bids. Testing for anchoring would consist on testing for a significant shift in the distribution of information provided by bracket respondents, where the shift would be shown to be a function of the starting bid. See Hurd et al (1998) for an example using an special module from the HRS, 1996.

$$\begin{aligned}
&P(Q1 = 1 | BR) \neq P(A > B1 | BR), \\
&\text{and} \\
&P(Q2k = 1 | BR) \neq P(A > B2k | BR, Q1 = k), \text{ for } k = 0, 1.
\end{aligned} \tag{18}$$

Deriving feasible sharp bounds on each of the expression in the right hand side of (18), thus feasible sharp bounds on  $E[g(y) | A]$ , requires plausible assumptions as to how the anchoring phenomena in the data affects information from partial respondent. The anchoring effect is a phenomena well documented in the psychology and economic literature (seminal examples are Jacowitz and Kahneman (1995), Rabin (1998) and Hurd et al. (1998)). In general it is explained by suggesting that a bid creates a fictitious belief in the respondent's mind. Faced with a question which relates to an unknown quantity, the respondent treats the question as a problem solving situation, the given bid becomes an 'anchor' and it is thus used as a cue to solve the problem. This can result in responses which are endogenous to the design of the unfolding bracket. For example, in the case of a continuous variable, the result can be a significant shift in the distribution of the categorical answers. If one aims to study the distribution of such variables, it is important to account for the possible bias created by the anchoring effect. In case of bounding intervals, this implies a modification of expressions (12) and (14), as well as a new set of conditions with respect to the bounding intervals on  $E[g(y) | A, BR]$ . Hurd et al. (1998) model the anchoring effect suggesting that respondents to (10) compare  $A$  to  $B1 + \varepsilon$ , where  $\varepsilon$  is the perception error. Whereas in Hurd et al. (1998)  $\varepsilon$  is assumed to be normally distributed with zero mean and independent of  $A$ , Vazquez-Alvarez et al. (2000) relax their parametric set up to a more flexible semi-parametric assumption where  $med(\varepsilon | A, BR) = 0$ . Hurd et al. (1998) provide an explanation for the anchoring phenomena in the data, but it might not be the most intuitively appealing way to model anchoring. Perhaps a more plausible alternative to model anchoring is given by Jacowitz and Kahneman (1995)<sup>12</sup>. In their experimental study they find that, if a high anchor is used, respondents too often report that the amount exceeds the anchor. This can be interpreted as  $P(Q1 = 1 | BR) \geq P(A \geq B1 | BR)$  if  $B1$  is large. Jacowitz and Kahneman (1995) report that this finding is not symmetric for their case study, and could well be reversed if the amounts have a natural upper instead of lower bound. An operational version of the phenomenon discussed by Jacowitz and Kahneman (1995) for one-bid unfolding bracket design would be

---

<sup>12</sup> See footnote 3 and Appendix B for alternative models of anchoring.

$$\begin{aligned}
P(Q1 = 1 | BR) &\geq P(A \geq B1 | BR) \quad \text{if} \quad P(Q1 = 1 | BR) \leq 0.5 \\
P(Q1 = 1 | BR) &\leq P(A \geq B1 | BR) \quad \text{if} \quad P(Q1 = 1 | BR) \geq 0.5
\end{aligned} \tag{19}$$

In the case of two bids the complement to (19) to model anchoring according to the Jacowitz and Kahneman's (1995) assumption would be given by

$$\begin{aligned}
P(Q2k = 1 | BR, Q1 = k) &\geq P(A \geq B2k | BR, Q1 = k) \quad \text{if} \quad P(Q2k = 1 | BR, Q1 = k) \leq 0.5 \\
P(Q2k = 1 | BR, Q1 = k) &\leq P(A \geq B2k | BR, Q1 = k) \quad \text{if} \quad P(Q2k = 1 | BR, Q1 = k) \geq 0.5
\end{aligned} \tag{20}$$

Thus, bounds allowing for this particular model on the anchoring effect are attained if we substitute (12) and (14) by a particular (data-dependent) interpretation of (19) and (20). Anticipating our empirical example, let's assume that in a two-bid unfolding bracket design, the data suggest that both B1 and B21 are 'large' and B20 is a small bid.<sup>13</sup> With this, bounds on  $P(A | BR)$  are derived such that,

$$\begin{aligned}
(i) \quad &P(A > B1 | BR) \leq P(Q1 = 1 | BR) \\
(ii) \quad &P(A > B21 | BR, Q1 = 1) \leq P(Q21 = 1 | BR, Q1 = 1) \\
(iii) \quad &P(A > B20 | BR, Q1 = 0) \geq P(Q20 = 1 | BR, Q1 = 0)
\end{aligned} \tag{21}$$

furthermore,

$$(MON) \quad P(A \leq B_j | Z) \leq P(A \leq B_j)$$

where 'MON' refers to a monotonic assumption with respect to any sub-space defined by Z, and complements (21). Given (21)-(iii),

---

<sup>13</sup> In our particular data set, where the number of bracket respondents equals 320, only 0.36 (<0.5) answer 'yes' when faced with B1, and of these only 0.26 (<0.5) answer 'yes' when faced with B21. With such data evidence we bounds will be derived assuming that B1 and B21 are 'large'. On the other hand, 204 of the 320 answer 'no' to the initial B1, and of these, 170 answer 'yes' to the second bid. Then,  $P(Q20=1 | BR, Q1=0) = 0.83 (>0.5)$ , suggesting that B20 is 'small'.

$$\begin{aligned}
P(A \leq B20 \mid BR) &= P(A \leq B20 \mid Q1 = 1, BR)P(Q1 = 1 \mid BR) \\
&\quad + P(A \leq B20 \mid Q1 = 0, BR)P(Q1 = 0 \mid BR) \\
&\leq P(Q20 = 0 \mid Q1 = 0, BR)P(Q1 = 1 \mid BR) \\
&\quad + P(A \leq B20 \mid BR)P(Q1 = 0 \mid BR). \\
&\text{since} \\
P(A \leq B20 \mid Q1 = 0, BR) &\leq P(Q20 = 0 \mid Q1 = 0, BR) \quad \text{by (21) – (iii),} \\
&\text{and} \\
P(A \leq B20 \mid Q1 = 0, BR) &\leq P(A \leq B20 \mid BR) \quad \text{by (MON).}
\end{aligned} \tag{22}$$

Therefore,

$$0 \leq P(A \leq B20 \mid BR) \leq P(Q20 = 0 \mid BR, Q1 = 0).$$

The lower bound equals zero because (21) provides no information to bound  $P(A \leq B20)$  from below. With respect to  $P(A \leq B1)$ ,

$$\begin{aligned}
P(A > B1 \mid BR) &\leq P(Q1 = 1 \mid BR) \quad \text{by (21) – (i),} \\
&\text{therefore,} \\
P(Q1 = 0 \mid BR) &\leq P(A \leq B1 \mid BR) \leq 1.
\end{aligned} \tag{23}$$

In this case, (21) lacks information on the upper bound of  $P(A \leq B1)$  so that the only thing we know is that it cannot exceed 1. Finally, for  $P(A \leq B21)$ ,

$$\begin{aligned}
P(A > B21 \mid BR, Q1 = 1) &\leq P(Q21 = 1 \mid Q1 = 1) \quad \text{by (21) – (iii),} \\
\Rightarrow \\
P(A \leq B21 \mid BR, Q1 = 1) &\geq P(Q21 = 0 \mid Q1 = 1), \\
\Rightarrow \\
P(A \leq B21) &\geq P(Q21 = 0 \mid Q1 = 1) \quad \text{by (MON)}
\end{aligned} \tag{24}$$

therefore,

$$P(Q21 = 0 \mid Q1 = 1) \leq P(A \leq B21) \leq 1.$$

Expressions (22) – (24) define a bounding intervals for  $P(A \mid BR)$  for regions of the distribution with partitions defined by the design of the unfolding brackets such that

$$\begin{aligned}
&\text{for } [0, B_{20}), \quad 0 \leq P(A | BR) < P(Q_{20} = 0 | BR, Q_1 = 0) \\
&\text{for } [B_{20}, B_1), \quad P(Q_{20} = 0 | BR, Q_1 = 0) \leq P(A | BR) < P(Q_1 = 0) \\
&\text{for } [B_1, B_{21}), \quad P(Q_1 = 0) \leq P(A | BR) < P(Q_{21} = 0 | BR, Q_1 = 1) \\
&\text{for } [B_{21}, \infty), \quad P(Q_{21} = 0 | BR, Q_1 = 1) \leq P(A | BR) < 1
\end{aligned} \tag{25}$$

Expression (25) provides bounds on  $P(A | BR)$ , such that  $L_2(B_j) \leq P(A | BR) \leq U_2(B_{j+1})$ , under the anchoring model according to Jacowitz and Kahneman (1995). These will change under any other assumption of anchoring. Appendix B summarises results as derived in Vazquez-Alvarez, R., B. Melenberg and A. vanSoest (2000), where expressions analogous to (25) make reference to models of anchoring according to Hurd et al. (1998) and Herringes and Shogren (1996). Bounds under alternative models of anchoring are competing bounds, that is, in theory nothing tells us that one set might be tighter than the other, and it is only empirically when this can be tested. On the other hand, bounds on  $P(A | BR)$  given by (15) derive tighter (sharp) bounds on  $E[g(y) | A]$  than those attained by applying expression (25) since the assumptions behind (25) are weaker than those underlying (15).

The next step before attaining an expression analogous to (17) for Case 2, is to incorporate the assumption by Jacowitz and Kahneman (1995) on  $E[g(y) | A, BR]$ . For easiness of exposition, express  $E[g(y) | A, BR]$  as  $P(y = 1 | A, BR)$ , where anticipating our empirical example we can think of  $y_i = 1$  if individual  $i$  smokes, and  $y_i = 0$  otherwise. As before, although bracket respondent provide only partial information on  $P(y = 1 | A, BR)$ , the measure  $P(y = 1 | BR)$  is identified by the sampling process. With this bounds on  $P(y = 1 | A, BR)$  are possible applying similar arguments as with expression (2)-(6) so that an interval in the distribution of  $A$  given by  $[B_j, B_{j+1}]$ ,

$$\begin{aligned}
&P(Y = 1 | BR, A \in [B_j, B_{j+1}]) \in \Psi(p_b) \cap \left\{ \frac{P(Y = 1 | BR) - (1 - p_b)\psi}{p_b}, \quad \psi \in \Psi(p_b) \right\} \\
&\text{for} \\
&p_b \in [L_2(B_j), U_2(B_{j+1})].
\end{aligned} \tag{26}$$

In (26),  $\psi = P(Y = 1 | BR, A \in [B_j, B_{j+1}]) \in \Psi(p_b)$ , and, therefore,  $\psi$  is also affected by anchoring via the conditioning set. However, at worst we know that  $\psi \in (0,1)$ , thus,

$$P(Y = 1 | BR, A \in [B_j, B_{j+1}]) \in \left\{ \max \left[ 0, \frac{P(Y = 1 | BR) - (1 - p_b)}{p_b} \right], \min \left[ 1, \frac{P(Y = 1 | BR)}{p_b} \right] \right\}$$

for

$$p_b \in [L_2(B_j), U_2(B_{j+1})]. \quad (27)$$

Thus, even if we have no knowledge on  $\psi$ , allowing such measure to vary over all possible values implies that (27) will nest the bounds which allow for anchoring. Further to having full information on  $P(Y = 1 | BR)$ , we also know that applying Bayer's Theorem,

$$P(Y = 1 | A \in [B_j, B_{j+1}]) = \frac{P(A \in [B_j, B_{j+1}] | Y = 1)P(Y = 1)}{P(A \in [B_j, B_{j+1}])} \quad (28)$$

Recall that (21) implies the model of anchoring by Jacowitz and Kahneman (1995), with such assumptions been equally applicable to all subspaces within the distribution of values of the covariate A, therefore, assumptions in (21) are equally applicable to the subspace defined by the condition  $(Y = 1)$ . For example, if (21)-(i) implied that  $P(A \geq B | BR) \leq P(Q1 = 1 | BR)$  for any B interval, the same assumption is applicable to the subspace of smokers in the bracket respondents, such that  $P(A \geq B | BR, Y = 1) \leq P(Q1 = 1 | BR, Y = 1)$ . Extending this to (i)-(iii) in (21), and applying a similar argument as with (22) – (24), it is possible to derive a set of bounds similar to those in (25) for  $P(A \leq B | BR, Y = 1)$  on the bids-dependent intervals such that,

$$\begin{aligned} \text{for } [0, B_{20}), \quad & 0 \leq P(A | BR, Y = 1) < P(Q_{20} = 0 | BR, Q_1 = 0, Y = 1) \\ \text{for } [B_{20}, B_1), \quad & P(Q_{20} = 0 | BR, Q_1 = 0, Y = 1) \leq P(A | BR, Y = 1) < P(Q_1 = 0 | BR, Y = 1) \\ \text{for } [B_1, B_{21}), \quad & P(Q_1 = 0 | BR, Y = 1) \leq P(A | BR, Y = 1) < P(Q_{21} = 0 | BR, Q_1 = 1, Y = 1) \\ \text{for } [B_{21}, \infty), \quad & P(Q_{21} = 0 | BR, Q_1 = 1, Y = 1) \leq P(A | BR, Y = 1) < 1 \end{aligned} \quad (29)$$

Applying expression (29) to the numerator in (28) leads to a bounding interval on  $P(Y = 1 | A \in [B_j, B_{j+1}], BR)$  such that,

$$P(Y = 1 | BR, A \in [B_j, B_{j+1}]) \in \left\{ \frac{L_3(B_j)P(Y = 1)}{p_b}, \frac{U_3(B_{j+1})P(Y = 1)}{p_b} \right\}$$

for

$$p_b \in [L_2(B_j), U_2(B_{j+1})] \quad (30)$$

Let bounds in (27) be written as  $P(Y = 1 | BR, A \in [B_j, B_{j+1}]) \in [LB_1(p_b), UB_1(p_b)]$ , and those in (30) we written as  $P(Y = 1 | BR, A \in [B_j, B_{j+1}]) \in [LB_2(p_b), UB_2(p_b)]$ . Then, under the assumption of anchoring as modelled by Jacowitz and Kahneman (1995), bounds on  $E[g(y) | BR, A \in [B_j, B_{j+1}]]$  will be given by

$$E[g(y) | BR, A \in [B_j, B_{j+1}]] \in [\min\{LB_1(p_b), LB_2(p_b)\}, \max\{UB_1(p_b), UB_2(p_b)\}] \quad (31)$$

where  $p_b \in [L_2(B_j), U_2(B_{j+1})]$ . With this, which allows for anchoring in the presence of partial information, bound on  $E[g(y) | A]$  have the same generic form as in (17) except that for  $[C_j, C_{j+1}] = [B_j, B_{j+1}]$ , we draw from (31) such that  $g_{0b}(p_b) = \min\{LB_1(p_b), LB_2(p_b)\}$  and  $g_{1b}(p_b) = \max\{UB_1(p_b), UB_2(p_b)\}$  for  $p_b \in [L_2(B_j), U_2(B_{j+1})]$ . Because (31) incorporates the anchoring assumption, bounds under (31) will be equal or wider than those under (17) which are tighter since the underlying assumption in (17) is stronger than that in (31), i.e., bounds in (17) assumes exogeneity with respect to the anchoring effect, while (31) accounts for anchoring thus leading to an increase in the estimated uncertainty interval due to regressors nonresponse.

### 3. Data

The empirical illustration in Section 3 draws from the 1996 wave of the Health and Retirement Survey (HRS). This is a longitudinal study conducted by the University of Michigan for the US National Institute of Ageing. It focuses mainly on aspects of health, retirement and economic status of US citizens born between 1931 and 1941, allowing for individuals and household information from a representative sample in this cohort. The data is collected every two years, starting in the summer of 1992, and is organised into four different sections, namely, demographics, health issues, assets & incomes, and employment status.

Initially the panel consisted on approximately 7,600 households. The 1996 wave has data from 6,739 households, representing 10,964 individuals. The respondents are the household representatives that satisfy the age criteria, and their partners, regardless of their age (second household respondents). All household representatives are asked to provide information on the four categories within the survey for themselves and in some cases for their spouses. Variables in the health section of the survey are answered individually by each member of the household, to provide both objective and subjective information on their health status, as well as on a variety of health habits, such as smoking, drinking or practice of regular exercise. Most of these variables are either categorical or binary in nature, and item nonresponse is a rear event, if at all. The empirical illustration in this paper aims at understanding smoking habits of individuals by gender and income level, thus we draw the following question from the health section of the questionnaire:<sup>14</sup>

*'Do you smoke cigarettes now?*

*'...yes'*

*'...cigars (if volunteered with this answer)'*

*'...pipe (if volunteered with this answer)'*

*'...no'*

*'...Don't know, Refuse'*

Each respondent classifies the answer in one of the above categories. Out of 10,964 individuals surveyed in 1996, only 1 was classified as 'don't know/Refuse', therefore, nonresponse is not a problem with respect to information on smoking behaviour.

The 'Income & Asset' section of the questionnaire is answered mostly by the household representative, with these amounting to 6,816 out of the total 10,964. They are asked to provide employment status and earned incomes for themselves and for their partners. We focus on information regarding household representative since this is the target population in the survey.<sup>15</sup>

---

14 This is the only question in the Health section which elicits information on the smoking habits of individuals. The question is not as complete since it is directed only to cigarette smokers. If we assume the question captures the smoking habits of the target population, we assume that all individuals who smoke pipe, cigars, etc., volunteer to defines their specific smoking habit.

15 Furthermore, because responses from second household respondent (often the spouse) are not directly answered by them, the quality of the information on these individuals is lower than information directly relevant to the household representative. For example, information on second household respondents suffer more item nonresponse

Initially, each household representative is asked if he or she worked for pay during the last calendar year. Of the 6,816 individuals, 4,145 declared to have worked for pay, 2,661 declare not to have worked for pay, and only 10 individuals (amounting to 0.2% of the weighted sample) answer with either ‘Don’t know’ or ‘Refuse’. With such distribution we can assume that nonresponse is not an issue at this level of information. The 4,145 individuals who declare to have worked for pay during the last calendar year are further asked if any of their earnings came from either wages or salaries (as opposed to earnings from self-employment). Of these, 3,602 declare that their incomes are from wages and/or salaries, and only 6 of the 4,145 are classified as ‘Don’t know’ or ‘Refuse’. Yet again it seems as if this categorical question is not affected by nonresponse. In this paper the sample is defined as the 3,602 individuals who declare to have earned some wages or salaries.<sup>16</sup> For these group nonresponse becomes an issue when they are asked to specify the exact amount of wages or salaries in an open ended question given as:

*‘about how much wage and/or salary income did you receive during the last calendar year?’*

*‘...any amount’ (in USA dollars)*

*‘...Don’t known’*

*‘...Refuse’*

Of the 3,602, there were 3,160 individuals who answered with an exact amount in US dollars, ranging from \$0,00 to \$350,000, with a mean of \$31,340 and standard deviation \$28,310. The median was \$25,370. Of the remaining, 438 individuals answered ‘don’t know’ (or ‘refuse’).<sup>17</sup> This translated into a 11.6% initial (weighted) nonresponse rate. Therefore, when

---

than those of the representative, while we cannot make the assumption that the reasons for nonresponse are the same. It might also be the case that answers to the open-ended question by respondent’s representatives on their spouses income is subject to large errors.

16 An alternative choice of sample would be to take the full sample of household representatives (6,816), and allow all those who report not to work for wages and salaries to have zero wages. However, we would be making a distinction between participants and non-participants, which is just another arbitrary definition of a sample. On the other hand, we only pick up those who declare to be asalariate (out of the 4,145) because nonresponse may affect employees differently to self-employed, who might also be given a completely different unfolding bracket design in case of initial nonresponse. Thus, our choice of sample guarantees that all bracket respondents react to the same bid.

17 There were a further 4 individuals who did declare an exact amount of wages/salaries, but the data collection agency decided to classify their responses as error data. These are neither non-response or full respondents and, therefore, we draw them out of the sample. Therefore, the final sample from which to base our illustration consists of 3,598 household representatives.

faced with a question which asks for an specific amount, nonresponse is sufficiently large for it not to be ignored. For this specific variable, the group of initial non-respondents were routed to a sequence of unfolding bracket questions as formulated in (10), with standing bid  $B1 = \$25,000$ . At this initial stage of the unfolding sequence, 119 individuals answered ‘don’t know’ (or ‘refuse’). Thus, the full nonresponse rate is 3.3%. The remaining 329 individuals make up the sample of bracket respondents.

**Table 2: Summary Statistics by different Sample definitions.**

|                             | All sample   | Males        | Females      | Smokers     |
|-----------------------------|--------------|--------------|--------------|-------------|
| <b>Size (%)</b>             | 3,602 (100)  | 1,177 (38.4) | 2,425 (61.6) | 797 (22.1)  |
| <b>Age (s.d)</b>            | 59.15 (3.08) | 59.20 (3.04) | 59.15 (3.10) | 58.70 (2.9) |
| <b>% Smokers (s.e)</b>      | 22.1 (0.7)   | 24.4 (1.3)   | 20.7 (0.8)   | --          |
| <b>Full Respondents (%)</b> | 3,160 (88.4) | 1,065 (91.2) | 2,095 (86.6) | 706 (89.0)  |
| <b>Mean income</b>          | 31,240       | 36,170       | 28,170       | 29,080      |
| <b>(s.d)</b>                | (28,310)     | (28,900)     | (27,450)     | (30,080)    |
| <b>Median</b>               | 25,370       | 32,000       | 23,000       | 24,000      |
| <b>Min/Max</b>              | 0/350,000    | 0/300,000    | 0/350,000    | 0/350,000   |
| <b>Initial Nonresponse.</b> |              |              |              |             |
|                             | 438 (11.5)   | 110 (8.8)    | 328 (13.20)  | 91 (11.0%)  |

Table 2 shows summary statistics for the selected sample of household representatives with non-zero wages and/or salaries. All means, standard deviations and percentages referred to the weighted sample.<sup>18</sup> The average age reflects the initial sample selection criteria for the HRS data set since all households representative are between 55 and 65 years of age. Approximately 22% of the target population smokes, with males showing a slightly higher probability of smoking than females. With respect to income, males are, on average, higher wage/salary earners than females, although such estimate takes into account only income from full respondents, and, therefore, might be a biased estimate. This is specially true since income nonresponse is higher for females than for males, both in absolute values and as respective percentages of their sub-populations. When analysing the population of smokers we see that these are slightly younger than the overall population, even if the difference is not significant, while their income

18 The use of weighs is needed since the HRS is based on various sub-samples representing different groups within the USA population, in particular, African American, Hispanic and other racial minorities. The use of cross-section weights is necessary to obtain a representative sample of the USA population using the different sub-samples. Appendix C examines the characteristics of these weights for the 2<sup>nd</sup> wave of the HRS (1996).

distribution shows a significant shift to the right relative to the male's income distribution.<sup>19</sup>

The unfolding sequence for the wages and salaries question consists of two steps. Those who answer 'yes' to the initial bid of \$25,000 were routed to a second question with bid B21=\$50,000. Those who answered 'no' to the first bid were also routed to a second question, this time with bid B20=\$5,000. In both cases the question was identical to (10) – only the bid changed. If an unfolding bracket design has more than one bid, two possible types of bracket respondents can emerge, namely complete (CBR) and incomplete (IBR) bracket respondents. This is because at each stage of the unfolding sequence respondents can choose to answer 'Don't know' or 'Refuse'.

**Table 3: Categories for Partial Respondents, according to information provided by bracket respondents to the variable annual labour income.**

| Group                                            | Anchor 1:<br>B1 | Answer to<br>Anchor 1 | Anchor 2:<br>B20/B21 | Answer to<br>Anchor 2 | Resulting<br>Categories | ALL           | MALE         | F'MALE        | S'MKRS       |
|--------------------------------------------------|-----------------|-----------------------|----------------------|-----------------------|-------------------------|---------------|--------------|---------------|--------------|
| <b>Complete<br/>Bracket<br/>response (CBR)</b>   | >\$25,000       | Yes                   | >\$50,000            | Yes                   | [\$50,000-<br>infinity) | 30<br>(10.7)  | 10<br>(10.4) | 20<br>(10.8)  | 2<br>(1.6)   |
|                                                  |                 |                       |                      | No                    | [\$25,000-<br>\$50,000) | 86<br>(27.8)  | 26<br>(35.5) | 60<br>(25.0)  | 19<br>(25.8) |
|                                                  |                 |                       |                      |                       |                         |               |              |               |              |
|                                                  |                 | No                    | >\$5,000             | Yes                   | [\$5,000-<br>\$25,000)  | 170<br>(48.6) | 37<br>(42.7) | 133<br>(50.9) | 42<br>(63.7) |
|                                                  |                 |                       |                      | No                    | [\$0-\$5,000)           | 34<br>(10.0)  | 6<br>(4.5)   | 28<br>(12.1)  | 6<br>(8.9)   |
|                                                  |                 |                       |                      |                       |                         |               |              |               |              |
| <b>Incomplete<br/>Bracket<br/>Response (ICB)</b> | >\$25,000       | Yes                   | >\$50,000            | DK/RF                 | [\$25,000-<br>infinity) | 9<br>(2.8)    | 3<br>(6.9)   | 6<br>(1.3)    | 0<br>(-)     |
|                                                  |                 | No                    | >\$5,000             | DK/RF                 | [\$0-\$25,000)          | 0             | 0            | 0             | 0            |

Table 3 shows the distribution of partial respondents, first between complete and incomplete sub-groups, and second, with respect to gender and smoking status. Most of those who are bracket

<sup>19</sup> Perhaps one striking fact in Table 2 is the overrepresentation of females in the sample. Since we are selecting on those who respond to earn wages, one could think that our selection is creating such unbalance between genders. However, as Appendix D in the paper shows, the distribution between males and females is similar in the original sample of household respondents in the survey. Since the HRS is supposed to represent the cohort of US Citizens

respondents are females, who are also the largest initial nonresponse group.

## 4. Estimation Methods and Results

### 4.1 Estimation Methods

The bounds in Section 2 are expressed in terms of population characteristics and can easily be estimated using the corresponding sample analogue. Section 4.2 applies bounds in Section 2 to estimate probabilities of the event ‘smoking’ at different points of the earnings distribution. The target population is a representative cohort of USA citizens born between 1931 and 1941 who claim to have been employed on the basis of wages/salaries during 1995. The data draws from the 1996 wave of the Health and Retirement Study as defined in Section 3.

The bounds are estimated separately for males and females. The results are then used to test for significant difference in the smoking habits between genders, conducting the test at different sections of the income distribution. The conditioning set is composed of two discrete variables, a dummy for gender and a binary variable identifying the even  $A \in [B_j, B_{j+1}]$ , where  $[B_j, B_{j+1}]$  is a range of wages. Thus in this empirical illustration sample estimates do not require smoothing parameters which would be the case if the conditioning set contained continuous variables (see Haedler and Linton (1994) for a detail account of nonparametric regression techniques).

The width between estimates of upper and lower bounds, which depends on  $p$ , for  $p = P(A | NR) \in [0,1]$ <sup>20</sup> – see expressions (8) and (17) in Section 2 –, reflects uncertainty due to item nonresponse. The component of the bounds are simple probabilities so that for known  $p = P(A | NR)$  it is straightforward to derive analytical expressions for their (pointwise) asymptotic distribution, and from these, estimate the sampling error. An equivalent method to estimate confidence bands is to use a naive bootstrap. In this paper such method is used to find confidence bands, re-sampling 500 times (with replacement) from the original data. The

---

born between 1931-1941, this shows that in such cohort the percentage of females dominates that of males.

<sup>20</sup> Estimates of (8) and (17) are the result of minimising with respect to  $P(A|NR)$  and  $P(A|NR) \& P(A|BR)$  respectively. In the case of (8) the procedure involves estimating the upper and lower bound over 100 partitions of the interval  $[0,1]$  which defines the range of possible values of  $P(A|NR)$ . In the case of (17) the minimisation and maximisation of lower and upper bounds, respectively, is also over 100 partitions of the interval  $[0,1]$  for the values of  $P(A|NR)$ , and for each of these 100 estimates, the bounds are estimated over 100 partitions of the possible values of the interval defined by  $P(A|BR)$ .

(pointwise) confidence intervals on the bounds are the 2.5% and 95.5% percentiles of the 500 estimates. Each of the figures presented in Section 4.2 shows the lower confidence band for the lower bound and the upper confidence band for the upper bound. The gap between these reflects both the uncertainty due to sampling error and the uncertainty due to item nonresponse.

## 4.2 Bounds on the Probability of Smoking

This section presents the result of estimating bounds as defined in Section 2, on the probability of smoking conditional on income and gender, where income refers to annual labour income. For each of the sub-populations of males and females, the results show sample estimates of the probability of smoking grouping individuals in four sub-categories according to whether information on their annual earnings can be classified in  $[0, \$5,000)$ ,  $[\$5,000, \$25,000)$ ,  $[\$25,000, \$50,000)$  or  $[\$50,000, \text{maximum}]$ . The classification reflects the structure set by the unfolding bracket design answered by initial non-respondents to wages and salaries (see Section 3, Table 3).

Under the assumption of exogeneity full respondents would be a representative sample of the population, thus one can estimate sample probabilities throwing away gender and other information on non-respondents to income. Effectively this implies that  $P(NR | A) = 0$  in (1), and, therefore,  $P(\text{smoking} | A, \text{gender}, FR)$  is equivalent  $P(\text{smoking} | A, \text{gender})$ . These estimates are shown in Table 4<sup>21</sup>, while Table 5 provides a graphic interpretation of the same results. Table 4 (row 1), shows that for the full population, the probability of smoking decreases as income increases; individuals in the lowest income are 7% more likely to be smokers than individuals in the highest brackets. Figure 1 in Table 5 illustrates graphically the same conditional (on income) monotonic decrease on the probability of smoking. Comparing rows 2 and 3 in Table 4 (cf. Figures 2 and 3) shows that the smoking pattern for the population as a whole is driven by the smoking pattern of females; for these, the probability of being a smoker decreases significantly as income increases, showing a 12.4% increase difference on the probability of smoking between low and high income brackets. On the other hand, estimates in Table 4 (row 2) show that the smoking habits of males are invariant to increasing levels of income.

---

21 Both test statistics in Table 4 are based on the absolute difference between two probabilities normalised by the corresponding standard error which equals the square root of the sum of the square of the standard errors for each probability (given independent samples). Under the hypothesis of zero difference between the probabilities, the test

**Table 4: Probability (s.e) of smoking by gender and income bracket: Random Nonresponse**

|                         | Wages in<br>[0-\$5,000) | Wages in<br>[\$5,000-\$25,000) | Wages in<br>[\$25,000-\$50,000) | Wages in<br>[\$50,000-max) | t-test(2)    |
|-------------------------|-------------------------|--------------------------------|---------------------------------|----------------------------|--------------|
| <b>All</b>              |                         |                                |                                 |                            |              |
| Point estimate<br>(s.e) | 0.248<br>(0.032)        | 0.246<br>(0.015)               | 0.215<br>(0.015)                | 0.178<br>(0.021)           | <b>1.83</b>  |
| <b>Males</b>            |                         |                                |                                 |                            |              |
| Point estimate<br>(s.e) | 0.210<br>(0.059)        | 0.248<br>(0.026)               | 0.251<br>(0.024)                | 0.211<br>(0.032)           | <b>0.015</b> |
| <b>Females</b>          |                         |                                |                                 |                            |              |
| Point estimate<br>(s.e) | 0.264<br>(0.038)        | 0.245<br>(0.018)               | 0.189<br>(0.018)                | 0.140<br>(0.025)           | <b>2.73</b>  |
| <b>t-test(1)</b>        | <b>0.769</b>            | <b>0.094</b>                   | <b>2.07</b>                     | <b>1.75</b>                |              |

Notes 1: All sample estimates are weighted using cross-section weights as provided by the HRS data set, 1996.

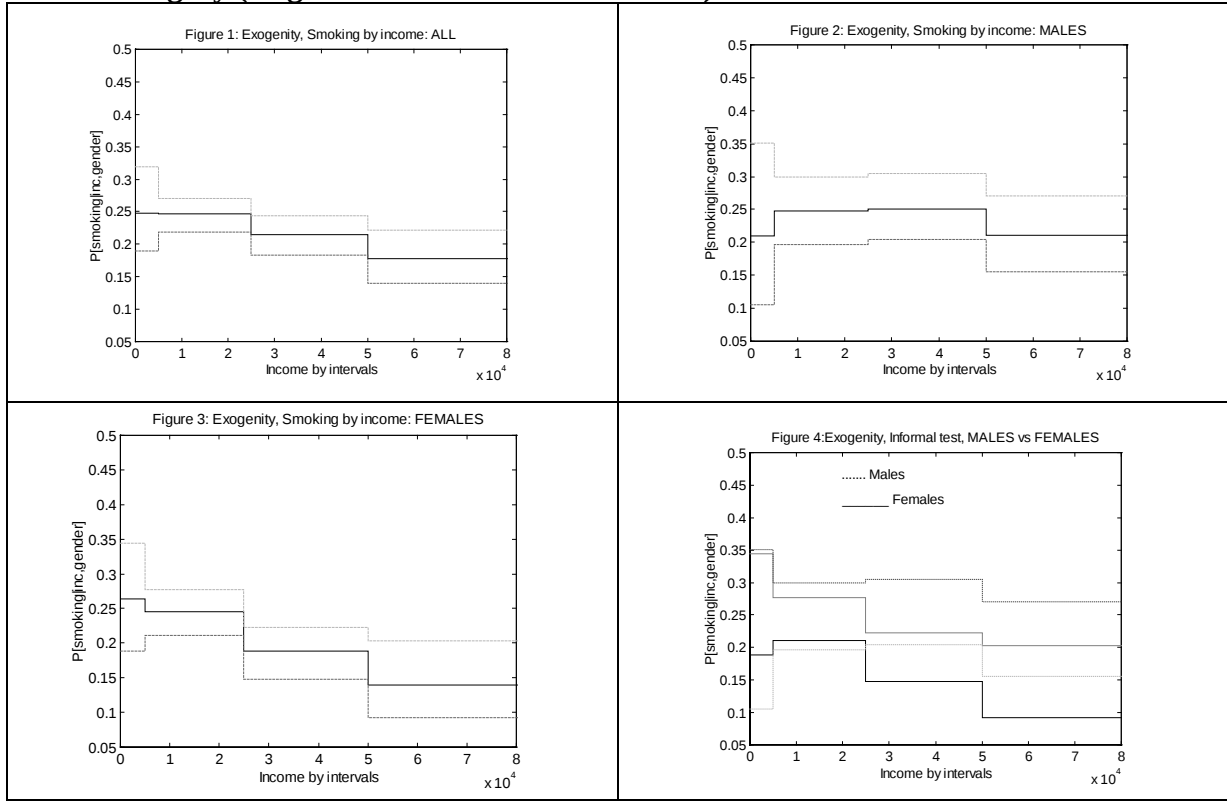
Notes 2: The t-test in Column 5 refers to testing entries for significant difference between the lowest and highest income bracket (column 2 versus column 5), whereas the t-test in Row 5 refers to testing for significant difference between males and females at each wage region (bracket).

With this, results in Table 4 (row 5) indicate a significant difference in the probability of smoking between males and females for individuals with annual earning of \$25,000 and above, with the probability that a male smokes been at least 6% higher than for a female. For annual labour income below \$25,000 the smoking habits of males and females is not significantly different, with approximately between 21% to 25% been smokers in both sub-populations.

---

is asymptotically standard normal.

**Table 5: Figures 1-4 show bounds on probabilities of smoking by gender and income category (wage income in 1996 USA dollars).**



Note 1: Figures 1 to 3 show pointwise probability estimates with a solid line, and 95% confidence bands with dotted lines.

Note 2: Figure 4 compares estimates of males and females using the 95% confidence interval of Figures 2 and 3.

Results so far have assumed that initial nonresponse to income happens at random. The next set of estimates relaxes this exogeneity assumption and shows the consequence of estimating bounds as given in (8). These bounds account for the initial sample (income) nonresponse rate. Table 6 shows these estimates at the different intervals of the income distribution and Table 7 illustrates the same results graphically. Each cells in Table 6 reports estimated upper and lower bounds as well as upper and lower confidence bands in the bounds. For example, allowing for nonresponse in income, the probability of smoking for a female in the lowest income bracket (\$0 to \$5,000) is bounded between 11.3% and 43.3%. If we further allow for uncertainty due to sampling error the estimate is between 7.1% and 54.5%, with 95% confidence. Comparing estimates in Columns 2 and 5 for any of the sample definitions (all, males or females) shows that the 95% confidence bands for the lowest income bracket always overlaps significantly (or nests) with those in the highest income bracket. Thus, for this particular sample, and once we allow for nonresponse on wages, the evidence cannot reject the null of no difference in the smoking patters

between low and high earners. This already contradicts the conclusions based on Tables 4-5, where the assumption of exogeneity lead to a positive and significant difference in smoking habits of low earners relative to high earners for both the sub-population of females and for the sample as a whole. Estimates of bounds in Table 6 also allow to compare the smoking habits between males and females. Table 6 shows that the estimated bounds between males and females overlap at all levels of income, with the upper bound on the probability of smoking for females always above that of the lower bound on the probability of smoking for males. Table 7 (Figure 8) shows this same result but with 95% confidence. The existence of an overlapping region between the estimated bounds in the two sub-populations is due to the existence of income nonresponse. Under nonrandom nonresponse, the true probabilities are unknown. What we know is that they could fall in the overlapping region, thus inspection of this overlap between estimated bounds becomes an informal test for the difference in the probabilities of the two sub-populations. This informal test suggest that the null of equality in the probabilities of smoking between genders cannot be rejected. This, again, contradicts the results in Table 4 which are based on the rather strong assumption of exogeneity. The ‘informal’ test can be formalised with what is effectively a sample dependent t-test, the result of which is shown in Table 6 (final row). In this particular case, since the estimated upper bound for the sample of females is always above the estimated lower bound on the probability of smoking for the sample of males, the null of no difference between the smoking habits of the two genders cannot be rejected if there is a positive and significant difference between pointwise estimates on the upper bound for females and the lower bound for males. Results in Table 6, final row, show that in fact the null of no difference on the probability of smoking between genders cannot be rejected throughout the distribution of income. Figure 8 (Table 7) further reinforces this argument by showing that, with 95% confidence, the bounds on the probability of smoking for males nests those of the estimates for the female population.

**Table 6: Probability (s.e) of smoking by gender and income bracket: Worst Case bounds in the presence of income nonresponse: No bracket response.**

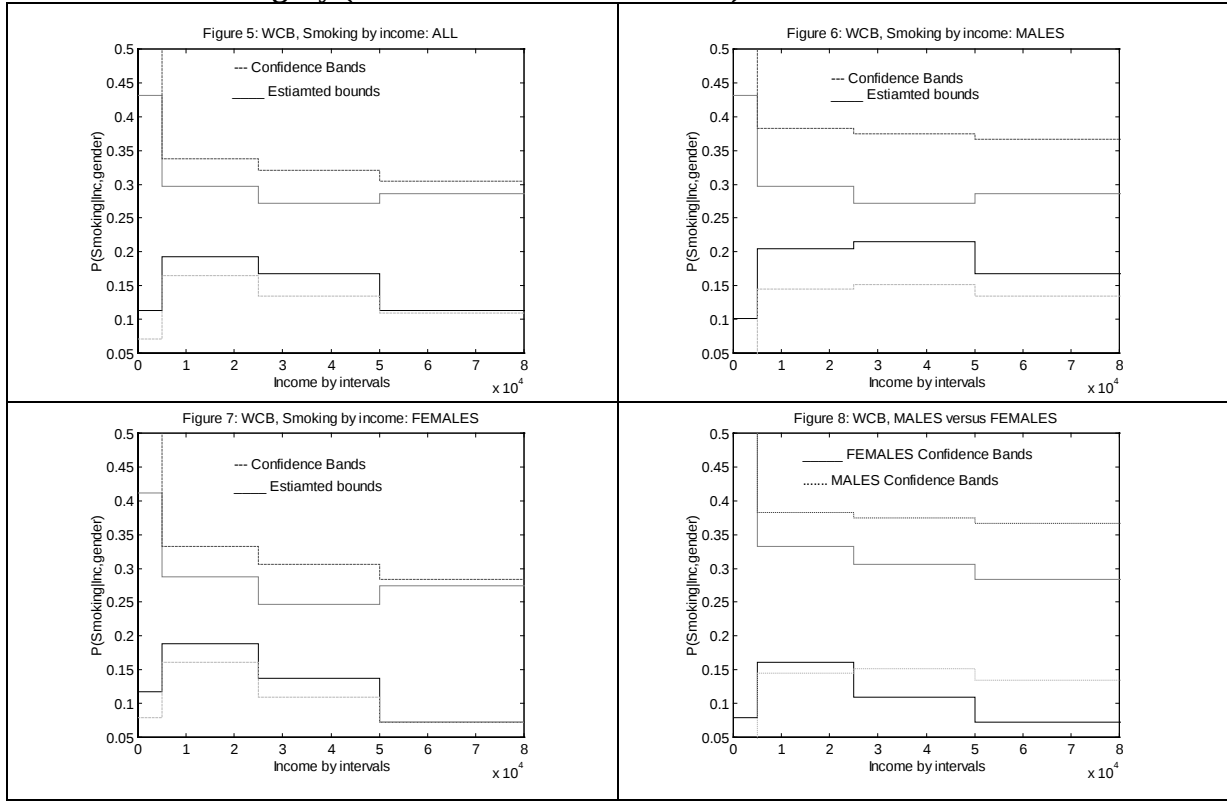
|                       | Wages in             | Wages in             | Wages in             | Wages in             | t-test(2) |
|-----------------------|----------------------|----------------------|----------------------|----------------------|-----------|
|                       | [0-\$5,000)          | [\$5,000-\$25,000)   | [\$25,000-\$50,000)  | [\$50,000-max)       |           |
| All                   |                      |                      |                      |                      |           |
| Estimated bounds      | 0.113 – 0.4331       | 0.193 – 0.297        | 0.167 – 0.272        | 0.114 – 0.286        | Overlap   |
| Confidence bands      | (0.071: 0.545)       | (0.165 :0.338)       | (0.135 : 0.320)      | (0.110 : 0.305)      |           |
| Males                 |                      |                      |                      |                      |           |
| Estimated bounds      | <u>0.102</u> – 0.470 | <u>0.204</u> – 0.315 | <u>0.215</u> – 0.305 | <u>0.167</u> – 0.295 | Overlap   |
| Confidence bands      | (0.031 : 0.622)      | (0.145 : 0.383)      | (0.152 : 0.375)      | (0.135 : 0.367)      |           |
| (s.e for lower bound) | (0.025)              | (0.021)              | (0.021)              | (0.022)              |           |
| Females               |                      |                      |                      |                      |           |
| Estimated bounds      | 0.117 – <u>0.412</u> | 0.188 – <u>0.287</u> | 0.187 – <u>0.246</u> | 0.073 – <u>0.274</u> | Overlap   |
| Confidence bands      | (0.079 : 0.514)      | (0.161 : 0.333)      | (0.109 : 0.306)      | (0.072 : 0.283)      |           |
| (s.e for lower bound) | (0.039)              | (0.019)              | (0.023)              | (0.022)              |           |
| t-test(1)             | 6.692                | 2.93                 | 2.66                 | 3.44                 |           |

Note 1: See Note1, Table 4

Note 2: The bracket number in the cells for males corresponds to the estimate of the standard error for the (underlined) estimated lower bound. As with the confidence intervals, this standard errors are attained using a bootstrap techniques which consists on re-sampling the original data 500 times. The bracketed number in the cell for females is also the estimate standard error, but this time corresponding to the (underlined) estimate of the upper bound.

Note 3: Inspection of estimated upper and lower bounds suggest an overlap between the identification regions of for the probability of smoking of males and females. The overlap occurs because at each point estimate (defined by the partition of the wage distribution), the upper bound on the probability of smoking for females,  $\hat{p}_f$  is always above that of the estimates of the lower bound on the probability of smoking for males,  $\hat{p}_m$ . The one-sided sample-dependent t-test in Row 5 is based on  $(\hat{p}_f - \hat{p}_m) / \hat{\sigma}$ , where  $\hat{\sigma}$  refers to the pool standard error estimate between independent population.

**Table 7: Figures 5-8 show Worst Case Bounds on probabilities of smoking by gender and income category (income in 1996 USA dollars).**



Note 1: Figures 5 to 7 show estimates of upper and lower bounds with two solid line, whereas the dotted lines are the estimated 95% upper confidence band on the upper bound and 95% lower confidence band on the lower bound.

Note 2: Figure 4 compares estimated regions of identification of the unknown probabilities of males and females using the 95% confidence interval of Figures 6 and 7.

Whereas estimates in Table 6 (and Table 7) allow for any type of nonrandom nonresponse in the variable wages, such estimated worst case bound do not account for partial information provided by bracket respondents. Expression (17) allows for such information thus leading to more informative bounds on the conditional probability of smoking. Table 8 shows analogous results to Table 6 but with information from partial respondents incorporated. Table 9 illustrates the results with Figure 12 comparing bounding intervals between genders with 95% confidence. Both Table 8 and Table 9 ignore the possible existence of the Anchoring effect on bracket respondents. Because estimates of (17) contain more information (and lower full nonresponse rate) than estimates on (8), the estimated regions between paired upper and lower bounds are narrower in Table 6 (Table 7) than those in Table 8 (Table 9). This results in bounds

which are more informative with respect to the unknown probabilities of smoking for either the full sample, or any of the gender dependent sub-samples. For example, estimates in Table 6 suggest that, in the presence of wage nonresponse, a female wage earner who earns between \$5,000 and \$25,000 has an estimated 16.1% to 33.3% chance of been a smoker, with 95% confidence. Once the bounds account for partial response information the same probability is now bounded between 19.0% and 28.4%, also with 95% confidence. This implies a 45% improvement in the identification region of the unknown probability.

**Table 8: Probability (s.e) of smoking by gender and income bracket: Worst Case bounds with bracket respondents. Case of No Anchoring.**

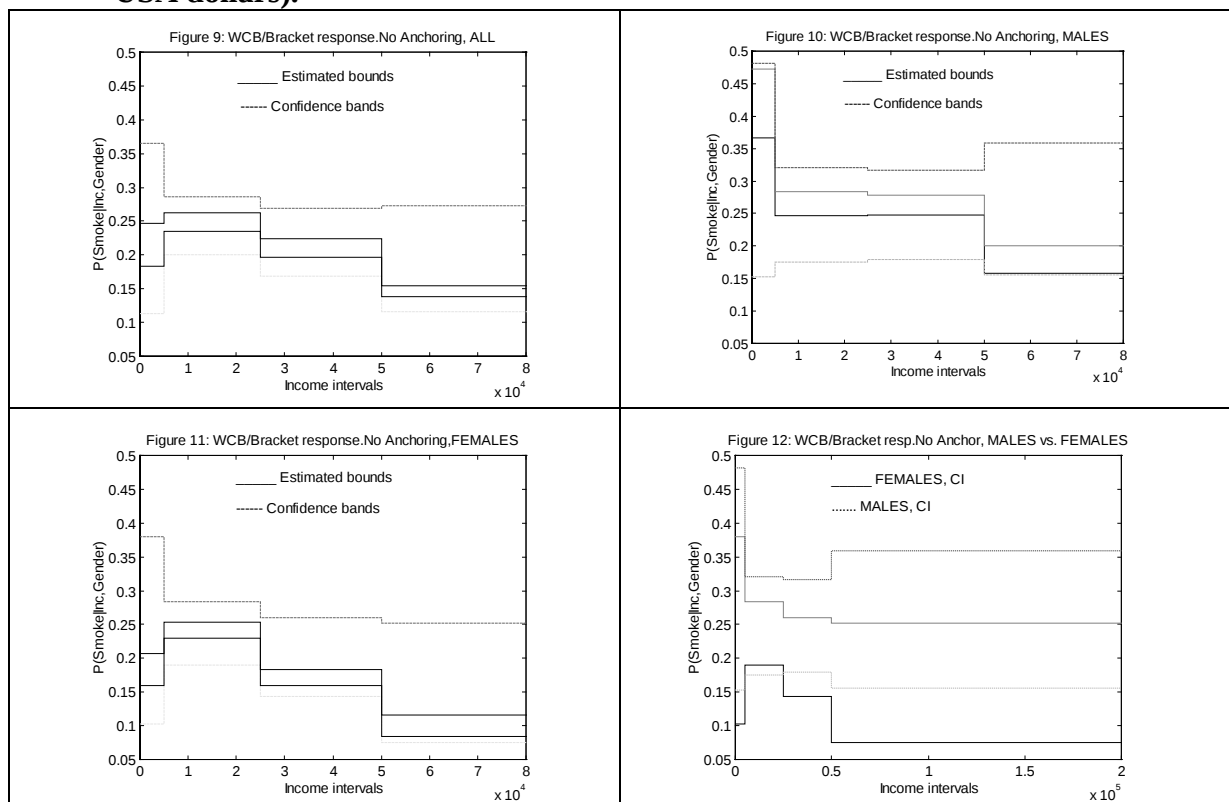
|                       | Wages in<br>[0-\$5,000) | Wages in<br>[\$5,000-\$25,000) | Wages in<br>[\$25,000-\$50,000) | Wages in<br>[\$50,000-max) | t-test(2)      |
|-----------------------|-------------------------|--------------------------------|---------------------------------|----------------------------|----------------|
| <b>All</b>            |                         |                                |                                 |                            |                |
| Estimated bounds      | 0.183 – 0.247           | 0.235 – 0.262                  | 0.197 – 0.224                   | 0.138 – 0.154              | <b>Overlap</b> |
| Confidence bands      | (0.114 : 0.365)         | (0.201 : 0.286)                | (0.169 : 0.269)                 | (0.116 : 0.273)            |                |
| <b>Males</b>          |                         |                                |                                 |                            |                |
| Estimated bounds      | <u>0.367</u> – 0.472    | <u>0.247</u> – 0.283           | <u>0.248</u> – 0.278            | <u>0.158</u> – 0.200       | <b>Overlap</b> |
| Confidence bands      | (0.153 : 0.481)         | (0.175 : 0.320)                | (0.179 : 0.317)                 | (0.155 : 0.359)            |                |
| (s.e for lower bound) | (0.082)                 | (0.028)                        | (0.027)                         | (0.042)                    |                |
| <b>Females</b>        |                         |                                |                                 |                            |                |
| Estimated bounds      | 0.160 – <u>0.207</u>    | 0.230 – <u>0.253</u>           | 0.160 – <u>0.183</u>            | 0.084 – <u>0.116</u>       | <b>Overlap</b> |
| Confidence bands      | (0.103 : 0.380)         | (0.190 : 0.284)                | (0.144 : 0.260)                 | (0.075 : 0.252)            |                |
| (s.e for lower bound) | (0.058)                 | (0.018)                        | (0.023)                         | (0.035)                    |                |
| <b>t-test(1)</b>      | <b>-1.59</b>            | <b>0.180</b>                   | <b>-1.83</b>                    | <b>-0.768</b>              |                |

See Note 1 and Note 2, Table 6

The two one sided test in Table 8 are based on the same sample inspection as in Table 6. Column 6 in Table 8 shows that the identification region on the probability of smoking for the higher earners is nested in the identification region of low earners, thus, the data cannot reject the null of equality between high and low income earners within sub-samples. On the other hand, allowing for bracket respondents reduces the identification region of the unknown probabilities for each sub-sample and at each income level. The result is a substantial reduction of the overlapping region between the bounding intervals of males and females, relative to the overlap observed in the case where partial information with bracket respondents is ignored (i.e., estimates in Table 6). The final row in Table 8 shows the result of the sample dependent one-sided t-test (see Note 2, Table 6) which tests the null of no difference between males and females with

respect to the probability of smoking. The test shows that, for all income intervals, the difference is either close to zero or negative, thus demonstrating that once we allow for partial information there is no overlapping region between the two sets of estimated bounds: the probability that a male smokes is well above the probability that a female smokes, and the null of no difference between the smoking habits of the two genders is not supported by the data. The illustration of results in Table 8 are given Table 9, Figures 9 to 12: comparison of these figures with those in Table 7 shows that for any of the sub-samples considered, the horizontal distance between sets of upper and lower bounds is narrower, thus each set becomes more informative with respect to the unknown probability at each income level. Figure 12, relative to Figure 6, shows that the distance between the upper confidence band on the upper bound for females is now horizontally closer to the lower confidence band on the lower bound estimate for the sub-sample of males. This reflects, allowing for 95% confidence, the reduced overlap between regions of identification for the unknown probabilities.

**Table 9: Figures 9-12 show Worst Case Bounds allowing for Bracket Respondents, on the probabilities of smoking by gender and wage income category (wage income in 1996 USA dollars).**



See Note 1 and Note 2, Table 7.

Estimates in Table 8 (Table 9) ignore the possibility that responses from bracket respondents might be subject to the anchoring effect. In the final set of estimates, Table 10 shows estimates of expression (17) modified according to (25) and (31). Thus, these set of estimates allows for a specific interpretation of anchoring based on Jacowitz and Kahneman (1995).<sup>22</sup>

**Table 10: Probability (see) of smoking by gender and wage bracket: Worst Case bounds with bracket respondents. Allow for the Anchoring Effect.**

|                       | Wages in<br>[0-\$5,000) | Wages in<br>[\$5,000-\$25,000) | Wages in<br>[\$25,000-\$50,000) | Wages in<br>[\$50,000-max) | t-test(2)      |
|-----------------------|-------------------------|--------------------------------|---------------------------------|----------------------------|----------------|
| <b>All</b>            |                         |                                |                                 |                            |                |
| Estimated bounds      | 0.171 – 0.396           | 0.213 – 0.295                  | 0.173 – 0.260                   | 0.114 – 0.231              | <b>Overlap</b> |
| Confidence bands      | (0.119 : 0.469)         | (0.183 : 0.338)                | (0.142 : 0.304)                 | (0.110 : 0.261)            |                |
| <b>Males</b>          |                         |                                |                                 |                            |                |
| Estimated bounds      | 0.154– 0.388            | 0.218 – 0.315                  | 0.216 – 0.296                   | 0.173 – 0.264              | <b>Overlap</b> |
| Confidence bands      | (0.052 : 0.567)         | (0.163 : 0.380)                | (0.155 : 0.353)                 | (0.141 : 0.324)            |                |
| (s.e for lower bound) | (0.040)                 | (0.023)                        | (0.022)                         | (0.023)                    |                |
| <b>Females</b>        |                         |                                |                                 |                            |                |
| Estimated bounds      | 0.177 – 0.397           | 0.210 – 0.286                  | 0.145 – 0.287                   | 0.075 – 0.201              | <b>Overlap</b> |
| Confidence bands      | (0.123 : 0.470)         | (0.177 : 0.336)                | (0.115 : 0.233)                 | (0.061 : 0.239)            |                |
| (s.e for lower bound) | (0.035)                 | (0.019)                        | (0.020)                         | (0.019)                    |                |
| <b>t-test(1)</b>      |                         |                                |                                 |                            |                |

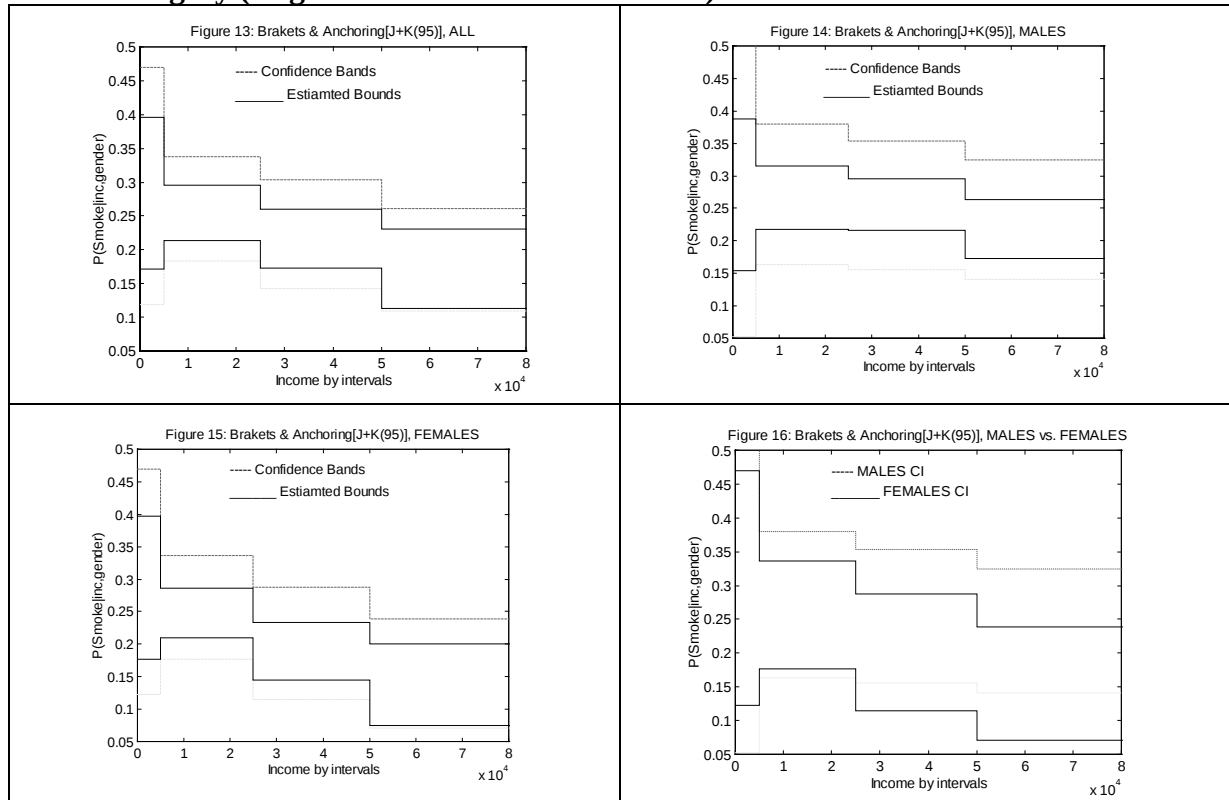
See Note 1 and Note 2, Table 6.

Comparing estimates in Table 8 to those in Table 10 shows that allowing for anchoring increases uncertainty as it widens the distance between estimated pairs of bounds. This is because estimates which allow for anchoring imply a weaker assumption on the sampling process. Following a previous example, but allowing for anchoring effects, the probability of smoking for a female with wages in the \$5,000 – \$25,000 interval is bounded between 17.7% and 33.6%. Although this represents an 8% improvement with respect to the information provided by worst case bounds without partial information (see Table 6), relative to bounds where anchoring was ignored, the identification region has widened and the information on the unknown probability for this particular example is reduced by approximately 40% (i.e., compare bounds in Table 8 to those in Table 10). Table 11, analogous to Table 7 and Table 9, illustrates the information in Table 10. The effect of allowing for anchoring when estimating bounds on the probability of smoking is clear when we compare figures between Tables 9 and 11, since for any given sample

<sup>22</sup> Refer to Appendix B for alternative interpretations of the Anchoring Effect.

definition the bounds are always wider for any of the income intervals, both in terms of point estimates as well as allowing for 95% confidence bands.

**Table 11: Figures 13-16 show Worst Case Bounds allowing for Bracket Respondents and the Anchoring Effect, on the probabilities of smoking by gender and wage income category (wage income in 1996 USA dollars).**



See Note1 and Note2, Table 7.

## 5. Conclusions

The approach by Horowitz and Manski (1998) deals with regressors nonresponse allowing for a flexible and intuitive tool to estimate parameters of interests, while avoiding the often strong (and non-testable) assumptions associated with parametric and semi-parametric methods. Their method in the presence of regressors (item) nonresponse allows for the identification of parameters of interest up to a bounding interval, as oppose to the identification of a point estimate. How informative a bounding interval depends on the degree of sample nonresponse. In this paper bounds in the presence of regressors nonresponse are extended in two directions. First, the paper shows how to derive bounds in the presence of partial respondents to categorical questions. These questioning strategy is often employed by survey designers to reduce item

nonresponse that arises from an initial open-ended question. Second, the paper deals with a type of bias known as the anchoring effect which the psychological literature has shown to arise in the presence of certain types of categorical questions, namely, unfolding bracket designs.

The theoretical framework is illustrated by testing for significant difference in the smoking pattern of males and females for different levels of wages. Thus, nonresponse occurs with respect to wages declared by individual who claim to have worked for wages and salaries over a particular calendar year. The data draws from the second wave of the Health and Retirement Study ((HRS, 1996), which has as target population the cohort of USA citizens born between 1931 and 1941.

If estimates of the probability of smoking assume random nonresponse, difference in the smoking behaviour between males and females is detected for those with wages equal to or above \$25,000, with evidence suggesting that at such level of income males have a significantly higher probability of smoking than females. Relaxing the rather strong assumption of exogeneity leads to estimates of bounding intervals for each of the gender sub-samples. The bounds create a region of identification for the unknown probability of smoking in the presence of income nonresponse. Because the two estimated sets of bounds for males and females overlap, and the overlapping region is not significantly different than zero, a relax of the exogeneity assumption implies that the null of equality in the probability of smoking between males and females cannot be rejected. Such conclusion affects all income intervals. Likewise, the data cannot reject the equality in the probability of smoking between high and low earners for any of the sub-samples considered. Once bounds incorporate information from partial respondents to a categorical question the result changes under the null of no anchoring effects. However, if anchoring effects are allowed for, the presence of partial information does not imply an improvement on the identification region for the unknown probabilities of smoking, for either of the sub-samples considered, relative to the identification provided by worst case bounds without bracket respondents.

## Reference

Green, D., K. Jacowitz, D. Kahneman and D. McFadden (1998), Referendum contingent valuation, anchoring, and willingness to pay for public goods, *Resource and Energy Economics* 20, 85-116.

Heckman, J.J., (1979), Sample selection bias as a specification error, *Econometrica*, 47, 153-161.

Herriges, J. and J. Shogren (1996), Starting point bias in dichotomous choice valuation with follow-up questioning, *Journal of Environmental Economics and Management*, 30, 112-131.

Horowitz, J. and C.F Manski (1995), Identification and Robustness with contaminated and corrupted data, *Econometrica*, 63, 281-302.

Horowitz, J. and C.F Manski (1998), Censoring of outcomes and regressors due to survey nonresponse: identification and estimation using weights and imputation, *Journal of Econometrics*, 95, 415-442.

Hurd, M., D. McFadden, H. Chained, L. Gan, A. Merrill and M. Roberts, (1998), Consumption and savings balances of the elderly: Experimental evidence on survey response bias, in D. Wise (ed.), *Frontiers in the Economics of Ageing*, University of Chicago Press, Chicago, 353-387.

Jacowitz, K. and D. Kahneman (1995), Measures of anchoring in estimation tasks, *Personality and Social Psychology Bulletin*, 21, 1161- 1166.

Manski, C.F., (1989), Anatomy of the selection problem, *Journal of Human Resources*, 24, 343-360.

Manski, C.F., (1990), Nonparametric bounds on treatment effects, *American Economic Review, Papers & Proceedings*, 80, 319-323.

Manski, C.F., (1995), Identification problems in the social sciences, *Harvard University Press*.

Rabin, M. (1998), Psychology and Economics, *Journal of Economic Literature*, March, 11-46.

Vazquez-Alvarez, R., B. Melenberg and A. vanSoest (2001), Nonparametric bounds in the

presence of item nonresponse, unfolding brackets, and Anchoring, CentER discussion Paper 2001-67, Tilburg University (in second revision for the *JASA, P&P*).

Vella, F. (1998), Estimating models with sample selection bias, *Journal of Human Resources*, 33, 127-172.

# Appendix A

Section 2 extends the framework in Horowitz and Manski (1998) to allow for bracket respondents in a bounding interval that accounts for regressors nonresponse. The paper by Horowitz and Manski (1998) derives bounds for a further two cases, that of joint covariates and dependent variable nonresponse, and the case with mixed nonresponse (i.e., only some covariates and the dependent variable suffer nonresponse). This appendix provides some guidelines and a summary of how to deal with joint and mixed nonresponse in case of the presence of a bracket response sub-population and the anchoring effect.

To motivate this section we look at two possible examples. Example number 1 assumes we want to estimate  $P(\text{smoking} | \text{income}, \text{gender})$ , where the dependent variable is a binary outcome while the covariate income is continuous, and it is assumed that nonresponse can affect both. Example number 2 would be if we want to estimate  $P(\text{income}_w \leq t | \text{income}_h \leq t)$ , that is, the distribution of wife's income conditional on husband's income ('w' and 'h' stand for wife and husband, respectively). The first example illustrates a situation where the possibility of partial information affects only the regressor. In the second example bracket responses can affect both the dependent variable and the regressor.

## A.1 Joint censoring of outcomes and regressors

If nonresponse is unit nonresponse, this means that a nonrandom percentage of the sample refuses to participate in the survey, so that nonresponse affects all questions. If the survey is at individual level, unit nonresponse implies nonresponse for all variables, including the binary variable 'smoking' and the variable income. If unit nonresponse is at household level information would be missing for all household earners, such that nonresponse for husbands/wives would automatically imply nonresponse for their spouses. This is the case treated by Horowitz and Manski (1998). In such case the sub-sample of initial non-respondents are full non-respondents, since the strategy of providing unfolding bracket designs to elicit partial information is not relevant to the problem of missing information. Thus, bounds as derived in Section 3 of Horowitz and Manski (1998) apply.

## A.2 The mixed case: Censoring of outcomes and regressors

When item nonresponse affects more than one variable relevant in the estimation process, three possible nonresponse combinations arise within one sample. Bounds derivation need to account for all three possibilities. One of these might be that some individuals do not respond to the dependent variable of interest while regressor's information is complete. With respect to example number 1 this would imply observing income for some individuals who do not declare their smoking status, whereas for the example number 2 it would imply to observe income for some husbands whose wives decide to provide no information on income. The second possibility is the reverse, with some individuals who do not respond to the one or regressors while, for this same group, information on the dependent variable is fully observed; Section 2 assumed this was the only possible situation with respect to nonresponse. Finally, some individuals might be such that information is not observed either for the outcome variable or (at least one) regressor. Although this could mean unit nonresponse, it might also be the case that respondents participate in the survey but do not provide information on those particular variables of interest, for example, while they might declare their gender and other social-economic variables important for the measure of interest, they do not provide information to either income or smoking status. There is an important distinction between this possibility and those who are defined as unit nonresponse, since as long as people are willing to be active participant in the survey, and in the case of continuous variables, it might always be possible to get partial information for the missing value, and, therefore, bounds need to account for such sub-population of bracket response information. On the other hand, bracket response is not relevant in the case of unit nonresponse.

Together with the three nonresponse combinations above, any sample will also comprise full respondents. To provide guidelines on deriving bounds in the mixed case, consider the partition of the population into four groups, namely full respondents (G1), outcome non-respondents (G2), regressors non-respondents (G3) and non-respondents to both outcome and regressors (G4). We assume that in  $(A, x)$ , the conditioning set, only one covariate ( $A$ ) is affected by nonresponse. With this, the measure of interest  $P(y | A, x)$  can be partition as follows:

$$P(y | A, x) = P(y | A, x, G1) \times P(G1 | A, x) + P(y | A, x, G2) \times P(G2 | A, x) + P(y | A, x, G3) \times P(G3 | A, x) + P(y | A, x, G4) \times P(G4 | A, x) \quad (A.2.1)$$

where  $P(y | A, x, G2)$ ,  $P(y | A, x, G3)$ ,  $P(y | A, x, G4)$ ,  $P(G3 | A, x)$  and  $P(G4 | A, x)$  are not identified by the data. Without bracket respondents the anchoring effect is not relevant. In such case Section 5 in Horowitz and Manski (1998) provide a particular example of how to bound (A.2.1) where the sub-population  $G1$  is not considered. However, if we consider a situation where initial non-respondents might be routed to an unfolding bracket design, bounds need to consider both the possibility of bracket respondents and that of the anchoring effect. In what follows we distinguish between two types of outcome, binary (so that  $y = (0,1)$ ) and continuous  $y$ .

### **Case 1: Binary Outcome**

When nonresponse affects the outcome but this is a binary variable, the strategy of providing initial non-respondents (to the outcome) with a categorical question is not relevant, so that bracket response is not an issue in the presence of a binary outcome (as would be the case in Example 1). For this reason, the only information on  $P(y | A, x, G2)$  and  $P(y | A, x, G4)$  allowed by the data generating process is that  $P(y | A, x, G2) \in (0,1)$  and  $P(y | A, x, G4) \in (0,1)$ .<sup>23</sup> On the other hand, the data might be more informative about the measure  $P(y | A, x, G3)$  than simply bounding it between the  $(0,1)$  interval, because for the sub-population  $G3$  the problem of missing information affects the conditioning set and not the outcome. It is easy to see that framework of Section 2 is applicable to  $P(y | A, x, G3)$ . Allow for the sub-population of  $G3$  to be partition between bracket respondents (BR) and full non-respondents (NR), where response refers to the variable  $A$  in the conditioning set. Then  $P(y | A, x, G3)$  can be partition as follows:

$$P(y | A, x, G3) = P(y | A, x, G3, BR)P(BR | A, x, G3) + P(y | A, x, G3, NR)P(NR | A, x, G3) \quad (A.2.2)$$

---

<sup>23</sup> In this appendix the variable 'A' is considered to be a continuous variable where bracket response applies. In the case of  $P(y|A,x,G4)$ , given that for this sub-group the problem of nonresponse affects both the outcome and variable 'A', we can consider the case where the binary 'y' is missing and information on 'A' is in the form of bracket response to an unfolding bracket design, thus the data reveals  $A \in [B_j, B_{j+1}]$ . In this case we still have that  $P(y|A,x,G4) \in (0,1)$  because full knowledge does not change the information on the probability of the outcome.

where,

$$P(BR | A, x, G3) = \frac{P(A | BR, x, G3)P(BR | x, G3)}{P(A | BR, x, G3)P(BR | x, G3) + P(A | NR, x, G3)P(BR | x, G3)}$$

and

$$P(NR | A, x, G3) = \frac{P(A | NR, x, G3)P(NR | x, G3)}{P(A | BR, x, G3)P(BR | x, G3) + P(A | NR, x, G3)P(BR | x, G3)} \quad (A.2.3)$$

Therefore, bounds on  $P(y | A, x, G3)$  follow easily applying Section 2 to (A.2.2) and (A.2.3). If we ignore anchoring, expression  $P(y | NR, x, G3, A)$  and  $P(y | BR, x, G3, A)$  in (A.2.2) are given by (4) and (16), respectively. Likewise,  $P(A | BR, x, G3)$  in (A.2.3) can be bounded as in (15) while the unknown  $P(A | NR, x, G3)$  is in the (0,1) interval. Modification of (17) to incorporate the four sub-populations (G1 to G4) leads to a bounding interval on the left hand side of (A.2.1). Allowing for anchoring bias, and assuming the Jacowitz and Kahneman (1995) model of anchoring, implies following a similar set of guidance but using expressions (25) and (31) instead of (15) and (16), respectively, while expression (4) and the interval (0,1) remain relevant for the measures  $P(y | NR, x, G3, A)$  and  $P(A | NR, x, G3)$ , respectively.

## **Case 2: Continous Outcome**

In this case bracket response and, therefore, anchoring, can affect not just the  $P(y | A, x, G3)$ , but also  $P(y | A, x, G2)$ ,  $P(y | A, x, G4)$ , and  $P(G4 | A, x)$ .

The case of how to bound  $P(y | A, x, G2)$  is extensively treated in Vazquez-Alvarez, R., B. Melenberg and A.vanSoest (2000), both allowing and not allowing for anchoring, and with alternative explanations of anchoring, including that provided by Jacowitz and Kahneman (1995). On the other hand, because  $P(y | A, x, G3)$  is affected by nonresponse only in the conditioning set, the same arguments as in Case 1 apply to bound such expression, either allowing or not for the anchoring effect. Thus, in Case 2, it only remains to indicate how to provide bounds on  $P(y | A, x, G4)$  and  $P(G4 | A, x)$ . This situation is the same as in example 2, where the interest was to find the distribution of a wife's income conditional on the earnings of the spouse. Using this example, if income is missing for both spouses, but they remain part of the sample – say, they provide information on any of the other variables in  $x$ , such as gender – their information

cannot be treated as unit nonresponse. They might provide partial information following an unfolding bracket design. Ignoring the anchoring effect, bounds on  $P(y | A, x, G4)$  will be attained starting by expressing a partition for  $P(y | A, x, G4)$  similar to that given in (A.2.2), although in this case the partition would have to account for more sub-populations than the basic bracket and full nonresponse. To simplify the exposition we can assume that a household respondent answers to both outcome and regressors, with those who provide partial information does so for all initial missing values. With this,

$$P(y | A, x, G4) = P(y | A, x, G4, BR)P(BR | A, x, G4) + P(y | A, x, G4, NR)P(NR | A, x, G4) \quad (A.2.4)$$

where the measure  $P(y | A, x, G4, NR)$  is bounded in the (0,1) interval. For bracket respondents the probability  $P(y | A, x, G4, BR)$  will be an estimate given by the events  $y \in [B_j, B_{j+1}]$  and  $A \in [K_g, K_{g+1}]$  where  $[B_j, B_{j+1}]$  ( $[K_g, K_{g+1}]$ ) defines a possible category in as many as  $j$  ( $g$ ) categories as defined by the unfolding bracket design for the variable  $y$  ( $A$ ). In case where anchoring effects are accounted for, partial information needs to be incorporated allowing for a particular model of anchoring. In the case of Section 2, this would imply allowing for expression estimating analogous expression to (25) that might take into account combined anchoring for both the covariate and the regressor. In the case where  $[B_j, B_{j+1}] = [K_g, K_{g+1}]$  this is simplified to be identical to (25).

## Appendix B

This section draws from Vazquez-Alvarez, R., B. Melenberg and A. vanSoest (2000) to summarise the effect of two alternative models of anchoring to that exposed in Section 2. The two alternatives are from Hurd et al. (1998) and Harriges and Shogren (1996), respectively.

In the Hurd et al. (1998) paper, the main assumption is that individuals who are driven through an unfolding bracket design react to the bid by comparing such ‘clue’ to the unknown amount (for which they have failed to provide information in the first place). However, the comparison made by individuals might be biased by a ‘perception error’ made when bracket respondents compare the bid to the otherwise unknown amount. In Hurd et al. (1998), the perception error is modelled following a normal distribution with zero mean and a variance which is uncorrelated over subsequent bids in the unfolding design. Vazquez-Alvarez, R., B. Melenger and A. vanSoest (2000) use a zero median assumption on the perception error, thus allowing for a more flexible flexible approach while making the Hurd et al. (1998) assumption operation within the bounding interval approach. Vazquez-Alvarez, R., B. Melenger and A. vanSoest (2000) show who this semi-parametric assumption, together with a monotonic assumption similar to that of expression (21), can lead to the following set of bounds on  $P(A | BR)$ :

$$\begin{aligned}
 &\text{for } [0, B_{20}), \quad 0 \leq P(A | BR) \leq \min[1, 2P(Q_{20} = 0 | BR, Q_1 = 0)] \times \min[1, 2P(Q_1 = 0)] \\
 &\text{for } [B_{20}, B_1), \quad \min[1, 2P(Q_{20} = 0 | BR, Q_1 = 0)] \times \min[1, 2P(Q_1 = 0)] \\
 &\quad \leq P(A | BR) < \min[2P(Q_1 = 0), P(Q_1 = 0) + 2P(Q_{21} = 0 | Q_1 = 1)P(Q_1 = 1)] \\
 &\text{for } [B_1, B_{21}), \quad \min[2P(Q_1 = 0), P(Q_1 = 0) + 2P(Q_{21} = 0 | Q_1 = 1)P(Q_1 = 1)] \\
 &\quad \leq P(A | BR) < \min[P(Q_1 = 0), \min[1, 2P(Q_{21} = 0 | Q_1 = 1)P(Q_1 = 1)]] \\
 &\text{for } [B_{21}, \infty), \quad \min[P(Q_1 = 0), \min[1, 2P(Q_{21} = 0 | Q_1 = 1)P(Q_1 = 1)]] \leq P(A | BR) < 1
 \end{aligned}
 \tag{B.1}$$

If we substitute expression (B.1) for expression (25), and the equivalent expression in (29) for the sub-space of smokers, the model of anchoring of Hurd et al. (1998) substitutes that of Jacowitz and Kahneman (1995), thus allowing for the anchoring effect in a bounding interval.

Another alternative model of anchoring considered in Vazquez-Alvarez, R., B. Melenberg and A. vanSoest, (2000) is that of Herriges and Shroges, (1996). In this case the anchoring effect

is thought to kick in only after the first bid. They assume that response bias is only due to the effect of the first bid on subsequent bids. Thus, once the respondent faces a second bid – in our example either B20 or B21 –, their answer is a reflection of how they perceive B2k as compared to B1, the initial bid. Therefore, the respondent takes the initial bid as information on ‘A’, the unknown amount. Vazquez-Alvarez, R., B. Melenberg and A.vanSoest (2000) make this assumption operation with their expression (HS), Appendix A, to show that under the Herriges and Shogren (1996) model of anchoring,  $P(A | BR)$  has the following bounds:

$$\begin{aligned}
&\text{for } [0, B_{20}), & 0 \leq P(A | BR) \leq P(Q1 = 0) \\
&\text{for } [B_{20}, B_1), & P(Q1 = 0) \leq P(A | BR) < P(Q1 = 0) \\
&\text{for } [B_1, B_{21}), & P(Q1 = 0) \leq P(A | BR) < P(Q1 = 0) \\
&\text{for } [B_{21}, \infty), & P(Q1 = 0) \leq P(A | BR) < 1
\end{aligned} \tag{B.2}$$

As was the case with expression (B.1), substitution of (25) (and the related (29)) by (B.2) makes the an alternative anchoring model, this case that of Herriges and Shogren (1996), operational within a set of bounding interval while allowing for anchoring.

## Appendix C

Section 3 showed that applying a particular sample selection criteria to the 1996 wave of the Health and Retirement Study lead to a sample where the percentage of females is significantly higher than males, where these are population percentage due to the use of (cross-section representative) weights. It is possible to think that the selection criteria creates a sample imbalance between genders. To show this is not the case this appendix provides summary statistics of the distribution between males and females for the data before selection.

The HRS started in 1992, collecting data every two years and up to the last wave in the year 2000. The target population are individuals born between 1931 and 1941, thus, selected households are those with a member fulfilling such criteria. These individuals are called ‘household representatives’. Any other household member is also interviewed and referred to ‘second household respondents’, with such individuals been usually spouses, family relatives living in the same household or carers. The 1996 wave interviewed a total of 6,816 households, thus the sample representative of the target population amounted to 6,816 individuals. The share of these between males and females was of 2,262 and 4,554 respectively. Using cross-section representative weights this share implies that the target population is distributed such that 36.6% are males and 63.4% are females. Therefore, this already shows that females are a significantly larger percentage than males. In applying our selection criteria the first step is to select those who answer ‘yes’ to the question ‘Did you work for pay in the last calendar year (i.e., 1995)?’, a question which is asked to all households representatives. Based on those who answer ‘yes’, the final selection criteria implies selecting those who answer ‘yes’ to the question ‘Did you receive wages/salaries in the last calendar year (i.e., 1995)?’. Using cross-section population weights, Tables C.1 and C.2 shows the share between genders with respect to the possible answers to the above two questions. These estimates show that the distribution between genders for the final selected sample – employed for the empirical estimates in Section 4 – maintains a similar share between genders as that observed in the original 6,816 household respondents.

- (1) Original sample = 6,816 household respondents representing the target population of individuals born between 1931 and 1941. Estimates of the Population share between males and females is 36.6% and 63.4%, respectively.
- (2) Selection Criteria 1: Select individuals who answer YES to the question '*Did you work for pay during the last calendar year?*':

**Table C1: Distribution between genders and response to the first selection criteria.**

|                          | All individuals | Males | Females |
|--------------------------|-----------------|-------|---------|
| <b>Answer YES</b>        |                 |       |         |
| Total                    | 4,145           | 1,405 | 2,740   |
| Weighted percentage      | 100.0           | 30.9  | 69.1    |
| <b>Answer NO</b>         |                 |       |         |
| Total                    | 2,661           | 850   | 1,811   |
| Weighted percentage      | 100.0           | 32.6  | 67.4    |
| <b>Answer Don't know</b> |                 |       |         |
| Total                    | 10              | 7     | 3       |
| Weighted percentage      | 100.0           | 68.0  | 32.0    |

- (3) Selection Criteria 2: Select individuals who, having answered YES to the first selection criteria, answer YES to the question '*Did you receive pay in the form of wages/salaries during the last calendar year?*':

**Table C2: Distribution between genders and response to the second selection criteria.**

|                          | All individuals | Males | Females |
|--------------------------|-----------------|-------|---------|
| <b>Answer YES</b>        |                 |       |         |
| Total                    | 3,602           | 1,177 | 2,425   |
| Weighted percentage      | 100.0           | 38.4  | 61.6    |
| <b>Answer NO</b>         |                 |       |         |
| Total                    | 706             | 288   | 418     |
| Weighted percentage      | 100.0           | 44.5  | 57.5    |
| <b>Answer Don't know</b> |                 |       |         |
| Total                    | 6               | 1     | 5       |
| Weighted percentage      | 100.0           | 24.0  | 76.0    |