

An Effective Procedure to Improve Covariates Balance for Multiple Groups in Randomized Trials

Na Li¹ and Zhen Meng^{2,*}

¹School of Applied Science, Beijing Information Science and Technology University, Beijing 100192, China

²School of Statistics, Capital University of Economics and Business, Beijing 100070, China

*mengz@cueb.edu.cn

ABSTRACT

In clinical trials, it is important to compare the differences in treatment effects among multiple groups to determine the effectiveness of certain treatments for a disease. Statistical inferences are typically made assuming that the sample is representative, with subjects randomly assigned to each treatment group and minimal allocation differences. To address these allocation differences, various randomization methods have been proposed, including complete randomization, rerandomization, pairwise sequential randomization, and k -resolution sequential randomization (k -RSR). In this paper, we propose an extension of the k -RSR method called multiple k -resolution sequential randomization (M_k -RSR). M_k -RSR aims to minimize the integrated measurement Mahalanobis distance among multiple groups with equal group size. While the original k -RSR method was developed for two groups, M_k -RSR is designed for two or more groups. By using M_k -RSR, researchers can achieve a balanced allocation and obtain a reasonable estimate of treatment effect. Theoretical results demonstrate that M_k -RSR is more likely to achieve the optimal value compared to existing methods. Extensive simulation studies and an example analysis are conducted to illustrate the superiorities of M_k -RSR.

Additional information

SUPPLEMENT: PROOFS FOR THEOREM 1-3

Proof of Theorem 1. Suppose that M_{Integr0} represents the minimum value of integrated measurement Mahalanobis distance for the fixed units $X_1, X_2, \dots, X_n$. Let $\Omega_1, \Omega_2, \dots, \Omega_K$ denote the set of units in treatment 1, 2, ..., K groups respectively for the realization of M_{Integr0} . Clearly, $\Omega_1 \cup \Omega_2 \cup \dots \cup \Omega_K = \{X_1, X_2, \dots, X_n\}$. We assume that $\Omega_1, \Omega_2, \dots, \Omega_K$ are unique without generality.

Let k_1 and k_2 be two integers with $k_1 < k_2$ and assume $(n \bmod k_1 K) = 0$ and $(n \bmod k_2 K) = 0$. We can define an event A_1 that randomly sequentially assign $k_1 K$ units to K groups with equal group size for unit sequence of $X_{i_1}, X_{i_2}, \dots, X_{i_n}$ to obtain $\Omega_1, \Omega_2, \dots, \Omega_K$, where $i_j \in \{1, 2, \dots, n\}$ and are not equal mutually. Then we have

$$P(A_1) = \frac{\left(\frac{(k_1 K)!}{(k_1!)^K} \right)^{\frac{n}{k_1 K}} \left(\frac{n!}{K!} \right)^K}{n!}. \quad (1)$$

In the same way, we consider another event A_2 for randomly sequentially assigning k_2 units in each resolution to both groups with equal group size for unit sequence of $X_{i_1}, X_{i_2}, \dots, X_{i_n}$ to obtain $\Omega_1, \Omega_2, \dots, \Omega_K$. Thus,

$$P(A_2) = \frac{\left(\frac{(k_2 K)!}{(k_2!)^K} \right)^{\frac{n}{k_2 K}} \left(\frac{n!}{K!} \right)^K}{n!}. \quad (2)$$

There are totally $\frac{n}{k_1 K}$ resolutions for the M_k -RSR with $k = k_1$, and let $Q_i, i = 1, 2, \dots, n/(k_1 K)$ represents the probability for assigning k_1 units in the i th resolution correctly. Then

$$P(Q_i = q_{[j]}) = \frac{1}{C}, \quad j = 1, \dots, C. \quad (3)$$

Consequently, $EQ_i = \frac{1}{C}, i = 1, 2, \dots, n/(k_1 K) - 1, C = \frac{(k_1 K)!}{(k_1!)^K}$ and

$$P(M(k_1 K) = M_{\text{Integr0}}) = \frac{\left(\frac{n!}{K!} \right)^K \left(\frac{(k_1 K)!}{(k_1!)^K} \right)^{\frac{n}{k_1 K}}}{n!} \times Q_1 \dots Q_{\frac{n}{(k_1 K)} - 1} Q_{\frac{n}{k_1 K}}. \quad (4)$$

Once $n/(k_1 K) - 1$ resolutions have been allocated correctly, the final resolution can be allocated correctly. That is $Q_{\frac{n}{k_1 K}} = q_{[1]}$. Hence,

$$P(M(k_1 K) = M_{\text{Integr0}}) = \frac{\left(\frac{n}{K}!\right)^K \left(\frac{(k_1 K)!}{(k_1!)^K}\right)^{\frac{n}{k_1 K}}}{n!} \times Q_1 \dots Q_{\frac{n}{(k_1 K)}-1} Q_{\frac{n}{k_1 K}} \quad (5)$$

and

$$E[P(M(k_1 K) = M_{\text{Integr0}})] = \frac{\left(\frac{n}{K}!\right)^K \frac{(k_1 K)!}{(k_1!)^K}}{n!} \times q_{[1]} \quad (6)$$

Similarly, we get

$$E[P(M(k_2 K) = M_{\text{Integr0}})] = \frac{\left(\frac{n}{K}!\right)^K \frac{(k_2 K)!}{(k_2!)^K}}{n!} \times q_{[1]} \quad (7)$$

According to the Stirling's approximation with $k! \approx \sqrt{2\pi k} \left(\frac{k}{e}\right)^k$, we get

$$E[P(M(k_2 K) = M_{\text{Integr0}})] > E[P(M(k_1 K) = M_{\text{Integr0}})], k_2 > k_1. \quad (8)$$

Proof of Theorem 2. We prove the case of $k=2, K=3$ and the proof of other k is similar. Suppose that $(n \bmod kK) = 0$ and let $Z_i = V^{-1/2} X_i$ where $V^{-1/2}$ is the Cholesky square root of V^{-1} , then $\text{Cov}(Z_i) = I_p, i = 1, 2, \dots, n$. Denote $Z = (Z_1, Z_2, \dots, Z_n)^\top$, $Z_{n+6} = (Z_{n+1}, Z_{n+2}, \dots, Z_{n+6})^\top$. Let B is a $K \times K$ matrix, and the design principle of which is as follows: each row randomly selects two positions, one of which is 1, the other is -1, the rest are 0, and each row must be different. Similarly, let A is a $K \times n$ matrix, each row of A randomly selects k elements to be 1, and all the others to be 0, and there is only one 1 in each column, and the k rows add up to n ones. A_1 is also a $K \times n$ matrix, the design rules are the same as for matrix A . Here is given a special case of B, A and A_1 as follows:

$$B = \begin{bmatrix} 1 & -1 & 0 \\ 1 & 0 & -1 \\ 0 & 1 & -1 \end{bmatrix}, A = \begin{bmatrix} 1 & \dots & 1 & 0 & \dots & 0 & 0 & \dots & 0 \\ 0 & \dots & 0 & 1 & \dots & 1 & 0 & \dots & 0 \\ 0 & \dots & 0 & 0 & \dots & 0 & 1 & \dots & 1 \end{bmatrix}_{K \times n}, A_1 = \begin{bmatrix} 1 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 1 \end{bmatrix}.$$

By the above assumption and define, we get

$$M_{\text{Integr}} = (BAZ)^\top BAZ. \quad (9)$$

We further define

$$\begin{aligned} y_n &= BAZ, \\ \Delta_{n+6} &= BA_1 Z_{n+6}. \end{aligned} \quad (10)$$

Then $\{y_n, y_{n+6}, y_{n+12}, \dots\}$ is a Markov process and $y_{n+6} = y_n + \Delta_{n+6}$. Define the function $F(y_n) = y_n^\top y_n$. Because of the

randomness of A_1 , there are 90 possible outcomes for $A_1 Z_{n+6}$, writing $\mathcal{Z}_{[1]} = \begin{bmatrix} Z_{n+1} + Z_{n+2} \\ Z_{n+3} + Z_{n+4} \\ Z_{n+5} + Z_{n+6} \end{bmatrix}$, $\mathcal{Z}_{[2]} = \begin{bmatrix} Z_{n+1} + Z_{n+3} \\ Z_{n+2} + Z_{n+4} \\ Z_{n+5} + Z_{n+6} \end{bmatrix}, \dots$, $\mathcal{Z}_{[90]} = \begin{bmatrix} Z_{n+5} + Z_{n+6} \\ Z_{n+3} + Z_{n+4} \\ Z_{n+1} + Z_{n+2} \end{bmatrix}$. Then we have

$$\mathbb{E}[F(y_{n+6} | y_n)] - F(y_n) = \mathbb{E}[y_{n+6}^\top y_{n+6} | y_n] - y_n^\top y_n = 2\mathbb{E}[y_n^\top \Delta_{n+6} | y_n] + \mathbb{E}[\Delta_{n+6}^\top \Delta_{n+6} | y_n]. \quad (11)$$

In the equation above, $\mathbb{E}[\Delta_{n+6}^\top \Delta_{n+6} | y_n]$ is a positive constant. To calculate $\mathbb{E}[y_n^\top \Delta_{n+6} | y_n]$, let $(w)^+$ denote the module

of a real number (or a real vector) w . Thus, we have

$$\begin{aligned}
\mathbb{E} \left[y_n^\top \Delta_{n+6} \mid y_n \right] &= \mathbb{E} \left[(BAZ)^\top BA_1 Z_{n+6} \mid y_n \right] \\
&= \mathbb{E} \left[((BAZ)^\top B)(A_1 Z_{n+6}) \mid y_n \right] \\
&= \mathbb{E} \left[((BAZ)^\top B) \mathcal{Z}_{[1]} q_{[1]} + ((BAZ)^\top B) \mathcal{Z}_{[2]} q_{[2]} + \dots + ((BAZ)^\top B) \mathcal{Z}_{[90]} q_{[90]} \mid y_n \right] \\
&= \mathbb{E} \left[(q_{[j_1]} - q_{[j_2]}) \left(((BAZ)^\top B) \mathcal{Z}_{[j_1]} \right)^+ + (q_{[j_3]} - q_{[j_4]}) \left(((BAZ)^\top B) \mathcal{Z}_{[j_3]} \right)^+ + \dots \right. \\
&\quad \left. + (q_{[j_{89}]} - q_{[j_{90}]}) \left(((BAZ)^\top B) \mathcal{Z}_{[j_{89}]} \right)^+ \mid y_n \right] \\
&= (q_{[j_1]} - q_{[j_2]}) \left((BAZ)^\top B \right)^+ \mathbb{E} \left[(\mathcal{Z}_{[j_1]})^+ \mid y_n \right] \mathbb{E} \left[(\cos \theta_{j_1})^+ \mid y_n \right] + \\
&\quad (q_{[j_3]} - q_{[j_4]}) \left((BAZ)^\top B \right)^+ \mathbb{E} \left[(\mathcal{Z}_{[j_3]})^+ \mid y_n \right] \mathbb{E} \left[(\cos \theta_{j_3})^+ \mid y_n \right] + \dots + \\
&\quad (q_{[j_{89}]} - q_{[j_{90}]}) \left((BAZ)^\top B \right)^+ \mathbb{E} \left[(\mathcal{Z}_{[j_{89}]})^+ \mid y_n \right] \mathbb{E} \left[(\cos \theta_{j_{89}})^+ \mid y_n \right],
\end{aligned} \tag{12}$$

it's worth mentioning that θ_{j_i} is the angle between $(BAZ)^\top B$ and $\mathcal{Z}_{[j_i]}$, and $i = 1, \dots, 90$ are not equal and belong to $\{1, 2, \dots, 90\}$. If $\mathcal{Z}_{[1]} = (Z_{n+1} + Z_{n+2}, Z_{n+3} + Z_{n+4}, Z_{n+5} + Z_{n+6})^\top$, then it means allocating Z_{n+1}, Z_{n+2} to treatment 1 group, Z_{n+3}, Z_{n+4} to treatment 2 group and Z_{n+5}, Z_{n+6} to treatment 3 group that results in the smallest Mahalanobis distance $d_{[1]}$. In this way, $\mathcal{Z}_{[l]}$ corresponds to the l th smallest Mahalanobis distance $d_{[l]}, l = 1, 2, \dots, 90$. Finally, if $i_1 < i_2$, then $q_{[j_{i_1}]} - q_{[j_{i_2}]} < 0$, due to the geometric property of the M_k -RSR. Hence, from the last equation, we can get "the drift condition" as in Qin et al(2016). Thus $M(k) = O_p(n^{-1})$.

Proof of Theorem 3. Denote the effects of K treatment groups by $\mu_1, \mu_2, \dots, \mu_K$, respectively. Let the response of the i th unit be $y_i, i = 1, 2, \dots, n$. Naturally, μ_j can be estimated by $\bar{\mu}_j = \sum_{i=1}^n I_{ij} y_i / \sum_{i=1}^n I_{ij}, j = 1, 2, \dots, K$. Hence, the treatment effect difference of the s th treatment group compared with the treatment 1 is $\tilde{\zeta} = \bar{\mu}_1 - \bar{\mu}_s$, where $s = 2, 3, \dots, K$. Suppose that $(n \bmod kK) = 0$ and let $Z_i = V^{-1/2} X_i$ where $V^{-1/2}$ is the Cholesky square root of V^{-1} , then $\text{Cov}(Z_i) = I_p, i = 1, 2, \dots, n$. By the assumption of normality, $Z_i \sim N(0, I_p)$. Define

$$Y = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{bmatrix}, Z = \begin{bmatrix} Z_1^\top \\ Z_2^\top \\ \vdots \\ Z_n^\top \end{bmatrix}, \mathbf{I} = \begin{bmatrix} I_1 \\ I_2 \\ \vdots \\ I_n \end{bmatrix} = \begin{bmatrix} I_{11} & I_{12} & \dots & I_{1K} \\ I_{21} & I_{22} & \dots & I_{2K} \\ \vdots & \vdots & \ddots & \vdots \\ I_{n1} & I_{n2} & \dots & I_{nK} \end{bmatrix}, \tag{13}$$

and $\tilde{Z} = [\mathbf{I}; Z]$, $\beta = (\beta_1, \beta_2, \dots, \beta_p)^\top$, $\beta^* = (\beta_1^*, \dots, \beta_p^*)^\top = (V^{-1/2})^\top \beta$, $\mu = (\mu_1, \dots, \mu_K)^\top$ and $\theta = (\mu^\top, \beta^\top)^\top$, $\theta^* = (\mu^\top, \beta^{*\top})^\top = (\mu_1, \dots, \mu_K, \beta_1^*, \dots, \beta_p^*)^\top$.

The true model can be rewritten as

$$Y = \tilde{X} \theta + \varepsilon = \mathbf{I} \mu + X \beta + \varepsilon = \mathbf{I} \mu + Z \beta^* + \varepsilon = \tilde{Z} \theta^* + \varepsilon. \tag{14}$$

Assume that $L = \underbrace{(1, 0, \dots, 0, -1, \dots, 0)}_K^\top$, where the first entry is 1, and the s th entry is $-1, s = 2, 3, \dots, K$. Although $\tilde{\zeta}$ can be obtained by running the regression $Y = \mathbf{I} \mu + \varepsilon$, the true model is $Y = Z \theta^* + \varepsilon = \mathbf{I} \mu + Z \beta^* + \varepsilon$. Especially, we can write $\tilde{\zeta}$ as

$$\begin{aligned}
\bar{\zeta} &= \frac{\sum_{i=1}^n I_{i1} y_i}{\sum_{i=1}^n I_{i1}} - \frac{\sum_{i=1}^n I_{is} y_i}{\sum_{i=1}^n I_{is}} \\
&= L^\top \left(\frac{\mathbf{I}^\top \mathbf{I}}{n} \right)^{-1} \frac{\mathbf{I}^\top \mathbf{Y}}{n} \\
&= L^\top \left(\frac{\mathbf{I}^\top \mathbf{I}}{n} \right)^{-1} \frac{\mathbf{I}^\top (\tilde{Z}\theta^* + \varepsilon)}{n} \\
&= L^\top \left(\frac{\mathbf{I}^\top \mathbf{I}}{n} \right)^{-1} \frac{\mathbf{I}^\top (\mu + Z\beta^* + \varepsilon)}{n} \\
&= L^\top \left[\mu + \left(\frac{\mathbf{I}^\top \mathbf{I}}{n} \right)^{-1} \frac{\mathbf{I}^\top (Z\beta^* + \varepsilon)}{n} \right] \\
&= \mu_1 - \mu_s + L^\top \left(\frac{\mathbf{I}^\top \mathbf{I}}{n} \right)^{-1} \frac{\mathbf{I}^\top (Z\beta^* + \varepsilon)}{n}.
\end{aligned} \tag{15}$$

When $n \rightarrow \infty$, we have

$$\frac{\mathbf{I}^\top \mathbf{I}}{n} \xrightarrow{p} \begin{bmatrix} 1/K & & & \\ & 1/K & & \\ & & \ddots & \\ & & & 1/K \end{bmatrix} = M. \tag{16}$$

Further, we can define

$$\begin{aligned}
A &= L^\top M^{-1} \left[\frac{\mathbf{I}^\top (Z\beta^* + \varepsilon)}{n} \right], \\
B &= L^\top \left[\left(\frac{\mathbf{I}^\top \mathbf{I}}{n} \right)^{-1} - M^{-1} \right] \left[\frac{\mathbf{I}^\top (Z\beta^* + \varepsilon)}{n} \right],
\end{aligned} \tag{17}$$

so that $\bar{\zeta} = A + B$.

For A , by some algebra, we can show

$$A = \frac{K}{n} \left[\sum_{j=1}^p \sum_{i=1}^n (I_{i1} - I_{is}) Z_{ij} \beta_j^* + \sum_{i=1}^n (I_{i1} - I_{is}) \varepsilon_i \right]. \tag{18}$$

For the first term on the right, we have

$$\sum_{j=1}^p \sum_{i=1}^n (I_{i1} - I_{is}) Z_{ij} \beta_j^* = \sum_{j=1}^p \beta_j^* \left[\sum_{i \in \{i: I_{i1}=1\}} Z_{ij} - \sum_{i \in \{i: I_{is}=0\}} Z_{ij} \right]. \tag{19}$$

where $\{i : I_{i1} = 1\}$ and $\{i : I_{is} = 0\}$ represent the 1th and s th treatment groups. From the proof of Theorem 2 of Qin et al (2016), we understand that $\sum_{i \in \{i: I_{i1}=1\}} Z_{ij} - \sum_{i \in \{i: I_{is}=0\}} Z_{ij}$ is a stationary process under the proposed method. Therefore,

$$\begin{aligned}
&\sum_{i \in \{i: I_{i1}=1\}} Z_{ij} - \sum_{i \in \{i: I_{is}=0\}} Z_{ij} = O_p(1), \\
&\sum_{j=1}^p \beta_j^* \left[\sum_{i \in \{i: I_{i1}=1\}} Z_{ij} - \sum_{i \in \{i: I_{is}=0\}} Z_{ij} \right] = O_p(1).
\end{aligned} \tag{20}$$

In addition, note that $(I_{i1} - I_{is})^2 = 1$ if and only if $I_{i1} = 1, I_{is} = 0$ or $I_{i1} = 0, I_{is} = 1$, so $\sum_{i=1}^n (I_{i1} - I_{is})^2 \varepsilon_i^2 = \frac{2n}{K} \varepsilon_i^2$, we have

$$\begin{aligned} \text{Var} \left(\frac{K}{n} \sum_{i=1}^n (I_{i1} - I_{is}) \varepsilon_i \right) &= \mathbb{E} \left(\frac{K^2}{n^2} \sum_{i=1}^n (I_{i1} - I_{is})^2 \varepsilon_i^2 \right) \\ &= \mathbb{E} \left(\frac{K^2}{n^2} \frac{2n}{K} \varepsilon_i^2 \right) \\ &= \frac{2K \sigma_\varepsilon^2}{n}. \end{aligned} \quad (21)$$

Therefore,

$$\sqrt{n}A \xrightarrow{D} N(0, 2K \sigma_\varepsilon^2). \quad (22)$$

Similarly, we will prove that $\sqrt{n}B \xrightarrow{p} 0$. First, notice that

$$\left(\frac{\tilde{\mathbf{I}}^\top \tilde{\mathbf{I}}}{n} \right)^{-1} - M^{-1} \xrightarrow{p} 0. \quad (23)$$

Therefore, proving $\sqrt{n}B \xrightarrow{p} 0$ is equivalent to prove

$$\frac{\tilde{\mathbf{I}}^\top (Z\beta^* + \varepsilon)}{\sqrt{n}} = O_p(1). \quad (24)$$

First note that

$$\frac{\tilde{\mathbf{I}}^\top (Z\beta^* + \varepsilon)}{\sqrt{n}} = \frac{1}{\sqrt{n}} \begin{bmatrix} \sum_{j=1}^p \sum_{i=1}^n I_{i1} Z_{ij} \beta_j^* + \sum_{i=1}^n I_{i1} \varepsilon_i \\ \sum_{j=1}^p \sum_{i=1}^n I_{i2} Z_{ij} \beta_j^* + \sum_{i=1}^n I_{i2} \varepsilon_i \\ \vdots \\ \sum_{j=1}^p \sum_{i=1}^n I_{iK} Z_{ij} \beta_j^* + \sum_{i=1}^n I_{iK} \varepsilon_i \end{bmatrix}. \quad (25)$$

Since

$$\begin{aligned} \frac{1}{\sqrt{n}} \left(\sum_{j=1}^p \sum_{i=1}^n I_{i1} Z_{ij} \beta_j^* + \sum_{i=1}^n I_{i1} \varepsilon_i \right) &= \frac{1}{\sqrt{n}} \left[\left(\sum_{j=1}^p \sum_{i=1}^n Z_{ij} \beta_j^* + \sum_{i=1}^n \varepsilon_i \right) + \right. \\ &\quad \left. \left(\sum_{j=1}^p \sum_{i=1}^n (I_{i1} - 1) Z_{ij} \beta_j^* + \sum_{i=1}^n (I_{i1} - 1) \varepsilon_i \right) \right]. \end{aligned} \quad (26)$$

By central limit theorem, we have

$$\frac{1}{\sqrt{n}} \left(\sum_{j=1}^p \sum_{i=1}^n Z_{ij} \beta_j^* + \sum_{i=1}^n \varepsilon_i \right) = O_p(1). \quad (27)$$

In addition,

$$\left(\sum_{j=1}^p \sum_{i=1}^n (I_{i1} - 1) Z_{ij} \beta_j^* + \sum_{i=1}^n (I_{i1} - 1) \varepsilon_i \right) = \frac{\sqrt{n}A}{K}. \quad (28)$$

Since $\sqrt{n}A$ converges to a normal distribution,

$$\frac{1}{\sqrt{n}} \left(\sum_{j=1}^p \sum_{i=1}^n (I_{i1} - 1) Z_{ij} \beta_j^* + \sum_{i=1}^n (I_{i1} - 1) \varepsilon_i \right) = O_p(1). \quad (29)$$

Therefore,

$$\frac{1}{\sqrt{n}} \left(\sum_{j=1}^p \sum_{i=1}^n I_{i1} Z_{ij} \beta_j^* + \sum_{i=1}^n I_{i1} \varepsilon_i \right) = O_p(1). \quad (30)$$

Similarly, we have

$$\frac{1}{\sqrt{n}} \left(\sum_{j=1}^p \sum_{i=1}^n I_{is} Z_{ij} \beta_j^* + \sum_{i=1}^n I_{is} \varepsilon_i \right) = O_p(1), s = 2, \dots, K. \quad (31)$$

Hence,

$$\frac{\tilde{\mathbf{I}}^\top (Z\beta^* + \varepsilon)}{\sqrt{n}} = O_p(1). \quad (32)$$

Therefore, $\sqrt{n}B \xrightarrow{p} 0$, and added $\sqrt{n}A \xrightarrow{D} N(0, 2K\sigma_\varepsilon^2)$, by Slutsky's theorem, we have

$$\sqrt{n}(\tilde{\xi} - (\mu_1 - \mu_s)) \xrightarrow{D} N(0, 2K\sigma_\varepsilon^2). \quad (33)$$

For $\hat{\xi}$, we can obtain their asymptotic distributions in similar ways.