

Design of Experiments (DOE): Recent Advances and Development

Poly. E. Chigbu

Department of Statistics,
University of Nigeria, Nsukka

Abstract

The role of Design of Experiment (DoE) or Experimental Design in researches in all disciplines is very significant, since its applicability ensures reliability of findings in any study. Design of experiment is concerned with planning experiments with specific objectives in order to obtain maximum amount of information from available resources. In the area of DoE, two phases are usually involved: design and analysis phases. Simplicity and efficiency are required of every experimental design, which are achievable by adhering to these three basic and pertinent principles of experimental design: replication, randomization and blocking. An experimental design is specified by describing exactly the method of randomly assigning (randomizing) the treatments of treatment structure to the experimental units of the block structure. Recent advances have shown that randomization in experimental design is not only the process of randomly assigning treatments to experimental units of the block structure by means of random process, but also that of ensuring that the validity of the conclusions drawn from experiment are free from the biases of the experimenter. This is achieved by introducing the notion of Hinkelmann and Kempthorne (2008) called design-random variables which are used to obtain a derived linear model (randomization model). Randomization model is a model that does not require any assumption concerning the form of the response being a function of the values of the factors influencing it as it is the case of the normal model, where the assumption of normality of the response variable is made evident.

1.0 OVERVIEW

1.1 Preamble

Understanding and appreciating the meanings of terminologies: Statistics, Combinatorics and “objects” would go a long way in making us appreciate the recent advances and development we are going to dwell on in this presentation. The word “objects” simply means *configurations* which implies array (arrangement of symbols or letters or things in rows and columns), as in, for example, semi-Latin square, quasi-semi-Latin square, etc.

But in the mean time, we bear in mind the following definition of *configuration* which is quite relevant here and which is also given in The Concise Oxford Dictionary: “form, shape, or figure resulting from an arrangement of parts or elements in some manner” [1].

1.2 Basic Concepts

Before discussing the main subject matter of this lecture, let us take some time to briefly discuss its rudiments.

As you may be aware, the word *statistics* is usually used in two senses: first, it means numerical data relating to any field of endeavour, be it the Arts/Humanities, Sciences/Engineering or what have you; second, it refers to the scientific process for collecting, understanding, analyzing and interpreting numerical and non-numerical data. For the purpose of the second, let us consider that there exist two sets of people on earth: the “Scientists” and the Statisticians. While a “Scientist” is anybody who requires his/her statistical problem to be solved, the Statistician is anybody that has undergone the required standards of training in statistics for becoming and possesses the required skills of a Statistician; which also lends credence to statistical consulting.

Corresponding author: **Poly. E. Chigbu**, E-mail: -, Tel. +234 8033930082

Statistics is also an inductive science, which attempts to generalize concepts based on particular cases; deals with the whole based on information from a part; and draws inferences about populations on the basis of samples. A special ingredient for achieving these is **randomization** (a process whereby every member or item of a population of interest has equal chance of being observed as a member or item of an observable sample or part under investigation).

Statistics as a discipline studied at both undergraduate and postgraduate levels at many tertiary level institutions in different parts of the world has among others the stress area known as Design and Analysis of Experiments often called Experimental Design or Design of Experiments (DOE). Design of Experiments was developed in the 1920s by R.A. Fisher at the Rothamsted Experimental Station near Cambridge, UK and since then it has been generally and widely applied in all disciplines [20].

While, the *analysis* phase of DOE involves the illustration of techniques, which enable the experimenter to analyze the experimental information in, say, a first- or second-order response (regression) model or even an analysis-of-variance, the *design* phase involves the presentation and illustration of experimental layouts for the fitting of the models under study.

Historically, DOE and another stress area known as *Regression analysis* have developed separately. But it turns out that one of the best ways of appreciating the power of designed experiments is by first understanding regression analysis. Regression models are a statistical way of characterizing relationships between variables. A regression model can be defined in words as:

$Y = \text{function of } X + \text{random variation};$

Symbolically written as:

$$Y = f(X) + \varepsilon$$

where Y represents a response variable called the dependent variable, X represents a predictor variable called an independent variable, $f(X)$ represents the systematic or repeatable part of the relationship between X and Y , ε represents the variation in Y which is not related to X or to any other measurable variable.

However, there are two types of variables (responses) in DOE: *controlled* and *uncontrolled*. The controlled variables refer to the experimental factors which are selected by the experimenter for comparison while the uncontrolled variables are measurements or observations that are recorded but cannot be controlled by the experimenter.

The main essence of DOE therefore is the designing of efficient and/or “best” experiments since life itself is full of experiments and experimentation on daily basis and whatever experiment that is being designed would need to be analyzed in some context by adopting certain appropriate techniques or criteria for interpretation. DOE entails the Statistics and Combinatorics aspects studied as Statistical and Combinatorial designs, respectively.

We all know that the term *statistical* is an adjective of the word *statistics*. However, the statistical design studies generally involve

- (i) the determination of the relationships or otherwise between variables, attributes, criteria, properties or factors as might from time to time be represented by one regression (design) model or another;
- (ii) the comparison of variables, attributes, etc, of an idealized representative model with the view of determining their effects in the representative model and/or the closeness of this model to a real-life situation.

Contextually, it entails the hierarchical classification of a collection of configurations or design arrays, which possess a certain class of properties satisfying a particular design model, based on the comparison of variances of their treatments/treatment contrasts’ estimates or some other conditions which not only relate to the variances but on the nature of the incidence of treatments to the experimental units (plots).

The statistical design is usually presented from the standpoint of the general linear model, wherefrom least squares estimators are developed and discussed using possibly the notation of matrices. It also involves the use of designs (factorial and fractional factorial) to fit regression models with their attendant analysis-of-variance (ANOVA).

The term *combinatorial* is indeed a mathematical adjective pertaining to the word *combination*, which simply relates to the combination of items. The combinatorial design basically entails the exploitation of certain mathematical properties (or patterns) of configurations (designs), which are analogous to standard statistical basis of comparison or judgment (popularly called optimality criteria), in the making of good or “best” choices. On the whole, it involves studying the patterns of the application of treatments to plots in an experimental layout. In this regard, it can then be easily said that a particular pattern (graphically or otherwise or according to the level of treatment concurrences in plots) is better than another.

From the foregoing, it can easily be seen that while the “Scientists” most times apply already-constructed designs (known in the literature as Classical designs) to their problems for analyses and interpretation, the Statisticians would most of the times be bothered with, among others, conceiving and constructing new designs (often called Optimal designs) for experimentation.

Indeed, the way and manner inputs (treatments) of an experiment are combined in application to plots determines the efficacy of the experiment and the fruitfulness of the output (yield). There is no gainsaying the fact that the terms *treatment* and *plot* commonly used in DOE today originated from the interest of early British/European researchers like F. Yates and R.A. Fisher in this area whose interest was mainly on agricultural experimentation that could lead to bumper harvest amidst famine after the World wars of the 1930's and 1940's.

Thus, beyond the earlier senses Statistics had been described it is on the whole, a science of the “best” decision-making in experimentation. While a popular dictum states that Classical statistics dwells on making valid decisions under uncertainty, it can be stated here that in recent times, the aspect of statistics popularly considered dwells on making optimum/optimal choices under many possibilities.

2.0 Some Relevant Terminologies and Their Meanings

2.1 Experiment

Even though an experiment could simply be literally called a test [2], an experiment is a set of procedures which are carried out under a set of conditions, and which may occur repeatedly for a result: see, for example, Arua et al [3].

In an experiment, one or more variables (or factors) are deliberately changed in order to observe the effect the changes have on one or more response variables.

2.2 Experimental Unit (Plot)

This is the smallest thing to which a treatment could be applied. However, when a response is measured from this smallest unit, it is called an *observational unit* [4].

In a real-life situation, the meaning of an experimental unit may vary from experiment to experiment. For example, while in an experiment involving the growing of varieties of a crop in genuine plots in a field, the *experimental units* are the *genuine plots*, the *experimental units* are the *sub-plots* in an experiment involving the growing of the same varieties of crop in *whole-plots* with fertilizers applied to *sub-plots* [4].

2.3 Experimental Treatment

An *Experimental Treatment* is the entire description or totality of what is applicable to a plot at a given place or time, which gives rise to a measurable observation.

In the two real-life experiments of section 2.2, the *experimental treatments* are respectively, *variety of crop* and *variety-fertilizer combination*.

2.4 Design

Let Ω denote the set of plots of an experiment while T denotes a whole set of its treatments. A *design* can easily be said to be the assignment of treatments to plots, i.e. it is a function, f , which maps the elements of Ω to the elements of T ($f: \Omega \rightarrow T$). This implies that a plot $\omega \in \Omega$, say, gets treatment, $f(\omega)$, during experimentation.

Bailey [5] also defines *design* as a mapping or function, f , of the experimental units (plots), Ω , to treatments, T , in such a way that, treatments are allocated to plots. This definition seems to be circular in meaning, but it works well enough in practice.

In this regard, therefore, the usual aim of designing an experiment is to choose f such that certain combinatorial properties or patterns are satisfied or exploited, as the case might be.

An (experimental design) is specified by describing exactly the method of randomly assigning the treatments (*treatment structure*) to the experimental units in their natural setting before the application of treatments (*block structure*).

2.4.1 Treatment Structure: The treatment structure of an experimental design consists of the set of treatments or treatment combinations that the experimenter wants to study and make inferences on. They could be levels of treatment (one-factor treatment structure) or levels of two or more treatments combined together (two-factor or higher-factor treatment structure).

2.4.2 Block Structure: The block structure consists of the grouping of the experimental units into homogeneous blocks, which already exist before the application of treatments. The block structure of an experiment can range from a Completely Randomized Design (CRD) structure with no blocking criterion to block structures with multiple blocking criteria.

2.5 Design of Experiments

This is the process of planning and designing experiments with specific objectives, so that appropriate data can be analyzed by statistical methods, resulting in *valid* conclusions.

It begins as soon as an experiment has been formulated and ends when all data have been collected for analysis. The statistical approach to experimental design is necessary if we wish to draw meaningful conclusions from the data collected.

Every experimental design is expected to be as much as possible simple and efficient; simple in the sense that the simplest possible experimental design is chosen among many possible designs in order to achieve the same specified objectives while efficiency implies that the experiment should be designed in such a way that time, cost, personnel and experimental materials are saved/minimized.

Thus, “the essence of DOE therefore is the designing of efficient and/or “best” experiments since life itself is full of experiments and experimentation on daily basis and whatever experiment that is being designed would need to be analyzed in some context by adopting certain appropriate techniques or criteria for interpretation” [6]

3.0 Basic Principles of DOE

There are three basic principles of DOE which are applied to reduce or remove experimental bias: *Replication*, *Randomization* and *Blocking*.

- i. **Replication:** This is the repetition of the entire or some portion of an experiment under the same conditions. It helps in obtaining an estimate for the experimental error (i.e., variations in response from treatments that are treated or handled alike) and also permits the experimenter to obtain a more precise estimate of the experimental factor.
- ii. **Randomization:** Contextually, this is the process of randomly allocating treatments to experimental units by means of random devices. It is important because it guards against bias from the experimenter.
- iii. **Blocking:** It involves the grouping of heterogeneous experimental units into sub-groups with homogeneous units such that units within the same sub-group have the same characteristics, to a large extent. Indeed, blocking helps to reduce the variations that exist among the experimental units of a given experiment.

4.0 Statistics and Statistical Design

Some attempts have been made above to define Statistics literally, technically and contextually. Here, we further state that in every statistical design, experimenters are usually interested in knowing about the effect of treatments applied and the plots on which they are applied. Usually, an experimenter has more control on the set of treatments applied than on the plots that receive these treatments.

Another usual issue of interest to experimenters is the allocation of treatments to plots: the randomization of this activity gives the statistical validity of the experiment, which forms the basis of any statistical design.

Thus, each design consists of two sets and a function between them. The sets are:

- a set that consists of treatments, denoted by T , say, and
- a set of experimental units, denoted by Ω .

These two sets are always finite.

On the whole, therefore, suppose Ω is a set of “things” (persons, animals, plants, machines, small portions of land, etc) upon which data are to be measured; a *factor* on Ω is a function, f , on Ω where we are interested not so much in the values of the function, f , as in which “things” have the same value of f .

Example: In a medical experiment involving a set, Ω , of persons, $f(\text{person})$ is the drug given to that person. To compare (the effect of) drugs, therefore, we need to know who and who had the same drug. The set of “things” with one given value of f is a subset of Ω . All such subsets form a *partition* of Ω and every “thing” is in one and only one such subset [6]. In this regard, we shall always assume that the yield on plot ω , y_ω , is given by

$$y_\omega = p_\omega + t_{f(\omega)}, \quad (1)$$

where $t_{f(\omega)}$ is a constant depending on the treatment, $f(\omega)$, applied to plot ω , p_ω is a random variable, depending on ω [6, 7].

Equation (1) is regression-based which leads to modeling statistical (regression-based) designs solved by the usual traditional regression analysis principles. This forms the basis of the Statistics aspect of our discussion here.

5.0 Combinatorics And Combinatorial Design

The term *Combinatorics* was defined by Street and Street [8] as “the branch of mathematics which deals with the problems of selecting and arranging objects in accordance with certain specified rules”. However, in studying Combinatorics, we always deal with Configurations.

Furthermore, Berge [9] regarded Combinatorics as that which counts, enumerates (constructs and classifies), examines and investigates the existence of Configurations with certain specified properties or characteristics. In doing all these, the knowledge of the concept of isomorphism and isomorphism classes is quite essential.

Let us not be bored with the details of technical mathematical terms, but briefly speaking and contextually, the *isomorphism* of two configurations implies their *sameness*.

Hence, the Combinatorics aspects of the work pertaining to DOE specifically involve the application of group and graph theory ideas in the construction and classification of Designs. These involve, for instance, selecting and arranging the treatments of the design, as in the case of semi-Latin squares, based on the elements of the *Symmetric group*, S_n , and in accordance with some definitions and constraints of the design.

A typical example is with the construction of semi-Latin squares where eventually different configurations and combinatorial patterns of the squares are grouped into two kinds of isomorphism classes based on their graphs ((treatment)variety-concurrence), family of the number of treatment-pairings which occur different number of times within the blocks of the squares known as Combinatorial parameters' method: see Preece and Freeman [10] for statistical design properties and the permutation sets associated with the arrangement of the treatments in the squares.

At this juncture, giving an illustrative example would help out in quickly fixing the ideas of combinatorial design with those of Replication, Randomization and Blocking. Thus, an experiment was conducted in an experimental area consisting of two fields, each divided into three strips of land, to compare three different varieties of corn: Red, White and Yellow, in combination with four amounts of phosphate fertilizer: 0, 70, 140 and 210 (all in kg/hectare). Each strip was made up of four plots. The measured yield or response of this experiment was the total weight of the starch harvested from each plot.

The cultivation of the varieties of corn in the farms was mechanized such that varieties were sown on whole-strips and not in small areas due to practical possibilities. On the other hand, it was very convenient to apply the phosphate fertilizers to smaller areas of land within a whole-strip, here known as plots [4]. A typical layout of this experiment is given in Figure 1.

140	0	210
70	140	70
0	70	140
210	210	0
↑	↑	↑
Red	Yellow	White

0	70	140
70	140	0
210	0	210
140	210	70
↑	↑	↑
Yellow	Red	White

Figure 1: A typical layout of the Cultivation of Varieties of Corn.

Considering Figure 1 for the above illustrative example, we notice the *existence of pattern* and the *lack of it*. The existence of pattern is depicted by the fact that each amount of fertilizer is applied to one plot per strip while each variety is applied to one strip per field. This pattern and studies pertaining to it manifests the combinatorial design investigation for this experiment. On the other hand, the existence of lack of pattern is depicted by, noting that there exists neither any systematic order in allocating the varieties to strips in each field nor any systematic order in the allocation of the amounts of fertilizer to plots in each strip. Having mentioned that Figure 1 is a typical layout of the experiment, it means that there could be other layouts, which derive their credence from some *randomization*. However, the lack of pattern of the above illustration, which is elicited by randomization subtly make for the statistical design discussed earlier. However, due to randomization, more layouts like those of Figure 1 could be made manifest in order to increase efficiency as lower variances would manifest with larger number of *Experimental units* or *design points*. Each of the two fields in this experiment has the characteristics of a *block*; hence *blocking* of the entire experiment could be seen to have been exploited.

Hence, Combinatorial design, as an integral aspect of the study of DOE [11] is “a way of choosing, from a given finite set, a collection of subsets with particular properties” [8].

The general algebraic connections between the analysis procedures of Statistical and Combinatorial designs are articulated in Chigbu [6].

5.1 ILLUSTRATIONS WITH THE SEMI-LATIN SQUARE

5.1.1 Definition of semi-Latin square

With reference to the various ways the term semi-Latin square is defined [10, 12, 13, 14], we hereby define a semi-Latin square as follows: An $(n \times n)/k$ semi-Latin square is an arrangement of nk symbols (treatments) in an $(n \times n)$ square array such that each row-column intersection contains k symbols and each symbol occurs once in each row and each column [7, 19].

	1	...	N
1	(1, ..., k)	...	(1, ..., k)
.	.	...	.
.	.	...	.
.	.	...	.
N	(1, ..., k)	...	(1, ..., k)

Figure 2: An array for an $(n \times n)/k$ semi-Latin square

In Figure 2, $(1, \dots, k)$ indicates the size, k , of each block but not the actual symbols of a semi-Latin square. For convenience, a row-column intersection of a semi-Latin square is also called a *block* as long as no ambiguity of usage is introduced by doing this.

When $n = 3$ and $k = 2$, a typical semi-Latin square would be of the form in Fig. 3.

1 2	3 4	5 6
3 4	5 6	1 2
5 6	1 2	3 4

Figure 3: A typical semi-Latin square for $n = 3$ and $k = 2$,

where the symbols from 1 to 6 represent the six treatments of the square.

A typical layout and square (configuration) for this use of semi-Latin square are shown in Figures 4 and 5, respectively.

		Housewives					
W	1	-	-	-	-	-	-
E	2	-	-	-	-	-	-
E	3	-	-	-	-	-	-
K	4	-	-	-	-	-	-

Figure 4: Layout of $(4 \times 4)/2$ semi-Latin for the Consumer Testing Experiment

		Housewives							
W	1	1	2	3	4	5	6	7	8
E	2	3	4	1	2	7	8	5	6
E	3	5	6	7	8	1	2	3	4
K	4	7	8	5	6	3	4	1	2

Figure 5: A $(4 \times 4)/2$ semi-Latin square

Each position marked with a dash in Figure 4 accommodates a particular Vacuum Cleaner while the symbols in each row-column intersection of Figure 5 represent different Vacuum Cleaners. Also Figure 5 is a typical configuration as there are other possible configurations for the same description and this is why the configuration or arrangement of choice would possess some inherent good statistical qualities, which depend on the nature of concurrences of the symbols in the row-column intersection of the array.

6. Recent Advances And Developments In DOE

6.1 Combinatorial Aspects

In defining the semi-Latin square, it is also conventional to disregard the order in which k symbols are written in any row-column intersection. This is simply because in adopting the semi-Latin square for the purpose of experimental design, individual “little” columns have always been assumed to have no statistical role. A semi-Latin square can therefore be thought of as not so much as an $(n \times nk)$ (n rows and nk columns) Latin rectangle in one combinatorial sense, but rather as an $(n \times n)$ square array with each row-column intersection containing k unordered symbols [10]. On the other hand, the individual “little” columns might be judged to have some statistical role to play in the design of experiments. Though some extensions are possible, much depends on the roles attached to the rows and sets of “little” columns of the squares. From the foregoing reasoning, therefore, different structures of designs, which are not semi-Latin squares per se, may arise. Each structure necessitates a particular way of viewing its isomorphism. These structures could be called *quasi-semi-Latin squares*.

A *quasi-semi-Latin square* is here defined as a combinatorial object whose entries are *ab initio* arranged as in the semi-Latin square formation without any regard to any other block structure apart from the one associated with the usual semi-Latin square but which actually has a peculiar blocking system as might eventually be defined as the case may be [15].

6.2 Randomization Aspects

It has been shown from recent research that randomization in experimental design is not only the process of randomly assigning treatments to experimental units of the block structure of the design, but also to ensure the validity of the conclusions drawn from experiment are free from bias. That is, it provides the basis for making inferences without requiring assumptions about the distribution of the experimental error. This is achieved by introducing the notion of Hinkelmann and Kempthorne [16] called *design-random variables* which are used to obtain a derived linear model known as *randomization model*.

Randomization models take into account the randomization process employed in assigning treatment combinations to the experimental units. They do not require any assumption concerning the form of the response being a function of the values of the factors influencing it as it is in the case of **normal model**; where the assumption of normality of the response variable is made evident. Mewhort [17] referred to randomization model as a **chance model**. The chance model used in a randomization test is based on action carried out in a particular experiment, namely, the experimenter's assignment of experimental treatments to experimental units.

The randomization model is essentially based on two assumptions: the way in which the investigation is carried out and the way in which observed values for a given individual change when it is moved from one group to another. The first assumption is that each possible arrangement of an individual's treatment in the experimental pattern is as likely as any other. This implies that each of the treatments has equal chance as likely as any other. This can be effected by using tables of random numbers, or some equivalent procedure and is called randomization; while the second assumption states that the values observed for a particular individual in a group differ from what would be obtainable if the individual were in another group: see Neter et al [18] for numerical illustration. Any experimental design analysis based on this procedure is known as *randomization-based analysis*. In the randomization-based analysis, there is no assumption of a normal distribution and independence; based on these, the randomization model has the following advantages: we are more confident of the nature of the random variables representing residual variation and at any time we can carry out a test of significance without introducing the assumption of normality of distribution.

7.0 Concluding Remarks

Combinatorial design studies are manifested from studies of patterns or arrangements of experimental units (naturally or otherwise). Of course, Randomization gives a lot of credence to these studies especially as it concerns all possible patterns of a given design.

In spite of the fact that randomization-based analysis has the slight disadvantage that its exposition involves tedious algebra and extensive time, experimenters are encouraged to use it since its applicability helps in overcoming the assumption of normality which must be made in order to adopt ANOVA procedures in analysis.

REFERENCES

- [1] Fowler, H.W. and Fowler, F.G. (1990), *The Concise Oxford Dictionary of Current English*, Oxford University Press.
- [2] Montgomery, D.C. (1991), *Design and analysis of experiments*, John Wiley & Sons.
- [3] Arua, A.I., Chigbu, P.E., Chukwu, W.I.E, Ezekwem, C.C. and Okafor, F.C. (2000), *Advanced Statistics for Higher Education (Vol. 1)*, The Academic Publishers, Nsukka.
- [4] Bailey, R.A. (2008), *Design of Comparative Experiments*, Cambridge Series in Statistical and Probabilistic Mathematics (No. 25).
- [5] Bailey, R.A. (1989), *Designs: mappings between structured sets*, Surveys in Combinatorics, 1989 (ed. J. Siemons), London Mathematical Society Lecture Notes Series 141, Cambridge University Press, Cambridge, 22 – 51.
- [6] Chigbu, P.E. (1998), *Block Designs: efficiency factors and optimality criteria for comparison*, Linco Press Ltd., Enugu.
- [7] Chigbu, P.E. (1995), *Semi-Latin squares: methods for enumeration and comparison*, Ph.D. Thesis, University of London.
- [8] Street, A.P. and Street, D.J. (1987), *Combinatorics of experimental design*, Oxford University Press.
- [9] Berge, C. (1971), *Principles of Combinatorics*, Academic Press Inc., New York and London.
- [10] Preece, D.A. and Freeman, G.H. (1983), *Semi-Latin squares and related designs*, J. Roy. Statist. Soc. Ser. B, 45, 267 – 277.

- [11] Bailey, R.A. (1991), Strata for randomized experiments, *J. Roy. Statist. Soc. Ser. B*, 27 – 66.
- [12] Bailey, R.A. (1988), Semi-Latin squares, *Journal of Statist. Planning and Inference* 18, 299 – 312.
- [13] Bailey, R.A. (1990), An efficient semi-Latin square for twelve treatments in blocks of size two, *Journal of Statist. Planning and Inference* 26, 263 – 266.
- [14] Bailey, R.A. (1992), Efficient semi-Latin squares, *Statistica Sinica* 2, 413 – 437.
- [15] Chigbu, P.E. and Oladugba, A.V. (2008), On block structures, null analyses of variance and expectations of mean squares of some quasi-semi-Latin squares, *preprint*.
- [16] Hinkelmann, K. and Kempthorne, O. (2008). *Design and Analysis of Experiments*. New York: Wiley.
- [17] Mewhort, D. J. K. (2005). A comparison of randomization with the F-test when error is skewed. *Behaviourial Research Method* 37, No. 34, 26-435.
- [18] Neter et al (2005). *Applied Linear Statistical Models*. WCB/McGraw-Hill.
- [19] Chigbu, P. E. (2009). Semi-Latin Squares and Related “Objects”: Statistics and Combinatorics Aspects, *Inaugural Lecture Series* No 43, University of Nigeria, Nsukka.
- [20] Fisher, R. A. (1925). *Statistical Methods for Research Workers*. Oliver and Boyd, Edinburgh.