

Use of incomplete block designs in agricultural experiments

V.K. GUPTA¹ AND RAJENDER PARSAD²

Indian Agricultural Statistics Research Institute, New Delhi 110 012

Received : March 2012; Revised accepted : May 2012

ABSTRACT

The purpose of this article is to demonstrate the application of incomplete block designs with factorial structure in agronomic experiments, particularly the crop sequences experiments. The paper also describes the usefulness of incomplete split-plot designs in factorial experiments with many sub-plot treatments.

In designed experiments the two main sources of variability in the data generated are the treatments and the experimental units. Variability due to treatments arises because it is expected that the treatment effects are different. This is a deliberate attempt on the part of the experimenter to generate variability. However, the variability in the experimental material needs to be accounted for through designing the experiment properly and including this source of variability in the model. If this is not taken care of, then the experimental error would be high. Block designs are useful for elimination of heterogeneity, i.e., when there is only one source of variability in the experimental material. A block design is said to be complete if it contains every treatment once in each block. A randomized complete block (RCB) design is an example of a complete block design. A complete block design is also an orthogonal design in the sense that the covariance between the block effects and the treatments effects is zero. In such designs, the adjusted sum of squares due to treatments (or blocks) is same as unadjusted sum of squares due to treatments (or blocks). There is no loss of information in estimating the treatment effects and the treatments effects are estimated with full efficiency. A block design is said to be incomplete if it has at least one block that does not contain all the treatments in it. Balanced Incomplete Block (BIB) designs and Partially Balanced Incomplete Block (PBIB) designs are examples of incomplete block designs. Incomplete block designs are non-orthogonal designs.

There is confounding of treatment effects with the block effects. As such there is a loss of precision in the estimation of treatment effects. However, this loss is more

than compensated by the reduction in intra block variance due to the reduction in the block size. In block designs the idea is to form blocks in a manner such that the experimental units within a block are as homogeneous as possible, so that between blocks variation can be isolated from the experimental error. This results into reduction of the error component, which in turn enhances the precision of the comparison of treatment effects.

Although it is well known that incomplete block designs help in making precise comparisons of the treatment effects by controlling the experimental error, these designs do not find favour with the agricultural experiments because the experimenter needs to demonstrate the treatment effects at one place, which is not possible by using an incomplete block design. In order to circumvent this problem, recourse is made to another useful concept in the context of incomplete block designs, which is known as resolvability of block designs. An incomplete block design is said to be *$\hat{a}$ -resolvable* if it is possible to rearrange the blocks into different groups in such a way that the blocks in each group form a complete replication in the sense that all the treatments appear $\hat{a}$ times in each group. Designs with $\hat{a} = 1$ are known as *1-resolvable* or simply *resolvable* designs. For example, suppose that an experimenter is interested in conducting an experiment with $v = 10$ treatments in $r = 3$ replications and with $b = 6$ blocks of size $k = 5$ each. The resolvable design (d) is the following:

Replication I		Replication II		Replication III	
Block 1	Block 2	Block 1	Block 2	Block 1	Block 2
1	2	1	2	1	2
3	4	3	4	4	3
5	6	6	5	5	6
7	8	8	7	8	7
9	10	10	9	9	10

¹Corresponding Author Email: vkgupta_1751@yahoo.co.in

¹ICAR National Professor, IASRI, New Delhi; ²Head, Division of Design Experiments

In this design the blocks have been so arranged that two blocks together form a complete replication. The randomization procedure for these designs is very simple. It consists of the following steps:

1. Randomize the treatment numbers (or labels).
2. Randomize the replications.
3. Randomize the blocks within replication.
4. Randomize the treatments within blocks.

The randomized layout of the design (d) in the Example is the following:

Replication I		Replication II		Replication III	
Block 1	Block 2	Block 1	Block 2	Block 1	Block 2
1	2	1	2	1	2
10	7	7	5	2	7
1	5	3	9	8	1
8	9	10	4	6	3
3	2	6	8	10	9
6	4	2	1	4	5

Resolvable block designs offer a great deal of flexibility in the sense that, by using them, experiments can be conducted with one replication at a time, with different replications applied at different places or at different points of time. For example, field trials with large number of crop varieties cannot always be laid out in a single location or a single season. Therefore, it is desired that variation due to location or time periods may be controlled along with controlling variation within location or within time periods. This can be handled by using resolvable block designs. Here, locations or time periods may be taken as replications and the variation within a location or a time period can be taken care of by blocking within the replication.

In general, for any resolvable incomplete block design with parameters as number of treatments = $v = pq$, number of blocks = $b = rp$, replication = r , block size = q , and number of blocks per replication = p , the ANOVA would be the following:

ANOVA		ANOVA for design d	
Source	DF	Source	DF
Treatments	$pq - 1$	Treatments	9
Blocks	$rp - 1$	Blocks	5
Replications	$r - 1$	Replications	2
Blocks (Replications)	$r(p - 1)$	Blocks (Replications)	3
Error	$rpq - rp - pq + 1$	Error	15
Total	$rpq - 1$	Total	29

The advantage of using a resolvable incomplete block design in place of a randomized complete block (RCB) design is that the variation between blocks within replica-

tion is taken away from the error, leading thereby into the reduction of experimental error without the requirement of any additional resources for conducting the experiment. For details on resolvable designs, the reader may refer to Parsad et al. (2007) or visit www.iasri.res.in/design/.

When the treatment structure is factorial in nature in the sense that the v treatments are in fact all possible combinations of the levels of two factors, one factor A at p levels and the other factor B at q levels, then the treatments sum of squares can be split up into components like main effects and two factor interactions. The ANOVA would then be

ANOVA

Source	DF
Treatments	$pq - 1$
Main Effect A	$p - 1$
Main Effect B	$q - 1$
Interaction AB	$(p - 1)(q - 1)$
Blocks	$rp - 1$
Replications	$r - 1$
Blocks (Replications)	$r(p - 1)$
Error	$rpq - rp - pq + 1$
Total	$rpq - 1$

It is indeed possible that the factor A has also factorial structure and the factor B has also factorial structure. For example, the p levels of factor A may be all possible combinations of the levels of s factors, i.e., $p = f_1 \times f_2 \times \dots \times f_s$, and the q levels of factor B may be all possible combinations of the levels of t factors, i.e., $q = h_1 \times h_2 \times \dots \times h_t$. In this case we have number of treatments $v = pq = f_1 \times f_2 \times \dots \times f_s \times h_1 \times h_2 \times \dots \times h_t$. The ANOVA would also get modified accordingly and the treatment sum of squares would be split up into the sum of squares due to main effects, two factor, three factor, and so on and $s+t$ factor interactions.

The construction of such designs with factorial structure is algebraically complicated because of the fact that it is an incomplete block design. There would be confounding of effects with the blocks within replication. It is quite possible that even the main effects are confounded. So there would be loss of efficiency on all the effects, main effects and interactions. Thus there is a need to construct designs in such a way that there is no loss of efficiency on the main effects and there is a control on the efficiency of the lower order interactions, particularly two factor interactions.

Orthogonal block designs with factorial structure are an answer to this problem. Among this class of designs are

the popular Extended Group Divisible designs. But a problem with these designs is that generally there is a loss of efficiency on the main effects with control on the efficiency of two factor interactions. These designs are very useful in crop sequence experiments and have been used very effectively by many experimenters.

Crop Sequence Experiments

In Indian National Agricultural Research System (NARS), many experiments are conducted for cropping systems research wherein a crop sequence of two crops is grown; one in *kharif* season followed by another in *rabi* season. More than two crops may also be grown in crop sequence experiment, but we focus our attention to only two crops grown in the sequence. The extension to three crops would follow naturally. In these experiments two different sets of treatments are applied in succession; one set applied in *kharif* crop and the other set applied in *rabi* crop. The observations are recorded in both the crop seasons. In two-crop sequence experiments, the interest of the experimenter is in direct effects of treatments applied in *kharif* and *rabi* season and residual effects of *kharif* treatments. The interaction between the residual effect of *kharif* crop treatments and the direct effect of *rabi* crop treatments may or may not be of interest to the experimenter. These experiments may be run in a block design to account for the variability in the experimental material.

Crop sequence experiments run in block designs can be viewed as block designs with factorial structure of treatments. For example if p treatments are applied in *kharif* crop and q treatments are applied in *rabi* crop, then on the application of q treatments in *rabi* crop, a maximum of pq treatment combinations may appear and the treatment structure is factorial in nature. These experiments are classified into two broad categories. *Category I* experiments are those in which all the possible pq treatment combinations applied in both the crops are taken for experimentation; in other words, one has a complete factorial experiment run in a block design. In these experiments, the interaction between the residual effect of *kharif* crop treatments and the direct effect of *rabi* crop treatments are also of interest to the experimenter. *Category II* experiments are similar to the category I experiments with the difference that after the application of q treatments in the second crop, all the pq treatment combinations do not appear. In other words, it is a fractional factorial plan rather than a complete factorial experiment run in a block design. Further, the interaction between residual effect of *kharif* treatments and direct effect of *rabi* treatments is not of interest to the experimenter.

Category I Experiments

In these experiments, all the mn treatment combinations are run in a block design and the interactions are of interest to the experimenter. An experiment is to be conducted for evaluation of sulphur sources at varying rates in aerobic rice–wheat cropping system for improved productivity and soil health. During *kharif*-rice, treatments comprising of combinations of two sources of sulphur, viz., a_0 = Gypsum and a_1 = Phosphogypsum and three levels of sulphur as $b_0 = 0$, $b_1 = 30$ and $b_2 = 60$ kgs. per hectare are to be applied. During the *rabi*-wheat three rates of sulphur to be applied through respective sources are $c_0 = 0$, $c_1 = 15$ and $c_2 = 30$ kg per hectare. The study needs to be conducted for two years. The interest of the experimenter is to study (i) the direct effect of sulphur sources and levels of sulphur on *kharif*-rice, (ii) the residual effect of sulphur treatments applied to *kharif* rice on succeeding *rabi*-wheat, (iii) the direct effect of sulphur applied to *rabi*-wheat, (iv) the interaction between residual effects of sulphur treatments applied during *kharif*-rice and the direct effect of sulphur treatments applied in *rabi*-wheat, (v) the residual effect of sulphur applied to *rabi*-wheat on succeeding *kharif*-rice in the second year, (vi) the cumulative effect of sulphur applied to *kharif*-rice and *rabi*-wheat in the first year on the *kharif*-rice and *rabi*-wheat of second year.

Split-plot design

For the first year, the *kharif*-rice has 6 treatment combinations denoted as $1 = a_0b_0$, $2 = a_0b_1$, $3 = a_0b_2$, $4 = a_1b_0$, $5 = a_1b_1$, $6 = a_1b_2$ and three *rabi* treatments denoted as c_0 , c_1 , c_2 . The experiment is generally run as follows: The experiment for *kharif*-rice treatments is run as RCB design in three replications (rows as replications):

Replication I	1	6	3	2	5	4
Replication II	3	5	1	6	4	2
Replication III	4	1	6	3	2	5

For application of *rabi*-wheat treatments, the plots in each replication of the design in *kharif* season are subdivided into three sub-plots and the three *rabi* treatments are applied randomly to these subplots in each plot. The layout of the design for *Rabi* season is the following:

The design obtained after the application of *rabi*-wheat treatments becomes a split plot design in three replications with main plot treatments representing residual effect of *kharif*-rice treatments and sub-plot treatments representing the direct effect of *rabi*-wheat treatments. This design also gives the interaction of the residual effect of *kharif* treatments and the direct effect of *rabi* treatments. It is assumed that the residual effects persist for one season only.

Replication I																	
Main-plot 1			Main-plot 2			Main-plot 3			Main-plot 4			Main-plot 5			Main-plot 6		
1	1	1	6	6	6	3	3	3	2	2	2	5	5	5	4	4	4
c_0	c_2	c_1	c_2	c_0	c_1	c_1	c_0	c_2	c_0	c_2	c_1	c_2	c_1	c_0	c_1	c_0	c_2
Replication II																	
Main-plot 1			Main-plot 2			Main-plot 3			Main-plot 4			Main-plot 5			Main-plot 6		
3	3	3	5	5	5	1	1	1	6	6	6	4	4	4	2	2	2
c_1	c_0	c_2	c_0	c_1	c_2	c_1	c_2	c_0	c_2	c_1	c_0	c_2	c_0	c_1	c_0	c_2	c_1
Replication III																	
Main-plot 1			Main-plot 2			Main-plot 3			Main-plot 4			Main-plot 5			Main-plot 6		
4	4	4	1	1	1	6	6	6	3	3	3	2	2	2	5	5	5
c_0	c_2	c_1	c_1	c_0	c_2	c_2	c_1	c_0	c_1	c_0	c_2	c_2	c_1	c_0	c_2	c_1	c_0

The treatments applied to *kharif* crop have a residual effect on *rabi* crop and the treatments applied to *rabi* crop have a residual effect on the following *kharif* crop. This experiment, as like many such experiments, is continued for two years. It could have been continued even for three or more years. In the second year the *kharif*-rice treatments are again applied on main plots (*same randomization*) and observations are recorded sub-plot wise. The sub-plot differences give the residual effects of *rabi* treatments and main plots give the direct effects of *kharif* treatments and the procedure is continued.

Incomplete split-plot design

When experimenting with treatments having factorial structure with two factors at levels p and q respectively, it is often of interest to evaluate the set of treatments (levels) of the second factor under different conditions of the levels of the first factor. In other words, the interaction between the two set of treatments is important. The split-plot design is a design that may be used for the factorial experiments when it is difficult to change the levels of one of the factors. These designs require separate randomization of the whole-plot and sub-plot treatments. Inevitably, an increasing number of the levels of the second factor (sub-plot treatments) will make the variation between sub-plots within main-plots larger, and therefore, the comparisons of sub-plot treatments will be more uncertain. Some types of incomplete blocks are then often used within each main-plot to increase the efficiency of comparisons within main-plots. For details on incomplete split-plot designs reference may be made to Rees (1969), Mejza and Mejza (1984), Hering and Mejza (2002) and Kritensen (2012).

Incomplete split plot designs are a practical compromise to the adoption of an RCB and split-plot design. Incomplete split-plot design may require guard-area between

incomplete blocks. The main-plot and sub-plot comparisons in incomplete split-plot design are more efficient compared to same in split plot design.

There are many methods of construction of incomplete split plot designs. These can be obtained by making use of group divisible designs as well as alpha designs. Kritensen (2012) found alpha designs to be very useful in construction of incomplete split plot designs and gave four methods of constructing these designs. Alpha designs are available at Design Resources Server (www.iasri.res.in/design). Some examples of incomplete split plot designs, obtained by using four different methods of Kritensen (2012), are given below:

Consider the problem of constructing an incomplete split-plot design with 3 ($=p$) main-plot treatments labelled as “a”, “b”, “c” and 8 ($=q$) sub-plot treatments labeled as numerals 1, 2, 3, . . . , 7, 8. Consider an α -design in 8 ($=q$) treatments arranged in 2 ($=s$) blocks of size 4 ($=k$) each. Obviously, $v = sk = q$. The α -design has 2 ($=r$) replications of treatments. For getting an α -design one may refer to www.iasri.res.in/design/.

An α -design with $v = 8$, $s = 2$, $k = 4$, $r = 2$, with rows as replications is

Replication I	1, 3, 5, 7	2, 4, 6, 8
Replication II	1, 3, 6, 8	2, 4, 5, 7

This α -design can be used in different ways to obtain incomplete split-plot design. Since there are 3 main-plot treatments, the 6 ($=sp$) incomplete blocks formed for both the replications from the 24 ($=pq$) combinations are the following:

Method 1: The incomplete split plot design obtained after the randomization of the 6 blocks in each replication is the following:

Method 2: The incomplete split plot design is obtained

Replication I																							
<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>
1	3	5	7	2	4	6	8	1	3	5	7	2	4	6	8	1	3	5	7	2	4	6	8
Replication II																							
<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>
1	3	6	8	2	4	5	7	1	3	6	8	2	4	5	7	1	3	6	8	2	4	5	7
Replication I																							
<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>
2	4	6	8	1	3	5	7	1	3	5	7	2	4	6	8	2	4	6	8	1	3	5	7
Replication II																							
<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>
1	3	6	8	2	4	5	7	1	3	6	8	2	4	5	7	1	3	6	8	2	4	5	7
Replication I																							
<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>
1	3	5	7	1	3	5	7	1	3	5	7	2	4	6	8	2	4	6	8	2	4	6	8
Replication II																							
<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>
2	4	5	7	2	4	5	7	2	4	5	7	1	3	6	8	1	3	6	8	1	3	6	8

after the randomization of the blocks within each group of identical blocks. The design is the following:

In other methods, we begin with an $\hat{a}$ -design in 6 ($= rt$) replications. An $\hat{a}$ -design with $v = 8$, $r = 6$, $s = 2$, $k = 4$ is the following:

Replication I	1, 3, 5, 7	2, 4, 6, 8
Replication II	1, 3, 6, 8	2, 4, 5, 7
Replication III	1, 4, 5, 8	2, 3, 6, 7
Replication IV	1, 4, 6, 7	2, 3, 5, 8
Replication V	1, 3, 5, 8	2, 4, 6, 7
Replication VI	1, 2, 3, 4	5, 6, 7, 8

Method 3:

The blocks in the first two replications are assigned to the main-plot treatment *a*, the blocks in 3rd and 4th replication are assigned to the main-plot treatment *b* and the blocks in the last two replications are assigned to the main-plot treatment *c*. These six replications are then converted into two replications by joining the blocks of replications 1, 3 and 5 and then of replications 2, 4 and 6. The following layout is generated:

Next the 6 blocks are randomized independently in each replication. The final randomized incomplete split-plot design is the following:

Method 4: In this method once again like Method 3, form 6 blocks in two replications by joining blocks in replications 1, 3 and 5 and 2, 4 and 6 independently like in Method 3. After this combine the blocks with different levels of main plots. This results into the following layout:

Now randomize the blocks within each group in each replication. The final layout of the design is the following:

Finally, in all the methods, it is assumed that the main-plot and sub-plot treatment levels are randomized.

The analysis of data generated will depend upon the type of randomization done or the method of construction.

This is so because the model is defined according to the method of construction.

Block designs with factorial structure

The use of a split-plot design for crop sequence experiments has the following limitations: (i) the treatment effects in sub-plots and the interaction between residual effects of *kharif* treatments and the direct effects of *rabi* treatments are estimated with high precision but the precision for main plot treatments is generally low, (ii) each block is a complete replicate and with large number of treatment combinations it is difficult to maintain homogeneity, (iii) because of two error components the combined analysis of experiments conducted over years is a problem particularly when the error variances are heterogeneous, (iv) the identification of best treatment combination is a problem.

In general, the designing problem in crop sequence experiments is more like an asymmetrical factorial experiment with two factors at p and q levels respectively, run in an incomplete (preferably resolvable) block design. Therefore, block designs having orthogonal factorial structure (OFS) with balance or balanced confounded factorial designs are useful in this situation (Parsad 2007; Gupta and Parsad, 2009). It is well known that an extended group divisible (EGD) design, if existent, has OFS with balance. Therefore, an EGD design may be useful, effective and efficient alternative to the split plot design used and described above. *Kronecker Product* type designs could also be used in this situation. Kronecker Product designs also have an OFS with balance. One easy way to obtain block designs having OFS with balance is through *Kronecker Product* of designs having OFS with balance. The efficiency of the design obtained in terms of main effects and

Replication I																							
<i>a</i> 1	<i>a</i> 3	<i>a</i> 5	<i>a</i> 7	<i>a</i> 2	<i>a</i> 4	<i>a</i> 6	<i>a</i> 8	<i>b</i> 1	<i>b</i> 4	<i>b</i> 5	<i>b</i> 8	<i>b</i> 2	<i>b</i> 3	<i>b</i> 6	<i>b</i> 7	<i>c</i> 1	<i>c</i> 3	<i>c</i> 5	<i>c</i> 8	<i>c</i> 2	<i>c</i> 4	<i>c</i> 6	<i>c</i> 7

Replication II																							
<i>a</i> 1	<i>a</i> 3	<i>a</i> 6	<i>a</i> 8	<i>a</i> 2	<i>a</i> 4	<i>a</i> 5	<i>a</i> 7	<i>b</i> 1	<i>b</i> 4	<i>b</i> 6	<i>b</i> 7	<i>b</i> 2	<i>b</i> 3	<i>b</i> 5	<i>b</i> 8	<i>c</i> 1	<i>c</i> 2	<i>c</i> 3	<i>c</i> 4	<i>c</i> 5	<i>c</i> 6	<i>c</i> 7	<i>c</i> 8

Replication I																							
<i>b</i> 1	<i>b</i> 4	<i>b</i> 5	<i>b</i> 8	<i>a</i> 1	<i>a</i> 3	<i>a</i> 5	<i>a</i> 7	<i>c</i> 2	<i>c</i> 4	<i>c</i> 6	<i>c</i> 7	<i>a</i> 2	<i>a</i> 4	<i>a</i> 6	<i>a</i> 8	<i>c</i> 1	<i>c</i> 3	<i>c</i> 5	<i>c</i> 8	<i>b</i> 2	<i>b</i> 3	<i>b</i> 6	<i>b</i> 7

Replication II																							
<i>c</i> 1	<i>c</i> 2	<i>c</i> 3	<i>c</i> 4	<i>a</i> 1	<i>a</i> 3	<i>a</i> 6	<i>a</i> 8	<i>b</i> 1	<i>b</i> 4	<i>b</i> 6	<i>b</i> 7	<i>a</i> 2	<i>a</i> 4	<i>a</i> 5	<i>a</i> 7	<i>c</i> 5	<i>c</i> 6	<i>c</i> 7	<i>c</i> 8	<i>b</i> 2	<i>b</i> 3	<i>b</i> 5	<i>b</i> 8

Replication I																							
<i>a</i> 2	<i>a</i> 4	<i>a</i> 6	<i>a</i> 8	<i>c</i> 2	<i>c</i> 4	<i>c</i> 6	<i>c</i> 7	<i>b</i> 2	<i>b</i> 3	<i>b</i> 6	<i>b</i> 7	<i>b</i> 1	<i>b</i> 4	<i>b</i> 5	<i>b</i> 8	<i>a</i> 1	<i>a</i> 3	<i>a</i> 5	<i>a</i> 7	<i>c</i> 1	<i>c</i> 3	<i>c</i> 5	<i>c</i> 8

Replication II																							
<i>b</i> 1	<i>b</i> 4	<i>b</i> 6	<i>b</i> 7	<i>a</i> 1	<i>a</i> 3	<i>a</i> 6	<i>a</i> 8	<i>c</i> 1	<i>c</i> 2	<i>c</i> 3	<i>c</i> 4	<i>b</i> 2	<i>b</i> 3	<i>b</i> 5	<i>b</i> 8	<i>c</i> 5	<i>c</i> 6	<i>c</i> 7	<i>c</i> 8	<i>a</i> 2	<i>a</i> 4	<i>a</i> 5	<i>a</i> 7

Replication I																							
<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>a</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>b</i>	<i>c</i>	<i>c</i>	<i>c</i>	<i>c</i>
1	3	5	7	1	4	5	8	1	3	5	8	2	4	6	8	2	3	6	7	2	4	6	7

Replication II																							
<i>a</i> 1	<i>a</i> 3	<i>a</i> 6	<i>a</i> 8	<i>b</i> 1	<i>b</i> 4	<i>b</i> 6	<i>b</i> 7	<i>c</i> 1	<i>c</i> 2	<i>c</i> 3	<i>c</i> 4	<i>a</i> 2	<i>a</i> 4	<i>a</i> 5	<i>a</i> 7	<i>b</i> 2	<i>b</i> 3	<i>b</i> 5	<i>b</i> 8	<i>c</i> 5	<i>c</i> 6	<i>c</i> 7	<i>c</i> 8

Next the groups (pairs of blocks) are randomized within each replication. The following layout is generated:

Replication I																							
<i>a</i> 2	<i>a</i> 4	<i>a</i> 6	<i>a</i> 8	<i>b</i> 2	<i>b</i> 3	<i>b</i> 6	<i>b</i> 7	<i>c</i> 2	<i>c</i> 4	<i>c</i> 6	<i>c</i> 7	<i>a</i> 1	<i>a</i> 3	<i>a</i> 5	<i>a</i> 7	<i>b</i> 1	<i>b</i> 4	<i>b</i> 5	<i>b</i> 8	<i>c</i> 1	<i>c</i> 3	<i>c</i> 5	<i>c</i> 8

Replication II																							
<i>a</i> 1	<i>a</i> 3	<i>a</i> 6	<i>a</i> 8	<i>b</i> 1	<i>b</i> 4	<i>b</i> 6	<i>b</i> 7	<i>c</i> 1	<i>c</i> 2	<i>c</i> 3	<i>c</i> 4	<i>a</i> 2	<i>a</i> 4	<i>a</i> 5	<i>a</i> 7	<i>b</i> 2	<i>b</i> 3	<i>b</i> 5	<i>b</i> 8	<i>c</i> 5	<i>c</i> 6	<i>c</i> 7	<i>c</i> 8

Replication I			Replication II			Replication III		
Block 1	Block 2	Block 3	Block 1	Block 2	Block 3	Block 1	Block 2	Block 3
$a_0 b_0 c_0$	$a_0 b_0 c_2$	$a_0 b_0 c_1$	$a_0 b_0 c_0$	$a_0 b_0 c_1$	$a_0 b_0 c_2$	$a_0 b_0 c_2$	$a_0 b_0 c_1$	$a_0 b_0 c_0$
$a_0 b_1 c_1$	$a_0 b_1 c_0$	$a_0 b_1 c_2$	$a_0 b_1 c_2$	$a_0 b_1 c_0$	$a_0 b_1 c_1$	$a_0 b_1 c_0$	$a_0 b_1 c_2$	$a_0 b_1 c_1$
$a_0 b_2 c_2$	$a_0 b_2 c_1$	$a_0 b_2 c_0$	$a_0 b_2 c_1$	$a_0 b_2 c_2$	$a_0 b_2 c_0$	$a_0 b_2 c_1$	$a_0 b_2 c_0$	$a_0 b_2 c_2$
$a_1 b_0 c_1$	$a_1 b_0 c_0$	$a_1 b_0 c_2$	$a_1 b_0 c_2$	$a_1 b_0 c_0$	$a_1 b_0 c_1$	$a_1 b_0 c_1$	$a_1 b_0 c_2$	$a_1 b_0 c_0$
$a_1 b_1 c_2$	$a_1 b_1 c_1$	$a_1 b_1 c_0$	$a_1 b_1 c_1$	$a_1 b_1 c_2$	$a_1 b_1 c_0$	$a_1 b_1 c_0$	$a_1 b_1 c_1$	$a_1 b_1 c_2$
$a_1 b_2 c_0$	$a_1 b_2 c_2$	$a_1 b_2 c_1$	$a_1 b_2 c_0$	$a_1 b_2 c_1$	$a_1 b_2 c_2$	$a_1 b_2 c_2$	$a_1 b_2 c_0$	$a_1 b_2 c_1$

All the main effects are estimated with full efficiency. The two-factor interaction efficiencies are F1F2 = 1.000, F1F3 = 1.000, F2F3 = 0.825.

interactions is, however, an important guiding factor in the choice of a design. From experimenter's consideration, main effects should be estimated with full efficiency or high efficiency (as close to one as possible). Considerable reduction in block size is possible through these designs. In some situations it is possible to get a design with smaller number of replications. It is also possible to carry out the combined analysis of the data generated through these experiments. Even the best performing treatment can be picked easily. These alternatives are valid even when the treatments in both or any crop also have a factorial structure.

Further, for the purpose of demonstrating the performance of all the treatments in the field one has to look for resolvable EGD or resolvable Kronecker Product type design for crop sequence experiments. Obviously for a resolvable block design with v treatments arranged in b blocks and replication of treatments r , it is possible to divide the b blocks into r groups containing $m = b/r$ blocks of size k each such that each group is a complete replication. The randomization of a resolvable design has already been explained above.

Continuing on the example discussed above, an alternative design is the following (resolvable) incomplete block design with factorial structure for $v = 18$ ($3^2 \cdot 2$), $b = 9$, $r = 3$, $k = 6$.

Gupta et al. (2011) gave a unified approach of obtaining resolvable incomplete block designs for factorial experiments. These designs have orthogonal factorial structure, have balance, estimate all main effects with full efficiency and have control over the interaction efficiencies. The parameters of the designs obtained are: number of treatments, $v = pq = f_1 \times f_2 \times \dots \times f_s \times h_1 \times h_2 \times \dots \times h_t$, $p = g_1 z$, $q = g_2 z$, $p > q$, number of blocks, $b = rz$, replication = r , block size = $g_1 g_2 z$, and number of blocks per replication = z . The method can generate designs for any number of replications. A catalogue of designs obtained through this method for number of levels of any factor at most 12 and three replications along with the layout of the

design is available at <http://iasri.res.in/design/factorial/factorial%20files/catalogues.htm> and maintained at IASRI, New Delhi. For this design the ANOVA will have the following form:

The ANOVA for the design in the Example will be the following:

ANOVA	
Source	DF
Treatments	17
Main Effect A	5
Sources of sulphur (A_1)	1
Doses of sulphur (A_2)	2
Interaction ($A_1 A_2$)	2
Main Effect B	2
Interaction AB	10
Interaction $A_1 B$	2
Interaction $A_2 B$	4
Interaction $A_1 A_2 B$	4
Blocks	8
Replications	2
Blocks (Replications)	6
Error	28
Total	53

Another Example

An experiment on Rice-Wheat sequence needs to be conducted. During the *kharif*-rice, 5 herbicidal treatments, denoted as a_0, a_1, a_2, a_3, a_4 are to be applied and on the following *rabi*-Wheat, 4 herbicidal treatments, combination of two levels of each of the two factors, denoted as $b_0 c_0, b_0 c_1, b_1 c_0, b_1 c_1$, are to be given. The purpose of the experiment is (i) to compare direct effects of *kharif* and *rabi* treatments, (ii) residual effects of *kharif* treatments, and (iii) interaction between residual effects of *Kharif* and direct effects of *Rabi* treatments.

The design recommended is a resolvable EGD design

ANOVA for resolvable design (YEAR I)

ANOVA	
Source	DF
Treatments	$pq - 1$
Main Effect A	$p - 1$
Main Effect B	$q - 1$
Interaction AB	$(p - 1)(q - 1)$
Blocks	$rz - 1$
Replications	$r - 1$
Blocks (Replications)	$r(z - 1)$
Error	$rpq - rz - pq + 1$
Total	$rpq - 1$

Kharif Season		Rabi Season	
Source	DF	Source	DF
Between Blocks	5	Between Blocks	5
<i>kharif</i> Treatment	4	Residual (<i>kharif</i> Treatments)	4
Error	50	Direct (<i>rabi</i> Treatments)	3
Total	59	Main Effect B	1
		Main Effect C	1
		Interaction BC	1
		Residual * Direct	12
		Interaction AB	4
		Interaction AC	4
		Interaction ABC	4
		Error	35
		Total	59

ANOVA for resolvable design (YEAR II)

Kharif Season		Rabi Season	
Source	DF	Source	DF
Between Blocks	5	Between Blocks	5
Residual (<i>rabi</i> Treatments)	3	Residual (<i>kharif</i> Treatments)	4
Main Effect B	1	Direct (<i>rabi</i>)	3
Main Effect C	1	Main Effect B	1
Interaction BC	1	Main Effect C	1
Direct (<i>kharif</i> Treatments)	4	Interaction BC	1
Residual * Direct	12	Residual * Direct	12
Interaction AB	4	Interaction AB	4
Interaction AC	4	Interaction AC	4
Interaction ABC	4	Interaction ABC	4
Error	35	Error	35
Total	59	Total	59

Treatment Chickpea (Safflower)		Safflower (Chickpea)
1	No Phosphorus	No Phosphorus
2	100% Recommended P	100% Recommended P
3	50% Recommended P	100% Recommended P
4	50% Recommended P	50% Recommended P
5	50% Recommended P+PSB	50% Recommended P+PSB
6	No Phosphorus	100% Recommended P
7	5 ton FYM/ha	100% Recommended P
8	PSB + 5 ton FYM/ha	100% Recommended P
9	100% Recommended P	50% Recommended P
10	100% Recommended P	No Phosphorus
11	100% Recommended P	5 ton FYM/ha
12	100% Recommended P	PSB + 5 ton FYM/ha

with $v = 5^2 2^2 (=20)$, $b = 6$, $r = 3$ and $k = 10$. This design has a factorial treatment structure with three factors, one at 5 levels and the other two at two levels each. In this design, the main effects are estimated with full efficiency. The efficiencies of the two factor interactions are: F1F2 = 1.0000, F1F3 = 0.9166, F2F3 = 0.9600.

Replication I		Replication II		Replication III	
Block 1	Block 2	Block 1	Block 2	Block 1	Block 2
$a_0 b_0 c_0$	$a_0 b_0 c_1$	$a_0 b_0 c_1$	$a_0 b_0 c_0$	$a_0 b_0 c_0$	$a_0 b_0 c_1$
$a_1 b_0 c_1$	$a_1 b_0 c_0$	$a_1 b_0 c_0$	$a_1 b_0 c_1$	$a_1 b_0 c_1$	$a_1 b_0 c_0$
$a_2 b_0 c_1$	$a_2 b_0 c_0$	$a_2 b_0 c_0$	$a_2 b_0 c_1$	$a_2 b_0 c_0$	$a_2 b_0 c_1$
$a_3 b_0 c_0$	$a_3 b_0 c_1$	$a_3 b_0 c_1$	$a_3 b_0 c_0$	$a_3 b_0 c_1$	$a_3 b_0 c_0$
$a_4 b_0 c_0$	$a_4 b_0 c_1$	$a_4 b_0 c_1$	$a_4 b_0 c_0$	$a_4 b_0 c_0$	$a_4 b_0 c_1$
$a_0 b_1 c_1$	$a_0 b_1 c_0$	$a_0 b_1 c_0$	$a_0 b_1 c_1$	$a_0 b_1 c_1$	$a_0 b_1 c_0$
$a_1 b_1 c_1$	$a_1 b_1 c_0$	$a_1 b_1 c_0$	$a_1 b_1 c_1$	$a_1 b_1 c_0$	$a_1 b_1 c_1$
$a_2 b_1 c_0$	$a_2 b_1 c_1$	$a_2 b_1 c_1$	$a_2 b_1 c_0$	$a_2 b_1 c_1$	$a_2 b_1 c_0$
$a_3 b_1 c_0$	$a_3 b_1 c_1$	$a_3 b_1 c_1$	$a_3 b_1 c_0$	$a_3 b_1 c_0$	$a_3 b_1 c_1$
$a_4 b_1 c_1$	$a_4 b_1 c_0$	$a_4 b_1 c_0$	$a_4 b_1 c_1$	$a_4 b_1 c_1$	$a_4 b_1 c_0$

The layout of the design is

The analysis of the data generated from such an experiment is:

If same randomization on same experimental area is followed in second year, then one can analyze the data as

The best treatment can be picked through the usual block design analysis.

Remark 1: Generally there is a tendency to report SE_m , the standard effort of the mean. In statistical terms this has no meaning because the mean is not estimable and its estimate is not unique. Instead, one should report SE_d , the standard error of the difference of two estimated treatment effects, because the pair wise difference between any two treatment effects is estimable and so its estimate is unique.

Remark 2: Generally we work out the critical difference for testing the significance of the pair wise treatment differences. The null hypothesis tested in such cases is that any two treatments have same effect. But when the treatment structure is factorial in nature and the treatments are all possible combinations of the levels of m factors, one has to be careful in using the critical difference. When we talk about the critical difference for a factorial main effect, we actually mean the critical difference for testing the pair wise difference between any two levels of that factor averaged over levels of remaining $m - 1$ factors. Similarly, the critical difference for the interaction effect involving any f factors means the critical difference for testing the pair wise difference between the treatments combinations of any two levels of those f factors averaged over levels of remaining $m - f$ factors.

Category II experiments

Some experiments are conducted to develop suitable integrated nutrient supply system of a crop sequence or crop rotation. One set of treatments is applied to the *kharif* crop and another set of treatments is applied to the *rabi* crop. In these experiments, the treatment combinations are smaller than pq (fractional factorial) and the estimation of interactions between the residual effects of *kharif* treatments and the direct effects of *rabi* treatments is not of interest to the experimenter. Another experimental situa-

tion, in which the treatment combinations are smaller than mn is dryland farming where the *kharif* is generally left as fallow and experiments on crop rotations are conducted: One Crop for the 1st Year (*rabi* season) and Another *rabi* Crop for the second year. Due to fallow *kharif* season, it is expected that direct effect of treatments applied in second crop year do not interact with the residual effect of treatments applied in first year.

Category II Experiment: An example

An experiment was conducted on Oilseeds (Safflower) with an objective to find out the better method of phosphorus (P) management in safflower based cropping system to increase P-use efficiency. This is a crop rotation experiment with one crop for each year in *rabi* season. The rotation is Chickpea - Safflower (Series-I) and Safflower - Chickpea (Series-II). Since it takes two years to complete one cycle, the experiment was conducted in two series so that at the end of two years, we have 2 cycles of data. RCB design was used with the following 12 treatments and 2 replications.

The data generated on the crop sequences is bivariate and is converted into an economic index (gross returns, net returns), energy equivalent or protein equivalent, etc. Converted data is analyzed as per RCB design procedure. This will give the cumulative effect of the treatments applied in two seasons.

The experimenter, however, is interested to compare the direct effects of treatments applied to first and second crops respectively and residual effect of the treatments applied to first crop using the data from second crop. This is not answered by converting the bivariate data into an economic index and analyzing it as RCB design.

From the treatment structure of the treatment combinations in the sequence, it is seen that there are five distinct treatments in each crop. But instead of 36 possible combinations, only 12 distinct treatment combinations are tried in the sequence. The 6 distinct treatments for individual crops are :

Phosphorus	Distinct Treatments	Repetitions in <i>kharif</i> crop	Repetitions in <i>rabi</i> crop
No Phosphorus	: T_1	2	2
50% Recommended P	: T_2	2	2
100% Recommended P	: T_3	5	5
50 % Recommended P+ PSB	: T_4	1	1
5 ton FYM/ha	: T_5	1	1
PSB + 5 ton FYM/ha	: T_6	1	1

At the first instance the data for both *rabi* and *kharif* crop may be analyzed by general block design ANOVA for 6 treatments run in 2 replications (blocks) of size 12 each

such that the i^{th} treatment is replicated r_{ij} times in the j^{th} block, with r_{ij} as given in Table above.

The experimenter is interested in comparing the direct effects of treatments applied to first and second crops respectively and residual effect of the treatments applied to first crop using the data from second crop.

To test the equality of residual effects of phosphorus treatments of chickpea in safflower in series-I and phosphorus treatments of safflower on chickpea in series-II, and the direct effects of phosphorus treatments applied to safflower in Series-I and phosphorus treatments applied to chickpea in Series-II, the treatments can be divided into two sets viz.,

	First Set	Second Set
Series-I	Chickpea	Safflower
Series-II	Safflower	Chickpea

Consider a row-column design with the replications of the original design as rows and the treatments applied to the Chickpea crop (first set of treatments) as the columns. The cells of the array have the treatments applied to the Safflower (second set of treatments). According to this structure, we get the following table from the Table of treatment combinations in the sequence:

Series-I	First Set					
	1	2	3	4	5	6
Replication - I	1, 3	3, 2	3, 2, 1, 5, 6	4	3	3
Replication - II	1, 3	3, 2	3, 2, 1, 5, 6	4	3	3

The treatments listed in the cells are the second set of treatments. The residual effects of the treatments applied to the first crop are the effects of first set of treatments and direct effects of phosphorus treatments applied to second crop are the effects of second set of treatments. The above hypothesis can be tested using the model:

$Y = \text{General Mean} + \text{Replication effect} + \text{effect of First Set of Treatments} + \text{effect of Second Set of Treatments} + \text{Error}.$

The coefficient matrix of reduced normal equations for first set of treatments (C_f) and the coefficient matrix of reduced normal equations for second set of treatments (C_s) are obtained as:

$$C_f = C_s = \begin{bmatrix} 2.6 & -0.4 & -1.4 & 0.0 & -0.4 & -0.4 \\ -0.4 & 2.6 & -1.4 & 0.0 & -0.4 & -0.4 \\ -1.4 & -1.4 & 3.6 & 0.0 & -0.4 & -0.4 \\ 0.0 & 0.0 & 0.0 & 0.0 & 0.0 & 0.0 \\ -0.4 & -0.4 & -0.4 & 0.0 & 1.6 & -0.4 \\ -0.4 & -0.4 & -0.4 & 0.0 & -0.4 & 1.6 \end{bmatrix}$$

The design adopted is disconnected in first set of treatments and also in second set of treatments. Similar phenomenon was observed for Series -II. Therefore, it is not possible to estimate all the possible paired comparisons of first set and second set of treatments. This implies that the treatment structure (the fractional factorial) chosen for experimentation is not proper. So far answering the questions for which the experiment was conducted, the choice of fraction is not proper. In the above set up, however, if the combination of 50% P + PSB and 50% P + PSB is removed from both the replications and instead 50% P + PSB and 5 ton FYM/ha is given in replication I and 5 ton FYM/ha and 50% P + PSB is given in replication II, then the design becomes treatment connected and it would be possible to estimate all the effects. In this way the total number of treatments and the total number of observations generated from the design remains the same. For the resulting design we get the following table from the Table of treatment combinations in the sequence:

Series-I	First Set					
	1	2	3	4	5	6
Replication - I	1, 3	3, 2	3, 2, 1, 5, 6	5	3	3
Replication - II	1, 3	3, 2	3, 2, 1, 5, 6	-	3, 4	3

It can easily be seen that the design is connected for first set of treatments as well as for second set of treatments.

DISCUSSION

In this article the problem of designing experiments for two different types of crop-sequence experiments has been discussed. These experiments were being run as split plot design or RCB design, but in the later case the questions to be answered were not being answered. It has been shown that designing crop-sequence experiments is equivalent to designing a factorial experiment run in an incomplete block design. The two types of experiments differ in the sense that in one case all the possible treatment combinations are run and the interaction between the residual effect of first set and the direct effect of second set of treatments is of interest to the experimenter. In the second case a fraction of the treatment combinations is tried

and the interaction is not of interest to the experimenter. It has been shown that block designs with orthogonal factorial structure and balance are useful and appropriate for category I experiments. EGD designs and many Kronecker Product type designs have been found to be extensively useful. The choice of a proper design, however, is important because one would like the main effects to be estimated with full efficiency or maximum efficiency and the efficiency of the interaction should also be high.

For category II experiments the choice of a proper fraction is very important. As seen an improper choice of fraction would lead to a disconnected design and all the questions of the experimenter cannot be answered. There is a connection between row-column designs and block designs with two sets of treatments applied in succession. This connection in fact unifies the research efforts made in the literature in two different directions viz. row-column designs and block designs for two sets of treatments applied in succession. Now with the help of a simple mapping one can easily obtain a block design for two sets of treatments applied in succession from a row column design and vice versa (see e.g. Parsad *et al.*, 2003).

These designs have been used for experimentation in Division of Agronomy, IARI, New Delhi, PDCSR, Modipuram, AICRP on Oilseeds, Solapur, AICRP on CSR, Indore, PAU, Ludhiana, etc.

For category I experiments, catalogue of efficient resolvable incomplete block designs with orthogonal factorial structure and with full efficiency on main effects and run in three replications has been prepared and is available at design resources server hosted at www.iasri.res.in/design/ for the benefit of the experimenters. The catalogue provides a randomized layout of the design. The steps of analysis of data generated would also be available at the Design resources server. The layout of designs for Category II experiments and for incomplete split-plot design would soon be made available at design resources server.

REFERENCES

- Gupta V.K., Nigam A.K., Parsad Rajender, Bhar, L.M. and Behera Subrat Keshori. 2011. Resolvable block designs for factorial experiments with full main effects efficiency. *Journal of the Indian Society of Agricultural Statistics* **65**(3): 303–15.
- Gupta V.K. and Parsad Rajender 2009. Block designs for cropping systems research. *NAAS News* **8**(2): 3–7.
- Hering F. and Mejza S. 2002. An incomplete split-plot design generated by GDPBIBD (2)s. *Journal of Statistical Planning and Inference* **106**: 125–34.
- Kristensen Kristian 2012. Incomplete split plot designs based on $\hat{\alpha}$ -designs: a compromise between traditional split-plot designs and randomized complete block design. *Euphytica* **183**: 401–13.
- Mejza I. and Mejza S. 1984. Incomplete split-plot designs. *Statistics and Probability Letters* **2**: 327–32.
- Parsad Rajender, Gupta, V.K., Batra P.K., Satpati S.K. and Biswas Pabitra. 2007. *Monograph on $\hat{\alpha}$ -designs*. Indian Agricultural Statistics Research Institute, New Delhi.
- Parsad, Rajender, Gupta, V.K. and Prasad, N.S.G. 2003. Structurally incomplete row-column designs. *Commun. Statist.-Theory &*