

Evaluation of RR-BLUP Genomic Selection Models that Incorporate Peak Genome-Wide Association Study Signals in Maize and Sorghum

Brian Rice and Alexander E. Lipka*

Dep. of Crop Sciences, 1102 S Goodwin Ave, Univ. of Illinois Urbana Champaign, Urbana, IL 61801.

ABSTRACT Certain agronomic crop traits are complex and thus governed by many small-effect loci. Statistical models typically used in a genome-wide association study (GWAS) and genomic selection (GS) quantify these signals by assessing genomic marker contributions in linkage disequilibrium (LD) with these loci to trait variation. These models have been used in separate quantitative genetics contexts until recently, when, in published studies, the predictive ability of GS models that include peak associated markers from a GWAS as fixed-effect covariates was assessed. Previous work suggests that such models could be useful for predicting traits controlled by several large-effect and many small-effect genes. We expand this work by evaluating simulated traits from diversity panels in maize (*Zea mays* L.) and sorghum [*Sorghum bicolor* (L.) Moench] using ridge-regression best linear unbiased prediction (RR-BLUP) models that include fixed-effect covariates tagging peak GWAS signals. The ability of such covariates to increase GS prediction accuracy in the RR-BLUP model under a wide variety of genetic architectures and genomic backgrounds was quantified. Of the 216 genetic architectures that we simulated, we identified 60 where the addition of fixed-effect covariates boosted prediction accuracy. However, for the majority of the simulated data, no increase or a decrease in prediction accuracy was observed. We also noted several instances where the inclusion of fixed-effect covariates increased both the variability of prediction accuracies and the bias of the genomic estimated breeding values. We therefore recommend that the performance of such a GS model be explored on a trait-by-trait basis prior to its implementation into a breeding program.

Abbreviations: GBS, genotyping-by-sequencing; GEBV, genomic estimated breeding values; GS, genomic selection; GWAS, genome-wide association study; LD, linkage disequilibrium; MAS, marker-assisted selection; MLM, mixed linear model; QTL, quantitative trait loci; QTN, quantitative trait nucleotide; RR-BLUP, ridge-regression best linear unbiased prediction; RR, ridge regression; SNP, single nucleotide polymorphism.

core ideas

- Augmenting RR-BLUP models with peak GWAS markers can hypothetically boost prediction accuracy
- We conducted a simulation study in maize and sorghum to test the performance of such models
- For most of the simulated traits, we observed a decrease in prediction accuracy
- These augmented models tended to yield greater variability in prediction accuracy
- An increase of bias in predicted breeding values from these models was noted

There are two prominent strategies for prediction of agronomically important crop traits from genotypes: marker-assisted selection (MAS) and GS. Both rely on the statistical analysis of genetic markers to quantify the contribution of each marker to phenotypic variability (Bernardo, 2010). Typically, MAS predicts trait values using only a small number of markers linked to large-effect quantitative trait loci (QTL) (Collard and Mackill, 2008), while GS uses all available markers across the genome to generate the predicted breeding value (Meuwissen et al., 2001). Phenotypic traits controlled by a few genes of large effect, referred to as Mendelian traits, are ideal candidates for MAS, where molecular markers

Citation: Rice, B., and A.E. Lipka. 2019. Evaluation of RR-BLUP genomic selection models that incorporate peak genome-wide association study signals in maize and sorghum. *Plant Genome* 12:180052. doi: 10.3835/plantgenome2018.07.0052

Received 25 July 2018. Accepted 5 Dec. 2018. *Corresponding author (alipka@illinois.edu).

This is an open access article distributed under the CC BY-NC-ND license (<http://creativecommons.org/licenses/by-nc-nd/4.0/>). Copyright © Crop Science Society of America 5585 Guilford Rd., Madison, WI 53711 USA

linked to such genes can be used to supplement phenotypic selection (Lynch and Walsh, 1998). A method of obtaining linked markers, known as association mapping, uses historical recombination to quantify statistical associations between a trait of interest and genetic markers (Lipka et al., 2015). The rapidly decreasing genotyping cost and increase in available genetic data have resulted in the widespread use of the GWAS (Guo et al., 2018; Jardim et al., 2018). Given that a Google scholar search (conducted on 18 Oct. 2018) for papers published in 2018 containing the key word GWAS yielded 11,000 results, it is apparent that association studies continue to be widely used in many research endeavors.

While MAS is useful for phenotypic prediction of Mendelian traits, many agronomic traits of interest are complex, meaning that they are governed by many genetic components of various effect sizes (often small) (Barton et al., 2017). When the underlying genomic sources contributing to a given trait consists of up to thousands of small-effect genes, then selection based on one or a few genetic markers will theoretically be ineffective (Xu and Crouch, 2008). To observe substantial selection gains, a more complex method than MAS is required. One such approach is GS, which is based on the infinitesimal model conferring that a trait value is a result of the linear combination of additive genetic and nongenetic sources (Fisher, 1919). First suggested by Meuwissen et al. (2001), GS takes into account the effect of all available genetic markers for prediction of breeding values instead of only those passing a significance threshold, which, according to the infinitesimal model, approximates the genomic underpinnings of a complex trait. Conducting GS therefore requires estimation of the effect of each marker, treated predominantly in practice as additive, although methods for including dominance (Technow et al., 2012), epistatic (Jiang and Reif, 2015), and genotype-by-environment (Cuevas et al., 2016) effects are becoming available to the research community. Because of the high dimensionality of genetic data, fitting GS models requires consolidation with the fact that the number of markers (p) available in a typical study exceeds the number of individuals (n) (de los Campos et al., 2013). Consequently, when a GS model that considers the additive effects of each of these markers is fitted to such large p -small n data, there will be an infinite number of maximum likelihood estimates of these effects (Gianola, 2013). One of the most common approaches to overcome this issue is to use the RR-BLUP GS model (Meuwissen et al., 2001), which incorporates all marker information to predict a line's genomic estimated breeding values (GEBV) while simultaneously implementing a penalization function to restrict the values that each marker's predicted additive contributions can equal. Although other approaches including Bayesian methods (reviewed in Gianola, 2013) and the least absolute shrinkage and selection operator (Tibshirani, 1996) are also widely used in GS to address the large p -small n issue, it has been shown that RR-BLUP is

of equal or superior prediction accuracies and requires lower computational time (Heslot et al., 2012; Resende et al., 2012; Riedelsheimer et al., 2012).

Both MAS and GS have historically been used separately with the optimal approach depending on the genetic architecture of the trait and number of markers available (Spindel et al., 2015). While some traits are simply inherited (controlled by few large-effect components) and others complex (many small-effect components), the reality is often that a mixture of large- and small-effect genomic components contribute to the phenotype (Mackay, 2001). Thus when penalized approaches such as RR-BLUP are used in GS models, the penalty is applied equally to all markers tagging both small- and large-effect genomic components. Therefore, it becomes possible for the contributions of the large-effect components to not be completely accounted for in the GS model, potentially resulting in lower prediction accuracies (Bernardo, 2014). When this is taken into consideration the question becomes are ridge regression and similar penalties grossly underestimating the contribution of large effect QTL to the overall phenotype? Bernardo (2014) showed in a simulation study that when major genes are known, including them in the model as fixed-effect covariates can increase prediction accuracy especially when they explain a substantial amount of phenotypic variance. The author suggested that when a gene explains >10% of genetic variance, it should be included as a fixed-effect covariate in RR-BLUP.

The practicality of this and any GS approach depends on knowledge of the genetic architecture of the traits of interest, including heritability, the number of underlying causative mutations, and their effect sizes (Huang and Mackay, 2016). Given the wide variety of complexity of genetic architectures reported in recent studies (Campbell et al., 2017; Divilov et al., 2018; Muqaddasi et al., 2017; da Silva Romero et al., 2018), an RR-BLUP model augmented with unpenalized fixed-effect marker covariates could theoretically accelerate the breeding cycles of many crop and livestock species.

When the approach described in Bernardo (2014) is applied to real data the exact location and sizes of large-effect genes are often unknown. In the absence of such information, GWAS results (in particular markers exhibiting peak associations with a trait of interest) can instead be used as fixed-effect covariates in a GS model. One of the first studies to explore such an approach was Zhang et al. (2014), where it was demonstrated that incorporating fixed-effect covariates identified as peak GWAS signals (available in public databases) into a GS model outperformed BayesB and genomic best linear unbiased prediction for nine of 11 traits in a rice

(*Oryza sativa* L.) diversity panel as well as two of three traits in cattle (*Bos taurus*). However, this improvement was marginal, with 0.1 to 1% higher prediction accuracies over competing models. One possible explanation of these findings is that the peak-associated GWAS markers used in this study were from public data. This could be an issue, since GWAS results can be population

specific because of differences in LD (Caldwell et al., 2006). Building off these previous studies, Spindel et al. (2016) suggested a method where GWAS is conducted on a training set and markers passing a threshold are set as fixed effects in the RR-BLUP models. This approach, named GS + de novo GWAS, outperformed six alternate GS and MAS approaches when used to analyze four traits in rice. Excitingly, GS + de novo GWAS yielded at least an ~10% increase in prediction accuracy over the other tested approaches for two of these traits. Arruda et al. (2016) conducted a similar procedure in wheat (*Triticum aestivum* L.) for six traits related to Fusarium head blight, where the standard RR-BLUP model was augmented with fixed-effect QTL. When a single fixed-effect covariate corresponding to major QTL *Fhb-1* prediction was included in the GS model, an increase in prediction accuracy of 3 to 14% was observed over an RR-BLUP model with no fixed effects. Similar results were seen when other independently published QTL were used as fixed-effect covariates. Finally, Raymond et al. (2018) investigated the incorporation meta-GWAS

results into a GS model that is equivalent to RR-BLUP. Although higher prediction accuracies for bull stature were observed, such gains were also accompanied by increases in the bias in the GEBV. To our knowledge, the previous five studies are the extent to which incorporating such covariates into RR-BLUP or similar models has been conducted. None indicated any significant penalty to incorporating fixed-effect covariates in GS (beyond a potential increase in bias of GEBV) and suggest that this approach should be tested in other species and traits that cover a variety of genomic and trait architectures.

Although these previous studies have indicated the potential of including fixed-effect marker covariates in GS models, none have done so using RR-BLUP in a maize or sorghum diversity panel. Evaluation of an RR-BLUP model that incorporates fixed-effect markers

in diversity panels from these two species is needed to

quantify its ability to predict diverse lines and recommend

its use as a tool for introgression of new genetic variation into breeding populations. While this model is hypothesized to increase prediction accuracy over

RR-BLUP with no fixed effects (referred to hereon as RR-BLUP) for complex traits with a few large-effect genomic components, it still requires evaluation across traits with genetic architectures ranging from simple to complex for a complete analysis. Therefore, the purpose of this work was to explore the performance of GS + de novo GWAS over a wide variety of genetic architectures underlying simulated traits in maize and sorghum diversity panels. Trait architecture in this study is determined by the following three aspects: (i) narrow sense heritability (h^2),

(ii) number of quantitative trait nucleotides (QTNs),

Materials And Methods

Simulation of Phenotypic Data

Phenotypes were simulated using publicly available single nucleotide polymorphism (SNP) data for 281 inbred lines in maize (Flint-Garcia et al., 2005) and 320 inbred lines in sorghum (Morris et al., 2013). Maize genotypes were collected using the Illumina MaizeSNP50 Bead-Chip resulting in 51,742 SNPs as described in Cook et al. (2012) (available at panzea.org/genotypes). Sorghum genotypes were collected using genotyping-by-sequencing (GBS) techniques (Elshire et al., 2011) as described in Bouchet et al. (2017) resulting in 90,441 SNPs (sorghum marker data available at datadryad.org//resource/doi:10.5061/dryad.gm073). Within each species, the respective marker data were used to simulate traits that represented a wide range of genetic architectures. The specific genomic contributions of each simulated trait varied accordingly by the narrow-sense heritability (h^2), the number of underlying QTNs, and additive effect sizes

of each QTN. The first step in the procedure for simulating these traits was to randomly select a set of markers to be QTNs. After phenotypic values were simulated, such markers were not considered for inclusion as fixed-effect covariates to reflect the reality that true causal mutations underlying phenotypic variation are often not genotyped (i.e., these markers were not considered as fixed effect regardless if they were identified as a peak GWAS signal). Next, the additive effect size of each QTN were assigned in two different configurations, which are described below. Genetic components of the phenotypic values

were thus determined by the function $\hat{a}_{ij} = \sum_{j=1}^s x_{ij} Q_j$ where

s is the number of simulated QTN, x_{ij} is the genotypic state of the j th QTN at the i th individual (coded numerically as -1, 0, 1), and Q_j is the assigned additive effect.

Lastly, environmental effects were randomly drawn from

a normal distribution with a mean $m = 0$ and variance $s_e^2 = s_a^2 h^2 - s_a^2$, where s_a^2 is the additive effect

variance (calculated as the variance of $\hat{a}_{ij} = \sum_{j=1}^s x_{ij} Q_j$) and h^2 is

and

(iii) additive effect size of each QTN. We hypothesized that the success of this model, defined as increase in mean prediction accuracies for 50 replications of five-fold cross-validation compared with RR-BLUP, would be dependent on the genetic architecture of a given trait.

the narrow-sense heritability. For each distinct genetic architecture simulated in each species, a total of 50 phenotypic replications were simulated. A summary of the spectrum of the simulated genetic architectures can be found in Table 1.

Type 1 and 2 Traits

To accommodate a wide range of biologically relevant genetic architectures, our simulated traits were subdivided into two categories that we named Type 1 and Type 2 traits. The former represents a biological trait, where s number of mutations affecting the trait have occurred. As time from when the mutation occurs increases, the effect size of the mutation diminishes, driving the population phenotypic standard deviation to zero (Fisher, 1930; Orr, 1998). To model this for Type 1 traits, the largest QTN effect size (Q) is first selected from the boundaries of zero to one (not including zero and one). The remaining simulated QTN effect sizes were

Table 1. The parameter settings for the number (*s*) of simulated quantitative trait nucleotides (QTNs), the additive effect size of the largest QTN (*Q*), and narrow-sense heritability (*h*²) that were considered in the simulation studies. (Sorghum and maize traits were simulated with the same parameters)

No. of QTNs (<i>s</i>)	Additive effect size of largest QTN (<i>Q</i>)	Heritability (<i>h</i> ²)
1†, 2‡, 3, 5, 10, 25, 100	0.1†, 0.3‡, 0.5, 0.9	0.1, 0.5, 0.9

† Parameter only present in Type 1 traits.

‡ Parameter only present in Type 2 traits.

then assigned in a geometric series where the effect size of the *i*th QTN was *Qⁱ*. For example, if three QTN (i.e., *s* = 3) were selected and *Q* = 0.9, then the effect sizes are 0.9, 0.9², and 0.9³. Using this model, a total of 54 Type 1 traits in maize and 54 in sorghum were simulated.

The second category of traits, called Type 2 traits, is one in which a relatively new mutation is present and thus its effect size is large relative to the others. An example of this is the maize *Sos1* mutant, which is a major effect dominant mutation controlling inflorescence; evidence suggests that this mutation arose after the domestication of maize from teosinte [*Zea mays* L. subsp. *mexicana* (Schrader) H. H. Iltis] (Doebley et al., 1995). To simulate Type 2 traits, one QTN with a large additive effect size *Q* is selected, while the remaining simulated effect sizes follow the same geometric series previously described for the Type 1 traits, starting with an effect size of 0.1. Thus if *s* = 3 and *Q* = 0.9, then the effect sizes would be 0.9, 0.1, and 0.1², respectively. Collectively, a total of 54 Type 2 traits were simulated in both maize and sorghum. It has been hypothesized that RR-BLUP models with fixed-effect covariates will outperform the standard RR-BLUP model for traits with genetic architectures similar to these Type 2 traits (Bernardo, 2014). Together, Type 1 and 2 simulated traits provide a wider range of representative phenotypes to evaluate GS + de novo GWAS than either could provide alone.

Genomic Selection

The RR-BLUP model (Meuwissen et al., 2001) was used to conduct GS as a baseline model for comparison to a model with fixed-effect covariates included. The RR-BLUP model (Model 1) is described as follows:

$$y_i = m + \sum_{k=1}^p x_{ik} b_k + e_i \tag{1}$$

where *y_i* is the observed phenotypic value of the *i*th individual, *m* is the grand mean, *x_{ik}* is the genotype at the *k*th marker of the *i*th individual, *p* is the total number of markers, *b_k* is the estimated random additive effect of the *k*th marker ~N(0, *s*²), and *e* is the error term ~N(0, *s*²). The BLUP of each *b_k* following ridge regression (RR) penalty (Hoerl and Kennard, 1970):

$$J(b) = \sum_{k=1}^p b_k^2 \tag{2}$$

where all terms are the same as those described for Eq. [1]. Intuitively, this penalty restricts the

values the BLUP of each *b_k* can take on. The model was implemented for analysis in R using the package rrBLUP (Endelman, 2011).

GWAS and Criteria for Markers Treated as Fixed Effects

The approach to conduct GWAS on the simulated data has been previously described (Lipka et al., 2013). Briefly, the unified mixed linear model (MLM; Yu et al., 2006) was fitted at each marker for each simulated trait. In both species, this model included the first three principal components from a principal component analysis of the markers to account for spurious associations arising from population structure; the scree plots used to determine that three principal components sufficiently account for subpopulation structure in both the maize and sorghum diversity panels are presented in Supplemental Fig. S1 and S2. In addition, kinship (i.e., additive genetic relatedness) matrices obtained from the method of Loiselle et al. (1995) were used to account for spurious associations arising from familial relatedness. All analyses were conducted using the genome association and prediction integrated tool R package (Lipka et al., 2012). The Benjamini and Hochberg (1995) procedure was used to control the false discovery rate at 5%. Because the purpose of the GWAS conducted in this research was to identify markers that could accurately predict the values of a given simulated trait, all of the evaluated markers were ordered by the degree of statistical association with the trait (i.e., from smallest to largest *P*-value). The top *m* associated markers (Table 2) from each analyses with the strongest associations with the traits were then carried on to the next phase of the analysis.

Genomic Selection plus De Novo Genome-Wide Association Study

After the GWAS was conducted on a given trait, the top *m* associated markers (Table 2) were included as fixed effects in the following model (Model 2):

Table 2. Number of peak-associated markers (*m*) included as fixed-effect covariates in the ridge-regression best linear unbiased prediction (RR-BLUP) model, which depended on the number (*s*) of simulated underlying quantitative trait nucleotides (QTNs). Sorghum and maize traits were evaluated with the same values of *m*.

No. of QTN (<i>s</i>)	No. of fixed effects evaluated (<i>m</i>)
1, 2, 3	1, 2, 3, 5, 10, 25
5	1, 2, 3, 5, 10, 25, 50
possible	

10	1, 2, 3, 5, 10, 25, 50, 100
<u>25, 100</u>	<u>1, 2, 3, 5, 10, 25, 100</u>

$$y_i = m + \sum_{j=1}^m \mathbf{a}_{ij} x_{ij} + \sum_{k=1}^p \mathbf{b}_k x_{ik} + e_i \quad [3]$$

where x_{ij} is the genotype at the j th marker of the i th individual, m is the number of top associated markers considered for inclusion as fixed-effect covariates, $\mathbf{a}_j$ is the fixed additive effect of the j th marker, x_{ik} is the genotype at the k th marker of the i th individual, p is the total number of markers, $\mathbf{b}_k$ is the estimated random additive marker effect of the k th marker $\sim N(0, s^2)$, and e is the residual error term $\sim N(0, s^2)$. Because the RR penalty is not used for the estimation of these fixed additive effects, no restrictions are placed on the numerical value of these estimates. Thus, peak markers tagging sufficiently large-effect QTN that are incorporated into Model 2 could hypothetically boost trait prediction accuracies over those from the standard RR-BLUP model.

Cross-Validation Scheme

A five-fold cross-validation scheme (described in Owens et al., 2014) was implemented to assess the prediction accuracies of the tested statistical models. Within each species, this procedure randomly subdivided the individuals into five subsets (i.e., folds), each with approximately the same number of individuals. All of the GS and GWAS models previously described were fitted in the four of five folds (i.e., the training set), and then the predictive ability of Models 1 and 2 was evaluated in the fifth fold (the validation set). This scheme was repeated five times so the GEBVs of the individuals within each fold was predicted once using each statistical model. Accuracy was reported as the Pearson correlation coefficient r between the simulated phenotypes in the validation set and the GEBVs predicted from models fitted in the corresponding training set. To enable a direct comparison of prediction accuracy across all models and genetic architectures, the same folds were used for all analyses conducted within each species.

To adequately compare prediction accuracies between Models 1 and 2, 50 replications of simulated phenotypic values for each of the 108 genetic architectures considered in each species were analyzed. For this study the values considered for m (i.e., the number of peak-associated SNPs to be included as fixed-effect covariates in Model 2) increased depending on the number of simulated QTN described in Table 1. For every replication and value of m , the mean and standard deviation of the prediction accuracy across the five folds was calculated. Subsequently, the mean and standard deviation of all replications for each model and level of m were calculated. Dunnett's mean comparison test (Dudewicz et al., 1975; Dunnett, 1955) was conducted to determine if any particular number of fixed covariates m yielded significantly higher prediction accuracies than the standard RR-BLUP Model 1 at an experimentwise type I error rate of $\alpha = 0.05$. All analysis conducted were identical for maize and sorghum.

Data Analysis

Simulations and all data analysis was conducted using the R software package (R Development Core Team, 2018). All materials used in these analyses are publicly available (<https://github.com/ricebrian/RR-BLUP-with-fixed-effects>).

results

Genomic Selection plus De Novo Genome-Wide Association Studies Tended to Yield Lower Prediction Accuracies

A total of 50 replications of phenotypic data for 108 genetic architectures were simulated using 281 maize inbred genotypes and again in 320 sorghum inbred genotypes. Using these simulated traits, the predictive ability of GS + de novo GWAS (Model 2) using various numbers of fixed-effect covariates was compared with that of Model 1 (i.e., a standard RR-BLUP model). For each replicate phenotype and number of fixed-effect covariates, the Pearson correlation coefficient r between predicted GEBVs and simulated trait values in a five-fold cross-validation scheme was calculated, and the mean value of r across the five folds was subsequently used to quantify the prediction accuracy.

We observed that the inclusion of fixed-effect covariates in an RR-BLUP model had a tendency to decrease instead of increase prediction accuracies (Fig. 1). Of the 216 different genetic architectures that were explored in this simulation study (i.e., 108 in maize and 108 in sorghum), the inclusion of at least $m = 1$ fixed-effect covariates increased prediction accuracy relative to the standard RR-BLUP Model 1 for only 60 of these genetic architectures (summarized in Table 3). Even in these instances where prediction accuracy did increase, the variability of mean prediction accuracies (across 50 phenotypic replications) also tended to drastically increase compared with Model 1, as illustrated in Fig. 2 through 5 and the online supplemental material (<https://github.com/ricebrian/RR-BLUP-with-fixed-effects>). Our results also suggest that the inclusion of at least one fixed-effect covariate in an RR-BLUP model can increase the bias of the predicted GEBVs (Table 4). The narrow-sense heritability (h^2), number of underlying QTN (s), additive effect size of the largest QTN (Q), and prediction accuracy for the tested number of fixed-effect covariates m are reported in Supplemental Table S1.

Instances where Prediction Accuracy Increased

Interestingly, 39 of the 60 genetic architectures where the prediction accuracy increased were in sorghum. Furthermore, 35 of these 60 genetic architectures were Type 2, meaning that these traits were simulated with one large-effect QTN in addition to a series of small-effect QTN. Finally, the greatest increases in prediction accuracy of Model 2 over Model 1 was observed in Type 2 traits simulated in sorghum with $s = 100$ QTNs, $h^2 = 0.9$, and up to $m = 3$ fixed-effect covariates

included in Model 2 (Fig. 1).

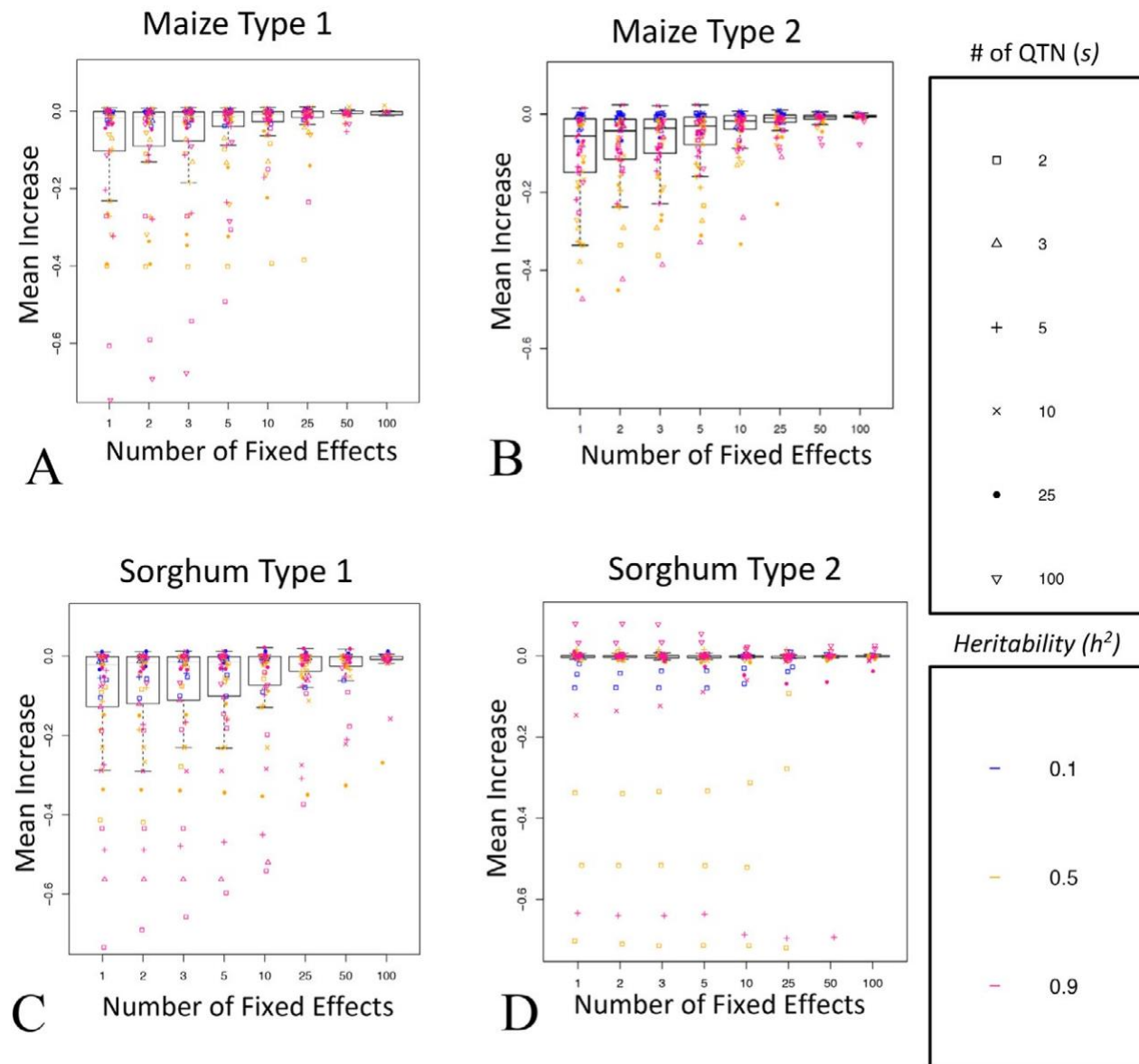


Fig. 1. Mean increase in prediction accuracy (y-axis) after including peak markers from a genome-wide association study as a fixed-effect covariate in the ridge-regression best linear unbiased prediction model (x-axis). For each type of simulated trait and species, this mean increase in prediction accuracy across all 50 replications is plotted as box and whisker plots at varying levels of m , the number of fixed effects. (A) Maize Type 1; (B) Maize Type 2; (C) Sorghum Type 1; (D) Sorghum Type 2. Each point in the mean change from model (1) where color denotes the trait's narrow sense heritability (h^2) and symbol indicates the number of simulated quantitative trait nucleotides (s).

Maize Type 1

An increase in prediction accuracy after including at least $m = 1$ fixed-effect covariate in Model 2 was observed for 15 of the 54 Type 1 traits simulated in maize (Supplemental Table S1). Seven of these traits had low heritability ($h^2 = 0.1$) and all 15 had moderate- to large-effect QTN ($Q = 0.5$ or 0.9). The total number of QTN in each of these 15 traits were both small and large, suggesting that the number of QTN underlying these traits did not have a substantial impact on the predictive ability of Model 2. Finally, a statistically significant increase in prediction of accuracy of Model 2 over Model 1 was observed in six of these 15 traits. The results for the maize Type 1 traits are summarized in Supplemental Table S1 and Fig. 2.

Maize Type 2

A total of six of the 54 Type 2 traits simulated in maize yielded higher prediction accuracies when at least one fixed-effect covariate was included in Model 2 (Supplemental Table S1). The genetic architectures of these six traits were similar to the 15 previously mentioned Type 1 maize traits presented in Supplemental Table S1. That is, all but one of these six Type 2 traits had a low heritability of $h^2 = 0.1$, and the effect size of the large-effect QTN was either $Q = 0.5$ or $Q = 0.9$. Although the total number of QTN simulated in these six traits ranged from 2 to 100, the two Type 2 maize traits where statistically significant increases in prediction accuracy was observed for Model 2 both had $s = 10$ QTN (Supplemental Table S1; Fig. 3C,D).

Table 3. Summary of the simulated genetic architectures where the inclusion of at least one peak-associated marker from a genome-wide association study (conducted in a training set) as a fixed-effect covariate in a ridge-regression best linear unbiased prediction (RR-BLUP) model resulted in higher prediction accuracies than a standard RR-BLUP model. The count and percentage for each parameter setting, held constant at all other parameters, are presented.

Parameter	Level	Count successful	Percentage successful
Species†	Maize	21	19.4
	Sorghum	39	36.1
Trait†	Type 1	25	23.1
	Type 2	35	32.4
No. of simulated quantitative trait nucleotides (QTNs) (<i>s</i>)‡	1§	2	11.1
	2¶	4	22.2
	3	11	30.5
	5	6	16.6
	10	16	44.4
	25	10	27.8
	100	11	30.5
Additive effect of the largest QTN (<i>Q</i>)#	0.1§	6	8.3
	0.3¶	10	31.3
	0.5	21	29.1
	0.9	23	31.9
Narrow-sense heritability (<i>h</i> ²)††	0.1	21	29.2
	0.5	13	18.1
	0.9	26	36.1

† Percentages are out of 108.

‡ Percentages are out of 36 except for bold italicized values where percentages are out of 18.

§ Parameter only present in Type 1 traits

¶ Parameter only present in Type 2 traits

Percentages are out of 72 except for bold italicized values where percentages are out of 36.

†† Percentages are out of 72.

Sorghum Type 1

Similar to the Type 1 maize traits, the inclusion of at least $m = 1$ fixed-effect covariates in Model 2 resulted in higher prediction accuracies for 10 of the 54 Type 1 traits simulated in sorghum (Supplemental Table S1). Six of these traits had statistically significant increases for at least one setting of m . However, in contrast to maize, seven of these 10 Type 1 sorghum traits had high heritability (i.e., $h^2 = 0.9$). The number of QTN simulated among these 10 traits ranged $s = 3$ to 100, suggesting that like the simulated maize Type 1 traits, the number of underlying QTN in the Type 1 sorghum traits does not appear to contribute to the predictive ability of Model 2. Finally, the QTN effect sizes of these 10 traits ranged from small ($Q = 0.1$) to large ($Q = 0.9$), which indicates that the QTN effect sizes do not substantially impact the predictive ability of Model 2. The results of six Type 1 sorghum traits are presented in Fig. 4.

Sorghum Type 2

An increase in prediction accuracy with the inclusion of at least $m = 1$ fixed-effect covariate in Model 2 was observed in 29 of the 54 Type 2 traits

sorghum (Supplemental Table S1). The genetic architectures of these 29 traits varied greatly with respect to heritability, effect size of the largest QTN, and the number of simulated QTN. The trait with the greatest increase of prediction accuracy after including at least one fixed-effect covariate to Model 2 had high heritability ($h^2 = 0.9$), large number of QTN ($s = 100$), and a moderately large-effect QTN ($Q = 0.5$) (Fig. 5C). Of the 29 Type 2 sorghum traits where increased prediction accuracies were observed, statistically significant increases in prediction accuracy after including at least $m = 1$ fixed effect were observed in 17 traits (Supplemental Table S1). The traits with the smallest significantly significant increase in prediction accuracy are also presented in Fig. 5B.

discussion

Genomic prediction is a powerful tool for predicting phenotypic performance in maize and sorghum. Recent studies have suggested the inclusion of markers assumed to contribute greater to phenotypic variance as fixed-effect explanatory variables in an RR-BLUP model could increase prediction accuracy (Arruda et al., 2016; Bernardo, 2014; Spindel et al., 2016). To evaluate the potential of such a GS + de novo GWAS approach in maize and sorghum, 108 traits were simulated using marker data from maize and sorghum diversity panels (Table 1). Although we did observe several instances where including fixed-effect covariates into the RR-BLUP model did improve prediction accuracy, we more often noted that the inclusion of such fixed-effect covariates decreased prediction accuracy (Fig. 1; Supplemental Table S1). Moreover, we also observed increases in the variance of prediction accuracies (Fig. 2–5) and the bias of GEBVs (Table 4) after including these fixed-effect covariates.

Observed Advantages and Disadvantages of Including Fixed-Effect Covariates

The core aim of this simulation study was to adequately evaluate Model 2 to give recommendations on its usage. A grid search of genetic architectures using the Type 1 and Type 2 trait definitions were simulated using maize and sorghum genotypes. Type 2 traits were simulated to represent phenotypes that had a relatively larger QTN contributing to variance compared with the remaining QTN. Given the percentage variance explained of large-effect QTN of most of the simulated traits was >0.1 (Supplemental Table S2), the majority of the traits simulated in this work had genetic architectures that loosely adhered to recommended ideal genetic architectures mentioned Bernardo (2014) where including fixed-effect covariates has greatest potential to increase prediction accuracy. Consistent with the conclusions from Bernardo (2014), the most frequent observation of increased prediction accuracies using Model 2 occurred in Type 2 traits with one large-effect QTN (Table 3; Fig. 1). Thus, our results suggest that using a GS + de novo GWAS model would be best suited for certain traits that have

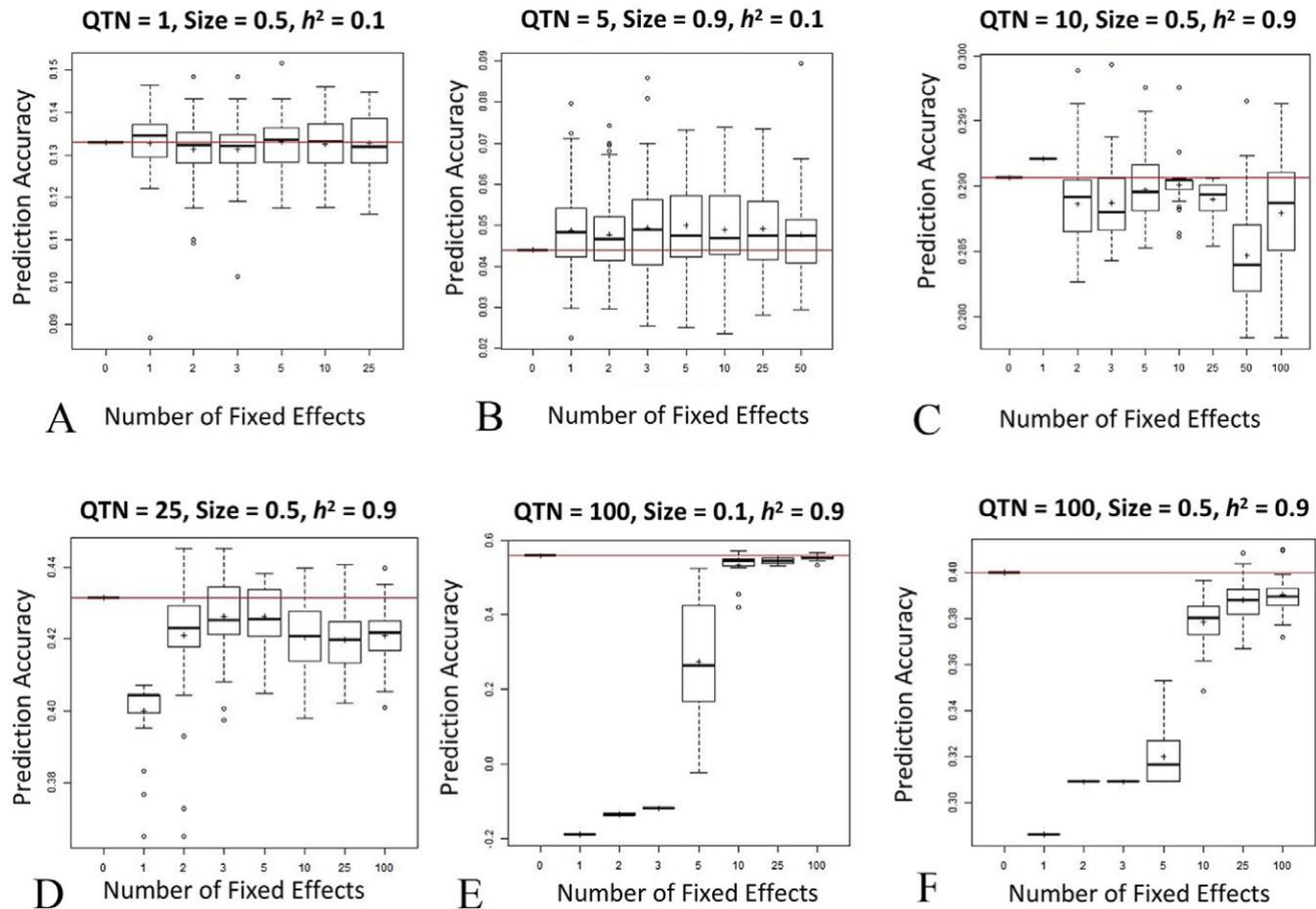


Fig. 2. Distribution of mean prediction accuracy for various genetic architectures in six maize Type 1 traits. The simulated genetic architecture is listed in the individual titles (A–F) where: QTN, number of simulated quantitative trait nucleotides; size, additive effect size of the largest QTN; h^2 , narrow-sense heritability. The x-axis is the number of fixed effects included in model (2). The y-axis is the mean Pearson correlation from the five-fold cross validation. In each graph, the mean five-fold cross-validation prediction accuracies of all 50 replications are used to generate these box plots. The orange line on each graph depicts the median prediction accuracy of the 50 replications when no markers are included in the model as fixed-effect covariates.

genetic architectures similar to the Type 2 traits that were simulated. One example of such a class of traits is disease resistance, which often have both large-effect resistant genes and a complex background of small-effect polymorphisms (Poland and Rutkoski, 2016).

The number of occurrences where low heritable traits had increased prediction accuracy could potentially be promising. Increasing prediction accuracy when heritability is low often requires models that include genotype-by-environment interactions to explain more variation (Brachi et al., 2011; Sukumaran et al., 2018). Modeling such interactions can be resource intensive because it requires multiple years and locations for planting. A GS + de novo GWAS approach offers an alternative or complement to this for germplasm screening of traits with low heritability where, for example, a 1% increase in accuracy could accelerate genetic gains. Given that the maximum increase in prediction accuracy with the inclusion of at least one fixed-effect covariate among the simulated traits with $h^2 = 0.1$ was 1.2%, the use of Model 2 in breeding programs could be beneficial.

Even though we observed very specific cases where the inclusion of at least one fixed-effect covariate resulted in an increased prediction accuracy relative to a standard RR-BLUP model, for the majority of the simulated genetic architectures, we found that such a GS + de novo GWAS approach actually led to a decrease in prediction accuracy. Given the promising results presented in previous studies where similar approaches were tested (Arruda et al., 2016; Bernardo, 2014; Spindel et al., 2016; Zhang et al., 2014), the frequency of how often we observed these negative results was surprising. One potentially important contributing factor underlying these results is differential LD patterns between the markers included as fixed-effect covariates and the QTN (s) that they are tagging. To approximate the realistic situation that causal mutations are unlikely to be included in marker data sets for GWAS and GS, we removed all markers selected to be QTN from consideration for GS + de novo GWAS. Thus all markers included as fixed-effect covariates were, at best, in high LD with the QTNs underlying the simulated traits. It is plausible

that the

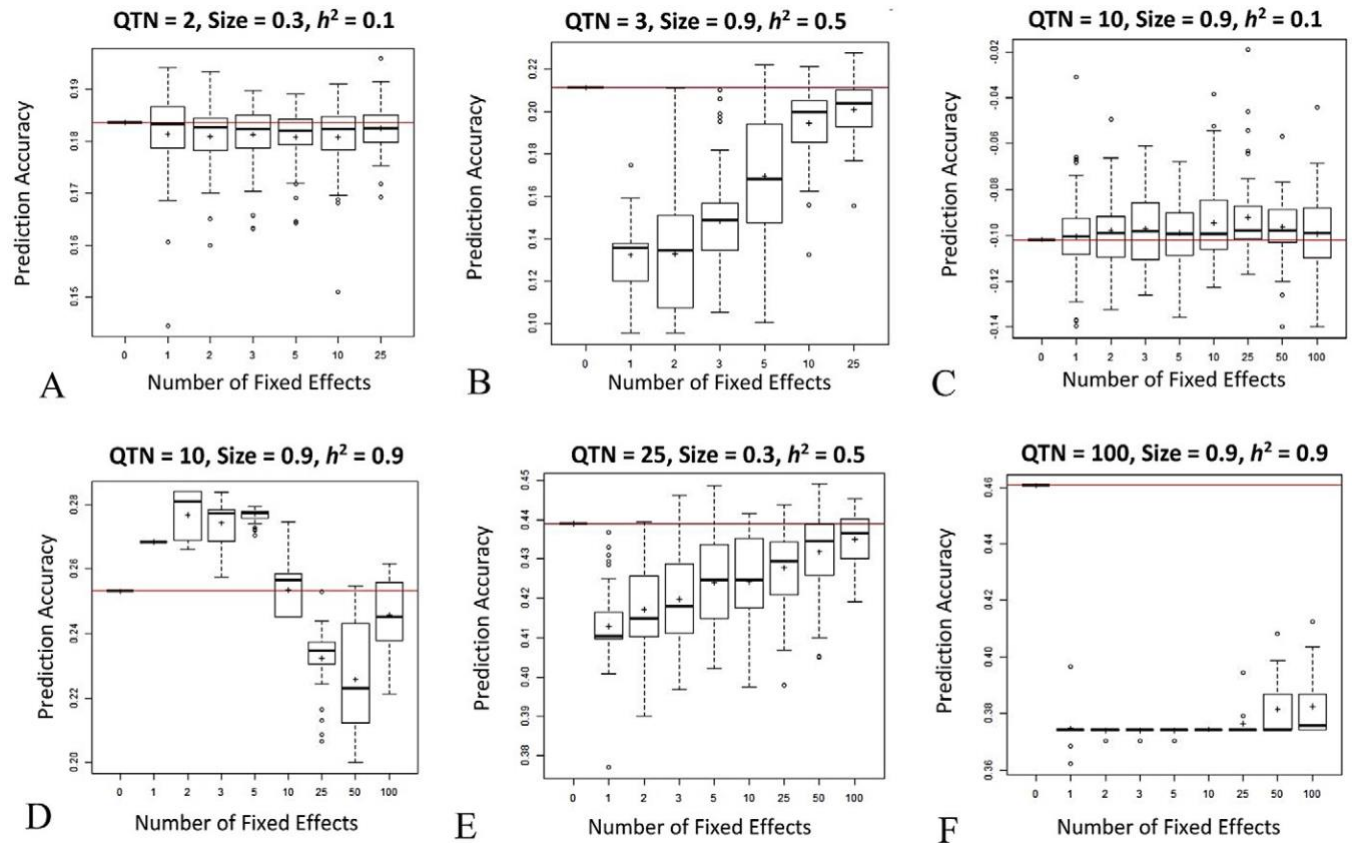


Fig. 3. Distribution of mean prediction accuracy for various genetic architectures in six maize Type 2 traits. The simulated genetic architecture is listed in the individual titles (A–F) where: QTN, number of simulated quantitative trait nucleotides; size, additive effect size of the largest QTN; h^2 , narrow-sense heritability. The x-axis is the number of fixed effects included in model (2). The y-axis is the mean Pearson correlation from the five-fold cross validation. In each graph, the mean five-fold cross-validation prediction accuracies of all 50 replications are used to generate these box plots. The orange line on each graph depicts the median prediction accuracy of the 50 replications when no markers are included in the model as fixed-effect covariates.

strength of LD between such markers and QTNs differ between the training and validation sets. This differential LD could result in the strength of the marker–trait associations to differ between these two sets. Interestingly, when a subset of the analysis was reran with the QTNs considered for inclusion as a fixed-effect covariate, the median prediction accuracies either decreased or did not change (depending on the number of fixed-effect covariates considered), while for low-heritable traits the variability of prediction accuracies increased (Supplemental Fig. 3E–9E). Collectively, these findings suggest that a marker strongly associated with a trait in the training set may have a substantially weaker association in the validation set; inclusion of such a marker as a fixed-effect covariate may offer either no advantage or even a disadvantage over a standard RR-BLUP model with respect to prediction accuracy.

Four of the traits analyzed by Spindel et al. (2016) reported to have increases in prediction accuracy when fixed effects were included in the RR-BLUP model were compared with our simulation study (Supplemental Table S3). There were also increases in prediction accuracy for simulation settings we evaluated that were similar to the genetic architecture of the traits from Spindel

et al. (2016), although our observed increases were not as large. Given this comparison, it is important to note that Spindel et al. (2016) performed their analysis using a structured rice breeding population and the peak SNPs considered for inclusion as fixed effects were the result of a binning procedure. In contrast, the analysis presented in this study considered all markers (except for the simulated QTN) as possible candidates for fixed-effect covariates regardless of their physical proximity to each other.

Our results also highlight two other potential disadvantages of undertaking a GS + de novo GWAS approach. First, we observed that including fixed-effect covariates in an RR-BLUP model typically resulted in an increased variability in mean prediction accuracy, as visualized in Fig. 2 through 5. This finding suggests that, when applied to data similar to the ones that we studied, a standard RR-BLUP model is more likely to produce stable and consistent prediction accuracies across minor perturbations of the data. A second disadvantage we noted was that the inclusion of at least one fixed-effect covariate could potentially introduce a bias in the GEBVs (Table 4). This result is consistent with those from a study that assessed the predictive ability of an a priori

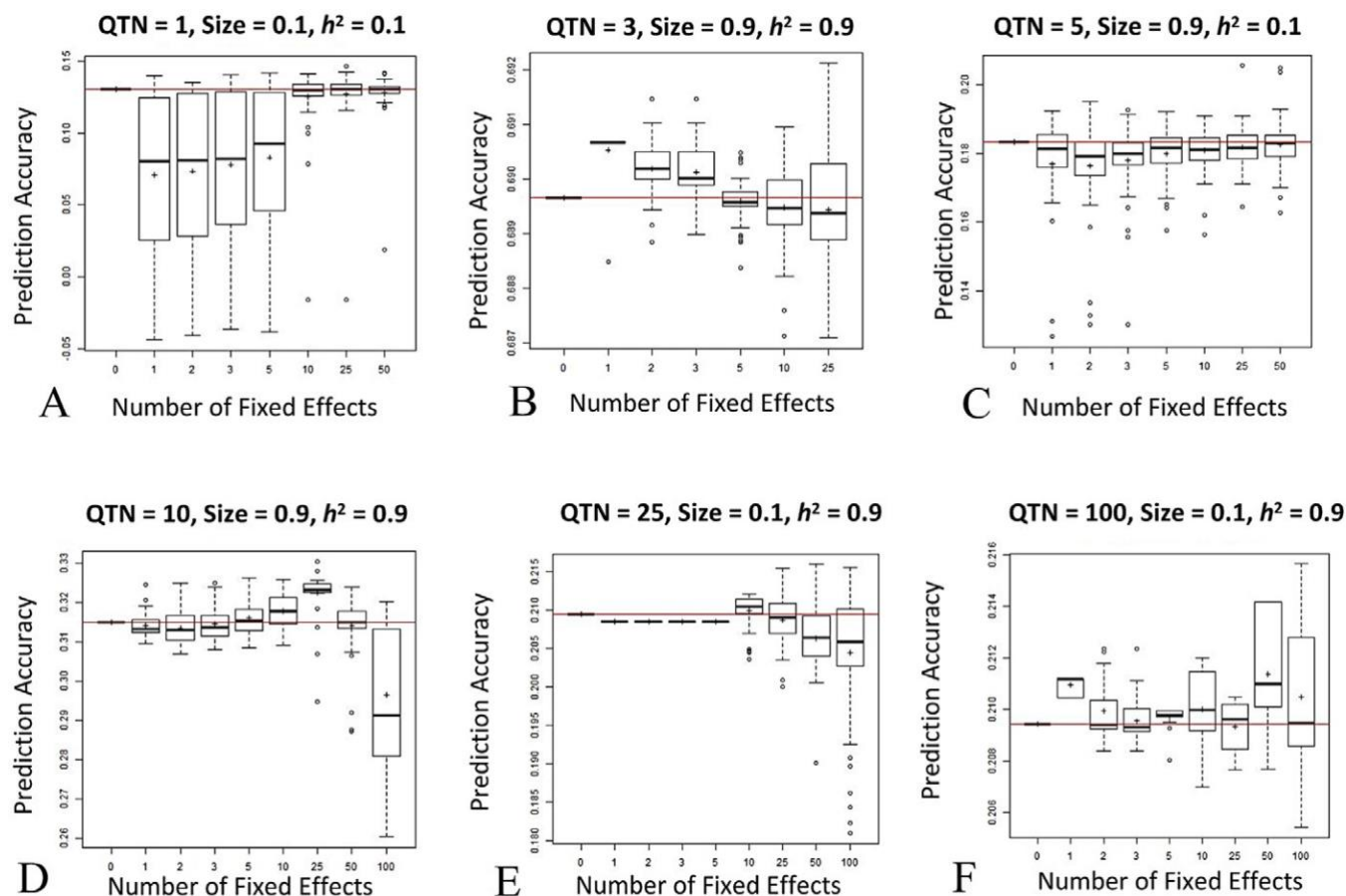


Fig. 4. Distribution of mean prediction accuracy for various genetic architectures in six sorghum Type 1 traits. The simulated genetic architecture is listed in the individual titles (A–F) where: QTN, number of simulated quantitative trait nucleotides; size, additive effect size of the largest QTN; h^2 , narrow-sense heritability. The x-axis is the number of fixed effects included in model (2). The y-axis is the mean Pearson correlation from the five-fold cross-validation. In each graph, the mean five-fold cross-validation prediction accuracies of all 50 replications are used to generate these box plots. The orange line on each graph depicts the median prediction accuracy of the 50 replications when no markers are included in the model as fixed-effect covariates.

(Raymond et al., 2018) and suggest that incorporating fixed-effect covariates into an RR-BLUP model may yield GEBVs that are substantially larger or smaller than a corresponding observed breeding value. We therefore conclude that any potential of a GS + de novo GWAS approach to increase prediction accuracies needs to be weighed against possible increases in the variability of prediction accuracies and bias of GEBVs.

Performance in Maize versus Sorghum

Although the scope of the work presented here is not extensive enough to extrapolate far beyond the two diversity panels that we studied, we did note that Model 2 out-performed Model 1 more often in sorghum than maize (Table 3). Perhaps the most notable difference between these species is their breeding patterns. Maize is a natural outcrossing species resulting in a LD pattern that decays on average at 2 kb (Remington et al., 2001). This results in relatively smaller LD blocks than sorghum, a natural inbreeding species, which decays around 150 kb (Morris et al., 2013). It is expected that inbreeding species present larger LD patterns than outcrossing ones (Flint-Garcia et

al., 2003). Higher LD means a stronger chance of detecting a marker closely associated with the causal mutation. A second major difference between the two data sets used for simulation is the genotyping method used to obtain the markers. Genotyping arrays, like the one used to collect the 55k maize genotypes, only include a fraction of SNPs present in a restricted set of lines (Brachi et al., 2011). For sorghum, GBS data were used; in contrast to a SNP array, GBS has the potential to capture near complete genomic data in any species (Andolfatto et al., 2011; Elshire et al., 2011). Recent evidence in wheat suggests GBS may have an advantage because of ascertainment bias associated with SNP arrays (Elbasyoni et al., 2018). Arguably though, the most notable difference between genotypic data sets was the number of markers available (sorghum = 90,441; maize = 51,741) and the number of individuals in the panel (sorghum = 320; maize = 281). Given the approximate size of the sorghum and maize genomes (respectively 800 million and 2.3 billion bp) (McCormick et al., 2018, Schnable et al., 2009), this translates to an average marker density of 78846 bp per marker in sorghum and 44,452 bp per marker in maize. A denser marker coverage of the genome

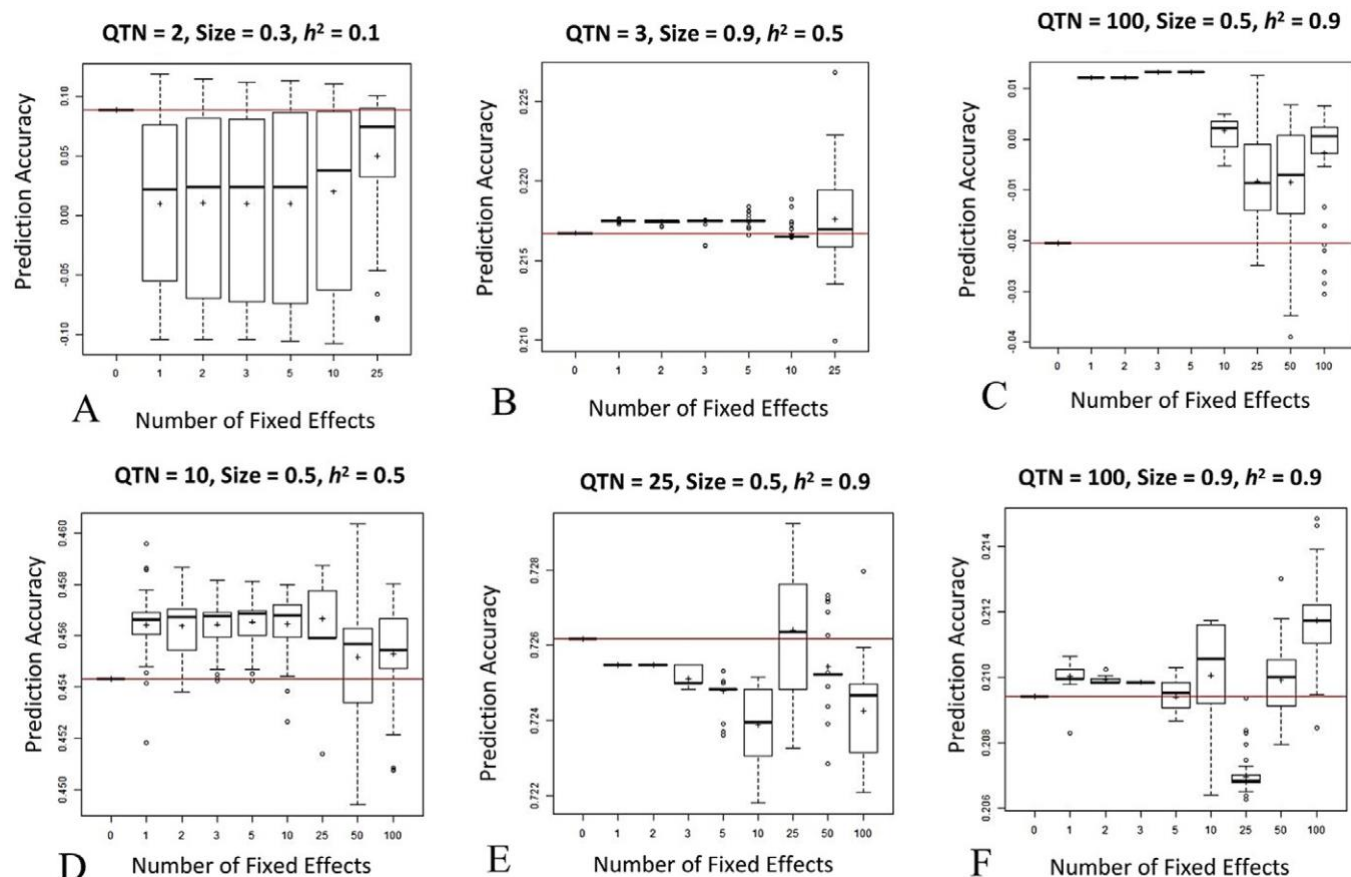


Fig. 5. Distribution of mean prediction accuracy for various genetic architectures in six sorghum Type 2 traits. The simulated genetic architecture is listed in the individual titles (A–F) where: QTN, number of simulated quantitative trait nucleotides; size, additive effect size of the largest QTN; h^2 , narrow-sense heritability. The x-axis is the number of fixed effects included in model (2). The y-axis is the mean Pearson correlation from the five-fold cross-validation. In each graph, the mean five-fold cross-validation prediction accuracies of all 50 replications are used to generate these box plots. The orange line on each graph depicts the median prediction accuracy of the 50 replications when no markers are included in the model as fixed-effect covariates.

raises the chance of a detecting associations as there is a relationship between LD and physical distance, albeit dependent on the population under study. This relationship has been demonstrated for the complex trait of height

Table 4. Intercept and slope estimates from a fitted simple linear regression model with observed breeding values as the response variable and the genomic estimated breeding values (GEBVs) as the explanatory variable. The predicted GEBVs were obtained from ridge-regression best linear unbiased prediction (RR-BLUP) models with the number of fixed-effect covariates (presented in the leftmost column) from a genome-wide association study (GWAS) conducted in a corresponding training set. The results presented here are from a sorghum Type 2 trait with a total of 100 quantitative trait nucleotides (QTNs), effect size of the largest QTN equal to 0.3, and a narrow-sense heritability of 0.9. The standard errors of the intercept and slope estimates are provided in parentheses.

No. offixed-effect covariates	Interceptestimate	Slope estimate
0	0.29 (0.01)	0.98 (0.07)
1	0.28 (0.01)	1.69 (0.25)
5	0.29 (0.01)	1.40 (0.15)
10	0.29 (0.01)	1.18 (0.09)
25	0.29 (0.01)	1.13 (0.09)
100	0.29 (0.01)	0.98 (0.07)

in maize (Peiffer et al., 2014). In addition, an increase in the number of individuals in a data set boosts the power to detect significant associations and improves the estimation of marker effect sizes in the training set (Zhong et al., 2009). Thus it would be interesting to see if the observed performance of Model 2 in sorghum could be achieved in maize if a data set with a larger sample size, greater marker density, and GBS data were used.

Recommendations for Future Research

Even though our results suggest that adding peak-associated markers as fixed-effect covariates to an RR-BLUP model is more likely to decrease rather than increase prediction accuracy, it could also serve as a foundation for future research. For example, instead of using predetermined numbers of peak-associated markers from GWAS (conducted in the training populations) as fixed-effect covariates, future studies could only consider GWAS markers that are statistically significantly associated with the target traits as fixed-effect covariates. Although an exploratory analysis of the use of such an approach on a subset of our simulated traits yielded similar results to the work presented already in this manuscript

(Supplemental Fig. S3–S9), further research into methodologies for determining which markers to include as fixed-effect covariates is warranted.

Another result that deserves attention was the scenarios where, as new fixed effects were added to Model 2, prediction accuracies first decreased sizably and then recovered to levels of Model 1 and in a few instances outperformed it (e.g., Fig. 2E). We put forth two hypotheses to explain this finding. The first is that a smaller number of fixed-effect covariates insufficiently tagged the additive effects of the masked QTN alleles. Thus, a greater number of markers in LD with a large-effect QTN needed to be included in the model to sufficiently capture the signal. The second hypothesis is that the effects of certain markers that were not included as fixed-effect covariates were being underestimated as a result of the RR penalty. This hypothesis could be tested by setting markers as fixed effects not by order of significance but in a model building process that searches for the subset of fixed effects obtaining the highest prediction accuracy. Indeed, the purpose of genomic selection is to allow for all markers to predict trait values regardless of significance or effect size. Therefore, it is possible that when small-effect, nonsignificant markers are included as fixed effects, prediction accuracies could increase. These two hypotheses suggest that criteria other than statistical significance of marker-trait associations from a GWAS in a training set be explored to identify potential markers to include as fixed-effect covariates in Model 2. Given that the use of peak-associated markers from a rudimentary stepwise model selection-based GWAS as fixed-effect covariates yielded prediction accuracies that were either equivalent or slightly higher than fixed-effect covariates identified through unified MLM-based GWAS (Supplemental Fig. S7–S9), future study of stepwise model selection or other marker-selection criteria (e.g., associations with RNA levels or protein expression levels) is justified.

Considering both the positive and negative findings presented in this work, it will be essential for future research on this topic to go beyond the limited range of data and simulation settings that we explored. That is, the largest number of QTN simulated was 100, which may or may not represent traits with highly polymorphic genetic backgrounds. All effects we considered were additive, and thus this study did not include the simulation of epistatic or dominance effects. Given that maize and sorghum are both diploid grass species with well-annotated genomes (Paterson et al., 2009; Schnable et al., 2009), species without as many available genetic markers may not be suitable for this approach if an insufficient amount of them are in LD with causal mutations. However, this could possibly become a nonissue in the near future because of the increasing availability of extensive sequence data (Goodwin et al., 2016; Poland and Rife, 2012).

conclusions

With the current wealth of available genomic data in maize and sorghum, it is theoretically possible to tailor

GS models with specific markers as fixed-effect covariates so that they reflect the genomic sources of a target trait as accurately as possible. Although we observed a maximum of 7.9% increase in prediction accuracy with such a model, we more frequently noted that the GS + de novo GWAS approach yielded lower prediction accuracies than the standard RR-BLUP model. We therefore do not recommend the universal implementation of GS + de novo GWAS for predicting the breeding values of all possible traits. Instead, we suggest that the merits of such an approach be investigated on a trait-by-trait basis, in particular through a *k*-fold cross validation scheme similar to what we used here. If, for a given trait, it turns out that the highest prediction accuracies are obtained with $m = 0$ fixed-effect covariates, then GS should be performed through a standard RR-BLUP model, for example. We ultimately feel that this research highlights the disadvantages as well as the advantages of GS + de novo GWAS, and we encourage trait-specific consideration of this approach before it is implemented into breeding programs.

Author Contributions

BR and AEL conceived the experiment and BR conducted the analyses and summarized the results. BR and AEL interpreted the results and wrote the manuscript.

Conflict of Interest

The authors declare no conflict of interest.

ACKNOWLEDGMENTS

We acknowledge and thank Patrick Brown, Matthew Hudson, and Marcus Olatoye for the feedback and helpful suggestions for this presented work. This work was funded by the Lawrence E. Schrader & Elfriede Massier Plant Physiology Fellowship, USDA National Institute of Food and Agriculture project accession number 1005564 and University of Illinois startup funds.

REFERENCES

- Andolfatto, P., D. Davison, D. Erezylmaz, T.T. Hu, J. Mast, T. Sunayama-Morita, et al. 2011. Multiplexed shotgun genotyping for rapid and efficient genetic mapping. *Genome Res.* 21:610–617. doi:10.1101/gr.115402.110
- Arruda, M.P., A.E. Lipka, P.J. Brown, A.M. Krill, C. Thurber, G. Brown-Guedira et al. 2016. Comparing genomic selection and marker-assisted selection for *Fusarium* head blight resistance in wheat (*Triticum aestivum* L.). *Mol. Breeding* 36:84. doi:10.1007/s11032-016-0508-5
- Barton, N.H., A.M. Etheridge, and A. Veber. 2017. The infinitesimal model: Definition, derivation, and implications. *Theor. Popul. Biol.* 118:50–73. doi:10.1016/j.tpb.2017.06.001
- Benjamini, Y., and Y. Hochberg. 1995. Controlling the false discovery rate—A practical and powerful approach to multiple testing. *J. Roy. Stat. Soc. B Met.* 57:289–300.
- Bernardo, R. 2010. *Breeding for quantitative traits in plants*, 2nd ed. Stemma Press, Woodbury.
- Bernardo, R. 2014. Genomewide selection when major genes are known. *Crop Sci.* 54:68–75. doi:10.2135/cropsci2013.05.0315
- Bouchet, S., M.O. Olatoye, S.R. Marla, R. Perumal, T. Tesso, J.M. Yu, et al. 2017. Increased power to dissect adaptive traits in global sorghum diversity using a nested association mapping population. *Genetics* 206:573–585. doi:10.1534/genetics.116.198499
- Brachi, B., G.P. Morris, and J.O. Borevitz. 2011. Genome-wide association studies in plants: The missing heritability is in the field. *Genome Biol.* 12:232. doi:10.1186/gb-2011-12-10-232

- Caldwell, K.S., J. Russell, P. Langridge, and W. Powell. 2006. Extreme population-dependent linkage disequilibrium detected in an inbreeding plant species, *Hordeum vulgare*. *Genetics* 172:557–567. doi:10.1534/genetics.104.038489
- Campbell, M.T., Q. Du, K. Liu, C.J. Brien, B. Berger, C. Zhang, et al. 2017. A comprehensive image-based phenomic analysis reveals the complex genetic architecture of shoot growth dynamics in rice (*Oryza sativa*). *Plant Genome* 10. doi:10.3835/plantgenome2016.07.0064
- Collard, B.C., and D.J. Mackill. 2008. Marker-assisted selection: An approach for precision plant breeding in the twenty-first century. *Philos. Trans. R. Soc. Lond. B Biol. Sci.* 363:557–572. doi:10.1098/rstb.2007.2170
- Cook, J.P., M.D. McMullen, J.B. Holland, F. Tian, P. Bradbury, J. Ross-Ibarra, et al. 2012. Genetic architecture of maize kernel composition in the nested association mapping and inbred association panels. *Plant Physiol.* 158:824–834. doi:10.1104/pp.111.185033
- Cuevas, J., J. Crossa, V. Soberanis, S. Perez-Elizalde, P. Perez-Rodriguez, G.L. Campos, et al. 2016. Genomic prediction of genotype × environment interaction kernel regression models. *Plant Genome* 9. doi:10.3835/plantgenome2016.03.0024
- da Silva Romero, A.R., F. Siqueira, G.G. Santiago, L.C.D. Regitano, M.D. de Souza, R.A.D. Torres, et al. 2018. Prospecting genes associated with navel length, coat and scrotal circumference traits in Canchim cattle. *Livest. Sci.* 210:33–38. doi:10.1016/j.livsci.2018.02.004
- de los Campos, G., J.M. Hickey, R. Pong-Wong, H.D. Daetwyler, and M.P.L. Calus. 2013. Whole-genome regression and prediction methods applied to plant and animal breeding. *Genetics* 193:327–345. doi:10.1534/genetics.112.143313
- Divilov, K., P. Barba, L. Cadle-Davidson, and B.I. Reisch. 2018. Single and multiple phenotype QTL analyses of downy mildew resistance in interspecific grapevines. *Theor. Appl. Genet.* 131:1133–1143. doi:10.1007/s00122-018-3065-y
- Doebley, J., A. Stec, and C. Gustus. 1995. teosinte branched1 and the origin of maize: Evidence for epistasis and the evolution of dominance. *Genetics* 141:333–346.
- Dudewicz, E.J., J.S. Ramberg, and H.J. Chen. 1975. New tables for multiple comparisons with a control (Unknown variances). *Biom. J.* 17:13–26. doi:10.1002/bimj.19750170103
- Dunnett, C.W. 1955. A multiple comparison procedure for comparing several treatments with a control. *J. Am. Stat. Assoc.* 50:1096–1121. doi:10.1080/01621459.1955.10501294
- Elbasyoni, I.S., A.J. Lorenz, M. Guttieri, K. Frels, P.S. Baenziger, J. Poland, et al. 2018. A comparison between genotyping-by-sequencing and array-based scoring of SNPs for genomic prediction accuracy in winter wheat. *Plant Sci.* 270:123–130. doi:10.1016/j.plantsci.2018.02.019
- Elshire, R.J., J.C. Glaubitz, Q. Sun, J.A. Poland, K. Kawamoto, E.S. Buckler, et al. 2011. A robust, simple genotyping-by-sequencing (GBS) approach for high diversity species. *PLoS One* 6:e19379. doi:10.1371/journal.pone.0019379
- Endelman, J.B. 2011. Ridge regression and other kernels for genomic selection with R Package rrBLUP. *Plant Genome* 4:250–255. doi:10.3835/plantgenome2011.08.0024
- Fisher, R.A. 1919. XV.—The correlation between relatives on the supposition of Mendelian inheritance. *Earth Environ. Sci. Trans. R. Soc. Edinburgh* 52:399–433. doi:10.1017/S0080456800012163
- Fisher, R.A. 1930. *The genetical theory of natural selection: A complete variorum edition*. Clarendon Press, Oxford. doi:10.5962/bhl.title.27468
- Flint-Garcia, S.A., J.M. Thornsberry, and E.S. Buckler IV. 2003. Structure of linkage disequilibrium in plants. *Annu. Rev. Plant Biol.* 54:357–374. doi:10.1146/annurev.arplant.54.031902.134907
- Flint-Garcia, S.A., A.C. Thuitet, J.M. Yu, G. Pressoir, S.M. Romero, S.E. Mitchell, et al. 2005. Maize association population: A high-resolution platform for quantitative trait locus dissection. *Plant J.* 44:1054–1064. doi:10.1111/j.1365-3113X.2005.02591.x
- Gianola, D. 2013. Priors in whole-genome regression: The Bayesian alpha-bet returns. *Genetics* 194:573–596. doi:10.1534/genetics.113.151753
- Goodwin, S., J.D. McPherson, and W.R. McCombie. 2016. Coming of age: Ten years of next-generation sequencing technologies. *Nat. Rev. Genet.* 17:333–351. doi:10.1038/nrg.2016.49
- Guo, J., W. Shi, Z. Zhang, J. Cheng, D. Sun, J. Yu, et al. 2018. Association of yield-related traits in founder genotypes and derivatives of common wheat (*Triticum aestivum* L.). *BMC Plant Biol.* 18:38. doi:10.1186/s12870-018-1234-4
- Heslot, N., H.P. Yang, M.E. Sorrells, and J.L. Jannink. 2012. Genomic selection in plant breeding: A comparison of models. *Crop Sci.* 52:146–160. doi:10.2135/cropsci2011.06.0297
- Hoerl, A.E., and R.W. Kennard. 1970. Ridge regression: Biased estimation for nonorthogonal problems. *Technometrics* 12:55. doi:10.1080/00401706.1970.10488634
- Huang, W., and T.F.C. Mackay. 2016. The genetic architecture of quantitative traits cannot be inferred from variance component analysis. *PLoS Genet.* 12:e1006421. doi:10.1371/journal.pgen.1006421
- Jardim, J.G., B. Guldbrandtsen, M.S. Lund, and G. Sahana. 2018. Association analysis for udder index and milking speed with imputed whole-genome sequence variants in Nordic Holstein cattle. *J. Dairy Sci.* 101:2199–2212. doi:10.3168/jds.2017-12982
- Jiang, Y., and J.C. Reif. 2015. Modeling epistasis in genomic selection. *Genetics* 201:759–768. doi:10.1534/genetics.115.177907
- Lipka, A.E., M.A. Gore, M. Magallanes-Lundback, A. Mesberg, H.N. Lin, T. Tiede, et al. 2013. Genome-wide association study and pathway-level analysis of tocopherol levels in maize grain. *G3: Genes, Genomes, Genet.* 3:1287–1299. doi:10.1534/G3.113.006148
- Lipka, A.E., C.B. Kandianis, M.E. Hudson, J. Yu, J. Drnevich, P.J. Bradbury, et al. 2015. From association to prediction: Statistical methods for the dissection and selection of complex traits in plants. *Curr. Opin. Plant Biol.* 24:110–118. doi:10.1016/j.pbi.2015.02.010
- Lipka, A.E., F. Tian, Q.S. Wang, J. Peiffer, M. Li, P.J. Bradbury, et al. 2012. GAPIT: Genome association and prediction integrated tool. *Bioinformatics* 28:2397–2399. doi:10.1093/bioinformatics/bts444
- Loiselle, B.A., V.L. Sork, J. Nason, and C. Graham. 1995. Spatial genetic structure of a tropical understory shrub, *PSYCHOTRIA OFFICINALIS* (RuBIACEAE). *Am. J. Bot.* 82:1420–1425. doi:10.1002/j.1537-2197.1995.tb12679.x
- Lynch, M., and B. Walsh. 1998. *Genetics and analysis of quantitative traits*. Sinauer, Sunderland, MA.
- Mackay, T.F.C. 2001. The genetic architecture of quantitative traits. *Annu. Rev. Genet.* 35:303–339. doi:10.1146/annurev.genet.35.102401.090633
- McCormick, R.F., S.K. Truong, A. Sreedasyam, J. Jenkins, S. Shu, D. Sims, et al. 2018. The *Sorghum bicolor* reference genome: Improved assembly, gene annotations, a transcriptome atlas, and signatures of genome organization. *Plant J.* 93:338–354. doi:10.1111/tpj.13781
- Meuwissen, T.H.E., B.J. Hayes, and M.E. Goddard. 2001. Prediction of total genetic value using genome-wide dense marker maps. *Genetics* 157:1819–1829.
- Morris, G.P., P. Ramu, S.P. Deshpande, C.T. Hash, T. Shah, H.D. Upadhyaya, et al. 2013. Population genomic and genome-wide association studies of agroclimatic traits in sorghum. *Proc. Natl. Acad. Sci. USA* 110:453–458. doi:10.1073/pnas.1215985110
- Muqaddasi, Q.H., J. Brassac, A. Borner, K. Pillen, and M.S. Roder. 2017. Genetic architecture of anther extrusion in spring and winter wheat. *Front. Plant Sci.* 8:754. doi:10.3389/fpls.2017.00754
- Orr, H.A. 1998. The population genetics of adaptation: The distribution of factors fixed during adaptive evolution. *Evolution* 52:935–949. doi:10.1111/j.1558-5646.1998.tb01823.x
- Owens, B.F., A.E. Lipka, M. Magallanes-Lundback, T. Tiede, C.H. Diepenbrock, C.B. Kandianis, et al. 2014. A foundation for provitamin A biofortification of maize: Genome-wide association and genomic prediction models of carotenoid levels. *Genetics* 198:1699–1716. doi:10.1534/genetics.114.169979
- Paterson, A.H., J.E. Bowers, R. Bruggmann, I. Dubchak, J. Grimwood, H. Gundlach, et al. 2009. The *Sorghum bicolor* genome and the diversification of grasses. *Nature* 457:551–556. doi:10.1038/nature07723

- Peiffer, J.A., M.C. Romay, M.A. Gore, S.A. Flint-Garcia, Z. Zhang, M.J. Millard, et al. 2014. The genetic architecture of maize height. *Genetics* 196:1337–1356. doi:10.1534/genetics.113.159152
- Poland, J.A., and T.W. Rife. 2012. Genotyping-by-sequencing for plant breeding and genetics. *Plant Genome* 5:92–102. doi:10.3835/plantgenome2012.05.0005
- Poland, J., and J. Rutkoski. 2016. Advances and challenges in genomic selection for disease resistance. *Annu. Rev. Phytopathol.* 54:79–98. doi:10.1146/annurev-phyto-080615-100056
- Raymond, B., A.C. Bouwman, C. Schrooten, J. Houwing-Duistermaat, and R.F. Veerkamp. 2018. Utility of whole-genome sequence data for across-breed genomic prediction. *Genet. Sel. Evol.* 50:27. doi:10.1186/s12711-018-0396-8
- R Development Core Team. 2018. R: A language and environment for statistical computing. R Foundation for Statistical Computing, Vienna, Austria.
- Remington, D.L., J.M. Thornsberry, Y. Matsuoka, L.M. Wilson, S.R. Whitt, J. Doebley, et al. 2001. Structure of linkage disequilibrium and phenotypic associations in the maize genome. *Proc. Natl. Acad. Sci. USA* 98:11479–11484. doi:10.1073/pnas.201394398
- Resende, M.F.R., P. Munoz, M.D.V. Resende, D.J. Garrick, R.L. Fernando, J.M. Davis et al. 2012. Accuracy of genomic selection methods in a standard data set of loblolly pine (*Pinus taeda* L.). *Genetics* 190:1503–1510. doi:10.1534/genetics.111.137026
- Riedelsheimer, C., F. Technow, and A.E. Melchinger. 2012. Comparison of whole-genome prediction models for traits with contrasting genetic architecture in a diversity panel of maize inbred lines. *BMC Genomics* 13:452. doi:10.1186/1471-2164-13-452
- Schnable, P.S., D. Ware, R.S. Fulton, J.C. Stein, F. Wei, S. Pasternak, et al. 2009. The B73 maize genome: Complexity, diversity, and dynamics. *Science* 326:1112–1115. doi:10.1126/science.1178534
- Spindel, J.E., H. Begum, D. Akdemir, B. Collard, E. Redona, J.L. Jannink, et al. 2016. Genome-wide prediction models that incorporate de novo GWAS are a powerful new tool for tropical rice improvement. *Heredity* 116:395–408. doi:10.1038/hdy.2015.113
- Spindel, J., H. Begum, D. Akdemir, P. Virk, B. Collard, E. Redona, et al. 2015. Genomic selection and association mapping in rice (*Oryza sativa*): Effect of trait genetic architecture, training population composition, marker number and statistical model on accuracy of rice genomic selection in elite, tropical rice breeding lines. *PLoS Genet.* 11:e1004982. doi:10.1371/journal.pgen.1004982
- Sukumaran, S., M. Lopes, S. Dreisigacker, and M. Reynolds. 2018. Correction to: Genetic analysis of multi-environmental spring wheat trials identifies genomic regions for locus-specific trade-offs for grain weight and grain number. *Theor. Appl. Genet.* 131:999. doi:10.1007/s00122-018-3066-x
- Technow, F., C. Riedelsheimer, T.A. Schrag, and A.E. Melchinger. 2012. Genomic prediction of hybrid performance in maize with models incorporating dominance and population specific marker effects. *Theor. Appl. Genet.* 125:1181–1194. doi:10.1007/s00122-012-1905-8
- Tibshirani, R. 1996. Regression shrinkage and selection via the Lasso. *J. Roy. Stat. Soc. B Met.* 58:267–288.
- Xu, Y.B., and J.H. Crouch. 2008. Marker-assisted selection in plant breeding: From publications to practice. *Crop Sci.* 48:391–407. doi:10.2135/cropsci2007.04.0191
- Yu, J.M., G. Pressoir, W.H. Briggs, I.V. Bi, M. Yamasaki, J.F. Doebley, et al. 2006. A unified mixed-model method for association mapping that accounts for multiple levels of relatedness. *Nat. Genet.* 38:203–208. doi:10.1038/ng1702
- Zhang, Z., U. Ober, M. Erbe, H. Zhang, N. Gao, J.L. He et al. 2014. Improving the accuracy of whole genome prediction for complex traits using the results of genome wide association studies. *PLoS One* 9:e93017. doi:10.1371/journal.pone.0093017
- Zhong, S., J.C. Dekkers, R.L. Fernando, and J.L. Jannink. 2009. Factors affecting accuracy from genomic selection in populations derived from multiple inbred lines: A Barley case study. *Genetics* 182:355–364. doi:10.1534/genetics.108.098277