



A PURELY SEQUENTIAL PROCEDURE IN ETHICAL ALLOCATION FOR NORMAL RESPONSE VARIABLE

V. N. KADAM¹ and H. S. PATIL²

¹ & ² Department of Statistics, S. B. Zadbukey Mahavidyalaya Barsi-413401
Maharashtra, India.

Email: gholap.vidya86@gmail.com, hspatil1960@yahoo.com

Communicated : 03.12.18

Accepted : 25.01.18

Published: 30.01.19

ABSTRACT:

In the present work, we formulate a purely sequential procedure for Normal response variable depending on the average responses at each stage under the assumption that the treatment variances are equal. It gives larger allocation to better treatment with minimum sample size. Such procedures are used to reduce the cost/ time under certain level of precision.

Keywords: Ethics; Clinical trial; Parallel group design; Superiority trials; Null and alternative hypotheses; Sample size.

INTRODUCTION:

In a sequential clinical trial some information is collected and is examined for it to be good enough to terminate the experiment and take a decision or continue to gather additional information. Such procedures are used to reduce the cost/ time, under certain level of precision. For the details Ghosh et.al.(1997).

In the era of competitions and as a need, invention for new medicines has become inevitable in clinical experiments. However how far a new medicine is superior to one existing is a matter of great concern. This is essentially can be validated by comparing the two drugs from ethical point of view, a better drug be given to a larger proportion of experimental units (patients).

Suppose units (patients) arrive sequentially in a clinical trial and are to be randomized to one of the two treatments. Without any loss of generality assume that a higher value of response variable indicates a more favorable situation. Assigning patients equally to both treatments cannot be recommended from an ethical point of view. Because half of the patients receives an inferior treatment. Therefore in practice it is desirable to have skewed allocations towards the better treatment. In the literature some procedures have been proposed and their performances are evaluated. It is desirable to satisfy certain optimal criteria:

(I) Proportion of the units treated by better treatment is as large as possible and/or

(II) The probability of error (selecting a treatment which is not better) is as less as possible

When planning a trial, an essential step is the calculation of the minimum sample size required to meet the given objectives of the study. Estimating the number of participants is important issue for the planning of clinical trials. A study with small sample size wastes resources and it has ethical implications. A study with large sample size runs the more individuals than necessary, receiving an inferior treatment. Flight and Julious (2015a) has highlighted the key general components required to estimate the sample size of a clinical trial. The first step in a sample size calculation is to establish the trial objective. The choice of endpoint is important in the sample size calculation.

MATERIALS & METHOD:

The effect size is main factor in the estimation of sample size. Suppose two treatments differ by an amount d . This amount d is the effect size. This is also known as a clinically important difference or the minimum value worth detecting. Estimate of the population variability is a component of a sample size calculation. Flight and Julious (2015b) provides a practical guide for applying steps to superiority parallel group clinical trials, where the primary endpoint can be assumed to be normally distributed. In this paper we highlight on superiority clinical trial where the primary end point can be assumed to be normally distributed.

RESULTS AND DISCUSSION:

In superiority trial, null and alternative hypotheses are given as bellow

H_0 : The two treatments are not different with respect to the mean response $\mu_A = \mu_B$.

H_1 : The two treatments are different with respect to the mean response $\mu_A \neq \mu_B$.

That is we want to test that the two means are equal against an alternative that they differ by an amount

d , where d is effect size. Let $f(\mu) = \mu_A - \mu_B =$ difference in the population means and T is the difference in the sample means. Assuming that the data from the clinical trial are sampled from a Normal population, then using standard notation, $T \sim N(f(\mu), \text{Var}(T))$ giving

$$\frac{T - f(\mu)}{\sqrt{\text{Var}(T)}} \sim N(0,1)$$

This statistical test is called a two - tailed test with each tail allocated an equal amount of the type I error (i.e. $\alpha/2$). The sum of these tails is equal to the overall type I error (α). Let $Z_{1-\alpha/2}$ denote the $(1-\alpha/2)$ 100 percentage point of a standard normal distribution.

Thus, an upper 2-tailed, α level critical region for a test of $f(\mu) = 0$ is

$$|T| > Z_{1-\alpha/2} \sqrt{\text{Var}(T)}$$

(1.1)

For this critical region, we needs to test it against an alternative that $f(\mu) = d$, for specified power $(1-\beta)$.

$$d - Z_{1-\beta} \sqrt{\text{Var}(T)} = Z_{1-\alpha/2} \sqrt{\text{Var}(T)}$$

(1.2)

where β is the overall type II error level and $Z_{1-\beta}$ is the 100 $(1-\beta)$ percent point of the standard Normal distribution. Therefore we can write

$$\text{Var}(T) = \frac{d^2}{(Z_{1-\beta} + Z_{1-\alpha/2})^2}$$

(1.3)

Here assumption is that the variances in each group are equal i.e. $\sigma_A^2 = \sigma_B^2 = \sigma^2$. So $\text{Var}(T)$ can be derived as

$$\text{Var}(T) = \frac{\sigma^2}{n_A} + \frac{\sigma^2}{n_B}$$

(1.4)

where $\text{Var}(T)$ will be unknown and depends on the sample size. Substitute equation (1.4) in equation (1.3) we get,

$$\frac{d^2}{\sigma^2 \left(\frac{1}{n_A} + \frac{1}{n_B} \right)} = (Z_{1-\beta} + Z_{1-\alpha/2})^2$$

Now we use sample variance estimate S^2 instead of σ^2 . Therefore we use t- statistics instead of Z- statistics for inference.

$$\frac{d^2}{\sigma^2 \left(\frac{1}{n_A} + \frac{1}{n_B} \right)} \geq (Z_{1-\beta} + t_{1-\alpha/2, n_A+n_B-2})^2$$

(1.5)

By convention, type I error and type II error are fixed at rate of 0.05 and 0.1 or 0.2 respectively. We think not in terms of the type II error but in terms of the power of a clinical trial. Power is one minus the probability of a type II error (i.e. probability of rejecting the H_0 when it is false). Power is usually set to be between 80% and 90%. The minimum power should be 80%. Trials should be designed to have the power as high as possible, preferably at least 90%. A 90% powered study is less sensitive to the assumptions in the sample size calculation than a 80% powered study. For details see Julious (2004).

Our aim is to find the minimum sample size for a fixed total numbers of patients N otherwise the cost of trial becomes infinity. So we propose following purely sequential procedure.

2. A purely sequential procedure:

Let $X_{ik} \sim N(\mu_k, \sigma^2)$ $k = A, B$ and σ^2 is unknown. The following is the proposed procedure.

1. Allocate A and B to m (≥ 2) units. Compute $\bar{x}_{m_A}$ and $\bar{x}_{m_B}$.

2. Allocate treatment B to next unit if $\bar{x}_{m_A} < \bar{x}_{m_B}$ otherwise allocate treatment A.

Let n_A and n_B be the current number of units allocated to A and B respectively such that $n_A + n_B \leq N$. Stop for the first time, if

$$\frac{(\bar{x}_{n_A} - \bar{x}_{n_B})^2}{S^2(1/n_A + 1/n_B)} \geq (Z_{1-\beta} + t_{1-\alpha/2, n_A+n_B-2})^2 \quad (2.1)$$

where β is the overall type II error, $Z_{1-\beta}$ is the $100(1-\beta)$ percent point of the standard Normal

distribution, $t_{1-\alpha/2, n_A+n_B-2}$ is the critical value of t distribution with n_A+n_B-2 degree of freedom (d.f) at $(1-\alpha/2)$ level of significance and

$$S^2 = \frac{1}{n_A + n_B - 2} \left[\sum_{i=1}^{n_A} (x_{iA} - \bar{x}_{n_A})^2 + \sum_{i=1}^{n_B} (x_{iB} - \bar{x}_{n_B})^2 \right]$$

, an unbiased estimator of σ^2 .

3. Once stopped allocate remaining $N - n_A - n_B$ units to treatment A (B) if $\bar{x}_{n_A} \geq (<) \bar{x}_{n_B}$.

Since N is finite, the rule (2.1) terminates with probability 1.

3. Simulation study: In this section, we compute n_A and n_B by simulation based on 10,000 repetition

with $N = 80$ and 1000 , $\mu_B = 1$, $m = 2$, $\mu_A = 1, 1.4, 1.8, 2.2, 2.6, 3$ by using R program. Take 90% power (i.e.

$\beta = 0.1$ therefore $Z_{1-\beta} = Z_{1-0.1} = 1.282$) and a two-sided 5% type I error rate (i.e. $\alpha = 0.05$), following results are obtained.

CONCLUSION:

Remark 4.1: From above table as μ_A increases, sample size of treatment B (n_B) decreases.

Remark 4.2: From above table as μ_A increases, total sample size ($n_A + n_B$) decreases and hence reduces the cost.

REFERENCE :

- Flight L, Julious S. A. (2015a): Practical Guide to Sample Size Calculations: An Introduction. *Pharmaceutical Statistics*, 15(1): 68-74.
- Flight L, Julious S. A. (2015b): Practical Guide to Sample Size Calculations: Superiority Trials. *Pharmaceutical Statistics*, 15(1): 75-79.
- Ghosh M., Mukhopadhyay N. and Sen, P.K. (1997): *Sequential Estimation*. New York: Wiley.
- Julious S. A. (2004): Tutorial in Biostatistics: Sample Sizes for Clinical Trial with Normal Data. *Statistics in medicine* 23:1921-1986.
- Patil H.S. (2010): Sequential Procedures for Some Non-regular and Regular Families of Distributions. *Ph.D. Thesis, Shivaji University, Kolhapur*.

Table 3.1

μ_A	Sample sizes by purely sequential procedure			
	N=80		N=1000	
	n_A	n_B	n_A	n_B
1	39	38	483	475
1.4	56	20	728	201
1.8	63	9	791	74
2.2	62	4	746	21
2.6	53	3	612	6
3	41	2	445	3