

ORIGINAL RESEARCH ARTICLE

Crop Breeding & Genetics

Boosting predictive ability of tropical maize hybrids via genotype-by-environment interaction under multivariate GBLUP models

Matheus Dalsente Krause^{1,2}  | Kaio Olímpio das Graças Dias¹  | Jhonathan Pedroso Rigal dos Santos¹  | Amanda Avelar de Oliveira¹  |
Lauro José Moreira Guimarães³  | Maria Marta Pastina³  |
Gabriel Rodrigues Alves Margarido¹  | Antonio Augusto Franco Garcia¹ 

¹ Dep. of Genetics, Luiz de Queiroz College of Agriculture, Univ. of São Paulo, PO Box 83, Piracicaba, São Paulo 13400-970, Brazil

² Dep. of Agronomy, Iowa State Univ., Ames, IA 50011, USA

³ Embrapa Milho e Sorgo, Rod. MG 424, Km 65, Sete Lagoas, Minas Gerais 35701-970, Brazil

Correspondence

Antonio Augusto Franco Garcia, Dep. of Genetics, Luiz de Queiroz College of Agriculture, Univ. of São Paulo, PO Box 83, Piracicaba, São Paulo 13400-970, Brazil. Email: augusto.garcia@usp.br

Assigned to Associate Editor Michael Casler.

Funding information

Conselho Nacional de Desenvolvimento Científico e Tecnológico, Grant/Award Number: Productivity scholarship for A.A.F.G.; Fundação de Amparo à Pesquisa do Estado de São Paulo, Grant/Award Number: 2016/12977-7; Coordenação de Aperfeiçoamento de Pessoal de Nível Superior, Grant/Award Number: Master's scholarship for M.D.K

Abstract

Genomic selection has been implemented in several plant and animal breeding programs and it has proven to improve efficiency and maximize genetic gains. Phenotypic data of grain yield was measured in 147 maize (*Zea mays* L.) single-cross hybrids at 12 environments. Single-cross hybrids genotypes were inferred based on their parents (inbred lines) via single nucleotide polymorphism (SNP) markers obtained from genotyping-by-sequencing (GBS). Factor analytic multiplicative genomic best linear unbiased prediction (GBLUP) models, in the framework of multi-environment trials, were used to predict grain yield performance of unobserved tropical maize single-cross hybrids. Predictions were performed for two situations: untested hybrids (CV1), and hybrids evaluated in some environments but missing in others (CV2). Models that borrowed information across individuals through genomic relationships and within individuals across environments presented higher predictive accuracy than those models that ignored it. For these models, predictive accuracies were up to 0.4 until eight environments were considered as missing for the validation set, which represents 67% of missing data for a given hybrid. These results highlight the importance of including genotype-by-environment interactions and genomic relationship information for boosting predictions of tropical maize single-cross hybrids for grain yield.

Abbreviations: BLUP, best linear unbiased prediction; CV1, Cross-Validation 1; CV2, Cross-Validation 2; E-BLUE, empirical best linear unbiased estimation; FA, factor analytic; GBLUP, genomic best linear unbiased prediction; GBS, genotyping-by-sequencing; GE, genotype-by-environment; GEBV, genomic estimated breeding value; GS, genomic selection; MAF, minor allele frequency; MET, multi-environment trial; SNP, single nucleotide polymorphism; VCOV, variance-covariance.

1 | INTRODUCTION

Future prospects of population growth and rising demand for agriculture products increase even further the role of releasing stable cultivars worldwide. Phenotyping in multi-environment trials (METs) plays a major role in accessing the performance of lines across target breeding regions (Oakey et al., 2016) and is one of the most resource-demanding stages in the breeding program (Cobb et al., 2019; Fritsche-Neto, Gonçalves, Vencovsky, & De Souza Junior, 2010). In METs, it is possible to detect and quantify the differential response of hybrids triggered by different environmental factors across environments, which is known as genotype-by-environment (GE) interaction. For maize (*Zea mays* L.) cultivated in the tropics, the evaluation in METs is even more critical for selecting stability, especially due to biotic and abiotic stresses frequently found in those regions inducing GE interaction responses.

Among the many challenges faced by breeders, the recommendation of stable and adaptable hybrids across environments is a key component in breeding (Elias, Robbins, Doerge, & Tuinstra, 2016). This is especially challenging due to the difficulties in understanding if the performance of the hybrid is due to its pure genetics or GE interaction. Groups of environments with high genetic correlation unleashed by the expression of common genes across environments will have fewer crossover interactions, which can result in few changes in the rank of genotypes across environments. More significant crossover interactions are expected in environments displaying low correlation caused by the variability of physiological factors critical for plant development, such as water availability, temperature, radiation, and disease pressure, causing reranking of genotypes across environments (Van Eeuwijk, Bustos-Korts, & Malosetti, 2016; Yan, 2016). Therefore, GE interaction is an important component for hybrid evaluation and posterior cultivar recommendation in the target breeding region, and its understanding requires models capable of integrating realistic scenarios observed in breeding programs in order to facilitate this decision-making process.

The concept of genomic selection (GS) was introduced by Meuwissen, Hayes, and Goddard (2001) for livestock and can be defined as a form of marker-assisted selection in which models using markers covering the whole genome are used to predict genomic breeding values (GEBVs) of individuals. The selections are therefore based on GEBVs. Genomic selection has been implemented in a range of breeding programs (Jonas & de Koning, 2016) and has been proven to facilitate the rapid selection of superior genotypes and to accelerate breeding cycles (Crossa et al., 2017), becoming an important tool to increase the annual rate of genetic gains (Heffner, Lorenz, Jannink, & Sorrells, 2010; Hickey, Chiurugwi, Mackay, & Powell, 2017). A

key point for the success of GS is the large availability of cost-effective high-throughput genotyping technologies (Crossa et al., 2017), resulting in large-scale genomic information for most crops (Bernardo, 2008; Krchov & Bernardo, 2015).

Breeding values can be predicted by linear mixed models that take into account the relationship between individuals, obtained by known pedigrees (Fisher, 1918; Wright, 1921) that rely on the concept of identity-by-descent (IBD), or by marker-based relationship estimated by identity-by-state (IBS) (Powell, Visscher, & Goddard, 2010; VanRaden, 2008), resulting in GEBVs. Both methods can be incorporated into linear mixed models for hybrids prediction, or even combined information from the pedigree and genomic relationship matrices. When the marker-based matrix is incorporated, the model is known as genomic best linear unbiased prediction (GBLUP) (Oakey et al., 2016). The advantage of using a marker-based matrix is that it does not require pedigrees, which might have errors and could not be available, and it captures the Mendelian sampling under the absence of inbreeding depression and assortative mating (Burgueño, de los Campos, Weigel, & Crossa, 2012; Powell et al., 2010).

The GE interaction can also be incorporated into a mixed model to predict performance of untested genotypes in one or more target environments (Burgueño et al., 2012; Crossa et al., 2016; Lopez-Cruz et al., 2015). The first model proposed for prediction of maize single-cross hybrids based on best linear unbiased prediction (BLUP) was implemented by Bernardo (1994) with balanced data. After, Bernardo (1995) used BLUP to predict single-cross hybrids performance with unbalanced yield trial data (missing hybrids). In both studies, a limited number of molecular markers were available and included in the models. Besides that, at that time, no effective variance-covariance (VCOV) structures were available that could take advantage of correlated environments by including the GE interaction and, at the same time, handle unbalanced data.

Unbalanced data from METs is routine in plant breeding programs, resulting in challenges for data analysis (Dawson et al., 2013). The process of selection naturally discards lines with poor performance, and on the other hand, new entries are added every year (Piepho, Möhring, Melchinger, & Büchse, 2008). Also, some unbalancing and heterogeneous quality of the data is due to different degrees of replication in different trials (e.g., preliminary or advanced trials; Lado, Barrios, Quincke, Silva, & Gutiérrez, 2016). Historically, joint analysis of variance and linear regression models were used in METs to analyze and quantify GE interactions (Elias et al., 2016). In these traditional models, genetic effects were assumed to have normal distribution with a unique genetic variance component (Dias, Gezan, Guimarães, Nazarian, et al., 2018; Piepho

et al., 2008; Smith, Cullis, & Thompson, 2001), which is not realistic assumption.

One strategy to use more realistic models is by including GE interaction by modeling the genetic VCOV matrix across environments (Burgueño et al., 2012). A common approach is to include the unstructured matrix, which allows a specific genetic variance for each environment and different pairwise covariances between environments. However, given the number of parameters to be estimated, fitting this model became computational prohibitively and impractical (Kelly, Smith, Eccleston, & Cullis, 2007). An alternative way to overcome this limitation is by including a factor analytic (FA) structure (Piepho, 1997, 1998; Smith et al., 2001). It requires a reduced number of parameters to be estimated and has been used in several breeding programs due its good applicability over the unstructured VCOV structure (Kelly et al., 2007; Oakey et al., 2016).

Most previous results of GBLUP models were fitted for single-environment predictions (Cuevas et al., 2016; Zhang et al., 2015). However, high levels of predictive accuracies been found in GS models that incorporate both genomic information from marker-based matrix and GE interaction (Acosta-Pech et al., 2017; Jarquín et al., 2014). In this context, the goals of this work were (a) to predict the performance of untested tropical maize single-cross hybrids for grain yield within environments using GBLUP models in the framework of METs, and (b) to investigate the usefulness of genomic relationship information in combination with different VCOV structures for genetics and residuals effects, under different levels of unbalanced trials.

2 | MATERIALS AND METHODS

2.1 | Experimental data

The dataset was obtained by the Brazilian Agricultural Research Corporation (Embrapa) Maize and Sorghum. Yield data were collected at eight locations in Brazil in 2012 at two conditions: two different crop seasons, “safrsummer” and “safrinha-winter.” Not all locations meet both conditions. The combination of locations and crop seasons were designated as “environment,” giving a total of 12 environments (Supplemental Figure S1). The trait under consideration is grain yield, in tons per hectare ($t\ ha^{-1}$), adjusted to 13% grain moisture.

In the first crop season (planting data from September to November, safrsummer), plants have favorable growing conditions, the result of the increase in temperature and rainfall, plus a reduced intensity of plant disease and insect pests. In the second crop season (planting data from January to March, safrinha-winter), these conditions are the

opposite. From the end of January to the following months, the intensity of rainfall and average temperature decreases, and field crops have to face the spore load plus pest infestations not efficiently controlled from the first crop season.

The dataset comprises 152 maize hybrids split into three trials, evaluated side by side in the field, in all environments. The first two trials (T1 and T2) and the third trial (T3) contain 60 and 32 hybrids each, respectively. Each trial was augmented by four common checks (commercial maize cultivars) and arranged as a balanced lattice square of eight by eight (T1 and T2) and six by six (T3), with two replications. The first two trials (T1 and T2) had 120 hybrids from an intermediate stage of hybrids evaluation, and the third trial (T3) had 32 hybrids from an advanced stage of the maize breeding program.

Among the 152 maize hybrids, 149 are single crosses, two are three-way crosses, one is a double cross, and four are commercial checks, with only the single crosses being under consideration for genomic prediction. The single crosses were obtained from 144 inbred lines, classified as dent (64 lines) and flint (77 lines) heterotic groups, and also another group C (three lines), which performs well when crossed with both dent and flint sources. Four lines were used as testers from the opposite heterotic group to synthesize the majority of single-cross hybrids.

2.2 | Genotypic data

The inbred lines used as parents were genotyped with the standard genotyping-by-sequencing (GBS) protocol (Elshire et al., 2011) by the Genomic Diversity Facility at Cornell University (Ithaca, NY, USA). Tags were aligned to the B73 reference genome (AGPv3; Schnable et al., 2009). Standard quality controls were applied to the data, removing all non-biallelic markers, and single nucleotide polymorphisms (SNPs) were discarded if at least one of the following was true: minor allele frequency (MAF) < 5%; >20% missing genotypes; and inbreeding coefficient < 0.8. The SNPs were called using the GBS pipeline available in the software TASSEL version 5 (Glaubitz et al., 2014). After filtering, missing data were imputed using Beagle 4.1 (Browning & Browning, 2016). The number of SNPs per chromosome ranged from 1,951 (chromosome 10) to 5,024 (chromosome 1), and the final number of SNPs after filtering was 29,515.

For each SNP, the genotypes of the single-cross hybrids were inferred based on the genotype of their parents (inbred line) based on Mendelian laws in the software R version 3.4.3 (R Core Team, 2017). One of the 144 inbred lines used as parents was not genotyped, resulting in the availability of genotypic information of 147 single-cross hybrids instead of 149 hybrids. Principal

components analysis (PCA) of SNP matrix of the 143 inbred lines was performed in the software TASSEL version 5 (Glaubitz et al., 2014) to verify the consistency of heterotic groups.

2.3 | Single-environment trial analyses

The following single-environment trial model was fitted to estimate genetic parameters, the accuracy of the trial, and the empirical best linear unbiased estimations (E-BLUEs) of hybrids in each environment:

$$\mathbf{y} = \mu \mathbf{1}_n + \mathbf{X}_1 \boldsymbol{\beta} + \mathbf{X}_2 \mathbf{g} + \mathbf{Z} \mathbf{b} + \boldsymbol{\varepsilon} \quad (1)$$

where $\mathbf{y}$ is a $n \times 1$ vector of phenotypes for m hybrids and j replicates; μ is the overall mean; $\boldsymbol{\beta}$ is a $j \times 1$ vector of fixed effects of replicates; $\mathbf{g}$ is a $m \times 1$ vector of fixed effects of hybrids; $\mathbf{b}$ is a $rj \times 1$ vector of random effects of blocks within replications, with $\mathbf{b} \sim N(\mathbf{0}, \sigma_b^2 \mathbf{I}_{rj})$; and $\boldsymbol{\varepsilon}$ is a $n \times 1$ vector of residuals, with $\boldsymbol{\varepsilon} \sim N(0, \sigma_e^2 \mathbf{I}_n)$. $\mathbf{X}_1$, $\mathbf{X}_2$, and $\mathbf{Z}$ are incidence matrices for their respective effects, with dimensions of $n \times j$, $n \times m$, and $n \times rj$, respectively. $\mathbf{I}_{rj}$ and $\mathbf{I}_n$ are identity matrices of their corresponding dimensions, and $\mathbf{1}_n$ is a vector of ones with dimension $n \times 1$.

The generalized measure of heritability was estimated using $\hat{H}^2 = 1 - [\text{PEV}/(2\sigma_g^2)]$, where PEV (prediction error variance) is the mean variance of the difference between two genetic effects and σ_g^2 is the genetic variance (Cullis, Smith, & Coombes, 2006). The genetic variance was estimated by assuming hybrids as independently and identically normal distributed with mean zero and variance σ_g^2 in Model 1. The coefficient of variation (CV, %) was estimated based on $\text{CV} (\%) = (\sigma_e/\mu) \times 100$, where σ_e is the square root of residual variance component (σ_e^2) and μ is the mean of grain yield of each trial within environment. The significance of variance components was assessed by the likelihood ratio test (LRT) assuming $\alpha = .05$.

All statistical models of this section and the section below were fitted using the package ASReml-R version 3 (Butler, Cullis, Gilmour, & Gogel, 2009) by solving the mixed-model equations proposed by Henderson (1950). Variance components were estimated using residual maximum likelihood (REML; Patterson & Thompson, 1971).

2.4 | Genomic prediction in multi-environment trials

In order to predict the performance of single-cross maize hybrids under METs, we fitted different models formulation differing by their genetic (Σ_g) and residual (Σ_r) VCOV structures between environments. The following generic

model was fitted:

$$\mathbf{y} = \mu \mathbf{1}_n + \mathbf{X} \boldsymbol{\beta} + \mathbf{Z}_1 \mathbf{b} + \mathbf{Z}_2 \mathbf{g} + \boldsymbol{\varepsilon} \quad (2)$$

where $\mathbf{y}$ is a $n \times 1$ vector of phenotypes for m hybrids across s environments and q trials, with $n = \sum_{i=1}^s n_i$, in which n_i is the number of plots within environment s ; $\boldsymbol{\beta}$ is a $f \times 1$ vector of fixed effects of environments, trials within environments, replicates within trials within environments, checks, and checks within environments (single-cross hybrids without marker information, three-way and double cross hybrids were considered as checks); $\mathbf{b}$ is a $v \times 1$ vector of random effects of blocks within replications within trials within environments, where $\mathbf{b} \sim \text{MVN}(\mathbf{0}, \sigma_b^2 \mathbf{I}_v)$; $\mathbf{g}$ is a $m \times 1$ vector of random effects of hybrids within environments, where $\mathbf{g} \sim \text{MVN}(\mathbf{0}, \Sigma_g)$; and $\boldsymbol{\varepsilon}$ is a $n \times 1$ vector of independent residuals within environments, where $\boldsymbol{\varepsilon} \sim \text{MVN}(\mathbf{0}, \Sigma_r)$. $\mathbf{X}$, $\mathbf{Z}_1$, and $\mathbf{Z}_2$ are incidence matrices for their respective effects, with dimensions of $n \times f$, $n \times v$, and $n \times ms$. $\mathbf{I}_v$ is an identity matrix of its corresponding dimensions, $\mathbf{1}_n$ is a vector of ones with dimension $n \times 1$, and MVN stands for multivariate normal distribution.

The VCOV structures for random effects of hybrids and residuals within environments were defined as $\Sigma_g = \mathbf{G} \otimes \mathbf{A}$ and $\Sigma_r = \mathbf{I} \otimes \mathbf{R}$, respectively. In Σ_g , the additive genomic matrix $\mathbf{A}$ refers to the genetic relatedness between hybrids (described below). It can also assume that hybrids are unrelated, with $\Sigma_g = \mathbf{G} \otimes \mathbf{I}$. In addition to the hybrid's relatedness, Σ_g was also modeled for MET, embracing three different VCOV structures: independent environments with equal variance ($\mathbf{G} = \mathbf{I}$); independent environments with unequal variances ($\mathbf{G} = \mathbf{D}$); and correlated environments with unequal variances ($\mathbf{G} = \mathbf{FA}$). $\mathbf{I}$, $\mathbf{D}$, and $\mathbf{FA}$ stand for identity, diagonal, and factor analytic VCOV matrices, respectively. The residual VCOV structure in MET was defined as $\Sigma_r = \mathbf{I} \otimes \mathbf{R}$, where homogeneous residual variance ($\mathbf{R} = \mathbf{I}$) or heterogeneous residual variances ($\mathbf{R} = \mathbf{D}$) across environments were allowed for all combinations of genetic effects in Σ_g .

The additive genomic relationship matrix $\mathbf{A}$ was computed based on SNP markers from GBS, following the methodology described by VanRaden (2008) as

$$\mathbf{A} = \frac{\mathbf{Z}\mathbf{Z}'}{2\sum p_i(1 - p_i)} \quad (3)$$

where $\mathbf{Z} = \mathbf{M} - \mathbf{P}$, in which $\mathbf{M}$ is the incidence matrix for markers considering two alleles (A and a) for a given i th marker locus, coded as 0, 1, and 2 for AA, Aa, and aa, respectively, and $\mathbf{P}$ is derived from observed allele frequencies expressed as $\mathbf{P} = 2p_i$, where p_i is the MAF of locus i . It was estimated using the package AGHmatrix (Amadeu

TABLE 1 Goodness of fit for models divided into three categories, (a) factor analytic (FA) models without incorporating the genomic relationship matrix (Models 1–8), (b) FA models with genomic relationship information (Models 9–16), and (c) models assuming no correlation between environments but including genomic relationship information (Models 17–20)

Model		Covariance structure			Selection criteria		
Number	Code	Σ_g	Σ_r	Nu. Par. ^a	AIC	BIC	$\bar{v}^b$
BLUP							
1	I _{FA(1)-I}	FA(1) $\otimes$ I	I $\otimes$ I	25	4,704.08	4,866.09	48.2
2	I _{FA(2)-I}	FA(2) $\otimes$ I	I $\otimes$ I	36	4,673.45	4,904.00	61.7
3	I _{FA(3)-I}	FA(3) $\otimes$ I	I $\otimes$ I	46	4,669.74	4,962.60	70.8
4	I _{FA(4)-I}	FA(4) $\otimes$ I	I $\otimes$ I	55	4,682.76	5,031.70	74.7
5	I _{FA(1)-D}	FA(1) $\otimes$ I	I $\otimes$ D	36	4,376.63	4,607.18	51.7
6	I _{FA(2)-D}	FA(2) $\otimes$ I	I $\otimes$ D	47	4,340.85	4,639.94	64.2
7	I _{FA(3)-D}	FA(3) $\otimes$ I	I $\otimes$ D	57	4,327.86	4,689.27	76.0
8	I _{FA(4)-D}	FA(4) $\otimes$ I	I $\otimes$ D	66	4,333.35	4,750.83	80.4
GBLUP							
9	A _{FA(1)-I}	FA(1) $\otimes$ A	I $\otimes$ I	25	5,186.90	5,348.91	95.3
10	A _{FA(2)-I}	FA(2) $\otimes$ A	I $\otimes$ I	36	5,181.70	5,412.26	97.7
11	A _{FA(3)-I}	FA(3) $\otimes$ A	I $\otimes$ I	46	5,174.77	5,467.63	99.9
12	A _{FA(4)-I}	FA(4) $\otimes$ A	I $\otimes$ I	55	5,275.04	5,623.99	99.9
13	A _{FA(1)-D}	FA(1) $\otimes$ A	I $\otimes$ D	36	4,707.40	4,937.96	99.9
14	A _{FA(2)-D}	FA(2) $\otimes$ A	I $\otimes$ D	47	4,736.43	5,035.53	99.9
15	A _{FA(3)-D}	FA(3) $\otimes$ A	I $\otimes$ D	57	4,784.38	5,145.79	99.9
16	A _{FA(4)-D}	FA(4) $\otimes$ A	I $\otimes$ D	66	4,782.31	5,199.80	99.9
17	A _{I-I}	I $\otimes$ A	I $\otimes$ I	2	5,494.34	5,513.03	–
18	A _{I-D}	I $\otimes$ A	I $\otimes$ D	13	4,979.53	5,066.77	–
19	A _{D-I}	D $\otimes$ A	I $\otimes$ I	13	5,509.51	5,596.75	–
20	A _{D-D}	D $\otimes$ A	I $\otimes$ D	24	4,997.61	5,153.38	–

Note. I, identity matrix; FA(*k*), factor analytic matrix of order *k*; D, diagonal matrix; A, genomic relationship matrix from molecular markers; Σ_g , genetic variance-covariance structure; Σ_r , residual variance-covariance structure; BLUP, best linear unbiased prediction; GBLUP, genomic best linear unbiased prediction. Genomic predictions in multi-environment trials were performed with models in which both Akaike information criterion (AIC) and Bayesian information criterion (BIC) values are in bold.

^aNumber of variance components estimated for each model.

^bPercentage of genetic variance accounted for FA models.

et al., 2016) in the software R version 3.4.3 (R Core Team, 2017).

For FA models in Σ_g , estimations of genetic variance ($\hat{G}_e$) and correlation matrices ($\hat{C}$) between environments were obtained for the by $\hat{G}_e = (\hat{\Lambda}\hat{\Lambda}' + \hat{\Psi})$ and $\hat{C} = \hat{D}\hat{G}_e\hat{D}$, respectively, where $\hat{\Lambda}$ is a $s \times k$ matrix of loadings (*k*) for all environments (*s*), $\hat{\Psi}$ is a $s \times s$ diagonal matrix of specific variances of each environment, and $\hat{D}$ is a diagonal matrix of the inverse of the square root of the diagonal values of $\hat{G}_e$. *k* refers to the number of multiplicative terms or factors of the FA structure (see Section 2.5). For more specialized literature on the FA VCOV structure, please check Piepho (1998), Smith et al. (2001), and Smith, Ganeshalingam, Kuchel, and Cullis (2015).

To evaluate the impact of modeling genetics and residuals covariance structures, the models were classified based on their assumptions as follows: Models 1–8 (first category) assume unrelated hybrids with unique genetic vari-

ance and correlated environments with heterogeneous genetic variances; Models 9–16 (second category) assume related hybrids with unequal genetic variances, based on the genomic relationship matrix A, and correlated environments with heterogeneous genetic variances; and Models 17–20 (third category) assume related hybrids with equal or unequal genetic variances but independent environments. For residuals within environments, homogeneous residual variance ($R = I$) or heterogeneous residual variances ($R = D$) across environments for all combinations of genetic effects were assumed (Table 1).

2.5 | Model selection

Three criteria were used to select the best-fit models: (a) the goodness of fit via Akaike information criterion (AIC; Akaike, 1974), (b) Bayesian information criterion (BIC;

Schwarz, 1978), and for FA models, (c) the overall percentage of genetic variance ($\bar{v}$) accounted by each k factor, defined as $\bar{v} = 100\text{tr}(\hat{\mathbf{A}}\hat{\mathbf{A}}')/\text{tr}(\hat{\mathbf{A}}\hat{\mathbf{A}}' + \hat{\mathbf{\Psi}})$, where “tr” is the trace of the matrix and the other terms were previously defined (Smith et al., 2015). The FA models were fitted until $\bar{v}$ exceeded the ad hoc value of 70%. For the first category of models (BLUP), the best k order of FA structure was selected to go forward with genomic predictions. For the third category of models, regardless the best-fit model, predictive accuracy was accounted for in all models (Models 17–20) to quantify the influence of modeling VCOV matrices on genomic predictions. To verify the advantage of FA structure under $\mathbf{G}$, the predictive accuracies of these models were compared with those of models that did not take into account information from correlated environments, hybrids within environments, or both (Table 1).

Models were also compared based on their predictive accuracy, computed via Pearson correlation between genetic estimated breeding value (GEBVs) and adjusted means from single-environment trial analysis estimated as E-BLUEs. Two distinct cross-validation strategies, CV1 and CV2, were implemented as proposed by Burgueño et al. (2012). In the first case (CV1), hybrids from validation set were deleted in all environments and predictions were performed based on the phenotypic information from relatives, through the genomic relationship matrix $\mathbf{A}$. The second strategy (CV2) highlights the situation where hybrids are phenotyped in some environments but missing in others. Predictions in this scenario take into account information from correlated environments if FA structure is included in Σ_g , and also information from relatives evaluated in multiple environments if the genomic relationship matrix $\mathbf{A}$ is also accounted for.

For both CV1 and CV2, a fivefold cross-validation procedure replicated 10 times was implemented to achieve the predictive accuracies, in which all single-cross hybrids with genotype (147) were randomly split into five non-overlapping groups, with four of them being training sets (80%) and one being a validation set (20%, around 30 hybrids), considered as not phenotyped. Therefore, all results are based on 20% missing hybrids within environment. Permutation of these five groups led to five possible training and validation datasets. To avoid bias due to the small sample size of the validation set in some scenarios, the GEBVs were predicted in all folds, and the Pearson correlation between observed and predicted values was estimated at the end of the cross-validation with the results from all permutations at once (Zhou, Vales, Wang, & Zhang, 2017).

Predictions in the CV2 scenario were accomplished as follows: (a) initially, one environment selected at random was considered as missing data for the validation set and predictions were performed for this environment based

on the training data sets, as explained above. Then, predictions were recorded in sequence for each one of the five permutations between training and validation sets. Next, the Pearson correlation between the vector of predictions and adjusted means was calculated. This process was repeated 10 times. After, (b) two environments selected at random were considered as missing data for the validation set (the same group of hybrids considered as missing within a given environment) and predictions were performed for these two environments, using the same approach described for one environment. This procedure was followed for all levels of missing environments: (a) one environment considered as missing data for the validation set, then (b) two environments, (c) three environments, (d) four environments, up to 11. Given that the data set has 12 environments, the last level of missing environments for the validation set in the CV2 scenario is with 11 missing environments at random. When all environments (12) were considered as missing data for the validation set, it was called CV1.

2.6 | Hybrid rank

The adjusted means of single-cross hybrids from single-environment trial analysis were used to access a hybrid's rank within each environment. For the top and bottom 20% hybrids, these ranks were compared with the ranks computed with GEBV for both CV1 and CV2 scenarios to access the coincidence index between ranks within environments. This procedure was also computed with ranks across environments.

3 | RESULTS

3.1 | Model selection

The AIC criterion for models from first category ranged from 4,327.86 [Model 7, $I_{\text{FA}(3)\text{-D}}$] to 4,704.08 [Model 1, $I_{\text{FA}(1)\text{-I}}$] and for the second category from 4,707.40 [Model 13, $A_{\text{FA}(1)\text{-D}}$] to 5,275.04 [Model 12, $A_{\text{FA}(4)\text{-I}}$]. Within each category, the inclusion of diagonal structure for residuals effects, allowing heterogeneous variances across environments, reduced both values of AIC and BIC criteria for the same k factor of FA models. However, the inclusion of the $\mathbf{A}$ matrix increased the values of both AIC and BIC, regardless of residual modeling (Table 1).

The percentage of genetic variance accounted for FA models ($\bar{v}$) ranged from 48.2 to 80.4% for the first category of models, and from 95.3 to 99.9% for the second category of models. Likewise for the AIC and BIC criteria, when $\Sigma_r = \mathbf{I} \otimes \mathbf{D}$, the $\bar{v}$ always increased for the same k

factor. The inclusion of the $\mathbf{A}$ matrix also increased $\bar{v}$. In the first category of models, the lowest AIC value was for Model 3, modeled with $\Sigma_g = \mathbf{FA}(3) \otimes \mathbf{I}$ and $\Sigma_r = \mathbf{I} \otimes \mathbf{D}$. The $\bar{v}$ of this model was 76%, superior to the ad hoc cut-off value of 70%. For the second category of models, that included genomic information, both Models 11 and 15, with $\Sigma_g = \mathbf{FA}(3) \otimes \mathbf{A}$, explained 99.9% of $\bar{v}$.

The BIC criterion for FA models always selected models with $k = 1$ and $\Sigma_r = \mathbf{I} \otimes \mathbf{D}$, being not informative to select FA multiplicative mixed models. Therefore, based on the AIC criterion and $\bar{v}$ in the first category of models (BLUE), Models 7 [$\mathbf{I}_{\mathbf{FA}(3)-\mathbf{D}}$] and 3 [$\mathbf{I}_{\mathbf{FA}(3)-\mathbf{I}}$] were selected to go forward for genomic prediction in METs. From the second category (GBLUP), although Models 15 [$\mathbf{A}_{\mathbf{FA}(3)-\mathbf{D}}$] and 11 [$\mathbf{A}_{\mathbf{FA}(3)-\mathbf{I}}$] with $k = 3$ do not have the lowest AIC values for their category, they were selected in order to make a fair comparison with selected models from the first category (BLUP) (Table 1).

For the third category of models, which accounts for genomic information with unrelated environments, Model 18 ($\mathbf{A}_{\mathbf{I}-\mathbf{D}}$) presented the lowest values for both AIC and BIC selection criteria. However, regardless of the best-fit model for this category, all third category models were tested for genomic prediction in METs in order to investigate the influence of modeling genetics and residuals effects in the predictive accuracy.

3.2 | Estimates of genetic parameters

Based on the likelihood ratio test (LRT), genetic variances were significantly greater than zero ($\sigma_g^2 > 0$) for most of the trials within environments, with exception of T1 within Environment 5 and T3 within Environments 1 and 4. For T3 within Environments 1 and 4, the CVs were >13%. The generalized measure of heritability ($\hat{H}^2$) ranged from .38 to .90, being zero for trial T3 within Environment 1, where the genetic variance component was estimated as equal to zero. Single-environment trial analysis also revealed that, at the same location, environments in which trials were sown in the first crop season were more productive than those sown in the second crop season (Table 2). This is an expected outcome due to the environmental differences between crop seasons in tropical areas (safra and safrinha).

Genetic correlations estimated from FA models varied considerably between pairs of environments. Models 11 [$\mathbf{A}_{\mathbf{FA}(3)-\mathbf{I}}$] and 15 [$\mathbf{A}_{\mathbf{FA}(3)-\mathbf{D}}$], that included the genomic additive relationship matrix $\mathbf{A}$, presented in general higher correlations than Models 3 [$\mathbf{I}_{\mathbf{FA}(3)-\mathbf{I}}$] and 7 [$\mathbf{I}_{\mathbf{FA}(3)-\mathbf{D}}$], in which hybrids were considered nongenetically related to each other. For example, the lowest value of pairwise correlation between environments for both Models 11 [$\mathbf{A}_{\mathbf{FA}(3)-\mathbf{I}}$] and 15 [$\mathbf{A}_{\mathbf{FA}(3)-\mathbf{D}}$] was 0.21, and for Models 3 [$\mathbf{I}_{\mathbf{FA}(3)-\mathbf{I}}$] and 7

[$\mathbf{I}_{\mathbf{FA}(3)-\mathbf{D}}$] was 0.08 and 0.06, respectively. Residuals modeling changed the magnitude of correlations for these models, being slightly higher for models that allowed homogeneous residuals variance across environments ($\Sigma_r = \mathbf{I} \otimes \mathbf{I}$). Overall, the estimated genetic and additive correlations between environments were reasonably high, with an average pairwise correlation of 0.51, 0.47, 0.61, and 0.58 for Models 3 [$\mathbf{I}_{\mathbf{FA}(3)-\mathbf{I}}$], 7 [$\mathbf{I}_{\mathbf{FA}(3)-\mathbf{D}}$], 11 [$\mathbf{A}_{\mathbf{FA}(3)-\mathbf{I}}$], and 15 [$\mathbf{A}_{\mathbf{FA}(3)-\mathbf{D}}$], respectively. Based on the average correlation between a given environment and all the others, Environments 6, 7, and 11 had the lowest average values, and Environments 5 and 9 the greatest averaged values of correlation (Figure 1).

Principal component analysis showed good heterotic group consistency of the 143 inbred lines used as parents of the single-cross hybrids (Figure 2). Using SNP markers information from GBS, the genotypes of the single-cross hybrids were inferred and, due the good consistency of inbred lines heterotic groups, most hybrids were not closely related (Figure 3). The four inbred lines used as testers produced 48, 38, 23 and 20 single-cross hybrids, respectively. Within each one of these groups, hybrids are half-sibs and their expected average coefficient of relationship is 0.25 (Lynch & Walsh, 1998). From the genomic relationship matrix $\mathbf{A}$, on average, these coefficients were 0.27, 0.29, 0.36, and 0.30, respectively.

3.3 | Predictive accuracy

When hybrids from validation set were considered as missing in all environments (CV1), models from second and third categories, that included genomic information, presented similar results. Models 11 [$\mathbf{A}_{\mathbf{FA}(3)-\mathbf{I}}$], 17 ($\mathbf{A}_{\mathbf{I}-\mathbf{I}}$), and 19 ($\mathbf{A}_{\mathbf{D}-\mathbf{I}}$), with $\Sigma_r = \mathbf{I} \otimes \mathbf{I}$, had average predictive accuracies of 0.261, 0.264, and 0.232, respectively. Models 15 [$\mathbf{A}_{\mathbf{FA}(3)-\mathbf{D}}$], 18 ($\mathbf{A}_{\mathbf{I}-\mathbf{D}}$), and 20 ($\mathbf{A}_{\mathbf{D}-\mathbf{D}}$), with $\Sigma_r = \mathbf{I} \otimes \mathbf{D}$, had average predictive accuracies of 0.273, 0.274, and 0.262, respectively (Figure 4, Supplemental Tables S1–S4). Hence, for the CV1 scenario, when the validation set was considered as missing in all environments, models that allowed heterogeneous residuals variance across environments performed slightly better. On the other hand, models from the first category were not able to predict in the CV1 scenario. These models do not borrow information from relatives within environments through the additive genomic relationship matrix $\mathbf{A}$, given that $\Sigma_g = \mathbf{FA}(3) \otimes \mathbf{I}$.

The accommodation of genetic correlation between environments through $\mathbf{FA}$ almost doubles the predictive accuracy of the models until five missing environments at random, independently of residuals modeling. Taking Model 20 ($\mathbf{A}_{\mathbf{D}-\mathbf{D}}$) as the baseline model and from 1 to 11 missing environments at random (all levels of CV2),

TABLE 2 Results from single-environment trial analysis (Model 1) for grain yield (GY), CV, generalized measure of heritability ($\hat{H}^2$), and genetic (σ_g^2), block (σ_b^2), and residual (σ_e^2) variance components

Environment	Trial	GY t ha ⁻¹	CV %	$\hat{H}^2$	σ_g^2	σ_b^2	σ_e^2
Env 1	T1	0.82	8.92	0.64	0.83	0.00ns ^a	0.93
	T2	11.14	7.08	0.82	1.43	0.00ns	0.62
	T3	11.25	13.12	0.00	0.00ns	0.04ns	2.14
Env 2	T1	3.78	8.99	0.81	0.24	0.00ns	0.12
	T2	3.71	11.90	0.78	0.34	0.00ns	0.20
	T3	4.18	8.70	0.72	0.19	0.02ns	0.13
Env 3	T1	9.47	11.50	0.57	0.86	0.19	1.19
	T2	9.74	11.10	0.72	1.53	0.00ns	1.17
	T3	9.94	10.96	0.64	1.05	0.00ns	1.19
Env 4	T1	8.40	13.50	0.44	0.53	0.04ns	1.29
	T2	8.68	10.60	0.75	1.32	0.02ns	0.85
	T3	8.92	13.22	0.35	0.41ns	0.40	1.39
Env 5	T1	6.10	14.30	0.32	0.19ns	0.02ns	0.76
	T2	6.57	14.60	0.38	0.30	0.12ns	0.93
	T3	7.43	12.01	0.46	0.34	0.00ns	0.80
Env 6	T1	7.33	11.60	0.72	1.05	0.30	0.72
	T2	7.14	16.60	0.51	0.78	0.27	1.41
	T3	6.21	13.94	0.61	0.68	0.33	0.75
Env 7	T1	8.35	12.90	0.78	2.19	0.15ns	1.16
	T2	8.30	13.50	0.76	1.97	0.00ns	1.26
	T3	7.04	13.11	0.75	1.37	0.12ns	0.85
Env 8	T1	11.84	8.20	0.53	0.56	0.08ns	0.94
	T2	10.83	9.73	0.51	0.60	0.09ns	1.11
	T3	12.26	7.97	0.53	0.54	0.01ns	0.96
Env 9	T1	8.46	12.30	0.72	1.27	0.16	0.89
	T2	8.37	12.30	0.75	1.29	0.14ns	0.82
	T3	8.11	10.26	0.74	1.03	0.00ns	0.74
Env 10	T1	7.69	12.20	0.67	0.35	0.06	0.31
	T2	7.32	11.90	0.86	0.86	0.07ns	0.26
	T3	8.36	9.60	0.90	1.56	0.01ns	0.36
Env 11	T1	4.57	11.50	0.46	0.40	0.00ns	0.94
	T2	4.27	9.17	0.61	0.49	0.04ns	0.59
	T3	6.26	10.77	0.66	0.75	0.01ns	0.76
Env 12	T1	4.08	9.16	0.72	0.20	0.03	0.14
	T2	4.00	12.80	0.62	0.22	0.02ns	0.26
	T3	7.32	7.51	0.73	0.45	0.09ns	0.30

^a ns, Variance component statistically equal to zero ($\sigma_g^2 = 0$ or $\sigma_b^2 = 0$), based in the likelihood ratio test (LRT) with $\alpha = .05$.

respectively, Models 3 [$I_{FA(3)-I}$] and 7 [$I_{FA(3)-D}$] had average predictive accuracies of 70.30, 62.92, 62.38, 61.05, 52.11, 58.66, 40.77, 36.35, 37.78, 12.63, and -19.67% superior or inferior to the baseline model. Therefore, for first category models that did not account genomic information, borrowing information from correlated environments increased

the predictive accuracy up to 50% over the baseline model until six missing environments at random. Six environments represent a reduction of 50% in phenotyping. For Models 11 [$A_{FA(3)-I}$] and 15 [$A_{FA(3)-D}$], from the second category, predictive accuracies were 71.38, 70.46, 68.49, 70.87, 53.11, 71.07, 57.15, 50.89, 56.88, 34.13, and 19.53% superior

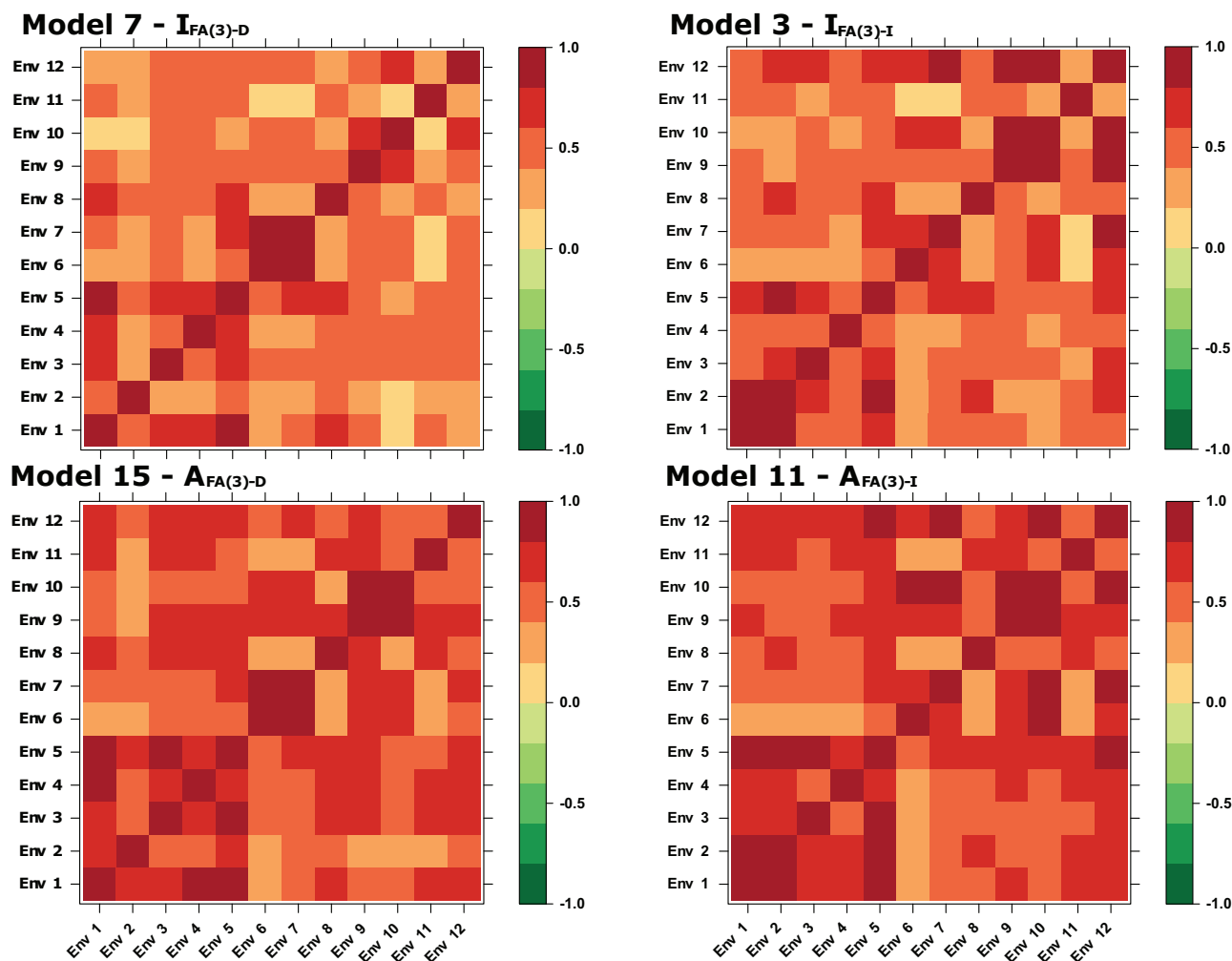


FIGURE 1 Heatmap of genetic and additive correlations between environments for factor analytic (FA) models without incorporating the genomic relationship matrix, Model 7 [$I_{FA(3)-D}$] and Model 3 [$I_{FA(3)-I}$]; and for FA models with genomic relationship information, Model 15 [$A_{FA(3)-I}$] and Model 11 [$A_{FA(3)-D}$]

to the baseline model, respectively. Hence, for the second category models that accounted for genetic correlation between environments and genomic information, predictions were up to 50% superior to the reference model until nine missing environments at random, representing a reduction of 75% in phenotyping (Figure 4, Supplemental Tables S1–S4).

Models from first and second categories had similar performance from one to five missing environments at random, and as the number of missing environments increased, models from the second category that also explored genomic information performed better. From 6 to 11 missing environments at random, respectively, the GBLUP Models 11 [$A_{FA(3)-I}$] and 15 [$A_{FA(3)-D}$] had, on average, superior performance of 7.83, 11.63, 10.66, 13.87, 19.01, and 48.8%, compared with BLUP Models 3 [$I_{FA(3)-I}$] and 7 [$I_{FA(3)-D}$], from the first category (Figure 4, Supplemental Tables S1–S4).

Second category Models 11 [$A_{FA(3)-I}$] and 15 [$A_{FA(3)-D}$] were able to keep predictions up to 0.4 until eight missing environments at random, and first category Models 3 [$I_{FA(3)-I}$] and 7 [$I_{FA(3)-D}$] were able to keep predictions up to 0.4 until five missing environments at random. This highlights the influence of missing environments for models that do not account genomic information (Figure 4, Supplemental Tables S1–S4).

In general, as the number of missing environments became larger, the predictive accuracies got smaller for FA models, mainly for BLUP models that did not account genomic information. Third category models, not modeled with FA structure but with the genomic relationship matrix **A**, had similar predictive accuracies across all levels of missing environments, including the CV1 scenario. For these models, heterogeneous residual variances across environments performed slightly better in all levels of missing environment. For example, in a pairwise

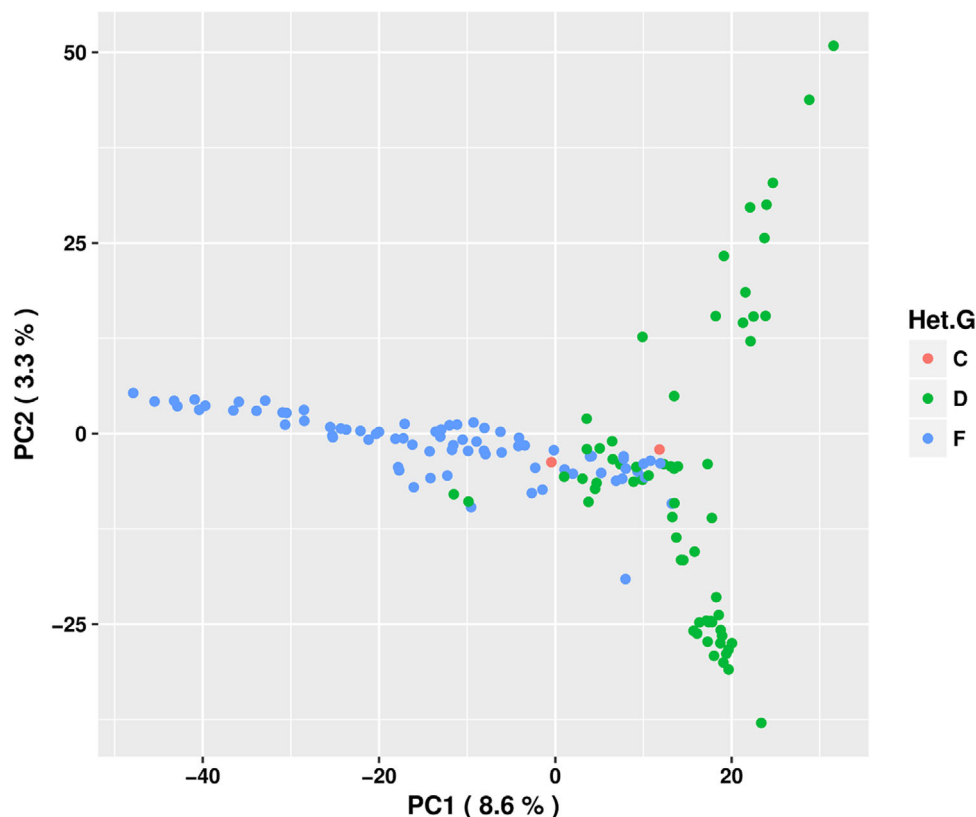


FIGURE 2 Scatter plot of two first principal components (PCs) from principal component analysis (PCA), based on 29,515 single nucleotide polymorphism (SNP) markers for 143 inbred lines used as parents of maize single-cross hybrids. The legend shows Heterotic Groups C (Group C), D (dent), and F (flint)

comparison between Models 19 (A_{D-I}) and 20 (A_{D-D}), the latter had an advantage of, on average, 14.3% in predictive accuracy. For Models 17 (A_{I-I}) and 18 (A_{I-D}), this advantage decreased to 3.6% for the latter. In this category, Model 18 (A_{I-D}) had the lowest values of both AIC and BIC, but in terms of prediction, no differences were found between this model and Model 20 (A_{D-D}), which also allowed for heterogeneous residual variances. Overall, Models 17 (A_{I-I}), 18 (A_{I-D}), 19 (A_{D-I}), and 20 (A_{D-D}) had predictive accuracies ranging from 0.240 to 0.284, 0.258 to 0.280, 0.211 to 0.253, and 0.242 to 0.276, respectively (Figure 4, Supplemental Tables S1–S4).

The predictive accuracies of models that borrowed information from correlated environments were superior within all environments. This superiority was less evident within Environments 6, 7, and 11. These environments had the lowest values of average correlation among themselves and the other environments. On the other hand, Environments 5 and 9 had the highest values of average correlation, and the differences in predictive accuracy between models with and without **FA** structure were more evident (Figure 5).

3.4 | Changes in hybrid rank

For selection across environments, BLUP Models 3 [$I_{FA(3)-I}$] and 7 [$I_{FA(3)-D}$] had the highest values of coincidence index of all models and levels of CV2, for both the top and bottom 20% hybrids. For these two models, until 10 missing environments at random, with one exception, all values of coincidence were >80% in comparison with the baseline rank obtained from single-environment trial analysis. The GBLUP Models 11 [$A_{FA(3)-I}$] and 15 [$A_{FA(3)-D}$] from the second category, with the exception of three values, presented values of coincidence >80% until eight missing environments at random, also for both the top and bottom 20% hybrids.

In the most challenging scenario where hybrids from validation set were removed from all environments (CV1), GBLUP models from the second category modeled with **FA** presented coincidence index ranging from 37 to 47% for both the top and bottom 20% hybrids, when selecting across environments. On the other hand, GBLUP models from the third category that did not account for correlation between environments presented, on average, 20 and

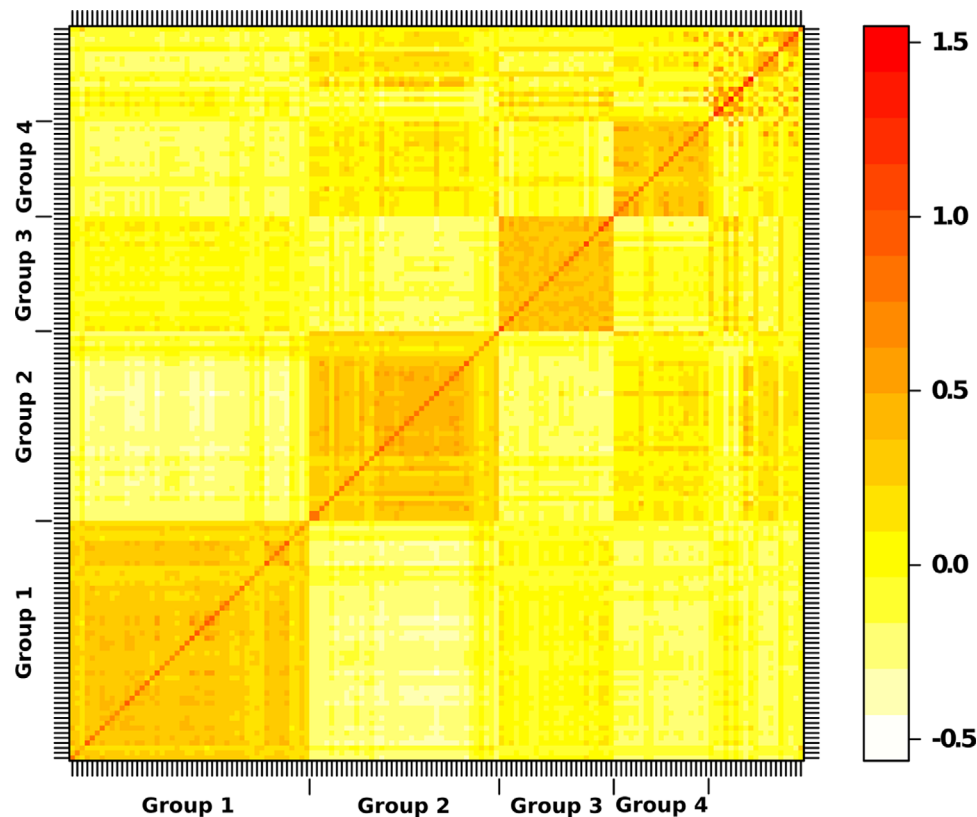


FIGURE 3 Heatmap of genomic relationship matrix **A** for 147 maize single-cross hybrids ordered by four groups of half-sibs, each group synthesized by a given tester. From Group 1 to 4, the sample size consists of 48, 38, 23, and 20 single-cross hybrids, with an average relatedness coefficient of 0.27, 0.29, 0.36, and 0.30, respectively. The remaining hybrids, without label, were synthesized by other testers

16% coincidence index for the top and bottom 20% hybrids, respectively, across all levels of missing environment (Supplemental Figure S2).

For selections within environments, models from the first and second categories, modeled with **FA** and hence including correlation between environments, also showed advantage over models that ignored it. When one environment was missing at random, in the first level of CV2, BLUP Models 3 [$I_{FA(3)-I}$] and 7 [$I_{FA(3)-D}$] presented values of coincidence index >50% up to nine missing environments at random, for both the top and bottom 20% hybrids. For second category GBLUP Models 11 [$A_{FA(3)-I}$] and 15 [$A_{FA(3)-D}$], the coincidence index was >50% up to seven missing environments at random. The GBLUP from third category that ignored genetic correlation between environments had similar performance of coincidence index for both selections across and within environments (Supplemental Figure S3).

4 | DISCUSSION

The inclusion of **FA** structure in Σ_g resulted in better goodness of fit of models, highlighting the importance of taking

into account information from correlated environments. For biological reasons, some correlation between yield performance of hybrids across environments is expected in a breeding program, rather than homogeneous variances and the absence of correlations. On the other hand, models with a **FA** structure in **G** and with homogeneous structure in **R** showed increasing values for both AIC and BIC criteria. These results indicate that, in terms of model selection criteria, modeling **G** under a **FA** structure along with heterogeneous residual variances across environments in **R** tends to result in a better representation of VCOV structures for leveraging GE trends across environments.

The BIC criterion always penalized FA models with high k order. Factor analytic models can be considered as nested models—for example, from $k = 1$ to $k = 4$ [**FA**(1), **FA**(2), **FA**(3), **FA**(4), and so on], as explained by Sorensen and Gianola (2002). The BIC is a well-defined criterion for non-nested models. Also, BIC penalizes most models with more parameters, which can lead to the choice of models that may underfit. Smith et al. (2015) observed the same pattern of BIC criterion to select FA models. The percentage of genetic variance ($\bar{v}$) accounted by each k factor, as expected, increased as k became greater (Table 1). Therefore, in the GS context for METs using FA multiplicative

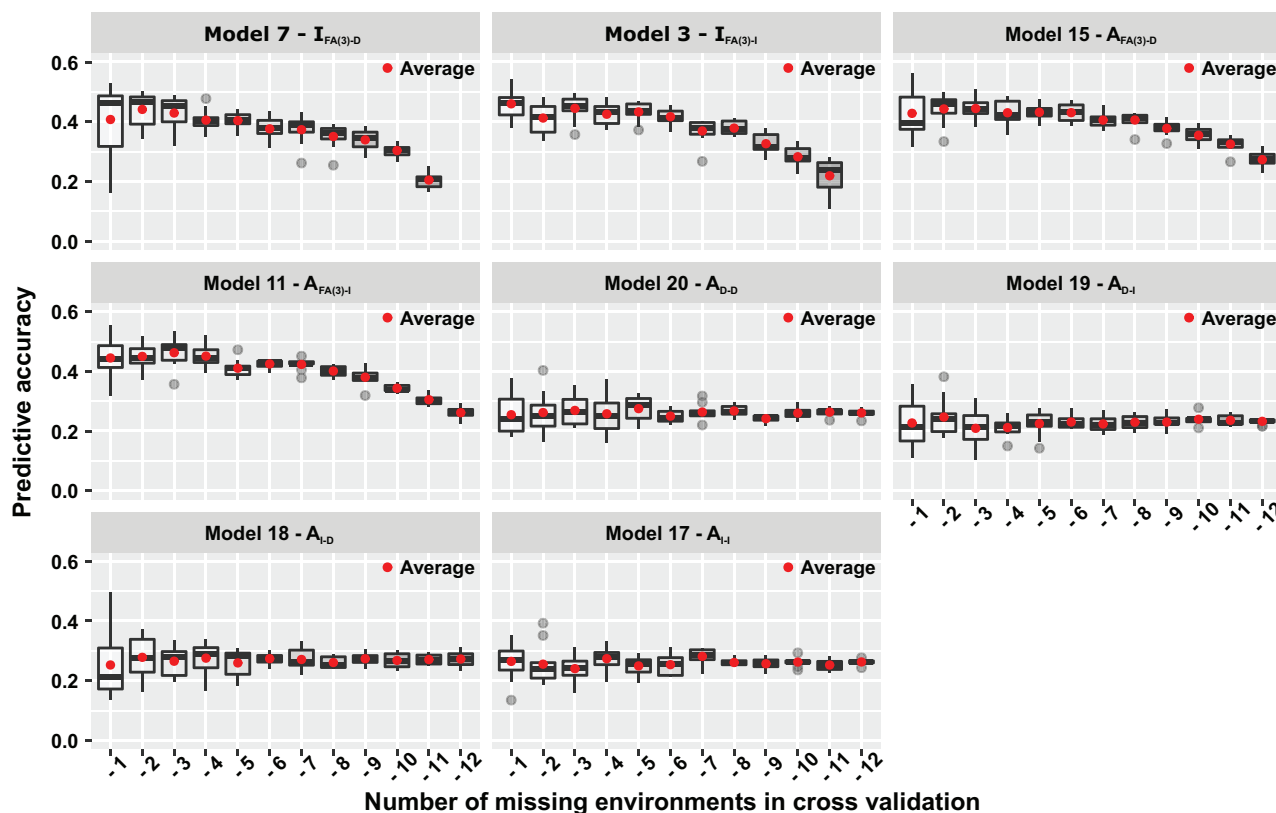


FIGURE 4 Mean predictive accuracies across environments of 10 times fivefold cross-validation, for factor analytic (FA) models without incorporating the genomic relationship matrix (Models 3 [$I_{FA(3)-I}$] and 7 [$I_{FA(3)-D}$]), for FA models with genomic relationship information (Model 11 [$A_{FA(3)-I}$] and Model 15 [$A_{FA(3)-D}$]), and for models assuming no correlation between environments but including genomic relationship information (Model 17 [A_{I-I}], Model 18 [A_{I-D}], Model 19 [A_{D-I}], and Model 20 [A_{D-D}])

mixed models, it is important to consider more than one criterion to select the best fit model, as also pointed out by Pastina et al. (2012) and Ferrão, Ferrão, Ferrão, Francisco, and Garcia (2017).

Considering residual heterogeneous variance resulted in improvements in the goodness of fit for all models. It is realistic to assume that each environment presented its own source of variation that cannot be explained by assuming a common normal distribution due to climate conditions, plant diseases, or any other sources of variation not considered by the model. However, residual modeling with homogeneous or heterogeneous variances did not improve the predictive accuracies for models under FA in G. Similar results were found by Burgueño et al. (2012), where genetic effects were more important than residuals effects.

The gain in predictive accuracy obtained in CV2 over CV1 with FA models is directly related to the ability of these models to borrow information from correlated environments. Phenotypic performance is expected to vary according to the environment due to the GE interaction (Ferrão, Ferrão, et al., 2017; Zhang et al., 2015). As a consequence, the relative performance and rank of genotypes may vary according to the environment. Our results

showed that models that allow heterogeneous variance components and genetic correlations across environments are more realistic and capable of capturing these patterns, thus improving maize single-cross hybrids prediction in the framework of METs. Therefore, prediction of newly lines (yet-to-be phenotyped), a situation created by the CV1 scenario, was more challenging than predicting single-cross hybrids that were evaluated in some environments but missing in others (CV2) under FA models.

The magnitude of correlations between environments is also an important parameter to be considered, which can directly influence the ability of models to borrow information from correlated environments. The matrices of genetic and additive correlation across environments for models with and without genomic information (BLUP and GBLUP), respectively, confirmed the high association between environments. Similar results were found by Burgueño et al. (2012), Crossa et al. (2014), and Lopez-Cruz et al. (2015), where including information from METs increased the prediction accuracy of models.

Several studies included FA structure to account correlations between environments (Burgueño, Crossa, Cotes, Vicente, & Das, 2011; Cullis, Jefferson, Thompson, &

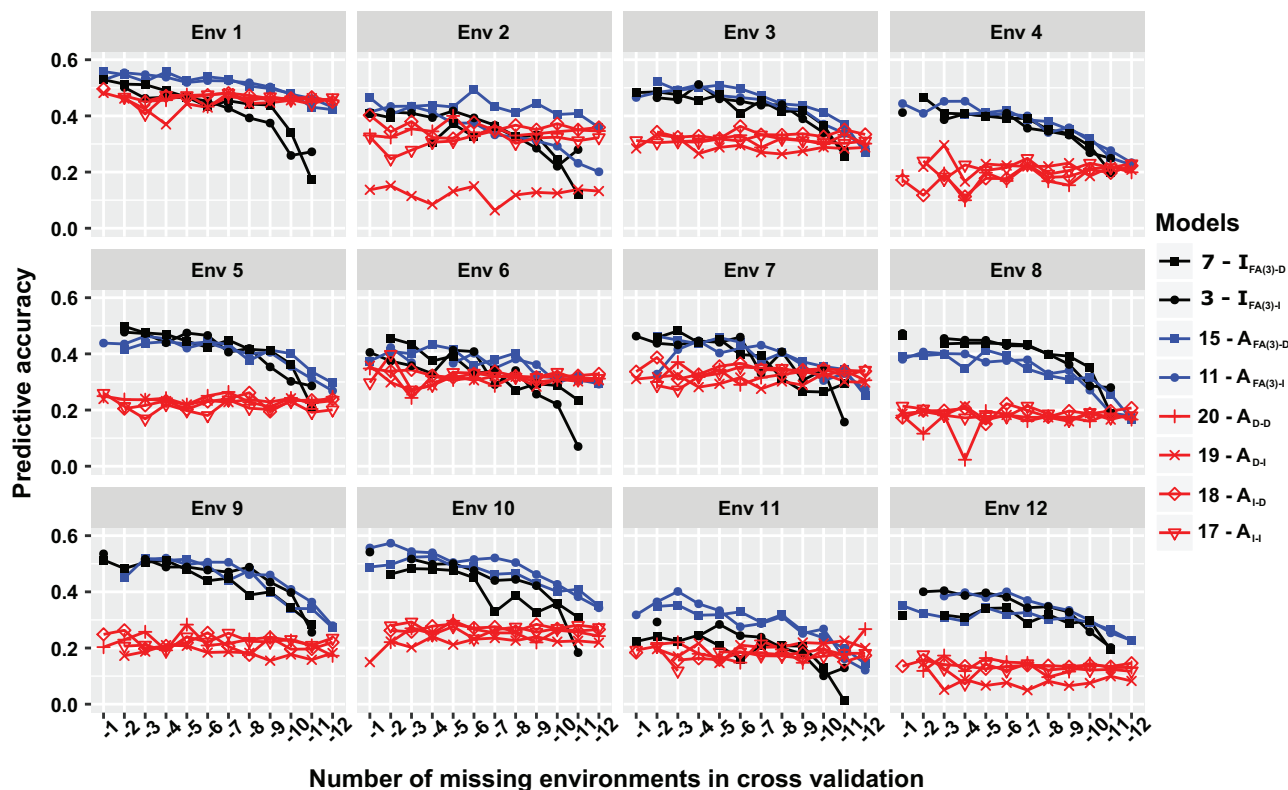


FIGURE 5 Mean predictive accuracies across environments of 10 times fivefold cross-validation, for all models within environments. From left to right, “Env 1” to “Env 12” are abbreviations for Environments 1–12, respectively

Smith, 2014; Dias, Gezan, Guimarães, Parentoni, et al., 2018; Kelly et al., 2007; Smith et al., 2015). Specifically for genomic prediction and selection, it has been used for wheat (*Triticum aestivum* L.; Burgueño et al., 2012; Dawson et al., 2013; Rutkoski et al., 2015), barley (*Hordeum vulgare* L.; Malosetti, Bustos-Korts, Boer, & Van Eeuwijk, 2016; Oakey et al., 2016), and maize (Dias, Gezan, Guimarães, Nazarian, et al., 2018; Schulz-Streeck et al., 2013). Phenotyping in an MET is routine in plant breeding programs, and although the FA structure is an approximation to the unstructured VCOV matrix, it provides reliable information to access the performance of single-cross maize hybrids within and across environments (Kelly et al., 2007; Smith et al., 2001). Our results emphasized the great flexibility of FA models to handle low, moderate, and high levels of missing data in the framework of METs, as also pointed out by Elias et al. (2016).

Cost reduction and improved selection are examples of how GS can reshape breeding programs (Hickey et al., 2017), but its application depends on the ability of models to predict real situations faced in the breeding programs (Ferrão, Ferrão, et al., 2017; Ferrão, Ortiz, & Garcia, 2017). Our results emphasized that even in high levels of unbalancing, models that account correlation between environments and genomic information can be a valuable tool to predict breeding values. As an example, consider the cost

of GBS genotyping in a sequencing coverage ($\times$) of $2x$ as US\$25.00 (Gorjanc et al., 2017) per inbred line, and the cost of one maize yield-trial plot as \$13.00 (Tech Services, 2018). Then, the needed budget for a breeding program with similar data as presented in this study would be \$3,575 for genotyping the 143 inbred lines used as parents, and \$45,864 for phenotyping the 147 single-cross hybrids in 12 environments. Using a GBLUP model that embraces the correlation between environments plus genomic information, it was shown that until eight missing environments at random (66% of missing data for the validation set), predictions of untested single-cross hybrids were up 0.40 with an average coincidence index of at least 80 and 50% for selections across and within environments, respectively. Hence, a reduction of breeding costs by 8.33% or \$3,822 can be achieved if hybrids were predicted just in one environment. This amount is sufficient to cover the costs of inbred lines genotyping (\$3,575). For the following levels of missing environment, the reduction in the amount of needed budget is linear. For hybrids predicted in two environments, the reduction would be by 16.67% (or \$7,644); for three environments, the reduction would be by 25.00%; and so on, until a reduction by 66.67% (or \$30,576) for prediction at eight environments.

Regardless of the level of missing data for genomic prediction, any reduction of the total budget of hybrid

phenotyping could be allocated to optimize the breeding program. An interesting way to allocate the saved budget is the development of newly synthetic populations for the synthesis of novel inbred lines. For example, the cost for producing a newly synthetic population obtained from 10 inbred lines—including the cost of labor, time demanded, and nursery space for crosses—is on average \$1,200 (Dr. David Benson, personal communication, 2018). Then, if hybrids were predicted in one environment, the saved budget could be used to produce three newly synthetic populations, or to cover the costs of inbred lines genotyping, as mentioned above. Other possibilities to allocate the saved budget could be (a) to be used in the evaluation of more hybrids at the intermediate stage of the breeding process, and therefore increase the intensity of selection; (b) to obtain genotypic data for newly inbred lines and hence predict the performance of newly developed single-cross hybrids; or just (c) to reduce costs. It was also noted by Krchov and Bernardo (2015) that once GS is implemented in the breeding process, the reduction in the amount of phenotyping labor leads to a better quality of the field data collected, enhancing the effectiveness of selection.

Maize is a well-known allogamous species, and single-cross hybrids express heterosis (Hallauer, Carena, & Miranda Filho, 2010). This means that it is also worthwhile to investigate the inclusion of a dominance relationship matrix into GBLUP models. If the dominance component is important regarding the observed genetic variation, the use of both additive plus dominance values could boost predictions (Oakey et al., 2016). However, results of this approach showed no improvement in hybrids predictions or in hybrids ranking (data not shown), although some exciting results have been reported in the literature (Dias, Gezan, Guimarães, Nazarian, et al., 2018; dos Santos, Vasconcellos, Pires, Balestre, & Von Pinho, 2016). Some bottlenecks associated with this result could be that the estimation of dominance effects requires specific mating designs (Nazarian & Gezan, 2016), and confounding between additive and nonadditive genetic components (Lee, Goddard, Visscher, & Van Der Werf, 2010; Moghaddar & van der Werf, 2017; Muñoz et al., 2014; Nazarian & Gezan, 2016). An increase in the number of evaluated hybrids could overcome this issue. It is also worth mentioning that selection and assortative mating could result in a low accurate estimation of nonadditive variance components (Hill, Goddard, & Visscher, 2008).

In summary, we obtained encouraging genomic prediction accuracies of tropical maize single-cross hybrids by accounting for genetic correlation between environments and genomic information in GBLUP models. The approach used in our manuscript can be expanded to other crops in which METs play an important role in the breeding process. Future research on the integration of optimized

training sets and crop growth models (Bustos-Korts et al., 2019; Rincen et al., 2017) that combine ecophysiological and genetics modeling seems to be a promising way to deal with unbalanced trials and to better understand and characterize the dynamics of GE in the framework of genomic predictions.

5 | CONCLUSION

This study demonstrated that (a) the inclusion of **FA** structure boosted the predictive accuracy of untested maize single-cross hybrids in the framework of MET, regardless residuals modeling; (b) models that included correlation between environments plus genomic information achieved higher predictive accuracy in elevated levels of missing environments over models that assumed unrelated hybrids; and (c) high levels of predictive accuracy were found with moderated to low levels of missing environments.

ACKNOWLEDGMENTS

This research was supported by CAPES (Coordination for the Improvement of Higher Education Personnel), Embrapa (Brazilian Agricultural Research Corporation) Maize and Sorghum. Also, A.A.F.G received a productivity scholarship from CNPq (National Council for Scientific and Technological Development), and K.O.G.D received a grant from FAPESP (Fundação de Amparo à Pesquisa do Estado de São Paulo, Grant 2016/12977-7). The authors wish to thank Rafael Storto Nalin, Rodrigo Rampazo Amadeu, and Gabriel de Siqueira Gesteira for their contributions to the initial development of the manuscript.

CONFLICTS OF INTEREST

The authors declare that there are no conflicts of interest regarding the publication of this manuscript.

ORCID

Matheus Dalsente Krause  <https://orcid.org/0000-0003-2411-9287>

Kaio Olímpio das Graças Dias  <https://orcid.org/0000-0002-9171-1021>

Jhonathan Pedrosa Rigal dos Santos  <https://orcid.org/0000-0001-9030-702X>

Amanda Avelar de Oliveira  <https://orcid.org/0000-0002-5595-6697>

Lauro José Moreira Guimarães  <https://orcid.org/0000-0003-2413-2481>

Maria Marta Pastina  <https://orcid.org/0000-0002-9112-3457>

Gabriel Rodrigues Alves Margarido  <https://orcid.org/0000-0002-2327-0201>

Antonio Augusto Franco Garcia  <https://orcid.org/0000-0003-0634-3277>

REFERENCES

- Acosta-Pech, R., Crossa, J., de los Campos, G., Teyssèdre, S., Claustres, B., Pérez-Elizalde, S., & Pérez-Rodríguez, P. (2017). Genomic models with genotype $\times$ environment interaction for predicting hybrid performance: An application in maize hybrids. *Theoretical and Applied Genetics*, 130, 1431–1440. <https://doi.org/10.1007/s00122-017-2898-0>
- Akaike, H. (1974). A new look at the statistical model identification. *IEEE Transactions on Automatic Control*, 19, 716–723. <https://doi.org/10.1109/TAC.1974.1100705>
- Amadeu, R. R., Cellon, C., Olmstead, J. W., Garcia, A. A. F., Resende, M. F. R., & Muñozal, P. R. (2016). AGHmatrix: R package to construct relationship matrices for autotetraploid and diploid species: A blueberry example. *The Plant Genome*, 9(3). <https://doi.org/10.3835/plantgenome2016.01.0009>
- Bernardo, R. (1994). Prediction of maize single-cross performance using RFLPs and information from related hybrids. *Crop Science*, 34, 20–25. <https://doi.org/10.2135/cropsci1994.001183X003400010003x>
- Bernardo, R. (1995). Genetic models for predicting maize single-cross performance in unbalanced yield trial data. *Crop Science*, 35, 141–147. <https://doi.org/10.2135/cropsci1995.001183X003500010026x>
- Bernardo, R. (2008). Molecular markers and selection for complex traits in plants: Learning from the last 20 years. *Crop Science*, 48, 1649–1664. <https://doi.org/10.2135/cropsci2008.03.0131>
- Browning, B. L., & Browning, S. R. (2016). Genotype imputation with millions of reference samples. *American Journal of Human Genetics*, 98, 116–126. <https://doi.org/10.1016/j.ajhg.2015.11.020>
- Burgueño, J., Crossa, J., Cotes, J. M., Vicente, F. S., & Das, B. (2011). Prediction assessment of linear mixed models for multi-environment trials. *Crop Science*, 51, 944–954. <https://doi.org/10.2135/cropsci2010.07.0403>
- Burgueño, J., de los Campos, G., Weigel, K., & Crossa, J. (2012). Genomic prediction of breeding values when modeling genotype $\times$ environment interaction using pedigree and dense molecular markers. *Crop Science*, 52, 707–719. <https://doi.org/10.2135/cropsci2011.06.0299>
- Bustos-Korts, D., Malosetti, M., Chenu, K., Chapman, S., Boer, M. P., Zheng, B., & van Eeuwijk, F. A. (2019). From QTLs to adaptation landscapes: Using genotype-to-phenotype models to characterize G $\times$ E over time. *Frontiers in Plant Science*, 10, 1–23. <https://doi.org/10.3389/fpls.2019.01540>
- Butler, D., Cullis, B. R., Gilmour, R., & Gogel, B. J. (2009). *ASReml-R reference manual* (Release 3). Tech. Rep. Brisbane, QLD, Australia: Queensland Department of Primary Industries.
- Cobb, J. N., Juma, R. U., Biswas, P. S., Arbelaez, J. D., Rutkoski, J., Atlin, G., ... Ng, E. H. (2019). Enhancing the rate of genetic gain in public-sector plant breeding programs: Lessons from the breeder's equation. *Theoretical and Applied Genetics*, 132, 627–645. <https://doi.org/10.1007/s00122-019-03317-0>
- Crossa, J., De Los Campos, G., Maccaferri, M., Tuberosa, R., Burgueño, J., & Pérez-Rodríguez, P. (2016). Extending the marker $\times$ environment interaction model for genomic-enabled prediction and genome-wide association analysis in durum wheat. *Crop Science*, 56, 2193–2209. <https://doi.org/10.2135/cropsci2015.04.0260>
- Crossa, J., Pérez, P., Hickey, J., Burgueño, J., Ornella, L., Cerón-Rojas, J., ... Mathews, K. (2014). Genomic prediction in CIM-MYT maize and wheat breeding programs. *Heredity*, 112, 48–60. <https://doi.org/10.1038/hdy.2013.16>
- Crossa, J., Pérez-Rodríguez, P., Cuevas, J., Montesinos-López, O., Jarquín, D., de los Campos, G., & Varshney, R. K. (2017). Genomic selection in plant breeding: Methods, models, and perspectives. *Trends in Plant Science*, 22, 961–975. <https://doi.org/10.1016/j.tplants.2017.08.011>
- Cuevas, J., Crossa, J., Soberanis, V., Pérez-Elizalde, S., Pérez-Rodríguez, P., De Los Campos, G., ... Burgueño, J. (2016). Genomic prediction of genotype $\times$ environment interaction kernel regression models. *Plant Genome*, 9(3). <https://doi.org/10.3835/plantgenome2016.03.0024>
- Cullis, B. R., Jefferson, P., Thompson, R., & Smith, A. B. (2014). Factor analytic and reduced animal models for the investigation of additive genotype-by-environment interaction in outcrossing plant species with application to a *Pinus radiata* breeding programme. *Theoretical and Applied Genetics*, 127, 2193–2210. <https://doi.org/10.1007/s00122-014-2373-0>
- Cullis, B. R., Smith, A. B., & Coombes, N. E. (2006). On the design of early generation variety trials with correlated data. *Journal of Agricultural, Biological, and Environmental Statistics*, 11, 381–393. <https://doi.org/10.1198/108571106X154443>
- Dawson, J. C., Endelman, J. B., Heslot, N., Crossa, J., Poland, J., Dreisigacker, S., ... Jannink, J. -L. (2013). The use of unbalanced historical data for genomic selection in an international wheat breeding program. *Field Crops Research*, 154, 12–22. <https://doi.org/10.1016/j.fcr.2013.07.020>
- Dias, K. O. d. G., Gezan, S. A., Guimarães, C. T., Nazarian, A., da Costa e Silva, L., Parentoni, S. N., ... Pastina, M. M. (2018). Improving accuracies of genomic predictions for drought tolerance in maize by joint modeling of additive and dominance effects in multi-environment trials. *Heredity*, 121, 24–37. <https://doi.org/10.1038/s41437-018-0053-6>
- Dias, K. O. d. G., Gezan, S. A., Guimarães, C. T., Parentoni, S. N., de Oliveira Guimarães, P. E., Carneir, N. P., ... Pastina, M. M. (2018b). Estimating genotype $\times$ environment interaction for and genetic correlations among drought tolerance traits in maize via factor analytic multiplicative mixed models. *Crop Science*, 58, 72–83. <https://doi.org/10.2135/cropsci2016.07.0566>
- dos Santos, J. P. R., Vasconcellos, R. C. De Castro, Pires, L. P. M., Balestre, M., & Von Pinho, R. G. (2016). Inclusion of dominance effects in the multivariate GBLUP model. *PLOS ONE*, 11(4). <https://doi.org/10.1371/journal.pone.0152045>
- Elias, A. A., Robbins, K. R., Doerge, R. W., & Tuinstra, M. R. (2016). Half a century of studying genotype $\times$ environment interactions in plant breeding experiments. *Crop Science*, 56, 2090–2105. <https://doi.org/10.2135/cropsci2015.01.0061>
- Elshire, R. J., Glaubitz, J. C., Sun, Q., Poland, J. A., Kawamoto, K., Buckler, E. S., & Mitchell, S. E. (2011). A robust, simple genotyping-by-sequencing (GBS) approach for high diversity species. *PLOS ONE*, 6(5). <https://doi.org/10.1371/journal.pone.0019379>
- Ferrão, L. F. V., Ferrão, R. G., Ferrão, M. A. G., Francisco, A., & Garcia, A. A. F. (2017). A mixed model to multiple harvest-location trials applied to genomic prediction in *Coffea canephora*. *Tree Genetics & Genomes*, 13(5). <https://doi.org/10.1007/s11295-017-1171-7>
- Ferrão, L. F. V., Ortiz, R., & Garcia, A. A. F. (2017). Genomic selection: State of the art. In H. Campos & P. D. S. Caligari (Eds.),

- Genetic improvement of tropical crops* (pp. 19–54). Cham, Switzerland: Springer. https://doi.org/10.1007/978-3-319-59819-2_2
- Fisher, R. A. (1918). The correlation between relatives on the supposition of mendelian inheritance. *Transactions of the Royal Society of Edinburgh*, 52, 399–433. <https://doi.org/10.1017/S0080456800012163>
- Fritsche-Neto, R., Gonçalves, M. C., Vencovsky, R., & De Souza Junior, C. L. (2010). Prediction of genotypic values of maize hybrids in unbalanced experiments. *Crop Breeding and Applied Biotechnology*, 10, 32–39. <https://doi.org/10.12702/1984-7033.v10n01a05>
- Glaubit, J. C., Casstevens, T. M., Lu, F., Harriman, J., Elshire, R. J., Sun, Qi, & Buckler, E. S. (2014). TASSEL-GBS: A high capacity genotyping by sequencing analysis pipeline. *PLOS ONE*, 9(2). <https://doi.org/10.1371/journal.pone.0090346>
- Gorjanc, G., Dumasy, J.-F., Gonen, S., Gaynor, R. C., Antolin, R., & Hickey, J. M. (2017). Potential of low-coverage genotyping-by-sequencing and imputation for cost-effective genomic selection in biparental segregating populations. *Crop Science*, 57, 1404–1420. <https://doi.org/10.2135/cropsci2016.08.0675>
- Hallauer, A. R., Carena, M. J., & Miranda Filho, J. B. (2010). *Quantitative genetics in maize breeding*. New York: Springer.
- Heffner, E. L., Lorenz, A. J., Jannink, J. L., & Sorrells, M. E. (2010). Plant breeding with genomic selection: Gain per unit time and cost. *Crop Science*, 50, 1681–1690. <https://doi.org/10.2135/cropsci2009.11.0662>
- Henderson, C. R. (1950). Estimation of genetic parameters. *Annals of Mathematical Statistics*, 21, 309–310.
- Hickey, J. M., Chiurugwi, T., Mackay, I., & Powell, W. (2017). Genomic prediction unifies animal and plant breeding programs to form platforms for biological discovery. *Nature Genetics*, 49, 1297–1303. <https://doi.org/10.1038/ng.3920>
- Hill, W. G., Goddard, M. E., & Visscher, P. M. (2008). Data and theory point to mainly additive genetic variance for complex traits. *PLOS Genetics*, 4(2). <https://doi.org/10.1371/journal.pgen.1000008>
- Jarquín, D., Crossa, J., Lacaze, X., Du Cheyron, P., Daucourt, J., Lorgeou, J., ... de los Campos, G. (2014). A reaction norm model for genomic selection using high-dimensional genomic and environmental data. *Theoretical and Applied Genetics*, 127, 595–607. <https://doi.org/10.1007/s00122-013-2243-1>
- Jonas, E., & de Koning, D. J. (2016). Goals and hurdles for a successful implementation of genomic selection in breeding programme for selected annual and perennial crops. *Biotechnology & Genetic Engineering Reviews*, 32, 18–42. <https://doi.org/10.1080/02648725.2016.1177377>
- Kelly, A. M., Smith, A. B., Eccleston, J. A., & Cullis, B. R. (2007). The accuracy of varietal selection using factor analytic models for multi-environment plant breeding trials. *Crop Science*, 47, 1063–1070. <https://doi.org/10.2135/cropsci2006.08.0540>
- Krchov, L.-M., & Bernardo, R. (2015). Relative efficiency of genomewide selection for testcross performance of doubled haploid lines in a maize breeding program. *Crop Science*, 55, 2091–2099. <https://doi.org/10.2135/cropsci2015.01.0064>
- Lado, B., Barrios, P. G., Quincke, M., Silva, P., & Gutiérrez, L. (2016). Modeling genotype $\times$ environment interaction for genomic selection with unbalanced data from a wheat breeding program. *Crop Science*, 56, 2165–2179. <https://doi.org/10.2135/cropsci2015.04.0207>
- Lee, S. H., Goddard, M. E., Visscher, P. M., & Van Der Werf, J. H. J. (2010). Using the realized relationship matrix to disentangle confounding factors for the estimation of genetic variance components of complex traits. *Genetics, Selection, Evolution*, 42. <https://doi.org/10.1186/1297-9686-42-22>
- Lopez-Cruz, M., Crossa, J., Bonnett, D., Dreisigacker, S., Poland, J., Jannink, J.-L., ... de los Campos, G. (2015). Increased prediction accuracy in wheat breeding trials using a marker $\times$ environment interaction genomic selection model. *G3: Genes, Genomes, Genetics*, 5, 569–582. <https://doi.org/10.1534/g3.114.016097>
- Lynch, M., & Walsh, B. (1998). *Genetics and analysis of quantitative traits*. Sunderland, MA: Sinauer Associates.
- Malosetti, M., Bustos-Korts, D., Boer, M. P., & Van Eeuwijk, F. A. (2016). Predicting responses in multiple environments: Issues in relation to genotype $\times$ environment interactions. *Crop Science*, 56, 2210–2222. <https://doi.org/10.2135/cropsci2015.05.0311>
- Meuwissen, T. H. E., Hayes, B. J., & Goddard, M. E. (2001). Prediction of total genetic value using genome-wide dense marker maps. *Genetics*, 157, 1819–1829. <https://doi.org/10.1129/0733>
- Moghaddar, N., & van der Werf, J. H. J. (2017). Genomic estimation of additive and dominance effects and impact of accounting for dominance on accuracy of genomic evaluation in sheep populations. *Journal of Animal Breeding and Genetics*, 134, 453–462. <https://doi.org/10.1111/jbg.12287>
- Muñoz, P. R., Resende, M. F. R., Gezan, S. A., Resende, M. D. V., de los Campos, G., Kirst, M., ... Peter, G. F. (2014). Unraveling additive from nonadditive effects using genomic relationship matrices. *Genetics*, 198, 1759–1768. <https://doi.org/10.1534/genetics.114.171322>
- Nazarian, A., & Gezan, S. A. (2016). Integrating nonadditive genomic relationship matrices into the study of genetic architecture of complex traits. *Journal of Heredity*, 107, 153–162. <https://doi.org/10.1093/jhered/esv096>
- Oakey, H., Cullis, B., Thompson, R., Comadran, J., Halpin, C., & Waugh, R. (2016). Genomic selection in multi-environment crop trials. *G3: Genes, Genomes, Genetics*, 6, 1313–1326. <https://doi.org/10.1534/g3.116.027524>
- Pastina, M. M., Malosetti, M., Gazaffi, R., Mollinari, M., Margarido, G. R. A., Oliveira, K. M., & Garcia, A. A. F. (2012). A mixed model QTL analysis for sugarcane multiple-harvest-location trial data. *Theoretical and Applied Genetics*, 124, 835–849. <https://doi.org/10.1007/s00122-011-1748-8>
- Patterson, H. D., & Thompson, R. (1971). Recovery of inter-block information when block sizes are unequal. *Biometrika*, 58, 545–554. <https://doi.org/10.1093/biomet/58.3.545>
- Piepho, H.-P. (1997). Analyzing genotype-environment data by mixed models with multiplicative terms. *Biometrics*, 53, 761–766. <https://doi.org/10.2307/2533976>
- Piepho, H. P. (1998). Empirical best linear unbiased prediction in cultivar trials using factor-analytic variance-covariance structures. *Theoretical and Applied Genetics*, 97, 195–201. <https://doi.org/10.1007/s001220050885>
- Piepho, H. P., Möhring, J., Melchinger, A. E., & Büchse, A. (2008). BLUP for phenotypic selection in plant breeding and variety testing. *Euphytica*, 161, 209–228. <https://doi.org/10.1007/s10681-007-9449-8>
- Powell, J. E., Visscher, P. M., & Goddard, M. E. (2010). Reconciling the analysis of IBD and IBS in complex trait studies. *Nature Reviews Genetics*, 11, 800–805. <https://doi.org/10.1038/nrg2865>

- R Core Team. (2017). *R: A language and environment for statistical computing*. Vienna: R Project for Statistical Computing.
- Rincint, R., Kuhn, E., Monod, H., Oury, F. X., Rousset, M., Allard, V., & Le Gouis, J. (2017). Optimization of multi-environment trials for genomic selection based on crop models. *Theoretical and Applied Genetics*, 130, 1735–1752. <https://doi.org/10.1007/s00122-017-2922-4>
- Rutkoski, J., Singh, R. P., Huerta-Espino, J., Bhavani, S., Poland, J., Jannink, J. L., & Sorrells, M. E. (2015). Efficient use of historical data for genomic selection: A case study of stem rust resistance in wheat. *Plant Genome*, 8(1). <https://doi.org/10.3835/plantgenome2014.09.0046>
- Schnable, P. S., Ware, D., Fulton, R. S., Stein, J. C., Wei, F., Pasternak, S., ... Wilson, R. K. (2009). The B73 maize genome: Complexity, diversity, and dynamics. *Science*, 326(5956), 1112–1115. <https://doi.org/10.1126/science.1178534>
- Schulz-Streeck, T., Ogutu, J. O., Gordillo, A., Karaman, Z., Knaak, C., & Piepho, H.-P. (2013). Genomic selection allowing for marker-by-environment interaction. *Plant Breed*, 132(6), 532–538. <https://doi.org/10.1111/pbr.12105>
- Schwarz, G. (1978). Estimating the dimension of a model. *Annals of Statistics*, 6, 461–464. <https://doi.org/10.1214/aos/1176344136>
- Smith, A., Cullis, B. R., & Thompson, R. (2001). Analyzing variety by environment data using multiplicative mixed models and adjustments for spatial field trend. *Biometrics*, 57, 1138–1147. <https://doi.org/10.1111/j.0006-341X.2001.01138.x>
- Smith, A. B., Ganesalingam, A., Kuchel, H., & Cullis, B. R. (2015). Factor analytic mixed models for the provision of grower information from national crop variety testing programs. *Theoretical and Applied Genetics*, 128, 55–72. <https://doi.org/10.1007/s00122-014-2412-x>
- Sorensen, D., & Gianola, D. (2002). *Likelihood, Bayesian, and MCMC methods in quantitative genetics*. New York: Springer.
- Tech Services. (2018). *Pricing brochure TSI 2018 test sites*. Bluffton, IN: Tech Services. Retrieved from <http://techservicespro.com/test-locations/>
- Van Eeuwijk, F. A., Bustos-Korts, D. V., & Malosetti, M. (2016). What should students in plant breeding know about the statistical aspects of genotype $\times$ environment interactions? *Crop Science*, 56, 2119–2140. <https://doi.org/10.2135/cropsci2015.06.0375>
- VanRaden, P. M. (2008). Efficient methods to compute genomic predictions. *Journal of Dairy Science*, 91, 4414–4423. <https://doi.org/10.3168/jds.2007-0980>
- Wright, S. (1921). Systems of mating. I. The biometric relations between parent and offspring. *Genetics*, 6, 111–126.
- Yan, W. (2016). Analysis and handling of G $\times$ E in a practical breeding program. *Crop Science*, 56, 2106–2118. <https://doi.org/10.2135/cropsci2015.06.0336>
- Zhang, X., Pérez-Rodríguez, P., Semagn, K., Beyene, Y., Babu, R., López-Cruz, M. A., ... Crossa, J. (2015). Genomic prediction in biparental tropical maize populations in water-stressed and well-watered environments using low-density and GBS SNPs. *Heredity*, 114, 291–299. <https://doi.org/10.1038/hdy.2014.99>
- Zhou, Y., Vales, M. I., Wang, A., & Zhang, Z. (2017). Systematic bias of correlation coefficient may explain negative accuracy of genomic prediction. *Briefings in Bioinformatics*, 18, 744–753. <https://doi.org/10.1093/bib/bbw064>

SUPPORTING INFORMATION

Additional supporting information may be found online in the Supporting Information section at the end of the article.

How to cite this article: Krause MD, Dias KOdG, dos Santos JPR, et al. Boosting predictive ability of tropical maize hybrids via genotype-by-environment interaction under multivariate GBLUP models. *Crop Science*. 2020;60:3049–3065. <https://doi.org/10.1002/csc2.20253>