

ADDIS ABABA UNIVERSITY
COLLEGE OF NATURAL SCIENCES
OFFICE OF GRADUATE STUDIES
DEPARTMENT OF STATISTICS



**ASSESSMENT OF PERFORMANCE OF
COMMONLY USED EXPERIMENTAL DESIGNS
(RANDOMIZED COMPLETE BLOCK DESIGN,
LATTICE DESIGN AND ALPHA-LATTICE
DESIGN) IN THE ETHIOPIAN AGRICULTURAL
RESEARCH SYSTEM**

October 2010

**ADDIS ABABA UNIVERSITY
COLLEGE OF NATURAL SCIENCES
OFFICE OF GRADUATE STUDIES
DEPARTMENT OF STATISTICS**

**ASSESSMENT OF PERFORMANCE OF COMMONLY USED
EXPERIMENTAL DESIGNS (RANDOMIZED COMPLETE BLOCK
DESIGN, LATTICE DESIGN AND ALPHA-LATTICE DESIGN) IN THE
ETHIOPIAN AGRICULTURAL RESEARCH SYSTEM**

By

Bacha Edosa

A Thesis submitted to the Office of Graduate Programs of Addis Ababa
University in partial fulfilment of the requirement for the Degree of Masters of
Science in Statistics

ADDIS ABABA UNIVERSITY
COLLEGE OF NATURAL SCIENCES
OFFICE OF GRADUATE STUDIES
DEPARTMENT OF STATISTICS

ASSESSMENT OF PERFORMANCE OF COMMONLY USED
EXPERIMENTAL DESIGNS (RANDOMIZED COMPLETE BLOCK
DESIGN, LATTICE DESIGN AND ALPHA-LATTICE DESIGN) IN THE
ETHIOPIAN AGRICULTURAL RESEARCH SYSTEM

By

Bacha Edosa

Approved by the Board of Examiners:

Department Head

Emmanuel A. Johannes (Dr.)

Examiner

Prof. Geben Wewochko

Examiner

Signature

[Signature]

Signature

[Signature]

Signature

Acknowledgements

First and foremost, I would like to acknowledge the one and true God, my Lord Jesus Christ for blessing me with the ability and guidance to succeed in every aspect. My heartiest thanks go to my advisor Dr. Girma Taye for his considerable contribution to the topics and direction of this thesis and his invaluable guidance, conversations, stretched patience, encouragement and support of various kinds through some difficult times. Without him this journey would never have begun.

I am highly indebted to the staff members of Department of Statistics, Addis Ababa University for their kind assistance in many ways.

I would like to gratefully acknowledge Ato Getachew Girmay, Dr. Abiy Girmay, Ato Faji Suleiman, Ato Hailemariam Mamo, Ato Abera Regasa and Mulu Chala for their invaluable support for my thesis. I am very much beholden to my best friends Dibaba Bayisa, Merga Belina and Bona Adugna for their incessant support and encouragement throughout the masters program.

My warm thanks go to my family members whose endless encouragement, moral and financial support were the sources of inspiration to me, during my entire academic career, including the period I was pursuing the M.Sc. study.

My thanks is also extended to the Oromia Finance and Economic Development Bureau for giving me an opportunity to participate in the graduate program.

Table of contents

<u>Contents</u>	<u>pages</u>
List of Tables.....	v
List of Figures.....	vi
ABSTRACT.....	vii
CHAPTER ONE.....	1
1. INTRODUCTION.....	1
1.1 Background of the study.....	1
1.2 General objective of the study.....	3
1.3 Application of the results.....	3
1.4 Limitation of the study.....	3
CHAPTER TWO.....	4
2. LITERATURE REVIEW.....	4
CHAPTER THREE.....	12
3. DATA AND METHODOLOGY.....	12
3.1 Data.....	12
3.2 Methodology.....	13
3.2.1 Introduction.....	13
3.2.2 Common experimental designs and treatment structures.....	15
3.2.2.1 Randomized Complete Block Design (RCBD).....	15
3.2.2.2 Missing observations in RCBD.....	17
3.2.2.3 Lattice Designs.....	18
3.2.2.3.1 Balanced Lattice.....	19
3.2.2.3.2 Partially Balanced Lattice.....	19
3.2.2.4 Alpha-Lattice Designs.....	20
3.2.2.5 Split-plot treatment structure.....	21
3.2.3 Fixed Effects Linear Model.....	22
3.2.4 Linear Mixed Effects Model.....	23

3.2.4.1	Estimation Methods Linear in Mixed effects Model.....	23
3.2.4.1.1	Analysis of Variance.....	23
3.2.4.1.2	Maximum Likelihood.....	24
3.2.4.1.3	Restricted Maximum Likelihood.....	24
3.2.5	Analysis of Covariance.....	25
3.2.6	Testing of hypotheses.....	26
3.2.7	Model Diagnostics.....	27
3.2.7.1	The normality Assumption.....	27
CHAPTER FOUR.....		28
4. STATISTICAL DATA ANALYSIS AND RESULTS.....		28
4.1 Testing for Normality Assumption.....		28
4.2 Analysis of the Linear Fixed Effect Model.....		28
4.2.1 Randomized Completely Block Design Analysis.....		29
4.2.2 Analysis of RCBD with missing observations.....		32
4.2.3. Analysis of Lattice Design.....		37
4.2.4. Analysis of Alpha Lattice Design.....		37
4.2.4 Analysis of Covariance.....		38
4.2.5 Split-plot Experiment Analysis.		39
4.3 Linear Mixed effects model Analysis		41
4.3.1 Analysis of Randomized Complete Block Design.....		41
4.3.2 Analysis of Lattice Design.....		42
4.3.3 Analysis of Alpha Lattice Design.....		43
4.4 Discussion.....		44
CHAPTER FIVE.....		46
5. CONCLUSIONS AND RECOMMENDATIONS.....		46

5.1 Conclusions.....46

5.2 Recommendations.....46

REFERENCES.....47

APPENDICES.....51

 Appendix I51

 Appendix II54

List of Tables

Table 1: Analysis of Variance for a RCBD

Table 2: Analysis of Variance for a Lattice Design

Table 3: Analysis of Variance for an Alpha-lattice Design

Table 4: Analysis of Variance for Split-plot

Table 5: ANOVA of RCBD for Bread Wheat Variety in 2005/06 of BW99RVII data set

Table 6: ANOVA of RCBD for Bread Wheat Variety in 2006/07 of BW99RVII data set

Table 7: ANOVA of RCBD for Durum Wheat Variety of DW99RVTR data set

Table 8: ANOVA of Durum Wheat for location one year 2004/05 of DW99RVTR data set with one missing value

Table 9: Analysis of variance of RCBD with two missing observations of durum wheat variety

Table 10: Analysis of variance of RCBD with three, four, five, six and eight missing observations of durum wheat variety compromise

Table 11: Efficiency of lattice design for Maize National Variety Trial of Melkasa, Dhera, Ziway, Mieso, and Kobo research centres in 2009/08 as compared to RCBD.

Table 12: Efficiency of Alpha lattice as compared to RCBD in 2005/06 Season

Table 13: Analysis of Covariance of RCBD on durum wheat variety trial

Table 14: ANOVA of Split-plot for Integrated Disease Management (IDM) on potato variety of 2006 in Jeldu, Degem and Wolmer sites

Table 15: Covariance Parameter estimates of RCBD for BW99RVII data set

Table 16: Covariance Parameter estimates of lattice design for melkasa Research centre

Table 17: Covariance Parameter estimates of alpha lattice design for 2A3099AW and 2A3499AW data sets

Table 18: Tests for Normality of RCBD in DW99RVTRII data set

Table 19: Tests for Normality of RCBD in DW99RVTR data set

Table 20: Tests for Normality of lattice design, alpha lattice design and Covariance analysis

Table 21: Tests for Normality of Split-plot

Table 22: The missing values and its estimated values using estimation approach in RCBD

Table 23: Mean and Standard error (SE) of Durum wheat variety with and without covariate



List of Figures

Figure 1: QQ-Plot of Residuals from RCBD in 2006/05 of BW99RVII data set

Figure 2: QQ-Plot of Residuals from RCBD in 2007/06 of BW99RVII data set

Figure 3: QQ-Plot of Residuals from lattice design

Figure 4: QQ-Plot of Residuals from Alpha lattice design of 2A3499AW data set

Figure 5: QQ-Plot of Residuals for Split plot of data set from Jeldu site

Figure 6: QQ-Plot of Residuals for ANCOVA

ABSTRACT

The choices of experimental design as well as of statistical analysis are of great importance in field experiments. To improve production and productivity of crop yield it is important to employ tested and proven experimental design. The aim of the study is to evaluate the precision of commonly used experimental designs (randomized complete block design, lattice design and alpha-lattice design) in the Ethiopian agricultural research system. Both fixed and linear mixed effect models were used to compare the efficiency of experimental designs. Different variety trials conducted at different locations by the National and Regional Agricultural Research Organization in 2004/05–2008/09 years (seasons) were used in this study. The results indicate that on average the randomized complete block design (RCBD) was more efficient than completely randomized design (CRD). For large missing values in RCBD, the use incomplete block approach is more efficient than using missing-estimated approach. Also use of covariance analysis for damaged yield in RCBD increased the relative efficiency of RCBD compared to CRD than the analysis without covariate. The lattice design was more efficient to increase the precision of agricultural field experiments than RCBD. The relative efficiency of alpha lattice design compared to RCBD was 101.9% and 102.4% for both data sets respectively. This indicates that researchers have in general correctly controlled the local variability during experimentation. However, the cases where blocking was not significant can be attributed to technical issues such as orientation and direction of blocking which should be improved by researches in similar experiments. Based on the results we conclude that RCBD is more efficient than CRD and lattice and alpha lattice designs are more efficient than RCBD. In order to increase the precision of agricultural field experiments researchers are advised to use RCBD for small number of treatments, lattice design and alpha lattice designs for large number of treatments.

CHAPTER ONE

1. INTRODUCTION

1.1. Background of the study

The choices of experimental design as well as of statistical analysis are of great importance in field experiments. These are necessary to be done correctly in order to obtain the best possible precision of results. The experimental design could be considered as the golden state of any experiment due to its aim to ensure that the experimenter is able to detect the treatment effects that are of interest by using the available resources to obtain the best possible precision. The choice of design as well as the choice of statistical analysis can make a huge difference (Leonardo Journal of Sciences, 2009). Precision is the ability of an experiment to detect a true treatment effect. It can be improved by increased replication, treatment selection, improved technique to reduce the variability among units treated alike, increasing the size of experimental units (within limits), the use of covariance, and the employment of a more efficient experimental design and method of analysis (Little and Hills, 1978).

Design of experiments forms the backbone of any research endeavor in the discipline of agriculture and allied sciences. Dr. MS Swaminathan in his message to International Conference on Statistics and Informatics in Agricultural Research said, "It is the effective use of the tools of statistical design of experiments that paved the way for the green revolution." This statement spells out the importance of the subject of Design of Experiments for improving the quality of agricultural research (Parsad and Gupta, 2000).

The foundations of the statistical approach to experimentation were laid by R.A. Fisher in the early 1930s. The subject evolved in agriculture but is now applicable in almost all sciences, engineering and arts. The aim of an experiment is to compare a number of treatments on the basis of the responses produced in the experimental material. The confidence and accuracy with which treatment differences can be assessed will depend to a large extent on the size of the experiment and on the inherent variability in the experimental material.

Design of experiments is an essential component of research in agriculture. It is through the data collected from designed experiments that valid inferences are drawn. In order to make research globally competitive, it is essential that sound statistical methodologies be adopted in the collection and analysis (Sharma, 2000). The subject of design of experiments deals with the statistical methodology needed for making inferences about the treatments effects on the basis of responses (univariable or multivariable) collected through planned experiments.

Experimental design is also the area of statistics concerned with designing an investigation to meet the study objectives, as well as the assumptions for statistical inferences.

In general, experimental design encompasses the rules and procedures used in initiating, organizing, and analyzing experiments. Several experimental designs are being used in the agricultural research system. Randomized complete block designs have been widely used in agricultural field experiments compared to other known designs such as the Latin square and families of incomplete block designs. The major reasons for the popularity of this design are its computational convenience, flexibility and efficiency given the required assumptions that are fulfilled. In the pre-computer era these properties were highly required. A review done to assess the type of design used in early days shows the importance of randomized complete block design (RCBD) in the history of field experimentation. Most of these studies show that the layout of these field trials were RCBD with or without split-plot. These studies include Cox (1950), Panse (1958), and Parsad and Gupta (2000)

The majority of research trials in Sub-Saharan Africa are also planned as RCBD regardless of the size of experiments. Based on a survey conducted on the type of design used in field trials in Ethiopia, it was found that about 78 percent of the designs used in breeding and agronomy research programmes is RCBD with and without split-plot. Such designs are often applied at all levels of research, including nursery, regional variety trial, and national variety trials. This resulted in squeezing a large number of treatment combinations into a complete block, which led to loss of homogeneity within the block. This invalidated assumptions for tests of hypothesis affect the analysis result (Girma, 2005).

Despite the efforts made to develop Ethiopian agriculture over the years the problem of hunger, famine, and malnutrition and land degradation still linger and present the greatest threat to the survival of the nation. Furthermore, deforestation accelerated soil erosion, and land degradation are now serious problems in Ethiopia. As a result crop and livestock yields are generally very low and the recurrent drought has aggravated the situation. To improve production and productivity of crop yield it is important to employ tested and proven experimental design. By selecting appropriate experimental designs, one would increase the precision of agricultural experiments enabling the generation and dissemination of appropriate technology.

The common experimental designs and treatment structures being used in Ethiopian agricultural research system are randomized complete block design, lattice design, alpha-lattice design, factorial experiments, and split-plot treatment structure (Girma, 2002 and Mandefro, 2005). One may ask several questions: why all these experimental designs are

used? Could there be any confusion in their use? Are their relative importance studied and documented? Which one is more appropriate in agricultural experiments? To answer these questions, assessment of their performance in agricultural experiments is important. In the Ethiopia scenario, however, there is no known study that documented the relative performance of these designs.

The aim of this study is, therefore, to identify commonly used designs, assess their performances in Ethiopian agricultural research system and provide recommendations. The initial technique in the analysis of most experimental data is the analysis of variance. Initially this analysis of variance technique was developed for considering differences between means, but later it came to be adapted to estimating variance components.

1.2 General objective of the study

The purpose of this study is to assess and document commonly used experimental designs in the National Agricultural Research System (NARS) setting, and evaluate their relative performance.

Specific objectives of the study are:

- To assess the commonly used experimental designs in Ethiopian agricultural experiments.
- To compare the efficiency of the commonly used experimental designs by Ethiopian Agricultural Research Organization and selecting of the efficient designs.
- To estimate missing values and analyze the RCB design using these estimated missing values
- Recommend suitable design(s) for use in the various research scenarios.

1.3 Application of the results

- The result of the study will help to identify the appropriate experimental designs for field experiments.
- The result will help to improve the precision of agricultural field experiments through design and analysis.
- The result will also help as a basis for further studies in this area.

1.4 Limitation of the study

- ❖ Lack of literature in our country related to the subjects under the study.
- ❖ Lack of access to sufficiently large data sets that cover different agro-ecology and over several years.

CHAPTER TWO

2. LITERATURE REVIEW

The main purpose of field experiments (Hatfield, 2000) is to compare effectiveness of different treatments. Precision and accuracy are paramount, but valid assessment of error is also important. Yield is influenced by non-treatment factors, such as pests and soil fertility. If ignored, this extraneous variation leads to erroneous comparisons. Proper field design and statistical analysis will help minimize this problem. Classical methods for controlling extraneous variation include replication, blocking, and randomization. Whereas the first two help to increase precision in the experiment, the last one is used to decrease bias. Most agronomic field experiments are based on the concepts of replication, local control (blocking) and randomization (Atkinson and Bailey, 2001)

Replication is used for the purpose of increasing precision by reducing standard error and increases representation since wider area is used. Without replication no estimation of experimental error is possible (Blau and Han, 2007). Whatever the source of the experimental error, replication of the experiment steadily decreases the error associated with the difference between the average results for two treatments, provided that precautions (such as randomization) have been taken to ensure that one treatment is no more likely to be favoured in any replicate than another, so that the errors affecting any treatment tend to cancel out as the number of replication increased (Cochran and Cox, 1957). Randomization ensures that no treatment is consistently favoured or discriminated being placed under best or unfavourable conditions, thereby avoiding bias. Randomization also ensures independence among observations, which is a necessary condition for validity of assumption to provide significance tests and confidence intervals. The use of randomization has been justified in many ways. Its basic purpose is to remove bias from the estimation of treatment effects, and to equalize the error overall treatment differences (Atkinson and Bailey, 2001). Randomization is often considered the best protection and assurance against malicious manipulation of plot layout. Randomization is also believed to better justify the assumption of normal errors (Van ES and Sellman, 2007).

Blocking is grouping of plots into blocks of more or less uniform plots. So, plots within the same block are homogeneous. Blocking can be used to eliminate its effect on the comparison among treatments and without blocking no blocking effects is taken into consideration (Blau and Han, 2007). Some reasons for blocking are:

- i) Differences among blocks could be removed from the experimental error to increase the precision of an experiment.
- ii) Treatments are compared most under the same conditions.
- iii) It increases information from an experiment as blocks could be placed at different locations.

The objective of experimental design is to select and group the experimental material so that the experimental error in the experiment is reduced. Thus, the experimental units on which the treatments are to be compared should be as much alike as possible so that a smaller significant difference between two treatments can be detected (Milliken and Johnson, 1992).

Design of experiments had its origin in supplying layout plans of experiments for comparison among a number of experimental treatments in regard to some of their responses when these are applied to a set of experimental units under certain conditions. Consideration of precision of such comparison led to the necessity of suitably grouping the experimental units and then allotting the treatments to these groups (Yigrem Tefera, 2007).

The randomized complete block (RCB) design is the most frequently used method of controlling variability. The principal advantages of RCBD (Mandefro, 2005) are: it is more precise than completely randomized design (CRD) when block effect is significant, any number of treatments and replications can be included, the statistical analysis is easy, and it provides information on the uniformity of experimental units. This design is one of the most frequently used designs as it provides a required precision with limited cost. In this design, if each treatment is applied to exactly one unit in the block and comparisons are only drawn between treatment responses from the same block, extraneous variability should be reduced. However, the size of within block error increases with the number of treatments in the block. This has a direct influence on the precision of treatment contrast. As the size of block increases, variance per unit for variety contrast increases and ultimately leads to inefficient estimates of precision. Effective control of error variance usually requires relatively small blocks (Eskridge, 2008). Alternative designs that increase efficiency and precision of field trials by accounting for variation are families of incomplete block designs (IBD) which utilize small blocks. Efficiency is the ratio, in percent, of the error variance of the randomized block design to the error variance of the design being compared (Snyder, 1962). Thus, if an incomplete block design had six replications and an efficiency of 150 percent, nine replications of randomized block would be needed for equal accuracy.

Incomplete block designs (IBD) occur as balanced or partially balanced. In balanced incomplete block designs all pairs of treatments occur together within a block the same

number of times. Since each block does not contain all treatments, block and variety effects are confounded (Pearce, 1983).

Incomplete block designs such as lattice designs provide more precise estimates when the homogeneity condition does not hold. Lattice designs (Yates, 1936) are extensively used in agricultural field experiments especially for varietals trials. These designs are resolvable, but the requirement that the number of treatments be a complete square is a limitation. A block design is resolvable if the blocks can be partitioned into replicates, defined as sets of blocks with the property that each treatment is assigned to one unit in each set.

Van Es and Sellman (2007) reported that, RCBD is the popular design for field experiments. Of the 414 agronomic field experiments in USA, the majority (72%) were implemented as RCB designs. Completely randomized, randomized incomplete block, split block and Latin square designs were rarely used. They also described that the vast majority (96.7 %) of agronomic field experiments conducted by agronomists are implemented through RCB designs for their simplicity and intuitive layout.

Yates (1939, 1940) has pointed the advantages of resolvable designs. On page 325 of his 1940 paper he noted that

"cases will arise in which the use of ordinary randomized blocks will be more efficient than the use of incomplete blocks, whereas lattice designs can never be less efficient than ordinary randomized blocks."

This advantage of lattice designs is shared by all other resolvable incomplete block designs. Yates (1940) further stated that

"incomplete block designs which cannot be arranged in complete replications are likely to be less value in agriculture than ordinary lattice designs. Their greatest use is likely to be found in dealing with experimental material in which the block size is definitely determined by the nature of the material."

In the class of equally replicated designs with v treatments, b blocks and a common block size k , a balanced incomplete block (BIB) design whenever existent, is the most efficient design for making tests versus control comparisons according to various efficiency criteria. But in the unrestricted class of all connected designs, the limiting efficiency of a BIB design with respect to most efficient design is only 50% for tests versus control comparisons (Parsad and Gupta, 2000).

Campbell and Bauer (2007) reported on improving the precision of cotton performance trials conducted on highly variable soils of the south-eastern USA coastal plain. They provided that on average the relative efficiency of the randomized complete block analysis relative to

the completely randomized design was for lint weight per boll (146 percent) and lint yield (118 percent). The relative efficiency of the randomized complete block design relative to the completely randomized design for lint weight per boll and lint yield ranged from 97 to 368 percent and 98 to 140 percent, respectively.

Parsad and Gupta (2000) showed that the simplest and most commonly used block design when the treatments are at several levels of a single factor by the agricultural researcher is a randomized complete block design (RCBD). In a RCBD every treatment appears in every block precisely once. A RCBD is the most efficient design because there is no loss of information in estimating treatment contrasts as well as block contrasts. However, when the number of treatments in experiments increases, the blocks for a RCBD become large and it is not possible to maintain homogeneity within the blocks. This may result in large intra-block variances and large coefficients of variation. To circumvent this problem, recourse is made to incomplete block designs. An incomplete block design is one in which at least one block does not contain all the treatments. Incomplete block designs are non-orthogonal and their efficiency factor as compared to RCBD is less than 1 under the assumption of equal intra-block variances. But this is offset by the reduced intra-block variance in the incomplete block designs. Consequently, the paired comparisons are estimated more precisely in the case of incomplete block designs and hence the coefficient of variation also gets reduced.

The randomized block, Latin square, and other complete block (Masood, *et al.* 2008) types of experiments are inefficient for large number of treatments, because of their failure to adequately minimize the effect of soil heterogeneity. Generally, the greater the heterogeneity within blocks, the poorer the precision of variety effect estimates. Incomplete block designs are arranged in relatively small blocks that contain fewer varieties than the total number of varieties to be compared. Consequently, there is a gain in precision due to use of small blocks. As far as the layout of the experiment is concerned the incomplete block designs are no more difficult than randomized blocks. Some extra planning is involved in drawing up and randomizing the experimental plan. Randomized complete block design (RCBD) is affordable when the block size contains small varieties/treatments. It is always useful to use incomplete block design when the number of varieties increases. Because of large number of treatments, the homogeneity among experimental units within a large block cannot be maintained. As a result, estimate of experimental error is inflated and results are low in precision.

The usual approach through local control by blocking is inefficient and a lot of research has been carried out which suggest new methods of local control in field experiments (Gleeson

and Cullis (1987), Cullis and Gleeson (1991), Kempton *et al.* (1994)). Alpha designs introduced by Patterson and Williams (1976) are now routinely used for statutory field trials in the United Kingdom (Patterson and Silvey, 1980) and are widely used for breeding and varietal trials in Australia and elsewhere. They are more flexible than lattice designs and can accommodate any number of varieties.

Additional improvement is possible through modeling field variability using spatial features of the field layout. It has been advocated by Wu and Dutilleul (1999) that use of incomplete blocking is generally more effective in reducing the unexplained structured variation in comparison with complete blocking. The advantage of alpha designs is that they are easy to construct, and can be constructed in cases where balanced incomplete block designs and lattice designs do not exist. The early alpha designs were aimed primarily at controlling variation down the columns of plots in the field. This is often adequate when plots are long and narrow. Patterson and Hunter (1983) have demonstrated the value of alpha designs in such circumstances in terms of gain in efficiency.

Masoon *et al.* (2006) compared efficiency of alpha lattice design. The results indicated that alpha lattice design improved the efficiency 8 to 9 percent as compared to RCBD. Yau (1997) reported the use of alpha lattice design in international yield trials of different crops and found average efficiency 18 percent higher than the RCBD.

Lattice designs are an important class of incomplete block designs. Lattice designs for $v = k^2$ treatments are also PBIB designs based on Latin-square association scheme, where k is the size of the block. Balanced lattice designs originate from a group of lattice designs and are available only when the square root of the number of treatments or varieties is either a prime number or a prime power.

Gunjaca, *et al* (2005) studied the efficiency of alpha lattice designs included 152 data sets in Croatian variety trials of cereal and non-cereal variety trials. They found that the maximum relative efficiency of alpha lattice design compared to RCBD in cereal and non-cereal varieties were 1.37 and 1.55, respectively. Raza and Masood (2009) compared the relative efficiency of the lattice design to the RCBD on three data sets. They found that for the three data sets alpha lattice design increased the precision of experiments by 26%, 17% and 55%, respectively.

Masood *et al.* (2008) reported that the efficiency of alpha lattice design and randomized complete block design (RCBD) was compared in two research trials conducted at the National Agricultural Research Centre, Islamabad, Pakistan on two different crops to assess the efficiency of each in minimizing experimental error, coefficient of variation and error

mean square for yield variable. Alpha designs are used for field trials because they provide better control on experimental variability among the experimental units under field conditions. For this purpose, two research trials on wheat and potato crops using alpha lattice design were conducted at the National Agricultural Research Centre, Islamabad. The results show improvements in the precision level (in terms of decline in the error mean square (EMS), coefficient of variation (CV) and standard error of difference). The CV calculated for wheat and potato preliminary yield trials are 7.70 and 13.6 for alpha lattice design and 8.54 and 16.39 for RCBD, respectively. The MSE calculated for these trials are 0.95 and 2.07 for respectively and 1.18 and 3.01 for RCBD, respectively. The standard error of difference for alpha lattice design is 218 and 0.8312, and 343 and 1.41 for RCBD. The value of relative efficiency 1.24 and 1.46 indicates that the use of alpha lattice design instead of RCBD increased experimental precision by 24 and 46 percent respectively. In the study at three sites (NARC, BARI and AARI), Idrees and Khan (2009) showed that general lattice design (alpha lattice design) was on average more efficient than complete block analysis in reducing the error mean square. Maximum average efficiencies obtained were 2.23 at NARC, 3.64 at BARI and 2.33 at AARI.

Alves *et al.* (2009) compared the efficiency of RCBD, alpha-design, and row-column design in genotypic mass selection. Their result indicated greater efficiency for alpha-design and row-column design, enabling more precise estimates of genotypic variance, greater precision in the prediction of genetic gain and consequently greater efficiency in genotypic mass selection.

Spring Bread Wheat Germplasm Improvement for Increased Yield and Yield Stability in West Asia and North Africa (WANA) project (2000) reported that to deal with the highly variable experimental and environmental conditions in WANA, efforts were made to identify and utilize efficient experimental designs. In 1999/2000 season, incomplete design and alpha-lattice design were used instead of randomized block design in all yield trials conducted at Tel Hadya, Breda, and Terbol stations. Results obtained from yield trials conducted at International Center for Agricultural Research in the Dry Areas (ICARDA) main research stations revealed that the use of alpha-lattice design exhibited relative efficiency of 120 – 206 % over RCBD. In both international yield trials distributed in 2000, the usually used RCBD was replaced by alpha-lattice design to improve the efficiency of identifying superior cultivars (<http://www.icarda.cgiar.org/>).

Patterson and Hunter (1983) reported the efficiency of alpha designs relative to other incomplete block designs. The efficiency factor is used in the selection of a design from a

group of incomplete block designs. Design efficiency is computed as a ratio of average variance of pair-wise variety difference of complete block design to the average variance of pair-wise variety difference for an incomplete block designs. If the resulting value is greater than 1.0 the latter design is more precise than the former. Using a large collection of experiments Patterson and Hunter have shown that alpha designs on the average produced a 30 percent gain in efficiency over designs which did not use incomplete block designs. They also reported that the use of generalized lattice designs (alpha lattice designs) instead of complete block designs in 244 cereal variety trials grown in UK has resulted in average reduction of 30% in variances of varietal yield differences. The lattice designs were most effective when the number of varieties was more than 50 but worthwhile reductions in variance, averaging about 24% were obtained in trials with fewer than 20 varieties.

Shunmugasundaram *et al.* (1980) compared a simple lattice design with an augmented design to select sugarcane clones. They found that the two designs closely agreed on the top 10% of the 86 sugarcane clones under selection. The relative efficiency of lattice design was much higher than that of the augmented design.

Analysis of covariance is an extension of the analysis of variance technique which controls the experimental error by taking into consideration the dependence of response variable on the covariate variables (Gharde, 2000). He also added that use of analysis of covariance should be considered in experiments in which blocking cannot adequately reduce the experimental error.

Parsad and Gupta (2000) compared the analysis of variance without using covariate and with using the covariate in RCBD. They obtained that the analysis with covariate highly reduced the error mean square and coefficient of variation than analysis without covariate.

The analysis of variance imposes some restrictions on the experimental design toward increasing accuracy. The equal number of replicates per treatment is required if all factor levels have equal importance, maximizing the precision of mean comparisons (Neter *et al.* 1990). Furthermore, the F-test is relatively insensitive to small departures from the assumption of equal sample variances if the sample sizes are equal (Montgomery, 1991).

Kerr (2003) compared three common designs using a mixed model analysis of variance. There are numerous applications of mixed models in designed experiments. One of the simplest applications is a two-way mixed model, with a continuous response variable and one random effect and one fixed effect.

Several multiple comparison tests are available and care is required in choosing the appropriate one. But multiple-comparison methods give more detailed information about the

differences among the means and enable to control error rates for a multitude of comparisons. Bower (1998) described multiple comparison applied to experiments with CRD and RCBD. The tests involve calculation of a range value, or values, via the error mean square, in which the difference between a pair of means must exceed for significance.

Gomez and Gomez (1984) also described that pair comparison is the simplest and most commonly used comparison (planned and unplanned pair) in agricultural research. The two most commonly used test procedures for pair comparisons in agricultural research are the least significant difference (LSD) test which is suited for a planned pair comparison, and Duncan's multiple range test (DMRT) which is applicable to an unplanned pair comparison. The procedure for applying the DMRT is similar to that for the LSD test; DMRT involves the computation of numerical boundaries that allow for the classification of the difference between any two treatments mean as significant or non-significant. However, unlike the LSD test in which only a single value is required for any pair comparison at a prescribed level of significance, the DMRT requires computation of a series of values, each corresponding to a specific set of pair comparisons.

CHAPTER THREE

3. DATA AND METHODOLOGY

3.1. Data

The data were yield obtained from preliminary bread wheat, durum wheat, maize, and potato trials. The trials were conducted by different National and Regional Agriculture Research Centers using RCBD, lattice design, alpha lattice design and split-plot from 2004/05 - 2008/09 years (seasons). Bread wheat variety trial data is collected using RCBD at six research centers (locations) in 2005/06 and 2006/07 seasons was given the name BW99RVII data set. The bread wheat variety trials were arranged in a RCBD with four replications (blocks) in each location of the two seasons. In each replication 20 entries were randomly assigned. Durum wheat variety trials data set with the name DW99RVTR was collected at four locations in 2004/05, 2005/06 and 2006/07 seasons. The durum wheat variety trial in DW99RVTR data set was arranged in a RCBD with four replications in each season of the four locations. The experiments on durum wheat crop was laid out with 19 entries, 4 replications (blocks) consisting of 19 entries in each replication (block).

The maize variety trial was conducted using 5×5 triple lattice design (i.e., a partially balanced lattice design with three replications) at Melkasa, Dhera, Ziway, Mieso and Kobo National Agricultural Research Centers in 2008/09. The experiment on maize crop was laid out with 3 replications, 25 entries, 5 blocks and 5 plots per block. Another two maize variety trials, 2A3099AW and 2A3499AW, were conducted using alpha lattice design at Awasa in the year 2006/07. 2A3099AW was laid out with 2 replications, 100 entries, 10 blocks and 10 plots per block, while 2A3499AW was laid out with 2 replications, 81 entries, 9 blocks and 9 plots per block. In the year 2006, split-plot experiment used for integrated disease management (IDM) using main plot as fungicide and sub-plot as six potato varieties/clones. Three replications were used with a plot size of 12.15 meter squares (m²) at Jeldu, Degem and Wolmera locations. SAS software (2008) was used for statistical analysis (proc glm and proc mixed).

3.2. Methodology

3.2.1 Introduction

In a designed experiment, treatments are likely to be administered on experimental units under the same condition. However, a difference among experimental units is inevitable to occur and this is called an experimental error (residual). This error is primarily the basis for deciding whether an observed difference is real or just due to chance. In other words responses from each treatment are obtained from different units called replications and they are essential for the estimation of experimental error. Replications also help to improve precision of an experiment by reducing the standard error of a mean or of a difference between means. Replication is the repetition of experiments under identical conditions; but in the context of experimental designs, it refers to the number of distinct experimental units under the same treatment (Akindele, 2008). Replication, with randomization, will provide a basis for estimating the error variance. In the absence of randomization, any amount of replication may not lead to a true estimate of error. The greater the number of replications, the greater is the precision in the experiment.

The control of experimental error is another aspect of experiment that needs attention. Intuitively one can anticipate an increase in treatment difference produced if there is a sizeable reduction in experimental error. To realize this, one way of controlling error is by blocking or putting together similar experimental units in the same group and randomly assigning all treatments into each block separately and independently. The purpose of randomization is to prevent systematic and personal biases from being introduced into the experiment by the experimenter (Parsad and Gupta, 2000).

The main technique adopted for the analysis and interpretation of data collected from an experiment is the analysis of variance technique that essentially consists of partitioning the total variation in an experiment into components ascribable to different sources of variation due to the controlled factors and error. Analysis of variance (ANOVA) clearly indicates a difference among the treatment means.

A design for agricultural trials must provide valid error terms and sufficient precision for the effects of interest. As Drane (1989) states the manner in which the experiment is designed and executed determines what constitutes the experimental unit, the proper error terms in the analysis of variance, and whether replication is either possible or desirable. Thus it is usual in evaluating a design to identify its experimental units which summarizes the randomization that is to be employed in the experiment.

The designed experiments in this study are analyzed by analysis of variance (ANOVA) with the following two purposes.

1. To provide a subdivision of a total variation between the experimental units into separate components, each component representing a different sources of variation, so that the relative importance of the different sources can be assessed.
2. To give estimate of the underlying variation between units that provides the basis for inference about the effects of treatments.

This is a measure of experimental error which provides the basis for interval estimates and significance tests. The variance or, more correctly, the mean square associated with each of the other sources of variation may be compared with the experimental error mean square. This comparison provides an F statistic for testing the significance of the difference among means for the particular variance source. In addition, the analysis of variance provides information from which standard errors of means and differences may be computed, and from which interval estimates may be constructed. A fundamental assumption of ANOVA is that there exists a scale on which the various effects are additive (Kerr and Churchill, 2001).

Another method used in this study to compare the performance of one design over the other design is relative efficiency. Efficiency is measured by the variances of the estimated treatment differences which depend on the design and the within-block variation, and it is estimated by the residual mean square. The efficiency of one design of experiment analysis over another is usually measured in terms of reduced error variance, expected error mean square, or average standard error of the difference between treatments means (Edme, *et al.* 2007).

In this study, the efficiency of designs are compared in different research trials conducted at National and Regional Agricultural Research Centers to assess the efficiency of each in minimizing experimental error, coefficient of variation (CV) and error mean square for yield. The coefficient of variation (CV) affects the degree of precision with which the treatments are compared and is a good index of the reliability of the experiment. It is an expression of the overall experimental error as percentage of the overall mean; thus, the higher the CV value, the lower is the reliability of the experiment.

3.2.2 The common experimental designs and treatment structures

The common experimental designs and treatment structures being used in Ethiopian agricultural research system are randomized complete block design, lattice design, alpha-lattice design, factorial experiments, and split-plot treatment structure (Girma, 2002 and Mandefro, 2005). Due to lack of data on it, these factorial experiments were not considered in this study.

3.2.2.1 Randomized Complete Block Designs

The most basic type of statistical design for making inferences about treatment means is the completely randomized design (CRD), where all treatments under investigation are randomly allocated to the experimental units. The CRD is appropriate for testing the equality of treatment effects when the experimental units are relative by homogeneous with respect to the response variable. When the experimental units are heterogeneous, the notion of blocking is used to control the extraneous sources of variability.

Randomized complete block design (RCBD) is one of the most widely used experimental designs in agricultural research. RCBD is the simplest and most commonly used block design when the treatments are the several levels of a factor and also it is the most efficient design because there is no loss of information in estimating treatment contrasts as well as block contrasts (Parsad and Gupta, 2000). A randomized complete block design is a restricted randomization design in which the experimental units are first sorted into homogeneous groups, called blocks, and the treatments are then assigned at random within blocks. The major reason for grouping plots (experimental units) into uniform blocks is to reduce plot-to-plot variation and to improve the precision of the experiment. Failure to adequately block a field can result in unacceptably large error variance and/or biased estimates of treatment effects (Eskridge, 2008).

Blocking is an effective method for improving the efficiency of an experiment by eliminating systematic variations due to non homogeneities of experimental units. The primary purpose of blocking is to reduce experimental error by eliminating the contribution of known sources of variation among experimental units. This is done by grouping the experimental units into blocks such that variability within block is minimized and variability among blocks is maximized. At all stages during the experiment, the techniques applied within a block should be as uniform as possible, thus keeping experimental error within blocks as small as possible. Differences between blocks are permitted to be large, but are not of major concern in the analysis since the comparisons of treatments and the computation of experimental error is

done within blocks. Blocking will be effective only if the error variance among units within blocks is smaller than the error variance over all units. After the experimental units have been blocked, treatments are then randomly assigned to the units within the blocks. The randomization process for a RCBD is applied separately and independently to each of the blocks. A separate randomization is used in each block and every treatment appears in every block precisely once. In complete block design, each block consists of one complete replication of the set of treatment.

Blocking stratifies experimental units into homogenous groups, and in field trials, plots are normally blocked according to their proximity to each other. Blocking will increase treatment precision only if plots are blocked according to one or more varying external factors. If an experimental area is homogenous, blocking may actually decrease the precision of estimating treatment effects. This results from a larger means square error (MSE) term in the ANOVA since error degrees of freedom are reduced without a comparable reduction in error sum of squares (SSE). In this situation, a completely randomized design (CRD) would more precisely estimate treatment effects than a RCBD (Karcher *et al*, 2003).

The statistical model for RCBD is

$$Y_{ij} = \mu + \alpha_i + \beta_j + \varepsilon_{ij}, \quad i = 1, 2, \dots, a; \quad j = 1, 2, \dots, b \quad [1]$$

where Y_{ij} is the i^{th} observation in the j^{th} block, μ is an overall mean, α_i is the effect of the i^{th} treatment, β_j is the effect of j^{th} block, and ε_{ij} is a random error component.

Assumptions:

- μ and α_i are fixed parameters
- β_j may be fixed or random effects.
- ε_{ij} is distributed normally and independently with mean 0 and σ^2 ; that is, ε_{ij} is distributed iid $N(0, \sigma^2)$.

The RCBD analysis of variance with t number of treatments, r replications and b number of blocks is given as follow:

Table 1: Analysis of Variance for a RCBD

Source of variation	Degree of Freedom (Df)	Sum of Square(SS)	Mean Square (MS) *
Treatment	t-1	SS _{Treatment}	MS _{Treatment}
Block	b-1	SS _{Block}	MS _{Block}
Error	(t-1)(b-1)	SS _{Error}	MS _{Error}
Total	tb-1	SS _{Total}	

*The MS for each source of variation is obtained by dividing each SS by its corresponding Df.

In this study, treatment effects are assumed to be fixed and block effects are assumed to be fixed or random accordingly. The interaction between treatment and block were found to be insignificant. Hence, in our analysis we haven't included the interaction effect.

A measure of relative efficiency was computed by comparing the error mean square for the RCBD relative to that of a CRD. To calculate the relative efficiency of the RCB design, the RCB error mean square was divided by the CRD error mean square. The error mean square was estimated for a CRD following the model

$$MSE_{CRD} = \frac{(b-1)MS_{block} + r(b-1)MSE_{RCBD}}{bt-1} \quad [2]$$

where MS_{block} and MSE_{RCBD} are the block and error mean squares of RCBD and t, b , and r are the number of treatments, blocks, and replications, respectively. Karcher *et al* (2003) reported the relative efficiency value less than 1 indicates that a CRD is a more efficient design, while values close to 1 suggest that the two designs yield similar results. Values greater than 1 suggest that the RCBD is more efficient design than CRD. Each estimated relative efficiency value represents the multiple of replications needed in a CRD to achieve the same level precision in estimating treatment effects as in the RCBD.

3.2.2.2 Missing observations in RCBD

When using the RCBD, sometimes an observation in one of the blocks might be missing. This may happen because of carelessness or error or for reasons beyond our control such as unavoidable damage to experimental units by rodent, water losing, etc. A missing observation introduces a new problem into the analysis since treatments are no longer orthogonal to blocks; that is, every treatment does not occur in every block. In this case we consider two methods of analysis when some of the data are missing in RCBD. The first method is estimating the missing value and the analysis of variance is then carried out as usual except that one degree of freedom is subtracted from the total sum square and the error sum square. With a single missing unit, the first step is to calculate a value for the unit as

$$\hat{y}_{ij} = \frac{ry_i + ty_j - y_{..}}{(r-1)(t-1)} \quad [3]$$

where $\hat{y}_{ij}$ is estimate of the missing observation, y_i is the sum of the remaining values in the block with the missing observation, y_j is the sum of the remaining observations on the treatment with the missing value, and $y_{..}$ is the grand total of all available observations; r and t are the numbers of replicates and treatments respectively.

When there are several missing values, for units $a, b, c, d, \dots$, we first assign initial values for $b, c, d, \dots$, and use formula [3] to find an approximation for a . The most commonly used

initial value for each missing observation is the average of its marginal means (Gomez and Gomez, 1984):

$$\bar{y}_{ij} = \frac{\bar{t}_i + \bar{b}_j}{2} \quad [4]$$

where $\bar{y}_{ij}$ is the initial value of the i^{th} treatment and the j^{th} replication, $\bar{t}_i$ is the mean of all observed values of i^{th} treatment, and $\bar{b}_j$ is the mean of all observed values of the j^{th} replication. With this approximated value and the values previously assumed for $c, d, \dots$, we again use formula to insert an approximation for b . After a complete cycle of these operations, a second approximation is found for a and so on until the new approximations are not materially different from those found previously. The analysis of variance is then completed for each missing unit; m degree of freedom is subtracted from the total and error sum of squares, where m is the total number of missing values. The standard error of the mean difference ($s_{\bar{d}}$) between non-missing varieties, missing and non-missing varieties, and missing varieties estimated as (Gomez and Gomez, 1984):

$$s_{\bar{d}} = \sqrt{\frac{2s^2}{r}}, \quad \text{between non-missing varieties} \quad [5]$$

$$s_{\bar{d}} = \sqrt{s^2 \left[\frac{2}{r} + \frac{t}{r(r-1)(t-1)} \right]}, \quad \text{between missing and non-missing varieties} \quad [6]$$

$$s_{\bar{d}} = \sqrt{s^2 \left(\frac{1}{r'_A} + \frac{1}{r'_B} \right)}, \quad \text{between missing varieties} \quad [7]$$

where s^2 is the error mean square, r is the number of replications, t is the number of treatments, r'_A and r'_B are the effect numbers of replications assigned to treatments A and B. For a treatment to be compared, say treatment A, a replication is counted as 1 if it has data for both treatments, as $\frac{1}{2}$ if it has data for treatment A but not treatment B, and as 0 if it does not have data for treatment A. The second method is to assume incomplete block design and using incomplete blocks analysis. This is done by using unequal observations in each block.

3.2.2.3 Lattice Designs

Lattice designs are an important class of incomplete block designs, originally introduced by Yates (1936). Lattice designs are the most commonly used in agricultural research. Lattice designs are frequently used in situations where there are a large number of treatment combinations. Lattice designs facilitate the comparison of a large number of treatments which are assigned to incomplete blocks within replications. In lattice designs the number of

treatments must be an exact square. The number of units in each block is the square root of the number treatments. Lattice designs used for single-factor experiments having a large number of treatments and reasonably small block size are maintained to ensure that each block does not lose its homogeneity due to the large size. And also each block does not contain all treatments. The two most commonly used lattice designs are:

- i. Balanced lattice
- ii. Partially balanced lattice

3.2.2.3.1 Balanced Lattice

In balanced lattice design, the number of treatments must be a perfect square and the block size is equal to the square root of the number of treatments. The number of replications in balanced lattice design is one more than the block size. Incomplete blocks are combined in groups to form separate replication. The special feature of the balanced lattice, as distinguished from other lattices, is that every pair of treatments occurs once in the same incomplete block. Consequently, all pairs of treatments are compared with the same degree of freedom. However, if there are b blocks, there must be $b+1$ replications to achieve balance. This restriction in the number of replications and treatments makes the design less practical in terms of resources.

3.2.2.3.2 Partially Balanced Lattice

The number of replications required for balanced lattice becomes very large as the number of treatments increases. For this reason it is not usually practical to use balanced lattices for blocks with more than about seven units per block. In the interest of economy, then, the scientist is forced to accept a partially balanced design with fewer replications than would be required for full balance.

In partially balanced lattice designs, the number of replications is not restricted, but the number of treatment must be a perfect square and the block size is equal to the square root of the number of treatments. However, not all treatments occur together in the same block. This leads to differences in precision with which some comparisons are made relative to other comparisons. The names of the sub categories of partially balanced lattice design follow the number of replications. For example, the balanced lattice with two replications is called simple lattice, with three replication triple lattice, with four replications quadruple lattice and so on. The pattern of statistical analysis is the same for simple, triple, and quadruple lattices. The partially balanced lattice is more efficient than RCBD in general, but is inflexible in the number of entries it can handle and has more restrictions on field layout; these aspects

prevent it from being used more widely (Yau, 1997). Lattice design analysis of variance with r replications, k block size and $t = k^2$ number of treatments is given in Table 2.

Table 2: Analysis of Variance for a Lattice Design

Source of variation	Degree of Freedom (Df)	Sum of Square(SS)	Mean Square (MS)
Replication	$r-1$	Replication SS	Replication MS
Block (adj.)	$r(k-1)$	Block (adj.) SS	Block (adj.) MS
Treatment (unadj.)	k^2-1	Treatment (unadj.) SS	Treatment (unadj.) MS
Intrablock error	$(k-1)(rk-k-1)$	Intrablock error SS	Intrablock error MS
Treatment (adj.)	k^2-1	Treatment (adj.) SS	Treatment (adj.) MS
Total	rk^2-1	Total SS	

*The MS for each source of variation is obtained by dividing each SS by its corresponding Df.

The intrablock error SS is estimated as:
 Intrablock error SS = Total SS – Replication SS – Treatment (unadj.) SS – Block (adj.)

The efficiency of the partially balanced lattice design relative to a comparable RCBD is computed as (Gomez and Gomez, 1984):

$$R.E.= \left[\frac{Block(adj.)SS + Intrablock\ error\ SS}{r(k-1) + (k-1)(rk-k-1)} \right] \left[\frac{100}{EMS} \right] \tag{8}$$

where SS= sum of square, EMS= Error mean square, r is number of replications and k is the block size.

3.2.2.4 Alpha-Lattice Designs

The basic feature of alpha-lattice is the same as other lattice designs except that the number of blocks within replication and the number of treatments within block are not the same. That means it bridges the gap between RCBD and lattice designs. That is it has an additional feature in that the number of treatments should not necessarily be a perfect square. The development of alpha-lattice designs removed the restrictions on the number of treatments to be considered and its relation with block size required for lattice designs. An alpha-lattice design analysis of variance with t number of treatments, b number of blocks within replication, and r number replications is given the following Table 3.

Table 3: Analysis of Variance for an Alpha-lattice Design

Source of variation	Degree of Freedom (Df)	Sum of Square(SS)	Mean Square (MS) *
Replication	r-1	Replication SS	Replication MS
Block(replication)	rb-r	Block(replication) SS	Block(replication) MS
Treatment (adj.)	t-1	Treatment (adj.) SS	Treatment (adj.) MS
Error	rt-rb-t + 1	Error SS	Error MS
Total	rt-1	Total SS	

*The MS for each source of variation is obtained by dividing each SS by its corresponding Df.

The relative efficiency (R.E.) of an alpha lattice design compared with a RCBD is estimated (Masood *et al.*, 2008) as:

$$R.E. = \frac{\text{Error Mean Square in RCBD}}{\text{Error Mean Square in Alpha lattice Design}} \times 100 \quad [9]$$

3.2.2.5 Split-plot treatment structure

In two-factor experiments where it is desired to favor one factor over the other or where practical situation dictates application of one factor on larger plots, the split-plot treatment structure is used. Split-plot treatment structure is an experimental design in which the level of one factor could be applied to larger plots while the levels of another factor are applied to smaller plots. The levels of a specified factor are assigned to larger plots at random that are grouped into blocks. The larger plots are divided into smaller plots within the larger plots. The larger plots are called main-plots or whole-plots, while the smaller units are called sub-plots.

In split-plot treatment structure with two factors, randomization is done at two levels independently, one for the levels of the factor to be assigned in the main-plots and the other one for the sub-plots. The levels of the main-plot factor are randomly assigned to the main-plots within the blocks using separate randomization for each block. The levels of the split-plot factor are randomly assigned to the sub-plots within the main-plots using a separate randomization in each main-plot.

The main-plot factor effects are estimated from the large units, while the sub-plot factor effects and the interaction of the sub-plot factor with the main-plot are estimated from the sub-plots. There are two experimental errors, one coming from the main-plot and the other one from the sub-plot, as they differ in size.

The analysis of variance for split-plot treatment structural involves two steps. The first step is main-plot analysis and the second step is the sub-plot analysis. In the analysis of variance the main plot factor has to be tested against the interaction main plot factor $\times$ replication (the main plot error), whereas the sub-plot factor is tested against the residual. Because the main plot factor is tested with less precision and with only a low number of degrees of freedom for the error term, usually only large differences become significant. Let A denote the main plot factor and B , the subplot factor with a as the number of main plot treatments, b as the number of sub-plot treatments, and r as the number of replications. The analysis of variance for split-plot treatment structure is given as follows:

Table 4: Analysis of Variance for Split-plot

Source of variation	Degree of Freedom (Df)	Sum of Square(SS)	Mean Square (MS) *
Replication	$r-1$	Replication SS	Replication MS
Main plot (A)	$a-1$	Main plot (A) SS	Main plot (A) MS
Error (main plot)	$(r-1)(a-1)$	Error (main plot) SS	Error (main plot) MS
Sub-plot (B)	$b-1$	Sub-plot (B) SS	Sub-plot (B) SS
$A \times B$	$(a-1)(b-1)$	$A \times B$ SS	$A \times B$ MS
Error (sub-plot)	$a(r-1)(b-1)$	Error (sub-plot) SS	Error (sub-plot) MS
Total	$rab-1$	Total SS	

*The MS for each source of variation is obtained by dividing each SS by its corresponding Df.

3.2.3 Fixed Effects Linear Model

Fixed effects are those factors whose levels (or values) are selected by a non random process or whose levels consist of the entire population of possible levels. That is, fixed effects arise when the levels of effect constitute the entire population in which researchers are interested. All levels of interest are in the data set. Thus, a fixed effects linear model is a model containing only fixed effects. The fixed effects models inferences are not about any hypothesized population of studies, but about the particular collection of studies that is observed.

The fixed effect linear model in matrix form is written as

$$y = X\beta + \varepsilon \quad [10]$$

where y denotes the vector of observed data, X is a known matrix associated with the fixed effects, β is the unknown fixed-effects parameter vector, and ε is the random error vector of independent and identically distributed with mean vector 0 and variance-covariance matrix $\sigma^2 I$ (I is the identity matrix). As a result the variance of y is

$$\text{Var}(y) = \sigma^2 I. \quad [11]$$

3.2.4 Linear Mixed Effects Model

Mixed models are models in which some factors are fixed effects and other factors are random effects. Random effects are those factors whose levels consist of a random sample of levels from a population of possible levels. Linear mixed effects model is an important method of statistical analysis, which can realize the fitting of linear mixed effects model for data that obey the normal distribution. This model can be applied to balanced and unbalanced designs.

A linear mixed effects model can be expressed in matrix notation as

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\boldsymbol{\gamma} + \boldsymbol{\varepsilon} \quad [12]$$

where, $\mathbf{y}$ represents an $n \times 1$ vector of observed data, $\boldsymbol{\beta}$ is an unknown $p \times 1$ vector of fixed effects parameters with a known $n \times p$ design matrix $\mathbf{X}$, $\mathbf{Z}$ is a known $n \times q$ design matrix, $\boldsymbol{\gamma}$ is an unknown $q \times 1$ vector of random-effects parameters that is normally distributed with mean $\mathbf{0}$ and variance-covariance matrix $\mathbf{G}$, and $\boldsymbol{\varepsilon}$ is an $n \times 1$ random error vector which normally distributed with mean $\mathbf{0}$ and variance-covariance matrix $\mathbf{R}$. The name mixed model comes from the fact that the model contains both fixed-effects vector of parameters $\boldsymbol{\beta}$, and random-effects vector of parameters $\boldsymbol{\gamma}$. It is generally assumed that $\boldsymbol{\gamma}$ and $\boldsymbol{\varepsilon}$ have independent multivariate normal distributions with mean vector $\mathbf{0}$ and covariance matrices $\mathbf{G}$ and $\mathbf{R}$.

Therefore, the variance of $\mathbf{y}$ in the mixed effects model is

$$\mathbf{V} = \mathbf{ZGZ}' + \mathbf{R}$$

3.2.4.1 Estimation Methods in Linear Mixed effects Model

The common estimation methods in linear mixed model are Analysis of Variance (ANOVA), Maximum Likelihood (ML) and Restricted Maximum Likelihood (REML). The variance components in the linear mixed effects model may be estimated by ML or REML. ML estimates do not take into account the estimation of fixed effects and so are biased downward. REML estimation corrects for the presence of these nuisance parameters by maximizing the likelihood of linearly independent error contrasts to obtain unbiased estimates (Morrell, 1998).

3.2.4.1.1 Analysis of Variance

ANOVA is a technique for analyzing experimental data in which one or more response variables are measured under various conditions identified by one or more classification variables. The analysis of variance has proved useful in the statistical analysis of experiments

in the estimation of components of variation, in the estimation of the variance of estimates of treatment comparisons, and in making tests of significance.

The assumptions required for ANOVA are

- ❖ Normally distributed data (i.e., experimental errors are normally distributed).
- ❖ Independence of errors.

3.2.4.1.2 Maximum Likelihood

Maximum likelihood is an estimation procedure of finding the value of the vector of parameters which maximize the likelihood function for a given data set. That is, the method is based on maximizing the likelihood function.

The likelihood function is

$$L = (2\pi)^{-1/2N} |\mathbf{V}|^{-1/2} \exp\left[-\frac{1}{2}(\mathbf{y} - \mathbf{X}\boldsymbol{\beta})'\mathbf{V}^{-1}(\mathbf{y} - \mathbf{X}\boldsymbol{\beta})\right] \quad [13]$$

Maximizing L with respect to elements of $\boldsymbol{\beta}$ and the variance components (the σ_i^2 s that occur in $\mathbf{V}$) leads to equations that have to be solved to yield ML estimators of $\boldsymbol{\beta}$ and of σ^2 . Maximizing L can be achieved by maximizing the logarithm of L of [13], which shall be denoted by:

$$l = \log L = -\frac{1}{2}N \log 2\pi - \frac{1}{2}|\mathbf{V}| - \frac{1}{2}(\mathbf{y} - \mathbf{X}\boldsymbol{\beta})'\mathbf{V}^{-1}(\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) \quad [14]$$

To maximize l with respect to $\boldsymbol{\beta}$ and σ_i^2 , differentiate [14] with respect to $\boldsymbol{\beta}$ and σ_i^2 then equate to zero and solve the resulting equations for $\boldsymbol{\beta}$ and the σ_i^2 s.

These equations can be written in a variety of ways and can be solved iteratively. Despite the numerical difficulties involved in solving these equations for obtaining ML estimators of variance components, it is preferred over ANOVA method.

3.2.4.1.3 Restricted Maximum Likelihood

REML is a popular method of estimation for variance parameters in linear mixed models (Knight, 2008). The REML is a procedure applies the ML to the linearly transformed response data vector. The transformation is done in such a way that linearly transformed data vector contains none of the fixed effects. In REML:

- i) The variance components are estimated without being affected by fixed effects.
- ii) The degrees of freedom for fixed effects are taken into account implicitly, whereas with ML they are not.
- iii) REML does not fixed effects.

The REML estimation procedure does not, however, include estimating β . On the other hand the REML estimations with balanced data provide solutions that are identical to ANOVA estimators which are unbiased and have attractive minimum variance properties. In this sense REML is said to take account of the degrees of freedom involved in estimating the fixed effects, whereas ML estimators do not. REML is the preferred and most common method of estimation for variance parameters in linear mixed models as it corrects for the bias in ML estimation by accounting for the loss of degrees of freedom attributable to estimation of the fixed effects (Knight, 2008).

Generally, the fixed and mixed ANOVA models have identical sums of squares, but they differ in the expected mean squares and the choice of the appropriate statistical test. For the fixed model, the error mean square is an appropriate term for testing hypotheses about any source of variation, but for the mixed model the choice of a suitable error term is more difficult particularly when main effects are tested. Random effects ANOVAs are distinguished from fixed effects ANOVAs by which error mean squares are used as the denominator for F-tests. Under a fixed effects model, there is only one true error term in the model, and the corresponding mean square is used as the denominator for all tests.

3.2.5 Analysis of covariance

It is well known that in designed experiments the ability to detect existing differences among treatments increases as the size of experimental error decreases. A good experiment attempts to incorporate all possible means of minimizing the experimental error. Besides proper experimentation, a proper data analysis also helps in controlling experimental error. In situations where blocking alone may not be able to achieve adequate control of experimental error, proper choice of data analysis may help a great deal. By measuring one or more covariates, the characters whose functional relationships to the character of primary interest are known and the analysis of covariance can reduce the variability among experimental units by adjusting their values to a common value of the covariates. Experiments may be planned very well initially by taking all possible local controls into consideration. But, unforeseen circumstances may occur after the trial, and some experimental units may be affected by rodents, water lodging, pests, etc. Thus, experiments can be modified during data analysis without altering the initial design so that environmental influences are eliminated from the estimates of the treatment effects (Girma, 2002).

We have to recognize that an initial measurement x that could be made on each experimental unit might be strongly correlated with the yield variable y . The accuracy of the experiment

can thus be improved by adjusting the value of the y variable by the initial variable x ; often referred as a concomitant variable or covariate. Covariance analysis should be considered in experiments in which blocking could not adequately reduce experimental error (Mandefro, 2005). This is true because blocking is done before the start of the experiment; it can be used only to cope with sources of variation that are known or predictable. The identification of covariates is an important task in the application of covariance analysis. Assigning the covariates is highly determined by the purpose for which the covariance technique is applied. Most of time analysis of covariance is used in CRD, RCBD and split-plot experiment (Gomez and Gomez, 1984). In this study we considered analysis of covariance for RCBD.

The model for the analysis of covariance for two-way classified data with k treatments in r blocks of the randomized block design (Das and Giri, 1986) is

$$y_{ij} = \mu + t_i + b_j + bx_{ij} + e_{ij} \quad [15]$$

where y_{ij} is the observation corresponding to the i^{th} treatment in the j^{th} block; $i = 1, 2, \dots, k$, $j = 1, 2, \dots, r$; μ is the fixed effect of general mean; t_i is the fixed effect of i^{th} treatment; b_j is the fixed effect of j^{th} block; x_{ij} is the observation on the covariate corresponding to y_{ij} ; b is the regression coefficient of y on x ; e_{ij} 's are the error components which are assumed to be independently and identically distributed with zero mean and a constant variance σ^2 .

Covariance analysis is essentially an extension of the analysis of variance and hence, all the assumptions for a valid analysis of variance apply here as well. In addition, covariance analysis requires that:

- ✓ The relationship between the primary character of interest y and the covariate x is linear
- ✓ The strength of the relation between y and x remains the same in each experimental group
- ✓ The treatment and covariate are independent.

3.2.6 Testing of hypotheses

The testing of hypotheses concerning the fixed and random effects in a general balanced mixed model is fairly simple. Under the assumptions of normality the sums of squares in the analysis of variance table are independently distributed as scalar multiples of either central or non-central chi-squared varieties. On this basis, tests of hypotheses are usually carried out by using appropriate ratios of mean squares as F-statistics. The hypothesis that the treatments have equal effects is tested by F-test where F is the ratio of mean square treatment to mean square error (MST/MSE).

In testing of hypothesis we consider treatment effects and treatment variations. In fixed effects, our interest is to test whether all the treatments have the same effects or not.

The null hypothesis is, $H_0: \tau = \tau_o$

Alternative hypothesis is, $H_a: \tau \neq \tau_o$.

The second one is, in random effects our interest is to test whether there are treatment variations or not. That is,

$$H_0: \sigma_t^2 = 0$$

$$H_a: \sigma_t^2 \neq 0.$$

3.2.7 Model Diagnostics

The interpretation of data based on analysis of variance of fixed and mixed effects models is valid only when the assumptions are satisfied. As a result, it is necessary to detect the deviations and apply the appropriate remedial measures.

3.2.7.1 The Normality Assumption

The assumption of normality implies that the distribution of the variable to be analyzed by ANOVA is normal in the population from which units are sampled. Shapiro-Wilk test and Kolmogorov-Smirnov test are the formal tests of normality. Since the Kolmogorov-Smirnov test is appropriate for only large data (sample), Shapiro-Wilk test is used in this study. The Shapiro-Wilk test has been adapted in SAS to test for normality in sample sizes up to 2000 (Hettasch, 2009). The null-hypothesis of the Shapiro-Wilk test is that the residuals are normally distributed, therefore p-values that are larger than 0.05 indicate that values are normally distributed at the 5% level of significance. If the test is significant, the assumption of normality is violated. In this case, transforming the data will frequently correct the problem.

Also, the simplest check for normality involves plotting the empirical quantiles of the residuals against the expected quantiles. This is known as the normal QQ-plot. Thus, QQ-plots are useful for diagnosing violations of the normality assumption. In this method, observed value and expected value are plotted on a graph. If the scatter plots deviate from a straight line, then the data are not normally distributed.

CHAPTER FOUR

4. STATISTICAL DATA ANALYSIS AND RESULTS

4.1 Testing for Normality assumption

The normality of data was tested using the Shapiro-Wilks test and by QQ-plots (*Appendix I and II*). Using the results in (*Appendix I*) Tables 18 to 21, we perform the test of normality to assess whether the Shapiro–Wilks test of normality distribution of crop variety was significant or not. We want to test the null hypothesis which states that the crop variety is normally distributed against the alternative that the crop variety is not normally distributed at 5% level of significance. If we were to reject the null hypothesis based on the significance of the Shapiro-Wilks statistic, we would conclude that crop variety is not normally distributed and we would transform the data. If we fail to reject the null hypothesis, we conclude that the assumption holds.

For RCBD, the durum wheat variety of all locations of 2005/06 and 2006/07 seasons, the computed values of the Shapiro-Wilks statistic are not significant ($p > 0.05$) and as a result the normality assumption of durum wheat variety is satisfied (Table 18). Also, the computed values of the Shapiro-Wilks statistic of durum wheat variety in all locations of DW99RVTR data set were not significant at 0.05 significance level. These showed that durum wheat variety was normally distributed (Table 19). Thus for RCBD the normality assumption in all data sets of durum wheat varieties were satisfied. The normality assumption of the analysis of covariance was satisfied ($p > 0.05$). Also the covariate and treatment effects were independent and the regression slope (b) is equal to -9.572.

From Table 20 we can see that the computed values of the Shapiro-Wilk statistic of maize variety in all locations of lattice design are not significant ($p > 0.05$). As a result the maize variety trial in the National Agricultural Research centers of Melkasa, Dhera, Ziway, Mieso and Kobo are normally distributed. In alpha-lattice design of the both data sets the normality assumption was satisfied (Table 20). In split-plot analysis of the potato varieties for Jeldu, Degem and Wolmer sites the normality assumptions were also satisfied (Table 21).

4.2 Analysis of the Linear Fixed Effects Model

In linear fixed effects model analysis both treatment and block (replication) were fixed and there is no interaction between treatments and blocks. The CRD and RCBD error variances (error mean squares), and lattice design and alpha lattice design error effective variance were calculated for each location and season of the data sets. To compare the performance of the designs, relative precision ratios of the CRD to RCBD, RCBD to lattice design and RCBD to

alpha lattice design were determined. These were calculated by dividing CRD error variance by the RCBD error variance and RCBD error variance by the alpha lattice effective error variance. The relative precision ratio of the lattice design to RCBD was calculated by Proc lattice using SAS software.

4.2.1 Randomized Complete Block Design Analysis

As can be seen in Table 5 below, the bread wheat variety and block effects are significant ($p < 0.05$) in all locations (location one to six). The relative efficiency of RCBD to CRD in location 1,2,3,4,5 and 6 were 1.78, 1.17, 1.26, 1.44, 1.17, and 1.53, respectively (Table 5).

Table 5: ANOVA of RCBD for Bread Wheat Variety in 2005/06 of BW99RVII data set

Location	Source	Df	SS	MS	F value	Pr>F	RE
1	Block	3	605.578	201.859	20.570	0.000*	1.78
	treatment	19	775.587	40.820	4.160	0.000*	
2	Block	3	287.227	95.742	4.560	0.006*	1.17
	treatment	19	1203.787	63.357	3.010	0.000*	
3	Block	3	374.249	91.416	6.790	0.001*	1.26
	treatment	19	813.212	42.801	3.180	0.000*	
4	Block	3	740.574	246.858	11.610	0.000*	1.44
	treatment	19	1020.225	53.696	2.530	0.004*	
5	Block	3	669.703	223.234	4.350	0.008*	1.17
	treatment	19	3182.410	167.495	3.270	0.000*	
6	Block	3	749.554	249.851	13.980	0.000*	1.53
	treatment	19	1465.307	77.121	4.320	0.000*	

*: significant ($p < 0.05$) df: Degree of freedom SS= sum of square MS= mean of square RE= relative efficiency

The results in Table 5 indicate that the use of the RCBD instead of a CRD in the six locations of BW99RVII data set increased experimental precision by 78, 17, 26, 44, 17, and 53 percent, respectively. This means, in all of the six locations we would have required approximately about two times as many replicates with CRD to obtain the same efficiency as is obtained by RCBD.

From Table 6 we can see that bread wheat variety effects were significant in locations 1, 3, 4, and 5 and insignificant in locations 2 and 6. The block effects was significant ($p < 0.05$) in locations 1 and 5. In locations 1 and 5 the relative efficiency of RCBD to CRD is 1.53 and 1.12, respectively. This shows that the use of RCBD increased experimental precision by 53 and 12 percent (Table 6). Blocking effect of bread wheat variety in locations 2, 3, 4, and 6 is

insignificant ($p>0.05$) and the relative efficiency of these locations is nearly one. This indicates that, except for the first and the fifth locations, the relative efficiency for the rest of the locations are nearest to one the RCBD is almost as efficient as the CRD. Thus, for those locations blocking seems to be insignificant and unnecessary.

Table 6: ANOVA of RCBD for Bread Wheat Variety in 2006/07 of BW99RVII data set

Location	Source	Df	SS	MS	F value	Pr>F	RE
1	block	3	1033.102	344.367	15.010	0.000*	1.530
	treatment	19	1986.950	104.576	4.560	0.000*	
2	block	3	114.227	38.076	1.750	0.168	1.030
	treatment	19	344.099	18.110	0.831	0.663	
3	block	3	16.835	5.611	0.410	0.745	0.980
	treatment	19	1314.447	69.181	5.080	0.000*	
4	block	3	63.819	21.273	0.500	0.683	0.980
	treatment	19	2089.942	109.997	2.590	0.003*	
5	Block	3	275.736	91.912	4.180	0.010*	1.120
	treatment	19	2438.329	128.333	5.840	0.000*	
6	Block	3	49.522	16.507	0.840	0.480	0.990
	treatment	19	622.321	32.754	1.660	0.073	

*: significant ($p<0.05$ df =Degree of freedom SS= sum of square MS= mean of square RE= relative efficiency

From Table 7 we can see that durum wheat variety effects of location one were significant ($p<0.05$) in the three seasons (years). And also the block effects of the three were significant at 0.05 of significance level. So, the relative efficiency of RCBD compared to CRD in 2004/05, 2005/06, and 2006/07 are 1.342, 1.072, and 1.123 respectively (Table 7). This shows that the using of RCBD increased experimental precision by 34.2, 7.2, and 12.3 percent in the three seasons, respectively.

Table 7: ANOVA of RCBD for Durum Wheat Variety of DW99RVTR data set

Location One							
Year(season)	Source	Df	SS	MS	F value	Pr>F	R.E
2004/05	Block	3	634.144	211.381	9.560	0.000*	1.342
	treatment	18	842.036	46.780	2.120	0.018*	
2005/06	Block	3	226.968	75.656	2.800	0.049*	1.072
	treatment	18	926.075	51.449	1.900	0.036*	
2006/07	Block	3	149.751	49.917	4.07	0.0112*	1.123

	treatment	18	440.914	24.495	2.000	0.026*	
Location Two							
2004/05	Block	3	140.224	46.741	1.780	0.162	1.031
	treatment	18	526.215	29.234	1.110	0.366	
2005/06	Block	3	226.333	75.444	8.310	0.000*	1.293
	treatment	18	312.742	17.375	1.910	0.034*	
2006/07	Block	3	230.319	76.773	3.410	0.024*	1.097
	treatment	18	361.575	20.088	0.890	0.589	
Location Three							
2004/05	Block	3	131.335	43778	1.240	0.306	1.009
	treatment	18	1259.571	69.976	1.980	0.028 *	
2005/06	Block	3	41.526	13.842	0.620	0.607	0.985
	treatment	18	1169.185	64.955	2.900	0.001 *	
2006/07	Block	3	51.187	17.062	0.660	0.582	0.986
	treatment	18	573.175	31.843	1.230	0.275	
Location Four							
2004/05	Block	3	271.636	90.546	6.400	0.001*	1.216
	treatment	18	396.479	22.027	1.560	0.106	
2005/06	Block	3	85.331	28.444	2.370	0.080	1.055
	treatment	18	395.08	21.949	1.830	0.045*	
2006/07	Block	3	361.773	120.591	6.500	0.005*	1.220
	treatment	18	839.438	46.635	2.510	0.005*	

*: significant ($p < 0.05$) df= Degree of freedom SS= sum of square MS= mean of square RE= relative efficiency

The results in Table 7 also indicate that there were significant differences among durum wheat variety effects in the majority of seasons except location 2 (2004/05 and 2006/07), location 3 season 2006/07 and location 4 season 2005/04. In location 1, location 2 (2005/6, 2006/07) and location 4 (2004/05, 2006/07) blocking was significant. RCBD is efficient as CRD for 2004/05 season and less efficient for 2005/06 and 2006/07 seasons of location three. On average RCBD was more efficient than CRD for durum wheat variety of DW99RVTR data set.

4.2.2 Analysis of RCBD with missing observations

In durum wheat variety trial of 2004/05 year (season) observations in different blocks of RCBD were missing. The analysis was performed by estimating and replacing the missing values and considering the missing values as incomplete block.

From durum wheat variety of DW99RVTR data set, unfortunately the observation in block (replication) one of treatment nine was missing due to measuring error. By using formula (Equation [3]) the missing value $x= 39.405$ was estimated. Then using of this value as a replacement for the missing value, the analysis of variance was performed and the result given in Table 8.

Table 8: ANOVA of Durum Wheat for location one year 2004/05 of DW99RVTR data set with one missing value

Analysis of variance with estimated missing value					
Source	Df	SS	MS	Sign.	CV
Total	74	2686.284			13.523
Block	3	654.994	218.331	0.000	
Treatment	18	841.773	46.765	0.017	
Error	53	1189.517	22.444		
Analysis of variance for Incomplete Block					
Total	74	2667.409			13.523
Block	3	633.035	211.012	0.000	
Treatment	18	841.703	46.761	0.020	
Error	53	1189.517	22.444		

From Table 8 we can see that the error mean square (EMS) and coefficient of variation (CV) of estimated missing value and the incomplete block analysis were the same. Also in both methods the block and treatment effects were significant. This shows that a single missing value may not affect the results of analysis considerably.

For more than one missing observations, the iterative procedure is used to estimate the missing values. Using this iterative procedure the missing values of RCBD with two, three, four, five, six and eight observations were estimated (Table 22, *Appendix I*). Then the estimated values were replaced and the analysis was preformed as usual except that the numbers of missing observations were deducted from error and total sum squares of degree of freedom.

Table 9 below shows that the treatment and block effects of RCBD with two missing observations replaced by their estimates were significant at 5% of the significance level for missing of different treatments from the same block, different treatments from different blocks and the same treatments from different blocks. Also, for the analysis of the missing values as incomplete block design (IBD), the treatment and block effects were significant. The error mean square (EMS) of RCBD with two missing observations which have been replaced by estimates were equal to the EMS of incomplete block design (Table 9), once again showing no considerable effect. But the coefficient of variation (CV) of RCBD with estimated missing values were less than the CV of IBD. The position of missing values affect the standard error of mean difference, for example, standard error of the mean difference between missing varieties is largest when missing observations of the same treatments from different blocks and also, treatment mean square of analysis as incomplete block and with estimated missing observations are largest when missing of the same treatments from different blocks. The average standard error of the mean difference between unequal replication varieties which was analyzed by IBD was less than the standard error of the mean difference between missing varieties and larger than that between missing and non-missing varieties and missing varieties.

Table 9: Analysis of variance of RCBD with two missing observations of durum wheat variety

Analysis of variance with estimated missing values																
Source		Missing of different treatments from the same block					Missing of different treatments from different blocks					Missing of the same treatments from different blocks				
		df	MS	Sign.	CV	RE	df	MS	Sign.	CV	RE	df	MS	Sign.	CV	RE
Block		3	198.920	0.000	13.213	1.318	3	224.525	0.000	12.675	1.394	3	222.795	0.000	12.63	1.392
Treatment		18	47.028	0.013			18	49.667	0.005			18	51.370	0.004	1	
Error		52	22.199				52	20.710				52	20.613			
Standard error of the mean difference	b/n non-missing varieties	3.332					3.218					3.210				
	b/n missing & non-missing varieties	3.613					3.490					4.145				
	b/n missing varieties	3.847					4.070					4.542				
Analysis of variance for the missed values as IBD																
Block		3	186.345	0.000	13.500	1.296	3	217.197	0.000	12.932	1.380	3	220.819	0.000	12.915	1.389
Treatment		18	46.944	0.019			18	48.903	0.008			18	49.081	0.008		
Error		52	22.199				52	20.710				52	20.613			
Average Standard error of the mean difference b/n unequal replication									3.514							
b/n = between																

As we see from Table 10 below, the treatment effects and block effects were significant for three, four, five, six and eight missing values. The analysis was performed by replacing the estimated values for the missing observations. In the analysis of variance considering the missing values as IBD, except for five missing values the treatment effects were not significant. This shows that as the number of missing observation increases, the difference between estimated missing and considering the design as incomplete block is getting larger. Also as the number of missing observation increases the estimation method (of missing values) becomes imprecise leading to biased EMS. But since only available information is used in IBD the results are more acceptable. Naturally as more observations are missing the system might lose integrity and the test statistic becomes less stable, but in missing-estimated approach, as the number of missing observations increases the treatment effect level of significance gets smaller showing strong difference between varieties. On the other hand, in the incomplete block approach, as the number of missing observation increases the level of significance for variety difference was relatively larger, this contradicts the former result. Therefore, presence of missing values affects the efficiency of designs. The RCBD error mean square (EMS) and coefficient of variation (CV) for five, six and eight missing values which was analyzed by replacing the missing observations with estimated values were relatively greater than the error mean square (EMS) and coefficient of variation (CV) of analysis performed as IBD but not for three and four missing values (Table 10). This shows that as the number of missing observation increases the EMS and CV obtained from missing estimated approach is greater than the incomplete block approach. As the number of missing values increases, the relative efficiency of the analysis with estimated missing values relatively less than the relative efficiency of analysis as IBD. Also, the standard error of the mean difference between complete and missing, and between missing varieties of missing estimated approach was greater than that of incomplete block approach. This is shows design efficiency is affected by missing values. Therefore, when we compare these results with the true condition values (i.e., where there are no missing values) such as EMS (22.117), CV (13.416) and standard error of the mean difference (3.325), as the number of missing values increases in RCBD use of IBD approach is more efficient than that of missing-estimated approach (Table 10).

Table 10: Analysis of variance of RCBD with three, four, five, six and eight missing observations of durum wheat variety

Missing values	ANOVA with estimated missing values								ANOVA for the missed values as IBD							
	Sign.		EMS	CV	R.E	Standard error of the mean difference			Sign.		EMS	CV	R.E	S.E of the mean difference between		
	Block	Treatment				b/n non-missing varieties	b/n missing & non-missing varieties	b/n missing varieties	Block	Treatment				unequal replication of varieties (missing & non-missing varieties)	b/n missing varieties	
Three	0.000	0.027	21.890	13.035	1.120	3.308	3.587	4.185	0.000	0.063	21.890	13.406	1.101	3.573	3.820	
Four	0.000	0.025	21.958	12.950	1.120	3.313	3.593	4.191	0.000	0.067	21.958	13.443	1.110	3.579	3.826	
Five	0.000	0.064	25.625	13.797	0.891	3.579	3.882	4.133 ¹ 4.528	0.000	0.047	23.393	13.763	1.052	3.694	3.949	
Six	0.000	0.036	24.940	13.929	0.924	3.531	3.829	4.079 ¹ 4.467	0.001	0.053	23.844	13.578	1.036	3.729	3.987	
Eight	0.000	0.018	24.903	14.205	1.001	3.529	3.827	4.076 ¹ 4.463	0.001	0.078	24.834	13.136	1.043	3.806	4.069	

[1] Standard error of the mean difference between missing varieties in the same block

[2] Standard error of the mean difference between missing varieties (incomplete varieties)

4.2.3 Analysis of Lattice Design

The coefficient of variation (CV) and the error mean square (EMS) of lattice design calculated for maize variety trial in all research centers were comparatively low as compared to RCBD (Table 11).

Table 11: Efficiency of lattice design for Maize National Variety Trial of Melkasa, Dhera, Ziway, Mieso, and Kobo research centers in 2008/09 as compared to RCBD.

Location (Research centre)	Lattice design		RCBD		Efficiency relative to RCBD
	CV	EMS	CV	EMS	
Melkasa	9.235	19.153	10.391	24.247	112.140
Dhera	25.679	71.296	26.505	75.956	101.290
Ziway	35.659	48.508	39.345	59.056	109.060
Mieso	21.523	39.477	24.911	52.884	117.100
Kobo	22.596	38.405	23.028	39.887	100.500

The efficiency relative to RCBD from Table 11 indicates that the use of lattice design instead of RCBD increased experimental precision by 12.1, 1.3, 9.1, 17.1, and 0.5 percent in research centre Melkasa, Dhera, Ziway, Mieso, and Kobo respectively.

4.2.3 Analysis of Alpha lattice Design

Table 12: Efficiency of Alpha lattice design as compared to RCBD in 2005/06 Season

Design	Data Set					
	2A3099AW			2A3499AW		
	CV	EMS	P value for block	CV	EMS	P value for block
Alpha lattice design	22.636	182.729	0.018	15.638	137.924	0.011
RCBD	24.627	216.285	0.047	17.505	172.823	0.031
Relative Efficiency	1.019			1.024		

The significance of blocking within replication (group) in both designs for the two data sets indicates that blocking was effective in reducing experimental error (Table 12). The relative efficiency was 101.91% for 2A3099AW data set implying that the use of alpha lattice design increased experimental precision by 1.91% compared to RCBD (Table 13). The coefficient of variation (CV) and error mean square (EMS) of alpha lattice design (22.64 and 182.73) are

comparatively low as compared to RCBD (24.63 and 216.29). The relative efficiency was 102.22% implying that the use of alpha lattice design for 2A3499AW data set increased the precision of experiments by 2.44% compared to RCBD (Table 10). For 2A3499AW data set, the coefficient of variation (CV) and error mean square (EMS) of alpha lattice design (15.64 and 137.92) comparatively low as compared to RCBD (17.51 and 172.82). The value of relative efficiency greater than one for both data sets shows that alpha lattice design was relatively more efficient than RCBD.

4.2.4. Analysis of Covariance

In the analysis of covariance the DW99RVTR data set of durum wheat variety trial for location one which was conducted using RCBD in 2004/05 season was used. In this trial some of the durum wheat varieties on some of the plots are assumed to be affected by rodents or water lodging or any other factor. As a result, using covariance analysis as an adjustment for damaged varieties was required. So, we have to construct another variable known as a covariate which can be used as an adjustment factor. To simulate situation of damage to the experimental values of selected plots were purposely altered. The changes made to plots were made arbitrary but it was assumed that the damage takes place at different level. Therefore, when values were altered some of them were left as they are originally some to be highly changed, some slightly changed and others somewhere in between. In principle, a quick assessment of yield condition in the experimental field is necessary to coin covariates. The level of precision in adjustment depends on how well the covariate is created. If each value is taken into consideration so that several scales are adopted, the covariate will be able to account for the systematic variability taking place in the field. For this particular data, a four-scale rating of the yield condition was used. Scale number 1 is given to the plots not damaged at all, scale number 2 for plots with slightly damaged, scale number 3 for intermediate damage of plots and 4 for severely-damaged plots. Numeric values 1, 2, 3 and 4 were assigned to the scales, respectively.

As can be seen from Table 13 below, in the analysis of RCBD, the treatment effect and blocking of durum wheat variety with covariate was significant but in the analysis without covariate, both effects are non-significant. Use of covariate analysis dramatically reduced the error mean square (EMS) and coefficient of variation (CV). This is because the covariate effect is highly significant showing that the covariate created from the assessment of the field indeed has taken care of the pattern created by the damaged plots. Thus, the systematic variability which was left in the error has now been removed reducing the error sum of square considerably. The mean of adjustment for damaged yields by use of covariate was larger

than the mean of unadjusted (without covariate). Also, the standard error of mean yield with covariate (adjusted mean) and without covariate (unadjusted mean) was 2.511 and 4.108, respectively (Table 23, *Appendix I*). The relative efficiency of RCBD compared to CRD was 1.211 and 1.052 for analysis with covariate and without covariate respectively. This shows that using of covariate improved the precision of RCBD.

Table 13: Analysis of Covariance of RCBD on durum wheat variety trial

Analysis with covariate						Analysis without covariate					
source	df	MS	Sign.	CV	RE	source	df	MS	Sign.	CV	R.E
Block	3	155.994	0.001	15.154	1.211	Block	3	155.994	0.087	24.958	1.052
Treatment	18	78.959	0.001			Treatment	18	78.959	0.318		
Covariate	1	2326.784	0.000			Error	54	67.518			
Error	53	24.890									

4.2.5. Split-plot Experiment Analysis

Table 14 shows that there were no significant effects among fungicide (sprayed with fungicide and unsprayed) at the 0.05 significance level in Jeldu, Degem and Wolmer sites. We can also see from Table 14 that the Potato varieties/clones were highly significant ($p<0.05$) in all sites (Jeldu, Degem and Wolmer).

The error mean squares of potato (subplot) in the three sites were less than the error mean squares of fungicide (main plot). This indicated that the subplot (potato) variety was more efficient than main plot (fungicide).

Table 14: ANOVA of Split-plot for Integrated Disease Management (IDM) on potato variety of 2006 in Jeldu, Degem and Wolmer sites

Source	Jeldu				Degem				Wolmer			
	df	SS	MS	Sign.	df	SS	MS	Sign.	Df	SS	MS	Sign.
Replication	2	375.921	187.960	0.3718	2	107.844	53.922	0.546	2	68.310	34.155	0.898
Fungicide(main plot)	1	1.960	1.960	0.907	1	375.060	375.060	0.138	1	898.173	898.173	0.226
Error fungicide	2	222.457	111.229		2	129.495	64.747		2	602.947	301.473	
Potato(subplot)	5	15224.164	3044.833	0.000	5	15555.102	3111.020	0.000	5	16198.149	3239.630	0.000
Fungicide*potato	5	88.702	17.740	0.842	5	535.800	107.160	0.018	5	621.356	124.271	0.022
Error Potato	56	2437.216	43.522		56	2000.755	35.728		56	2421.331	43.238	

4.3. Linear mixed effects model analysis

In this section, the block effects are assumed to be random and treatment effects were considered to be fixed. Here some of the data sets that have been used in the fixed effects model were taken and analyzed as a linear mixed effects model. Particularly we have taken BW99RVII, Milkasa, 2A3099AW and 2A3499AW for RCBD, Lattice design and Alpha lattice design, respectively. The variance components from the mixed model are estimated using the method of restricted maximum likelihood (REML). Here, unlike the fixed effects model assessing the performance of experimental designs using relative efficiency is not possible. We have assessed the performance of experimental design under linear mixed effects model using their corresponding variance components.

4.3.1 Analysis of Randomized Complete Block Design

The variance components of treatment and residual summarized in Table 15 was estimated by REML method. The estimated covariance parameters of residual (Table 15) of bread wheat variety in 2005/06 season of all the locations for RCBD were smaller than the residual variance of CRD. Also, the residual variance of location one, two and five in 2006/07 season were less than that of CRD. This indicates that RCBD is more efficient than CRD in a situation where there is systematic variability. But in location three, four and six of 2006/07 season the residual variance of RCBD was the same as CRD implying that, for these three locations both RCBD and CRD are equally efficient. The variance component for block of 2006/07 season at locations three, four and six were approximated to zero. This shows that blocking was ineffective in these trials. Hence, in future experiments to be conducted in same locations, CRD may be used.

Table 15: Covariance Parameter estimates of RCBD for BW99RVII data set

Year(season)	Location	Covariance Parameter Estimates		
		Covariance Parameter	Estimate	
			CRD	RCBD
2005/06	One	block		9.6023
		Residual	19.4155	9.8132
	Two	block		3.7362
		Residual	24.7540	21.0177
	Three	block		3.8980
		Residual	17.3539	13.4559

	Four	block		11.2797
		Residual	32.5444	21.2647
	Five	block		8.5984
		Residual	59.8655	51.2672
	Six	block		11.5990
		Residual	29.4713	17.8724
2006/07	One	block		16.0713
		Residual	39.0119	22.9406
	Two	block		0.8138
		Residual	22.6134	21.7996
	Three	block		0
		Residual	13.216	13.216
	Four	block		0
		Residual	41.375	41.375
	Five	block		3.4974
		Residual	25.4619	21.9645
	Six	block		0
		Residual	19.6001	19.6001

4.3.2 Lattice design analysis

As it can be seen in Table 16 below, the covariance parameter estimates of block within replication and residual of maize variety trial in Melkasa Research Centre was 1.910 and 24.247; 6.113 and 19.153 for RCBD and lattice design, respectively. The residual covariance of maize variety of RCBD was greater than the lattice design residual covariance. The variance component of block for RCBD was less than the variance component of block for lattice design. This indicates there is some evidence that lattice design was more efficient than RCBD.

Table 16: Covariance Parameter estimates of lattice design for Melkasa Research Centre

Covariance Parameter Estimates		
Covariance Parameter	Estimate	
	RCBD	Lattice design
block(replication)	1.9104	6.113
Residual	24.247	19.153

4.3.3 Alpha lattice design analysis

Table 17 shows that the covariance parameter estimates of block within replication and residual of alpha lattice design in 2A3099AW and 2A3499AW data sets were 20.681 and 39.892 for blocks and 152.720 and 138.530 for residuals, respectively. The residual variance estimated for alpha lattice design was comparatively low as compared to RCBD for both data sets. The variance component of block for RCBD was less than the variance component of block for alpha lattice design for both of data sets. This shows that alpha lattice design is more efficient than RCBD for this data set.

Table 17: Covariance Parameter estimates of alpha lattice design for 2A3099AW and 2A3499AW data sets

Covariance Parameter Estimates				
Cov Parameter	2A3099AW data set		2A3499AW data set	
	Estimate		Estimate	
	RCBD	Alpha lattice	RCBD	Alpha lattice
block(replication)	8.6582	20.6814	8.2019	39.8921
Residual	170.90	152.72	172.82	138.53

4.4 Discussion

From the fixed effects model analysis of RCBD on bread wheat variety trial, in 2005/06 season the block effect was high and the relative efficiency of all the locations of this season were greater than one. The majority of block effects in 2006/07 of bread wheat variety trial were low in all locations. As a result, the estimated relative efficiency of RCBD compared with CRD was less than or near 1.0, except at locations one and five for this season. On average for the two seasons of bread wheat variety trial, the relative efficiency of RCBD compared to CRD is greater than one. This result agrees with that of Karcher *et al* (2003). Therefore, RCBD is more efficient than CRD for the bread wheat variety trial considered in this study. In the durum wheat variety trial, about 55% the block effects were not significant. For the three seasons of location one, the block effects were significant and the relative efficiency was greater than one. For this location-blocking has effectively increased the precision of the durum wheat variety experiment. Due to technical or cultural practices, the blocking effect was low (not significant) in certain trials and the technical problem results from wrong block orientation and direction which render blocking ineffective (Girma, 2005). We have also shown that in RCBD when the number of missing value increases the error mean square, standard error of the mean difference and the coefficient of variation of estimated missing approach were greater in the analysis performed as IBD. As a result using IBD is more efficient than the estimated missing approach of RCBD.

Results from the lattice design analysis indicate that the coefficient variation (CV) and the error mean square (EMS) of national maize variety trial was calculated for Melkasa, Dhera, Ziway, Mieso, and Kobo research centers. The coefficient of variation and error mean squares of the randomized complete block design was greater than that of lattice design. The lattice design efficiency relative to RCBD (Table 11) show that in all trials considered, lattice design was more efficient than RCBD.

The results of experiments in the 2A3099AW and 2A3499AW data sets show that there are smaller error mean squares (EMS) under alpha design than RCBD. The coefficient of variation (CV) of alpha lattice design is comparatively low as compared to RCBD. Low value of CV indicates a good index of reliability. The relative efficiency indicates how much more efficient the alpha lattice design is as compared to RCBD: if the value of relative efficiency is greater than one then the alpha lattice design results in smaller error variance. The value of relative efficiency greater than one for both the data sets of the experiments show that alpha lattice design was relatively more efficient than RCBD (Table 12). Alpha lattice design

provided better control of within replicate variation and also to experimental error than RCBD. This study showed that alpha lattice design is more efficient than RCBD. This was similar with the result of Patterson and Hunter (1983), Yau (1997), and Masood et al. (2008). The analysis of covariance for RCBD on durum wheat variety trial shows that the error mean squares and coefficient of variation of analysis with covariate, that was used to adjust the damaged yield, were smaller than the analysis of variance (without covariate). Also, the adjusted mean of yield by covariate was greater than the mean of without covariate while the standard error of adjusted mean yield by covariate was smaller than the standard error of unadjusted mean yield. Use of covariance analysis increased the relative efficiency of RCBD compared to CRD than the analysis of variance (without covariate) for damaged yield (Table 13).

From the linear mixed effects model analysis of RCBD on bread wheat variety trial, the residual variance of RCBD were less than the residual variance of CRD in 2005/06 season for the six locations. Also for locations one, two, and five the residual variances of RCBD were less than the residual variance of CRD in 2006/07 season. But for locations three, four, and six the residual variances of RCBD were the same as the residual variance of CRD (Table 15). This shows that RCBD is as efficient as the CRD in increasing the precision of the bread wheat variety experiments. The residual variance of lattice design was comparatively low as compared to RCBD, and the variation among block for RCBD was less than the variation among block within replication for lattice design in Melkasa research centre. This indicates that, the lattice design was more efficient than RCBD to improve the precision of maize variety experiments. The residual variances of alpha lattice design for 2A3099AW and 2A3499AW data sets were smaller than the residual variances of RCBD. And also, for the two data sets of the alpha lattice design, the variation among block within replication were greater than the variation among block of RCBD. This shows that alpha lattice design was more efficient than RCBD.

CHAPTER FIVE

5. CONCLUSIONS AND RECOMMENDATIONS

5.1. Conclusions

Our prime concern is identification of the more efficient experimental design for field experiments. We conclude the findings of this study as follows.

On average RCBD was more efficient than CRD for bread wheat and durum wheat variety trials. Due to technical or cultural practices, the blocking effect was low or insignificant in some locations of the bread wheat and durum wheat variety trials. It is therefore important to consider suitability of RCBD, but block orientation and direction should be carefully selected.

For missing values in RCBD, analysis as IBD is more efficient than missing-estimated approach. Also for damaged yield in RCBD, use of covariance analysis increases the precision of RCBD than analysis of variance without covariate.

For National maize variety trial, lattice design was more efficient than RCBD to increase the precision of the field experiments. Also, alpha lattice design was more efficient than RCBD to improve the precision of field experiments. Under land heterogeneity and in a situation where a large number of entries are tested, the chance that families of IBD are more appropriate is quite high and researches are encouraged to use one of these designs under such conditions.

In conclusion, RCBD, lattice design and alpha lattice design were the more efficient designs than the CRD for agricultural field experiments. Thus, researches must be cautious in using CRD in field experiments.

5.2. Recommendations

- To increase the precision of agricultural field experiments, RCBD, lattice design and alpha lattice design should be used based on available land resource and size of experiments.
- For a large number of missing values in RCBD, we recommend that the agronomist uses incomplete block approach instead of using missing-estimated approach.
- For some locations used in this study, block effects were insignificant. This can be due to technical issues such as orientation and direction of blocking. So, we recommend that further studies should be done in this area.

REFERENCES

- [1] Akindele, S.O. (2008): The Place of Biometrics in Forestry Research. Department of Forestry and Wood Technology, Federal University of Technology. Akure, Nigeria.
- [2] Alves, E. G., St. Aubyn, A. and Martins, A. (2009): Experimental designs for evaluation of genetic variability and selection of ancient grapevine varieties: a simulation study. *Heredity Journals*, 104(6):552-562.
- [3] Atkinson, A.C. and Bailey, R.A. (2001): One Hundred Years of the Design of Experiments on and off the Pages of "Biometrika". *Biometrika*, 88(1): 53-97.
- [4] Blau, G. and Han, S. (2007): Statistical Design of Experiments. <http://www.Pharmahub.org/>
- [5] Bower, J. A. (1998): Statistics for food science. Department of Applied Consumer Studies, Queen Margaret College, Edinburgh, UK.
- [6] Campbell, B.T. and Bauer, P.J. (2007): Improving the precision of cotton performance trials conducted on highly variable soils of the South-Eastern USA coastal plain. *Blackwell synergy-plant breeding*, **126**, 622-627.
- [7] Cochran, W.G. and Cox, G.M. (1957): *Experimental Designs*. John Wiley & Sons, Inc. New York.
- [8] Cox, V.G. (1950): *Experimental Design*. *Biometrics* **6**, 301-302.
- [9] Cullis, B.R., and Gleeson, A.C. (1991): Spatial Analysis of Field Experiments. An Extension to Two Dimensions. *Biometrics* **47**, 1449-1460.
- [10] Das, M.N. and Giri, N.C. (1986). *Design and analysis of experiments*. New Age international (P) limited.
- [11] Drane, J.W. (1989): *Compromises and Statistical Designs for Grazing Experiments in Grazing Research*. Crop Science Society of America and American Society of Agronomy: **16**, 69-83.
- [12] Edme, S.J., Tai, and J.D. Miller (2007): Relative Efficiency of Spatial Analysis for Non-replicated Early-stage Sugarcane Field Experiments. *Journal of the American Society of Sugar Cane Technologists*, **27**: 89-104.
- [13] Eskridge, K.M. (2008): *Field Trial Designs in Plant Breeding*. University of Nebraska, Lincoln, NE.
- [14] Gharde, Y. (2000): *Analysis of Covariance*. I.A.S.R.I., Library Avenue, New Delhi.

- [15] Girma T. A. (2002): Design and Analysis of Field Experiments in Agriculture. Ethiopian Agriculture Research Organization. Technical Manual No. 15, Addis Ababa.
- [16] Girma T. A. (2005): Using Spatial Modeling Techniques to Improve Data Analysis from Agricultural Field Trial. PhD Thesis of University of KwaZulu-Natal, Pietermaritzburg.
- [17] Gleeson, A.C., and Cullis, B.R. (1987): Residual Maximum Likelihood (REML) Estimation of Neighbour Model for Field Experiments. *Biometrics* **43**, 277-288.
- [18] Gomez, K.A. and Gomez, A.A. (1984): Statistical Procedures for Agricultural Research, 2nd ed. John Wiley and Sons, Inc.; New York
- [19] Gunjaca J., Renka H., Jambresic I. and Sindrak Z. (2005): Efficiency of Alpha Designs in Croatian Variety Trials. 27th Int. Conf. Information Technology Interfaces ITI, Cavtat, Croatia.
- [20] Hatfield, J.L. (2000): Precision Agriculture and Environmental Quality: Challenges for Research and Education. U.S. Department of Agriculture, in Ames, Iowa.
- [21] Hettasch, H.M. (2009): Applicability of best linear unbiased prediction (BLUP) for the selection of Ortets in Eucalyptus hybrid populations. Master Dissertation, University of Pretoria.
- [22] Idrees N. and Khan, M.I. (2009): Design improvement Using Uniformity Trials Experimental data. *Pakistan Journal Agricultural Science*. Vol. 46, Faisalabad
- [23] Karcher D.E., Richardson M.D., and Secks M.E (2003): Does Blocking Affect Experimental Efficiency on Sand-Based Putting Greens? *Crop Science Society of America*, **43**: 2295-2297.
- [24] Kempton, R.A., Seraphin, J.C., Sword, A.M. (1994): Statistical Analysis of Two Dimensional Variations in Variety Yield Trials. *Journal of Agricultural Science, Cambridge*, **122**, 335-342.
- [25] Kerr MK. (2003): Design Considerations for Efficient and Effective Microarray Studies. *Biometrics* **59**: 822-828.
- [26] Kerr, MK. And Churchill, GA. (2001): Statistical Design and the Analysis of Gene Expression Microarray Data. *Genet Res* **77**: 123-128.
- [27] Knight E. (2008): Improved Iterative Schemes for REML Estimation of Variance Parameters in Linear Mixed Models. School of Agriculture, Food and Wine, University of Adelaide.

- [28] Leonardo Journal of sciences (2009): Statistical Approaches in Analysis of Variance, Romania.
- [29] Little, T. M. and Hills, F. J. (1978): Agricultural Experimentation, Design and Analysis. John Wiley and Sons, Inc. New York.
- [30] Mandefro N. (2005): Statistical procedures for Designed Experiments. Ethiopian Agriculture Research Organization, Addis Ababa.
- [31] Masoon, A. A., Mujahid Y., Khan, M.I., and Abid, S. (2006): Improving precisions of Agricultural field Experiments. Journal of Sustainable Development, Vol.3. 11-13.
- [32] Masood, M. A., Farooq K., Mujahid Y. and Anwar, M. Z. (2008): Improvement in Precision of Agricultural Field Experiments through Design and Analysis. Pak. J. Life Soc. Sci., **6**(2): 89-91.
- [33] Milliken, G.A and Johnson, D.E (1992): Analysis of Messy Data. Volume I: Designed experiments. Chapman and Hall, New York, NY, USA.
- [34] Montgomery, D. C. (1991): Design and Analysis of Experiments. 3rd ed. John Wiley and Sons, Inc. New York.
- [35] Morrell H. C. (1998): Likelihood Ratio Testing of Variance Components in the Linear Mixed-Effects Model Using Restricted Maximum Likelihood. International Biometric Society. Biometrics, **54**, 1560-1568.
- [36] Neter, J., W. Wasserman, and M. H. Kutner. (1990): Applied linear statistical models: regressions, analysis of variance, and experimental designs, 3rd ed. Irwin, Burr Ridge, Illinois.
- [37] Panse, V.G. (1958): Status of Agricultural Experiments. F.A.O., Rome.
- [38] Parsad, R. and Gupta, V.K. (2000): Covariance Analysis. I.A.S.R.I., Library Avenue, New Delhi.
- [39] Patterson, H.D., and Hunter, E.A. (1983): The efficiency of incomplete block designs in National List and Recommended cereal variety trials. J. Agric. Sci., **101**, 427-433.
- [40] Patterson, H.D. and Williams, E.R. (1976): A New Class of Resolvable Incomplete Block Designs. Biometrika, **63**, 83-92.
- [41] Patterson, H. D. and Silvey, V. (1980): Statutory and Recommended List Trials of Crop Varieties in the United Kingdom. J. R. Statist. Soc. A. **143**, 219-253.
- [42] Pearce, S.C. (1983): The Agricultural Field Experiments. A Statistical Examination of Theory and Practice. John Wiley and Sons, Inc., New York.

- [43] Raza I. and Masood, M.A (2009): Efficiency of Lattice Design in Relation to Randomized Complete Block Design in Agricultural Field Experiments. *Pakistan J. Agric. Res.* **22**: 150-153.
 - [44] Sharma, V.K. (2000): *Balanced Incomplete Block Designs*. I.A.S.R.I., Library Avenue, New Delhi-110012.
 - [45] Shunmughasundaram, S., M.B. Batcha, K.V. Bhagyalaskshmi, and S.S. Narayanam (1980): Suitability of Augmented Design in Sugarcane Breeding Trials. *Proc. Int. Soc. Sugar Cane Technol.* 1168-1176.
 - [46] Snyder, E.B. (1962): *Lattice and Compact Family Block Designs in Forest Genetics*. U.S. Department of Agriculture, Gulfport, Mississippi.
 - [47] Vas Es, C.P. G. and Sellmann, M. (2007): Spatially-Balanced Complete Block Designs for field experiments. *Cornell University, Ithaca, United States. Geoderma* **140**, 346-352.
 - [48] Wu, T. and Dutilleul, P. (1999): Validity and Efficiency of Neighbor Analysis in Comparison with Classical Complete and Incomplete Block Analysis of Field Experiments. *Agron. J.* **91**: 721-731.
-
- [49] Yates, F. (1936): A New Method of Arranging Variety Trials Involving a Large Number of varieties. *J. Agric. Sci.* **26**, 424-455.
 - [50] Yates, F. (1939): The Recovery of Inter-block Information in Variety Trials Arranged in Three-dimensional Lattices. *Ann. Eugen.* **9**, 136-156.
 - [51] Yates, F. (1940): The Recovery of Inter-block Information in Balanced Incomplete Block Designs. *Ann. Eugen.* **10**, 317-325.
 - [52] Yau, S.K. (1997): Efficiency of alpha-lattice designs in international variety yield trials of barley and wheat. *The Journal of Agricultural Science*, **128**, 5-9.
 - [53] Yigrem Tefera (2007): *Application of Balanced Incomplete Block Design in PPS Sampling Without Replacement*. Unpublished Master Thesis, Addis Ababa University.

APPENDICES

Appendix I

Table 18: Tests for Normality of RCBD in DW99RVTRII data set

Location	Year			
	2006/05		2007/06	
	Shapiro-Wilk Test		Shapiro-Wilk Test	
	W-statistic	Sign.	W-statistic	Sign.
1	0.981372	0.295	0.977	0.170
2	0.981372	0.295	0.985	0.467
3	0.992694	0.935	0.986	0.527
4	0.982415	0.340	0.875	0.106
5	0.987121	0.606	0.991	0.861
6	0.988063	0.669	0.982	0.331

Table 19: Tests for Normality of RCBD in DW99RVTR data set

		DW99RVTR data set	
Location	Season (Year)	Shapiro-Wilk Test	
		Statistic	Sign.
One	2005/04	0.973	0.102
	2006/05	0.978	0.221
	2007/06	0.987	0.646
Two	2005/04	0.990	0.818
	2006/05	0.976	0.158
	2007/06	0.995	0.990
Three	2005/04	0.985	0.491
	2006/05	0.982	0.377
	2007/06	0.995	0.990
Four	2005/04	0.991	0.879
	2006/05	0.982	0.356
	2007/06	0.984807	0.500

Table 20: Tests for Normality of lattice design, alpha lattice design and Covariance analysis

Shapiro-Wilk Test		
	W-statistic	Sign.
Lattice design		
Melkasa	0.981	0.303
Dhera	0.988	0.702
Ziway	0.984	0.486
Mieso	0.985	0.541
Kobo	0.976	0.161
Alpha lattice design		
2A3099AW	0.982	0.099
2A3499AW	0.989	0.266
Covariance analysis		
	0.993	0.944

Table 21: Tests of Normality for Split-plot

Location	Shapiro-Wilk Test	
	W-statistic	Sign.
Jeldu	0.992	0.944
Degem	0.964	0.365
Wolmer	0.966	0.455

Table 22: Missing values and their estimates using estimation approach in RCBD

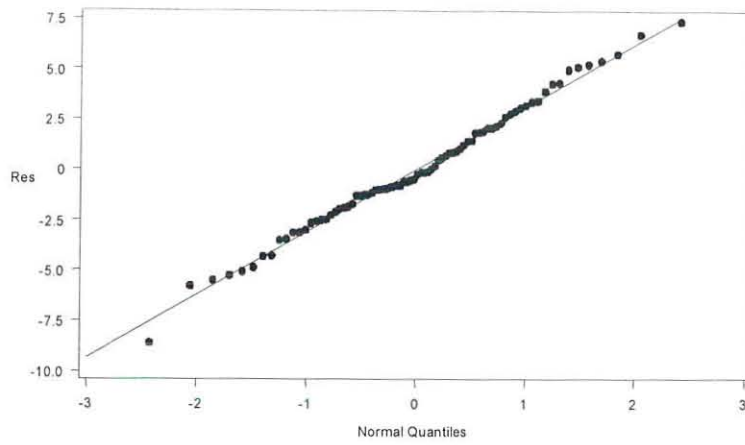
No of missing values	Treatment	Block	Estimated value
Two (from the same block)	9	1	39.005
	19	1	37.574
Two (from d/t block d/t treatment)	5	2	40.670
	1	3	32.775
Two (from d/t block but the same treatment)	5	2	41.530
	5	3	37.903
Three	2	1	40.876
	8	2	29.676

	6	4	32.720
Four	2	1	41.022
	8	2	29.768
	4	3	32.820
	6	4	32.810
	1	1	36.348
Five	3	2	39.671
	9	3	33.366
	10	4	36.546
	11	4	27.871
	1	1	36.529
Six	3	2	39.290
	5	3	33.609
	9	3	33.426
	10	4	36.447
	11	4	28.045
Eight	1	1	36.456
	7	1	38.148
	3	2	39.671
	15	2	34.619
	5	3	33.371
	9	3	33.123
	10	4	36.138
	11	4	27.733

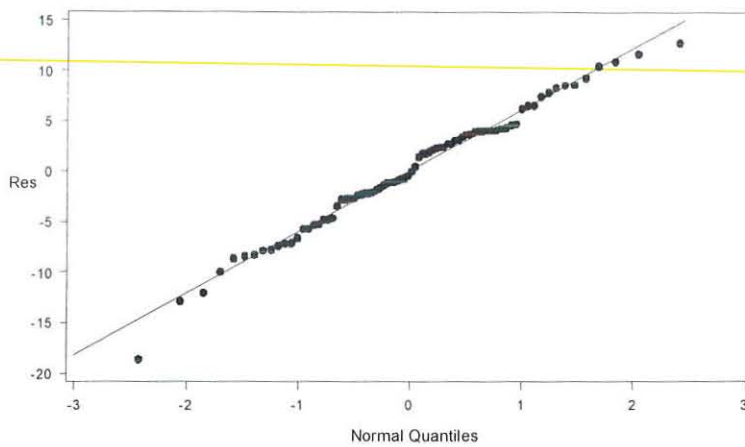
Table 23: Mean and Standard error (SE) of Durum wheat variety with and without covariate

Variety	With covariate		Without covariate		Actual value of yield (undamaged yield)	
	Adjusted mean	Standard Error	Unadjusted mean	Standard Error	Mean	Standard Error
1	29.510	2.511	26.739	4.108	34.526	2.351
2	36.930	2.511	38.945	4.108	38.945	2.351
3	34.908	2.511	32.137	4.108	38.380	2.351
4	33.872	2.511	35.887	4.108	35.887	2.351
5	33.121	2.511	30.350	4.108	35.591	2.351
6	36.319	2.511	38.335	4.108	38.335	2.351
7	28.312	2.511	27.934	4.108	32.215	2.351
8	28.072	2.511	27.694	4.108	27.694	2.351
9	31.274	2.511	30.897	4.108	34.696	2.351
10	35.789	2.511	35.411	4.108	39.460	2.351
11	32.100	2.511	26.936	4.108	30.152	2.351
12	32.372	2.511	34.387	4.108	34.387	2.351
13	35.745	2.511	37.760	4.108	37.760	2.351
14	30.587	2.511	32.602	4.108	32.602	2.351
15	32.132	2.511	26.968	4.108	32.869	2.351
16	31.110	2.511	33.125	4.108	33.125	2.351
17	30.613	2.511	32.628	4.108	32.628	2.351
18	39.206	2.511	41.222	4.108	41.222	2.351
19	33.556	2.511	35.571	4.108	35.571	2.351

Appendix II

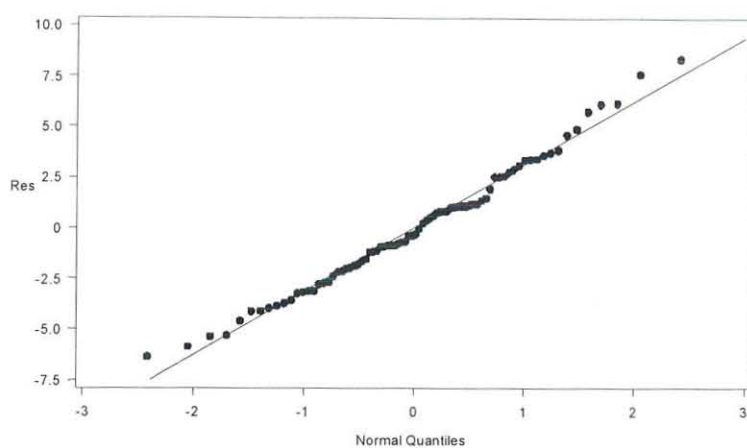


a. Location three

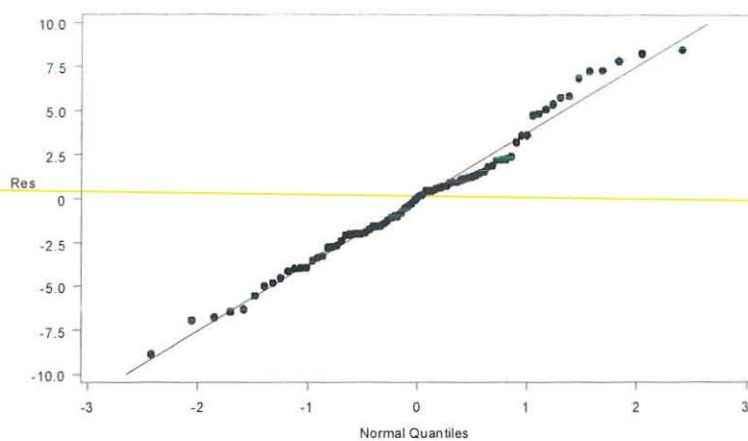


b. Location five

Figure 1: QQ-Plot of Residuals from RCBD in 2006/05 of BW99RVII data set

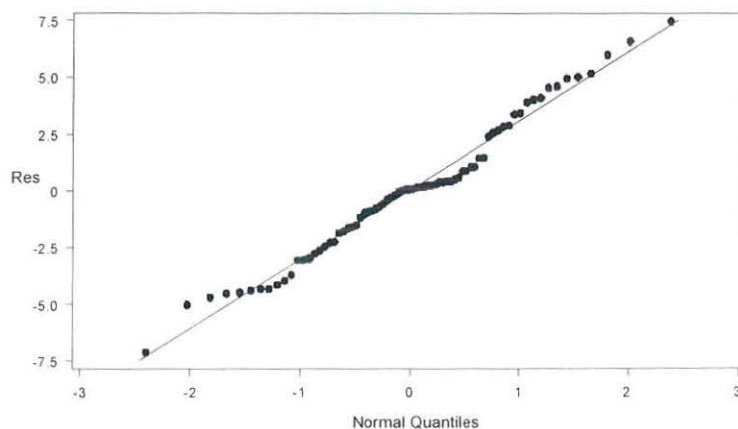


a. Location three

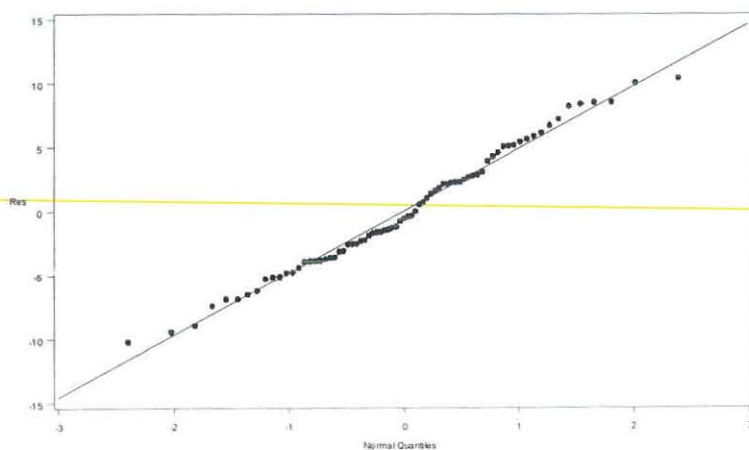


b. Location six

Figure 2: QQ-Plot of Residuals from RCBD in 2007/06 of BW99RVII data set



a. Melkasa research centre



b. Ziway research centre

Figure 3: QQ-Plot of Residuals from lattice design

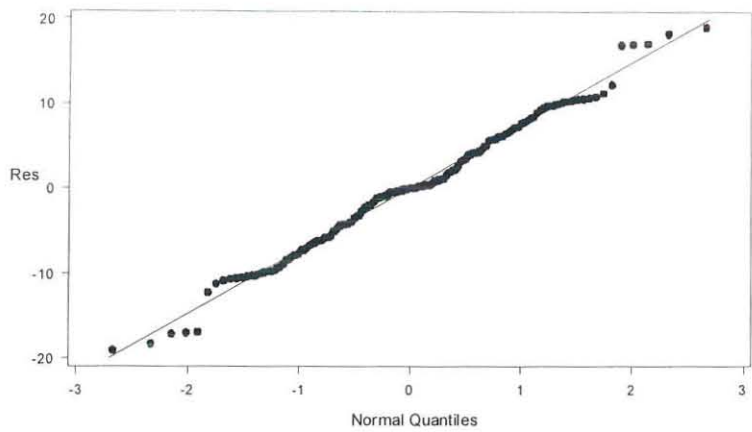


Figure 4: QQ-Plot of Residuals from Alpha lattice design of 2A3499AW data set

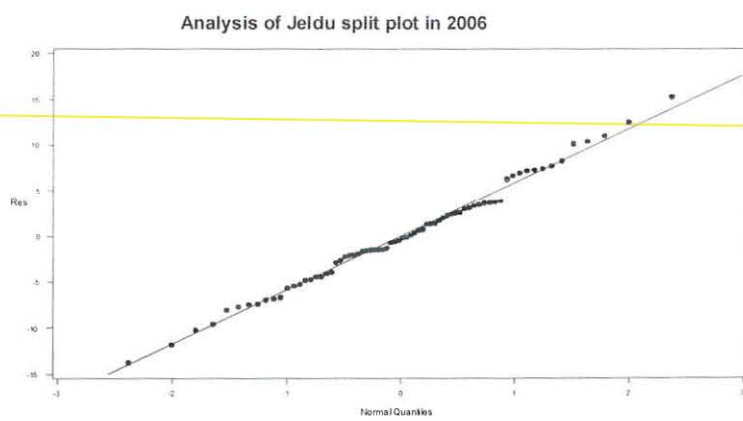


Figure 5: QQ-Plot of Residuals for Split plot of data set from Jeldu site

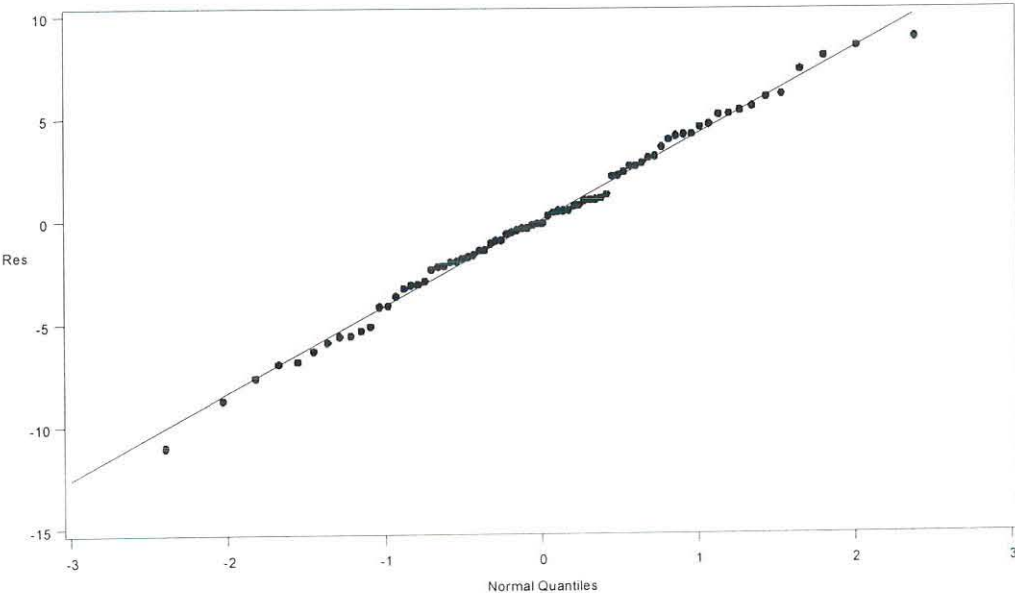


Figure 6: QQ-Plot of Residuals for ANCOVA

DECLARATION

I, the undersigned, declare that the thesis is my original work, has not been presented for degrees in any other University and all sources of material used for the thesis have been duly acknowledged.

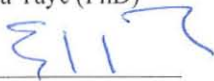
Name: Bacha Edosa

Signature:  _____

Date: 09/20/2020

This thesis has been submitted for examination with my approval as a University advisor.

Name: Girma Taye (PhD)

Signature:  _____

Data: 09/10/10