

A sprightly mathematical model in the presence of scrambled responses

M. Younus Bhat*, Tanveer A. Tarry, Saboor A. Mantoo

Department of Mathematical Sciences, Islamic University of Science and Technology, Kashmir, India

(Communicated by Ali Jabbari)

Abstract

The crux of this paper is to develop a new “Partial” randomized response model. Its properties are studied both theoretically as well as empirically. The proposed model is proved to be more efficient than the randomized response models studied by Eichhorn and Hayre [3] and the “Partial” randomized response model.

Keywords: Randomized response sampling, Estimation of proportion, sensitive quantitative variable, Bogus pipeline
2010 MSC: 62D05

1 Introduction

Social desirability response bias (SDB) is a major problem in survey research involving sensitive questions [2]. Warner [13] was the first to suggest an ingenious method to estimate the proportion of a sensitive character like induced abortion, drug used, etc., through a randomization device such as deck of cards, spinners etc. such that respondent's privacy would be protected. Randomized response technique is one of several methods to partially overcome SDB. Other methods involve use of bogus pipeline (BPL) [7] and a SBD scale [1]. A rich growth of literature on randomized response techniques can be found in Fox and Tracy [4], Zaizai et al. [14], Singh and Tarry [11], Singh et al. [9] and Singh and Singh [10]. Gupta and Thornton [6] have presented a comparison of BPL and RRT methods using survey data. They have shown that a “Partial” RRT is at least as effective in circumventing SDB as BPL, while being more friendly and portable. Below we give the “Full” RRT model due to Eichhorn and Hayre [3] and the “Partial” RRT model of Mangat and Singh [8].

2 The “Full” and “Partial” RRT Models

Eichhorn and Hayre [3] proposed a multiplicative model to gather information on quantitative sensitive variables like income, tax evasion, amount of drug used etc. In the “Full” RRT model of [3], each subject provides a scrambled response. This model works as follows. Let X be a sensitive quantitative variable of interest with an unknown mean of μ_x and an unknown variance of σ_x^2 . Let there be a deck of flash cards that follows a probability distribution S , independent of X , with a known mean of $\mu_s (= \theta)$ and a known variance of σ_s^2 . The respondent is asked to draw a card

*Corresponding author

Email addresses: gyounusg@gmail.com (M. Younus Bhat), tanveerstat@gmail.com (Tanveer A. Tarry), sabojimantoo@gmail.com (Saboor A. Mantoo)

from the deck and is requested to report the scrambled response which is the product of the true response and the number on the card, and divided by the mean of the scrambling variable. Therefore, the reported response Y is given by

$$Y = \frac{XS}{\theta}. \quad (2.1)$$

The expected response, therefore, is given by $E(Y) = \mu_x$. This suggests estimating μ_x by $\hat{\mu}_{xF}$, where $\hat{\mu}_{xF} = \bar{Y}$. The variance of $\hat{\mu}_{xF}$ is given by

$$Var(\hat{\mu}_{xF}) = Var(\bar{Y}) = \frac{Y}{n} = \frac{[\sigma_x^2 + C_s^2(\sigma_x^2 + \mu_x^2)]}{n} = \frac{\mu_x^2}{n} [C_x^2(1 + C_s^2)] C_s^2. \quad (2.2)$$

where $C_s^2 = \sigma_s^2/\theta^2$. In the ‘‘Partial’’ RRT model, a predetermined proportion of randomly selected respondents are asked to provide a true response and the rest provide a scrambled response, just as in the ‘‘Full’’ RRT model. Mangat and Singh [8] gave their ‘‘Partial’’ RRT model for the binary response (Yes/No) case, but it can be easily adapted for the quantitative response case also. If T is the proportion of respondents providing a true response, then the reported response is given by

$$Y = \begin{cases} X & \text{with probability } T \\ \frac{XS}{\theta} & \text{with probability } (1 - T). \end{cases}$$

The expected response is given by

$$E(Y) = \mu_x T + \frac{\mu_x \mu_s (1 - T)}{\theta} = \mu_x \text{ since } \mu_s = \theta.$$

This suggests estimating μ_x by $\hat{\mu}_{xP} = \bar{Y}$. Obviously $\hat{\mu}_{xP}$ is an unbiased estimator of μ_x since $\bar{Y}$ is an unbiased estimator of $E(Y)$. The variance of this estimator is given by

$$Var(\hat{\mu}_{xP}) = Var(\bar{Y}) = \frac{Y}{n} = \frac{[\sigma_x^2 + (1 - T)C_s^2(\sigma_x^2 + \mu_x^2)]}{n} = \frac{\mu_x^2}{n} [C_x^2 + (1 - T)C_s^2(1 + C_x^2)]. \quad (2.3)$$

From (2.3), we have

$$\begin{aligned} Var(\hat{\mu}_{xF}) - Var(\hat{\mu}_{xP}) &= \frac{1}{n} [\sigma_x^2 + C_s^2(\sigma_x^2 + \mu_x^2) - \sigma_x^2 - (1 - T)C_s^2(\sigma_x^2 + \mu_x^2)] \\ &= \frac{TC_s^2(\sigma_x^2 + \mu_x^2)}{n} \end{aligned} \quad (2.4)$$

which is always positive thus the variance of the estimator $\hat{\mu}_{xP}$ is smaller than the variance of the estimator of $\hat{\mu}_{xF}$.

In this paper we suggested an improved estimator for the population mean μ_x and its properties are studied. We have compared the proposed estimator with that of Eichhorn and Hayre [3] and the estimator $\hat{\mu}_{xP}$ based on ‘‘Partial’’ RRT model. It is found that the proposed estimator is more efficient than the estimator $\hat{\mu}_{xF}$ and $\hat{\mu}_{xP}$.

3 The Suggested ‘‘Partial’’ RRT Model

We have suggested the following Partial RRT model :

$$Y = \begin{cases} X & \text{with probability } T \\ w \left(\frac{XS}{\theta} \right) & \text{with probability } (1 - T). \end{cases} \quad (3.1)$$

where $w(0 < w < 1)$ is a constant to be determined such that variance of the estimator based on the model (3.1) is minimum. The expected response is given by

$$E(Y) = \mu_x T + \frac{w \mu_x \mu_s (1 - T)}{\theta} = \mu_x [T + w(1 - T)] \text{ since } \mu_s = \theta.$$

Thus an unbiased estimator of μ_x is given by

$$\hat{\mu}_{xPw} = \frac{\bar{Y}}{\{T + w(1 - T)\}}. \quad (3.2)$$

The variance of $\hat{\mu}_{xPw}$ is given by

$$V(\hat{\mu}_{xPw}) = \frac{V(\bar{Y})}{n[T + w(1 - T)]} \quad (3.3)$$

where

$$\begin{aligned} V(Y) &= E(Y^2) - (E(Y))^2 \\ &= \left[T\mu_x^2(1 + C_x^2) + (1 - T)w^2(1 + C_x^2)(1 + C_s^2)(\mu_s^2/\theta^2) - \mu_x^2 \{T + w(1 - T)\}^2 \right] \\ &= \mu_x^2 \left[T(1 + C_x^2) + (1 - T)w^2(1 + C_x^2)(1 + C_s^2) - \{T + w(1 - T)\}^2 \right] \\ &= \mu_x^2 \left[(1 + C_x^2)\{T + (1 - T)w^2(1 + C_s^2)\} - \{T + w(1 - T)\}^2 \right] \end{aligned} \quad (3.4)$$

Putting (3.4) in (3.3) we get the explicit expression of the variance of $\hat{\mu}_{xPw}$ as

$$\begin{aligned} V(\hat{\mu}_{xPw}) &= \frac{\mu_x^2}{n} \left[(1 + C_x^2) \frac{\{T + w^2(1 - T)(1 + C_s^2)\}}{\{T + w(1 - T)\}^2} - 1 \right] \\ &= \frac{\mu_x^2}{n} [C_x^2 + (1 + C_x^2)C_{(T)}^2] \end{aligned} \quad (3.5)$$

where

$$C_{(T)}^2 = \left[\frac{\{T + w^2(1 - T)(1 + C_s^2)\}}{\{T + w(1 - T)\}^2} - 1 \right]. \quad (3.6)$$

From (2.2) and (3.5) we have

$$V(\hat{\mu}_{xF}) - V(\hat{\mu}_{xP}) = \frac{\mu_x^2}{n} (1 + C_x^2)[C_s^2 - C_{(T)}^2]$$

which is possible if

$$\begin{aligned} C_{(T)}^2 &< C_s^2 \\ \text{i.e. if } &\frac{\{T + w^2(1 - T)(1 + C_s^2)\}}{\{T + w(1 - T)\}^2} - 1 < C_s^2 \\ \text{i.e. if } &\frac{\{T + w^2(1 - T)(1 + C_s^2)\}}{\{T + w(1 - T)\}^2} < (1 + C_s^2) \\ \text{i.e. if } &[(1 - T)(1 + C_s^2)(1 + w^2 - 2w) - C_s^2] < 0 \\ \text{i.e. if } &\left\{ 1 - \sqrt{\frac{C_s^2}{1 - T)(1 + C_s^2)}} \right\} < w < \left\{ 1 + \sqrt{\frac{C_s^2}{1 - T)(1 + C_s^2)}} \right\}. \end{aligned} \quad (3.7)$$

Thus, we established the following theorem.

Theorem 3.1. The proposed estimator $\hat{\mu}_{xPw}$ is better than Eichhorn and Hayre (1983) estimator $\hat{\mu}_{xF}$ if

$$\left\{ 1 - \sqrt{\frac{C_s^2}{1 - T)(1 + C_s^2)}} \right\} < w < \left\{ 1 + \sqrt{\frac{C_s^2}{1 - T)(1 + C_s^2)}} \right\}.$$

It is to be noted that if $\left\{1 - \sqrt{\frac{C_s^2}{1-T}(1+C_s^2)}\right\} < 0$ then the proposed estimator $\hat{\mu}_{xPw}$ would be more efficient than the estimator $\hat{\mu}_{xF}$ if

$$0 < w < \left\{1 + \sqrt{\frac{C_s^2}{1-T}(1+C_s^2)}\right\}. \quad (3.8)$$

Now from (2.3) and (3.5) we have

$$V(\hat{\mu}_{xP}) - V(\hat{\mu}_{xPw}) = \frac{\mu_x^2(1+C_x^2)}{n}[(1-T)C_s^2 - C_{(T)}^2]$$

which is possible if $C_{(T)}^2 < (1-T)C_s^2$, i.e. if $\frac{\{T + w^2(1-T)(1+C_s^2)\}}{\{T + w(1-T)\}^2} - 1 < (1-T)C_s^2$, i.e. if $1 - TC_s^2 + w^2\{1 + (2-T)C_s^2\} - 2w\{1 + (1-T)C_s^2\} < 0$, i.e. if

$$\frac{1 - TC_s^2}{\{1 + (2-T)C_s^2\}} < w < 1. \quad (3.9)$$

Thus we established the following theorem.

Theorem 3.2. The proposed estimator $\hat{\mu}_{xPw}$ is more efficient than the estimator $\hat{\mu}_{xF}$ if

$$\frac{1 - TC_s^2}{\{1 + (2-T)C_s^2\}} < w < 1.$$

4 Optimum choice of the scalar 'w'

Differentiating the variance of $\hat{\mu}_{xPw}$ at (3.5) with respect to 'w' we have

$$\frac{\partial V(\hat{\mu}_{xPw})}{\partial w} = \frac{\mu_x^2}{n} \left[0 + (1+C_x^2) \left\{ \frac{(0+2w(1-T)(1+C_s^2))}{\{T+w(1-T)\}^2} - \frac{(2(1-T)(T+w^2(1-T)(1+C_s^2)))}{\{T+w(1-T)\}^3} \right\} \right]. \quad (4.1)$$

Putting $\frac{\partial V(\hat{\mu}_{xPw})}{\partial w} = 0$, we have

$$2w(1-T)(1+C_s^2) - \frac{2(1-T)(T+w^2(1-T)(1+C_s^2))}{T+w(1-T)} = 0 \quad \text{or} \quad w = \frac{1}{(1+C_s^2)}. \quad (4.2)$$

Differentiating (4.1) with respect to w we have

$$\frac{\partial^2 V(\hat{\mu}_{xPw})}{\partial w^2} = \frac{2\mu_x^2(1+C-x^2)T(1-T)}{n[T+w(1-T)]^2} [3-2T+TC_s^2-2w(1-T)(1+C_s^2)] > 0$$

if $[3-2T+TC_s^2-2w(1-T)(1+C_s^2)] > 0$ i.e. if $w < \frac{1}{(1+C_s^2)} \left[1 + \frac{(1+TC_s^2)}{2(1-T)} \right]$.

It is observed from (4.2) that the equation $\frac{\partial V(\hat{\mu}_{xPw})}{\partial w} = 0$ yields $w = (1+C_s^2)^{-1}$ which is always less than upper bound of the inequality (4.3) i.e., $\frac{1}{(1+C_s^2)} \left[1 + \frac{(1+TC_s^2)}{2(1-T)} \right]$. Thus $(1+C_s^2)^{-1}$ is the optimum value of w which will minimize the variance of the proposed estimator $\hat{\mu}_{xPw}$ i.e. the optimum value of w is

$$w = (1+C_s^2)^{-1} = w_{opt} \text{ (say)}. \quad (4.3)$$

Substitution of (4.3) in (3.5) yields the minimum variance of the estimator $\hat{\mu}_{xPw}$ as

$$\min Var(\hat{\mu}_{xPw}) = \frac{\mu_x^2}{n} \frac{[C_x^2 + C_s^2(1-T+C_x^2)]}{(1+TC_s^2)}. \quad (4.4)$$

Substitution of (4.3) in (3.1) yields the optimum Partial RRT model as

$$Y^* = \begin{cases} X & \text{with probability } T \\ w \left(\frac{XS}{\theta(1+C_s^2)} \right) & \text{with probability } (1-T). \end{cases} \quad (4.5)$$

The expected value of Y^* is given by

$$E(Y^*) = TE(X) + (1-T) \frac{1}{\theta(1+C_s^2)} E(XS) = \mu_x \frac{(1+TC_s^2)}{(1+C_s^2)}. \quad (4.6)$$

Thus the unbiased estimator of μ_x is given by

$$\hat{\mu}_{xPo} = \frac{(1+C_s^2)}{(1+TC_s^2)} \bar{Y}^*. \quad (4.7)$$

The variance of the unbiased estimator $\hat{\mu}_{xPo}$ is given by

$$Var(\hat{\mu}_{xPo}) = \frac{(1+C_s^2)^2}{(1+TC_s^2)^2 n} V(Y^*). \quad (4.8)$$

The variance of Y^* is given by

$$\begin{aligned} Var(Y^*) &= E(Y^{*2}) - (E(Y^*))^2 = TE(X^2) + \frac{(1-T)}{\theta^2(1+C_s^2)^2} E(X^2S^2) - \frac{\mu_x^2(1+TC_s^2)^2}{(1+C_s^2)^2} \\ &= \frac{\mu_x^2(1+TC_s^2)^2}{(1+C_s^2)^2} [C_x^2 + C_s^2(1-T+C_x^2)] \end{aligned} \quad (4.9)$$

Putting (4.9) in (4.8), we get the variance of $\hat{\mu}_{xPo}$ as

$$Var(\hat{\mu}_{xPo}) = \frac{\mu_x^2 [C_x^2 + C_s^2(1-T+C_x^2)]}{n(1+TC_s^2)^2}. \quad (4.10)$$

It is observed from (4.4) and (4.10) that

$$\min Var(\hat{\mu}_{xPw}) = Var(\hat{\mu}_{xPo}). \quad (4.11)$$

Thus the estimator $\hat{\mu}_{xPo}$ is an optimum estimator in the class of estimator $\hat{\mu}_{xPw}$.

5 Efficiency Comparison

From (2.2), (2.3) and (4.10) we have

$$Var(\hat{\mu}_{xF}) - Var(\hat{\mu}_{xPo}) = \frac{TC_s^2\mu_x^2(1+C_x^2)(1+C_s^2)}{n(1+TC_s^2)^2} > 0. \quad (5.1)$$

$$Var(\hat{\mu}_{xP}) - Var(\hat{\mu}_{xPo}) = \frac{T(1-T)C_s^4\mu_x^2(1+C_x^2)}{n(1+TC_s^2)^2} > 0. \quad (5.2)$$

It is observed from (5.1) and (5.2) that the proposed optimum estimator $\hat{\mu}_{xPo}$ is more efficient than $\hat{\mu}_{xF}$ (Eichhorn and Hyarein [3] estimator) and $\hat{\mu}_{xP}$ in Gupta and Shabbir [5]. Further from (2.4) and (5.2) we have the inequality:

$$Var(\hat{\mu}_{xPo}) < Var(\hat{\mu}_{xP}) < Var(\hat{\mu}_{xF}) \quad (5.3)$$

The proposed optimum estimator $\hat{\mu}_{xPo}$ is more efficient than $\hat{\mu}_{xF}$ and $\hat{\mu}_{xP}$. It is interesting to note that the proposed optimum estimator $\hat{\mu}_{xPo}$ depends on the information (θ, C_s^2) associated with scrambled variable S which are known in advance. Thus the proposed estimator $\hat{\mu}_{xPo}$ is recommended advantageously for its use in practice.

6 Numerical illustration

Researchers in the domains of statistics, computer science, public policy, and other disciplines are creating new ideas and methodologies. To reduce the danger of exposure, further ideas and techniques have been created, including grouping, data swapping, synthetic data, l-diversity, and differential privacy. In our opinion, response randomisation is one of the most fundamental and effective strategies for preserving privacy and data confidentiality. In particular, we think there is plenty of room to create post-randomization techniques for the data secrecy approach. There has long been discussion about the benefits of randomizing real replies to preserve respondent confidentiality and privacy. However, it has not been widely utilized in actual surveys, presumably because there aren't enough practical privacy protections and good instructions for selecting the transition probabilities. Companies and computer scientists have paid close attention to randomized response approaches recently in a new context, namely for privacy protection when capturing data from diverse online activities. Newer studies have produced exact privacy principles, safeguards, and rigorous procedures for calculating transition probabilities. Several of such advancements have been examined by us. Newer studies have produced exact privacy principles, safeguards, and rigorous procedures for calculating transition probabilities. Several of such advancements have been examined by us. Partial RRT surveys are comparable in that both randomize true responses with predetermined probabilities and the transition probabilities control their mathematical properties. The similarities between RR surveys and Partial RRT, however, lie in the fact that both randomize genuine replies with a specified probability, and the mathematical features of both are governed by the transition probabilities.

Using (3.7) and (3.9), the range of the scalar ' w ' has been computed for various values of $T = 0.1(0.1)0.9$ and $C_s = 0.1(0.1)0.9$ for which the suggested estimator $\hat{\mu}_{xPw}$ is better than Eichhorn and Hayre's [3] estimator $\hat{\mu}_{xF}$ and the estimator $\hat{\mu}_{xP}$ cited in Gupta and Shabbir [5]. The computed value of the ranges w have been compiled in Tables 6.1 and 6.2. We have computed the percent relative efficiencies (PRE's) of $\hat{\mu}_{xPw}$ with respect to $\hat{\mu}_{xF}$ and $\hat{\mu}_{xP}$ for different values of $C_s = 1.00, 1.50, 2.00, 2.50, 3.00, C_x = 1.50, 2.00, 2.50, 3.00, T = 0.10, 0.20, 0.30, 0.40, 0.50$ and $w = 0.42, 0.44, 0.46$ by using the formulae:

$$PRE(\hat{\mu}_{xPw}, \hat{\mu}_{xF}) = \frac{Var(\hat{\mu}_{xF})}{Var(\hat{\mu}_{xPw})} \times 100 = \frac{\{C_x^2 + (1 + C_x^2)C_s^2\}}{\{C_x^2 + (1 + C_x^2)C_{(T)}^2\}}. \quad (6.1)$$

and

$$PRE(\hat{\mu}_{xPw}, \hat{\mu}_{xP}) = \frac{Var(\hat{\mu}_{xP})}{Var(\hat{\mu}_{xPw})} \times 100 = \frac{\{C_x^2 + (1 - T)C_s^2(1 + C_x^2)\}}{\{C_x^2 + (1 + C_x^2)C_{(T)}^2\}}. \quad (6.2)$$

Finding are displayed in Tables 6.3 and 6.4. We have further computed the percent relative efficiencies (PREs) of the proposed optimum estimator $\hat{\mu}_{xPo}$ with respect to $\hat{\mu}_{xF}$ and $\hat{\mu}_{xP}$ by using the following formulae:

$$PRE(\hat{\mu}_{xPo}, \hat{\mu}_{xF}) = \frac{Var(\hat{\mu}_{xF})}{Var(\hat{\mu}_{xPo})} \times 100 = \frac{\{C_x^2 + (1 + C_x^2)(1 + TC_s^2)\}}{\{C_x^2 + (1 - T + C_x^2)C_s^2\}}. \quad (6.3)$$

and

$$PRE(\hat{\mu}_{xPo}, \hat{\mu}_{xP}) = \frac{Var(\hat{\mu}_{xP})}{Var(\hat{\mu}_{xPo})} \times 100 = \frac{\{C_x^2 + (1 - T)C_s^2(1 + C_x^2)\} (1 + TC_s^2)}{\{C_x^2 + (1 - T + C_x^2)C_s^2\}}. \quad (6.4)$$

$C_s = 1.00, 1.50, 2.00, 2.50, 3.00, C_x = 1.50, 2.00, 2.50, 3.00$ and $T = 0.10, 0.20, 0.30, 0.40, 0.50$. Findings are shown in Tables 6.5 and 6.6. We have computed the range of w for different values of T and C_s and findings are made known in Tables 6.1 and 6.2. Table 6.1 and 6.2 put on display that there is enough scope of selecting the value of scalar ' w ' for obtaining the estimators superior than $\hat{\mu}_{xP}$ and $\hat{\mu}_{xF}$ respectively. To justify this we have computed the $PRE(\hat{\mu}_{xPw}, \hat{\mu}_{xP})$ and $PRE(\hat{\mu}_{xPw}, \hat{\mu}_{xF})$ for chosen values of w, C_x, C_s and T and findings are displayed in Tables 6.3 and 6.4. It is observed from Tables 6.3 and 6.4 that the values of $PRE(\hat{\mu}_{xPw}, \hat{\mu}_{xP})$ and $PRE(\hat{\mu}_{xPw}, \hat{\mu}_{xF})$ are greater than 100. It follows that the proposed estimator $\hat{\mu}_{xPw}$ is more efficient than the usual estimators $\hat{\mu}_{xP}$ and the estimator $\hat{\mu}_{xF}$ with substantial gain in efficiency. We have also worked out the percent relative efficiencies of the "Optimum" estimator $\hat{\mu}_{xPo}$ with respect to $\hat{\mu}_{xF}$ and $\hat{\mu}_{xP}$ and the results are shown in Tables 6.5 and 6.6. The findings of the Tables 6.5 and 6.6 undoubtedly show that the proposed optimum estimator $\hat{\mu}_{xPo}$ is more efficient than the estimator $\hat{\mu}_{xF}$ and $\hat{\mu}_{xP}$ with considerable gain in efficiency. Tables 6.5 and 6.6 also exhibit that the values of

$PRE(\hat{\mu}_{xPo}, \hat{\mu}_{xP})$ and $\hat{\mu}_{xF}$ and $\hat{\mu}_{xP}$ remain higher if the values of C_s increases. Thus, based on our theoretical and simulation results, the use of the proposed estimators $\hat{\mu}_{xPw}$ and $\hat{\mu}_{xPo}$ over Eichhorn and Hayre [3] and the estimator $\hat{\mu}_{xP}$ are recommended in practice.

Table 6.1 The range of w in which the proposed estimator $\hat{\mu}_{xPw}$ is better than the estimator $\hat{\mu}_{xF}$

C_s	Range of w				
	T				
1.5	0.10	0.20	0.30	0.40	0.50
	0.12 ~ 1.88	0.07 ~ 1.93	0.01 ~ 1.99	0.00 ~ 2.07	0.00 ~ 2.18
2.0	0.06 ~ 1.94	0.00 ~ 2.00	0.00 ~ 2.07	0.00 ~ 2.15	0.00 ~ 2.25
2.5	0.02 ~ 1.98	0.00 ~ 2.04	0.00 ~ 2.11	0.00 ~ 2.20	0.00 ~ 2.31
3.0	0.00 ~ 2.00	0.00 ~ 2.06	0.00 ~ 2.13	0.00 ~ 2.22	0.00 ~ 2.23

Table 6.2 The range of w in which the proposed estimator $\hat{\mu}_{xPw}$ is better than the estimator $\hat{\mu}_{xP}$

C_s	Range of w				
	T				
1.5	0.10	0.20	0.30	0.40	0.50
	0.15 ~ 1.00	0.11 ~ 1.00	0.07 ~ 1.00	0.02 ~ 1.00	0.00 ~ 1.00
2.0	0.07 ~ 1.00	0.02 ~ 1.00	0.00 ~ 1.00	0.00 ~ 1.00	0.00 ~ 1.00
2.5	0.03 ~ 1.00	0.00 ~ 1.00	0.00 ~ 1.00	0.00 ~ 1.00	0.00 ~ 1.00
3.0	0.01 ~ 1.00	0.00 ~ 1.00	0.00 ~ 1.00	0.00 ~ 1.00	0.00 ~ 1.00

Table 6.3 The PRE ($\hat{\mu}_{xPw}, \hat{\mu}_{xF}$)

C_s	C_o	w	PRE				
			T				
			0.10	0.20	0.30	0.40	0.50
1.00	1.50	0.40	125.90	154.66	186.59	222.08	261.65
		0.40	133.32	172.64	218.82	272.99	336.64
		0.40	137.57	183.84	240.59	310.23	396.11
		0.40	140.16	191.02	255.30	336.84	441.23
		0.40	123.80	149.32	176.59	205.67	236.65
		0.40	131.53	167.65	208.66	254.91	306.88
1.50	2.50	0.40	136.16	179.67	231.57	293.16	366.08
		0.40	139.05	187.63	247.67	321.72	413.25
		0.40	123.77	146.77	171.94	198.27	225.76
		0.40	136.60	165.15	203.69	246.35	293.27
		0.40	135.41	177.51	227.01	284.76	351.78
		0.40	138.46	185.84	243.71	314.07	399.50
2.00	3.00	0.40	122.22	145.41	169.51	194.46	220.23
		0.40	130.10	163.78	201.03	241.83	286.21
		0.40	134.99	176.32	224.52	280.24	344.22
		0.40	138.13	184.84	241.53	309.90	392.11
		0.40	121.89	144.62	168.10	192.27	217.09
		0.40	129.79	162.98	199.47	239.21	282.14
3.00	3.00	0.40	134.74	175.61	223.05	277.59	339.82
		0.40	137.92	184.24	240.22	307.42	387.77
		0.44	124.53	152.07	182.97	217.65	256.68
		0.44	130.52	166.80	209.83	260.94	321.85
		0.44	133.89	175.71	227.32	291.28	371.23
		0.44	135.93	181.32	238.84	312.31	407.41
1.00	1.50	0.44	122.57	147.07	173.56	202.11	232.84
		0.44	128.91	162.35	200.79	244.79	295.03
		0.44	132.65	172.10	219.57	276.62	345.29
		0.44	134.97	178.44	232.44	299.68	383.99
		0.44	121.60	144.66	169.16	195.08	222.40
		0.44	128.08	160.11	196.35	237.09	282.66
2.00	2.50	0.44	131.99	170.23	215.63	269.36	332.82
		0.44	134.45	176.92	229.10	293.25	372.37
		0.44	121.08	143.39	166.86	191.44	217.10
		0.44	127.62	158.89	193.97	233.01	276.22
		0.44	131.62	169.19	213.48	265.43	326.19
		0.44	134.16	176.06	227.26	289.72	366.09
3.00	3.00	0.44	120.77	142.64	165.53	189.36	214.09
		0.44	127.35	158.17	192.57	230.64	272.50
		0.44	131.40	168.57	212.20	263.12	322.32
		0.44	133.98	175.55	226.15	287.63	362.40
		0.46	123.81	150.67	180.96	215.14	253.81
		0.46	129.21	164.03	205.51	255.04	314.48
1.00	2.50	0.46	132.24	172.02	221.23	282.48	359.49
		0.46	134.05	177.01	231.48	301.24	391.97
		0.46	121.91	145.84	171.86	200.08	230.62
		0.46	127.68	159.84	196.99	239.80	289.07
		0.46	131.07	168.65	214.04	268.89	335.38
		0.46	133.15	174.35	225.60	289.68	370.51
2.00	3.00	0.46	120.97	143.52	167.61	193.25	220.45
		0.46	126.90	157.72	192.79	232.51	277.31
		0.46	130.45	166.90	210.38	262.12	323.74
		0.46	132.67	172.93	222.53	283.76	359.81
		0.46	120.47	142.28	165.38	189.72	215.28
		0.46	126.46	156.56	190.54	228.64	271.17
3.00	3.00	0.46	130.10	165.94	208.37	258.47	317.54
		0.46	132.40	172.14	220.82	280.52	354.03
		0.46	120.17	141.56	164.09	187.69	212.33
		0.46	126.20	155.88	189.21	226.39	267.62
		0.46	129.89	165.36	207.18	256.31	313.92
		0.46	132.23	171.67	219.81	278.59	350.61

Table 6.4 The PRE ($\hat{\mu}_{xPw}, \hat{\mu}_{xP}$)

C_x	C_i	w	PRE				
			T				
			0.10	0.20	0.30	0.40	0.50
1.00	1.50	0.40	115.69	129.35	140.79	149.40	154.61
		0.40	121.47	141.95	160.47	175.92	187.02
		0.40	124.84	149.89	173.76	195.33	212.72
		0.40	126.88	154.82	182.74	209.20	232.23
		0.40	123.99	147.32	169.83	187.59	201.29
1.50	2.00	0.40	114.34	126.49	136.88	142.76	146.16
		0.40	120.32	139.07	155.30	167.99	176.08
		0.40	123.99	147.32	169.83	187.59	201.29
		0.40	126.14	152.78	178.47	202.23	221.39
		0.40	113.71	125.11	133.89	139.77	142.48
2.00	2.50	0.40	119.72	137.62	152.77	164.23	171.07
		0.40	123.41	146.04	166.44	183.78	195.85
		0.40	125.74	151.71	176.57	198.70	216.06
		0.40	113.38	124.39	132.74	138.22	140.62
		0.40	119.39	136.84	151.42	162.25	168.48
2.50	3.00	0.40	125.13	145.33	165.33	181.73	192.97
		0.40	125.52	151.19	175.40	196.77	213.19
		0.40	113.18	123.96	132.88	137.34	139.56
		0.40	119.20	136.37	150.62	161.10	166.98
		0.40	122.96	144.91	164.56	180.53	191.29
3.00	3.50	0.40	125.39	150.74	174.71	195.63	211.51
		0.44	114.34	127.19	138.86	146.42	151.67
		0.44	118.92	137.15	153.88	168.16	178.81
		0.44	121.59	143.17	164.17	183.40	199.36
		0.44	123.05	146.97	170.86	193.96	214.43
1.00	1.50	0.44	113.19	123.58	133.74	140.29	143.81
		0.44	117.92	134.67	149.44	161.32	169.28
		0.44	120.71	141.11	160.27	177.01	189.86
		0.44	122.43	145.30	167.69	188.37	205.71
		0.44	112.63	123.32	131.72	137.51	140.37
1.50	2.00	0.44	117.41	133.43	147.27	158.06	164.89
		0.44	120.29	140.05	158.28	173.84	185.29
		0.44	122.10	144.42	165.88	185.52	201.38
		0.44	112.32	122.65	130.67	136.08	138.62
		0.44	117.13	132.75	146.10	156.33	162.60
2.00	2.50	0.44	120.06	139.45	157.20	172.13	182.86
		0.44	121.92	143.93	165.84	183.96	199.05
		0.44	112.14	122.26	130.86	135.26	137.63
		0.44	116.95	132.35	145.41	155.33	161.28
		0.44	119.92	139.10	156.56	171.12	181.45
3.00	3.50	0.44	121.80	143.63	164.47	183.04	197.67
		0.46	113.68	126.01	136.54	144.73	149.98
		0.46	117.73	134.87	150.71	164.36	174.71
		0.46	119.99	140.16	159.78	177.86	193.06
		0.46	121.35	143.47	165.89	187.09	206.30
1.00	1.50	0.46	112.59	123.53	132.44	138.88	142.44
		0.46	116.89	132.58	146.61	158.03	165.86
		0.46	119.27	138.29	156.23	172.06	184.41
		0.46	120.79	141.97	162.75	182.08	198.49
		0.46	112.05	122.34	130.52	136.23	139.14
1.50	2.00	0.46	116.32	131.43	144.60	155.00	161.76
		0.46	118.88	137.31	154.43	169.17	180.24
		0.46	120.49	141.17	161.22	179.52	194.59
		0.46	111.76	121.71	129.51	134.86	137.46
		0.46	116.05	130.89	143.51	153.40	159.62
2.00	2.50	0.46	118.67	136.77	153.44	167.61	178.02
		0.46	120.31	140.72	160.37	178.12	192.49
		0.46	111.59	121.34	128.93	134.07	136.50
		0.46	115.90	130.43	142.87	152.46	158.39
		0.46	118.54	136.45	152.85	166.69	176.72
3.00	3.00	0.46	120.21	140.46	159.86	177.79	191.24

Table 6.5 The PRE ($\hat{\mu}_{xP0}, \hat{\mu}_{xP}$)

C_x	C_i	PRE				
		T				
		0.10	0.20	0.30	0.40	0.50
1.00	1.50	127.73	157.92	190.93	227.17	267.14
		146.51	197.56	253.85	316.22	385.71
		170.39	247.96	333.87	429.55	536.75
		199.45	309.30	431.29	567.53	720.69
		125.45	152.16	180.22	209.74	240.83
1.50	2.00	143.77	189.97	238.79	290.48	345.28
		167.13	238.20	313.56	393.61	478.82
		195.59	296.97	404.69	519.35	641.67
		124.33	149.41	175.26	201.92	229.42
		142.37	186.21	231.58	278.57	327.27
2.00	2.50	165.43	233.27	303.65	376.72	452.63
		193.56	290.68	391.58	496.48	605.62
		123.73	147.95	172.67	197.89	223.65
		141.61	184.18	227.75	272.36	318.05
		164.49	230.59	298.35	367.83	439.11
2.50	3.00	192.42	287.23	384.52	484.39	586.94
		123.38	147.10	171.17	195.59	220.37
		141.15	182.99	225.52	268.78	312.77
		163.93	229.00	295.24	362.68	431.35
		191.74	285.19	380.37	477.36	576.19

Table 6.6 The PRE ($\hat{\mu}_{xP0}, \hat{\mu}_{xP}$)

C_x	C_i	PRE				
		T				
		0.10	0.20	0.30	0.40	0.50
1.00	1.50	113.01	123.17	130.08	133.26	132.14
		128.84	152.68	170.77	182.16	185.71
		149.76	191.84	225.00	247.73	258.13
		175.58	240.23	292.15	329.09	348.28
		110.44	117.78	121.79	122.21	118.75
1.50	2.00	125.45	145.12	158.51	165.05	164.15
		145.67	182.11	208.48	223.83	227.11
		170.81	228.18	270.86	297.42	306.25
		109.17	115.20	117.95	117.25	112.94
		123.73	141.38	152.63	157.14	154.55
2.00	2.50	143.55	177.21	200.42	212.60	213.08
		168.30	222.03	260.37	282.38	287.08
		108.49	113.84	115.94	114.70	110.00
		122.78	139.36	149.52	153.02	149.62
		142.37	174.53	196.11	206.69	205.85
2.50	3.00	166.89	218.67	254.72	274.43	277.15
		108.09	113.04	114.78	113.24	108.33
		122.22	138.17	147.70	150.63	146.81
		141.67	172.95	193.58	203.26	201.69
		166.06	216.67	251.40	269.81	271.43

References

- [1] D.P. Crowne and D. Marlowe, *A new scale of social desirability of psychopathy*, J. Consul. Psycho. **24** (1960), 349–354.
- [2] A.L. Edward, *The Social Desirability Variable in Personality Assessment and Research*, Dryden, New York, 1975.
- [3] B.H. Eichhorn and L.S. Hayre, *Scrambled randomized response methods for obtaining sensitive quantitative data*, J. Statist. Plann. Inf. **7** (1983), 307–316.
- [4] J.A. Fox and P.E. Tracy, *Randomized Response: A method of Sensitive Surveys*, SAGE Publications, Newbury Park, 1986.
- [5] S. Gupta and J. Shabbir, *Sensitivity estimation for personal interview survey questions*, Statistica **64** (2004), no. 4, 643–653.
- [6] S. Gupta and B. Thornton, *Circumventing social desirability response bias in personal interview surveys*, Amer. J. Math. Manag. Sci. **22** (2003), 369–383.
- [7] E.E. Jones and H. Sigall, *he bogus pipe line: A new paradigm for measuring affect and attitude*, Psycho. Bulle **76** (1971), 349–364.
- [8] N.S. Mangat and R. Singh, *An alternative randomized response procedure*, Biometrika **77** (1990), 439–442.
- [9] G.N. Singh, C. Singh, and A. Kummar, *A modified randomized device for estimation of population mean of quantitative sensitive variable with measure of privacy protection*, Comm. Stat. Simul. Comput. **51** (2022) 1867–1894.
- [10] G.N. Singh and C. Singh, *Proficient randomized response model based on blank card strategy to estimate the sensitive parameter under negative binomial distribution*, Ain Shams Eng. J. **13** (2012), no. 5, 101611.
- [11] H.P. Singh and T.A. Tarry, *A Stratified Unknown repeated trials in randomized response sampling.*, Comm. Korean Statist. Soc. **19** (2012), no. 6, 751–759.
- [12] H.P. Singh and T.A. Tarry, *An alternative to Kim and Warde’s mixed randomized response model*, Statist. Oper. Res. Trans. **37** (2013), no. 2, 189–210.
- [13] S.L. Warner, *Randomized response: a survey technique for eliminating evasive answer bias*, J. Ameri. Statist. Assoc. **60** (1965), 63–69.
- [14] Y. Zaizai, W. Jingyu, L. Junfeng, and W. Hua, *Ratio imputation method for handling item-nonresponse in Eichhorn model*, Assist. Statist. Appl. **3** (2008), no. 2, 89–98.