



Reinventing quantitative genetics for plant breeding: something old, something new, something borrowed, something BLUE

Rex Bernardo¹

Received: 23 January 2020 / Revised: 23 March 2020 / Accepted: 23 March 2020
© The Author(s) 2020. This article is published with open access

Abstract

The goals of quantitative genetics differ according to its field of application. In plant breeding, the main focus of quantitative genetics is on identifying candidates with the best genotypic value for a target population of environments. Keeping quantitative genetics current requires keeping old concepts that remain useful, letting go of what has become archaic, and introducing new concepts and methods that support contemporary breeding. The core concept of continuous variation being due to multiple Mendelian loci remains unchanged. Because the entirety of germplasm available in a breeding program is not in Hardy–Weinberg equilibrium, classical concepts that assume random mating, such as the average effect of an allele and additive variance, need to be retired in plant breeding. Doing so is feasible because with molecular markers, mixed-model approaches that require minimal genetic assumptions can be used for best linear unbiased estimation (BLUE) and prediction. Plant breeding would benefit from borrowing approaches found useful in other disciplines. Examples include reliability as a new measure of the influence of genetic versus nongenetic effects, and operations research and simulation approaches for designing breeding programs. The genetic entities in such simulations should not be generic but should be represented by the pedigrees, marker data, and phenotypic data for the actual germplasm in a breeding program. Over the years, quantitative genetics in plant breeding has become increasingly empirical and computational and less grounded in theory. This trend will continue as the amount and types of data available in a breeding program increase.

Introduction

The marriage between quantitative genetics and plant breeding, albeit nonexclusive, has reaped benefits for both during the last 100 years. Estimates of genetic variances and heritability in many plant species have increased our understanding of the inheritance and variation of important traits related to yield, quality, and adaptation (Gardner 1963; Matzinger 1963; Hallauer and Miranda 1988). Quantitative genetics principles have led to the design and refinement of breeding methods for continuous traits (Dudley and Moll 1969). Mapping studies have identified major quantitative trait

loci (QTL) that have been found useful in crop improvement, and that have helped elucidate the nature of quantitative variation (Kearsey and Farquhar 1998; Bernardo 2020).

That being said, developments in both plant breeding and quantitative genetics in the last several decades should cause us to pause and ponder how quantitative genetics can best be applied in contemporary plant breeding. The purpose of this *Opinion* article is to provide a framework for reflection, discussion, and constructive debate of reasons and ways to reinvent quantitative genetics within the context of current plant breeding. The ideas proposed herein apply mainly to this context and they might not apply to other fields, such as human genetics and animal improvement, where quantitative genetics has also played an important role.

Guest editor: Lindsey Leach

✉ Rex Bernardo
bernardo@umn.edu

¹ Department of Agronomy and Plant Genetics, University of Minnesota, 411 Borlaug Hall, 1991 Buford Circle, Saint Paul, MN 55108, USA

Origins and foundations of quantitative genetics

Plant breeders may sometimes think that quantitative genetics was developed for the purpose of enhancing plant and animal improvement. After all, most economically

important traits in crops and livestock species are quantitative rather than qualitative. However, quantitative genetics was developed not to provide a basis for artificial selection in different species, but to model how variation for continuous traits arises from unknown genes. Quantitative genetics originated more than a century ago in the aftermath of a heated dispute between two groups of eminent English scientists: the biometricians led by Karl Pearson and W.F.R. Weldon, and the Mendelists led by William Bateson (Provine 1971).

The field of biometry began with Francis Galton's work on how human characteristics, particularly height, were passed from parent to offspring (Galton 1869, 1889). Galton observed a regression toward mediocrity (now known as regression toward the mean) in that tall parents tended to have offspring shorter than them, and short parents tended to have offspring taller than them. For height, Galton estimated that the regression of offspring on mid-parent values was $2/3$. To illustrate, suppose the mean height of a father and mother exceeded the mean of the population by 9 cm. The height of their offspring would vary but, on average, was expected to exceed the population mean by $2/3 \times 9 = 6$ cm. Pearson later analyzed Galton's data as well as subsequent data (Pearson and Lee 1903), and found that the correlation of the height of fathers and mothers with the height of their sons and daughters had a mean of 0.51. Fraternal correlations were found to be slightly higher (mean of 0.53) than parent–offspring correlations. Overall, these studies by Galton and Pearson suggested that continuous variation in humans is at least partially inherited, although the mechanisms for the transmission of such traits remained a mystery.

These findings by the biometricians disagreed with Mendel's findings, which were rediscovered independently in 1900 by Hugo DeVries, Carl Correns, and Erich von Tschermak. Mendel's results indicated a particulate nature of inheritance that led to distinct classes, rather than a continuum of observations for a given trait. The absence of distinct phenotypic classes then suggested to the Mendelists that continuous traits are not heritable. DeVries went to the extent of considering the continuous versus discontinuous nature of variation as a criterion for a trait's transmissibility (Mather 1949).

In 1918, R.A. Fisher wrote a paper entitled “*The correlation between relatives on the supposition of Mendelian inheritance*” that sought to reconcile the conflicting views of the biometricians and Mendelists. In so doing, this seminal article established the foundations of quantitative genetics as we know it today. The “*on the supposition*” phrase in the title indicated the conjecture that quantitative traits are controlled by genes that behave in a Mendelian fashion, and the crux of Fisher's supposition was that continuous variation is the result of the cumulative effects

of many such genes. Fisher showed that such a model for quantitative variation could lead to the well-established correlations between relatives found by the biometricians. Fisher initially made two key assumptions to make the model tractable: the individuals mated at random, and each biallelic locus behaved independently of others. These two simplifying assumptions enabled genotype frequencies to be expressed in terms of allele frequencies and vice versa according to Hardy–Weinberg equilibrium, and allowed the cumulative effects across loci to be described in terms of the effect at an individual locus. At several junctures in his article, Fisher investigated the effects of multiple alleles, linkage, assortative mating, and epistasis.

A key feature of Fisher's derivations was the partitioning of effects into a transmissible and a non-transmissible portion. The transmissible or “additive” portion was due to the effects of individual alleles passed on by a parent to its offspring. These transmissible effects of alleles have since become known as the average effect of an allele (Fisher 1941). The variance due to the average effects of the two alleles at a locus in a diploid has since become known as additive variance (V_A). On the other hand, the non-transmissible portion of effects included dominance deviations as well as any environmental effects; to Fisher, dominance deviations were “*a residue which acts in much the same way as an arbitrary error introduced into the measurements.*” The variance due to dominance deviations has since become known as dominance variance (V_D). Fisher recognized the existence of effects due to interactions of alleles at different loci and treated such effects as higher-order extensions in a linear model: “*We may use the term Epistacy to describe such deviation, which although potentially more complicated, has similar statistical effects to dominance.*”

Fisher revisited Galton's concept of parent–offspring regression and derived this regression as being equal to $\frac{1}{2}V_A/(V_A + V_D)$. Because Fisher treated V_D as a residual value that was not any different from random error, we can modify this parent–offspring regression as being equal to $\frac{1}{2}V_A/V_P = \frac{1}{2}h^2$, where V_P is the phenotypic variance and $h^2 = V_A/V_P$ is the narrow-sense heritability. Returning to our height example, suppose the difference between the height of a father and the population mean is denoted by S . If the height of the father and mother is uncorrelated (i.e., mating is at random with respect to height), the expected deviation (from the population mean) of the height of the offspring is then equal to $\frac{1}{2}h^2S$.

We see this expression as clearly being related to the so-called breeder's equation of $R = h^2S$, where R is the response to selection and S is the selection differential (Lush 1937). The factor of $\frac{1}{2}$ is due to regression on only one of the two parents, and this situation is equivalent to mass selection without pollen control in plants, for which the genetic gain quantified via R is equal to $\frac{1}{2}h^2S$. This brief historical summary shows that although quantitative

genetics was initially developed as a basic rather than an applied science, quantitative genetics does provide a framework for designing breeding programs that can maximize genetic gain.

Why reinvent quantitative genetics for plant breeding?

Through the years, scientists have investigated the confluence of the science of quantitative genetics and the process of plant breeding. The 1940s saw investigations of the genetic basis of heterosis (Comstock and Robinson 1948; Robinson et al. 1949) and the development of recurrent selection procedures (Jenkins 1940; Hull 1945; Comstock et al. 1949) to address perceived limitations in line and hybrid development (Jenkins 1934). The 1950s and 1960s have been described as the golden era for quantitative genetics in plant breeding (Gardner 1977), given the research during these two decades on covariances between relatives (Kempthorne 1954; Cockerham 1956); mating designs to estimate V_A , V_D , and epistatic variance (Comstock and Robinson 1952; Cockerham 1963); empirical estimates of genetic variances in different plant species (reviewed by Gardner 1963; Matzinger 1963; Hallauer and Miranda 1988); genotype $\times$ environment interaction (Sprague and Federer 1951); stability analysis (Finlay and Wilkinson 1963; Eberhart and Russell 1966); and index selection for multiple traits (Brim et al. 1959; Pešek and Baker 1969). The 1970s saw the start of investigations of isozymes as markers for quantitative traits (Stuber and Moll 1972; Hamrick and Allard 1975). Long-overdue work started in the 1980s on formal methods to select parents in breeding programs (Dudley 1984). “Bandwagons” related to plant breeding from the 1990s to the present (Bernardo 2016) have included linkage mapping of QTL, association mapping, genomewide prediction (or genomic selection), phenomics, envirotyping, and gene editing.

Plant breeding in the 2020s is markedly different from plant breeding in decades past, and some of the older quantitative genetics approaches described in the preceding paragraph may have become irrelevant. Described below are three reasons for reinventing quantitative genetics as it applies to plant breeding: different expectations, unmet assumptions, and new tools.

What breeders expect from quantitative genetics has changed

A half-century ago, plant breeders expected quantitative genetics to provide answers to multiple questions (Dudley and Moll 1969) that can be distilled into three:

- (1) Which germplasm is the most promising?
- (2) What type of variety or cultivar should be developed?
- (3) What breeding method should be used?

Methods in quantitative genetics could indeed provide answers to each of the above questions. A set of germplasm that has a superior mean and a large V_A would constitute promising breeding germplasm, and information on the mean and genetic variance can be combined into a usefulness criterion (Schnell 1983) that estimates the mean of a selected proportion of the best individuals in a given population. A large amount of heterosis and V_D , as has been found for maize (*Zea mays* L.) yield (Hallauer and Miranda 1988), would indicate that hybrid or synthetic cultivars are suitable. Versions of the breeder’s equation for different types of recurrent selection procedures could be used, along with estimates of V_A , V_D , and nongenetic variance, to identify which breeding procedures would lead to the largest predicted gain.

The reality, however, is that sufficient answers to the first two questions are obtained without any detailed quantitative-genetic analysis whatsoever. For example, a large set of germplasm can be phenotyped to identify candidates with the best mean performance. Information on germplasm origins or, since the 1990s, data on molecular markers can be used to assess germplasm diversity. Hybrid cultivars are feasible if a hybrid outperforms its parents by a substantial amount, and if hybrid seeds can be produced in a cost-effective manner, with estimates of V_D being unnecessary. As for the third question, the predicted R differs across recurrent selection methods such as recurrent mass selection, half-sib recurrent selection, full-sib recurrent selection, or reciprocal recurrent selection. However, breeders of major crop species such as maize, wheat (*Triticum aestivum* L.), rice (*Oryza sativa* L.), soybean (*Glycine max* L.) Merrill, and tomato (*Solanum lycopersicum*) have preferred the expediency of a nonrecurrent line and hybrid development system with biparental crosses instead of longer-term recurrent selection with broadbase populations. Because the variance among recombinant inbreds in a biparental cross is constant at $2 V_A$, there is little or no difference in the gains from line development via different methods (pedigree breeding, bulk method, single-seed descent, or doubled haploids) as long as reliable phenotyping of lines is done at some point during the line development process.

What, then, do current plant breeders expect from quantitative genetics? Today’s plant breeders expect one primary outcome from quantitative genetics: to help identify which candidates have the best genotypic value for a given set of continuous traits, with genotypic value being defined as the expectation of the performance of the candidate in a

target population of environments. The candidates would be individual plants, partially inbred lines, or recombinant inbreds in a self-pollinated species such as wheat; test-crosses or hybrids in a cross-pollinated species such as maize; or individual clones in an asexually propagated species such as cassava (*Manihot esculenta*). This singular emphasis on finding lines, hybrids, or clones with the best genotypic value is a much more focused objective than the above three questions from 50 years ago (Dudley and Moll 1969). This singular emphasis suggests that some of the classical methods used in quantitative genetics to address broader sets of questions have outlived their usefulness.

Classical assumptions in quantitative genetics are unmet in plant breeding

The two main assumptions invoked by Fisher (1918)—random mating and independence of segregating loci—are typically unmet in plant breeding programs. The assumption of random mating is met in species for which recurrent selection is the main breeding procedure, because each cycle of recurrent selection is created by random mating a group of selected individuals in the previous cycle. Recurrent selection is common in forage species such as alfalfa (*Medicago sativa* L.) and perennial ryegrass (*Lolium perenne* L.) but, as mentioned above, not in row crops such as maize and wheat. The F_2 of a biparental cross between homozygous maize or wheat parents has the same 1:2:1 genotype ratio attained via random mating. Each F_2 population can therefore be treated as a random-mating population, and V_A , V_D , and h^2 can be estimated, but such estimates apply only within the given F_2 population.

In most situations, the entirety of germplasm within a breeding program does not comprise a random-mating population. Maize-breeding germplasm in the United States, for example, includes lines derived from key progenitors, such as A632, B37, B73, Iodent, LH82, Maiz Amargo, Minnesota 13, Mo17, and Oh43 (Troyer 1999; Mikel and Dudley 2006). Maize breeding has focused on developing newer versions of inbreds that maintain these lineages or that combine only a few of these lineages at a time. The preponderance of these key lineages indicates that the maize germplasm in the breeding program cannot, as a whole, be assumed as representative of a random-mating population.

The assumption of linkage equilibrium or independence among segregating loci is particularly problematic (Cockerham 1963). Multiple generations of random mating are needed to achieve linkage equilibrium, especially for closely linked loci. While geneticists are sometimes able to create mapping populations that have undergone multiple generations of random mating (e.g., intermated B73 $\times$ Mo17 maize population (Lee et al. 2002)), breeders do not have that luxury in cultivar development programs.

The failure in plant breeding to meet the two main assumptions that underlie classical quantitative genetics theory indicates a need to revisit such theory or find ways to circumvent the assumptions.

New tools and computing capabilities have emerged

Two technological developments have enabled new approaches in both quantitative genetics and plant breeding. First, cheap and abundant molecular markers in different plant species have allowed new types of analysis that were not possible prior to the 1980s. Breeders and geneticists were limited to the use of only up to a few dozen isozyme markers in the 1970s, but the number of available markers increased with the development of restriction fragment length polymorphism (RFLP) markers in the 1980s (Beckmann and Soller 1983). Other types of markers developed after RFLP markers included random amplified polymorphic DNA markers, amplified fragment length polymorphism markers, and simple sequence repeat markers. The number of markers increased drastically with the development of single nucleotide polymorphism (SNP) markers which, unlike previous marker systems, are amenable to high-throughput genotyping platforms (Syvänen 2005). The costs of SNP genotyping have decreased since the 2010s to the extent that in major crop species, the per-sample cost of genotyping (Ertiro et al. 2015) is less than the per-individual cost of multi-environment phenotyping (<http://techservicespro.com/>).

Second, developments in computers have enhanced data analysis and simulation. G.F. Sprague, a renowned maize breeder who developed the concepts of general and specific combining ability in the 1940s (Sprague and Tatum 1942), once told me that in his day, having a winter nursery was infeasible because yield data from the fall harvest could not be analyzed in time to select candidates to include in a winter nursery. When I started my career with a seed company in 1988, data from a given year's yield trials were used only for the purpose of selecting the best lines and hybrids that year. Such data were not yet viewed as a valuable contribution to a cumulative, historical data set useful for predicting the performance of candidates in future years. This mindset within the company started to change after I developed genomic best linear unbiased prediction (GBLUP) for single-cross performance in 1994, with the first GBLUP calculations being implemented on a Pentium 90 machine that had 64 megabytes of random access memory (Bernardo 1994). Today's computers allow large-scale data analysis, such as solving a system of equations with up to 2 million unknowns (Gray 2016), and statistical resampling procedures, such as bootstrapping and cross-validation (Efron 1980), that were not possible before.

Something old, something new, something borrowed, something BLUE

Reinventing quantitative genetics for plant breeding requires revisiting classical concepts and approaches in quantitative genetics, keeping what remains useful, letting go of what has become obsolete, and considering new ideas. One may question whether a reinvention is needed, because an alternative approach is to simply let the language of quantitative genetics evolve with both use and misuse, to the extent that some current practices might bear little pertinence to the original concepts. We as practitioners would still be able to effectively communicate with each other because we all uniformly follow the same practices. A premise in this *Opinion* article, however, is that if we wish our science to be precise and rigorous, we then should be precise in defining key concepts and rigorous in adhering to them. The stance herein is that it is better to invent new concepts to accommodate new developments rather than to force new developments into classical concepts that might not be able to accommodate them.

Described below are eight main ideas that could underlie nequantitative genetics for plant breeding. These ideas are organized according to the old English rhyme of “something old, something new, something borrowed, something blue.”

Old: multiple loci plus environmental effects

As postulated by Fisher (1918), continuous variation will continue to be modeled as being due to the joint effects of multiple Mendelian loci. No assumptions need to be made regarding the number of underlying QTL. For most traits, however, the number of QTL is expected to be large and their individual effects are expected to be small.

The core model for quantitative variation will remain $P = G + E$, where P is the phenotypic value, G is the genotypic value, and E is a residual value due to nongenetic effects. More specifically, the phenotypic value of genotype i in environment j is modeled as

$$Y_{ij} = \mu + g_i + e_j + (ge)_{ij} + \varepsilon_{ij} \quad (1)$$

where μ is the grand mean; g_i is the effect of genotype i ; e_j is the effect of environment j ; $(ge)_{ij}$ is the genotype $\times$ environment interaction effect associated with genotype i and environment j ; and ε_{ij} is the within-environment error. This classical linear model remains unchanged. However, there may be multiple ways to model the g_i , e_j , and $(ge)_{ij}$ components. For example, $(ge)_{ij}$ can itself be modeled as a multiplicative effect obtained as the product of an interaction score due to genotype i and an interaction score due to environment j (Gollob 1968; Gauch 1988).

Envirotyping can be used to model the e_j component, at least in part, as a function of environmental variables such as precipitation and temperature (Cooper et al. 2014).

Old: identify candidates with the best genotypic value

Plant breeders have always been interested in identifying and selecting candidates with the best genotypic value, and these efforts will intensify as the number of candidates increases due to expansion of breeding programs. Plant breeding has always been predictive in the sense that in the $P = G + E$ equation, the G component is predicted from the observed P . Predictions of G from P will obviously become more precise as E approaches zero. Methods to make phenotyping more accurate and precise, so that E approaches zero, will continue to be important.

Phenotypic data routinely generated in a breeding program and SNP markers have been found useful for predicting the genotypic values of other candidates (see Bernardo (2020) for a review). Such predictions are commonly made via GBLUP, in which SNP markers are used to estimate relatedness among individuals (Lynch 1988; VanRaden 2008), or via an approach such as ridge regression–best linear unbiased prediction (RR–BLUP), in which the effects of each SNP marker are calculated from a set of related individuals (Meuwissen et al. 2001). The GBLUP and RR–BLUP approaches are equivalent when the number of QTL is large, no major QTL are present, and the QTL are evenly distributed across the genome (Fernando 1998; Habier et al. 2007). Given that GBLUP was developed more than a quarter-century ago (Bernardo 1994) and genomewide prediction via RR–BLUP or Bayesian models was proposed nearly 20 years ago (Meuwissen et al. 2001), we must consider both approaches for predicting genotypic value as old.

Old: continue finding major QTL

Major QTL alleles, such as *Fhb1* for Fusarium (*F. graminearum*) head blight resistance in wheat (Anderson et al. 2007) and *Sub1* for flooding tolerance in rice (Septiningsih et al. 2009), will continue to be useful in cultivar development. A major QTL has an effect large and consistent enough to be meaningful in a breeding program, which implies that a QTL might be considered as major in one breeding program but not in another (Bernardo 2014a). Major QTL may be present for traits such as morphology, phenology, and tolerance to biotic and abiotic stresses, but are likely absent for a highly selected trait such as yield in elite germplasm.

The expected change in the mean value for the trait should be used as the criterion to assess whether or not a

marker-trait association represents a major QTL. The R^2 value should not be used as the criterion because a high R^2 value may correspond to too small a predicted change. For example, studies at the University of Minnesota identified a marker with $R^2 = 27\%$ for oil concentration in maize (Garcia 2008). The positive QTL allele was predicted to increase kernel oil concentration from 3.5 to 5.5%. Because high-oil maize hybrids with up to 8.0% oil have been commercially available (Lambert et al. 1998), the effect of the QTL is deemed too small for it to be considered as a major QTL in maize breeding (Bernardo 2020).

New: no need for a reference population in Hardy–Weinberg equilibrium

As mentioned earlier, Fisher's (1918) assumption of having a random-mating population is typically violated in plant breeding. It is high time to openly acknowledge this fact rather than to pretend, by estimating parameters such as V_A and h^2 in a nonrandom-mated population (Sughrue and Hallauer 1997) or diversity panel, that the assumption of a random-mating population is met. Fisher's assumption of random mating was a matter of necessity in 1918. In contrast, today's availability of SNP markers allows breeders to track the transmission of chromosomal segments, and random mating therefore does not need to be assumed. Furthermore, because breeders are mainly interested in identifying candidates with superior genotypic values, there is no need to have a reference population to which inferences would apply—regardless of whether or not such a reference population in Hardy–Weinberg exists.

To illustrate, suppose recombinant inbreds are developed by selfing from a (Parent 1 $\times$ Parent 2) F_2 and from a ((Parent 1 $\times$ Parent 2) $\times$ Parent 1) BC_1 . Whereas the F_2 population can be assumed to behave as a random-mating population because the expected genotype ratio of 1:2:1 at segregating loci is the same as that with random mating, the same assumption cannot be made for the BC_1 population because the expected genotype ratio is 1:1 at segregating loci, with one of the homozygotes not being recovered. At this point, the breeder is simply interested in identifying the best recombinant inbreds, regardless of whether an inbred was derived from the F_2 or the BC_1 . The breeder has no need to make inferences regarding the mean or V_A or h^2 in the F_2 or BC_1 , and it is therefore moot that one population is in Hardy–Weinberg equilibrium and the other is not.

The preceding example also brings to light the historical divergence between two schools of thought in quantitative genetics: the Edinburgh/Alan Robertson/Falconer (1960) school that emphasized arbitrary allele frequencies in random-mating, outbred populations versus the Birmingham/Mather (1949) school that focused on crosses between homozygous lines for which allele frequencies are expected to be $\frac{1}{2}$ at

segregating loci. While the Edinburgh school has become more prevalent in plant and animal breeding because of its emphasis on artificial selection with arbitrary allele frequencies, the Birmingham school has features that lend themselves naturally to plant breeding as practiced today. At the same time, allele frequencies of $\frac{1}{2}$ are easily accommodated in the Edinburgh framework, and the F_2 genotype frequencies are those expected with random mating. Both the Edinburgh and Birmingham schools therefore apply to F_2 populations encountered in plant breeding.

Not needing a reference population also circumvents the issue of whether the candidates are random members of a base population or are a fixed set of lines, clones, or hybrids that are not random members of any population. Suppose the candidates are a set of wheat cultivars and pre-commercial lines with diverse genetic backgrounds. In this situation, the variance component for lines (if the lines were random) is substituted by $\sum c_i^2/(n-1)$, where c_i is the fixed effect of the i th candidate and n is the number of candidates.

New: no need to define different types of genetic variances

Classical quantitative genetics as founded by Fisher (1918) focused on breeding value. However, breeding value is of minimal value to plant breeders for two reasons. First, plant breeders are much more interested in genotypic value than on breeding value. The genotypic value is that of the candidate itself, whereas the breeding value is reflected by the mean of the candidate's progeny when it is mated with random individuals. To illustrate a key difference between animal and plant breeding, the breeding value of a dairy bull (*Bos taurus*) is of utmost importance because a top bull is prized not for its own milk yield (which is zero), but it is prized for the superior milk yield of its female progeny. In contrast, the genotypic value of a wheat line or a cassava clone is of utmost importance because producers would grow the wheat or cassava cultivar itself, rather than the cultivar's progeny. The ability to genetically replicate plants but not animals therefore plays a key role in this distinction. The foregoing does not imply that breeders are uninterested in the performance of the progeny of an individual: on the contrary, selection of good parents is a key to success in plant breeding. What the foregoing implies is that individuals used as future parents are first and foremost selected on the basis of their superior genotypic value as individuals.

Second, breeding value is defined only when the mate is chosen at random and, as Falconer (1985) indicated, "*The concept of breeding value is shown to have no useful meaning when mating is not random.*" The assumption of random mating, as discussed earlier, is typically violated in plant breeding.

If we accept that the classical concept of breeding value is not meaningful to plant breeders, neither would V_A be meaningful because V_A is the variance among breeding values. The V_D and epistatic variance (V_I) would likewise not be meaningful because these variances are derived from the same framework as V_A . Therefore, quantitative genetics reinvented for plant breeding would retire the concepts of V_A , V_D , and V_I . On the other hand, plant breeding students will still need to learn about V_A , V_D , and V_I so that they can understand classical literature and communicate with quantitative geneticists who work with non-plant species, for which these concepts remain important.

What then should be used as a substitute for V_A , V_D , and V_I ? The logical alternative is to simply calculate the variance component due to the candidates. In other words, breeders can calculate V_{F2} among F_2 plants; V_{F3} among F_3 lines; V_{RI} among recombinant inbreds; V_{Clones} among clones; V_{HS} among half-sib families; V_{FS} among full-sib families; and V_{SX} among single crosses. Single crosses are typically made between parents from two heterotic groups that complement each other, and V_{SX} can be partitioned into V_{GCA1} for the general combining ability (GCA) of parents from the first heterotic group, V_{GCA2} for the GCA of parents from the second heterotic group, and V_{SCA} for specific combining ability effects. It will be important to not use the symbol V_G for any of these variance components because V_G is traditionally defined as the sum of $V_A + V_D + V_I$. Avoiding the V_G notation will therefore reduce confusion.

Calculating variance components such as V_{F2} or V_{Clones} or V_{SX} has the advantage of being a direct expression of the variation due to genetic effects among the candidates undergoing selection. As previously mentioned, for example, the variance component for recombinant inbreds is $V_{RI} = 2V_A$. The V_{RI} therefore quantifies the amount of genetic variation expressed among the candidates, whereas V_A by itself does not. Fisher's (1918) assumption of linkage equilibrium is unnecessary for defining and estimating variance components for candidates.

Mating designs such as the factorial, nested, and diallel designs have been developed to estimate V_A and V_D (Cockerham 1963). This proposal renders such mating designs obsolete. Variance components due to candidates, as listed above, can instead be estimated by restricted maximum likelihood methods (Dempster et al. 1977; Harville 1977) within the framework of mixed-model analysis, which is described at the end of this paper.

Borrowed: focus on reliability and the least significant difference

Retiring the concepts of V_A , V_D , and V_I also means retiring the concept of h^2 . This should be of little practical consequence because, as previously mentioned, h^2 measures the proportion

of transmissible genetic effects, whereas breeders are more interested in the performance of the candidates themselves rather the performance of their progeny. On the other hand, plant breeders will continue to be interested in the influence of nature versus nurture on quantitative traits. Broad-sense heritability, which is defined as $H = (V_A + V_D + V_I)/V_P$, has traditionally provided such measure. Calling H as a form of "heritability" is oxymoronic because dominance and non-additive types of epistatic effects captured in H are nonheritable from a parent to its offspring.

The concepts of h^2 and H have actually been recognized as being muddled in plants as far back as the 1960s (Hanson 1963). The definition of h^2 is straightforward in animals because an individual animal is both the selection unit and the basis for defining h^2 . In contrast, an individual plant is the selection unit in mass selection but not in other breeding procedures that have been used in plants. Suppose selection is conducted among half-sib families developed from a random-mating population of perennial ryegrass. The full amount of V_A is expressed among individual plants in the random-mating population. But when h^2 is estimated by rearranging the breeder's equation as $h^2 = R/S$, the numerator of the resulting h^2 has $\frac{1}{4}V_A$ instead of $(1)V_A$ because only a quarter of the V_A is expressed among half-sib families. Hanson (1963) concluded that the definition of heritability in plants "*becomes lost in a maze of confusion*" and he raised the need (but was reluctant) to consider an alternative to h^2 and H .

The concept of "reliability" can be borrowed as an alternative measure of the influence of nature versus nurture on phenotypic measurements (Bernardo 2020). Reliability is defined as the consistency of a test or measurement, and reliability has been widely used to gauge the quality of tests in educational, behavioral, and industrial settings (Cronbach 1951). For example, a university admission test is considered reliable if the same student obtains similar scores when taking different editions or versions of the test (assuming that the student's ability remains the same across retakes). There are several ways of measuring reliability, one of which is $V_{Subjects}/(V_{Subjects} + V_e)$, where $V_{Subjects}$ is the variance component due to subjects and V_e is the error variance. In plant breeding, reliability can be defined as the variance component due to the candidates divided by V_P .

Reliability, which we denote by i^2 , is then strictly tied to the selection unit used. If selection is among recombinant inbreds, reliability is $i^2_{RI} = V_{RI}/V_{P(RI)}$, where $V_{P(RI)}$ is the phenotypic variance among recombinant inbreds. If selection is among clones, reliability is $i^2_{Clones} = V_{Clones}/V_{P(Clones)}$. The V_P should be calculated on an entry-mean basis. Suppose that cassava clones are evaluated in e environments with r replications in each environment. The $V_{P(Clones)}$ is calculated as $V_{Clones} + (1/e)(V_{Clones \times Environments}) + (1/re)V_e$, where V_e is the within-environment error variance.

Repeatability has been used as a measure of consistency among repeated measurements (Falconer 1960). While reliability is similar to repeatability, the two are different because repeatability can refer to multiple measurements on the same individuals over space (e.g., measurements on multiple fruits from the same tree) or time (e.g., multiple harvests from the same plants), whereas reliability pertains to multiple measurements (e.g., in different environments) that are independent of each other (Bernardo 2020). In plants, repeatability has been commonly used to measure the consistency among multiple harvests in perennial forage species (Casler et al. 2008; Braz et al. 2015). Repeatability has also been proposed to refer to estimates of heritability when “a nonrandom sample of genotypes is evaluated” (Fehr 1987). Such proposal leads to confusion because it adds a second meaning different from the original concept of repeatability as described by Falconer (1960). The concept of reliability therefore fills a long-time void that neither heritability nor repeatability (Falconer 1960) have filled.

While a high i^2 indicates consistent measurements, information on what constitutes a statistically significant difference would also be helpful for the purposes of selection. For example, breeders might be aware that h^2 is about 0.40 for Fusarium head blight resistance, but they might be surprised to find that this level of h^2 could correspond to a line having 0% infection in one test and 7% in a different test. As a standard practice, it would be helpful to report estimates of both i^2 and the least significant difference, e.g., “Among recombinant inbreds, estimates of reliability and the least significant difference (in parentheses; $P=0.10$) were 0.50 (0.72 Mg ha⁻¹) for yield, 0.70 (0.85%) for protein percentage, and 0.40 (8.5%) for Fusarium head blight incidence.” Relaxed significance levels ($P=0.10$ or 0.20) are recommended, given the relative impacts of a Type I error versus a Type II error in cultivar evaluation (Carmen 1976).

Borrowed: simulation approach to design a breeding program

Plant breeding for major species is like a factory process in which raw materials (germplasm) are input into a manufacturing system (line, hybrid, or clone development) to have products (cultivars) that are marketed and distributed after rigorous testing and quality control (multi-environment trials) (Bernardo 2020). Individual components of a breeding program can be designed quite easily. For example, the parents crossed to form new breeding populations can be selected according to superior performance for multiple traits and SNP diversity between the parents. Estimates of $V_{\text{Clones} \times \text{Environments}}$ and V_e can be used to determine the number of environments and replications per environment needed to detect a given difference for a

quantitative trait as being statistically significant. Selection indices can be constructed to include information on the economic weights of different traits. However, these piecemeal approaches do not consider the entirety of a plant breeding program as a system of interdependent processes. Plant breeders would benefit from the availability of tools to design a plant breeding program as a whole.

Such tools will need to be borrowed from other fields. Operations research involves the use of advanced analytics to make better decisions and can be utilized to design breeding processes such as trait introgression (Cameron et al. 2017). Computer simulation has long been advocated for manufacturing systems (Carrie 1988). Instead of calculating R for individual crosses via the breeder’s equation, genetic gains for a breeding program as a whole can be simulated vis-à-vis the costs and time required. Simulation packages such as QU-GENE (Podlich and Cooper 1998), AlphaSim (Faux et al. 2016), and DeltaGen (Jahufer and Luo 2018) have been developed to model quantitative variation and selection, and such software has been used to compare breeding schemes (Wang et al. 2003; Jahufer and Luo 2018). In the future, simulation tools need to be able to incorporate existing germplasm, molecular marker data, and phenotypic data as input variables. This means that in a wheat breeding program, for example, the genetic entities in a simulation process would not be generic individuals, but instead would reflect the pedigrees of the actual wheat lines used in the program, along with their associated SNP and performance data for multiple traits.

The incorporation of simulation tools in plant breeding would require changes in a typical plant breeding curriculum, as well as increased collaborations with experts in operations research, data science, and simulation of manufacturing systems. Commercial breeding organizations would be apt to hire graduates with expertise in these areas.

BLUE: mixed-model analysis

Mixed-model analysis, which involves best linear unbiased estimation (BLUE) of fixed effects and BLUP of random effects (Henderson 1975, 1985), has proven to be an effective framework for analyzing phenotypic and SNP data routinely generated in a breeding program. Breeders of row crops conduct balanced experiments by evaluating the same set of candidates in each of several locations in one or more years. However, the entirety of phenotypic data across candidates, locations, and years are highly unbalanced. Mixed-model analysis provides two key advantages: it handles unbalanced data, and it can incorporate information from relatives through one or more covariance matrices of random genetic effects.

The linear model underlying mixed-model analysis can be Eq. 1 or an extension thereof. For example, the linear

model can be expanded to include fixed genetic effects due to major QTL (Bernardo 2014b), subpopulations (Yu et al. 2006), transgenes, or different types of cytoplasm, or fixed environmental effects such as level of nitrogen fertilizer or the prior crop grown in the field. The random genetic effects may correspond to lines, clones, or hybrids as in GBLUP, or to SNP markers as in RR-BLUP. Linkage among markers is reflected in the coefficient matrix in RR-BLUP, thereby circumventing the need to assume linkage equilibrium.

The BLUP approach was first used in animal breeding in 1970 to evaluate about 1200 Holstein dairy bulls in an artificial insemination program (Freeman 1991). The use of BLUP in plants started as a trickle, first via shrinkage of estimates of combining ability toward the mean (Melchinger et al. 1987) followed by utilization of information from relatives to predict single-cross performance (Bernardo 1994). It is encouraging that since the 2000s, mixed-model analysis has become familiar to plant breeding students and commonplace in plant breeding programs.

The use of mixed-model analysis decreases the role of genetics and increases the role of statistics in quantitative genetics. This tendency was already observed in animal breeding more than 30 years ago by the statistical geneticist Oscar Kempthorne, who gave the following characterization of animal breeding theory (Kempthorne 1988):

“With the idea that the genes of an individual are random ones of a population of genes with independence between loci, and the idea the environmental effects can be regarded as realizations of independent Gaussian random variables, we see that we have reduced the whole theory to what we in statistics call mixed linear model theory. The outcome is that what is called the theory of animal breeding is reduced to theory of a mixed linear model with fixed effects and independent Gaussian random effects.

So genetics has disappeared from ideation, except that use is made of coefficients of relationship or kinship or “de parenté” (Malécot 1948) or of parentage (Kempthorne 1957).”

The latter paragraph should be modified today because mixed-model analysis now uses SNP marker data in GBLUP rather than pedigree-based coefficients of coancestry via traditional BLUP. The shrinkage factor (λ) used in GBLUP, which is equivalent to $(1 - h^2)/h^2$ in BLUP, can be estimated by a grid search followed by cross-validation to find the λ value that maximizes the predictive ability of the model (de Vlaming and Groenen 2015). An estimate of h^2 is therefore not required and, in this context, identifying

the candidates with the best genotypic values becomes purely a statistics problem. Overall, reinvention of quantitative genetics as described in this *Opinion* article involves approaches that are computational and empirical, rather than being based on classical theory and assumptions. This trend will undoubtedly continue as increasing amounts of phenotypic and SNP data and additional types of data (e.g., climate or cultural management practices) become available to inform breeding decisions.

Acknowledgements Most of the ideas expressed herein were concocted over espressos while I was on a sabbatical leave at the University of Bologna in autumn 2019. I thank Professors Elisabetta Frascaroli, Roberto Tuberosa, and their colleagues for being very gracious hosts, and the Institute for Advanced Studies at the University of Bologna for providing sabbatical funding.

Compliance with ethical standards

Conflict of interest The author declares that he has no conflict of interest.

Publisher's note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this license, visit <http://creativecommons.org/licenses/by/4.0/>.

References

- Anderson JA, Chao S, Liu S (2007) Molecular breeding using a major QTL for Fusarium head blight resistance in wheat. *Crop Sci* 47 (S3): S112–S119
- Beckman JS, Soller M (1983) Restriction fragment length polymorphisms in genetic improvement: methodologies, mapping and costs. *Theor Appl Genet* 67:35–43
- Bernardo R (1994) Prediction of maize single-cross performance using RFLPs and information from related hybrids. *Crop Sci* 34:20–25
- Bernardo R (2014a) Essentials of plant breeding. Stemma Press, Woodbury, Minnesota
- Bernardo R (2014b) Genomewide selection when major genes are known. *Crop Sci* 54:68–75
- Bernardo R (2016) Bandwagons I, too, have known. *Theor Appl Genet* 129:2323–2332
- Bernardo R (2020) Breeding for quantitative traits in plants, 3rd edn. Stemma Press, Woodbury, Minnesota
- Braz TGS, Fonseca DM, Jank L, Cruz CD, Martuscello JA (2015) Repeatability of agronomic traits in *Panicum maximum* (Jacq.) hybrids. *Genet Mol Res* 14:19282–19294

- Brim CA, Johnson HW, Cockerham CC (1959) Multiple selection criteria in soybeans. *Agron J* 51:42–46
- Cameron JN, Han Y, Wang L, Beavis WD (2017) Systematic design for trait introgression projects. *Theor Appl Genet* 130:1993–2004
- Carmer SG (1976) Optimal significance levels for application of the least significant difference in crop performance trials. *Crop Sci* 16:95–99
- Carrie A (1988) Simulation of manufacturing systems. Wiley, New York
- Casler MD, Jung H-J, Coblentz WK (2008) Clonal selection for lignin and etherified ferulates in three perennial grasses. *Crop Sci* 48:424–433
- Cockerham CC (1956) Effect of linkage on the covariances between relatives. *Genetics* 41:138–141
- Cockerham CC (1963) Estimation of genetic variances. In: Hanson WD, Robinson HF (eds) *Statistical genetics and plant breeding*. National Academy of Sciences—National Research Council, Washington, DC, p 53–93
- Comstock RE, Robinson HF (1948) The components of genetic variance in populations of biparental progenies and their use in estimating the average degree of dominance. *Biometrics* 4:254–266
- Comstock RE, Robinson HF (1952) Estimation of the average dominance of genes. In: Gowen JW (ed) *Heterosis*. Iowa State College Press, Ames, p 494–516
- Comstock RE, Robinson HF, Harvey PH (1949) A breeding procedure designed to make maximum use of both general and specific combining ability. *Agron J* 41:360–367
- Cooper M, Messina CD, Podlich D, Totir LR, Baumgarten A, Hausmann NJ et al. (2014) Predicting the future of plant breeding: complementing empirical evaluation with genetic prediction. *Crop Pasture Sci* 65:311–336
- Cronbach LJ (1951) Coefficient alpha and the internal structure of tests. *Psychometrika* 16:297–334
- Dempster AP, Laird NM, Rubin DB (1977) Maximum likelihood from incomplete data via the EM algorithm. *J R Stat Soc B* 39:1–38
- de Vlaming R, Groenen PJF (2015) The current and future use of ridge regression for prediction in quantitative genetics. *Biomed Res Int* 2015:143712. <https://doi.org/10.1155/2015/143712>
- Dudley JW (1984) A method of identifying lines for use in improving parents of a single cross. *Crop Sci* 24:355–357
- Dudley JW, Moll RH (1969) Interpretation and use of estimates of heritability and genetic variances in plant breeding. *Crop Sci* 9:257–262
- Eberhart SA, Russell WA (1966) Stability parameters for comparing varieties. *Crop Sci* 6:36–40
- Efron B (1980) The bootstrap, the jackknife, and other resampling plans. Society for Industrial and Applied Mathematics, Philadelphia
- Ertiro BT, Ogugo V, Worku M, Das B, Olsen M, Labuschagne M et al. (2015) Comparison of kompetitive allele specific PCR (KASP) and genotyping by sequencing (GBS) for quality control analysis in maize. *BMC Genom* 16:908. <https://doi.org/10.1186/s12864-015-2180-2>
- Falconer DS (1960) *Introduction to quantitative genetics*. Oliver and Boyd, London
- Falconer DS (1985) A note on Fisher's 'average effect' and 'average excess'. *Genet Res* 46:337–347
- Faux AM, Gorjanc G, Gaynor RC, Battagin M, Edwards SM, Wilson DL et al. (2016) AlphaSim: software for breeding program simulation. *Plant Genome* 9. <https://doi.org/10.3835/plantgenome2016.02.0013>
- Fehr WR (1987) *Principles of cultivar development: theory and technique*, vol. 1. Macmillan: New York
- Fernando RL (1998) Genetic evaluation and selection using genotypic, phenotypic and pedigree information. In: *Proceedings of the 6th World Congress on Genetics Applied to Livestock Production*, Armidale, Australia, p 329–336
- Finlay KW, Wilkinson GN (1963) The analysis of adaptation in a plant-breeding programme. *Aust J Agric Res* 14:742–754
- Fisher RA (1918) The correlation between relatives on the supposition of Mendelian inheritance. *Trans R Soc Edinb* 52:399–433
- Fisher RA (1941) Average excess and average effect of a gene substitution. *Ann Eugen* 11:53–63
- Freeman AE (1991) C.R. Henderson: contributions to the dairy industry. *J Dairy Sci* 74:4045–4051
- Galton F (1869) *Hereditary genius*. Macmillan and Co., London
- Galton F (1889) *Natural inheritance*. Macmillan and Co., London
- Garcia NS (2008) Mapping QTLs for seed oil, starch, and embryo size in corn using korean high oil germplasm. MSc thesis, University of Minnesota, Saint Paul
- Gardner CO (1963) Estimates of genetic parameters in cross-fertilizing plants and their implications in plant breeding. In: Hanson WD, Robinson HF (eds) *Statistical genetics and plant breeding*. National Academy of Sciences—National Research Council, Washington, DC, p 225–252
- Gardner CO (1977) Quantitative genetics research in plants: past accomplishments and research needs. In: Pollak E, Kempthorne O, Bailey TB (eds) *Proceedings of the International Conference on Quantitative Genetics*, Iowa State University Press, Ames, p 29–37
- Gauch HG (1988) Model selection and validation for yield trials with interaction. *Biometrics* 44:705–715
- Gollob HF (1968) A statistical model which combines features of factor analytic and analysis of variance techniques. *Psychometrika* 33:73–115
- Gray A (2016) Invertastic: large-scale dense matrix inversion. ARCHER whitepaper. <https://www.archer.ac.uk/documentation/white-papers/invertastic/invertasticGray.pdf>
- Habier D, Fernando RL, Dekkers JCM (2007) The impact of genetic relationship information on genome-assisted breeding values. *Genetics* 177:2389–2397
- Hallauer AR, Miranda JB, Fo (1988) *Quantitative genetics in maize breeding*, 2nd edn. Iowa State University Press, Ames
- Hamrick JL, Allard RW (1975) Correlations between quantitative characters and enzyme genotypes in *Avena barbata*. *Evolution* 29:438–442
- Hanson WD (1963) Heritability. In: Hanson WD, Robinson HF (eds) *Statistical genetics and plant breeding*. National Academy of Sciences—National Research Council, Washington, DC, p 125–140
- Harville DA (1977) Maximum likelihood approaches to variance component estimation and to related problems. *J Am Stat Assoc* 72:320–338
- Henderson CR (1975) Best linear unbiased estimation and prediction under a selection model. *Biometrics* 31:423–447
- Henderson CR (1985) Best linear unbiased prediction of nonadditive genetic merits in noninbred populations. *J Anim Sci* 60:111–117
- Hull FH (1945) Recurrent selection for specific combining ability in corn. *J Am Soc Agron* 37:134–145
- Jahufer MZZ, Luo D (2018) DeltaGen: a comprehensive decision support tool for plant breeders. *Crop Sci* 58:1118–1131
- Jenkins MT (1934) Methods of estimating the performance of double crosses in corn. *J Am Soc Agron* 26:199–204
- Jenkins MT (1940) The segregation of genes affecting yield of grain in maize. *J Am Soc Agron* 32:55–63
- Kearsey MJ, Farquhar AGL (1998) QTL analysis in plants; where are we now? *Heredity* 80:137–142
- Kempthorne O (1954) The correlations between relatives in a random mating population. *Proc R Soc Lond (B)* 143:103–113
- Kempthorne O (1957) *An introduction to genetic statistics*. Wiley, New York

- Kempthorne O (1988) An overview of the field of quantitative genetics. In: Weir BS, Eisen EJ, Goodman MM, Namkoong G (eds) Proceedings of the Second International Conference on Quantitative Genetics, Sinauer Associates, Sunderland, Massachusetts, p 47–56
- Lambert RJ, Alexander DE, Han ZJ (1998) A high oil pollinator enhancement of kernel oil and effects on grain yields of maize hybrids. *Agron J* 90:211–215
- Lee M, Sharopova N, Beavis WD, Grant D, Katt M, Blair D et al. (2002) Expanding the genetic map of maize with the intermated B73 × Mo17 (IBM) population. *Plant Mol Biol* 48:453–461
- Lush JL (1937) Animal breeding plans. Iowa State College Press, Ames
- Lynch M (1988) Estimation of relatedness by DNA fingerprinting. *Mol Biol Evol* 5:584–599
- Malécot G (1948) Les mathématiques de l'hérédité. Masson, Paris
- Mather K (1949) Biometrical genetics. Methuen, London
- Matzinger DF (1963) Experimental estimates of genetic parameters and their applications in self-fertilizing plants. In: Hanson WD, Robinson HF (eds) Statistical genetics and plant breeding. National Academy of Sciences—National Research Council, Washington, DC, p 253–279
- Melchinger AE, Geiger HH, Seitz G, Schmidt GA (1987) Optimum prediction of three-way crosses from single crosses in forage maize (*Zea mays* L.). *Theor Appl Genet* 74:339–345
- Meuwissen THE, Hayes BJ, Goddard ME (2001) Prediction of total genetic value using genome-wide dense marker maps. *Genetics* 157:1819–1829
- Mikel MA, Dudley JW (2006) Evolution of North American dent corn from public to proprietary germplasm. *Crop Sci* 46:1193–1205
- Pearson K, Lee A (1903) On the laws of inheritance in man: I. inheritance of physical characters. *Biometrika* 2:357–462
- Pešek J, Baker RJ (1969) Desired improvement in relation to selection indices. *Can J Plant Sci* 49:803–804
- Podlich DW, Cooper M (1998) QU-GENE: a simulation platform for quantitative analysis of genetic models. *Bioinformatics* 14:632–653
- Provine WB (1971) The origins of theoretical population genetics. University of Chicago Press, Chicago
- Robinson HF, Comstock RE, Harvey PH (1949) Estimates of heritability and degree of dominance in corn. *Agron J* 41:353–359
- Schnell FW (1983) Probleme der Elternwahl-Ein Überblick. Arbeitstagung der Arbeitsgemeinschaft der Saatzuchtleiter. Verlag und Druck der Bundesanstalt für alpenländische Landwirtschaft, Gumpenstein, Austria, p 1–11
- Septiningsih EM, Pamplona AM, Sanchez DL, Neeraja CN, Vergara GV, Heuer S et al. (2009) Development of submergence-tolerant rice cultivars: the *Sub1* locus and beyond. *Ann Bot* 103:151–160
- Sprague GF, Federer WT (1951) A comparison of variance components in corn yield trials: II. error, year × variety, location × variety, and variety components. *Agron J* 11:535–541
- Sprague GF, Tatum LA (1942) General vs. specific combining ability in single crosses of corn. *J Am Soc Agron* 34:923–932
- Stuber CW, Moll RH (1972) Frequency changes of isozyme alleles in a selection experiment for grain yield in maize (*Zea mays* L.). *Crop Sci* 12:337–340
- Sughrue JR, Hallauer AR (1997) Analysis of the diallel mating design for maize inbred lines. *Crop Sci* 37:400–405
- Syvänen AC (2005) Toward genome-wide SNP genotyping. *Nat Genet* 37:s5–s10
- Troyer AF (1999) Background of U.S. hybrid corn. *Crop Sci* 39:601–626
- VanRaden PM (2008) Efficient methods to compute genomic predictions. *J Dairy Sci* 91:4414–4423
- Wang J, van Ginkel M, Podlich D, Ye G, Trethowan R, Pfeiffer W et al. (2003) Comparison of two breeding strategies by computer simulation. *Crop Sci* 43:1764–1773
- Yu J, Pressoir G, Briggs WH, Bi IV, Yamasaki M, Doebley JF et al. (2006) A unified mixed-model method for association mapping that accounts for multiple levels of relatedness. *Nat Genet* 38:203–208