

Bounds, health habits and anchoring income effects

Rosalia Vazquez-Alvarez ^a

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



This version: December, 2007*

Abstract

Surveys are often design so that initial non-respondents to some continuous amount can disclose partial information with follow-up questions. These questions are often based on prompting responses with a sequence of bids that classify the undisclosed amount within a category. Secondary variables may reduce the problem of nonresponse but are unlikely to eliminate the problem altogether, thus, identification of the population parameters remains problematic. Furthermore, eliciting partial information by ‘anchoring’ the answers to a set of bids may induce anchoring bias. This paper develops from Horowitz and Manski (1995, 1998) to derive bounding intervals with partial information that allow for anchoring effects according to the experimental finding in Jacowitz and Kahneman (1995). The bounds provide regions of identification allowing for any type of non-random nonresponse. The method is illustrated using the 1996 wave of the Health and Retirement Study to show that bounding intervals can be used to detect for differences in health habit between income sub-groups in the population.

Keywords: Identification, item nonresponse, survey design, anchoring effects, health habit formation

JEL – Classification : C13, C14, C81, I11

^b Rosalia Vazquez-Alvarez, Rosalia.Vazquez-Alvarez@unisg.ch, www.siaw.unisg.ch

* Address for correspondence: SIAW, University of St.Gallen, Bodanstrasse 8, St. Gallen, 9000, Switzerland. I am very grateful to seminar participants at ISSC, the MCC meeting in Madrid and seminar participants at the University of St.Gallen. I am also indebt to Bertrand Melenberg, Arthur van Soest, and Michael Lechner for providing very valuable advice, and to Charles F. Manski who originally suggested the ideas developed in this paper as an extension to Chapter 5 in my doctoral thesis.

1 Introduction

The quality of survey data is subject to the response behaviour of survey participants. Issues such as misreporting as well as unit and item nonresponse can deteriorate the quality of empirical studies that rely on such data sources. Item nonresponse occurs when survey participants provide information for most of the variables in the survey but decline to answer particular questions often associated with sensitive information, for example, variables related to income or wealth. This means that item nonresponse cannot be thought as a random event. An immediate consequence is that the sample of full respondents does not constitute a random sample for the population of interest. Using information from full respondents (i.e., assuming nonresponse exogeneity) can bias estimates of the population parameters when these rely on variables that suffer from item nonresponse as result of the selection problem (see the seminal work by Manski, 1990, 1994 or more recent surveys in item nonresponse by Mason et al., 2002).

There are alternative methods to that of exogeneity but most of these alternatives imply specific non-testable conjectures on the distribution of the missing data. A classic reference that provides a battery of methods to deal with nonresponse is Rubin (1987) where most of the suggestions imply the use of rather strong distributional assumptions leading to the imputation of the missing value. The problem is that these assumptions cannot be tested with the same data set that applies them in the process of estimation. Since the early 1990s Charles Manski has advocated for an approach to deal with survey nonresponse that consists on using the available information to locate all possible values of the parameter in question without imposing any distributional restriction on the unobserved non-respondents information (see Manski, 1989, 1990, 1994, 1995, 1997, 2000, but to mention a few). The location region or bounding interval provides a region of identification for the unknown population parameter. The usefulness of the identification region to draw inference from the sample to the population depends on the fraction of sample nonresponse. However, it is often the case that weak data-driven assumptions can tighten the bounds (for empirical illustrations that follow this approach, see for example, Ginther (2000), Manski and Pepper (2000), Pepper (2000) and Pepper (2003)).

Data collection strategies can also reduce the problem of item nonresponse even before the data reaches the user. For example, the Health and Retirement Study (HRS) in the USA or the Survey on Health, Aging and Retirement in Europe (SHARE) employ unfolding bracket designs to elicit partial information from initial item non-respondents. These designs drive individuals through a sequence of questions that sort individual's undisclosed amount inside a particular category in the support of the variable that suffers from nonresponse. This is done by prompting initial non-respondents to react to a sequence of J 'bids' and at each bid the respondent has to reveal if they perceive the otherwise undisclosed amount to be greater (or not) than the bid. The inclusion of bracket designs is an effective

way to reduce the problem of initial item nonresponse. In general, the use of categorical questions is motivated because it is thought that rather than the characteristics of the question, it is personal characteristics and cognitive factors (e.g., confidentiality reasons) what drives nonresponse behaviour (for issues relating to survey response behaviour see for example, Borgers and Hox, 2001, Juster and Smith, 1997, Mason et al., 2000, Rabin, 1998, or Tourangeau et al., 2000). The problem with unfolding bracket designs is that these may induce ‘anchoring effects’, a phenomenon well documented in the psychology literature (see, for example, Tversky and Kahneman, 1974, Swizer and Snizek, 1991, Jacowitz and Kahneman, 1995). In the context of the present paper, anchoring effects can be explained as follows: initial non-respondents are presented with questions that involve ‘bids’ (e.g., the bid is \$25,000 in the question ‘*is the amount greater than \$25,000?*’). Individuals may or may not be uncertain about the undisclosed amount (e.g., some individuals genuinely do not know the amount and other prefer not to disclose it). What is certain is that respondents are asked to provide an answer that involves a cognitive process (e.g., a complex memory task) subject to both their own unobserved knowledge and the external information in the form of a bid. Different sequences of bids can induce different sets of answers if the bids act as ‘anchors’ and responses are driven in the direction of the bids. Hurd et al. (1998) show the existence of anchoring on questions relating to wealth and consumption using an experimental module from the 1996 wave of the HRS. In their article identification of the anchoring effect is possible because the module implies randomizing different unfolding designs (i.e., different sets of bids) among respondents. However, social scientists are often faced with answers from partial respondents to one single sequence of bids. Thus, dealing with anchoring effects in estimation implies making parametric, semiparametric or data-driven assumptions on the process that underlines such effects.

This paper develops from Horowitz and Manski (1995, 1998) to derive bounding intervals with partial information that allow for anchoring effects according to the experimental findings in Jacowitz and Kahneman (1995). Clearly, the latter is not the only possibility to model anchoring effects. Alternatives are the well known parametric interpretation in Hurd et al. (1998), or alternative interpretations such as those in Cameron and Quiggin (1994), Green et al. (1998), Herriges and Shogren (1996) or O’Connor et al. (1999) (see Vazquez-Alvarez et al., 2001, for a discussion of this alternative models of anchoring). The choice of Jacowitz and Kahneman (1995) is motivated because they provide a sound intuitive account of the effect of anchoring on estimation tasks through the use of a well designed experiment. Their findings provide grounds to model anchoring using weak data-driven assumptions as advocated by Manski’s bounding interval approach.

The bounds are applied to the probability of following adverse health habits (AHB) conditioning the probability on annual earnings, using the publicly released 1996 wave of the Health and Retirement Study. An initial earning’s nonresponse of 12% precludes identification of the conditional (on earnings) probability of the health outcome AHB because these 12% cannot be classified in the support of earnings.

An unfolding bracket design reduces the initial nonresponse rate to 3.2% but at the expense of potentially inducing anchoring effects. Different sets of bounds are estimated to cover all types of non-random nonresponse. The bounds are employed as a tool to detect for differences in the probability of following adverse health habits comparing between males and females who have in common similar earning's characteristics.

The paper is organized as follows. Section 2 derives 3 sets of bounds: the first ignores the presence of partial respondents; the other two allow for partial respondents but differ in the interpretation of anchoring effects. Section 3 describes the estimation process. Section 4 describes the data used in the empirical example and illustrates the bounds derived in Section 2 using the estimates in two distinct ways; first, an imputed earnings variable (HRS 1996 imputed files) is checked for validity against one of the estimated identification regions, and second, the empirical bounds are used to test for differences in health habits between genders. Section 4 concludes.

2 Theoretical framework

2.1 Covariate nonresponse without partial information

Suppose we observe cross-sectional information on an outcome $Y \in \mathbb{R}^1$ and covariates $X \in \mathbb{R}^k$ for a random sample of size n representative of the target population. Our aim is to make inference on $E[g(y)|A] = E[g(y)|A \subseteq X]$, i.e., the expectation on some bounded function $g(y)$ of the outcome $Y = y \in \mathbb{R}$ conditional on X for some value $A = a, a \in X$ or partition $A \subseteq X$ in the support of X ; for simplicity, throughout the paper $A \subseteq X$ is used to indicate either possibility. The problem arises because whereas all individuals provide information on Y , the conditioning set X suffers from non-negligent non-random nonresponse for some (but not all) individuals in the sample.¹ Non-random nonresponse on X precludes bias-free inference from the sample to the population. Let FR and NR be mutually exclusive binary outcomes indicating, respectively, full response and full nonresponse on X such that $P(FR) + P(NR) = 1$. Drawing from Horowitz and Manski (1998) and applying Bayes Theorem implies the following interpretation on the measure of interest:

¹ All bounds are derived assuming that only one covariate in X suffers from nonresponse; this simplification helps to clarify the exposition of the concepts throughout the paper. Allowing for more than one covariate to suffer from nonresponse would imply that the sample NR becomes partition in some d possible nonresponse subsets so that $E[g(y)|A] = E[g(y)|A, FR] \times P(FR|A) + \sum_k E[g(y)|A, NR_k] \times P(NR_k|A)$ would explain the partition of $E[g(y)|A]$. The nature of the bounds does change with d , thus we proceed with $d = 1$.

$$\begin{aligned}
E[g(y)|A] = & E[g(y)|A, FR] \times \left(\frac{P(A|FR)P(FR)}{P(A|FR)P(FR) + P(A|NR)P(NR)} \right) \\
& + E[g(y)|A, NR] \times \left(\frac{P(A|NR)P(NR)}{P(A|FR)P(FR) + P(A|NR)P(NR)} \right)
\end{aligned} \tag{1}$$

Expression (1) cannot be identified because the sampling process is not informative on either $P(A|NR)$ – i.e., the probability of $A \subseteq X$ when X is not observed for all in the sample– or $E[g(y)|A, NR]$ – i.e., the expectation of $g(y)$ conditional on $A \subseteq X$ for those with missing information on X . On the other hand, the sampling process identifies $E[g(y)|NR]$ while $P(A|NR) \in [0,1]^2$. The proposal is to use this information to derive sharp bounds on (1), bounds that might be further tighten using plausible weak data-driven assumptions. First we notice that $E[g(y)|A, NR] \in E[g(y)|NR]$ since nonresponse may occur for $A \subseteq X$ but also for $\bar{A} \subseteq X$, where $\bar{A}$ is the complement space of A with respect to X . With this and drawing from Horowitz and Manski (1998) expression $E[g(y)|NR]$ can be partition as follows:

$$\begin{aligned}
E[g(y)|NR] &= E[g(y)|NR, A] \times p + E[g(y)|NR, \bar{A}] \times (1-p) \\
&\text{where} \\
p &= P(A|NR), \quad 1-p = P(\bar{A}|NR)
\end{aligned} \tag{2}$$

Proposition 1 in Horowitz and Manski (1995) provide the conditions to establish a set of restrictions on $E[g(y)|NR, A]$:

- (a) Assume p is known, where $0 < p < 1$
 - (b) Let Ψ be the space of all distributions on Y
 - (c) For $p \in (0,1)$, $\Psi(p) \subseteq \Psi$
- Then,

$$E[g(y)|NR, A] \in \Psi(p) \equiv \Psi \cap \left\{ \frac{E[g(y)|NR] - \psi \times (1-p)}{p}, \psi \in \Psi \right\} \tag{3}$$

Assumption (3) – (a) implies that nonresponse is informative about the complement space of A (i.e., $\bar{A}$). Since the sampling process does not reveal the true value for p , the extreme assumptions of either $p=0$ or $p=1$ would imply imposing types of exogeneity assumptions on expression (1).³

² The sampling process is also informative on $E[g(y)]$ but for $P(NR) > 0$ the implication is that $E[g(y)|NR] \in E[g(y)]$. Thus, using $E[g(y)]$ instead of $E[g(y)|NR]$ would not lead to sharp bounds on (1).

³ Assumption $p=0$ identifies all non-respondents with the event $\bar{A} \subseteq X$. This exogeneity assumption on nonresponse leads to identification of (1) since $E[g(y)|A] = E[g(y)|A, FR]$. Alternatively, assumption $p=1$ implies that all non-

Assumption (3)–(b) defines a space for all probability distributions associated with the dependent variable; for any transformation of Y (including Y itself) the mapping of the transformation is an element in Ψ . Finally, (3)–(c) follows immediately from the union in (3), that is, $\Psi_{11}(p)$ is the space of probabilities containing elements common to Ψ and $\{(E[g(y)|NR] - \psi \times (1-p))/p : \psi \in \Psi\}$ so that all $\psi_{11} \in \Psi_{11}(p)$ are elements in Ψ , therefore $\Psi_{11}(p) \subseteq \Psi$ is always true. Allowing for $P(FR) + P(NR) = 1$, the conditions in (3) imply a sharp restriction on $E[g(y)|NR, A]$, i.e., the restrictions use all information available from the sampling process with no prior assumptions.⁴

Expression (3) implies general p –dependent identification conditions for all feasible $\psi \in \Psi(p)$. These conditions are able to characterize bounds on $E[g(y)|A, NR]$ for specific $g(y)$ functions. Horowitz and Manski (1995) use their Proposition 1 to derive sharp bounds for $g(y) = y$ and $g(y) = P(Y \leq y)$, $y \in R$. In general, if $g(y)$ is a bounded function mapping Ψ into $\mathbb{R}$, for either continuous or discrete $\psi \in \Psi(p)$, the restrictions in (3) imply values $g_0(p)$ and $g_1(p)$ satisfying the conditions $g_0(p) \equiv \inf[\psi : \psi \in \Psi(p)]$ and $g_1(p) \equiv \sup[\psi : \psi \in \Psi(p)]$, respectively, thus characterizing the upper and lower bounds on $E[g(y)|NR, A]$:

Given

$$\begin{aligned} E[g(y)|NR, A] \in \Psi(p) &\equiv \Psi \cap \left\{ \frac{E[g(y)|NR] - \psi \times (1-p)}{p}, \psi \in \Psi(p) \right\} \\ \text{for } g_0(p) &\equiv \inf[\psi : \psi \in \Psi(p)], \quad g_1(p) \equiv \sup[\psi : \psi \in \Psi(p)] \\ &\Rightarrow \\ E[g(y)|NR, A] &\in \{g_0(p), g_1(p)\} \end{aligned} \tag{4}$$

For example, if y is a binary outcome, $\Psi(p) \subseteq \Psi \equiv [0,1]$ the function of interest becomes the indicator function $g(y) = I[y = 1]$, so that for known p the following characterization of (4) applies:

respondents are identified with the event $A \subseteq X$, so that $E[g(y)|A] = E[g(y)|A, NR]$ and again, expression (1) is identifies.

⁴ Proof that (3) implies a sharp restriction. From (2) and given $P_A = E[g(y)|NR, A]$ and $P_{\bar{A}} = E[g(y)|NR, \bar{A}]$, the direct product $\Psi \times \Psi$ explains the space of probabilities for all feasible pairs between P_A and $P_{\bar{A}}$ so that $(P_A, P_{\bar{A}}) \in \{(\psi_A, \psi_{\bar{A}}) \in \Psi \times \Psi : E[g(y)|NR] = pP_A + (1-p)P_{\bar{A}}\}$ applies. Thus, (3) gives feasible values of P_A , while feasible $P_{\bar{A}}$ values are $\psi_{\bar{A}}$ such that $(E[g(y)|NR] - \psi_{\bar{A}} \times (1-p))/p \in \Psi$. Since the sampling process reveals no more information than that implied by the probability space $\Psi \times \Psi$ the restriction in (3) is sharp for $P(FR) + P(NR) = 1$. This proof is a straight forwards adaptation of that in Horowitz and Manski (1995, page 297).

$$\begin{aligned}
& p = P(A | NR) \in (0,1), \quad p : \text{known}; \\
& P(y | NR, A) \in \Psi(p) \equiv [0,1] \cap \left\{ \frac{P(y | NR) - \psi \times (1-p)}{p}, \psi \in [0,1] \right\} \\
\Rightarrow & \\
& P(y | NR, A) \in \Psi(p) \equiv \left\{ \max \left(0, \frac{P(y | NR) - (1-p)}{p} \right), \min \left(1, \frac{P(y | NR)}{p} \right) \right\}
\end{aligned} \tag{4'}$$

Expression (4') characterizes the upper and lower bound on $E[g(y) | NR, A]$ for $g(y) = I[y=1]$. Since these derive from the restrictions imposed in (3) and these restrictions are sharp, bounds in (4) and (4') are sharp. Applying (4) to expression (1) implies identifying $E[g(y) | A]$ up to a bounding interval as follows:

$$\begin{aligned}
& E[g(y) | A, FR] \times \left(\frac{P(A | FR)P(FR)}{P(A | FR)P(FR) + P(NR)p} \right) + g_0(p) \times \left(\frac{P(NR)p}{P(A | FR)P(FR) + P(NR)p} \right) \\
& \leq E[g(y) | A] \leq \\
& E[g(y) | A, FR] \times \left(\frac{P(A | FR)P(FR)}{P(A | FR)P(FR) + P(NR)p} \right) + g_1(p) \times \left(\frac{P(NR)p}{P(A | FR)P(FR) + P(NR)p} \right)
\end{aligned} \tag{5}$$

The width between bounds in (5) equals $(g_1(p) - g_0(p)) \times \left(\frac{P(NR)p}{P(A | FR)P(FR) + P(NR)p} \right)$ and it is clearly p -dependent. For low nonresponse $P(NR)$ the bounding interval becomes more informative as $p \rightarrow 1$, since this would imply that $E[g(y) | NR, A]$ is well approximated by the expectation $E[g(y) | NR]$, a measure that is fully revealed by the sampling process.

Bounds in (5) are derived assuming knowledge of $p = P(A | NR)$ when in reality we only know that $p = P(A | NR) \in [0,1]$. Following Horowitz and Manski (1998) sharp bounds are obtained from (5) by minimising and maximising the lower and upper bounds, respectively, over the interval $[0,1]$:

Let $p \in [0,1]$ and allow for (4):

$$\begin{aligned}
& \inf_p \left\{ E[g(y) | A, FR] \times \left(\frac{P(A | FR)P(FR)}{P(A | FR)P(FR) + P(NR)p} \right) + g_0(p) \times \left(\frac{P(NR)p}{P(A | FR)P(FR) + P(NR)p} \right) \right\} \\
& \leq E[g(y) | A] \leq \\
& \sup_p \left\{ E[g(y) | A, FR] \times \left(\frac{P(A | FR)P(FR)}{P(A | FR)P(FR) + P(NR)p} \right) + g_1(p) \times \left(\frac{P(NR)p}{P(A | FR)P(FR) + P(NR)p} \right) \right\}
\end{aligned} \tag{5'}$$

Expression (5') provides an empirically implementable set of bounds on $E[g(y)|A]$ when X suffers from non-random non-negligent item nonresponse.⁵ In this paper we take (5') as the basis to derive bounding intervals that include the possibility of partial information obtained through an unfolding bracket design, with such information being potentially subject to anchoring effects.

2.2 Covariate nonresponse and partial information respondents

Surveys are often design so that initial non-respondents to some variable X are routed to secondary variables where individuals can disclose partial information on the missing variable, for example, disclosing information in a category within a range of categories. This is especially true when X refers to 'amounts' that could be subject to confidentiality problems and/or problems of missing information.⁶ To insert this possibility in (5') let the outcome PR indicate 'partial information' so that NR is now defined as 'full nonresponse' with respect to both an open-ended question and questions that elicit partial information. With this $P(FR) + P(PR) + P(NR) = 1$ applies and expression (1) is modified as follows:

$$\begin{aligned} E[g(y)|A] = & E[g(y)|A, FR] \times \left(\frac{P(A|FR)P(FR)}{P(A|FR)P(FR) + P(A|PR)P(PR) + P(A|NR)P(NR)} \right) \\ & + E[g(y)|A, PR] \times \left(\frac{P(A|PR)P(PR)}{P(A|FR)P(FR) + P(A|PR)P(PR) + P(A|NR)P(NR)} \right) \\ & + E[g(y)|A, NR] \times \left(\frac{P(A|NR)P(NR)}{P(A|FR)P(FR) + P(A|PR)P(PR) + P(A|NR)P(NR)} \right) \end{aligned} \quad (6)$$

Expression (6) suffers from a similar identification problem as in (1); neither $E[g(y)|A, NR]$ nor $P(NR|A) \in (0,1)$ are identified by the data. Moreover, partial information implies that both $E[g(y)|A, PR]$ and $P(PR|A) \in (0,1)$ are only 'partially' but not fully identified. Compared to (1), expression (6) may be more informative on $E[g(y)|A]$ because $P(NR)$ is now reduced by $P(PR) > 0$ thus reducing the width between upper and lower bounds that would have resulted given some initial nonresponse rate. Furthermore, when correct and unbiased, partial information on $A \subseteq X$ from initial non-respondents cannot worsen our understanding on the conditional characteristics of $g(y)$.

To understand the consequence of partial information in (5') we need to understand the consequence of bounding $E[g(y)|A, PR] \times \left(\frac{P(A|PR)P(PR)}{P(A|FR)P(FR) + P(A|PR)P(PR) + P(A|NR)P(NR)} \right)$ in the right

⁵ An underlying assumption in (5') is that of a positive density for the event $A \subseteq X$ for all A partitions in the support containing X .

⁶ See, for example, Juster and Smith (1997)

hand side of (6). Derivation of these bounds depends on the method employed to elicit partial information. This paper is concerned with the effect of unfolding brackets as designs used to elicit such information. Unfolding brackets are often used when individuals answer either ‘*don’t know*’ or ‘*refuse to respond*’ to an initially open-ended question such as ‘*how much earnings did you received during the last year?*’. Following up to the open-ended question, initial non-respondents are routed towards a questions such as ‘*Is the amount greater than B1 ?*’, where the initial bid $B1$ is a pre-specified feasible quantity in X . Respondents to this bracket design may answer to $B1$ with either ‘yes’, ‘no’ or yet again with ‘*don’t know*’ or ‘*refuse to respond*’. Individuals’ answers conditions the second bid in the design: answering ‘yes’ leads to a follow up question such as ‘*Is the amount greater than B2 ?*’ where $B2 > B1$, whereas answering ‘no’ to the bid $B1$ implies a similar question but with a bid $B2$ such that $B2 < B1$. Figure 1 shows the first part in the dynamics of an unfolding bracket design. In this example it is assumed that initial non-respondents to (say) ‘*earning received over the last calendar year*’ where routed towards an unfolding bracket design with $B1 = \$25,000$, $B2 | (B1 = \text{yes}) = \$50,000$ and $B2 | (B1 = \text{no}) = \$5,000$.

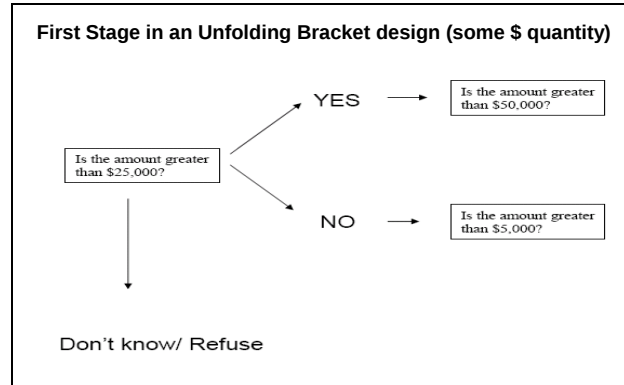


Table 1: Example (earnings): An Unfolding Bracket Design and resulting categories

	Anchor 1: B1	Answer to Anchor B1	Anchor 2: B20 or B21	Answer to Anchor 2	Resulting Bracketed categories
Complete Bracket Respondent(CBR)	> \$25,000?			YES	[50,000; infinity)
		YES	> \$50,000?	NO	[25,000; 50,000)
				YES	[5,000; 25,000)
		NO	> \$5,000	NO	[0; 5,000)
Incomplete Bracket Respondent (ICB)	> \$25,000?	YES	> \$50,000?	Don't know/ refuse	[25,000; infinity)
		NO	> \$5,000	Don't know/ refuse	[0; 25,000)

The results of Figure 1 are explained in Table 1 by illustrating the consequent categories of an unfolding bracket design assuming that individuals face (at most) two bids.⁷ This illustration is closely related to examples found in surveys such as ‘*The Survey on Health, Aging and Retirement in Europe*’ (SHARE, Europe), ‘*The Health and Retirement Study*’ (HRS, United States), or the *DNB Household Survey Data* (University of Tilburg, The Netherlands).

As described in Section 1, ‘anchoring effects’ may be induced when individuals’ answers are influenced by the wording in the question. To allow for this possibility we derive bounds on $E[g(y) | A]$ first ignoring and then allowing for potential anchoring effects. In both cases we centre our attention on complete bracket respondents (CBR) to an unfolding bracket design and came back to examine the effect of incomplete bracket respondents (IBR) at the end of the section.⁸

Partial information and no anchoring effects

Let κ be a binary outcome so that $\kappa=1$ indicates a ‘yes’ answer and $\kappa=0$ indicates a ‘no’ answer to a given question. An unfolding bracket design implies that a non-respondent to an initial open-ended question faces a sequence of questions of the following nature:

$$\begin{aligned} &\text{Is the amount greater than } B_j | (B_1 = \kappa, B_2 = \kappa, \dots, B(j-1) = \kappa)? \\ &\text{where} \\ &B_j: j\text{th bid, } j=1, \dots, J \end{aligned} \tag{7}$$

The bid in (7) implies that at each round of the design the amount faced by individuals is conditional on their previous answers: for CBR the indicator κ can take only two values (0,1). Thus, with J bids $\{\min(X), (B_j | B_1 = \dots = B(j-1) = \text{NO}), \dots, (B_j | B_1 = \dots = B(j-1) = \text{YES}), \max(X)\}$ defines all possible conditional bids in the support of X . Once the sequence is completed the result is a set of 2^J categories in the support of the variable that suffers from nonresponse. These 2^J categories are consecutively paired and each pair is associated with either one of the two possible answers ($\kappa=0,1$). To simplify notation let B_τ represent $B_\tau = (B_j | B_1 = \kappa, B_2 = \kappa, \dots, B(j-1) = \kappa)$ for any $j \in J$. Answering $\kappa=1$ to B_τ leads to a lower limit in one of the 2^J categories, say for category $A_{\tau, \kappa=1}$. Alternatively, if the answer is $\kappa=0$ the same bid acts as an upper limit in the preceding $A_{\tau, \kappa=0}$ category. With this,

⁷ The initial bid leads to two conditional bids so that the sequence implies 3 bids, a starting bid and two follow up conditional bids. We call it a ‘two bid’ sequence because from the respondent’s point of view each faces (at most) two bids.

⁸ Table 1 distinguishes between complete (CBR) and incomplete (IBR) bracket respondents. This distinction is only relevant at an empirical level, with both types of respondents leading to bounds that are qualitatively the same. Not completing the sequence (i.e., IBR) provide partial information as if individuals had faced a single bid B_1 ; so that X is partition such that $[0, B_1)$ and $[B_1, \max)$. On the other hand, the partitions given CBR are $[0, (B_2 | B_1 = \text{no}))$, $[(B_2 | B_1 = \text{no}), B_1]$, $[B_1, (B_2 | B_1 = \text{yes})]$ and $[(B_2 | B_1 = \text{yes}), \max)$. In deriving bounds with partial information what matters is how to model the information for any given partition, and not the number of partitions.

$X = \{A_{1,0} = [\min(X), B_1), A_{1,1} = [B_1, B_2), \dots, A_{\tau,0} = [B_{\tau i-1}, B_{\tau i}), A_{\tau,1} = [B_{\tau i}, B_{\tau i+1}), \dots, A_{\tau=J^2,1} = [B_{\tau}, \max(X))\}$ explains the partition of X into the 2^J categories. The partition shows that each category's lower limit corresponds to the upper bid of the preceding category. Table 1 illustrates this for $J = 2$ bids so that $X = \text{earnings} \in \{[\$0; \$5,000), [\$5,000; \$25,000), [\$25,000; \$50,000), [\$50,000; \max(X)]\}$ represents the $2^J = 4$ resulting categories.

Our aim is to interpret the information with respect to each of the $A_{\tau,\kappa}$ categories resulting from a sequence of 'yes' ($\kappa = 1$) and 'no' ($\kappa = 0$) answers in an unfolding bracket design. Let $Q\tau$ indicate the outcome to a sequence of κ answers to questions in (7) so that τ implies a route that leads to the j -th bid conditional on some combined set of answers to all other $(j-1)$ preceding bids. Thus, τ stands as a sequence of numbers so that $Q\tau, \tau \Rightarrow \{j, \kappa_{j-1}, \kappa_{j-2}, \dots, \kappa_1\}, \kappa_j = \{0,1\}, j \in [1, J]$ summarizes all possible combination of answers to the conditional bid $B_\tau = (Bj | B1, B2, \dots, B(j-1))$. For example, Q311 results from $j = 3; \tau \Rightarrow \{3, \kappa_2, \kappa_1\} = \{3, 1, 1\}$ to indicate that individuals face a 3rd bid ($B3$) following a 'yes' answer to (7) for each of the two sequential bids $B1$ and $B2$. If an answer is such that $Q311 = 1$, this means that the respondent perceives the otherwise undisclosed amount to be greater than $B3$ conditional on having perceived the undisclosed amount to be greater than $B1$ and $B2$. The event $Q311 = 0$ becomes the mutually exclusive event to $Q311 = 1$ for CBR, i.e., complete bracket respondents answer either $Q\tau = 1$ or $Q\tau = 0$ at each nodes of the design, so that $P(Q\tau = 1 | PR, Q_{-\tau}) + P(Q\tau = 0 | PR, Q_{-\tau}) = 1$ where $Q_{-\tau}$ stands for the path of answers leading to $Q\tau$ (and is mutual to both $Q\tau = 1$ or $Q\tau = 0$). Using these two mutually exclusive outcomes leads to the following representation:

$$\begin{aligned}
 \text{Let } X &= \{A_{1,0}, A_{1,1}, \dots, A_{\tau,0}, A_{\tau,1}, \dots, A_{\tau,0}, A_{\tau,1}\}, \text{ where} \\
 A_{\tau,0} &= [B_{\tau i-1}, B_{\tau i}) \quad A_{\tau,1} = [B_{\tau i}, B_{\tau i+1}) \\
 \text{Let } Q\tau &= \kappa, \quad \kappa = \{0,1\} \text{ be associated with } A_{\tau,\kappa} \\
 P(Q\tau = \kappa | PR, Q_{-\tau}) &= \\
 &P(Q\tau = \kappa | PR, Q_{-\tau}, A \in A_{\tau,\kappa})P(A \in A_{\tau,\kappa} | PR, Q_{-\tau}) \\
 &+ P(Q\tau = \kappa | PR, Q_{-\tau}, A \notin A_{\tau,\kappa})P(A \notin A_{\tau,\kappa} | PR, Q_{-\tau})
 \end{aligned} \tag{8}$$

Expression (8) allows for a direct interpretation of individual's answers to the bids, that is, we interpret $P(Q\tau = 1 | \dots)$ as opposed to directly looking at $P(A | \dots)$. Introducing $P(Q\tau = 1 | \dots)$ is crucial for us to deal with anchoring effects. This is because the bid $B_\tau = (Bj | B1, B2, \dots, B(j-1))$ associated with the

two *potential* outcomes $Q\tau = \kappa$, $\kappa = \{0,1\}$ may act as an anchor and affect the *actual* outcome. Thus, anchoring implies the possibility that even when the true (but undisclosed) amount is $A \notin A_{\tau,\kappa}$, for $A_{\tau,\kappa}$ associated with a given κ , the answer becomes $Q\tau = \kappa$. In the absence of anchoring effects the following applies:

$$\begin{aligned}
& \text{Let } Q\tau = \kappa, \quad \kappa = \{0,1\} \text{ be associated with } A_{\tau,\kappa} \\
& \text{No anchoring} \Rightarrow P(Q\tau = \kappa | BR, Q_{-\tau}, A \notin A_{\tau,\kappa}) = 0 \\
& \text{From (8),} \\
& P(Q\tau = \kappa | PR, Q_{-\tau}) = \\
& \quad 1 \times P(A \in A_{\tau,\kappa} | PR, Q_{-\tau}) + 0 \times P(A \notin A_{\tau,\kappa} | PR, Q_{-\tau}) \\
& \Rightarrow P(Q\tau = \kappa | PR, Q_{-\tau}) = P(A \in A_{\tau,\kappa} | PR, Q_{-\tau})
\end{aligned} \tag{9}$$

Thus, with no anchoring effects, expression (9) shows that unfolding brackets information identifies exactly the conditional probability of falling within each of the 2^J categories in X , with the number of categories endogenously determined according to respondent's answers.⁹ Expression (6) shows that in the presence of partial information identifying $E[g(y)|A]$ depends (in part) on identifying $P(A|PR)$ and $E[g(y)|A, PR]$ for $A \subseteq X$. Expression (9) reveals that for any given $A_{\tau,\kappa} \subseteq X$ the probability $P(A \in A_{\tau,\kappa} | PR)$ is identified by the sampling process under the assumption of no anchoring effects. Estimation will depend on the sample analogue for $P(Q\tau = \kappa | PR)$ where $Q\tau$ implies $Q_{-\tau}$. Likewise $E[g(y)|A_{\tau,\kappa} \subseteq X, PR]$ is also fully identified under the assumptions leading to (9). Let $p_{\tau,\kappa} = P(A \in A_{\tau,\kappa} | PR)$, $g_{\tau,\kappa} = E[g(y)|A \in A_{\tau,\kappa}, PR]$ and allow for expression (4) to determine the contribution from full non-respondents; bounds on (6) are given as follows:

⁹ This is in fact the only difference between the end result of an unfolding bracket design and the result of giving individuals a simultaneous range of categories from where to choose. Assuming no error (or induced anchoring effect) in classifying the missing value, individual's answers determine the categories within X , whereas with range cards the categories that partition X are predetermined.

For each $A_{\tau,\kappa} \subseteq X$,

Let $p \in [0,1]$, $p_{\tau,\kappa} = P(A_{\tau,\kappa} | PR)$, $g_{\tau,\kappa} = E[g(y) | A \in A_{\tau,\kappa}, PR]$,

and allow for (4) to define $\{g_0(p), g_1(p)\}$,

$\Rightarrow$

$$\inf_p \left\{ \begin{aligned} & E[g(y) | A, FR] \times \left(\frac{P(A | FR)P(FR)}{P(A | FR)P(FR) + P(PR)p_{\tau,\kappa} + P(NR)p} \right) \\ & + (g_{\tau,\kappa}) \times \left(\frac{P(PR)p_{\tau,\kappa}}{P(A | FR)P(FR) + P(PR)p_{\tau,\kappa} + P(NR)p} \right) \\ & + g_0(p) \times \left(\frac{P(NR)p}{P(A | FR)P(FR) + P(PR)p_{\tau,\kappa} + P(NR)p} \right) \end{aligned} \right\} \leq E[g(y) | A] \leq \sup_p \left\{ \begin{aligned} & E[g(y) | A, FR] \times \left(\frac{P(A | FR)P(FR)}{P(A | FR)P(FR) + P(PR)p_{\tau,\kappa} + P(NR)p} \right) \\ & + (g_{\tau,\kappa}) \times \left(\frac{P(PR)p_{\tau,\kappa}}{P(A | FR)P(FR) + P(PR)p_{\tau,\kappa} + P(NR)p} \right) \\ & + g_1(p) \times \left(\frac{P(NR)p}{P(A | FR)P(FR) + P(PR)p_{\tau,\kappa} + P(NR)p} \right) \end{aligned} \right\} \quad (10)$$

It has to be emphasised that partial information elicited with unfolding brackets is not informative for specific values of X so that bounds in (10) apply for the conditional event $A_{\tau,\kappa} \subseteq X$ but do not apply for conditional events such as $A = a \in X$. On the other hand, partial information through unfolding brackets is informative on $P(A \leq B | PR)$ provided that B matches any of the bids in $X = \{A_{1,0} = [\min(X), B_1], A_{1,1} = [B_1, B_2], \dots, A_{\tau,\kappa} = [B_{\tau i + \kappa - 1}, B_{\tau i + \kappa}], \dots, A_{1=J^2,1} = [B_{T,1}, \max(X)]\}$. If so, it is easily shown that the sampling process identifies the sampling analogue of $P(A \leq B | PR)$ as a cumulative step function and bounds in (10) apply for cumulative regions $(A \leq B_i)$.¹⁰ For $P(PR) > 0$ and given $P(FR) + P(NR) + P(PR) = 1$, the interval in (10) may be narrower than that implied by (5'). The problem is that expression (10) is derived under the restrictive assumption of no anchoring effects. Thus, bounds in (10) may be tighter than those in (5') not just because they include partial respondent's information but because they further include the assumption of no anchoring effects.

¹⁰ Notice at this point that deriving bounds as in (10) but applying (8) and (9) to the difference in probabilities $P(A < B_{\tau_2} | \dots) - P(A \leq B_{\tau_1} | \dots)$ for $A_{\tau} = [B_{\tau_1}, B_{\tau_2}]$, would not lead to sharp bounds. Instead, dealing directly with $P(A \in [B_{\tau_1}, B_{\tau_2}] | \dots)$ as we did in (8) and (9) implies that the bounds in (10) are sharp. Manski (1994, Footnote 2) already points out this problem when bounds apply to intervals in a continuous distribution. For our case the sampling process reveals $P(A \leq B_{\tau_2} | \dots)$ and not $P(A < B_{\tau_2} | \dots)$, where necessarily $P(A \leq B_{\tau_2} | \dots) > P(A < B_{\tau_2} | \dots)$.

Partial information and anchoring Effects

Anchoring effects implies that the bids given along a sequence of unfolding brackets can act as ‘cues’ or ‘anchors’ thus affecting individual’s answers towards the anchor (see Section 1). Being anchored by a bid at any point of the unfolding sequence implies that individuals may classify the otherwise undisclosed $A \subseteq X$ in $A_{\tau,\kappa}$ when in fact $A \notin A_{\tau,\kappa}$. This modifies expression (9):

$$\begin{aligned}
 & \text{Let } Q\tau = \kappa, \quad \kappa = \{0,1\} \text{ be associated with } A_{\tau,\kappa} \\
 & \text{Anchoring} \Rightarrow P(Q\tau = \kappa | BR, Q_{-\tau}, A \notin A_{\tau,\kappa}) \neq 0 \\
 & \text{From (8),} \\
 & \Rightarrow P(Q\tau = \kappa | PR, Q_{-\tau}) \neq P(A \in A_{\tau,\kappa} | PR, Q_{-\tau})
 \end{aligned} \tag{11}$$

The immediate effect of anchoring is that the fraction classified in a given category among all categories in X may provide bias information on the true fraction that would have been revealed if individuals had in fact reported the actual amount. Thus, anchoring implies that expression (11) substitutes (9) at the expense of braking the direct link between information provided by unfolding bracket respondents, i.e., $P(Q\tau = \kappa | PR, Q_{-\tau})$, $\kappa = \{0,1\}$, and the conditional distribution of X , i.e., $P(A | PR, Q_{-\tau})$.¹¹ Ignoring the possibility of anchoring may bias our estimates. For example, relative to the true distribution of X the sample analogue may be a shift of such distribution in the direction of the bids. Allowing for anchoring effects when deriving bounds implies modifying (10) to take into account such potential bias. Vazquez-Alvarez, Melenberg and vanSoest (2001) already provide guidelines to bound distribution of an outcome that suffers from nonresponse and allowing for various models of anchoring. In what follows we extend the findings in Vazquez-Alvarez et al. (2001) to interpret the concept of anchoring according to one of these models, namely, the model of anchoring according to the experimental findings in Jacowitz and Kahneman (1995). Choosing to interpret anchoring according to Jacowitz and Kahneman (1995) is purely for subjective reasons. For example, an alternative model would be that implied by the Hurd et al (1998) study. In the latter it is assumed that respondents are uncertain about the quantity $A \in X$ when faced with a question as given in (7) and they resolve this uncertainty by comparing A to $(Bj + \varepsilon)$ where ε is ‘the perception error’: the key issue is to make distributional assumptions on ε and to estimate the effect these assumption have on the distribution of X . Alternatively, Jacowitz and Kahneman (1995) conduct an experimental study to show that inserting either low or high anchors in estimation tasks significantly affects individual answer’s in the direction of the

¹¹ A test for anchoring effects requires that initial non-respondents are randomly assigned to different unfolding designs each with a different set of bids over the same support of X . Hurd et al (1998) perform a test of anchoring using a special module of the HRS (1996) and cannot reject the null of no anchoring. As in the present paper, most surveys rely on a single set of bids when eliciting partial information. In the light of evidence in Hurd et al (1998) the assumptions in (9) can be rejected.

anchors. Their empirical findings provide a plausible and more intuitive account of anchoring than the structural proposal of a perception error. Thus, in what follows potential anchoring bias is modelled in (10) allowing for the empirical findings of Jacowitz and Kahneman (1995). Allowing for anchoring effects using other modelling strategies (e.g., Hurd et al. (1998) or Herriges and Shogren (1996)) implies different sets of assumptions but does not change the qualitative nature of the bounds.¹²

In Vazquez-Alvarez et al. (2001) the effect of anchoring as discussed by Jacowitz and Kahneman (1995) is made operational only for a 2-bid unfolding design with the specific aim to bound the cumulative distribution function $P(A \leq B_r)$ where B_r stands for anyone of the 2 conditional bids. The present paper shows how to derive bounds on $E[g(y) | A_{r,x} \subseteq X, PR]$ allowing for Jacowitz and Kahneman (1995). Thus, the present paper extends Horowitz and Manski (1998) to account for anchoring effects in a bounding interval drawing from assumptions in Vazquez-Alvarez et al. (2001) but generalizing for unfolding designs with any number of J bids.

The empirical findings in Jacowitz and Kahneman (1995) suggest that in the presence of a high anchor respondents too often report that the otherwise unknown amount exceeds the anchor; the same study finds that a low anchor drives individual's answers to lower the median estimate of the quantity (when compared to the calibration group who are not subject to anchors). Finally, although both low and high anchors are found to be effective at pulling individuals' answers towards the anchor, Jacowitz and Kahneman (1995) find that the effect of a high anchor is significantly larger than that of a low anchor. The following set up implies an interpretation of the findings by Jacowitz and Kahneman (1995):

(i,1) Let the bid B_{ri} be perceived as a HIGH anchor: Respondents too often will report that the amount exceeds the anchor:

$$\Rightarrow B_{ri} \text{ 'high': } P(Q_{ri} = 1 | PR, Q_{-ri}) \geq P(A > B_{ri} | PR, Q_{-ri})$$

(ii,1) Let the bid B_{ri} be perceived as a LOW anchor: Respondents too often will report that the amount falls below the anchor:

$$\Rightarrow B_{ri} \text{ 'low': } P(Q_{ri} = 1 | PR, Q_{-ri}) \leq P(A > B_{ri} | PR, Q_{-ri})$$

(13)

To apply expression (13), one needs to make further assumptions so to distinguish between 'perceived' low and high anchors. The following weak data assumption distinguishes between high and low anchors within an unfolding bracket design:

¹² The choice is subjective by nature since all models imply a set of assumptions that cannot, by definition, be tested. All assumptions as to how anchoring kicks in are equally valid and Vazquez-Alvarez et al. (2001) provide the basis to allow for alternative anchoring assumptions drawing from Hurd et al. (1998) or Herriges and Shogren (1996).

Assumption.1

(i,2) B_{τ_i} is HIGH and acts as a HIGH ANCHOR if the following is observed:

$$P(Q_{\tau_i} = 1 | PR, Q_{-\tau_i}) \leq 0.5 \quad (14)$$

(ii,2) B_{τ_i} is LOW and acts as a LOW ANCHOR if the following is observed:

$$P(Q_{\tau_i} = 1 | PR, Q_{-\tau_i}) \geq 0.5$$

Notice that for Assumption 1 in (14) the only event that matters to determine the type of anchor is that associated with a ‘yes’ answer (i.e., $Q_{\tau} = 1$). Together, assumptions (13) and (14) can be used to place bounds on $P(A \in A_{\tau, \kappa} | \dots)$ with such bounds allowing for potential anchoring effects. However, since $A_{\tau, \kappa} = [B_{\tau_i}, B_{\tau_{i+1}})$, the effect that anchoring may have had on the sample classified in $A_{\tau, \kappa}$ is the combined effect that the two bids have had on individual’s answers. We start by looking at what happens when a bid B_{τ_i} is perceived as a high anchor:

B_{τ_i} is peceived as a HIGH anchor:

i.e.,

$$\begin{aligned} \text{Observing } & P(Q_{\tau_i} = 1 | PR, Q_{-\tau_i}) \leq 0.5 \\ \Rightarrow & P(A > B_{\tau_i} | PR, Q_{-\tau_i}) \leq P(Q_{\tau_i} = 1 | PR, Q_{-\tau_i}) \quad \text{from (13)} \\ & \text{where } P(A > B_{\tau_i} | PR, Q_{-\tau_i}) \text{ is unknown (given anchoring).} \end{aligned}$$

then,

$$\begin{aligned} & P(A > B_{\tau_i} | PR, Q_{-\tau_i}) \leq P(Q_{\tau_i} = 1 | PR, Q_{-\tau_i}) \quad (15) \\ \Rightarrow & P(A \leq B_{\tau_i} | PR, Q_{-\tau_i}) \geq P(Q_{\tau_i} = 0 | PR, Q_{-\tau_i}) \\ \Rightarrow & P(A \leq B_{\tau_i} | PR) \geq P(Q_{\tau_i} = 0 | PR, Q_{-\tau_i}) \quad \text{because } P(A | Z) \leq P(A) \\ \Rightarrow & P(Q_{\tau_i} = 0 | PR, Q_{-\tau_i}) \leq P(A \leq B_{\tau_i} | PR) \leq 1 \quad (a.1) \end{aligned}$$

Expression (15) shows an upper bound of 1 for the unknown probability $P(A \leq B_{\tau_i} | PR)$; this simply reflects the fact that the assumption $P(A > B_{\tau_i} | PR, Q_{-\tau_i}) \leq P(Q_{\tau_i} = 1 | PR, Q_{-\tau_i})$ provides no information to bound $P(A \leq B_{\tau_i} | PR)$ from above. We now look at what happens if B_{τ_i} is perceived as a low anchor:

B_{τ_i} is perceived as a LOW bid:

i.e.,

$$\begin{aligned} \text{Observing } & P(Q_{\tau_i} = 1 | PR, Q_{-\tau_i}) \geq 0.5 \\ \Rightarrow & P(A > B_{\tau_i} | PR, Q_{-\tau_i}) \geq P(Q_{\tau_i} = 1 | PR, Q_{-\tau_i}) \quad \text{from (13)} \\ & \text{where } P(A > B_{\tau_i} | PR, Q_{-\tau_i}) \text{ is unknown (given anchoring).} \end{aligned}$$

then,

$$\begin{aligned} & P(A > B_{\tau_i} | PR, Q_{-\tau_i}) \geq P(Q_{\tau_i} = 1 | PR, Q_{-\tau_i}) \\ \Rightarrow & P(A > B_{\tau_i} | PR) \geq P(Q_{\tau_i} = 1 | PR, Q_{-\tau_i}) \quad \text{because } P(A | Z) \leq P(A) \\ \Rightarrow & P(A \leq B_{\tau_i} | PR) \leq P(Q_{\tau_i} = 0 | PR, Q_{-\tau_i}) \end{aligned} \tag{16}$$

$$\Rightarrow 0 \leq P(A \leq B_{\tau_i} | PR) \leq P(Q_{\tau_i} = 0 | PR, Q_{-\tau_i}) \quad (a.2)$$

Expression (16) shows a low bound of 0 for the unknown $P(A \leq B_{\tau_i} | PR)$ because, similar to the case in (15), in the presence of a low bid B_{τ_i} the assumptions that allow for anchoring effects are only informative with regards to the upper bound on $P(A \leq B_{\tau_i} | PR)$. With (a.1) in expression (15) and (a.2) in expression (16) we can derive bounds on $P(A \in A_{\tau, \kappa} | BR)$:

Let $A_{\tau, \kappa} = [B_{\tau_i}, B_{\tau_{i+1}}]$;

The sampling process is informative on both, B_{τ_i} and $B_{\tau_{i+1}}$ up to a bounding interval defined in either (15/1.a) or (16/2.a) so that the following applies:

$$\begin{aligned} L(B_{\tau_i}) & \leq P(A \leq B_{\tau_i} | PR) \leq U(B_{\tau_i}) \\ L(B_{\tau_{i+1}}) & \leq P(A \leq B_{\tau_{i+1}} | PR) \leq U(B_{\tau_{i+1}}) \end{aligned} \tag{17}$$

$\Rightarrow$

$$L_{\tau, \kappa} \leq P(A \in A_{\tau, \kappa} | PR) \leq U_{\tau, \kappa}$$

where

$$\begin{aligned} L_{\tau, \kappa} &= L(A_{\tau, \kappa}) = \min\{L(B_{\tau_i}), L(B_{\tau_{i+1}})\} \\ U_{\tau, \kappa} &= U(A_{\tau, \kappa}) = \max\{U(B_{\tau_i}), U(B_{\tau_{i+1}})\} \end{aligned}$$

Compared to expression (9) the results in (17) show that allowing for anchoring effects $P(A \in A_{\tau, \kappa} | PR) = p_{\tau, \kappa}$ is only known up to an interval whose width equals $U(A_{\tau, \kappa}) - L(A_{\tau, \kappa})$. For $B_{\tau_{i+1}} > B_{\tau_i}$ the width is always positive.¹³ Knowledge of $p_{\tau, \kappa}$ up to an interval has implications on the

¹³This can be easily shown: Write $P(A \in A_{\tau, \kappa} | PR) \in \{\min[L_{\tau_i}, L_{\tau_{i+1}}], \max[U_{\tau_i}, U_{\tau_{i+1}}]\}$. By definition $(U_{\tau_i} - L_{\tau_i}) > 0$ and $(U_{\tau_{i+1}} - L_{\tau_{i+1}}) > 0$. Thus ambiguity can only occur in two situations, namely, when $P(A \in A_{\tau, \kappa} | PR) \in \{\min = L_{\tau_i}, \max = U_{\tau_{i+1}}\}$ or when $P(A \in A_{\tau, \kappa} | PR) \in \{\min = L_{\tau_{i+1}}, \max = U_{\tau_i}\}$. If both B_{τ_i} and $B_{\tau_{i+1}}$ are perceived as high bids (according to Assumption 1), the use of (16) and (17) shows that in the case of

properties of partial information with respect to $g_{\tau,\kappa} = E[g(y) | A \subseteq A_{\tau,\kappa}, PR]$ as given in (10). In the presence of partial information and anchoring effects the sampling process reveals $E[g(y) | PR]$ but does not reveal $E[g(y) | PR, A_{\tau,\kappa}]$. Thus, we need to apply (3) but with respect to $p_{\tau,\kappa} \in [L_{\tau,\kappa}, U_{\tau,\kappa}]$ and $\Psi(p_{\tau,\kappa})$, where the latter is also defined as a proper subset of Ψ . Thus, applying (3) for $p_{\tau,\kappa}, g_{\tau,\kappa}(p_{\tau,\kappa})$ and given that $E[g(y) | PR]$ as $E[g(y) | PR] = E[g(y) | PR, A_{\tau,\kappa}]p_{\tau,\kappa} + E[g(y) | PR, \bar{A}_{\tau,\kappa}](1 - p_{\tau,\kappa})$, the analogue to expression (4) but with respect to expression $E[g(y) | PR]$ implies the following bounds:

Given

$$\begin{aligned}
 E[g(y) | PR, A_{\tau,\kappa}] &\in \Psi(p_{\tau,\kappa}) \equiv \Psi \cap \left\{ \frac{E[g(y) | PR] - \psi_{\tau,\kappa} \times (1 - p_{\tau,\kappa})}{p_{\tau,\kappa}}, \psi_{\tau,\kappa} \in \Psi(p_{\tau,\kappa}) \right\} \\
 \text{for } g_0(p_{\tau,\kappa}) &\equiv \inf [\psi_{\tau,\kappa} : \psi_{\tau,\kappa} \in \Psi(p_{\tau,\kappa})], \quad g_1(p_{\tau,\kappa}) \equiv \sup [\psi_{\tau,\kappa} : \psi_{\tau,\kappa} \in \Psi(p_{\tau,\kappa})] \\
 &\Rightarrow \\
 E[g(y) | PR, A_{\tau,\kappa}] &\in \{g_0(p_{\tau,\kappa}), g_1(p_{\tau,\kappa})\}
 \end{aligned} \tag{18}$$

Bounds in (18) substitute those in (4) when dealing with partial (rather than full) nonresponse information and anchoring effects. Allowing for (17) and (18) in the bounding intervals implied by (10) leads to a new set of bounds:

$\{\min = L_{ri}, \max = U_{ri+1}\} = \{P(Q_{ri} = 0 | \dots), 1\}$ and $\{\min = L_{ri+1}, \max = U_{ri}\} = \{P(Q_{ri+1} = 0 | \dots), 1\}$; in both cases the width between upper and lower limits is positive. If both B_{ri} and B_{ri+1} are perceived as low bids, $\{\min = L_{ri}, \max = U_{ri+1}\} = \{0, P(Q_{ri+1} = 0 | \dots)\}$ and $\{\min = L_{ri+1}, \max = U_{ri}\} = \{0, P(Q_{ri} = 0 | \dots)\}$, so that once more the distance between paired upper and lower limits is positive. Finally, if B_{ri} and B_{ri+1} are perceived as low and high, respectively, $\{\min = L_{ri}, \max = U_{ri+1}\} = \{0, 1\}$ and $\{\min = L_{ri+1}, \max = U_{ri}\} = \{0, 1\}$.

For each $A_{\tau,\kappa} \subseteq X$,

Let $p \in [0,1]$, $p_{\tau,\kappa} \in [L_{\tau,\kappa}, U_{\tau,\kappa}]$,

allow for (4) and (18) to define $\{g_0(p), g_1(p)\}$ & $\{g_0(p_{\tau,\kappa}), g_1(p_{\tau,\kappa})\}$

$\Rightarrow$

$$\inf_{p, p_{\tau,\kappa}} \left\{ \begin{aligned} & E[g(y) | A, FR] \times \left(\frac{P(A|FR)P(FR)}{P(A|FR)P(FR) + P(PR)p_{\tau,\kappa} + P(NR)p} \right) \\ & + (g_0(p_{\tau,\kappa})) \times \left(\frac{P(PR)p_{\tau,\kappa}}{P(A|FR)P(FR) + P(PR)p_{\tau,\kappa} + P(NR)p} \right) \\ & + g_0(p) \times \left(\frac{P(NR)p}{P(A|FR)P(FR) + P(PR)p_{\tau,\kappa} + P(NR)p} \right) \end{aligned} \right\} \\ \leq E[g(y) | A] \leq \\ \sup_{p, p_{\tau,\kappa}} \left\{ \begin{aligned} & E[g(y) | A, FR] \times \left(\frac{P(A|FR)P(FR)}{P(A|FR)P(FR) + P(PR)p_{\tau,\kappa} + P(NR)p} \right) \\ & + (g_1(p_{\tau,\kappa})) \times \left(\frac{P(PR)p_{\tau,\kappa}}{P(A|FR)P(FR) + P(PR)p_{\tau,\kappa} + P(NR)p} \right) \\ & + g_1(p) \times \left(\frac{P(NR)p}{P(A|FR)P(FR) + P(PR)p_{\tau,\kappa} + P(NR)p} \right) \end{aligned} \right\} \quad (19)$$

Expression (19) completes all bounding intervals of interest. Bounds in (19) are sharp assuming anchoring effects and the implications in Jacowitz and Kahneman (1995).

Three sets of empirically implementable bounds have been considered; (5'), (10) and (19). Bounds in (5') are wider than those in (10); this is clear if we think that $P(NR | PR) < P(NR)$. However, both (5') and (10) are p -dependent, where (10) is minimized and maximized over similar values of p but different conditional domains (since (10) includes partial information). Thus, it is not possible to define with a specific expression the difference in width between sets of bounds. Comparing bounds (10) and (19) we see that minimization and maximization of bounds in (19) is with respect to both p and $p_{\tau,\kappa}$: Both, $p_{\tau,\kappa}$ and the expectations $g(p_{\tau,\kappa})$ in (19) imply intervals describing the uncertainty as result of anchoring effects on partial respondents. The uncertainty filters through and widens the interval between upper and lower bounds in (19) compared to the interval in (10) where $p_{\tau,\kappa}$ and $g(p_{\tau,\kappa})$ are assumed to be known with certainty.

Allow for partial information with incomplete bracket respondents (IBR)

Initial non-respondents who decide to provide partial information may not necessarily complete the sequence of J conditional bids. Table 1 illustrates the possibility of incomplete bracket respondents

(IBR) when individuals decide to stop providing information beyond the bid $B_1 = 25,000$. First it is easy to observe that in the case of IBR the resulting number of categories shrinks. In general, individuals who decide to stop responding (either ‘don’t know’ or ‘refuse’) at the $(J - d)th$ category have the chance to classify their otherwise undisclosed amount in as many as $2^{(J-d)}$ categories. Categories resulting from IBR may share values defining the upper and lower limits of categories formed by CBR. However, in terms of *categorical range*, comparing the two types of partial respondents implies that none of the $2^{(J-d)}$ categories from IBR match the 2^J categories from CBR, and vice versa. This means that it will never be possible to join CBR and ICB information to bound the expectation $E[g(y) | A]$ when A refers to some category in X . When A refers to cumulative ranges in X so that bounds apply to $E[g(y) | A \leq B_r]$ it is possible to join information from IBR and CBR since the conditional bids B_r (and not the range and location of some category) is what matters. To deal with this situation, simplify the notation in (5’), (10) and (19) so that in all three cases $LB_{r,\kappa} \leq E[g(y) | A_{r,\kappa}] \leq UB_{r,\kappa}$ express the implied bounding intervals. In the presence of both CBR and IBR the following applies:

$$E[g(y) | A \leq B_r, PR] = E[g(y) | A \leq B_r, CBR, PR] \times P(CBR | A \leq B_r, PR) + E[g(y) | A \leq B_r, IBR, PR] \times P(IBR | A \leq B_r, PR) \quad (20)$$

Assuming no anchoring effects means applying the consequence of (9) to $P(CBR | A \leq B_r, PR)$, $E[g(y) | CBR, A \leq B_r, PR]$, $P(IBR | A \leq B_r, PR)$ and $E[g(y) | IBR, A \leq B_r, PR]$, where (9) implies identification of the event $(A \leq B_r)$ without ambiguity; i.e., the sampling process identifies all elements in (20) such that $p_C = P(CBR | A \leq B_r, PR)$, $p_I = P(IBR | A \leq B_r, PR)$, $g_C(y) = E[g(y) | CBR, A \leq B_r, PR]$ and $g_I(y) = E[g(y) | IBR, A \leq B_r, PR]$, i.e., $E[g(y) | A \leq B_r, PR] = g_C(y)p_C + g_I(y)p_I = g(A \leq B_r)$ for the cumulative distribution of X . Allowing for appropriate changes to reflect the conditional distribution in the covariate set, bounds in (10) apply by substituting $g_{r,\kappa}$ for $g(A \leq B_r)$ and $p_{r,\kappa}$ for $p(A \leq B_r | PR)$ at each of the B_r conditional bids.

Allowing for anchoring, on the other hand, implies further partition for each of the elements in (20) and to apply (13) – (16) so that the bounding interval in (19) becomes operational. From (20), we see that $E[g(y) | h, A \leq B_r, PR]$ where $h = CBR$ or IBR can be dealt with using assumptions similar to those implied by (18) because only $E[g(y) | h, PR]$ is revealed by the sampling process. Similar conditions apply to $P(CBR | A \leq B_r, PR)$ because as with $E[g(y) | h, A \leq B_r, PR]$, only $P(CBR | PR)$ is identified by the sampling process. Both $E[g(y) | h, A \leq B_r, PR]$ and $P(h | A \leq B_r, PR)$ require assumptions to allow for

anchoring effects on $P(A \leq B_r | PR)$ as given in (15) and (16). Following all recommendations (20) can be bounded as follows:

$$\begin{aligned}
& \text{Let,} \\
& E[g(y) | A \leq B_r, CBR, PR] \in \{L(g, C), U(g, C)\} \\
& E[g(y) | A \leq B_r, IBR, PR] \in \{L(g, I), U(g, I)\} \\
& \text{and} \\
& P(CBR | A \leq B_r, PR) \in \{L(p, C), U(p, C)\} \\
& P(IBR | A \leq B_r, PR) \in \{L(p, I), U(p, I)\} \\
& \Rightarrow \\
& \{L(g, C) \times L(p, C) + L(g, I) \times L(p, I)\} \\
& \leq E[g(y) | A \leq B_r, PR] \leq \{U(g, C) \times U(p, C) + U(g, I) \times U(p, I)\}
\end{aligned} \tag{21}$$

Bounds in (21) substitute the interval $\{g_0(p_{r,k}), g_1(p_{r,k})\}$ in (19). Likewise, bounds on p_r where $p_r = P(A \leq B_r | PR) = P(A \leq B_r | PR, CBR)P(CBR | PR) + P(A \leq B_r | PR, IBR)P(IBR | PR)$ can be derived by applying both (15) and (16) to both $P(A \leq B_r | PR, CBR)$ and $P(A \leq B_r | PR, IBR)$, while anchoring effects do not impact $P(CBR | PR)$ or $P(IBR | PR)$. Thus, $P(A \leq B_r | PR) \in \{L(p_r), U(p_r)\}$; the latter, together with (21) completes the substitution in (19) thus accommodating joint information from both CBR and IBR, but only if the conditional set implies cumulative distribution with respect to the variable suffering from nonresponse.

Finally we notice that nonresponse can affect variables not just in the covariate set, but it may also affect the outcome variable or both outcome and covariates at the same time. Appendix 1 provides guidelines to extend the bounds in (5'), (10) and (19) for the case when either outcome or both outcome and covariates jointly suffer from nonresponse in the presence of anchoring effects. Thus, our analysis covers all possibilities as those described in Horowitz and Manski (1998).

3 Estimation

Bounds given by (5'), (10) and (19) in Section 2 are explained in terms of population characteristics. For specific $g(y)$ functions, each of these three sets of bounds can easily be estimated using the corresponding sample analogue. Anticipating the empirical illustration, let $g(y)$ be the indicator function for some binary outcome (e.g., $y = 1$ if smoker, 0 if otherwise). For the event $A \subseteq A_{r,k}$, the expectation $E[g(y) | A, r]$, $r = FR, PR, NR$ is obtained by taking sample averages. The same applies to all other probabilities in expressions (5'), (10) and (19). Including other covariates M so that $E[g(y) | A, r, M]$, $r = FR, PR, NR$ could change the estimation methods if M contains a mixture of

continuous and discrete variables. In this case bounds could be estimated non-parametrically, e.g., with the use of product Kernel estimation and an appropriate selection of bandwidths (for example, see Härdle, 1990, as reference for nonparametric regression). The empirical example below considers only one discrete covariate M (gender) so that kernel weights are not required. Thus, the illustration will be based on minimizing and maximizing over quantities estimated using weighted sample averages of a set of probabilities over each of the $A \subseteq A_{\tau, \kappa}$ partitions of a variable ‘earnings’, where the variable ‘earnings’ suffers from non-negligent nonresponse. The use of cross-sectional weights is a requirement for the sample to become representative of the underlying population (see Section 4.1 below).

An important point is to notice that for bounds (5’) and (10), the width between the upper and the lower bound is p – dependent where $p = P(A | NR)$, and for bounds in (19) the width is $p, p_{\tau, \kappa}$ – dependent where $p = P(A | NR)$ and $p_{\tau, \kappa} = P(A | PR)$. We minimize and maximize bounds in (5’) and (10) over values of $p \in (0,1)$ that we choose to partition in 1000 equidistant intervals. In the case of (19), $p \in (0,1)$ is equally partition into 1000 intervals while $p_{\tau, \kappa}$ is partition into 250 equidistant intervals.¹⁴ Estimates of the width between a set of upper and lower bounds describes uncertainty due to item nonresponse thus reflecting a population concept rather than a sampling concept. This means that for each of these bounding intervals we still need to estimate confidence bands to reflect finite sampling error. For any of the three bounding intervals, any pointwise distance in the support of X between the upper confidence band for the upper bound and lower confidence band for the lower bound reflects joint uncertainty due to both sampling error and error due to item nonresponse. This paper makes use of a naïve bootstrap procedure that consists on re-sampling 1000 times (with replacement) from the original data to estimate pairs of upper and lower 95% confidence bands for the upper bound and for the lower bound for each of the three sets of bounds. In all cases, the 95% confidence bands displayed in the empirical section are the lower confidence band for the lower bound and the upper confidence band for the upper bound.

4 Illustration: Health habits and earnings

4.1 The Data

We use the 1996 wave of the Health and Retirement Study (HRS) to illustrate the empirical implications of (5’), (10) and (19). The HRS is a longitudinal study conducted by the University of Michigan on behalf of the US National Institute of Aging. This biannual panel has been active since 1992 collecting extensive information on aspects of health, wealth, retirement and other important socio-economic conditions for a representative cohort of US citizens born between 1931 and 1941.

¹⁴ Different choices of partition do not change the results in any significant way.

We select of the 1996 wave for two reasons. First, the 1996 wave applies an unfolding bracket design to elicit further earning information in the event of initial earnings nonresponse. Second, the aim in this illustration is to understand the health behaviour of individuals over the distribution of earnings. Maximizing the number of individuals who are still earners implies picking a wave where all individuals in the representative sample are age 65 or below. Since the HRS weights only individuals born between 1931 and 1941, the 1996 appropriately picks up individuals with ages between 55 and 65, thus, the 1996 wave provides an appropriate source of information to illustrate the three sets of bounds in Section 2.

Initially (1992) the panel surveyed 7,600 households. The 1996 has data from 6,736 household resulting in 8,436 eligible individuals.¹⁵ Each household has a representative member who answers all questions relating to own earnings and earnings from other surveyed household members (e.g., the spouse). The first sample selection criterion consists on selecting only those who declare to be household representatives.¹⁶ There are 6,075 household respondents and of these 2,862 declare to have worked for wages and salary (i.e., we avoid self-employed) during the last calendar year (i.e., in 1995).¹⁷ These 2,862 individuals provide the basis for the empirical illustration which consists on studying the probability of being associated with an adverse health habit indicator by earning's profile. That is, we aim at bounding $E[g(y) | A \in A_{r,k}, M]$ where $E[g(y) | \dots] = P(H = 1 | \dots)$ and $H_i = 1$ for $i \in n$ indicates that the i th individual in the sample of size n is associated with at least one of these three events: smoking, heavy alcohol consumption, or obesity. All 2,862 sample units respond to questions relating to smoking behaviour and alcohol consumption, thus we are able to build indicators for these outcomes without nonresponse problems. The variable BMI¹⁸ measures relative weight. Only 5 of the 2,862 individuals fail to declare one or both of the items needed to construct the BMI. This paper does not deal with misreporting or measurement error, thus, we assume that all three health habit measures are free from self-reporting error.¹⁹ The covariate M indicates 'gender'.

¹⁵ The HRS collects information from representative respondents (i.e., born between 1931 and 1941) and their partners (or closely related household members). The sample is representative using cross sectional weights, and those who are not age or sample eligible are given a weight of zero. In 1996 only 8,436 individuals (of 27,109 thus far interviewed since 1992) have a positive weight assigned to them.

¹⁶ We are dealing with nonresponse. Not responding to 'own earnings' may be due to reasons completely different to not responding to 'other person's earnings'. For example, reporting earnings from a partner can be affected by the fact that a partner is (or not) present in the household at the time of the interview, itself possibly as consequence of working, or even health reasons.

¹⁷ At this point the problem of nonresponse is considered negligent. That is, 3,405 of the initial 6,075 individuals declared to have worked for pay during the last calendar year. When the 3,405 are asked if they worked for 'wages & salary' (i.e., for earnings as employees), 2,862 declare to be wage/salary earners, 549 declare self-employment and only 6 declare 'don't know'. Using cross sectional weights this amount to a 0.0018 probability of nonresponse, and this can be considered negligent nonresponse.

¹⁸ Body Mass Index (BMI) is the most widely used and best known obesity measure and consists on balancing the body mass against height. $BMI = (weight/height^2)$.

Table 2: Sample statistics for a selection of Health behaviour indicators by gender (standard errors in brackets)

	Full Sample n=2,862	Males N=1,036	Females n=1,826
% Smokers	22.1 (0.8)	24.4 (1.3)	20.7 (1.2)
% No alcohol consumption	49.9 (0.9)	39.1 (1.5)	56.6 (1.2)
% Moderate to low alcohol consumption	46.1 (0.9)	53.0 (1.6)	41.9 (0.5)
% heavy alcohol consumption	4.0 (1.2)	7.9 (1.2)	1.5 (0.3)
Average BMI index	27,1 (0.1)	27,4 (0.1)	26,9 (0.1)
% with BMI above 16 and below 25 (correct BMI)	36.7 (0.7)	28.2 (1.4)	41.9 (1.2)
% with BMI below 16 (underweight)	0.15 (0.9)	0.07 (0.8)	0.20 (0.9)
% with BMI equal or above 25 (overweight)	63.2 (0.8)	71.7 (1.4)	57.9 (1.2)
% with BMI equal or above 30 (obesity)	23.2 (0.3)	21.9 (1.3)	23.9 (1.0)
% with BMI equal or above 40 (morbid obesity)	2.2 (0.9)	1.5(0.4)	2.6 (0.3)
% rating health as good or very good	57.5 (1.2)	57.3 (2.0)	57.7 (1.6)
% rating health as good or very good if smoke	49.7 (1.3)	47.6 (2.1)	51.2 (1.6)
% rating health as good or very good if heavy alcohol	49.9 (1.3)	46.3 (2.1)	61.7 (1.5)
% rating health as good or very good if overweight	42.2 (1.2)	46.0 (2.1)	40.1 (1.6)

Note 1: Smokers are anyone that declares 'yes' when asked if they are current smokers. This dummy variable does not distinguish between type of smokers (e.g., cigarettes or pipe) or the intensity. Heavy alcohol consumers are those who declare to drink 4 or more days per week and, on average, 4 or more units of drink per drinking day (since 'binge' drinking is defined in the HRS as having +4 drinks per day on a drinking day). BMI is calculated with the standard formula (weight/tallness*tallness). According to the WHO guidelines a healthy weight is that of a BMI between +16 and below 25. A BMI below 16 indicates underweight and a BMI above 25 indicates clinically overweight. Obesity occurs at a BMI of 30 or above (www.who.ch). The HRS provides self-reported health status so that individuals declare if they feel their health to be very good, good, fair or very poor. The dummy 'health rating' by health habit groups individual into 2 groups according to 'very good' or 'good' and otherwise.

Note 2: All estimates are weighted estimates using the personal cross-sectional weights as provided by the HRS 1996 tracker file.

Table 2 shows sample statistics with reference to $H = \{smoking, alcohol, obesity\}$ for the full sample and for the sub-samples according to gender.²⁰ All estimates are weighted estimates. These summary statistics show that the 22% of the USA population born between 1931 and 1941 are smokers with the percentage increasing significantly for males (24.4%) relative to females (20.7%). The percentage of heavy alcohol drinkers in the population is low (in magnitude) with a significant difference between males and females (7.9% versus 1.5%, respectively). A weighted average BMI of 27,1 indicates the well documented problem of obesity in the Western World, in this case for the USA.²¹ The same figures show that some 20% of the population suffers from obesity or morbid obesity. These estimates also show that although males are more prone to be overweight (71.7%) than females (57.9%), the percentage of clinically obese among the overweight is almost identical between genders (21.2% for males and 23.9% for females). Finally we see that 57.5% of those in the overall sample rate their health as good or very good. However, if we look at sub-samples defined according to each of the adverse health habits the same estimate can drop by as much as 15%. The drop in 'good health rating' is larger for those with overweight (who may or may not be smokers or drinkers) than the drop observed when considering smokers and heavy alcohol consumers. Thus, it seems that feeling unhealthy is potentially associated with following bad health habits.

¹⁹ Many of the variables in the health section of the survey rely on self-report, and this is the case with smoking (and the amount), alcohol drinking (and the amount), body mass (weight) and tallness (height).

²⁰ Obesity is not so much a health habit as an indicator for adverse health habits (e.g., diet, sedentary life, etc)

²¹ For example, see The Economist, May 9th, 2003 for a related survey

Table 2 examines the health habits in the population of HRS respondents and compares these health habits between genders. What might be more interesting is to look at how these health habits vary as earnings vary, allowing for a comparison within gender for different earnings level and between genders for similar earning's outcome. The relation between health habit outcomes and labour market indicators has received a good deal of attention among social scientist. For example Zarkin et al. (1998) on the relation between alcohol use and earnings, Levine et al. (1997) on the relation between smoking and earnings, or more recently the study by van Ours (2004) that looks at the effect that both smoking and alcohol drinking may have on earnings. In the illustration that follows we are not concerned with causality. In fact, at the ages of 55 to 65 individual's earnings and health habits can be thought as outcomes at their steady state. The empirical illustration is simply concerned with understanding if health habits vary between groups who differ in terms of earnings, and to understand the consequence of making inference when earning's nonresponse is accounted for. Thus we deal with estimates of $P(H = 1 | A \in A_{\tau,k}, M)$ for a definition of H and different earning's categories $A_{\tau,k}$.

The variable $X = \text{earnings}$ is not fully observed for all individuals in the selected sample as it suffers from nonresponse for a non-negligent percentage. Initially all 2,862 respondents who have declared to have worked for wages and salaries over the last calendar year face the following open-ended question:

$$\begin{aligned}
 & \text{'about how much wage and/or salary income did you receive during the last} \\
 & \text{calendar year?'} \\
 & \quad \text{'.....any amount' (in USA dollars)} \\
 & \quad \text{'.....don't know'} \\
 & \quad \text{'.....refuse'}
 \end{aligned} \tag{22}$$

Out of the 2,862 individuals, 2,508 become full respondents (FR) because they answer with a specific amount to question (22), and the remaining 354 are the *initial* non-respondents (NR) who are then routed to an unfolding bracket design: this implies an initial (weighted) nonresponse probability of 0.12. The 354 initial non-respondents face question (7), all with the same bid of \$25,000 to start up the unfolding bracket design. At this point 86 are classified as full non-respondents (thus, NR is now an event associated with these individuals) since they declare 'don't know' or 'Refuse to respond' to the initial bid and, consequently, are not driven through the unfolding design by the interviewer. The rest of initial non-respondents are such that 255 complete the sequence of brackets (CBR) and 13 leave the sequence (IBR) before this is completed (i.e., they report either 'don't know' or 'Refuse to respond' at some bid before the design ends).²² As described in Section 2, considering earnings intervals implies that

²² It is worth noting at this point that survey response processes often distinguish between types of nonresponse, for example, nonresponse with 'don't know' (that might be due to true inability to retrieve information) and 'refuse' to respond, possibly

CBR and IBR become separate samples that need separate analysis, simply because IBR individuals define categories that differ in both width and range when comparing these to the categories resulting from CBR. The nature of the bounds in both cases would be identical so we dismiss the 13 individuals in the IBR group and deal with bounds with reference to the CBR sub-sample. The new sample of 2,849 is distributed such that $P(FR) = 0.89$ (2,508 units), $P(NR) = 0.031$ (86 units), and $P(PR) = 0.081$ (255 units), where all in the PR group are CBR. The distribution between genders is such that 1,032 are males and 1,817 are females. The unfolding bracket design in the 1996 HRS for the variable earnings starts with one initial bid of $B1 = \$25,000$ followed by either one or two conditional bids depending on the route followed by the individual along the sequence. Those who faced with $B1 = \$25,000$ answer to (7) such that $Q1 = 1$ are further routed to a second bid $B21 = \$50,000$; $B21$ indicating the 2nd bid conditional on having answered ‘yes’ to $B1$. If $Q21 = 1$ at $B21$ individuals are driven to a 3rd bid $B311 = \$100,000$ and the sequence stops. However, the sequence stops at $B21$ if $Q21 = 0$. Likewise, if $Q1 = 0$ individuals are routed to $B20 = \$5,000$ and this is also the last bid faced by partial respondents who declare their earnings to be below \$25,000. Table 3 illustrates the complete sequence and resulting categories. Table 4 complements Table 3 with the sample distribution of CBR among the categories for the full CBR sample and by gender sub-groups.

associated with the information that the interviewer tries to elicit (see, for example, Tourangeau et al, 2000). In our case we eliminate either ‘don’t know’ answers or ‘refuse’ in the case of IBR. The final 86 full non-respondents contribute equally to the uncertainty created by their no-information behaviour irrespective of what drives this behaviour, so we need not distinguish between ‘don’t know’ answers and ‘refuse to respond’ answers.

Table 3: Distribution of CBR to an unfolding bracket design for the variable 'annual earnings', Health and Retirement Study, (HRS, 1996)

Group	1 st BID; B1	Answer to bid	2 nd BID: B20 or B21	Answer to BID	3 rd BID: only for Q21:B311	Anchor 2: B20 or B21	Resulting Bracketed categories
CBR	>\$25,000?	YES Q1=1	>\$50,000?	YES Q21=1	>\$100,000	YES (end) Q311=1	[100,000; INF)
				NO (end) Q21=0		NO (end) Q311=0	[50,000; 100,000) [25,000; 50,000)
		NO Q1=0	>\$5,000?	YES (end) Q20=1			[5,000; 25,000)
				NO (end) Q20=0			[0; 5,000)

Note: This table is the complete version of the unfolding bracket design that served as example in Table 1. There is no natural upper bound to earnings other than to allow for it to be infinity (INF). However, restricting the example to cover employees, the natural lower bound for the lowest located category is earnings=0. Including self-employees would not had made possible this assumption since self-employed could in principle receive negative earnings in which case the lowest category would had been (-INF; 5,000).

Table 4: Distribution of initial Non-respondents between full non-respondents (NR) and complete Bracket respondents (CBR) by resulting categories on Annual Earnings (HRS, 1996)

Category	Full Sample n= 341	Males n= 90	Females n= 251
[100,000; INF)	n=4, %= 0.016 (0.007)	n=1, %= 0.015 (0.013)	n=3, %= 0.017 (0.008)
[50,000; 100,000)	n=19, %= 0.059 (0.013)	n=5, %= 0.046 (0.022)	n=14, %= 0.064 (0.015)
[25,000; 50,000)	n=69, %= 0.208 (0.022)	n=23, %= 0.246 (0.045)	n=46, %= 0.189 (0.025)
[5,000; 25,000)	n=134, %= 0.364 (0.026)	n=29, %= 0.308 (0.049)	n=105, %= 0.387 (0.031)
[0; 5,000)	n=29, %= 0.075 (0.014)	n=4, %= 0.032 (0.019)	n=25, %= 0.093 (0.018)
Full Non-respondents (NR)	n=86, %= 0.277 (0.024)	n=28, %= 0.343 (0.050)	n=58, %= 0.251 (0.027)

Note: See notes in Table 3. All percentages are with respect to the 341 members, or the 90 males or the 251 females in the sample of partial respondents excluding IBR. In each case the samples are weighted with personal level cross-sectional weights (HRS 1996, tracker file). Standard errors in brackets; 'n' indicates number of observed units per category.

Table 3 shows that the unfolding bracket design results in 5 categories with [\$100,000; INF) as the highest ranking and [0; \$5,000) as the lowest ranking category.²³ Table 4 shows the distribution of the initial non-respondents (341) according to resulting categories and resulting full nonresponse. Units of females outnumber units of males in the CBR group, but weighted probabilities show that relative to the complete sample of 341 there is an even representation of genders by categories. Full nonresponse is

²³ Just as a comparative note, it is worth mentioning that answers to the design from the 13 incomplete bracket respondents (IBR) result in the following set of categories: [\$0; \$25,000), [\$25,000; ∞) and [\$50,000; ∞), with 6, 5 and 2 units in each, respectively. Clearly, these categories are not compatible with categories in Table 3.

significantly higher for males than for females. At this point it is important to examine the probability of being at particular ‘nodes’ defined by Table 3, that is, we need to examine the sample properties of $P(Q1=1|PR)$, $P(Q21=1|Q1=1,PR)$, $P(Q311=1|Q21=1,PR)$ and $P(Q20=1|Q1=0,PR)$ so that Assumption 1 in (14) becomes operational for each of the bids $\{\$5,000, \$25,000, \$50,000, \$100,000\}$, a necessary condition to develop and estimate bounds in (19):

Table 5: Sample Probabilities for different nodes in Table 3 given by CBR to Annual Earnings, (HRS 1996)

Probability at nodes defined by Table 3	Full Sample n= 255	Males n= 62	Females n= 193	Conclusions with respect to HIGH or LOW ANCHOR
$P(Q1 = 1 PR)$	n=92/255 %= 0.283 (0.028)	n=29/62 %= 0.317 (0.059)	n=63/193 %= 0.217 (0.030)	Q1=1 → B1 (\$25,000) HIGH ANCHOR
$P(Q21 = 1 Q1 = 1, PR)$	n=19/92 %= 0.265 (0.048)	n=5/29 %= 0.192 (0.073)	n=14/63 %= 0.299 (0.058)	Q21=1 → B21 (\$50,000) HIGH ANCHOR
$P(Q311 = 1 Q21 = 1, PR)$	n=4/23 %= 0.217 (0.086)	n=1/6 %= 0.252 (0.178)	n=3/17 %= 0.206 (0.098)	Q311=1 → B311 (\$100,000) HIGH ANCHOR
$P(Q20 = 1 Q1 = 0, PR)$	n=134/163 %= 829 (0.029)	n=29/33 %= 0.906 (0.051)	n=105/130 %= 0.807 (0.035)	Q20=1 → B20 (\$5,000) LOW ANCHOR

Note: See notes in Table 3. All percentages are with respect to the corresponding sub-groups. The indicator 'n' shows the total units at each Q_i given the number of units in the corresponding space. The probabilities are based on weighted averages. Bracketed numbers show standard errors.

Thus, from Table (5) and assuming the Jacowitz and Kahneman (1995) assumption we conclude that for the case of CBR, the bids B1=\$25,000, B21=\$50,000 and B311=\$100,000 act as ‘high anchors’ and the bid B20=\$5,000 acts as a ‘low anchor’. This information drives the characterization of (19) in Appendix 2.

Finally, Table 6 examines the variable ‘earnings’ showing two distinct sets of information. The second row in Table 6 refers to weighted estimates using full respondents only thus assuming exogeneity with respect to nonresponse. The last row in Table 6 is based on the weighted sample when both non-respondents and partial respondents have had their missing values imputed. The imputed variable is constructed at source (HRS, 1996 imputed variables files) using imputation methods as described by the HRS,1996 documentation; it should be noted that this imputation methods makes use of partially provided information for those in the PR group. The empirical section checks the properties of the imputed variable against estimates of (5’). The point is that (5’) are bounds derived imposing no assumptions other than those implied by the natural upper and lower limits, therefore although the widths in (5’) cannot pin down a particular point estimate for $P(H = 1 | A \in A_{\tau,k}, M)$, with probability one the region contains the unknown population probability; if imputed values are not incorrectly imputed, sample estimates based on these should fall within the bounding interval implied by (5’). On the other hand bounds in (10) and (19)

imply, respectively, excluding or including the assumption of anchoring. Bounds given by (10) may exclude estimates based on the imputed values as consequence of strong anchoring effect which (10) ignores, thus, we cannot test the goodness of imputation by means of estimates on (10). Bounds in (19) relax the assumption of no anchoring thus widening the width between upper and lower bounds. However, estimates of (19) may result on a location shift when compared to the interval implied by (10). This depends on what is the effect that the Jacowitz and Kahneman assumption has on them. Thus, empirical estimate of $P(H = 1 | A \in A_{t,\kappa}, M)$ based on imputed values can be checked to see if they fall inside the estimated region given by (19), but any conclusions that this comparison may imply has to hold one of the assumptions (either anchoring effects or imputation procedures) as constant and valid.

Table 6: Sample statistics of gross annual earnings (standard errors in brackets)

	Full Sample	Males	Females
Earnings, FR only (n=2,508: male = 942, female = 1,566)			
Average (standard error)	\$31,342 (566)	\$36,202 (942)	\$28,148 (695)
Range	min=\$0 max=\$350,000	min=\$0 max=\$300,000	min=\$0 max=\$350,000
Imputed earnings (n=2,849: male = 1,032, female = 1,817)			
Average (standard error)	\$30,824 (526)	\$35,274 (879)	\$28,099 (649)
Range	min=\$0 max=\$350,000	min=\$0 max=\$300,000	min=\$0 max=\$350,000

Note 1: Estimates are based on the weighted sample. The variable 'earnings' refers to 'gross earnings' and is based on an initial open-ended question where individuals who declare to have worked over the last calendar year (i.e., 1995) are asked to declare the total gross amount of earnings received from such work. Amounts refer to nominal US dollars. The original variable draws from the publicly accessible HRS- 996 file named HR96J_H. The derived variable with the imputed amounts draws from the file named H96i_jh.

Table 5 suggests that males earn significantly more than females (the t-value on the difference is $\hat{t} = 6.9$). Imputation of missing earnings implies a slight drop in mean values for the full sample and for the sub-sample of males, but has weaker effects on mean earnings for females. The standard errors do not differ much before and after imputation, thus reflecting small noise induced from imputing non-respondents and partial respondents. This feature is expected when imputation is performed in the support of the full response sample.

4.2 Specification

Section 2 provides the general expression for bounds that allow for partial information when initial respondents are routed to an unfolding design with J possible conditional bids resulting in 2^J categories. It is also generalized for any $g(y)$ function. Table 3 shows that for the specific case under consideration J is at most 3 and in the case of CBR the number of resulting categories are $1 + 2^2 = 5$; the '1' accounts for the special case of $B3$ affecting only those who eventually classify their earnings in the upper tail of the earnings distribution. With 5 categories for this specific example the following

applies; $B: B1, B20, B21 \text{ and } B311 \Rightarrow Q: Q1, Q21, Q20 \text{ and } Q311$. The result is such that $X = \{A_{1,0}, A_{1,1}, A_{2,0}, A_{2,1}, A_{3,1}\} = \{[0, B20), [B20, B1), [B1, B21), [B21, B311), [B311, \infty)\}$, or in terms of currency, $X = \{[0, \$5,000), [\$5,000; \$25,000), [\$25,000; \$50,000), [\$50,000; \$100,000), [\$100,000, \infty)\}$. For each $A_{\tau,k}$ category in X , we want to estimate the sample analogue of $P(H=1 | A \in A_{\tau,k}, M)$ for different M gender sub-groups and for a health outcome H derived from information in Table 2. Appendix 2 characterize (5'), (10) and (19) for the particular case when the unfolding design follows the example as in Table 3 and for $g(y)$ defined as an indicator function of health habit behaviour.

4.3 *Bounds on the probability of following adverse health habits*

Studies that look at the relation between health habits (e.g., smoking, alcohol consumption, etc.) and issues such labour market participation, earnings, and other labour market outcomes, tend to disaggregate the sample between males and females thus assuming systematic differences between genders either with regards to health habits or outcomes (e.g., van Ours, 2004). Thus, an interesting question is to understand how different genders are with respect to adverse health habit behaviour comparing these differences between individuals with similar earnings.

This section compares the proportion of individuals associated with at least one of three adverse health habit indicators (smoking, heavy alcohol drinking and obesity) by earning's profile and allowing for gender difference.²⁴ Table 7 shows summary statistics for the variable 'adverse health behaviour indicator (AHB)'. When compared to females, males are 3.3% (p-value = 0.043) more likely to be associated with at least one of the three health habits indicators (i.e., to be associated with the event $AHB=1$).

²⁴ We are comparing 'annual earnings' between genders where the variable 'earnings' is assumed to measure similar intensity at work for all the members in the sample. This would require that the average number of months worked per calendar year is similar among those with similar earnings. We looked at months worked over 1995 and found that on average individuals work 10.8 months per year. The extreme values of months worked were between 2 and 12 but we also saw similarity number of months worked by earnings categories (e.g., those working below 4 months are associated with the lowest category of earnings). It is not possible for us to compare individuals by 'monthly averages' because this would distort the lower and upper limits for the categories and information from partial respondents would need to be disregarded.

Table 7: Sample statistics for the indicator ADVERSE HEALTH HABIT (standard errors in brackets)

	Full Sample N=2,849	Males N=1,032	Females n=1,817
% Smokers	0.22 (0.008)	0.25 (0.013)	0.21 (0.009)
% heavy alcohol consumption	0.04 (0.004)	0.08 (0.008)	0.02 (0.003)
% with BMI equal or above 30 (obesity)	0.23 (0.008)	0.22 (0.013)	0.24 (0.010)
% associated with an Adverse Health Habit indicator (AHB) (either smoke, engage in heavy alcohol drinking or are obese, or any combination of the three)	0.433 (0.010)	0.453 (0.015)	0.420 (0.012)

See notes in Table 2.

Our problem is that 12% in the representative sample do not declare annual earnings, although they have declared to have worked for pay during the corresponding calendar year. Empirically, we deal with nonresponse in two distinct ways; we allow for imputation to overcome the nonresponse problem, and we estimate bounding intervals that imply no assumptions or very weak data assumptions on the response behaviour of the sample. Table 8 shows estimates on the probability of AHB either as precise point estimates or allowing for bounding intervals to identify a region for the same probability. Point estimates are shown with standard errors. We choose to display the estimated bounds allowing for 95% confidence intervals, i.e., each cell in Table 8, Columns 2 to 4, contains upper and lower bounds that take into account total uncertainty as result of both sampling error and error due to earnings nonresponse.

Column 1 shows estimates based on imputed values. Column 2 shows estimates of bounds based on (5'). First we notice that estimates in Column 1 (for any given earning's category and for either gender) fall inside the estimated intervals (with 95% confidence) implied by bounds in (5'). Thus, imputation does not distort the sample estimates on the probability of AHB away from the region that contains the unknown population probability. We conclude that the imputation procedure followed by the data management of the HRS cannot be rejected, although our check is no guarantee that imputation has identified the underlying true distribution for the variable earnings.

Our aim is to use the bounding interval as a tool to test for gender difference in AHB controlling for difference in earnings. When comparing genders with regards to the (bounded) identification regions on AHB, three things can happen: (i) the two bounded regions (males and females) never overlap, (ii) the overlap is partial, or (iii) the overlap is total so that a 95% confidence region identified for females contains that identified for males, or vice versa. Total overlap within a category of earnings implies that the null of equality between genders with regards to the probability of AHB cannot be rejected: estimates of bounds (5') in Column 2 show this to be the case for the categories [\$50,000; \$100,000) and [\$100,000; \$max). For example, in the case of [\$50,000; \$100,000) the probability that females are associated with AHB is bounded between 12.7% and 60.3%, whereas for the same earning's category male's probability is bounded between 29.6% and 60.0%. Thus, allow for 95% confidence, we cannot

reject the null that both genders might share the same unknown population proportion with regards to the probability of AHB. For categories [\$0; \$5,000), [\$5,000; \$25,000) and [\$25,000; \$50,000) the identification region for males is above the analogue region for females, but they still overlap: for example, for the category [\$0; \$5,000) the probability of AHB for females is bounded between 17.1% and 69.5% whereas for males the same probability is bounded between 21.1% and 82.1%. Clearly the upper bound for females (69.5%) is above the lower bound on the region for males (21.1%) while the distance between upper and lower bounds (48.4%) and significantly different than zero, thus the overlap is significant. We conclude that for earners in the [\$0; \$5,000) category the null of gender equality with regards AHB cannot be rejected.²⁵ Similar conclusions would apply when comparing males and females for the earnings categories [\$5,000; \$25,000) and [\$25,000; \$50,000); in both cases there is a positive and significant overlap. The overall conclusion from Column 2 in Table 8 is that allowing for any type of earnings nonresponse (but ignoring partial respondents) the null of no difference between males and females with regards to the probability of AHB cannot be rejected for all levels of earnings. A slightly different picture emerges if we estimate bounds allowing for partial respondents but ignore anchoring effects: Column 3 shows estimates of bounds in (10). For the categories [\$0; \$5,000), [\$5,000; \$25,000) and [\$25,000; \$50,000) the identification regions on the probability of AHB for males is above that of the identification regions for females, although the overlap is only significantly different from zero for individuals in the [\$0; \$5,000) category where the overlap implies a positive and significant distance of 48.4% (p-value= 0.00). For the other two categories, [\$5,000; \$25,000) and [\$25,000; \$50,000), the positive distance is not significantly different than zero. Thus, with 95% confidence we conclude that males with earnings in such categories are significantly more likely to engage in AHB than females with similar earnings. The problem with bounds estimated according to (10) is that they impose the assumption of no anchoring. Column 4 estimates bounds in (19) that allow for anchoring bias as implied by the experiment in Jacowitz and Kahneman (1995). Relaxing the no anchoring assumption implies that bounds in (10) are wider than those in (19); thus, relaxing this assumption comes at the expense of reducing the power of the test (between genders). In fact, the qualitative conclusion with respect to which of the two genders is more likely to engage in AHB by earning's profile would be identical to the conclusion we obtain from estimates of bounds (5') in Column 2: irrespective of the level in annual earnings, we cannot reject the null of equality in the probability of AHB between males and females. If anything, and within

²⁵ The test consists on estimating if the *distance* between the two overlapping bounds (one will be an upper, the other a lower bound) is sufficiently large and significant to suggest that the null of possible equal probabilities of AHB cannot be rejected. Thus, the test is based on $t = (U - L) / \sqrt{s.e(L)^2 + s.e(U)^2}$, where U (L) is the point estimate on the upper (lower) 95% confidence band on the upper (lower) bound, and $s.e(U)^2$ ($s.e(L)^2$) stands as the square of the standard error on the estimated points, respectively. The test is always a 1 tail test because when bounded regions overlap, by definition the distance, between the Upper bound (of the lower region) and the Lower bound (of the higher region) is always positive. If the distance is negative this is because the identified regions do not overlap.

gender, the regions of identification using bounds in (5') seem to be a lower shift compared to the regions of identification using bounds in (19), and for some categories (e.g., [\$5,000; \$25,000]) the identified regions are almost identical. This latter situation implies that the minimization and maximization of lower and upper bounds in (19) implies that information from partial respondents does not help to tighten the bounds as compared to the situation when bounds ignore partial information (i.e., bounds in (5')).

Overall, it seems that allowing for different bounds that apply either weak data assumption or no assumptions at all, the comparison between genders on the probability of AHB by earning's categories implies that the null of no difference by genders cannot be rejected throughout the distribution of earnings. This is in contrast with a similar comparative exercise allowing for imputation of missing earnings: Column 1 shows that once missing values are imputed the exact classification of non-respondents and partial respondents among the 5 earnings categories changes the conclusions considerably: the results implies that for individuals with earnings in the regions [\$0, \$5,000), [\$50,000; \$100,000) and [\$100,000; \$max) the probability of AHB is significantly higher for males than for females. For the two middle earning's regions the difference in probability between genders is not significant. Thus, allowing for imputed values we detect gender differences in AHB with a u-shape explaining the positive different (in favour of males) with regards to the probability of AHB in the support of earnings: Although a check against bounds (5') implies that the imputation procedure cannot be rejected, our results cannot reject the potential of imputation bias. For example, it is noticeable that imputation of partial respondents employs unfolding bracket response information that is likely to be subject to anchoring effects (as documented in Hurd et al. (1998), for example). This bias would filter through when our inferences are based in Column 1 because imputation is what allows allocation of non-respondents and partial respondents among the $A_{\tau, \kappa}$ earning's categories, thus eventually affecting sample estimates of the probability $P(H = 1 | A \in A_{\tau, \kappa})$.

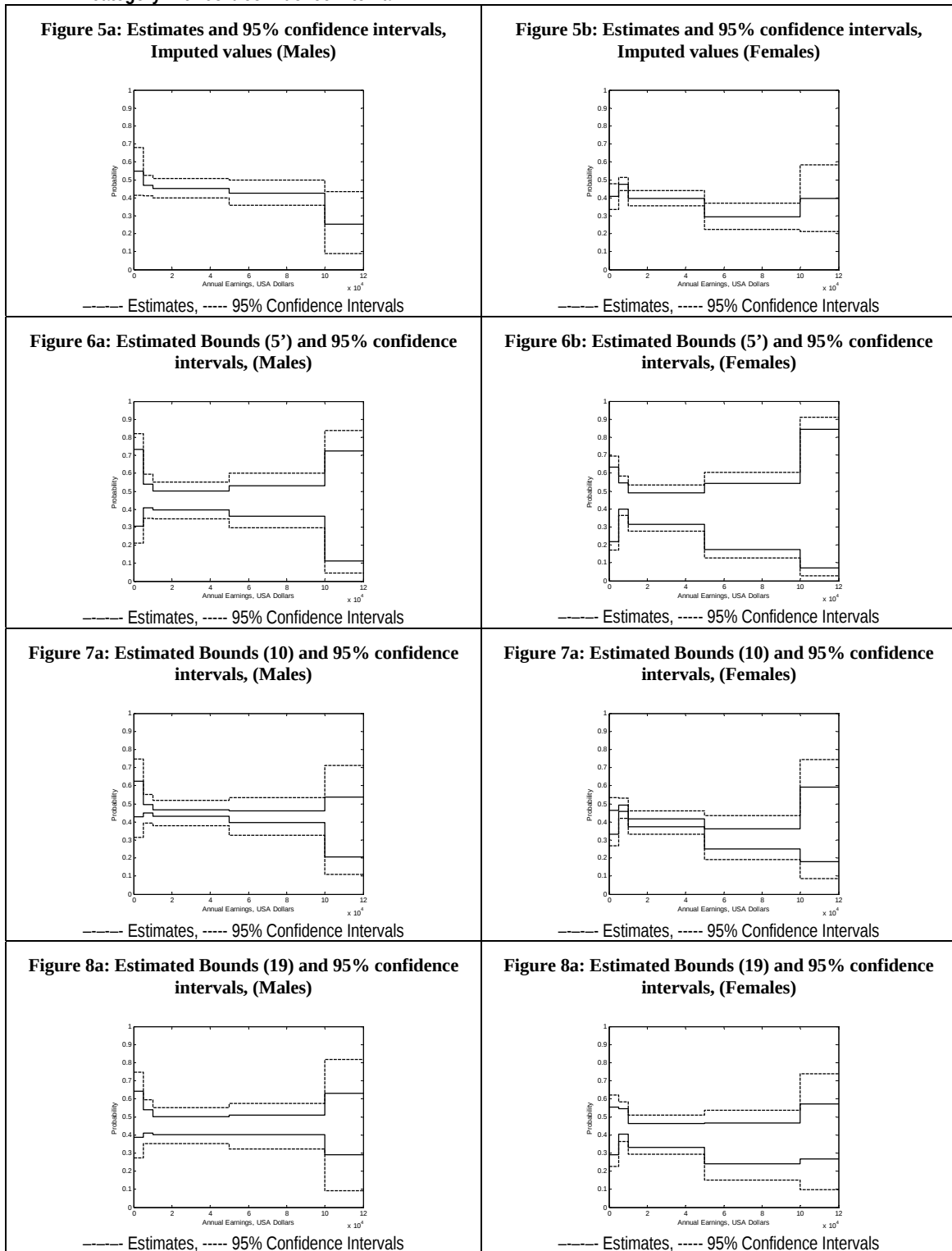
To complement Table 8, estimates based on imputed values and bounds (5'), (10) and (19) are plotted in Figures 5a/b to 8a/b. This figures show both point estimates of the bounds and corresponding 95% confidence intervals. The figures help to visually understand the effect of estimating alternative bounding intervals. For example, comparing Figures 7b and 8b (females) clearly shows the loss in power when tests between genders use bounds in (19) as opposed to bounds in (10).

Table 8: Probability of Adverse Health Habit (H=1) comparing males (n=1,032) against females (n=1,817). Estimates using imputation on missing values and bounds that allow for non-random nonresponse (standard errors in brackets)

		Column 1 $P(H = 1 A \in A_{\tau,K})$ with Imputation	Column 2 $P(H = 1 A \in A_{\tau,K})$ with bounds (5'')	Column 3 $P(H = 1 A \in A_{\tau,K})$ with bounds (10)	Column 4 $P(H = 1 A \in A_{\tau,K})$ with bounds (19)
$A \in [0, \$5000)$	Male (n=64)	0.547 (0.062)	<u>L=0.211</u> <u>U=0.821</u> (0.051) (0.071)	<u>L=0.313</u> <u>U=0.746</u> (0.060) (0.054)	<u>L=0.273</u> <u>U=0.748</u> (0.056) (0.054)
	Female (n=197)	0.407 (0.035)	L=0.171 <u>U=0.695</u> (0.027) (0.033)	L=0.268 <u>U=0.535</u> (0.032) (0.036)	L=0.228 <u>U=0.622</u> (0.034) (0.035)
	Magnitude of Distance / gap. (p-values)	$\Pr(M) > \Pr(F)$ 0.140 (0.024)	Bounds (M) above Bounds (F); lower overlap 0.484 (0.000)	Bounds (M) above Bounds (F); lower overlap 0.222 (0.003)	Bounds (M) above Bounds (F); lower overlap 0.349 (0.000)
$A \in [\$5000; \$25,000)$	Male (n=363)	0.469 (0.026)	<u>L=0.350</u> <u>U=0.596</u> (0.025) (0.026)	L=0.394 <u>U=0.551</u> (0.025) (0.026)	L=0.351 <u>U=0.596</u> (0.025) (0.026)
	Female (n=827)	0.476 (0.017)	L=0.363 <u>U=0.583</u> (0.017) (0.017)	<u>L=0.420</u> <u>U=0.531</u> (0.017) (0.017)	<u>L=0.365</u> <u>U=0.583</u> (0.017) (0.017)
	Magnitude of Distance / gap. (p-values)	$\Pr(F) > \Pr(M)$ 0.007 (0.419)	Bounds (M) above Bounds (F); lower overlap 0.233 (0.000)	Bounds from males nests bounds from females	Bounds from males nests bounds from females
$A \in [\$25,000; \$50,000)$	Male (n=382)	0.453 (0.015)	<u>L=0.346</u> <u>U=0.551</u> (0.025) (0.025)	<u>L=0.378</u> <u>U=0.518</u> (0.025) (0.026)	<u>L=0.353</u> <u>U=0.551</u> (0.024) (0.025)
	Female (n=575)	0.396 (0.020)	L=0.277 <u>U=0.533</u> (0.019) (0.021)	L=0.331 <u>U=0.461</u> (0.020) (0.021)	L=0.294 <u>U=0.511</u> (0.019) (0.021)
	Magnitude of Distance / gap. (p-values)	$\Pr(M) > \Pr(F)$ 0.057 (0.011)	Bounds (M) above Bounds (F); lower overlap 0.187 (0.000)	Bounds (M) above Bounds (F); lower overlap 0.083 (0.006)	Bounds (M) above Bounds (F); lower overlap 0.158 (0.000)
$A \in [\$50,000; \$100,000)$	Male (n=200)	0.426 (0.015)	<u>L=0.296</u> <u>U=0.600</u> (0.032) (0.035)	<u>L=0.327</u> <u>U=0.534</u> (0.033) (0.035)	<u>L=0.323</u> <u>U=0.575</u> (0.033) (0.035)
	Female (n=188)	0.295 (0.033)	L=0.127 <u>U=0.603</u> (0.024) (0.036)	L=0.191 <u>U=0.433</u> (0.029) (0.036)	L=0.152 <u>U=0.536</u> (0.020) (0.036)
	Magnitude of Distance / gap. (p-values)	$\Pr(M) > \Pr(F)$ 0.131 (0.000)	Bounds from females nests bounds from males	Bounds (M) above Bounds (F); lower overlap 0.106 (0.015)	Bounds (M) above Bounds (F); lower overlap 0.213 (0.000)
$A \in [\$100,000; \max(X)]$	Male (n=23)	0.252 (0.014)	<u>L=0.045</u> <u>U=0.839</u> (0.043) (0.077)	<u>L=0.109</u> <u>U=0.713</u> (0.065) (0.094)	L=0.091 <u>U=0.817</u> (0.060) (0.081)
	Female (n=30)	0.396 (0.089)	L=0.031 <u>U=0.911</u> (0.032) (0.052)	L=0.085 <u>U=0.743</u> (0.051) (0.080)	<u>L=0.099</u> <u>U=0.739</u> (0.054) (0.080)
	Magnitude of Distance / gap. (p-values)	$\Pr(F) > \Pr(M)$ 0.144 (0.055)	Bounds from females nests bounds from males	Bounds from females nests bounds from males	Bounds from males nests bounds from females

See Notes Table 1. Numbers in sub-groups by gender and earning's categories based on the imputed values. Numbers in brackets are standard errors. Magnitudes of difference between gender probabilities are presented together with corresponding p-values allowing for a two tail's rejection region.

Table 9: Figures 5a/5b to 8a/8b: Estimates and Bounds on the Probability of Adverse Health Behaviour by earnings category with 95% confidence interval



See Notes Table 6

5 Conclusions

Surveys are often design so that initial non-respondents to questions relating to some quantity A (e.g., earnings) are lead to secondary variables that may elicit partial information on the missing value, for example, asking individuals to complete an unfolding bracket design. Examples of such data collection strategies are found in well known publicly released data sets such as the Health and Retirement Study (HRS) or its counterpart in 15 European countries (The SHARE panel). The design consists on routing individuals through a sequence of questions each of which helps to sort individual's undisclosed amount inside a particular region in the support of A . This is done by prompting partial respondents to react to a bid B_j and reveal if the otherwise undisclosed amount is greater than B_j : the prompting can happen as many as J times along the design. But such effort to reduce the information loss due to nonresponse is unlikely to completely eliminate the problem of nonresponse. Thus, even in the presence of secondary variables the sampling process fails to identify the population parameters of interest. It is possible to solve the identification problem by means of assumptions on individual's response behaviour (e.g., random nonresponse or parametric imputation procedures as those defined in Rubin, 1987) at the expense of inducing bias on the estimated parameters in the direction of the assumptions. An alternative approach is to estimate a bounding interval on the parameter of interest where the width between the upper and the lower bound reflects uncertainty due to nonresponse (for an introduction to the literature see Manski, 1989, 1994 and reference therein).

Bounding intervals are derived free from specification assumptions and are often tighten with very weak data-driven assumptions. The problem arises if we aim at deriving bounds using partial information collected with unfolding bracket designs. It is well known that eliciting information with the use of prompts (e.g., bids) can bias individual's responses in the direction of the prompt (See Section 1). Thus, bounds that allow for information obtained with an unfolding bracket design have to be derived allowing for potential anchoring effects.

This paper draws from Horowitz and Manski (1995,1998) to derive bounds on a conditional outcome $E[g(y) | A]$ where the response behaviour to the information set A defines three sub-groups: full respondents, full non-respondents and partial respondents that provide information with an unfolding bracket design. It is possible for such design to induce anchoring bias, thus, one of the bounding intervals derived in this paper models the effect of anchoring with assumptions that follow the experimental findings in Jacowitz and Kahneman (1995). All bounds are generalized to cover J possible bids in the unfolding bracket design.

The empirical illustration uses data from the second wave of the Health and Retirement Study (HRS, 1996) to estimate bounds on the probability that individuals are associated with at least one of three adverse health habits indicators (smoking, heavy alcohol drinking and obesity) where $AHB=1$ identifies

such an adverse event.. We want to compare the probability of AHB between males and females that share similar regions in the distribution of earnings. The problem is that there is an initial 12% nonresponse rate. Thus, the sampling process fails to identify the probability of AHB for each of the genders by regions of earnings because 12% of the sample cannot be classified with respect to earnings. An unfolding bracket design applied to initial non-respondents reduces full nonresponse rate to 3.2% but may induce anchoring bias on earnings information provided by partial respondents. We estimate three distinct sets of bounds on the probability of AHB conditional on earnings and gender. One of the sets ignores partial respondent's information, the other accounts for such information but assumes that answers from partial respondents are not affected by anchoring, while the last set of bounds incorporates anchoring effects based on weak data-driven assumptions.

Whichever of the three sets of bounds we use, these cannot – in general – detect significant difference in the probability of AHB between males and females for any of the earning's categories. These results contrast with those obtained when missing earnings are imputed; in this case the probability of AHB is significantly higher for males than for females in at least three categories of earnings. The problem we have when making inference using imputed values is that such imputation may induce bias in the direction of the assumptions that underlying the procedure and this may filter through the final probability estimates on AHB. Bounding the parameter of interest, on the other hand, implies using a tool that is easily implementable, while inference using bounds are unquestionably free from all kinds of specification bias.

References

- Borgers, N and J. Hox, (2001), Item nonresponse in questionnaire research with children, *Journal of Official Statistics*, 17,2, 321-335
- Cameron, T. and J. Quiggin, (1994), Estimation using contingent valuation data from a dichotomous choice with follow-up questionnaire, *Journal of Environmental Economics and Management*, 27, 218-234
- Ginther, D., (2000), Alternative estimates of the effect of schooling on earnings, *Review of Economics and Statistics*, 82,1, 102 - 116
- 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
- Hädler, W., (1990), Applied nonparametric regression, *Econometric Society Monographs*, 19, Cambridge University Press, Cambridge
- Herriges, J. and C. 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. Manski, (1995), Identification and robustness with contaminated and corrupted data, *Econometrica*, 63,2, 281-302
- Horowitz, J. and C. Manski, (1998), Censoring of outcomes and regressors due to survey nonresponse: identification and estimation using weights and imputations, *Journal of Econometrics*, 84, 37-58
- Hurd, M., D. McFadden, H. Chand, 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 Aging*, 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
- Juster, T. and J. Smith (1997), Improving the quality of economic data: Lessons from the HRS and AHEAD, *Journal of the American Statistical Association*, 92, 1268-1278
- Lavine, B., A. Gustafson and D. Velenchik, (1995), More bad news for smokers? The effects of cigarette smoking on labour market outcomes, *NBER WP5270*
- Manski, C., (1989), Anatomy of the selection problem, *Journal of Human Resources*, 24, 343-360
- Manski, C., (1990), Non-parametric bounds on treatment effects, *American Economic Review*, Papers & Proceedings, 80, 319-323
- Manski, C., (1994), The selection problem in 'Advances in Econometrics', C. Sims (ed.), Cambridge University Press, 143-170
- Manski, C., (1995), Identification problems in the social sciences, Harvard University Press
- Manski, C., (1997), Monotone treatment response, *Econometrica*, 68, 1311-1334
- Manski, C., and J. Pepper, (2000), Monotone instrumental variables with an application to the returns to schooling, *Econometrica*, 68, 997-1010
- Mason, R., V. Lesser, and M. Traugott, Effect of item nonresponse on nonresponse error and inference, In Groves, R., D. Dillman, J. Eltinge and R. Little (eds.), *Survey Nonresponse*, Wiley, New York, 2002.
- O'Connor, R. M. Johansson and P. Johansson (1999), Stated preferences, real behaviour and anchoring: some empirical evidence, *Environmental and Resource Economics*, 13, 235-248.
- Pepper, J. (2000), The inter-generational transmission of welfare receipts: a non-parametric bounds analysis, *The Review of Economics and Statistics*, 82, 472-488
- Pepper, J. (2003), Using experiments to evaluate performance standards, what do welfare to work demonstrations reveal to welfare reformers?, *Journal of Human Resources*, 38,4, 860-880
- Rabin, M. (1998), Psychology and Economics, *Journal of Economic Literature*, 11-46.
- Switzer, F. and J. Snizek, (1991), Judgement processes in motivation: anchoring and judgement effects on judgement and behaviour, *Organizational Behaviour and Human Decision Process*, Vol. 49, 208-229
- Tourangeau, R., L. Rips and K. Rasinski (2000), The psychology of survey response, New York and Cambridge University Press
- Tversky, A. and D. Kahneman, (1974), Judgement under uncertainty: Heuristics and biases, *Science*, 185, 1124-1131
- Van Ours, J. (2004), A pint a day raises a man's pay, but smoking blows that gain away, *Journal of Health Economics*, 23, 863-886
- Vazquez-Alvarez, R., B. Melenberg and A. vanSoest (2001), Nonparametric bounds in the presence of item nonresponse, unfolding brackets, and Anchoring, *CentER discussion Paper* 01-67, (submitted, 2006)
- Wilson, T. D., C. Houston, K. Etling, and N. Brekke (1996), A new look at anchoring effects: basic anchoring and its antecedents, *Journal of Experimental Psychology*,. 135, 387-402
- Zarkin, G., M. French, T. Mroz and J. Bray, (1998), Alcohol use and wages: new results from the National Household Survey on drug abuse, *Journal of Health Economics*, 17, 53-68

Appendix 1: Mixed outcome and covariate nonresponse

Section 2 extends the framework in Horowitz and Manski (1998) to derive bounds for the case of covariates nonresponse given three response sub-samples (FR, PR, NR) and allowing for anchoring effects resulting from partial responses to an unfolding bracket design. In their paper Horowitz and Manski (1998) deal with two further situations: (i) bounds given that both outcome and *all* covariates suffer jointly from nonresponse (i.e., the case of unit nonresponse as opposed to item nonresponse) and (ii) bounds given that the outcome and ‘*some*’ of the covariates are subject to nonresponse.

(i) Unit Nonresponse

Unit nonresponse implies that both outcome Y and covariate X are missing for a non-negligent percentage of the sample. This prevents identification of $E[g(y) | A \in A_{\tau,k}]$, $A \in X$. To show this, let R indicate full response to either outcome or covariate and let UNR indicate unit nonresponse so that $P(R) + P(UNR) = 1$. Unit nonresponse indicates individuals who dismiss the survey altogether. Then, $E[g(y) | A \in A_{\tau,k}] = E[g(y) | A \in A_{\tau,k}, R]P(R | A \in A_{\tau,k}) + E[g(y) | A \in A_{\tau,k}, UNR]P(UNR | A \in A_{\tau,k})$, with the identification problem arising because neither $E[g(y) | A \in A_{\tau,k}, UNR]$ or $P(UNR | A \in A_{\tau,k})$ are identified by the data. But by definition unit non-respondents completely dismiss the survey and will not be put through an unfolding bracket design (i.e., $P(PR) = 0$). In this situation (1) applies to $E[g(y) | A \in A_{\tau,k}]$ as opposed to (6) and anchoring effects are not an issue. Thus, bounding intervals on $E[g(y) | A \in A_{\tau,k}]$ are identical to those derived with expression (14), Section 3 in Horowitz and Manski (1998).

(ii) Combining outcome and covariate item nonresponse

Here it is assumed that although individuals might fail at providing full information on both outcome Y and the covariate X , they are still participants in the survey (as opposed to completely dismissing the survey at the point of selection). For example, individuals might not answer *in full* to the variable earnings, and might not even answer to variables indicating health habits. Still, these individuals are participants and might provide partial information that cannot be dismissed. With this, Table A.1 explains all possible combinations between the different response behaviour with regards to both the outcome and the covariates:

Table A.1: Types of response behaviour with respect to outcome and covariates

Types	Full Response (FR)	Full Nonresponse (NR)	Partial Response (PR)	Derived Bounds?
1	(Y, X)			Fully identified. Bounds not required since the probability of nonresponse is zero
2	(Y, X)	(Y, X)		Unit nonresponse and combined possibilities of item nonresponse. Partial response does not apply. Bounds derived in Horowitz and Manski (1998)
3	(Y, X)	(Y)	(Y)	Item nonresponse affects the outcome variable and partial response is introduced. Bounds derived in Vazquez-Alvarez et al (2001)
4	(Y, X)	(X)	(X)	Item nonresponse affects the covariate. Bounds derived in this paper, Section 2
5	(Y, X)	(Y, X)	(Y, X)	<i>Item nonresponse might affect outcomes and covariates simultaneously, but it is not classified as 'unit nonresponse' because individuals are still participants in the survey for all other variables, including the possibility to provide partial information to both (if applicable) X and Y .</i>

Table A.1 shows all types of response behaviour always allowing for some meaningful percentage of the sample to response fully to both Y and X . Notice that even if $P(FR) = 0$ bounds would apply equally but to the full sample. The response combination described in ‘type 5’ is the only possibility thus far not accounted for when partial respondents matter and might induce anchoring bias. Deriving bounds when survey responses emulate ‘type 5’ implies accounting for 4 sub-samples of respondents: sub-sample $R1$ that responds in full to both (Y, X) , sub-sample $R2$ that responds to (Y) but are either non-respondents or partial respondents to (X) , sub-sample $R3$ that responds to (X) but are either non-respondents or partial respondents to (Y) , and sub-sample $R4$ that are either non-respondents or partial respondents to (Y, X) . With this, the following partition applies for all $A \in X$:

$$E[g(y)|X] = E[g(y)|X, R1] \times P(R1|X) + E[g(y)|X, R2] \times P(R2|X) + E[g(y)|X, R3] \times P(R3|X) + E[g(y)|X, R4] \times P(R4|X) \quad (A1.1)$$

From expression (A1.1) we know that $E[g(y)|X, R1]$ and $P(R1|X)$ are both identified by the sampling process. Bounds on $E[g(y)|X, R2]$ and $P(R2|X) \in (0,1)$ are identical to those derived in Section 2 that would apply uniquely to the event $R2$, either assuming full nonresponse (see expression (4) and allow for $p_2 = P(R2|X) \in (0,1)$ to vary), or allowing for partial respondents without anchoring effects (see, expressions (9) and define $(g_{\tau, \kappa}, p_{\tau, \kappa})$), or allowing for anchoring effects so that $(g_{\tau, \kappa}, p_{\tau, \kappa})$ are allowed to vary over a range of values given by expression (17) and (18).

Next, assume that Y is some continuous variable and our interest is to study the distribution of Y , that is $g(y) = I[Y \leq y]$: bounds on $E[g(y)|X, R3]$ are identical to those derived in Vazquez-Alvarez et al. (2001) while full response on X implies that $P(R3|X)$ is fully identified by the data and there is

no need to minimize or maximize with respect to possible ranging values in $P(R3|X)$. Assume, therefore, that $E[g(y)|X, R3] \in \{g_{r3,0}(y), g_{r3,1}\}$ represent the interpretation of bounds on $E[g(y)|X, R3]$ given by Vazquez-Alvarez et al. (2001) allowing for full nonresponse and partial information on the outcome Y .

Finally, bounds on $E[g(y)|X, R4]$ and $P(R4|X)$ (also assuming $g(y) = I[Y \leq y]$ for continuous Y) would imply accounting for possibly as many as 4 combinations between full and partial nonresponse in either Y or X . The idea here would be to combine the bounds derived in Vazquez-Alvarez et al. (2001) with bounds derived in Section 2 for the full and partial nonresponse problem in the covariates X . For example, let's assume that $R4$ implies that all information for Y and X comes in the form of partial responses subject to anchoring effects. For any value in the support of X , we know that partial information from Y implies a bounding interval so that $E[g(y)|X, R4]P(R4|X) \in \{g_{r4,0}(y, p_4), g_{r4,1}(y, p_4)\}$ is derived from Vazquez-Alvarez (2001); for each possible $y \in Y$ and for each value of $p_4 = P(R4|X) \in (0,1)$ we have an identification region ranging from $g_{r4,0}(y, p_4)$ to $g_{r4,1}(y, p_4)$: think that each point in this region is derived allowing for possible anchoring effects. Take only the lower bound $\{g_{r4,0}(y, p_4)\}$; the fact that partial responses to X are also subject to anchoring effects implies bounding $\{g_{r4,0}(y)\}$ with expression (19) to account for the type of available information on X so $\inf_{p_4, p_{4,\tau,\kappa}} \{g_{r4,0}(y, p_{4,\tau,\kappa}, p_4)\} \leq g_{r4,0}(y) \leq \sup_{p_4, p_{4,\tau,\kappa}} \{g_{r4,0}(y, p_{4,\tau,\kappa}, p_4)\}$. But $\{g_{r4,0}(y, p_4)\}$ is a lower bound such that $E[g(y)|X, R4]P(R4|X) \in \{g_{r4,0}(y, p_4), g_{r4,1}(y, p_4)\}$ applies. This means that the supreme value $\inf_{p, p_{\tau,\kappa}} \{g_{r4,0}(y, p_{\tau,\kappa}, p)\} \leq g_{r4,0}(y) \leq \sup_{p, p_{\tau,\kappa}} \{g_{r4,0}(y, p_{\tau,\kappa}, p)\}$ does not apply. Applying similar arguments to the upper bounding interval given by $g_{r4,1}(y, p_4)$, the following applies: $\inf_{p_4, p_{4,\tau,\kappa}} \{g_{r4,0}(y, p_{4,\tau,\kappa}, p_4)\} \leq E[g(y)|X, R4]P(R4|X) \leq \sup_{p_4, p_{4,\tau,\kappa}} \{g_{r4,1}(y, p_{4,\tau,\kappa}, p_4)\}$.

Combination lower bounds and upper bounds from each of the events $\{R1, R2, R3, R4\}$ implies summing up the probability for each outcome for any given region in X , and this leads to bounds for 'type 5' of nonresponse in Table A.1. An example that would apply empirically to a 'type 5' data situation would be if we want to estimate the distribution of individual's earnings conditional on the distribution of their partner's earnings (i.e., let $Y = \text{earnings}$, we might want to estimate $P(Y \leq y_i | Y \leq y_j)$ for married couples (i, j)). When Y is a discrete variable (e.g., $H=1$ as health indicator) bounds become simpler and similar to those described by 'type 3' but for subsets of response behaviour in Y . However, if the nature of a discrete outcome Y allows for partial information then

bounds on $g(y) = I[y = j]$ (for all j possible discrete values) are identical to those derived for Y continuous.

Appendix 2: Characterizing bounds for the empirical illustration

Characterising Bounds in (5')

For $g(y) = I(y = 1)$, bounds are characterized such that the following applies:

Let $y = 1$ indicate some health habit outcome:

$$\Rightarrow g(y) = I(y = 1):$$

$$\Rightarrow \text{For any } A_{\tau, \kappa} \in X$$

$$E[g(y) | A_{\tau, \kappa}, FR] = P(y = 1 | A_{\tau, \kappa}, FR) \quad (A2.1)$$

$$\text{with } p \in (0, 1),$$

and

$$g_0(p) = \max\left(0, \frac{P(y = 1 | NR) - (1 - p)}{p}\right), \quad g_1(p) = \min\left(1, \frac{P(y = 1 | NR)}{p}\right)$$

Characterising Bounds in (10)

For $g(y) = I(y = 1)$, bounds are characterized such that the following applies:

Let $y = 1$ indicate some health habit outcome:

$$\Rightarrow g(y) = I(y = 1):$$

$$\Rightarrow \text{For any } A_{\tau, \kappa} \in X$$

$$E[g(y) | A_{\tau, \kappa}, FR] = P(y = 1 | A_{\tau, \kappa}, FR) \quad (A2.2)$$

$$g_{\tau, \kappa} = P(y = 1 | A_{\tau, \kappa}, PR)$$

$$\text{with } p \in (0, 1), \text{ and } p_{\tau, \kappa} = P(A_{\tau, \kappa} \in X | PR)$$

and

$$g_0(p) = \max\left(0, \frac{P(y = 1 | NR) - (1 - p)}{p}\right), \quad g_1(p) = \min\left(1, \frac{P(y = 1 | NR)}{p}\right)$$

Characterising Bounds in (19)

For $g(y) = I(y = 1)$, characterizing bounds in (19) thus allowing for the Jacowitz and Kahneman (1995) assumption requires assessing if the bids described in Table 3 act as either high or low anchors. This is ascertained according to how the sample reacts to them. Table 5 (Section 4) determines that

$B1 \rightarrow \text{High Anchor}$, $B21 \rightarrow \text{High Anchor}$, $B311 \rightarrow \text{High Anchor}$ and $B20 \rightarrow \text{Low Anchor}$. With this, the following characterization of (19) applies:

Let $y=1$ indicate some health habit outcome:

$$\Rightarrow g(y) = I(y=1):$$

$\Rightarrow$ For any $A_{\tau,\kappa} \in X$,

$$E[g(y) | A_{\tau,\kappa}, FR] = P(y=1 | A_{\tau,\kappa}, FR)$$

with $p \in (0,1)$

and for $p_{\tau,\kappa} = P(A_{\tau,\kappa} \in X | PR)$, apply (15)–(17) such that

$$A_{\tau,\kappa} : [0, B20) \Rightarrow 0 \leq p_{\tau,\kappa} \leq P(B20=1 | Q1=0, PR)$$

$$A_{\tau,\kappa} : [B20, B1) \Rightarrow 0 \leq p_{\tau,\kappa} \leq 1$$

$$A_{\tau,\kappa} : [B1, B21) \Rightarrow \min[P(Q1=1 | PR), P(Q21=0 | Q1=1, PR)] \leq p_{\tau,\kappa} \leq 1$$

$$A_{\tau,\kappa} : [B21, B311) \Rightarrow \min[P(Q21=0 | Q1=1, PR), P(Q311=0 | Q21=1, PR)] \leq p_{\tau,\kappa} \leq 1$$

$$A_{\tau,\kappa} : [B311, \infty) \Rightarrow P(Q311=0 | Q21=1, PR) \leq p_{\tau,\kappa} \leq 1$$

and

$$g_0(p) = \max\left(0, \frac{P(y=1 | NR) - (1-p)}{p}\right), \quad g_1(p) = \min\left(1, \frac{P(y=1 | NR)}{p}\right) \quad (A23)$$

and

$$g_0(p_{\tau,\kappa}) = \max\left(0, \frac{P(y=1 | PR) - (1-p_{\tau,\kappa})}{p_{\tau,\kappa}}\right), \quad g_1(p_{\tau,\kappa}) = \min\left(1, \frac{P(y=1 | PR)}{p_{\tau,\kappa}}\right)$$