

Efficient pseudo Rao-Blackwellized estimator in ranked sampling

Bardia Panahbehagh[†] and Mahdi Salehi^{‡*}

[†]*Department of Mathematics, Faculty of Mathematics and Computer Science, Kharazmi University, Tehran, Iran*

[‡]*Department of Mathematics and Statistics, University of Neyshabur, Neyshabur, Iran*

Email(s): panahbehagh@khu.ac.ir, salehi.sms@neyshabur.ac.ir

Abstract. Ranked set sampling (RSS) and judgment post-stratification (JPS) are efficient sampling designs that rely on ranking units without requiring full measurement, typically through visual inspection or the use of auxiliary information. In their standard forms, including multi-ranker RSS, only a portion of the available ranking information is utilized, and the corresponding Rao–Blackwellized estimators are generally unavailable in closed form. In this paper, we propose two practical pseudo Rao–Blackwellization approaches that enhance RSS and JPS by incorporating ranking information from all selected sets. The resulting estimators admit closed-form expressions, make more complete use of the available ranking information for each measured unit, and mitigate the impact of inaccurate rankings from individual sets. We further extend the proposed framework to other ranked sampling designs, including median RSS and extreme RSS, and investigate the performance of the resulting estimators under both symmetric and asymmetric populations. Finally, we analyze two real data sets to illustrate the practical applicability of the proposed methods and to corroborate the findings from the simulation studies.

Keywords: Efficiency; Judgment; Median ranked set sampling; Post-stratification; Ranked set sampling; Rao-Blackwellization.

1 Introduction

Ranked set sampling (RSS) and judgment post-stratification (JPS), introduced by McIntyre (1952) and MacEachern et al. (2004), respectively, utilize ranking information to improve estimation accuracy. In these sampling schemes, ranking is typically performed without full measurement of the variable of interest; instead, it relies on visual inspection or expert judgment. Following its introduction for estimating pasture and forage yields, RSS was rigorously studied by Takahasi and Wakimoto (1968), and since then it has attracted considerable attention across various fields. For example, Yanagawa and Shirahata (1976) proposed a random allocation approach in RSS that coincides with JPS under an infinite population framework.

*Corresponding author

Received: 8 June 2026 / Accepted: 6 July 2026

DOI: 10.22067/sm.ps.2026.99259.1061

A substantial body of literature has compared RSS and JPS with simple random sampling (SRS). Notable contributions include [Stokes and Sager \(1988\)](#), [Kaur et al. \(1995\)](#), [Chen et al. \(2004\)](#), [Salehi et al. \(2016\)](#), and [Dey et al. \(2017\)](#). A key finding from these studies is that RSS is generally more efficient than, or at least as efficient as, SRS, depending on the quality of the ranking process. More recent developments have extended RSS to settings with missing data; for instance, [Bhushan and Kumar \(2024\)](#) proposed logarithmic imputation techniques under RSS to address non-response while preserving estimation efficiency.

[MacEachern et al. \(2004\)](#) introduced JPS as a post-stratification method applied to an SRS using auxiliary ranking information. While JPS typically improves upon SRS, it is generally less efficient than RSS for estimating population means. [Frey and Ozturk \(2011\)](#) showed that strata induced by ranking satisfy specific structural constraints not present in conventional stratification and proposed a Rao–Blackwellized estimator, albeit without a closed form. [Ozturk \(2012\)](#) further integrated ranking information from multiple rankers within JPS and RSS frameworks.

Subsequent work has explored theoretical and methodological extensions of JPS. [Ozturk \(2014\)](#) demonstrated that the efficiency of JPS-based quantile estimators lies between that of SRS and RSS in finite samples and approaches that of RSS as the sample size increases. [Ozturk \(2016\)](#) developed JPS designs for finite populations and proposed algorithms to approximate the Rao–Blackwellized mean estimator. [Matthews and Wolfe \(2017\)](#) introduced a unified framework encompassing both RSS and JPS methodologies. More recently, [Ozturk and Bayramoglu Kavlak \(2020\)](#) and [Ozturk and Kravchuk \(2021\)](#) proposed improved inference procedures by incorporating additional ranking information.

Despite its lower efficiency compared to RSS, JPS is often preferred in practice because it retains the structure of SRS, with ranking information used only at the estimation stage. This feature allows practitioners to apply standard inferential tools developed for SRS while still benefiting from auxiliary ranking information. Recent contributions, such as [Ozturk \(2023\)](#), continue to expand the applicability of rank-based methods in diverse experimental settings.

Both RSS and JPS involve multiple sets, each containing several units, raising an important question: how can the full ranking information across all sets be utilized optimally? [Panahbehagh et al. \(2018\)](#) addressed this issue by proposing a cost-efficient sampling strategy based on RSS. This work was extended by [Panahbehagh and Bruggemann \(2021\)](#) and [Panahbehagh et al. \(2021\)](#) to multivariate settings, where more than one unit can be selected from each set. As long as the number of selected units per rank remains smaller than the number of sets, the resulting design outperforms SRS.

In this paper, we propose a method that fully exploits the ranking information from all units across all sets, along with two closed-form estimators. In both RSS and JPS, ranking is typically based on easily observable auxiliary variables. While such variables provide substantial information, conventional approaches utilize only a limited portion of it. For example, [Ozturk \(2012\)](#) uses ranking information from a single set for each unit, whereas incorporating information from multiple sets can potentially yield greater efficiency.

The remainder of the paper is organized as follows. The next section briefly reviews rank-based sampling designs, including RSS and JPS. Section 2 introduces the proposed improving approaches. Section 3 presents the corresponding estimators of the population mean. A simulation study and further numerical explorations are considered in Section 4. Section 5 provides applications to two real datasets, and Section 6 concludes the paper.

1.1 A brief explanation of RSS and JPS

Suppose that the random variable X has probability density function (pdf) f_μ , with $E(X) = \mu$, $Var(X) = \sigma^2$, and $E(|X|^2) < \infty$. Our objective is to estimate the population mean μ by exploiting the ranking information provided by one or more auxiliary covariates, denoted collectively by z .

To derive an RSS of size $n = cM$, an SRS of size M is selected from the population. None of the units are measured at this stage, but they are instead ranked by judgment without actual measurement. In

this regard, visual inspection, expert opinion, or highly correlated concomitant variables may be used for judgment ranking. The unit, which seems to be the smallest one among the M units is measured, and it will be the first observation of the RSS, denoted by $X_{(1)11}$. Another SRS of size M is obtained, and the judgments are ranked similarly. The unit judgment ranked, which sounds to be the second smallest one, is measured as the second member of the RSS, denoted by $X_{(2)21}$. This procedure continues up to the unit having the largest (M^{th}) judgment rank is measured and completes the last observation of the RSS, denoted by $X_{(M)M1}$. Finally, the sample

$$(X_{(1)11}, X_{(2)21}, \dots, X_{(M)M1})$$

is called a one-cycle RSS of size M . Repeating the above process $c - 1$ times more, gives a c -cycle RSS containing $n = cM$ units, denoted by $\{X_{(i)ij}, i = 1, \dots, M, j = 1, \dots, c\}$, where $X_{(i)ij}$ is the i^{th} judgment ranked unit from the i^{th} SRS in the j^{th} cycle.

JPS is another rank-based sampling that was proposed by [MacEachern et al. \(2004\)](#). In this design, an SRS of size n , $X_1, \dots, X_n$, is taken from the population and actually measured. Then, a further SRS of size $M - 1$ is obtained for each X_i . These extra units are used only as auxiliary ones for judgment ranking of X_i . Hence, we will finally have the pairs of (X_i, R_i) , $i = 1, \dots, n$, where $R_i (\in \{1, \dots, M\})$ is the judgment rank of X_i among its corresponding SRS.

2 Pseudo Rao-Blackwellization of JPS and RSS

In this section we propose two methods to improve JPS and RSS by exploiting set information in a manner analogous to Rao-Blackwellization, but that yield closed-form, approximately unbiased estimators. The Improving JPS (IJPS) procedure for sample size n proceeds as follows. First, draw a simple random sample with replacement (SRS) of size n from f and fully measure its units, denoted $X_1, \dots, X_n$. Independently, select n sets, each an SRS of size $M - 1$, without full measurement, and use all sets to provide judgmental ranks for the measured sample units. Thus, for each unit X_k we obtain n ranks—one from each set—and form the pairs $(X_k, \mathbf{R}_k)$ for $k = 1, \dots, n$, where $\mathbf{R}_k = (R_{k1}, R_{k2}, \dots, R_{kn})$ and $1 \leq R_{ki} \leq M$ denotes the judgmental rank of X_k in the i th set (see Table 1).

Table 1: IJPS of size n ranking in M strata.

$\mathbf{R}$	X_1	X_2	$\dots$	X_n
1 st set	R_{11}	R_{21}	$\dots$	R_{n1}
2 nd set	R_{12}	R_{22}	$\dots$	R_{n2}
$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
n^{th} set	R_{1n}	R_{2n}	$\dots$	R_{nn}

For the improving RSS (IRSS) of size n , assume $n = cM$ where c is an integer. We implement a conventional RSS with a c -cycle of size M (see Table 2). Up to this point, we have a conventional RSS in which the final fully measured sample is given by

$$\{X_{(1)11}, X_{(2)21}, \dots, X_{(M)M1}, X_{(1)12}, X_{(2)22}, \dots, X_{(M)M2}, \dots, X_{(1)1c}, X_{(2)2c}, \dots, X_{(M)Mc}\},$$

where $X_{(i)ij}$ denotes the i th order statistic from the i th set in the j th cycle. We therefore have $n = cM$ sets, all of which can be utilized to extract additional information about the fully measured sample units. To this end, we construct a multi-part matrix (MPM), as illustrated in Table 2. The first part of this matrix consists of the data in Table 2 with the column corresponding to the smallest units ($X_{(1)ij}$) removed. The second part is obtained by removing the column corresponding to the second smallest units ($X_{(2)ij}$), and

Table 2: IRSS of size n ranking in M strata. Each row is ordered. Underbraces indicate selected samples for full measurement.

Cycle	Set	Ranks			
		1	2	$\dots$	M
1	1	$\underbrace{X_{(1)11}}$	$X_{(2)11}$	$\dots$	$X_{(M)11}$
	2	$X_{(1)21}$	$\underbrace{X_{(2)21}}$	$\dots$	$X_{(M)21}$
	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
	M	$X_{(1)M1}$	$X_{(2)M1}$	$\dots$	$\underbrace{X_{(M)M1}}$
2	1	$\underbrace{X_{(1)12}}$	$X_{(2)12}$	$\dots$	$X_{(M)12}$
	2	$X_{(1)22}$	$\underbrace{X_{(2)22}}$	$\dots$	$X_{(M)22}$
	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
	M	$X_{(1)M2}$	$X_{(2)M2}$	$\dots$	$\underbrace{X_{(M)M2}}$
	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
c	1	$\underbrace{X_{(1)1c}}$	$X_{(2)1c}$	$\dots$	$X_{(M)1c}$
	2	$X_{(1)2c}$	$\underbrace{X_{(2)2c}}$	$\dots$	$X_{(M)2c}$
	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
	M	$X_{(1)Mc}$	$X_{(2)Mc}$	$\dots$	$\underbrace{X_{(M)Mc}}$

Table 3: Part k of Multi-part matrix of data to indicate ranks of full measured units.

Part k						
$X_{(1)11}$	$X_{(2)11}$	$\dots$	$X_{(k-1)11}$	$X_{(k+1)11}$	$\dots$	$X_{(M)11}$
$X_{(1)21}$	$X_{(2)21}$	$\dots$	$X_{(k-1)21}$	$X_{(k+1)21}$	$\dots$	$X_{(M)21}$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
$X_{(1)M1}$	$X_{(2)M1}$	$\dots$	$X_{(k-1)M1}$	$X_{(k+1)M1}$	$\dots$	$X_{(M)M1}$
$X_{(1)12}$	$X_{(2)12}$	$\dots$	$X_{(k-1)12}$	$X_{(k+1)12}$	$\dots$	$X_{(M)12}$
$X_{(1)22}$	$X_{(2)22}$	$\dots$	$X_{(k-1)22}$	$X_{(k+1)22}$	$\dots$	$X_{(M)22}$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
$X_{(1)M2}$	$X_{(2)M2}$	$\dots$	$X_{(k-1)M2}$	$X_{(k+1)M2}$	$\dots$	$X_{(M)M2}$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
$X_{(1)1c}$	$X_{(2)1c}$	$\dots$	$X_{(k-1)1c}$	$X_{(k+1)1c}$	$\dots$	$X_{(M)1c}$
$X_{(1)2c}$	$X_{(2)2c}$	$\dots$	$X_{(k-1)2c}$	$X_{(k+1)2c}$	$\dots$	$X_{(M)2c}$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
$X_{(1)Mc}$	$X_{(2)Mc}$	$\dots$	$X_{(k-1)Mc}$	$X_{(k+1)Mc}$	$\dots$	$X_{(M)Mc}$

this process continues similarly. See Table 3 for the k th part of the MPM. Then, we use Part 1 to indicate n ranks for each of $X_{(1)1j}; j = 1, 2, \dots, c$, Part 2 to indicate n ranks for each of $X_{(2)2j}; j = 1, 2, \dots, c$, and so on. Therefore, for each fully measured unit, we obtain n ranks, which leads to a table analogous to Table 1, n observations, each with n ranks.

Then, we re-match each of $\{X_{(k)kj}; j = 1, 2, \dots, c\}$, units (which have rank k in their original sets), in all the $n - 1$ sets in Part k to create sets of size M and indicate $n - 1$ additional ranks for each of $\{X_{(k)kj}; j = 1, 2, \dots, c\}$. Therefore for each full measured unit we will have n ranks which leads to a table exactly like Table 1, n observations, each with n ranks. As a simple example consider $M = 3$ and $c = 2$, which leads to Table 4. Now each of $X_{(1)11}$ and $X_{(1)12}$ which will be fully measured, will receive 5 new

Table 4: An example of IRSS with $M = 3$ and $c = 2$.

Cycle	Set	Ranks		
		1	2	3
1	1	$X_{(1)11}$	$X_{(2)11}$	$X_{(3)11}$
	2	$X_{(1)21}$	$X_{(2)21}$	$X_{(3)21}$
	3	$X_{(1)31}$	$X_{(2)31}$	$X_{(3)31}$
2	1	$X_{(1)12}$	$X_{(2)12}$	$X_{(3)12}$
	2	$X_{(1)22}$	$X_{(2)22}$	$X_{(3)22}$
	3	$X_{(1)32}$	$X_{(2)32}$	$X_{(3)32}$

ranks with each of the sets in Part 1 of Table 5. Please note that ranking them in their original sets leads to rank 1 for them. With this procedure, although we expect these units to receive rank 1, but the rank of them will be adjusted using the other sets, leading to an increase in the accuracy of our judgment in estimating the mean parameter.

Table 5: The parts of the three-layer matrix associated with the example of Table 4.

Part 1		Part 2		Part 3	
$X_{(2)11}$	$X_{(3)11}$	$X_{(1)11}$	$X_{(3)11}$	$X_{(1)11}$	$X_{(2)11}$
$X_{(2)21}$	$X_{(3)21}$	$X_{(1)21}$	$X_{(3)21}$	$X_{(1)21}$	$X_{(2)21}$
$X_{(2)31}$	$X_{(3)31}$	$X_{(1)31}$	$X_{(3)31}$	$X_{(1)31}$	$X_{(2)31}$
$X_{(2)12}$	$X_{(3)12}$	$X_{(1)12}$	$X_{(3)12}$	$X_{(1)12}$	$X_{(2)12}$
$X_{(2)22}$	$X_{(3)22}$	$X_{(1)22}$	$X_{(3)22}$	$X_{(1)22}$	$X_{(2)22}$
$X_{(2)32}$	$X_{(3)32}$	$X_{(1)32}$	$X_{(3)32}$	$X_{(1)32}$	$X_{(2)32}$

In usual Rao-Blackwellization method, we would use more information involving all the possible combinations of all the units of the sets which will lead to complexity both in implementation and estimation. But based on the approaches proposed in this section, we will construct closed form of the estimators easily. Then the proposed methods utilize information inside the sets, more than usual approaches and less than Rao-Blackwellization, and that's why we name them pseudo Rao-Blackwellization approaches.

3 Estimators

JPS is more acceptable than RSS (because of working based on SRS) and its efficiency is asymptotically the same as RSS (MacEachern et al. (2004)). We therefore introduce an estimator for JPS in detail, and we will suffice to present only an estimator for RSS in this section.

Let $(X_k, \mathbf{R}_k); k = 1, 2, \dots, n$, be an IJPS of size n drawn from a population with mean μ . Suppose further that the function $I_{R_{ki}}^m$ is an indicator variable which equals 1 if $R_{ki} = m$ and 0 otherwise. Now, the fraction of cases where X_k takes the rank m , denoted by ω_{km} , is obtained as

$$\omega_{km} = \frac{1}{n_m} \sum_{i=1}^n I_{R_{ki}}^m, \quad n_m = \sum_{k=1}^n \sum_{i=1}^n I_{R_{ki}}^m,$$

where n_m counts the total number of observations of IJPS whose ranks equal m . Then, for estimating μ based on IJPS we propose the following estimator

$$\hat{\mu}_{\text{IJPS}} = \frac{1}{M} \sum_{m=1}^M \hat{\mu}_{[m]}, \quad (1)$$

where

$$\hat{\mu}_{[m]} = \begin{cases} \sum_{k=1}^n \omega_{km} X_k, & n_m > 0, \\ 0, & n_m = 0. \end{cases} \quad (2)$$

Conjecture 1. Under the proposed IJPS mechanism, for any two distinct measured units X_k and $X_{k'}$, $k \neq k'$, the corresponding weighted observations satisfy

$$\text{Cov}(\omega_k X_k, \omega_{k'} X_{k'}) \leq 0.$$

Motivated by the fixed-sum property of the weights and by the numerical evidence from our simulation study, we state the following conjecture.

The intuition behind Conjecture 1 is as follows. The measured observations X_k and $X_{k'}$ are independent, while their weights ω_k and $\omega_{k'}$ are negatively related because they are constructed from the same ranking information and satisfy the fixed-sum relation

$$\sum_{k=1}^n \omega_k = M$$

when all rank classes are non-empty. In particular, under exchangeability, this fixed-sum relation implies

$$\text{Cov}(\omega_k, \omega_{k'}) = -\frac{\text{Var}(\omega_k)}{n-1} \leq 0, \quad k \neq k'.$$

Although the weights are rank-based and therefore related to the measured values, extensive simulations did not show any contradiction to Conjecture 1. Across the parent distributions, sample sizes, and set sizes considered in this study, the estimated covariance terms were consistently non-positive.

Lemma 1. The random variable $\hat{\mu}_{\text{IJPS}}$ given by (1) is an asymptotically unbiased estimator of μ . More specifically, under the usual equal-probability rank-class argument,

$$E(\hat{\mu}_{\text{IJPS}}) = \mu - \left(1 - \frac{1}{M}\right)^{n^2} \mu,$$

and hence

$$E(\hat{\mu}_{\text{IIPS}}) - \mu = - \left(1 - \frac{1}{M}\right)^{n^2} \mu \longrightarrow 0 \quad \text{as } n \rightarrow \infty.$$

Moreover, if Conjecture 1 holds, then

$$\hat{V} = \frac{n}{M^2(n-1)} \sum_{k=1}^n \left(\omega_k X_k - \frac{1}{n} \sum_{k=1}^n \omega_k X_k \right)^2 \quad (3)$$

is an empirical estimator of an upper bound for the variance of $\hat{\mu}_{\text{IIPS}}$, where

$$\omega_k = \sum_{m=1}^M \omega_{km}.$$

If Conjecture 1 is not invoked, the same quantity may be viewed as a first-order empirical variance estimator, obtained by estimating the individual variance component and ignoring the covariance component.

Proof. Let $\mu_{[m]}$ denote the mathematical expectation of an observation whose judgmental rank is m . Then, for a non-empty rank class,

$$E(\hat{\mu}_{[m]} \mid n_m > 0) = \mu_{[m]}.$$

Therefore,

$$\begin{aligned} E(\hat{\mu}_{[m]}) &= \mu_{[m]} P(n_m > 0) \\ &= \mu_{[m]} \left(1 - \left(1 - \frac{1}{M}\right)^{n^2} \right). \end{aligned}$$

Consequently,

$$\begin{aligned} E(\hat{\mu}_{\text{IIPS}}) &= \frac{1}{M} \sum_{m=1}^M E(\hat{\mu}_{[m]}) \\ &= \left(1 - \left(1 - \frac{1}{M}\right)^{n^2} \right) \frac{1}{M} \sum_{m=1}^M \mu_{[m]} \\ &= \left(1 - \left(1 - \frac{1}{M}\right)^{n^2} \right) \mu. \end{aligned}$$

Since

$$\left(1 - \frac{1}{M}\right)^{n^2} \longrightarrow 0 \quad \text{as } n \rightarrow \infty,$$

the bias term vanishes, and hence $\hat{\mu}_{\text{IIPS}}$ is asymptotically unbiased.

For the variance, using (1) and (2), we can write

$$\hat{\mu}_{\text{IIPS}} = \frac{1}{M} \sum_{k=1}^n \omega_k X_k.$$

Therefore,

$$\begin{aligned} \text{Var}(\hat{\mu}_{\text{IIPS}}) &= \frac{1}{M^2} \text{Var} \left(\sum_{k=1}^n \omega_k X_k \right) \\ &= \frac{1}{M^2} \left[\sum_{k=1}^n \text{Var}(\omega_k X_k) + 2 \sum_{1 \leq k < k' \leq n} \text{Cov}(\omega_k X_k, \omega_{k'} X_{k'}) \right]. \end{aligned} \quad (4)$$

The first term in (4) is the individual variance component, whereas the second term is the covariance component. Under Conjecture 1, the covariance component is non-positive. Thus,

$$\text{Var}(\hat{\mu}_{\text{IJPS}}) \leq \frac{1}{M^2} \sum_{k=1}^n \text{Var}(\omega_k X_k).$$

By exchangeability,

$$\text{Var}(\hat{\mu}_{\text{IJPS}}) \leq \frac{n}{M^2} \text{Var}(\omega_k X_k).$$

Replacing $\text{Var}(\omega_k X_k)$ with its sample analogue gives

$$\hat{v} = \frac{n}{M^2(n-1)} \sum_{k=1}^n \left(\omega_k X_k - \frac{1}{n} \sum_{k=1}^n \omega_k X_k \right)^2.$$

If Conjecture 1 is not imposed, this expression still estimates the first-order individual variance component in (4). This completes the proof. $\square$

Using almost the same approach proposed for the IJPS, one can obtain an alternative estimator based on an IRSS of size $n = cM$ explained in Table 2. More specifically, if we denote $W_{ijm} = \sum_{l=1}^{Mc} I(R_{ijl} = m)$, where R_{ijl} is the rank of $X_{(i)ij}$ in the l^{th} set, then the mentioned estimator is defined as

$$\hat{\mu}_{\text{IRSS}} = \frac{1}{M} \sum_{m=1}^M \hat{\mu}_{(m)}, \quad (5)$$

where

$$\hat{\mu}_{(m)} = \frac{\sum_{i=1}^M \sum_{j=1}^c X_{(i)ij} W_{ijm}}{\sum_{i=1}^M \sum_{j=1}^c W_{ijm}}.$$

4 Further numerical studies

4.1 Examining various distributions

We conducted a simulation study under various distributions and defined efficiency as the ratio of the variance of the SRS estimator to the mean squared error (MSE) of the corresponding estimator. Specifically, let $\hat{\mu}$ be an estimator of μ obtained from one of the ranked sampling designs described above. We compute the relative efficiency as

$$RE = \frac{\frac{\sigma^2}{n}}{MSE(\hat{\mu})}, \quad (6)$$

where the numerator represents the variance of the sample mean under SRS of size n , and MSE denotes the Monte Carlo mean squared error. We considered a wide range of distributions, namely $N(10, 2)$, $\text{Uniform}(0, 1)$, $\text{Exp}(1)$, $\text{Weibull}(2)$, $\text{Beta}(7, 4)$, $\text{Gamma}(2, 1)$, $\text{Chisq}(2)$, and $\text{Poisson}(8)$, with $n = 12, 18, 24$ and $M = 3, 6$. The results are reported in Table 6. They show consistent efficiency gains for IJPS and IRSS compared with JPS and RSS, respectively, across all scenarios. Moreover, the efficiency gains tend to increase with respect to n and M in most cases, although the effect is more pronounced with respect to M . The improvements are also more substantial for symmetric distributions.

Results are shown in Table 6. They clearly show that there are efficiency gains for IJPS and IRSS over JPS and RSS, respectively, in all cases. Consequently, the results show that the gains in efficiency

increase as M and n increase; however, the increase in efficiency is more evident for M . Finally, the improvements are more significant for symmetric distributions.

To assess $\hat{V}$ in (3), we constructed 90% and 95% confidence intervals based on the tails of the standard normal distribution and the estimator $\hat{\mu}_{\text{JPS}}$ in (1):

$$\begin{aligned} CI_{90} &= \left(\hat{\mu}_{\text{JPS}} - 1.64\sqrt{\hat{V}_{\text{upper}}}, \quad \hat{\mu}_{\text{JPS}} + 1.64\sqrt{\hat{V}_{\text{upper}}} \right) \\ CI_{95} &= \left(\hat{\mu}_{\text{JPS}} - 1.96\sqrt{\hat{V}_{\text{upper}}}, \quad \hat{\mu}_{\text{JPS}} + 1.96\sqrt{\hat{V}_{\text{upper}}} \right) \end{aligned}$$

for $M = 3, 6$ and $n = 24$. The coverage rates of these confidence intervals for different distributions are reported in Table 6, where C_{90} and C_{95} denote the coverage probabilities (CPs) of CI_{90} and CI_{95} , respectively. These CPs were computed as the proportion of intervals containing the population parameter μ over 2000 iterations. In most cases, the CP exceeded the nominal level, which is expected when an upper bound is used for the variance. At the same time, the upper bound has two advantages: it applies to all considered distributions and nearly guarantees the nominal coverage probability. Moreover, larger sample sizes led to more reliable coverage.

4.2 Proposing another improving approach

The approach proposed in Section 3 improves both JPS and RSS designs regardless of whether the parent distribution is symmetric or asymmetric. As stated previously and illustrated in Table 3 for the algorithm that improves RSS, the auxiliary units from the $n - 1$ remaining sets whose judgment ranks coincide with those of a given fully measured unit are eliminated, and then the unbiased efficient estimator (5) is constructed. One may ask what happens if these auxiliary units are not eliminated. In that case we have $1 \leq R_{ijl} \leq M + 1$, and the estimator (5) reduces to

$$\hat{\mu} = \frac{1}{M+1} \sum_{m=1}^{M+1} \hat{\mu}_{(m)}. \quad (7)$$

Similarly, for JPS, if the $n - 1$ remaining sets contain their matching fully measured units, then the R_{ki} 's in Table 1 also lie in the interval $[1, M + 1]$, and estimator (7) with $\hat{\mu}_{[m]}$ instead of $\hat{\mu}_{(m)}$ applies to the JPS design as well. We refer to the method described in Section 3 as the “First Improving Approach” and to the method presented here as the “Second Improving Approach.”

4.3 Examining some more ranked sampling designs

In addition to this question, we are also interested in examining other ranked sampling designs, such as median RSS (MedRSS) and extreme RSS (ERSS). There are two main reasons for considering these variants of RSS: their simplicity and reduced ranking error (particularly in ERSS), as well as their strong performance under symmetric populations.

MedRSS was introduced by Muttalak (1997) as a simplified version of RSS. The following algorithm can be used to obtain a c -cycle MedRSS with set size M . If M is odd, observations are obtained by measuring the judgment medians from each SRS. For even set sizes, suppose that M simple random samples, each of size M , are drawn. Then, the units with the largest judgment ranks are measured from the first half of the samples, while the units with the smallest judgment ranks are measured from the second half. By repeating this procedure c times, the desired sample size $n = cM$ is obtained.

More specifically, the gathered observations are of the form

$$\left\{ X_{\left(\frac{M+1}{2}\right)ij}, \quad i = 1, \dots, M, \quad j = 1, \dots, c \right\},$$

Table 6: The values of the RE, the relative efficiency of the respective design to the SRS design, and the CP under different distributions (the simulation is carried out based on 2000 Monte Carlo replications).

n	M	CP	RE				CP	RE			
		f	JPS	IJPS	RSS	IRSS	f	JPS	IJPS	RSS	IRSS
12	3	$N(10, 2)$	1.51	2.33	1.95	2.35	$Beta(7, 4)$	1.42	2.25	1.97	2.36
	6		1.57	3.77	3.12	4.25		1.54	3.80	3.18	4.38
18	3		1.57	2.36	1.83	2.22		1.72	2.35	2.00	2.44
	6	C_{90}, C_{95}	1.83	4.1	3.08	4.36	C_{90}, C_{95}	1.89	4.21	3.2	4.54
24	3	0.92, 0.94	1.8	2.48	2.10	2.56	0.91, 0.93	1.70	2.55	1.99	2.42
	6	0.97, 0.98	2.22	4.33	3.30	4.79	0.97, 0.98	2.27	4.14	3.08	4.33
12	3	$Uniform(0, 1)$	1.46	2.48	1.95	2.36	$Gamma(2, 1)$	1.41	2.10	1.79	2.11
	6		1.61	4.15	3.39	4.53		1.48	3.26	2.89	4.06
18	3		1.64	2.42	1.98	2.47		1.49	2.09	1.71	2.08
	6	C_{90}, C_{95}	1.88	4.95	3.29	4.41	C_{90}, C_{95}	1.69	3.48	2.90	4.14
24	3	0.97, 0.98	1.84	2.68	2.02	2.55	0.96, 0.98	1.51	2.14	1.87	2.26
	6	0.96, 0.98	2.22	4.58	3.29	4.70	0.98, 0.99	1.98	3.59	2.83	3.93
12	3	$Exp(1)$	1.32	1.82	1.74	2.01	$Chisq(2)$	1.23	1.73	1.48	1.73
	6		1.41	2.91	2.57	3.59		1.22	2.52	2.13	2.92
18	3		1.36	1.89	1.70	2.01		1.56	1.84	1.52	1.77
	6	C_{90}, C_{95}	1.55	3.05	2.38	3.45	C_{90}, C_{95}	1.44	2.87	2.25	3.15
24	3	0.95, 0.97	1.49	1.92	1.62	1.93	0.94, 0.97	1.3	1.85	1.39	1.64
	6	0.98, 0.99	1.78	3.33	2.71	3.73	0.97, 0.99	1.74	2.77	2.41	3.37
12	3	$Weibull(2, 1)$	1.38	2.32	1.88	2.25	$Poisson(8)$	1.32	1.94	1.7	1.97
	6		1.56	3.72	3.03	4.13		1.29	2.86	2.47	3.34
18	3		1.57	2.26	1.90	2.30		1.41	1.87	1.56	1.86
	6	C_{90}, C_{95}	1.78	4.43	3.07	4.25	C_{90}, C_{95}	1.68	3.20	2.57	3.53
24	3	0.96, 0.98	1.76	2.56	1.98	2.44	0.96, 0.97	1.51	1.96	1.63	1.93
	6	0.99, 0.99	2.04	4.02	3.02	4.43	0.97, 0.98	1.81	3.33	2.58	3.59
12	3	$\omega N(100, 10)$	1.44	2.18	1.80	2.15	$\omega N(100, 10)$	1.35	1.90	1.76	2.06
	6	$+(1-\omega)N(150, 10)$,	1.59	3.98	3.47	4.51	$+(1-\omega)N(150, 10)$,	1.43	2.89	3.03	4.17
18	3	$\omega \sim Ber(0.5)$	1.65	2.40	1.87	2.26	$\omega \sim Ber(0.2)$	1.56	2.03	1.84	2.17
	6	C_{90}, C_{95}	2.04	4.42	3.25	4.32	C_{90}, C_{95}	1.68	3.23	2.70	3.78
24	3	0.91, 0.93	1.68	2.44	2.00	2.41	0.92, 0.94	1.56	2.05	1.71	2.08
	6	0.97, 0.98	2.24	4.46	3.02	4.14	0.97, 0.98	1.87	3.28	2.82	3.91

when M is odd, and

$$\left\{ X_{(\frac{M}{2})_{ij}}, i = 1, \dots, \frac{M}{2}, j = 1, \dots, c \right\} \cup \left\{ X_{(\frac{M}{2}+1)_{ij}}, i = \frac{M}{2} + 1, \dots, M, j = 1, \dots, c \right\},$$

when M is even. As mentioned earlier, if the required conditions of running the RSS are provided in a population, it would be logical to utilize it rather than SRS. But, if the set size M is not small, the ranking error would be non-negligible and this implies enhancing the variances of the estimators obtained based on the RSS. But then, inspecting the extreme values is not so difficult even in large sets. Therefore, [Samawi et al. \(1996\)](#) proposed ERSS which inspects only the extreme observations. Actually, in this sampling design, we select M random samples of size M and rank each sample separately with respect to a variable of interest by visual inspection. If the sample size M is even, we select from $\frac{M}{2}$ samples the smallest unit and from the $\frac{M}{2}$ remaining samples the largest unit for full measurement. If the sample size is odd, select from $\frac{M-1}{2}$ samples the smallest unit, from the other $\frac{M-1}{2}$ the largest unit and from one

sample the median of the sample for actual measurement. This cycle may be repeated c times in order to get $n = cM$ units. In summary, for an even M we have

$$\left\{ X_{(1)ij}, i = 1, \dots, \frac{M}{2}, j = 1, \dots, c \right\} \cup \left\{ X_{(M)ij}, i = \frac{M}{2} + 1, \dots, M, j = 1, \dots, c \right\},$$

and for an odd M we have

$$\begin{aligned} & \left\{ X_{(1)ij}, i = 1, \dots, \frac{M-1}{2}, j = 1, \dots, c \right\} \\ & \cup \left\{ X_{(M)ij}, i = \frac{M+1}{2} + 1, \dots, M-1, j = 1, \dots, c \right\} \\ & \cup \left\{ X_{(\frac{M+1}{2})Mj}, j = 1, \dots, c \right\}. \end{aligned}$$

Therefore in the next simulation scenario, we consider the population mean estimators on the basis of four ranked sampling schemes including RSS, JPS, MedRSS and ERSS of size $n = cM$, with $(c, M) = (3, 4), (4, 5), (4, 15), (5, 20), (5, 30), (6, 30)$, under the two approaches mentioned above as well as their existing conventional versions. In order to be closer, the standard normal distribution as well as the Gamma(2, 1) are selected as the parent distributions in order to involve both of symmetric and asymmetric populations. Two criteria are employed for the comparison purpose: the bias and the RE given by (6). Figures 1 and 2 depict the result. The following points are observed:

- The approaches of using full ranking information improve the ordinary methods of ranking. However, the first approach has a better performance than the second one according to the all criteria.
- Sampling design performance is heavily influenced by the parent distribution. For instance, the improved ERSS outperforms its competitors under normal distributions (representing the class of symmetric distributions). However, it becomes one of the least efficient designs, along with MedRSS, when the population follows an asymmetric distribution, such as the gamma distribution.
- Figures 1 and 2 confirm the asymptotic unbiasedness of the estimators in (1) and (5), which are obtained based on the improved JPS and improved JPS, respectively, under the first approach. However, the improved versions of ERSS and MedRSS produce unbiased estimators for μ only for populations with symmetric densities. This property was proved by Samawi et al. (1996) for the ordinary version of ERSS as well.
- The estimators obtained based on the second approach are clearly biased but still more efficient than their conventional counterparts.

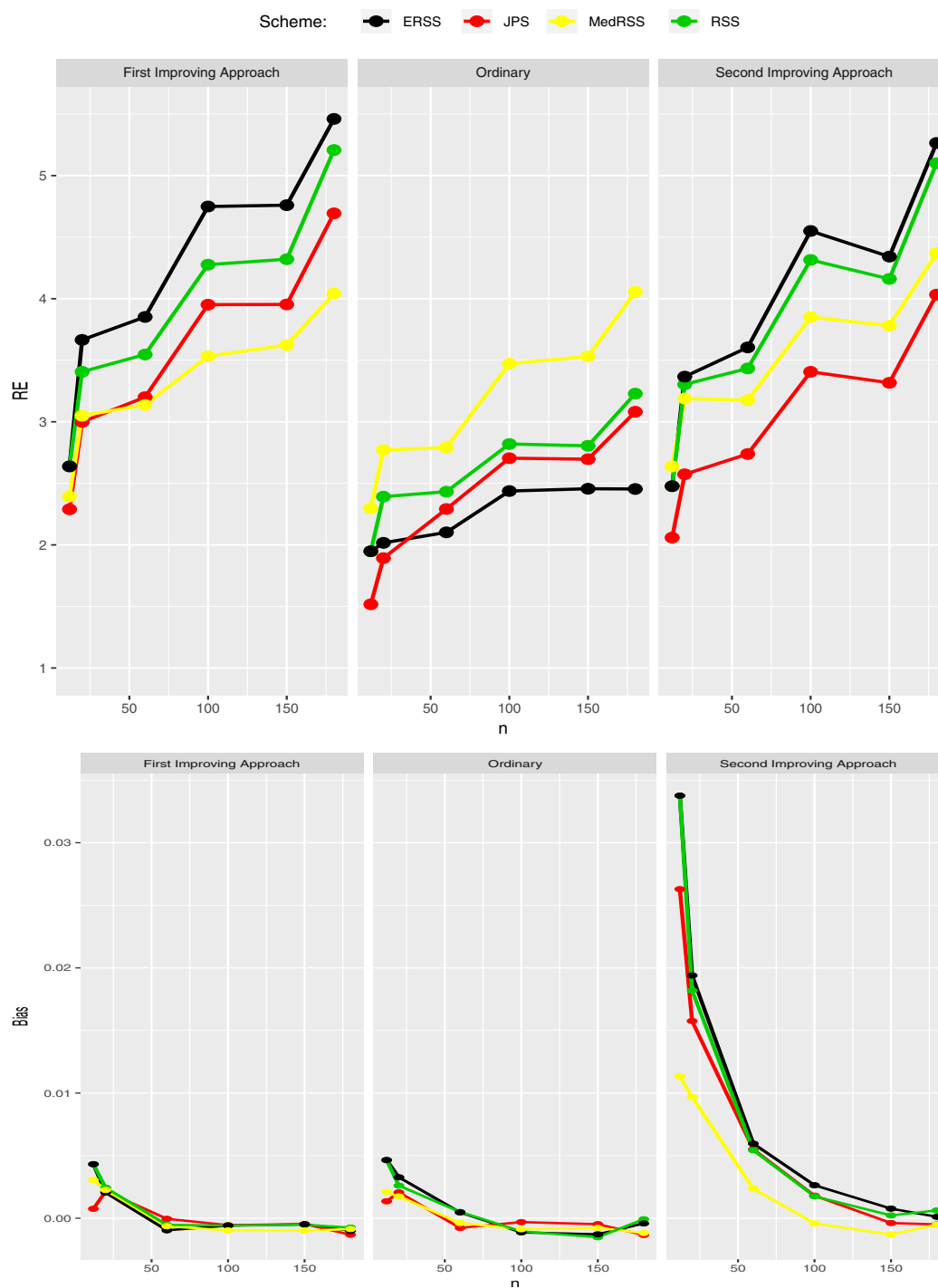


Figure 1: Comparison of the estimators obtained based on various sampling designs proposed via RE and bias criteria under the standard normal distribution.

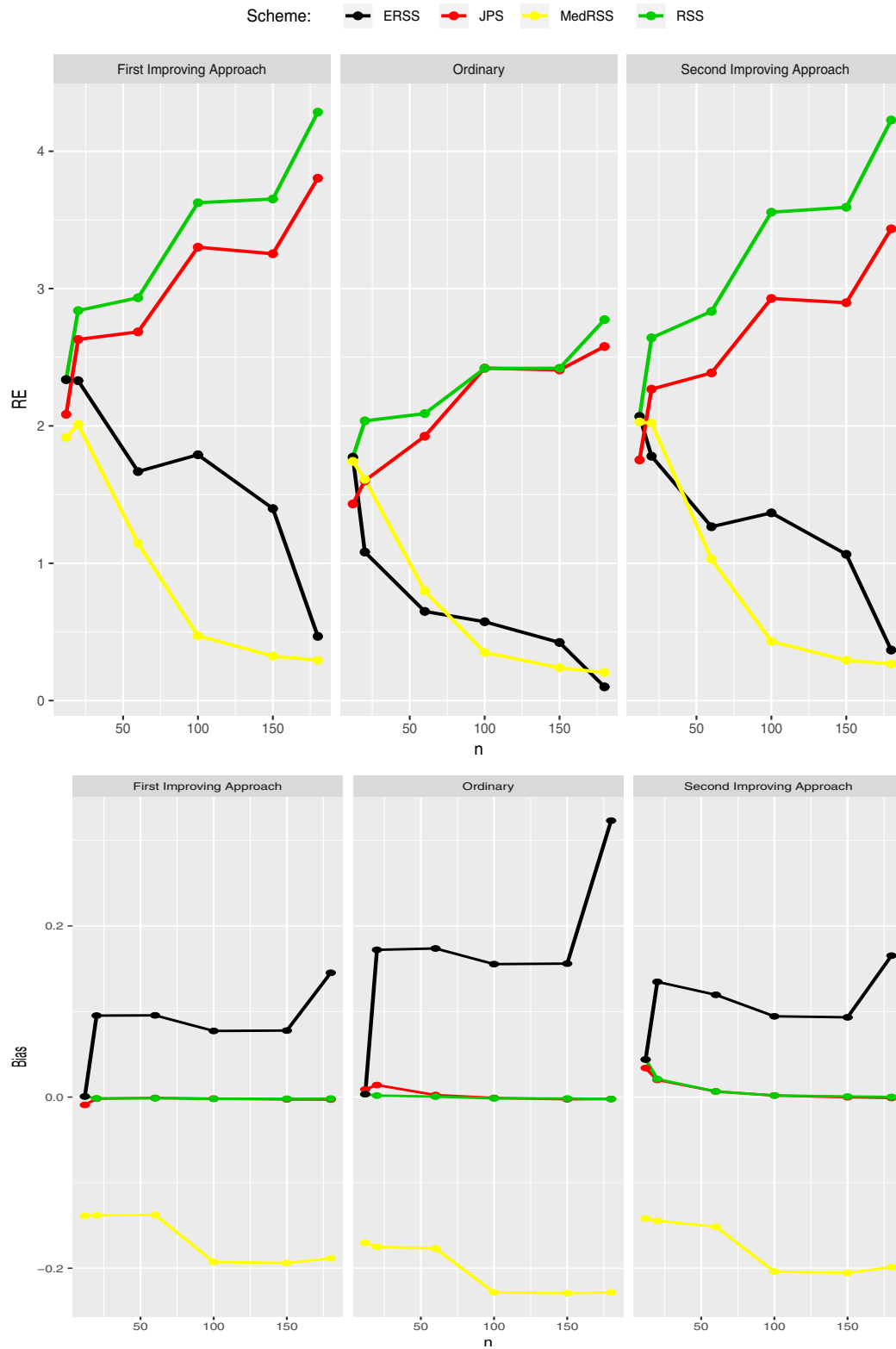


Figure 2: Comparison of the estimators obtained based on various sampling designs proposed via RE and bias criteria under the gamma distribution.

5 Real data analysis

Two real data sets are considered to illustrate the results in the previous sections. To check the chosen sample's sensitivity and obtain the efficiencies of the mean estimators derived based on various designs, the Monte Carlo technique with 2000 replications is employed. Moreover, the judgment ranking process is performed using a covariate variable highly correlated with the variable of interest. Then, the relative efficiencies are computed from (6). It is to be noted that here the variance that appeared in the numerator of (6) is substituted by the *MSE* of the estimator obtained based on the SRS.

5.1 Belgium municipalities data

We implemented the new method on Belgium Municipalities Data. The data set is available in the R package `sampling` under the name `belgianmunicipalities`. It contains different information about the Belgian population between 2001 and 2004 by municipalities. The data set contains 589 rows corresponding to each municipality of Belgium. We just considered two variables, the variable of interest x , which represents the taxable income in euros in 2001, and the auxiliary variable z , which represents the total population on July 1, 2004. As it is clearly observed from Table 7, the improved versions of the JPS and the RSS designs yield more efficient estimators. This fact confirms the results obtained in the preceding section.

Table 7: Belgium municipalities data. The numbers indicate efficiency of the respective design relative to SRS design.

n	M	JPS	IJPS	RSS	IRSS	n	M	JPS	IJPS	RSS	IRSS
12	3	1.09	1.24	1.25	1.36	10	5	1.02	1.68	1.36	1.72
	6	1.11	1.63	1.18	1.50		10	1.05	2.58	1.91	2.83
18	3	1.07	1.22	1.20	1.31	20	5	1.00	1.29	1.51	1.84
	6	1.08	1.59	1.60	2.01		10	1.05	2.14	1.81	2.58
24	3	1.05	1.24	1.20	1.30	50	5	1.33	1.37	1.31	1.62
	6	1.12	1.48	1.53	2.02		10	1.25	1.80	1.71	2.45
48	3	1.16	1.24	1.23	1.38	90	5	1.29	1.39	1.35	1.66
	6	1.32	1.56	1.47	1.86		10	1.49	1.79	1.45	1.96
24	4	1.16	1.48	1.30	1.55	42	7	1.10	1.50	1.52	1.92
	8	1.12	1.77	1.35	1.90		14	1.06	2.02	1.60	2.34
40	4	1.26	1.37	1.32	1.57	70	7	1.15	1.52	1.40	1.78
	8	1.20	1.64	1.49	1.96		14	1.36	2.00	1.67	2.33
56	4	1.14	1.39	1.11	1.36	84	7	1.33	1.56	1.51	1.90
	8	1.12	1.52	1.48	2.00		14	1.39	1.99	1.63	2.25
72	4	1.23	1.37	1.22	1.47	112	7	1.38	1.59	1.49	1.88
	8	1.25	1.55	1.38	1.82		14	1.43	1.90	1.67	2.16

5.2 Switzerland provinces data

For this part we used Standardized fertility measures and socio-economic indicators for 47 French-speaking provinces of Switzerland at about 1888. The data set is available in the R package `stats` under the name `swiss`. We selected two of them as the main and one as the auxiliary variables (see Figure 3):

x_1 : Examination; draftees receiving highest mark on army examination,

x_2 : Agriculture; of males involved in agriculture as an occupation,

z : Education; education beyond primary school for draftees.

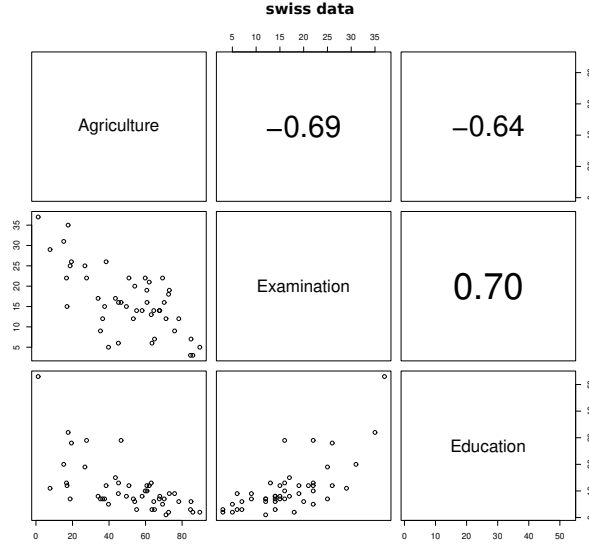


Figure 3: Relation between Switzerland provinces data.

Analogously to the simulation section, the analysis of the Switzerland provinces data in Table 8 suggests that IRSS and IJPS outperform the traditional RSS and JPS designs, respectively. Moreover, from Table 8 and Figure 3, it can be generally concluded that the improved estimators become more efficient as the judgment ranking becomes more precise (as we expected before). Also as an interesting point, as JPS has some bias and we used MSE to measure efficiency, in some cases JPS could not be as efficient as SRS. But we can see that the improved version was always more efficient than SRS.

Overall, the results in Tables 7 and 8 show that using more accurate ranking information can lead to noticeable gains in estimation efficiency in practical settings. For the Belgium municipalities data, the improved JPS and RSS methods consistently outperform their ordinary counterparts, with the gain becoming more highlighted as the set size M increases. For the Switzerland provinces data, the improved methods again provide more efficient estimates in most cases, especially when the auxiliary variable is strongly associated with the variable of interest. These findings indicate that the proposed improved designs are not only theoretically justified, but also practically useful for real data analysis, since they can generate more precise estimation without adding significant cost or complexity.

6 Conclusion

We have proposed simple and practical improvements to JPS and RSS that make use of the complete ranking information within each selected set, which is either already available or can be obtained at virtually no additional cost. The proposed methods are motivated by the Rao-Blackwellization principle and effectively exploit all available rank information, leading to more reliable rank judgments and, consequently, more efficient parameter estimation. They retain the simplicity of the original sampling designs while providing substantial gains in estimation efficiency. In particular, IJPS and IRSS offer a favorable

Table 8: Switzerland provinces data. The numbers indicate efficiency of the respective design relative to SRS design.

n	M	JPS	IJPS	RSS	IRSS	n	M	JPS	IJPS	RSS	IRSS
$z = \text{Education}, x = \text{Agriculture}$											
4	2	0.95	1.08	1.24	1.25	4	0.94	1.23	1.32	1.33	
6		0.94	1.12	1.20	1.22		0.98	1.33	1.35	1.39	
8		0.93	1.13	1.18	1.19		1.05	1.40	1.37	1.39	
10		1.01	1.16	1.13	1.13		1.03	1.35	1.30	1.33	
6	3	0.89	1.16	1.23	1.27	5	0.95	1.29	1.33	1.39	
9		0.92	1.18	1.27	1.31		1.02	1.45	1.40	1.47	
12		0.95	1.24	1.22	1.25		1.03	1.47	1.37	1.44	
15		1.02	1.26	1.18	1.24		1.05	1.44	1.52	1.60	
$z = \text{Education}, x = \text{Examination}$											
4	2	0.94	1.14	1.23	1.25	4	0.97	1.28	1.51	1.57	
6		0.97	1.17	1.24	1.26		1.02	1.39	1.38	1.48	
8		0.97	1.16	1.16	1.19		1.05	1.42	1.35	1.46	
10		1.01	1.22	1.16	1.20		1.12	1.43	1.43	1.55	
6	3	0.90	1.20	1.29	1.33	5	0.92	1.36	1.37	1.45	
9		0.91	1.19	1.21	1.24		1.02	1.45	1.52	1.63	
12		0.93	1.29	1.29	1.34		1.01	1.44	1.51	1.59	
15		1.03	1.30	1.34	1.40		1.04	1.43	1.49	1.57	

balance between improved precision and computational simplicity without introducing significant additional complexity. For symmetric or approximately symmetric populations, we recommend IERSS, which fully exploits the ranking information available in ERSS and generally provides the greatest efficiency gains.

Acknowledgements

The authors thank the reviewers for their valuable comments and suggestions, which have improved the quality of the paper.

References

- Bhushan, S. and Kumar, A. (2024). Novel logarithmic imputation methods under ranked set sampling. *Journal of Statistics and Management Systems*, **27**, 235–257.
- Chen, Z., Bai, Z. and Sinha, B. (2004). *Ranked Set Sampling: Theory and Applications*. Springer-Verlag, New York.
- Dey, S., Salehi, M. and Ahmadi, J. (2017). Rayleigh distribution revisited via ranked set sampling. *Metron*, **75**, 69–85.
- Frey, J. and Ozturk, O. (2011). Constrained estimation using judgment post-stratification. *Annals of the Institute of Statistical Mathematics*, **63**, 769–789.

- Kaur, A., Patil, G.P. and Sinha, A. K. (1995). Ranked set sampling: An annotated bibliography. *Environmental and Ecological Statistics*, **2**, 25–54.
- MacEachern, S.N., Stasny, E.A. and Wolfe, D. A. (2004). Judgment post-stratification with imprecise rankings. *Biometrics*, **60**, 207–215.
- Matthews, M. and Wolfe, D. A. (2017). Unified ranked set sampling. *Statistics and Probability Letters*, **122**, 173–178.
- McIntyre, G. (1952) A method for unbiased selective sampling using ranked set. *Australian Journal of Agricultural Research*, **3**, 385–390.
- Muttlak, H. A. (1997). Median ranked set sampling. *Journal of Applied Statistical Science*, **6**, 577–586.
- Panahbehagh, B., Bruggemann, R., Parvardeh, A., Salehi, M. and Sabzalian, M. R. (2018). An unbalanced ranked-set sampling method to get more than one sample from each set. *Journal of Survey Statistics and Methodology*, **6**, 285–305.
- Panahbehagh, B. and Bruggemann, R. (2021). Introduction into sampling theory, applying partial order concepts. In: Bruggemann, R., Carlsen, L., Beycan, T., Suter, C. and Maggino, F. (eds), *Measuring and Understanding Complex Phenomena*. Springer, Cham.
- Panahbehagh, B., Bruggemann, R. and Salehi, M. (2021). Sampling of multiple variables based on partially ordered set theory. *Journal of the Iranian Statistical Society*, **20**, 307–331.
- Samawi, H.M., Ahmed, M.S. and Abu-Dayyeh, W. (1996). Estimating the population mean using extreme ranked set sampling. *Biometrical Journal*, **38**, 577–586.
- Ozturk, O. (2014) Statistical inference for population quantiles and variance in judgment post-stratified samples. *Computational Statistics and Data Analysis*, **77**, 188–205.
- Ozturk, O. and Bayramoglu Kavlak, K. K. (2020). Statistical inference using stratified judgment post-stratified samples from finite populations. *Environmental and Ecological Statistics*, **27**, 73–94.
- Ozturk, O. and Kravchuk, O. (2021). Judgment post-stratified assessment combining ranking information from multiple sources, with a field phenotyping example. *Journal of Agricultural, Biological and Environmental Statistics*, **26**, 329–348.
- Ozturk, O. (2012). Combining ranking information in judgment post-stratified and ranked set sampling designs. *Environmental and Ecological Statistics*, **19**, 73–93.
- Ozturk, O. (2016). Statistical inference based on judgment post-stratified samples in finite population. *Survey Methodology*, **42**, 239–262.
- Ozturk, O. (2023). Ranked set sampling: A review of its development and applications. *Journal of Applied Statistics*, **50**, 2785–2814.
- Salehi, M., Ahmadi, J. and Dey, S. (2016). Comparison of two sampling schemes for generating record-breaking data from the proportional hazard rate models. *Communications in Statistics - Theory and Methods*, **45**, 3721–3733.
- Stokes, S. and Sager, T. (1988). Characterization of a ranked-set sample with application to estimating distribution functions. *Journal of the American Statistical Association*, **83**, 374–381.
- Takahasi, K. and Wakimoto, K. (1968). On unbiased estimates of the population mean based on the sample stratified by means of ordering. *Annals of the Institute of Statistical Mathematics*, **20**, 1–31.

- Yanagawa, T. and Shirahata, S. (1976). Ranked set sampling theory with selective probability matrix. *Australian Journal of Statistics*, **18**, 45–52.