



ARTICLE

Optimal designs in plant breeding experiments: a simulation study using wheat pedigree matrix

 Paula Ribeiro Santos Gouveia* and  Renata Alcarde Sermarini ¹

Department of Exact Sciences, Luiz de Queiroz College of Agriculture (ESALQ), University of Sao Paulo, Piracicaba-SP, Brazil.

*Corresponding author. Email: paulas@usp.br

(Received: August 12, 2024; Revised: October 31, 2024; Accepted: December 19, 2024; Published: May 30, 2025)

Abstract

Plant breeding programs involve the selection of new superior lines. However, for a large number of test lines, there are several limitations in the use of certain designs. Therefore, the success of these programs depends on an adequate experimental design that allows obtaining accurate estimates of genetic effects, increasing the efficiency of the experiment and controlling experimental variability. In addition, considering the dependence between genetic effects is desirable to ensure the validity and generalization of results, avoiding biased estimates and incorrect interpretations. For this purpose, using partially replicated designs (p -rep), in which a percentage, p , of test lines are replicated and the others not, can be a good option. Thus, a simulation study was conducted to evaluate designs for early phase wheat breeding experiments according to the optimality criterion C , considering the dependence or independence between test lines, comparing them in relation to the realized genetic gain and, consequently, the quality of the material selection, for a given experimental area and for $p = 20\%$, for different genetic variance values. It could be concluded that the differences between designs are small, and that they are more affected by the magnitude of the genetic variance assumed for data.

Keywords: Optimality criterion C ; Spatially optimized designs; Partially replicated design; Genetic gain; Quality of genetic selection; Relationship matrix A .

1. Introduction

In the early stages of plant breeding programs, trials allow the testing of a large number of lines, or test lines, with the aim of detecting superior ones. Such trials are generally characterized by limited resources, either in terms of genetic material or restrictions regarding the experimental area, which results in the use of unreplicated designs (Kempton, 1984).

According to Kempton (1984), initially, the designs most commonly used in these trials were grid-plot designs. These are row-column designs that include replicates of the standard varieties systematically occupying the check plots in a regular grid, and the unreplicated test lines. Standard varieties are usually varieties already present in the market and with desirable characteristics, which are used as a reference to evaluate the yields of the test lines and also as a way of dealing with possible trends in soil fertility.

Seeking to improve these tests, Federer (1956) proposed augmented designs. Such designs present the same two treatments groups, those that will be repeated, the standard varieties, used to explain part of the spatial variability, and unreplicated ones, or test lines. The augmented block design consists of distributing the replications of varieties in blocks and then, augmenting each block with unreplicated test lines. Others augmented designs were proposed by Federer (1961), Federer (2002), and Federer (2005), Federer & Raghavarao (1975), Federer *et al.* (1975), Lin & Poushinsky (1983) and Lin & Poushinsky (1985) and Federer & Crossa (2001).

Cullis *et al.* (2006) proposed an alternative to grid-plot and augmented designs, the partially replicated (p -rep) design. This new class of designs assumes replication for p percent of test lines, which totally or partially replace check plots, and the remaining test lines are not repeated. Originally, p -rep designs were obtained in such a way to achieve maximum genetic gain, a characteristic in common with optimal designs specifically designed to maximize the precision of estimates of the effects of interest. When working with optimal designs, some assumptions are made, since the following are necessary: the design model, the values of parameters and a search criterion (Shah & Sinha, 1989).

The grid-plot and p -rep designs were compared by some authors. Cullis *et al.* (2006) showed that p -rep designs improved the precision of line selection. In the work, the p -rep designs evaluated are spatially optimized (C-optimal), that is, given the model and its parameters, the design found is the one that presents the lowest average variance for simple contrasts of the effects of test lines.

Clarke & Stefanova (2011) carried out this comparison in uniformity trials. Moehring *et al.* (2014) simulated the genetic effects and allocated them according to four designs, two of which were grid-plot and p -rep, to represent triticales yields. Santos (2017) compared these designs in a specific study with sugarcane. Goes (2020) also compared these types of designs through simulation studies. It is worth mentioning that, in these works, p -rep designs presented better results for selecting superior test lines. Furthermore, these designs are widely used in most plant breeding programs in Australia (Cullis *et al.*, 2020) and are being introduced in Brazil.

There are several measures to evaluate and compare designs, two of which are the realized genetic gain (RGG), which is the ratio between the mean of the s superior EBLUPs (empirical best linear unbiased predictions) and the mean of the s superior true genetic effects, where s corresponds to the percentage of selection of the genetic material, which in this work was defined as $s = 15\%$ (Cullis *et al.*, 2006; Smith *et al.*, 2006; Santos, 2017; Goes, 2020; Sermarini *et al.*, 2020), and the selection success, defined by the percentage of truly superior test lines selected from EBLUPs (Sermarini *et al.*, 2020). Both the RGG and the selection success depend on the true genetic effects, so they can only be obtained through simulation studies.

There are other comparison measures that do not depend on simulation studies. For example, Müller *et al.* (2010) used the mean of the empirical variance, a measure that depends on the number of genotypes, including test lines and standard varieties, when present, and on the sum of the differences in the adjusted means of the genotype. Another proposal is the SE ratio, which is the ratio between the standard error of the comparison between test lines and checks, calculated from the design under study and also from a completely randomized design (Clarke & Stefanova, 2011). Piepho & Williams (2016) evaluated the use of relative efficiency, which is the harmonic mean of non-zero eigenvalues of the information matrix of test lines.

The p -rep designs are usually obtained by ignoring information about genetic relationship (pedi-

gree), although in plant breeding programs, it is often possible to access some information about the degree of genetic relationship of individuals included in the trial (Bueno Filho & Gilmour, 2003; Butler *et al.*, 2014). Therefore, given that the relatedness information is known, that is, information about the parents of each individual, it is possible to consider it when obtaining an optimal p -rep design.

Few authors use the kinship matrix $\mathbf{A}$ in the search for optimal designs. Bueno Filho & Gilmour (2003), considered correlated genetic effects for unresolvable incomplete block designs of size four to six treatments. The genetic effects in the design were analyzed in three situations of $\mathbf{A}$: $\mathbf{A}_1 = \frac{1}{2}\mathbf{Id}$, where $\mathbf{Id}$ denotes the identity matrix, $\mathbf{A}_2$ is a simple family structure, i.e., three families of two siblings (half-siblings) and $\mathbf{A}_3$ is a complex kinship, i.e., two treatments are related if they have a common parent. In all situations, the $\mathbf{A}$ -optimality criterion was used for which the search is made for designs with the minimum sum of variances of treatment effects estimates.

In another study, Bueno Filho & Gilmour (2007) investigated the effect of varying degrees of uncertainty in point estimates of a range of parameter settings of additive genetic variance expressed in terms of heritability on design selection. The authors reported that in some situations in plant breeding programs, it is very likely to have a sparse relationship matrix, $\mathbf{A}$, or discrete family structures. For some of these situations it is possible to find optimal designs that are robust to assumptions about heritability. However, for complex kinship structures, the choice of design can change drastically depending on previous point estimates.

Butler *et al.* (2014) studied three cases: variety selection in canola, estimation of the crossing value in sorghum genetic selection (hybrid) and estimation of the breeding value for forest improvement, with different selection objectives, genetic complexity and scale. They extend the work developed by Bueno Filho & Gilmour (2003) and also by Cullis *et al.* (2006). The extension concerns the specifications of the linear model, both in terms of genetic and nongenetic components.

Cullis *et al.* (2020) evaluated simulation studies for early stages, S1 and S2, based on a plant breeding case study, with the aim of reinforcing the importance of including information on genetic relationship. Selection in advanced stages in genetic breeding generally occurs sequentially. These stages are called S1, S2, S3 and S4. In the study, they evaluated 256 test lines from stage S1 with the aim of making selections for progression to test stage S2. The study also included four standard varieties. To this end, a mixed linear model for genetic and nongenetic effects in a p -rep design was adopted. Thus, the authors considered different p values and different values for the genetic and residual variance ratio. For different p values, the field scenarios varied the number of lines, between 22 and 28 lines, with a fixed number of 12 columns.

dos Santos (2023) investigated the effectiveness of spatially optimized designs in multienviromental trials to select the best test lines in plant breeding grain yield, regarding genetic gain and quality of genetic material selection. For this, a simulation study compared the grid-plot and p -rep designs. For the p -rep design, the p percentages of duplicated lines were varied for $p = 11\%$, 22% and 33% , and number of standard varieties (0, 5, 10, 15 and 20). The analysis was performed using mixed linear models considering joint and individual analyses, which incorporated spatial variation in plot errors. In both designs, the assumption of dependence and independence between line effects was considered.

It is observed that studies that evaluate p -rep designs in such a way to consider the relationship between test lines are limited. In this context, the objective of this study is to compare the p -rep designs with and without the inclusion of relationship matrix regarding the realized genetic gain and, consequently, the quality of the selection of the genetic material to be considered for a next phase of the plant breeding study for a given experiment size, based on linear mixed models and simulation studies. Here, spatially optimized p -rep designs were generated for an experimental area composed of 36 rows by 20 columns, assuming 598 test lines, with $p = 20\%$ and a selection percentage $s = 15\%$. In the design model, random treatment effects were assumed, different parameter values

for the genetic variance and independence (**Id**) or dependence (**A**) between treatment effects.

2. Materials and Methods

In this study, a dataset was used as a motivating example. This section presents this dataset, the designs and how they are obtained, the models for simulating data and the measures used to compare designs. The procedures used in this study are similar to those presented by Santos, 2017; Goes, 2020; Sermarini *et al.*, 2020; dos Santos *et al.*, 2024.

2.1 Material

The motivating example was a set of 598 wheat lines developed by CIMMYT (CIMMYT, 2021), the International Maize and Wheat Improvement Center, available in the package `lme4GS` (Perez-Rodriguez, 2021). The CIMMYT breeding program has conducted numerous international trials in a wide variety of wheat production environments. The environments represented in these trials were grouped into four basic sets of environments, comprising four main agroclimatic regions previously defined and widely used by CIMMYT. The dataset includes information on average grain yield and parentage, among others.

The file `wheat.Pheno` contains four columns: environments (`env`), replicates (`rep`), genotype identifiers (`GID`), and grain yields (`Yield`). The file `wheat.Pedigree` contains three columns, `gp1d1` and `gp1d2` correspond to the GIDs of parents 1 and 2, respectively, and `progeny`, which correspond to the GIDs of the progeny. Finally, the file `wheat.X` contains a matrix of dimensions 598×1279 , which corresponds to the Diversity Array Technology (DArT) markers encoded as 0 and 1.

2.2 Generated designs

The configuration for generating the designs in this work was a rectangular experimental area with 36 rows (n_r) by 20 columns (n_c), totaling 720 plots. A total of 598 test lines (n_t) were considered, as in the motivating example, in a p -rep design, with $p = 20\%$. Thus, 122 test lines were replicated and 476 were not.

The linear mixed model for generating p -rep designs was based on Gilmour *et al.* (1997):

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}_{gd}\mathbf{u}_{gd} + \mathbf{Z}_o\mathbf{u}_o + \boldsymbol{\epsilon}, \quad (1)$$

where $\mathbf{y}_{n \times 1}$ is the vector of the response variable being $n = n_r \times n_c$; $\boldsymbol{\beta}_{(q \times 1)}$ is the vector of fixed effects with design matrix $\mathbf{X}_{(n \times q)}$; $\mathbf{u}_{gd(n_t \times 1)}$ is the vector of random genetic effects with design matrix $\mathbf{Z}_{gd(n \times n_t)}$; $\mathbf{u}_{o((n_r+n_c) \times 1)}$ is the vector of nongenetic random effects with design matrix $\mathbf{Z}_o(n \times (n_r+n_c))$, and $\boldsymbol{\epsilon}_{n \times 1}$ is the vector of random errors. It was assumed that $(\mathbf{u}_o, \mathbf{u}_{gd}, \boldsymbol{\epsilon})$ are independent with Gaussian joint distribution with zero mean and variance-covariance matrix:

$$\sigma^2 \begin{bmatrix} \mathbf{D}_o(\gamma_o) & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \mathbf{D}_g(\gamma_{gd}) & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{R}(\boldsymbol{\phi}) \end{bmatrix}, \quad (2)$$

such that $\gamma = \sigma^2/\sigma^2$, γ_o contains the variance parameters for the row and column effects. Additionally, we assume $\mathbf{R}(\boldsymbol{\phi}) = \boldsymbol{\Sigma}_r(\boldsymbol{\phi}_r) \otimes \boldsymbol{\Sigma}_c(\boldsymbol{\phi}_c)$, where $\boldsymbol{\phi}_r$ and $\boldsymbol{\phi}_c$ are the correlation parameters for rows and columns, respectively, that characterize first-order separable autoregressive processes.

To generate the designs, $\sigma^2 = 1$, $\gamma_r = \gamma_c = 0.1$, $\boldsymbol{\phi}_r = \boldsymbol{\phi}_c = 0.75$ and four γ_{gd} values (0.5, 1, 2 and 3). The values adopted for the nongenetic parameters were assumed based on the works of Goes, 2020 and Sermarini *et al.*, 2020, which reported estimates for similar experimental conditions in different crops. These parameters reflect environment aspects, such as spatial correlation and dependence,

and are not directly associated with genetics or a specific crop. The designs were generated under two assumptions: the first in which the test lines were assumed unrelated, i.e., $\mathbf{D}_g = \gamma_{gd} \mathbf{I}_d$, and the second in which they were assumed related and $\mathbf{D}_g = \gamma_{gd} \mathbf{A}$, in which $\mathbf{A}$ is the pedigree matrix, totaling eight evaluated designs.

The pedigree matrix, $\mathbf{A}$, was obtained in R (R Core Team, 2021) using the *pedigree* package (Bates & Vazquez, 2014). The process began by assuming the $\mathbf{A}$ matrix to be an identity matrix, where all off-diagonal elements are zero and on-diagonal elements are one, representing that each individual is completely related to itself. For each individual, a_{ij} is calculated in the pedigree, checking whether it has information about its father (f) and mother (m). Based on this, the additive kinship coefficients are calculated according to Henderson (1976), where

$$a_{ij} = \begin{cases} 1 + 0.5a_{fm}, & \text{if } i = j \\ 0.5(a_{if} + a_{im}), & \text{if } i \neq j \end{cases}.$$

The $\mathbf{A}$ matrix is updated and this process continues until all individuals in the pedigree have been processed. The $\mathbf{A}$ matrix is symmetric, since the coefficient of kinship between i and j is the same as that between j and i .

The optimality criterion adopted was the C -optimal, which seeks a design with minimum average variance of pairwise differences of test lines effects (AVPD), so that,

$$\text{AVPD} = \frac{2}{n_t - 1} \left(\text{tr}(\mathbf{C}^{-1}) - \frac{1}{n_t} \mathbf{1}^T \mathbf{C}^{-1} \mathbf{1} \right),$$

where $\text{tr}()$ represents the trace of the matrix and $\mathbf{C}$ is the information matrix of test lines, which for the proposed model is given, according to Hooks *et al.* (2009), by:

$$\mathbf{C} = \mathbf{Z}_o^T (\mathbf{Z}_o \mathbf{D}_o (\sigma_o)^2 \mathbf{Z}_o^T + \mathbf{R}(\phi))^{-1} \mathbf{Z}_{gd} + \mathbf{D}_g^{-1} (\sigma_{gd})^2 - \mathbf{Z}_{gd}^T (\mathbf{Z}_o \mathbf{D}_o (\sigma_o)^2 \mathbf{Z}_o^T + \mathbf{R}(\phi))^{-1} \mathbf{X} (\mathbf{X}^T (\mathbf{Z}_o \mathbf{D}_o (\sigma_o)^2 \mathbf{Z}_o^T + \mathbf{R}(\phi))^{-1} \mathbf{X})^{-1} \mathbf{X}^T (\mathbf{Z}_o \mathbf{D}_o (\sigma_o)^2 \mathbf{Z}_o^T + \mathbf{R}(\phi))^{-1} \mathbf{Z}_{gd}.$$

One thousand moves were adopted for the design search. The package used was *odw* (Butler, 2022) for the R software (R Core Team, 2021).

2.3 Simulation study

After obtaining the designs, a simulation study was carried out to obtain data. Data were simulated for 32 scenarios setting as presented in the previous Section, with the same definitions, adding only the genomic relationship structure. Thus, the linear mixed model for data generation was based on Perez-Rodriguez (2021):

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}_{ga}\mathbf{u}_{ga} + \mathbf{Z}_{gg}\mathbf{u}_{gg} + \mathbf{Z}_o\mathbf{u}_o + \boldsymbol{\epsilon}, \quad (3)$$

where $\mathbf{y}_{n \times 1}$ is the vector of the response variable being $n = n_r \times n_c$; $\boldsymbol{\beta}_{(q \times 1)}$ is the vector of fixed effects with the design matrix $\mathbf{X}_{(n \times q)}$; $\mathbf{u}_{ga(n_t \times 1)} \sim \text{MN}(\mathbf{0}, \sigma_{ga}^2 \mathbf{A})$ is the vector of random additive genetic effects with design matrix $\mathbf{Z}_{ga(n \times n_t)}$, $\mathbf{A}$ is the pedigree matrix described in Section 2.2, σ_{ga}^2 is the associated variance parameter; $\mathbf{u}_{gg(n_t \times 1)} \sim \text{MN}(\mathbf{0}, \sigma_{gg}^2 \mathbf{G})$ is the vector of random genomic effects with design matrix $\mathbf{Z}_{gg(n \times n_t)}$, $\mathbf{G} = \mathbf{W}\mathbf{W}'/k$ is a genomic relationship matrix, $\mathbf{W}$ is the centered and standardized marker matrix, k is the variance parameter associated with the number of markers, σ_{gg}^2 is the variance parameter associated with the markers.

The scenarios evaluated are the combinations of eight experimental designs and four values for $\gamma_{ga} = \gamma_{gg}$ (0.5, 1, 2 and 3), as particular cases for the genetic variances. In all cases, $\sigma^2 = 1$,

$\gamma_r = \gamma_c = 0.1$ and $\phi_r = \phi_c = 0.75$ and $\gamma_{ga}\mathbf{A} + \gamma_{gg}\mathbf{G}$) were assumed. As mentioned in Section 2.2, to generate data, the values for the adopted nongenetic parameters are also obtained based on Goes, 2020 and Sermarini *et al.*, 2020.

Independent samples were generated for the genetic and nongenetic effects and for the errors, using the standard normal distribution. In this way, 1,000 vectors of size 598 were obtained for the genetic effects, 1,000 vectors of size 56 (36+20) for the nongenetic effects of rows and columns, and 1,000 vectors of size 720 (36x20) for the error, ϵ . The appropriate transformation for each vector was used, so that the variances for the genetic, nongenetic, and residual effects followed the desired assumption.

After data simulation, the models were fitted. Data were analyzed following the model analogous to that presented in Equation 3, with the same definitions. For this, the `asreml` (The VSNi Team, 2023) package for the R software (R Core Team, 2021) was used. For each set of simulated data, the following were recorded: information on the algorithm convergence, the genetic EBLUPs, to calculate the comparison measures and the estimates of the variance and correlation components.

2.4 Evaluation of designs

To compare the eight designs in each of the evaluated scenarios, RGG and selection success measures were calculated. RGG is used to evaluate the genetic gain achieved in each design. For each scenario of the generated designs, the EBLUPs of test lines are recorded and the RGG is calculated as the ratio between the average of the best s EBLUPs and the average of the best s of the true genetic effects, where s is the selection percentage, assumed to be 15%, a value commonly adopted by Brazilian plant breeding programs (Goes, 2020; Sermarini *et al.*, 2020; dos Santos *et al.*, 2024). Another measure calculated was the selection success (selection probability for $s = 15\%$), defined by the percentage of truly superior test lines selected from EBLUPs (Sermarini *et al.*, 2020).

3. Results and Discussion

In Figure 1, the layouts for the p -rep designs are presented, considering the four genetic variance values ($\gamma_{gd} = 0.5, 1, 2$ and 3), and the variance and covariance matrix assumed for the genetic effects, considering dependence (on the right) and independence (on the left). In general, for all designs considered, treatments are randomly distributed in the experimental area; however, some highlights are made. Regardless of the assumption about the test lines effects relationship, it was observed that the greater the genetic variance, the lower the frequency of clustered replicated test lines, that is, different replicated test lines occupy a smaller number of neighboring plots, and of these, occupying plots on the border of the experimental region. It is noted that, when assuming the genetic variance to be twice the residual variance ($\gamma_{gd} = 2$), the lowest number of clustered replicates was observed, the lowest number of replicates with at least one replicate occupying its neighboring plot and the lowest number of replicates occupying plots on the border of the experimental region.

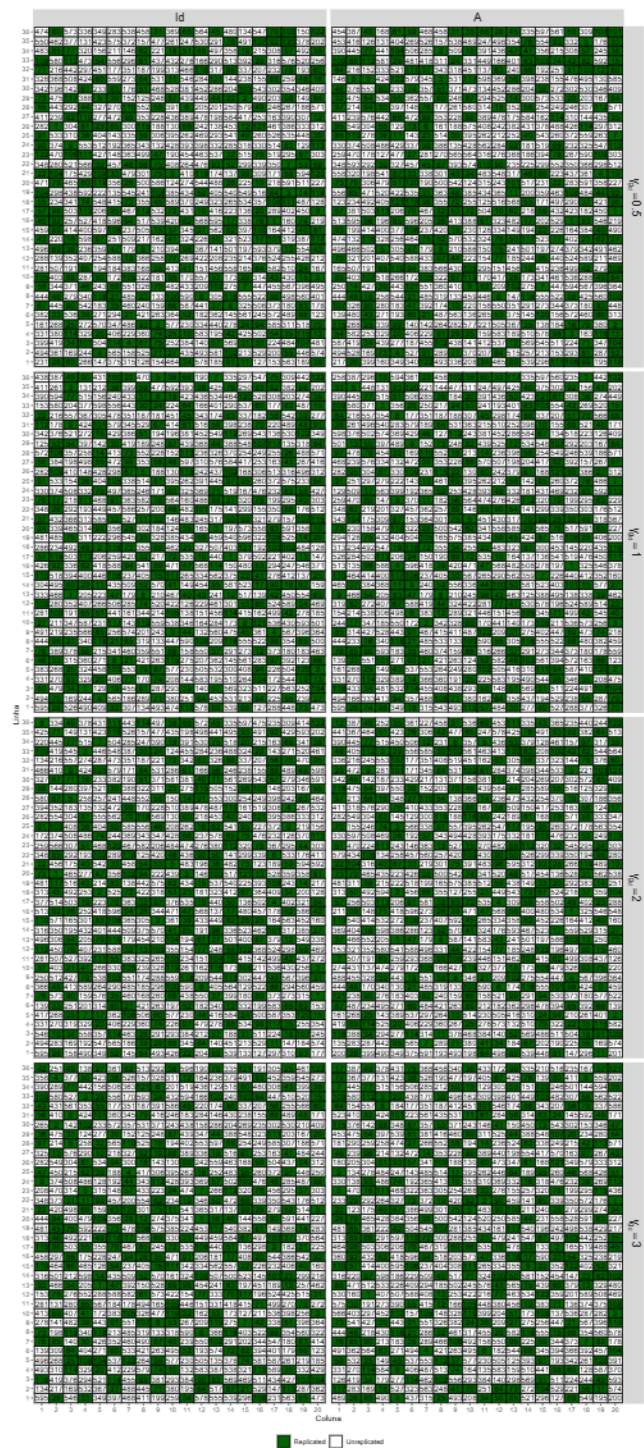


Figure 1. Layouts for the p -rep design generated considering the kinship matrix A (on the right) and for the p -rep design generated without considering the kinship matrix (on the left), for an experimental area of 36 rows by 20 columns and 598 test lines, with $p = 20\%$.

Regarding simulation results, algorithm convergence was observed for all 1,000 fitted models in each of the scenarios under the eight different designs.

In Table 1, when generating designs with the genetic variance being half of the residual variance, for the cases evaluated, it was observed that the average RGG is higher and more accurate when using spatially optimized designs in which the effects of test lines are assumed to be independent, not exceeding 0.3% when compared to the designs for which the matrix $\mathbf{A}$ was considered. Interestingly, the same was observed when $\gamma_{gd} = 2$. However, when the genetic and residual variances are equal ($\gamma_{gd} = 1$) and when the former is three times the latter ($\gamma_{gd} = 3$), the spatially optimized designs with $\mathbf{D}_g(\gamma_{gd}) = \gamma_{gd}\mathbf{A}$ showed better performance, but not exceeding 0.3% and 0.1% for the RGG, on average, for $\gamma_{gd} = 1$ and $\gamma_{gd} = 3$, respectively. Similar behavior was observed when comparing the designs in relation to selection success. However, for this measure, greater difference was observed between designs in relation to the assumption of genetic effects in relation to the precision of the selection success (Table 2).

Table 1. Mean and standard deviation of realized genetic gain (RGG) for designs generated with $\gamma_{gd} = 0.5, 1, 2, 3$ and assumed genetic variance for data, $\gamma_{ga} = \gamma_{gg} = 0.5, 1, 2, 3$, considering the kinship matrix, $\mathbf{D}_g = \gamma_{gd}\mathbf{A}$, and without considering the kinship matrix, $\mathbf{D}_g = \gamma_{gd}\mathbf{Id}$, for selection percentage $s = 15\%$.

	$\gamma_{gd} = 0.5$			
	$\gamma_{ga} = \gamma_{gg} = 0.5$	$\gamma_{ga} = \gamma_{gg} = 1$	$\gamma_{ga} = \gamma_{gg} = 2$	$\gamma_{ga} = \gamma_{gg} = 3$
$\mathbf{D}_g = \gamma_{gd}\mathbf{Id}$	1.513 (0.315)	2.239 (0.432)	3.269 (0.605)	4.053 (0.732)
$\mathbf{D}_g = \gamma_{gd}\mathbf{A}$	1.511 (0.318)	2.233 (0.439)	3.264 (0.610)	4.048 (0.736)
	$\gamma_{gd} = 1$			
	$\gamma_{ga} = \gamma_{gg} = 0.5$	$\gamma_{ga} = \gamma_{gg} = 1$	$\gamma_{ga} = \gamma_{gg} = 2$	$\gamma_{ga} = \gamma_{gg} = 3$
$\mathbf{D}_g = \gamma_{gd}\mathbf{Id}$	1.510 (0.318)	2.235 (0.435)	3.264 (0.607)	4.048 (0.736)
$\mathbf{D}_g = \gamma_{gd}\mathbf{A}$	1.514 (0.310)	2.239 (0.429)	3.268 (0.599)	4.056 (0.723)
	$\gamma_{gd} = 2$			
	$\gamma_{ga} = \gamma_{gg} = 0.5$	$\gamma_{ga} = \gamma_{gg} = 1$	$\gamma_{ga} = \gamma_{gg} = 2$	$\gamma_{ga} = \gamma_{gg} = 3$
$\mathbf{D}_g = \gamma_{gd}\mathbf{Id}$	1.510 (0.317)	2.236 (0.434)	3.267 (0.605)	4.058 (0.728)
$\mathbf{D}_g = \gamma_{gd}\mathbf{A}$	1.511 (0.318)	2.233 (0.436)	3.263 (0.606)	4.050 (0.734)
	$\gamma_{gd} = 3$			
	$\gamma_{ga} = \gamma_{gg} = 0.5$	$\gamma_{ga} = \gamma_{gg} = 1$	$\gamma_{ga} = \gamma_{gg} = 2$	$\gamma_{ga} = \gamma_{gg} = 3$
$\mathbf{D}_g = \gamma_{gd}\mathbf{Id}$	1.513 (0.320)	2.238 (0.439)	3.267 (0.611)	4.054 (0.733)
$\mathbf{D}_g = \gamma_{gd}\mathbf{A}$	1.514 (0.317)	2.238 (0.436)	3.269 (0.608)	4.057 (0.731)

As in Sermarini *et al.* (2020), the design generated with the same set of parameters assumed for data does not necessarily present better performance. In the present study, by not informing the relationship between test lines in the search for the optimal or near optimal design, it is recommended to adopt low genetic variance ($\gamma_{gd} = 0.5$). However, when searching for the design when such a relationship was declared, no consistency in results was observed, and a general recommendation is not appropriate.

In the study by Bueno Filho & Gilmour (2003), which evaluated three different assumptions for $\mathbf{A}$, the authors showed that, for selection purposes, genetic relationship plays an important role in the search for the optimal block design. The authors reported that some designs are quite robust to misspecification of the covariance structure. For simple genetic covariance structures, which were defined as $\mathbf{A} = \frac{1}{2}\mathbf{Id}$, which resembles the structure used in this work $\mathbf{D}_g = \gamma_{gd}\mathbf{Id}$, for generating the designs; the optimal design is in the class of optimal designs for unrelated treatments. For special covariance structures, which resemble the $\mathbf{D}_g = \gamma_{gd}\mathbf{A}$ structure used in this work, the authors state that it is possible to find an optimal design outside the class of optimal designs for unrelated

Table 2. Mean and standard deviation of selection success for designs generated with $\gamma_{gd} = 0.5, 1, 2, 3$ and assumed genetic variance for data, $\gamma_{ga} = \gamma_{gg} = 0.5, 1, 2, 3$, considering the kinship matrix, $D_g = \gamma_{gd} \mathbf{A}$, and without considering the kinship matrix, $D_g = \gamma_{gd} \mathbf{Id}$, for selection percentage $s = 15\%$.

$\gamma_{gd} = 0.5$				
	$\gamma_{ga} = \gamma_{gg} = 0.5$	$\gamma_{ga} = \gamma_{gg} = 1$	$\gamma_{ga} = \gamma_{gg} = 2$	$\gamma_{ga} = \gamma_{gg} = 3$
$D_g = \gamma_{gd} \mathbf{Id}$	61.846 (4.532)	66.559 (3.929)	71.106 (3.482)	73.297 (3.120)
$D_g = \gamma_{gd} \mathbf{A}$	61.599 (4.417)	66.442 (3.879)	70.954 (3.375)	73.163 (3.204)
$\gamma_{gd} = 1$				
	$\gamma_{ga} = \gamma_{gg} = 0.5$	$\gamma_{ga} = \gamma_{gg} = 1$	$\gamma_{ga} = \gamma_{gg} = 2$	$\gamma_{ga} = \gamma_{gg} = 3$
$D_g = \gamma_{gd} \mathbf{Id}$	61.608 (4.320)	66.558 (3.781)	70.917 (3.358)	73.215 (3.208)
$D_g = \gamma_{gd} \mathbf{A}$	61.582 (4.418)	66.442 (3.959)	70.938 (3.503)	73.275 (3.691)
$\gamma_{gd} = 2$				
	$\gamma_{ga} = \gamma_{gg} = 0.5$	$\gamma_{ga} = \gamma_{gg} = 1$	$\gamma_{ga} = \gamma_{gg} = 2$	$\gamma_{ga} = \gamma_{gg} = 3$
$D_g = \gamma_{gd} \mathbf{Id}$	61.626 (4.320)	66.494 (3.965)	71.021 (3.437)	73.327 (3.124)
$D_g = \gamma_{gd} \mathbf{A}$	61.477 (4.427)	66.477 (3.919)	70.917 (3.467)	73.154 (3.144)
$\gamma_{gd} = 3$				
	$\gamma_{ga} = \gamma_{gg} = 0.5$	$\gamma_{ga} = \gamma_{gg} = 1$	$\gamma_{ga} = \gamma_{gg} = 2$	$\gamma_{ga} = \gamma_{gg} = 3$
$D_g = \gamma_{gd} \mathbf{Id}$	61.438 (4.357)	66.406 (3.889)	70.841 (3.363)	73.239 (3.164)
$D_g = \gamma_{gd} \mathbf{A}$	61.503 (4.363)	66.450 (3.760)	70.844 (3.554)	73.315 (3.082)

treatments.

Regarding the assumption of dependence and independence between test lines, in the study carried out by dos Santos, 2023, no major differences were observed in the comparison measures when considering the dependence or independence between the test lines effects in the design model, corroborating results presented in this work. However, the author states that when the kinship matrix is used in the analysis model, there is a gain of at least one truly good test line selected by the model.

Additionally, when observing genetic gain and selection success, these were greater in cases where the ratio between genetic and residual variance for data were greater, a fact expected and also observed by Cullis *et al.* (2006), Clarke & Stefanova (2011), Santos (2017), Goes (2020) and Sermarini *et al.* (2020).

Finally, Figures 2 and 3 present the densities of the estimates of variance components of fitted models for optimal designs obtained considering the kinship matrix $\mathbf{A}$, and without considering the kinship matrix. For γ_{ga} , in Figure 2, in general, it was observed that the assumptions considered to generate the designs did not significantly influence the estimates of variance components. However, it is interesting to highlight that: (i) when assuming independence between the genetic effects when searching for an optimal design, the cases presenting the least bias in relation to the estimates of variance components were for $\gamma_{gd} = 2$ at $\gamma_g = 0.5$ and for $\gamma_{gd} = 3$ at $\gamma_g = 2$; (ii) when assuming the pedigree matrix of the genetic effects in the design model, the smallest bias in the estimates of variance components was observed when $\gamma_{gd} = 3$ and $\gamma_{ga} = \gamma_{gg} = 3$. The same occurs when analyzing γ_{gg} in Figure 3, which suggests that for higher variance values it is more appropriate to declare the information on kinship between treatments effects. In addition, there were cases in which the components presented overestimation and underestimation, with slight tendencies.

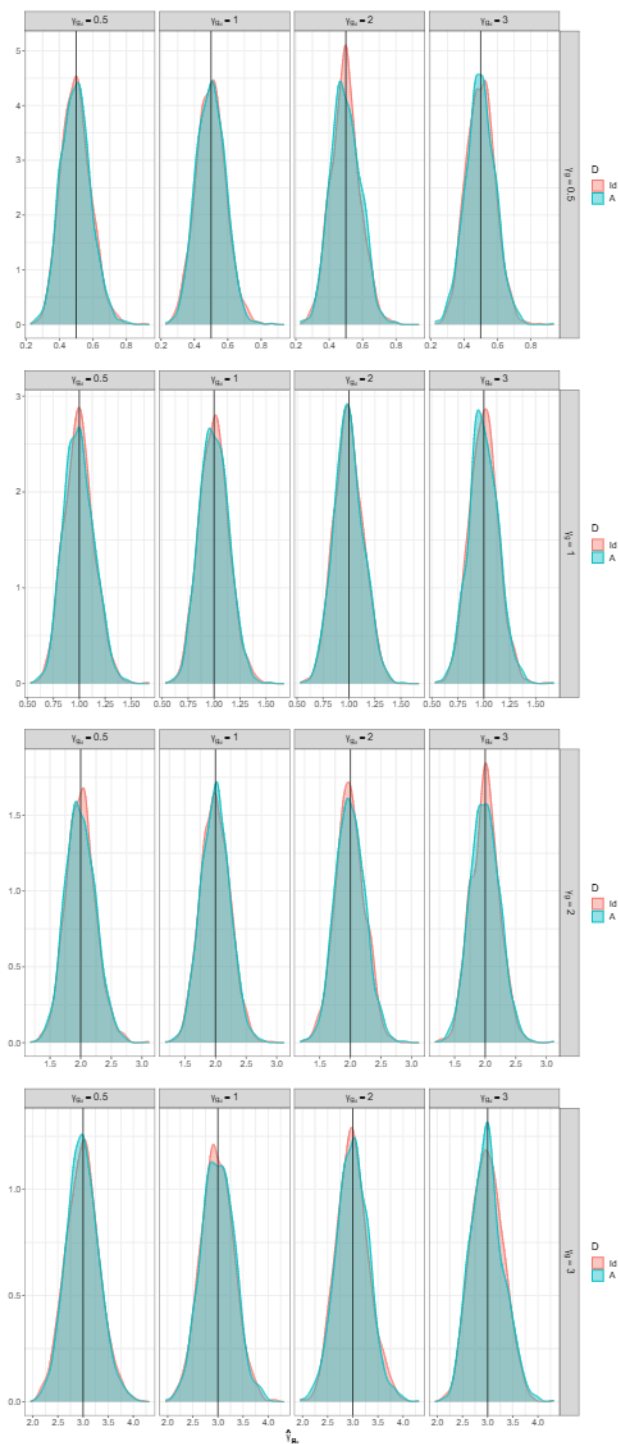


Figure 2. Estimates of parameter related to the genetic variance ($\hat{\gamma}_{gd}$) for the p -rep designs generated considering the kinship matrix **A** and for the p -rep designs generated without considering the kinship matrix, for the four γ_{gd} values. The vertical line represents the expected parameter value.

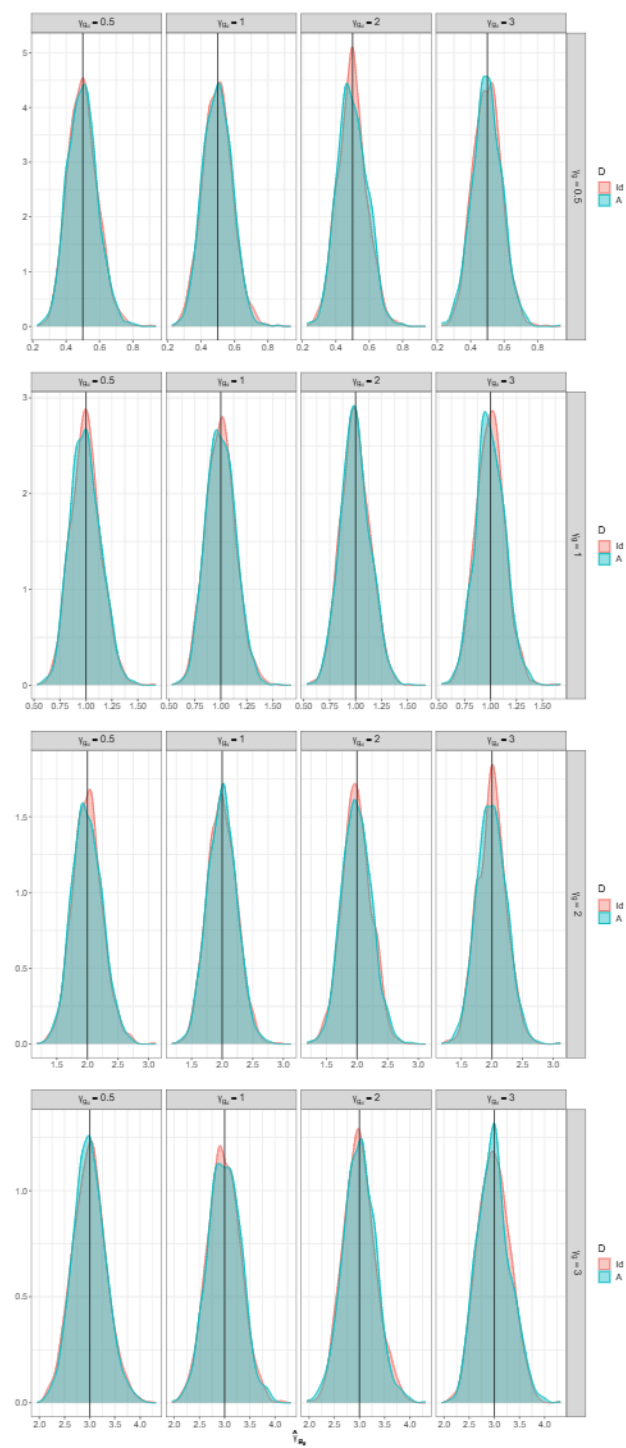


Figure 3. Estimates of parameter related to the genetic variance ($\hat{\gamma}_{gd}$) for the p -rep designs generated considering the kinship matrix $\mathbf{A}$ and for the p -rep designs generated without considering the kinship matrix, for the four of γ_{gd} values. The vertical line represents the expected parameter value.

4. Conclusions

In this study, spatially optimized p -rep designs for 598 test lines were compared in an experimental area of 36 rows by 20 columns, with the objective of evaluating genetic gain and selection quality through measures of realized genetic gain and selection success when assuming some characteristics for the designs. That is, four values for the generation of the designs and for the simulation study, γ_{gd} and $\gamma_{gs} = \gamma_{gg}$, being 0.5, 1, 2, 3; a percentage value for the repeated *test lines*, $p = 20\%$; a value for the selection percentage, $s = 15\%$ and a scenario for the values of the nongenetic effects. In addition, assumptions of dependence and independence between the treatment effects, in this case, the test lines, were considered.

No large differences were identified for realized genetic gain and selection success among the designs that were generated assuming the pedigree matrix (**A**) and the designs that considered independence (**Id**) between the treatment effects. The same was observed for selection success. Regarding the distribution of test lines in the experimental area, including or not the pedigree matrix did not interfere with randomness, but it is noted that for higher values of genetic variance, the frequency of test lines in neighboring plots and on the border of the experimental area is lower. When $\gamma_{gd} = 2$ is assumed, there is a smaller number of clustered replicates, a smaller number of replicates with at least one neighboring replicate, and a smaller number of replicates on the border.

It is highlighted that the results for the quality of selection are strongly affected by the characteristics of the data, since the higher the ratio between the genetic and residual variances, the higher the values for the RGG and the success of selection. This study has some limitations. The data were simulated based on a set of parameters used by plant breeding programs and the same experimental area size, the same percentage of repeated test lines, the same percentage of selection, a single scenario for the values of nongenetic effects and the genetic variances were supposed to be equal. In addition, the matrices **A** and **G** used in this work, extracted from CIMMYT, are for the wheat crop. Therefore, other results could be found if different scenarios were assumed, for example, for the percentage of repeated test lines, percentage of selection and higher values for the nongenetic variance parameters. Therefore, it is expected that experimental designs that provide more detailed information about treatments present a better performance in terms of precision and reliability of results, since greater amount of information allows a more robust analysis and a better identification of the treatment effects.

Acknowledgments

This work was supported by Conselho Nacional de Desenvolvimento Científico e Tecnológico (CNPq).

Conflicts of Interest

The authors declare no conflict of interest.

Author Contributions

Conceptualization: GOUVEIA, P.R.S.; SERMARINI, R.A. **Data curation:** GOUVEIA, P.R.S.; SERMARINI, R.A. **Formal analysis:** GOUVEIA, P.R.S.; SERMARINI, R.A. **Funding acquisition:** GOUVEIA, P.R.S.; SERMARINI, R.A. **Investigation:** GOUVEIA, P.R.S.; SERMARINI, R.A. : GOUVEIA, P.R.S.; SERMARINI, R.A. **Project administration:** GOUVEIA, P.R.S.; SERMARINI, R.A. **Software:** GOUVEIA, P.R.S.; SERMARINI, R.A. **Resources:** GOUVEIA, P.R.S.; SERMARINI, R.A. **Supervision:** GOUVEIA, P.R.S.; SERMARINI, R.A. **Validation:** GOUVEIA, P.R.S.; SERMARINI, R.A. **Visualization:** GOUVEIA, P.R.S.; SERMARINI, R.A. **Writing - original draft:** GOUVEIA, P.R.S.; SERMARINI, R.A. **Writing - review and editing:** GOUVEIA, P.R.S.; SERMARINI, R.A.

References

1. Bates, D. & Vazquez, A. I. *pedigreemm: Pedigree-based mixed-effects models* R package version 0.3-3 (2014). <https://CRAN.R-project.org/package=pedigreemm>.
2. Bueno Filho, J. S. S. & Gilmour, S. G. Block designs for random treatment effects. *Journal of Statistical Planning and Inference* **137**, 1446–1451. doi:10.1016/j.jspi.2006.02.002 (2007).
3. Bueno Filho, J. S. S. & Gilmour, S. G. On the Design of Field Experiments with Correlated Treatment Effects. *Biometrics* **59**, 375–381. doi:10.1007/s13253-014-0191-0 (2003).
4. Butler, D. G., Smith, A. B. & Cullis, B. R. On the Design of Field Experiments with Correlated Treatment Effects. *Journal of Agricultural, Biological, and Environmental Statistics* **19**, 541–557. doi:10.1007/s13253-014-0191-0 (2014).
5. Butler, D. *odw: Generate optimal experimental designs* R package version 2.1.3 (2022).
6. CIMMYT. <https://www.cimmyt.org/>, accessed 21 April of 2021. 2021.
7. Clarke, G. P. Y. & Stefanova, K. T. Optimal Design for Early-Generation Plant-Breeding Trials with Unreplicated or Partially Replicated Test Lines. *Australian & New Zeland Journal of Statistics* **53**, 461–480. doi:10.1111/j.1467-842X.2011.00642.x (2011).
8. Cullis, B. R., Smith, A. B., Cocks, N. A. & Butler, D. G. The Design of Early-Stage Plant Breeding Trials Using Genetic Relatedness. *Journal of Agricultural, Biological, and Environmental Statistics* **24**, 553–578. doi:10.1007/s13253-020-00403-5 (2020).
9. Cullis, B. R., Smith, A. B. & Coombes, N. E. On the Design of Early Generation Variety Trials With Correlated Data. *Journal of Agricultural, Biological, and Environmental Statistics* **11**, 381–392. doi:10.1198/108571106X154443 (2006).
10. Dos Santos, D. P. *Delineamentos ótimos para experimentos multi-ambientais de melhoramento genético de plantas* Tese de doutorado. PhD thesis (Universidade de São Paulo, São Paulo, Brasil, 2023). doi:10.11606/T.11.2023.rde-12072023-170331.
11. Dos Santos, D. P., Sermarini, R. A., dos Santos, A. & Demétrio, C. G. B. Optimal Designs in Plant Breeding Experiments: A Simulation Study Comparing Grid-Plot and Partially Replicated (p -Rep) Design. *Sugar Tech* **26**, 387–395. issn: 0974-0740. doi:10.1007/s12355-024-01375-3 (2024).
12. Federer, W. T. Augmented split block experiment design. *Agronomy Journal* **97**, 578–586. doi:10.2134/agronj2005.0578 (2005).
13. Federer, W. T. Construction and analysis of an augmented lattice square design. *Biometrical journal* **44**, 251–257. doi:10.1002/1521-4036(200203)44:2<251::AID-BIMJ251>3.0.CO;2-N (2002).
14. Federer, W. T. & Crossa, J. On the design and analysis of field experiments (2001).
15. Federer, W. T., Nair, R. C. & Raghavarao, D. Some augmented row-column designs. *Biometrics* **31**, 361–373. doi:10.2307/2529426 (1975).
16. Federer, W. T. & Raghavarao, D. On augmented designs. *Biometrics* **31**, 29–35. doi:10.2307/2529707 (1975).
17. Federer, W. T. Augmented (or hoonuiaku) designs. *Biometrics Unit Technical Reports* **55**, 191–208 (1956).
18. Federer, W. Augmented designs with one-way elimination of heterogeneity. *Biometrics* **17**, 447–473. doi:10.2307/2527837 (1961).

19. Gilmour, A. R., Cullis, B. R. & Verbyla, A. P. Accounting for Natural and Extraneous Variation in the Analysis of Field Experiments. *Journal of Agricultural, Biological, and Environmental Statistics* **2**, 269–293. doi:10.2307/1400446 (1997).
20. Goes, A. L. *Delineamentos ótimos para experimentos com cana-de-açúcar* Dissertação de Mestrado (Universidade de São Paulo, Piracicaba, 2020), 44.
21. Henderson, C. R. A Simple Method for Computing the Inverse of a Numerator Relationship Matrix Used in Prediction of Breeding Values. *Biometrics* **32**, 69–83. doi:10.2307/2529339 (1976).
22. Hooks, T., Marx, D., Kachman, S. & Pedersen, J. Optimality Criteria for Models with Random Effects. *Revista Colombiana de Estadística* **32**, 17–31. http://scielo.org.co/scielo.php?script=sci_arttext&pid=S0120-17512009000100002 (2009).
23. Kempton, R. A. The design and analysis of unreplicated field trials. *Vorträge für Pflanzenzüchtung* **7**, 219–242 (1984).
24. Lin, C. S. & Poushinsky, G. A modified augmented design for an early stage of plant selection involving a large number of test lines without replication. *Biometrics* **39**, 553–561. doi:10.2307/2531083 (1983).
25. Lin, C.-S. & Poushinsky, G. A modified augmented design (type 2) for rectangular plots. *Canadian Journal of Plant Science* **65**, 743–749. issn: 0008-4220. doi:10.4141/cjps85-094. <http://pubs.aic.ca/doi/abs/10.4141/cjps85-094> (1985).
26. Moehring, J., Williams, E. R. & Piepho, H.-P. Efficiency of augmented p-rep designs in multi-environmental trials. *Theoretical and Applied Genetics* **127**, 1049–1060. doi:10.1007/s00122-014-2278-y (2014).
27. Müller, B. U., Schützenmeister, A. & Piepho, H. P. Arrangement of check plots in augmented block designs when spatial analysis is used. *Plant Breeding* **6**, 581–589. doi:10.1111/j.1439-0523.2010.01803.x (2010).
28. Perez-Rodriguez, P. *lme4GS: Lme4 for Genomic Selection* R package version 0.1 (2021).
29. Piepho, H.-P. & Williams, E. R. Augmented Row–Column Designs for a Small Number of Checks. *Agronomy Journal* **108**. doi:10.2134/agronj2016.06.0325 (2016).
30. R Core Team. *R: A Language and Environment for Statistical Computing* R Foundation for Statistical Computing (Vienna, Austria, 2021). <https://www.R-project.org/>.
31. Santos, A. *Design and analysis of sugarcane breeding experiments: a case study* PhD thesis (Universidade de São Paulo, Piracicaba, 2017), 175. doi:10.11606/T.11.2017.tde-06102017-103933.
32. Sermarini, R. A., Brien, C., Demétrio, C. G. B. & Santos, A. Impact on genetic gain from using misspecified statistical models in generating p-rep designs for early generation plant-breeding experiments. *Crop Science*, 1–13. doi:10.1002/csc2.20257 (2020).
33. Shah, K. R. & Sinha, B. K. *Theory of Optimal Designs* (Springer-Verlag Berlin Heidelberg, Germany, 1989).
34. Smith, A. B., Lim, P. & Cullis, B. R. The design and analysis of multi-phase plant breeding experiments. *The Journal of Agricultural Science* **144**, 393. issn: 0021-8596. doi:10.1017/S0021859606006319. [http://www.journals.cambridge.org/abstract\[_\]S0021859606006319](http://www.journals.cambridge.org/abstract[_]S0021859606006319) (2006).
35. The VSNi Team. *asreml: Fits Linear Mixed Models using REML* R package version 4.2.0.276 (2023). www.vsnr.co.uk.