

30

Selection and Development of Pure Lines

Draft Version 27 May 2013

When dealing with autogamous species, plant breeders have used two broad approaches towards developing new strains. For most species (cereals, soybeans, and many vegetable crops), breeders have been concerned with the creation of pure line **cultivars** that are adapted to a specific set of environments and optimized for yield and other desirable characteristics. An important exception is maize, where (for both biological and historical reasons) the goal of breeders has generally been the construction of collections of pure lines whose hybrid seed is the commercial product. Chapter 31 examines selection under such **cross-breeding** schemes.

There are two broad issues when considering selection involving pure lines: choosing the best lines from a large collection and the optimal construction of such lines following the initial hybridization, and we consider these in turn. There is an enormous plant breeding literature dealing with these concerns and we will only introduce the basic issues and give a general overview of the methods and problems. For further treatment, see Allard (1960), Williams (1964), Simmonds (1979), Hallauer and Miranda (1981), Mayo (1987), and in particular Wricke and Weber (1986) and Stoskopf et al. (1993).

Suppose we have a collection of pure-breeding lines, either because each is completely homozygous or each is asexually (clonally) propagated, and we wish to select the best fraction p . Under this setting, offspring are genetically identical to their parents (ignoring mutation), and between-line selection has the advantage that elite genotypes are propagated intact across generations. Thus, it would appear that between-line selection is very straightforward. This is far from the truth, as there are a number of very subtle, yet critical, issues that needed to be considered. First, exact propagation of chosen genotypes actually has a key disadvantage relative to selection in segregating populations. If we choose inferior lines, we cannot recover more elite genotypes unless we cross these back to other lines. Hence, attempting to choose the best pure line (or lines) carries with it the risk that we will not improve the population or (more likely) not pick the best genotypes present in the collection of lines. The second broad issue is the actual formation of the lines, which entails two components: when to apply selection when advancing the collection of lines derived from the hybrid cross towards complete inbreeding and the choice of parents to start the lines.

We start by examining issues of selecting the best line from among a collection of lines in this section. The various approaches for selection during segregating lines (the **advancement of generations**) are discussed next, followed by issues in the choice of parents.

RESPONSE TO BETWEEN-LINE SELECTION: BASIC ISSUES

Recall from the generalized breeders' Equation 13.10) that the response in the offspring y from selection on parents x is $R_y = \bar{r}_x \sigma(x, y) / \sigma_x$. Since we are using pure-breeding lines, the parent and offspring have the same genotype, and $\sigma(x, y) = \sigma(\bar{p}_i, g_i) = \sigma_G^2$, the correlation between the mean phenotypic value of the line $\bar{p}$ and its genotypic value g . Here $\sigma(\bar{p}_i, g_i) = \sigma_G^2$, the

between-line genetic variance. As mentioned above, lines are scored by their mean values over a particular experimental design. If the sampling variance of the phenotypic mean of a line about its genotypic value is σ_ϵ^2 , then $\sigma_x^2 = \sigma^2(\bar{p}) = \sigma_G^2 + \sigma_\epsilon^2$, and the response becomes

$$R = \frac{\sigma_G^2}{\sigma(\bar{p})} \bar{i} = \frac{\sigma_G^2}{\sqrt{\sigma_G^2 + \sigma_\epsilon^2}} \bar{i} \quad (30.1)$$

As we discuss below, considerable control over σ_ϵ^2 can be exercised by the choice of experimental design (see Equations 30.3 and 30.4).

Multistage Selection: Finney's Rule

Lines are often measured over two (or more) growing seasons before selection, with the mean value of the line over seasons used for final selection. This approach partly accounts for genotype $\times$ year interactions, which can have a considerable effect of σ_ϵ^2 (see Example 30.2 below). Equation 30.1 gives the expected response for such a single cycle of selection. Alternatively, a breeder may choose to use several cycles, rather than a single cycle, of selection to choose the elite lines. The logic behind such **multistage selection** is that the initial cycle can cull out many of the inferior lines, allowing for greater replications (and hence higher precision) to discriminate among the remaining lines. How is multistage selection best carried out and what is the resulting response? We address the first question here, leaving the expected response until Chapter 38, as it requires some details from the theory of index selection.

The question to be answered for the optimal design is as follows: starting with g genotypes, what is the optimal way to select for a fraction α of these if we have the resources to measure a total of N plots over a total of s stages (cycles) of selection? Note that if α_i is the fraction selected in stage i ,

$$\alpha = \prod_{i=1}^s \alpha_i \quad (30.2a)$$

The optimal design reduces to the optimal choice of the α_i , and a simple approximate solution was suggested by Finney (1958): Select the same fraction each generation, with

$$\alpha_i = \alpha^{1/s} \quad (30.2b)$$

The amount of replication in stage i is set by N/g_i , the total number of plots available in that cycle (N) divided by the number of remaining genotypes (g_i).

Example 30.1. Suppose we have funds allocated to score a total of 1500 plots over three generations, and we wish to select the best 4 genotypes from an initial sample of 500. Here $\alpha = 4/500 = 0.008$, and Finney's rule gives:

$$\alpha_i = (0.008)^{1/3} = 0.2 \quad \text{and} \quad N_i = \frac{1500}{3} = 500$$

Thus in the first stage, each genotype is replicated once, and the 100 best ($500 \cdot 0.2$) are chosen. In the second stage, each genotype is now replicated 5 times ($500/100$), and the best 20 ($100 \cdot 0.2$) are chosen. Finally, in the third stage of selection, each genotype is replicated 25 times, with the best four chosen. Summarizing:

Stage	Genotypes	Replications/Genotype	α_i	No. of Selected Genotypes
1	$g_0 = 500$	$N/g_0 = 1$	0.2	$\alpha_1 g_0 = 100$
2	$g_1 = 100$	$N/g_1 = 5$	0.2	$\alpha_2 g_1 = 20$
3	$g_2 = 20$	$N/g_2 = 25$	0.2	$\alpha_3 g_2 = 4$

As mentioned, the exact solution for the expected response under multistage selection requires concepts from multivariate selection and index selection, and we defer further discussion until Chapter 38.

CONTROLLING THE RESIDUAL VARIANCE

A major advantage of either pure or clonal lines is that one can replicate identical genotypes, replacing individual measurements with group measurements, such as the yield in a plot. Further, individuals can be replicated over environments, such as different locations and/or different seasons/years. The use of replication at several levels (multiple individuals within a plot, plots within a major environment, plots or blocks of plots across major environments) reduces the error variance σ_e^2 in estimating the genotypic value of a line from its phenotypic mean. Reducing σ_e^2 increases the response to selection (Equation 30.1), and increases the probability of choosing the best genotypes in the sample of lines being tested. A second route for reducing the error variance is the use of appropriate plot designs within a given major environment to more fully account for the variance generated by both within- and between-plot microenvironmental values. The final route is by using appropriate statistical methodology to account for as much of the structure in genotype $\times$ environment interactions as possible. Using an appropriate model is often akin to a factor of 2-4 fold extra replications. We examine these three topics in turn.

Replication

Akin to our treatment in Chapter <check 8> with replication we can decompose the error variance σ_e^2 on the performance measurement of the genotype over the entire experimental design into variance components associated with different sources of error. By suitable allocation within a particular design, we can reduce the error variance. There is an extensive literature on this subject and we only cover some of the key areas here. Reviews of general agricultural experimental design can be found in Little and Hillis (1978), Pearce (1983), Gomez and Gomez (1984), Pearce et al. (1988), and Petersen (1994).

Since lines are typically scored as plot averages, rather than individual values, the sampling variance is often expressed in terms of the **plot variance** $\sigma_{w(p)}^2$, the variance among line members within a plot. If σ_e^2 is the within-plot variance for an individual plant about its line's genotypic value in that plot, then the variance for a plot scoring n individuals is $\sigma_{w(p)}^2 = \sigma_e^2/n$. If we replicate members of a pure line over r plots (of size n) in each of e distinct environments, then

$$\sigma_e^2 = \frac{\sigma_{G \times E}^2}{e} + \frac{\sigma_p^2}{er} + \frac{\sigma_e^2}{er n} = \frac{\sigma_{G \times E}^2}{e} + \frac{\sigma_p^2 + \sigma_{w(p)}^2}{er} \quad (30.3)$$

where $\sigma_{G \times E}^2$ is the line $\times$ environment interaction (genotype $\times$ environment) variance and σ_p^2 the between-plot environmental variance. This error variance decomposition allows us to intelligently design experiments to minimize the error variance for a particular set of

constraints. For example, if one has a total of $N = re$ plots per line, σ_ϵ^2 is minimized by using one replicate per environment, as the weighting on $\sigma_p^2 + \sigma_{w(p)}^2$ remains unchanged, while that on $\sigma_{G \times E}^2$ decreases with increasing e . Increasing the number n of plants per plot decreases the plot variance $\sigma_{w(p)}^2$ but does not reduce the effects of either the between-plot or line $\times$ environment variances. As mentioned in Chapter < **check 8** >, the line $\times$ environment is often further decomposed into location (ℓ), year (y), and location $\times$ year effects. If y years and ℓ locations are used,

$$\sigma_\epsilon^2 = \frac{\sigma_{G \times \ell}^2}{\ell} + \frac{\sigma_{G \times y}^2}{y} + \frac{\sigma_{G \times \ell \times y}^2}{y\ell} + \frac{\sigma_p^2 + \sigma_{w(p)}^2}{y\ell r} \quad (30.4)$$

The genotype $\times$ year interaction is a measure of **yield dependability** (the yield stability of a particular genotype over time), while $G \times \ell$ is a measure of its **wide adaptability** (or lack thereof). These useful descriptors of crop performance only make sense in the context of low to moderate $G \times \ell \times y$ interactions.

It is often observed that the year effects are much larger than the location effects, and that $G \times \ell \times y$ interactions are larger still (e.g., Rasmusson and Lambert 1961, Patterson et al. 1977, Talbot 1984). The relative impacts from interactions involving locations can be reduced by increasing the number of locations, which decreases the weighting on both $\sigma_{G \times \ell}^2$ and $\sigma_{G \times \ell \times y}^2$. Decreasing the impact from genotype by year interaction is more difficult, as this requires increasing the number of years a trial is run. Such an increase, however, is often quite advantageous. Talbot (1984), in a summary of 100 trials of over a dozens crops in the UK, found that two years of trials at six locations gives greater precision than a single year trail over 12 locations.

Example 30.2. Patterson et al. (1977) estimated the components of error variation associated with a large number of barley trails in the UK as follows:

	$\sigma_{G \times \ell}^2$	$\sigma_{G \times y}^2$	$\sigma_{G \times \ell \times y}^2$	$\sigma_p^2 + \sigma_{w(p)}^2$
Actual	0.0084	0.0322	0.0561	0.1101
Percentage	4.06	15.57	27.13	53.24

A 90-plot design with $\ell = 10$ locations, $y = 3$ years and $r = 3$ replications has an error variance of

$$\sigma_\epsilon^2 = \frac{0.0084}{10} + \frac{0.0322}{3} + \frac{0.0561}{30} + \frac{0.1101}{90} = 0.0147$$

and relative contribution from these sources of error on the total error variance σ_ϵ^2 becomes

	$G \times \ell$	$G \times y$	$G \times \ell \times y$	Plot variance
Percentage	5.73	73.18	12.75	8.34

Because of the few associated degrees of freedom, the effect of $G \times y$ is greatly magnified, and in this case accounts for almost 3/4 of the total error variance. If the number of years tested is increased to 4, and the number of locations reduced to 7 (giving a total of 84 plots), the error variance is reduced to 0.0126, 85.7% of the original error variance. With 5 years and 6 locations, and 6 years and 5 locations (both having 90 total plots, the variance is reduced to 0.0109 and 0.0101 (respectively) or 74% and 69% of the error variance under the original design.

Field Plot Designs

One source of error over which the investigator has considerable control is the plot variance (both within- and between-plots). Excellent reviews of plot design theory can be found in LeClerc et al. (1962) and LeClerc (1966). While appropriate field preparation can reduce some of this variation (for example by the use of mechanical seeding devices to ensure uniformity of planting or laser-guided plowing for a field of uniform height), other sources of variation, especially soil heterogeneity are generally beyond the direct control of the experimenter. The larger the plot size, the greater than chance that individuals within the plot experience significant soil heterogeneity. Historically, the trend has been a move to smaller and smaller plots. In the early 1900's, plot sizes for field crops in the United States ranged from 2 acres to 1/40 of an acre, with an average size around 1/10 of an acre (Taylor 1908). The work of Mercer and Hall (1911) lead to usage of much smaller plots (1/40 of an acre), while Woods and Stratton (1910) strongly prompted replication of smaller plots in place of a single measure on a larger plot. The **rod-row** plot, which involves rows of rod length (roughly 16 feet or 5 meters) as the basic unit of replication, appears frequently in the plant breeding literature. Rod-row and field plots show generally good agreement (LeClerc 1966). Even smaller plots can be used, such as **hill plots** (Bonnett and Bever 1947) which involve only a few seeds (typically 2-3 to 30-50, depending on the species being examined). A literature review by Weber (1984) finds a good correlation between yield estimating using hill and row plots. While the results in the literature are somewhat mixed, the majority point to an advantage of oblong plots over square plots (LeClerc 1966), consistent with the prediction (on theoretical grounds) of Dendrinis (1931) that under patchy soil heterogeneity, oblong plots generally have lower variation than square plots.

B	D	C	D	C
A	E	B	C	B
D	B	A	E	D
E	A	C	B	A
C	E	A	D	E

Completely Randomized Design

D	A	B	E	C
A	E	C	B	D
B	C	A	D	E
E	B	D	C	A
C	D	E	A	B

Latin Square

C	B	C	E	B
B	A	D	C	A
A	D	E	A	D
E	E	A	B	C
D	C	B	D	E

Blocks = Columns

D	B	A	E	C
C	A	B	D	E
A	D	C	E	B
C	A	B	D	E
B	D	C	A	E

Block = Rows

Randomized Blocks

Figure 30.1. Experimental designs for treating a series of lines (here A through E) replicated in a single location. Under a **completely randomized design**, lines are assigned at random over the entire grid. Under a **randomized block design**, the growing area is first partitioned into a series of blocks, and the position of lines within each block is assigned at random. Examples of using either rows or columns as blocks are given. Removal of the block effect decreases the residual error variance. Under a **Latin square design**, each line appears only once in each row and column, allowing for row and column plot effects to be estimated, further reducing the residual error variance.

When plots are replicated within a given macroenvironment, a number of experimental designs can be used to account for plot-to-plot variation. A few standard designs are given in Figure 30.1. Suppose the lines to be tested in a particular location are replicated as a series of plots over a grid. While one can plant individuals in a **completely randomized design** (where plots are assigned at random over the entire grid), this design typically results in higher variances associated with the line means over the entire replicate as plot-to-plot variation is not accounted for, and thus appears in the error variance for the line means. The use of a **randomized block design**, wherein the growing area is first assigned into **blocks** and then lines are randomized within blocks, partly accommodates between-plot variation by inclusion of a block effect (in a **complete block design**, each line is tested in each block.)

Removal of some of the between-plot variance decreases the residual error variance. Figure 30.1 shows examples of using rows as the blocks (and hence randomizing lines within each row) or conversely using columns as blocks. Other shapes can also be used for blocks as well. Under the **Latin square** design, each line appears only once in each row and column, allowing a further partition of plot variation into row and column effects (also giving rise to the popular game of Sudoku). Latin squares are more efficient, with complete block design being only 65-75% as efficient (Yates 1935, Garner and Well 1939, Ma and Harrington 1948). Since the number of replicates under a Latin square must equal the number of lines (or more generally be an integer multiple of the number of lines), this is not an efficient design when the number of lines exceed six to eight.

When a large number of lines are being tested, **incomplete block designs** (where not all lines appear in all blocks) are typically used, such as various **lattice designs**. The major complication with lattice designs is that they impose certain constraints on the number of lines that can be tested. For example, for a **simple lattice** design there must be $\sqrt{n}$ plots per block, implying only perfect square number of lines can be tested (i.e., 25, 36, 49, 64, 81, etc.), and the number of replicates must be a multiple of two. For a **triple lattice** design, again the number of tested lines must be a perfect square, but now the number of replicates are multiples of three. Details can be found in LeClerg et al (1962), Cochran and Cox (1957), or any standard experimental design text. In a review of over 80 lattice experiments, Ma and Harrington (1948) found that the double and triple lattice designs had relative efficiencies of 128% and 160% (respectively) over the comparable complete block designs.

Finally, a more recent approach to error control in field plot design is the use of spatial statistics to analyze neighboring plots in an attempt to control for local soil variation. Reviews can be found in Bartlett (1978), Wilkinson et al (1983), and Besag and Kempton (1986).

Accounting for $G \times E$

Another major source of error variation is genotype $\times$ environment interaction, which can occur in both obvious (across locations and years), and more subtle, forms. Examples of the later are differences in the seed density used, for example **space planting** to allow individual plants to be scored is rather different from using normal seed density. Likewise, a plot consisting of a single genotype has a very different competitive environment that a plot consisting of a mixture of genotypes. When multiple lines are being simultaneously tested, **guard rows** between strains are often included. While these are discarded from the analysis, they serve as a barrier to between-line competition which can alter yield and other traits. Likewise, a plot often consists of 3-5 rows, with the **border rows** also being excluded from the analysis.

On a larger scale, major differences in the environment can occur between locations and between years (or seasons) at the same locations. The final testing of lines generally involves **multiregional yield trials**, testing lines over multiple locations (often different countries), and over several years. Clearly, in such tests there is the strong potential for $G \times E$ interactions. If y_{ijk} denotes the average for the k th replicate of line i grown in environment j , then the

standard model for $G \times E$ is $y_{ijk} = \mu + G_i + E_j + I_{ij} + e_{ijk}$, where I_{ij} represents the interaction term (this is also denoted by GE_{ij}). Obviously, the more accurately we can estimate I_{ij} , the more precise our estimates of the average value of a line across environments becomes. A variety of linear model and regression-based approaches for approximating I_{ij} have been proposed (most of which are discussed in more detail in LW Chapter 22). The simplest is the **treatments model**, where we estimate all the I_{ij} 's individually by ANOVA. Assuming g genotypes and e environments (such as specific location-year combinations), the treatment model requires $g \cdot e - 1$ degrees of freedom. Given the large number of parameters to estimate, the individual accuracy suffers. A reduction in the number of parameters to estimate occurs if we use a regression estimator for I_{ij} , such as the Finlay-Wilkinson (or joint) regression of interaction on mean environment value, $I_{ij} = \beta_i E_j + \epsilon_{ij}$. In its most general form, this regression can be written as

$$I_{ij} = \beta_i E_j + \alpha_j G_i + \kappa G_i E_j + \epsilon_{ij}$$

In theory, one (or more) of these regression-based estimators could be used to fit data from a series of trials, and the model showing the smallest residual variance used for further analysis.

A different, and very promising, alternative to either the straight ANOVA or the regression estimators is the **AMMI**, or **Additive Main effects, Multiplicative Interactions**, model. In the literature this is also referred to as the **biplot model** (Gabriel 1971) and (more rarely) as **Factor Analysis of Variation**, or **FANOVA**, (Gollob 1968). The basic foundations are due to Gollob (1968), Mandel (1971), Gabriel (1971), Bradu and Gabriel (1978), and Kempton (1984). Its application to yield-testing is largely due to Gauch (1988), Zobel et al. (1988) and Gauch and Zobel (1989). Additional reviews of AMMI can be found in Gauch (1992), Crossa (1990), and Cooper and DeLacy (1994). Gauch and Zobel (1990) used the Expectation-Maximization (EM) algorithm (LW Appendix 4) to allow for missing data, referring to this as **EM-AMMI**. Chapter 43 examines these, and other powerful methods, for dealing with $G \times E$, so we only offer a few brief comments here.

Under AMMI, the main effects (G , E) are assumed to be additive, while a multiplicate approach is used to estimate I_{ij} . Suppose there are g genotypes replicated r times each over e environments. Under a standard ANOVA, only the r replicates for each particular $G \times E$ are used to estimate the interaction. Under AMMI, all the data ($g \cdot e \cdot r$ points) are used to estimate each interaction, as information from all other genotypes and environments influences the estimate of each particular interaction. To apply AMMI, the main effects are first estimated from a standard additive (i.e., no interaction) ANOVA. The residual is then fitted using a principal components analysis (Appendix 5), which attempts to find the series of orthogonal axis of a covariance matrix that explain the most variation. In particular, for the AMMI model, n terms are used to predict the interaction term,

$$I_{ij} = \sum_{k=1}^n \lambda_k \gamma_{Gk} \delta_{Ek} + \epsilon_{ij} \quad (30.5)$$

where λ_k is the eigenvalue associated with the k -th axis from a principal component analysis (PCA), and γ_{Gk} and δ_{Ek} are the genotypic and phenotypic scores on PC axis k (these are related to the eigenvectors associated with λ_k). The first PC axis accounts for the largest variation, the second the largest of the remaining variation, and so on. Both the additive ANOVA and joint regression approaches are special cases of AMMI, specifically AMMI0 and AMMI1 (Gollob 1968, Mandel 1971, Gabriel 1978, Bradu and Gabriel 1978, Kempton 1984, Zobel et al 1988).

A resolution on deciding of the order of the AMMI model (the number n of PCAs chosen) was offered by Gauch (1988) and Gauch and Zobel (1988). Two approaches can be used assess

model fit, **predictive** and **postdictive**. Under a predictive approach, one randomly partitions the data into a model fitting part and a model evaluation (or **cross validation**) part. A series of alternative models (AMMI of different order) are then fit and assessed by their ability to predict the model evaluation data. Under a **postdictive** approach, the model giving the best fit to the overall data is chosen, typically by adding additional components and using F-tests to judge whether the addition significantly improves model fit. As Gauch and Zobel point out, the postdictive fitting approach tends to overfit the data, capturing both model signal as well as any structure in the noise. Thus, they recommend a predictive approach to fitting AMMI. Gollub (1968), Gauch (1988) and Zobel et al. (1988) assign the appropriate degrees of freedom associated with the k th term in Equation 30.5 as

$$df_k = g + e - 1 - 2k \quad (30.6)$$

An important caveat to cross validation schemes was made by Piepho (1994). Under a completely randomized design, individual observations can be assigned to either the model fitting and cross validation data sets. However, if the data are in a randomized complete block design, care must be taken not to split observations from blocks into the two data sets. If observations from blocks are kept intact when assigning observations to the model-testing and model-fitting data sets, any block effects are simply added to the additive environmental effect. When observations from blocks are split, error is added to the interaction term, so that AMMI is examining $b_j(i) + I_{ij} + e_{ij}$ instead of $I_{ij} + e_{ij}$, where $b_j(i)$ is the mean block effect for the i th genotype in the j environment.

Cornelius et al (1992) and Cornelius (1993), motivated by the work of Goodman and Haberman (1990), developed an alternative approach to fitting the correct number of PC axes. Their F_{GH1} and F_{GH2} tests (based on the Wishart distribution, see their papers for details) tend to include more PC axes than the cross validation approach, which tends to underfit the model. They suggest if a cross validation scheme is used, that only one replicate for each $G \times E$ treatment be maintained for the model-testing data set, and the rest used in the model-fitting data set. They also show that the postdictive approach based on the F value significance of the PCs (which they refer to as the **Gollob test**) has a very high type I error and should be avoided.

Example 30.3. A series of regional soybean yield trials in upstate New York were analyzed by Zobel et al. (1988) and Gauch and Zobel (1988). The data examined involved 7 genotypes grown in 40 environments (certain combinations of 9 years and 9 sites), most with 4 replications per trial (average 3.73 replications). The resulting ANOVA for the treatments model is as follows:

Source	df	SS	MS	F
Total	1043	245,722,327		
Treatment	279	165,694,661	593,888	5.67
Environment	39	129,469,970	3,319,743	31.69
Genotype	6	9,722,960	1,620,493	15.47
$G \times E$	234	26,501,731	113,255	1.08
PC 1	44	18,632,502	423,446	4.04
Residual	190	7,869,229	41,417	0.40
Error	764	80,027,666	104,748	

The treatment consists of three sources: additive genotypic effects, additive environment effects, and $G \times E$ interactions. One immediately troubling feature is that the F statistic for the interaction term is not significant, yet the interaction sum of squares (SS) is almost three times the genotype SS. However, due to the large number of degrees of freedom, the

resulting $G \times E$ mean square (MS) is over a magnitude smaller than the genotype MS, which is highly significant. However, when the first principal component (PC 1) is extracted from the interaction terms, it accounts for 70% of the interaction SS while using only 44 df , and has a significant F statistic. The principal analysis in this case was very effective in capturing structure in the interaction terms.

Gauch (1988) found that a first-order (AMMI1) model provided the best predictive fit, while the first three PCs were significant in a postdictive analysis. Gauch and Zobel (1988) found that the AMMI1 model with two replicates was equal in prediction accuracy to a standard ANOVA with interactions with five replicates. Hence, simply by using an AMMI analysis resulted in a huge improvement, and a resulting decrease in the overall error variance.

As the above example points out, simply by using an AMMI analysis, one (for this case) receives an advantage equivalent to increasing the number of replicates by a factor of 2.5. In an analysis of international maize trials, Crossa et al. (1990) found that AMMI results in an increase of 4.3 in one trial and 2.6 in the other, and in an international wheat yield trial it had an increase of 3.7 (Nachit et al. 1992, Crossa et al. 1991). Thus, simply by using an alternative statistical analysis of already collected data, one can reduce the error variance considerably. This not only results in an increase in selection accuracy (Gauch and Zobel 1989, Crossa et al. 1990), but also has important implications for the efficiency of experimental designs to find the best line. Using the New York regional soybean trial data, Gauch and Zobel note that for a hypothetical multistage selection of the best line from 1,000 candidates, AMMI analysis reduces the required number of plots from 2,440 to 1,400 and the time from three years to two years.

Finally, mixed model approaches (such as BLUP) have also been suggested for treating $G \times E$, and again are more fully discussed in Chapter 43. Building on the work of Hill and Rosenberger (1985) and Stroup and Multize (1991) who examined BLUP for estimating the additive effects for yield trials, Piepho (1994) considered BLUP estimates of interaction ($G \times E$) effects as well. Piepho's approach is to treat the genotypes as random (as opposed to fixed) effects, so that $G \sim N(0, \sigma_G^2)$, $I \sim N(0, \sigma_{GE}^2)$, and $e \sim N(0, \sigma_e^2)$. Environments are treated as fixed effects. In this case, Piepho shows that the BLUP estimate for $\mu + G_i + I_{ij}$, the mean of genotype i in environment j , less the main effect of that environment, is

$$\text{BLUP}(\mu + G_i + I_{ij}) = h_g^2(\bar{z}_{.i} - \bar{z}_{..}) + h_{ge}^2(z_{ij} - \bar{z}_{.i} - \bar{z}_{j.} + \bar{z}_{..}) \quad (30.7a)$$

where for n_e environments,

$$h_g^2 = \frac{\sigma_{GE}^2 + n_e \sigma_G^2}{\sigma_{GE}^2 + \sigma_e^2 + n_e \sigma_G^2} \quad (30.7b)$$

is the repeatability of estimated genetic effects, and

$$h_{ge}^2 = \frac{\sigma_{GE}^2}{\sigma_G^2 + \sigma_e^2} \quad (30.7c)$$

is the repeatability of estimation interaction. An analysis of a German faba bean registration data showed that BLUP method outperformed AMMI in four of the five data sets.

CHOOSING THE BEST GENOTYPES

Selecting for the single best line from a collection of pure lines brings with it the risk that the best genotype is actually not the one chosen. We examine several issues here related to this

risk. First, what statistical method should be used to pick the “best” lines from a collection of lines. Second, we review the types of statistical errors that can be made. For choosing the best line, our major concern is controlling type III errors (picking inferior lines). We conclude by examining the tradeoff between adding new lines versus increasing the number of replicates per line.

Least Significant Difference (LSD) Tests

A problem that arises in selection among pure lines is determining which lines, from a potentially large number of comparisons, are significantly different from each other. A number of multiple comparison tests, such as Fisher’s least significant difference (LSD), Tukey’s significant difference, and Duncan’s multiple range (to name a few), have been proposed to deal with this issue. Simulation studies by Carmer and Swanson (1971, 1973) find that a simple modification of Fisher’s LSD test is at least as good as any other tested procedure and is the simplest to implement. Fisher’s original LSD test compares the absolute value of the difference for the pair of means under consideration with the test statistic

$$\text{LSD} = t_{\alpha, df} s_d$$

where $t_{\alpha, df}$ is the student’s t distribution cutoff for the (two-sided) significance level α chosen and df is the degrees of freedom associated with the standard error s_d of the difference between the two means being compared. The concern is that Fisher’s original test does not give an overall (experiment-wise) test at level α due to multiple comparisons being made. To correct for this, a preliminary F test for a significant treatment effect (based on the treatment mean squares divided by the error mean squares) is first performed. If this is not significant, there is no need at looking at all the pairwise comparisons. If it is significant, then pairs are tested using LSD as defined above.

Type I, II, and III Errors

Recall that a type I error occurs when we declare one line to be significantly better than another when in fact both are equivalent, and a type II error occurs when the lines are in fact different, but we lack the statistical power to distinguish between them (i.e. failure to reject the hypothesis that they are equal). Carmer (1973) notes that type I errors have no economic consequences (as the lines are identical), while some type II errors result in inferior lines being chosen. If two lines are not statistically significantly different from each other (but their true values are indeed different), then if we simply chose one at random, half the time we will chose the inferior line. If we instead simply chose the line with the largest value, the inferior lines are chosen less than half the time when we make a type II error. A **type III error** is judged to be the most serious, where an inferior line is chosen over a superior line. A type II error has an indirect component (choosing an inferior line when a type II error occurs) and a direct component (the inferior line is found to be significantly better than the superior line). Carmer (1973) argues that type III errors are really the ones of most concern, a point reiterated by Johnson et al. (1992) who suggests that fear of type III errors are a major barrier to the more widespread use of improved strains by farmers in developing countries. The probability γ of a type III error has two components: an indirect component (which Johnson et al. refer to as **implied** γ) from those type II errors where we chose the less superior line, and a direct component (referred to as **pure** γ) when the less superior line is judged significantly better than a superior line. If δ is the absolute value of the true difference in means, and σ_d is the standard error, then the (large-sample) probability β of a type II error is

$$\beta = \Pr \left(-[t_{\alpha, df} + \delta/\sigma_d] < t < t_{\alpha, df} - \delta/\sigma_d \right) \quad (30.8a)$$

where t is a student's t random variable with df degrees of freedom. If a line is chosen at random when there is no statistical difference between them, the implied γ is $\beta/2$. The "pure" γ is

$$\Pr\left(t > t_{\alpha, df} + \delta/\sigma_d\right) \quad (30.8b)$$

The type III error is thus

$$\gamma = \frac{1}{2} \Pr\left(-[t_{\alpha, df} + \delta/\sigma_d] < t < t_{\alpha, df} - \delta/\sigma_d\right) + \Pr\left(t > t_{\alpha, df} + \delta/\sigma_d\right) \quad (30.8c)$$

For control of Type II errors, Carmer (1976) suggests using LSD with a significance (type I error) of $\alpha = 0.2$ to 0.4 . While this will generate a larger number of type I errors, as mentioned it is really type III errors that are generally more critical to control in the choice of a best line.

Example 30.4. Suppose $\delta/\sigma_d = 1$ and we assume that the number of degrees of freedom is sufficiently large that we can approximate a t distribution with a unit normal z . In this case, for a LSD test with a type I error of $\alpha = 0.05$, the matching normal cutoff for a two sided test is $z = 1.96$. The resulting type II error is $\Pr(-2.96 < z < 0.96) = 0.830$, while the "pure" type III error is $\Pr(z > 2.960) = 0.002$, giving a total type III error of $(1/2) 0.830 + 0.002 = 0.417$. Of this, 99.6% is from type II errors. In a similar fashion, for α values of 0.1, 0.2, and 0.4, the total type III error is 37.2%, 31.3%, and 20.8% (respectively), which is still largely from Type II errors (98.8%, 96.4%, and 86%).

Probability of Choosing the Best Genotypes

Since lines are scored based on their phenotypes, rather than genotypes, sampling variance can result in a less than optimal line still having the largest phenotypic value. There is a tradeoff between replications per line and number of lines tested. The greater the number of replications, the more precise the estimate of the genotypic mean, reducing the risk of choosing a genotypically-inferior line. However, this also reduces the number of lines to be tested, and hence the selection intensity. We examine the optimal design given this tradeoff in the next section. A related issue we examine first is the important point stressed by Gauch and Zobel (1989, 1996) that adding new genotypes is not all good. Suppose there is a two percent chance of adding a superior genotype. An increase in 50 lines will (roughly) introduce one superior line, which is good. It will also introduce (on average) 49 inferior genotypes, which is bad, as these contribute noise from which the signal (the superior genotype) needs to be extracted.

To see how the addition of interior genotypes causes problems for selection, suppose we have 30 lines in a yield trial, all with true yields of (say) 1000 and the experimental design is such that the standard error for any particular mean from its true value is 150. The expected maximum value and range can be obtained from standard order statistics theory. Guach and Zobel (1989) give the following expected maximum value and expected range of values for N draws from a unit normal:

N	2	4	10	30	100	500	3000	20,000
E[Max]	0.56	1.03	1.53	2.04	2.51	3.04	3.54	4.02
E[Range]	1.12	2.06	3.06	4.08	5.02	6.08	7.08	8.04

Hence, we expect that the largest observed mean for these 30 identical lines has an expected value of $1000 + 2.04 \cdot 150 = 1306$, while the expected range is 612, from 694 to 1306. In such a setting, it would be difficult to select for a genotype with a true mean of (say) 1250, which is a 25% increase over the inferior lines. Indeed, we would choose the superior line in this case only 37.1% of the time.

This particular setting, where one line is superior and the other N lines are inferior by the same amount R (which we will scale in standard deviations), is called the **generalized least favorable** (or GLF) configuration, and represents the most difficult selection problem (Gibbons et al. 1977). The probability of selecting the superior line under the GLF configuration is a function of R and N , and representative values are plotted in Figure 30.2. In the example above, $R = 250/150 = 1.67$, and the probability of success (37.1%) was obtained from extrapolating from the table given in Gauch and Zobel (1989).

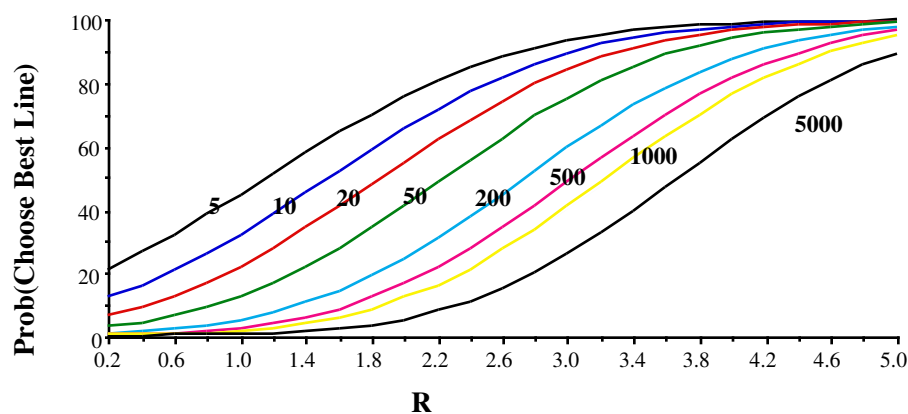


Figure 30.2. Probability of picking the best line under the GLF model (one superior line, N equally inferior lines). The difference between the superior and inferior lines, measured in standard deviations, is R . Data plotted from a more extensive table in Gauch and Zobel (1989).

Accuracy vs. Intensity in Selecting the Best Lines

As mentioned, with a fixed number of plots to be scored, there is a tradeoff between increasing the accuracy (by increasing the number of replications/line) and increasing the selection intensity. If we are restricted to a total of N plots and our goal is to select the best n lines, what is the optimal number of replications per line? This question has been examined by a number of workers (England 1977, Bos 1983, Gauch and Zobel 1996).

We can solve this problem as follows. Recall from Equation 30.1 that for a set series of lines, the between-line variance σ_G^2 is fixed, and response is maximized by maximizing $\bar{t}/\sigma(\bar{p})$, both of which change with the experimental design. For example, assume only a single environment is of interest and each line is replicated in r plots,

$$\sigma^2(\bar{p}) = \sigma_G^2 + \frac{\sigma_p^2 + \sigma_{w(p)}^2}{r} = \sigma_z^2 \left(H^2 + \frac{1 - H^2}{r} \right)$$

where the broad-sense heritability H^2 is defined using single plots as the reference design ($H^2 = \sigma_G^2 / (\sigma_G^2 + \sigma_p^2 + \sigma_{w(p)}^2)$). If r replicates are used, then the fraction saved is $p = n/(N/r) = nr/N$, where N/r is the number of distinct genotypes measured. The expected response

becomes

$$R(p, r) = \bar{i}_{N/r, n} \frac{\sigma_G^2}{\sigma_z \sqrt{H^2 + (1 - H^2)/r}} = \bar{i}_{N/r, n} \sigma_G \sqrt{\frac{H^2}{H^2 + (1 - H^2)/r}} \quad (30.9a)$$

Using Burrow's approximation (Equation 14.4b) to correct the selection intensity for finite sampling effects gives

$$R(p, r) = \left(\bar{i}_p - \left[\frac{1 - p}{2p(N/r + 1)} \right] \frac{1}{\bar{i}_p} \right) \sigma_G \sqrt{\frac{H^2}{H^2 + (1 - H^2)/r}} \quad (30.9b)$$

We can further simplify calculations by recalling the approximation given by Equation 14.3b for the (infinite-population) selection intensity given a fraction p saved,

$$\bar{i}_p \simeq 0.8 + 0.41 \ln(1/p - 1) = 0.8 + 0.41 \ln(N/nr - 1)$$

For fixed total plots N and final number of lines to be saved n , one can numerically compute the expected response for different replicate sizes to obtain the optimal value. More generally, suppose genotypes are examined over a total of e distinct environments, with r replicate plots/environment. Recalling Equation 30.3, Equation 30.9a becomes

$$R(p, r) = \left(\bar{i}_p - \left[\frac{1 - p}{2p(N/r + 1)} \right] \frac{1}{\bar{i}_p} \right) \frac{\sigma_G^2}{\sqrt{\sigma_G^2 + \frac{\sigma_{G \times E}^2}{e} + \frac{\sigma_p^2 + \sigma_{w(p)}^2}{N}}} \quad (30.9c)$$

With estimates of the appropriate variance components in hand, one can numerically find the values of e and r that maximize $\bar{i}/\sigma(\bar{p})$, and hence maximize response.

Example 30.5. Suppose $H^2 = 0.2$ and we wish to select the best ten genotypes ($n = 10$) from a total of $N = 200$ plots. Applying Equation 30.9b for different values of r gives

r	p	$\bar{i}$	$\sigma_G/\sigma(\bar{p})$	R/σ_G	R/R_{max}
1	0.05	1.29	0.447	0.576	0.806
2	0.10	1.15	0.577	0.666	0.932
4	0.20	1.01	0.707	0.714	0.998
5	0.25	0.96	0.745	0.715	1.000
8	0.40	0.84	0.816	0.685	0.958
10	0.50	0.77	0.845	0.651	0.910

Under these values of H^2 , N , and n , the optimal number of replications/line is 5, resulting in 40 lines being tested. If each line is tested in only a single plot (so that 100 lines are tested), the expected response is only 81% of the optimal.

Another way to look at the efficiency of a particular design is to examine what fraction of the total plots is required under the optimal design to give the same response. Suppose we measured 100 lines, each with two replicate plots. The expected response in this case is $0.666\sigma_G$. The same expected response can also be obtained using only 148 total plots, by examining 37 genotypes each with four replicate plots. Using the optimal design requires only 74% of the number of plots for the original (suboptimal) design.

Increasing the number of replicates always increases the accuracy of selection, bringing the expected response closer to the maximal possible response, $R_{max} = \bar{i} \sigma_G$ (which occurs when all genotypes are measured without error). However, even using the optimal number of replications imposed by N , the response is less than the maximal possible if all genotypes could be clearly distinguished. To see this, again consider the simple case of a line replicated as r plots in a single environment, implying Equation 30.1 can be expressed as

$$R = \frac{\sigma_G^2}{\sqrt{\sigma_G^2 + \sigma_e^2/r}} \bar{i} = \frac{1}{\sqrt{1 + (Ar)^{-1}}} \sigma_G \bar{i} \quad (30.10a)$$

$A = \sigma_G^2/\sigma_e^2 = H^2/(1 - H^2)$ can be considered as a signal (genotypic variance) to noise (environmental variance) ratio (Gauch and Zobel 1996). For the same selection intensity, the ratio of the expected response to the maximal possible response is

$$\frac{R}{R_{max}} = \frac{1}{\sqrt{1 + (1/Ar)}} \quad (30.10b)$$

which is plotted in Figure 30.3.

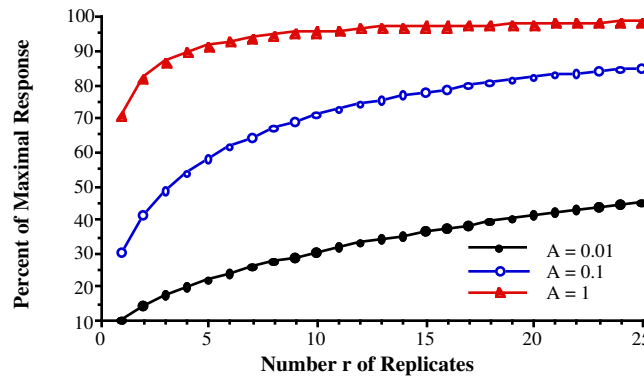


Figure 30.3. Percent of maximal selection response as a function of the number of replicates r and the signal/noise ratio $A = \sigma_G^2/\sigma_e^2$. When σ_e^2 is the environmental variance within a plot, then $A = H^2/(1 - H^2)$.

Example 30.6. Suppose the trait of interest has a broad-sense heritability of $H^2 = 0.2$ (when measured in a single plot of defined size), giving $A = 0.11$. If there is no replication ($r = 1$), Equation 30.10b shows that the expected response is only 31% of the possible response if we could indeed choose the elite genotypes. With 2, 5, and 10 replicate plots, these percentages increase to 42%, 60%, and 72%. To achieve 90, 95, and 99% of the possible response requires 39, 84, and 450 plots.

A key issue in the development of pure lines is how to advance the derived lines from a cross to complete homozygosity and in particular at what stage should selection be practiced. Historically, plant breeders have employed at least some selection during the segregating generations following a hybridization, but for low heritability traits (such as yield), such selection was usually not very effective. Pedigree and bulk selection (see below) are examples of such schemes that use either artificial or natural selection to cull undesirable genotypes as the lines are moving towards fixation. The effectiveness of such **early generations selection** is often questionable, leading to advancement of generations approaches that attempt to minimize selection during inbreeding to fixation (such as single seed descent and doubled haploids, DH), practicing selection only once the resulting lines are very close to completely inbred.

Outcrossing Hybrids Before Selfing

While most selection schemes start by immediately allowing the F_1 s to self (or otherwise inbreed such as by DH), the alternative is to allow one or a few generations of intermating before selfing. Such **interbreeding** schemes are generally not done, as intermating takes time and (above all) effort, as pollinations must be carefully controlled and is often done by hand. The potential advantage of intermating is that it allows recombination to break up favorable alleles tightly linked in repulsion (Baker 1968, Pederson 1974, Bos 1977, Stam 1977, Sneep 1977). Such was the situation for two studies with cotton. Miller and Rawlings (1967) found a negative genetic correlation between yield and fiber strength, and this correlation was reduced by several generations of intermating. Similar results were reported by Meredith and Bridge (1971). Conversely, if there are favorable alleles tightly linked in coupling, intermating would reduce the potential response relative to immediate selfing.

Pedigree Selection

The idea of **pedigree selection** to handle the segregating generations of an autogamous species as the lines inbred to complete homozygosity has its roots at the very beginning of modern plant breeding (e.g., Nilsson-Ehle 1910, Love 1927). It represents a compromise at trying to move self-pollinating lines along towards complete homozygosity while still allowing some selection to occur. Under pedigree selection, as individuals become more inbred, selection decisions move from individual plants (during the F_2 and perhaps F_3 generations) to family-based decisions and finally to extensive yield trials (with extensive replication) in later generations (the F_6 or higher) when resulting lines are sufficiently inbred. There are any number of variants of pedigree selection (e.g., Lupton and Whitehouse 1957, Riggs et al. 1981). One popular example is the F_2 **progeny** design introduced by Lupton and Whitehouse (1957), illustrated in Figure 30.4. Pedigree selection has been historically preferred over mass selection by breeders of autogamous crops as selection can be shifted towards family, rather than individual, performance, increasing the efficiency for low heritability traits. This scheme attempts to balance between being able to locate superior genotypes during early generations of selection while still maintaining sufficient genetic variation for selection in later generations.

Individual selection in early generations requires that plants are grown under space-planted conditions so that individuals can be scored (as opposed to the spacing used in many crops where individual scoring is essentially impossible). Family selection, on the other hand, often involves planting family members in a bulk, with the entire plot (as opposed to single individuals) being scored. Given its increased accuracy, one would like to start family selection as early as possible. For many crops, however, there is simply not sufficient seed from individual F_2 plants for family selection, and instead family selection starts by taking seed from F_3 plants from each selected F_2 plant. When plants produce sufficient seed, F_2

families can be tested, and pure lines selected from the superior families. Given replication and the generally large number of families to be tested, identification of superior F_2 families often requires considerable resources. One pedigree scheme turns this testing on its head — under the **pure-line family method** (or PLF) of Ivers and Fehr (1978), pure lines are used to identify the superior families. Here, individual F_2 plants are harvested and bulk progeny from each plant are selfed for several generations (typically the F_4 or F_5). A pure line is then extracted from each of these bulks and yield tested. For superior lines, additional members from the original F_2 family are chosen for yield testing.

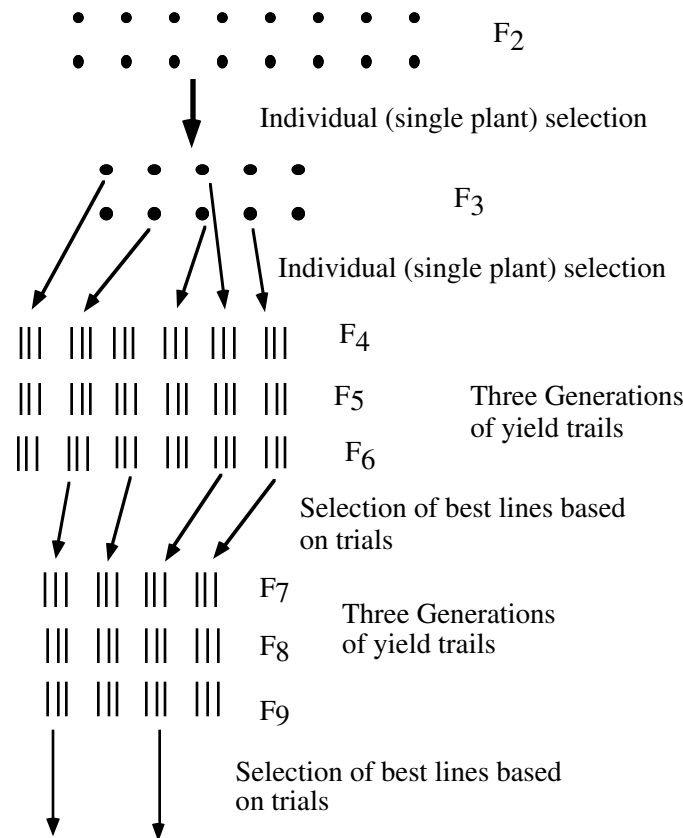


Figure 30.4. The F_2 progeny method of pedigree selection, proposed by Lupton and Whitehouse (1957). Single circles represent individual plants (usually space-planted). Each line represents a plot (such as a rod-row) of the progeny from a single parent. Families are replicated over several plots for trials, and are chosen from the average over a three generation period (e.g., F_4 to F_6). All reproduction is by strict selfing, and arrows indicate where selection has occurred.

Given the goal of developing pure lines that perform well over a number of environments, pedigree selection can involve **sequential selection**, where the different generations are cycled through a series of different environments to minimize the effects of $G \times E$. (e.g., Branch et al. 1991). An interesting example of the effects of selection in different environments is the work of St-Pierre et al. (1967), who examined selection of yield in barley on F_2 to F_5 plants in a different series of environments. Starting with the same F_1 , lines were

grown and selected over all possible yearly combinations involving two sites, Macdonald college in Montreal and La Pocatière (about 250 miles NE). This was done for 27 different F_1 families for each combination of environments, and the resulting F_7 and F_8 lines were yield-tested at both sites. Figure 30.5 summarizes their results. Notice that the top six performing lines were all selected at La Pocatière in the F_4 , while the six lowest performing lines were selected at Macdonald in the F_4 . The authors suggest that the stressed environment at La Pocatière during the year differentiates barley strains better than a lower stress (Macdonald) environment.

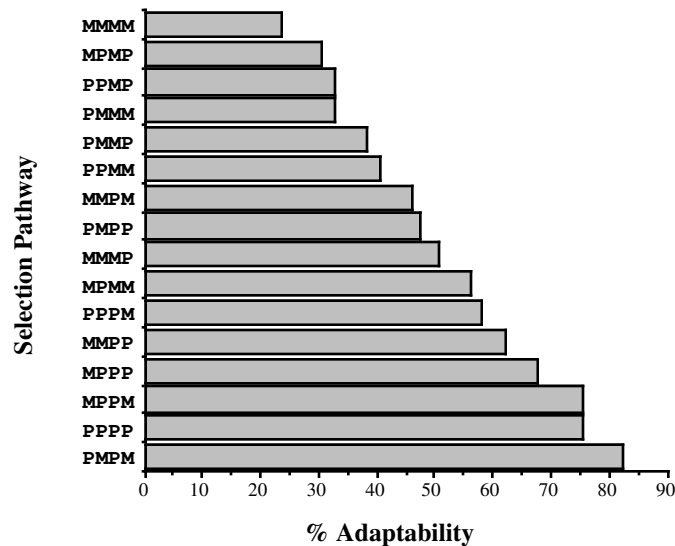


Figure 30.5. Results of the segregants of the same F_1 when grown under different environments. Seeds from selected barely plants each year were grown in two different locations, M = Macdonald college and P = La Pocatière. Thus MMPM implies that the resulting seed were selected in the F_2 and F_3 at Macdonald, the F_4 at La Pocatière, and the F_5 back at Macdonald. In the F_7 and F_8 , strains were yield-tested at both sites. The measure “percent adaptability” for a specific selection pathway (e.g., MMPP) is the percentage of times that plants selected by that pathway had a yield exceeding the mean yield of all progeny from their original F_1 family. After St-Pierre et al. (1967).

While pedigree selection schemes have historically been widely used for a number of autogamous crops, there are two major concerns with this approach — the effectiveness of individual selection and the effectiveness of early-generation selection.

Because individual selection for low heritability traits (such as yield) is very inefficient, one modification of pedigree selection is to focus on only highly heritability traits (such as height or disease resistance) in the early generations, postponing selection on other traits until later generations. For lower heritability traits, several approaches have been suggested to improve the effectiveness of individual selection. One is improved planting designs to reduce the error variance associated with soil heterogeneity, such as stratified mass selection (Gardner 1961), introduced in Chapter 13. Another approach is the **honeycomb design** of Fasoulas (1973), where individuals are planted in a honeycomb fashion and compared with their six equidistant neighbors. Wide spacing (e.g., 60cm in wheat) is used to avoid the effects of competition and the neighborhood comparison will reduce the effect of soil heterogeneity. Mitchell et al. (1982) used this design for early generation selection in Durum wheat yield.

While they found a slight improvement under this design (compared with standard mass selection), it was not deemed sufficiently large to be worth the additional resources. A second approach is to select on component characters which have higher heritabilities than the trait of interest. For example, Valentine (1979) examined yield and its components in Spring Barley. While the broad-sense heritability (for single plants) for yield was 5.3%, heritability was at least twice as high in several components (20.2% for grain weight, 12.9% for grains/plant, 21.2% for grains/tiller). Another approach is to asexually propagate individuals and then use replication to obtain much more precise estimates of their genotypic values. However, asexual propagation is not usually justifiable given the considerable investment in both time and resources.

Finally, a confounding factor is that individuals are often chosen via **visual selection**, where in the breeders makes a selection decision by a simple visual assessment as opposed to basing selection on detailed measurements. This is often not unreasonable, given the large amount of material that must be sorted. Visual selection of individuals for yield is not very accurate, and has been found to be only appropriate for rejecting the poorest lines (Hanson et al. 1962, Kwon and Torrie 1964 for soybeans; Atkins 1964 for Barley; Briggs and Shebeski 1970 for wheat). Indeed, Anderson and Howard (1981) found considerable differences between potato lines developed under pedigree selection when individual plants were chosen by different selectors. While Luedders et al. (1973) found that visual selection for soybean yield as accurate as F_4 or F_5 tests, this may be more a reflection on the low heritability of yield.

Even if we can alleviate the concerns of individual selection, a second and perhaps more serious problem with pedigree selection is that we are interested in the performance of completely inbred (F_∞) lines, and F_2 or F_3 lines are often very poor predictors of the final line performance. Thus, early-generation selection, whether individual- or family-based, may simply not be efficient. Because of both potential differences in planting designs (space-planting versus more seeding densities) and in the composition of the population, high yielding genotypes in early generations may not be high yielding under normal cropping conditions. In early generations, individual performance is scored in heterogeneous populations (mixtures of different genotypes), but for cropping conditions we seek the best performance in a monoculture. Thus early generations select for performance under **allogenotypic competition**, while crop conditions require performance under **autogenotypic competition**. Pedigree selection tends to select for lines that can do well in both mixed and pure cultures, because successive lines must survive both early generations selection and yield testing in later generations. It is thus unclear if lines with the highest performance in pure cultures are those that are also most favored by pedigree selection. We examine the effectiveness of early generation selection in more detail shortly.

While the pedigree approach has been very widely applied, one sobering appraisal of its success is offered by Kulshrestha (1989), who examined 630 wheat strains developed using pedigree selection by the Indian government's National Wheat Program. These lines were tested at a total of 271 locations from 1983 to 1987, and involved tests in nine agro-climatic zones. **Checks** (elite cultivars grown as controls) were also used and results reported compared to the best check. The results (summarized in Table 30.1) was that only 4.4 % (27 of 630) of the lines were significantly better than the best check, and only 14.8% (93 of 630) where numerically (but not significantly) superior. Fully 81% of all lines either no better or inferior to the best check.

Table 30.1. The success (and lack thereof) of the pedigree method of selection as applied to wheat in Indian. From Kulshrestha (1989).

No. of	No. of	No. cultivars compared to best check		
		Significantly	Numerically	Equal or

Zone	locations	cultivars	superior	superior	inferior
Northern Hills	30	50	—	16	34
Northern Plains	49	115	—	9	106
NW Plains	31	76	4	11	61
NE Plains	44	79	4	9	66
Far Eastern	18	73	—	11	62
Southeastern	10	55	5	13	37
Central	40	64	—	—	64
Pennisular	45	72	6	12	54
Southern Hills	4	46	8	12	26
Totals	271	630	27	93	510

Bulk Selection

Pedigree selection requires consider effort on the part of the breeder in tracking individual lines. An alternative approach is **bulk selection**, which involves growing selfed lines in bulk (all seeds are planted in a plot under normal seeding densities) and simply collecting and planting the resulting seed for several generations. Under this scheme, the resulting lines have survived natural selection for both competitive ability and yield. In modern planting breeding, this approach was apparently first used Nilsson-Ehle in 1908 to screen winter wheat. Suneson (1956), motivated by the work of Harlan and Martini (1929, 1938), brought this approach to prominence.

While considerably easier to implement than pedigree selection, bulk selection also potentially suffers the problem of early generation selection — the genotypes that compete the best in mixed stands may not be those that give the highest yields in pure mixtures. Further, in a segregating population, genotypes change over time and the advantage of one genotype may not be fully passed onto its offspring. Empig and Fehr (1971) note that segregation supplies new weak competitors (until the inbreeding becomes sufficiently high that little further segregation occurs), hence the population can change rather slowly over time unless there is strong selection for competition. Suneson (1956) notes that significant shifts in the mean may not occur until Generation 15 (Figure 30.6).

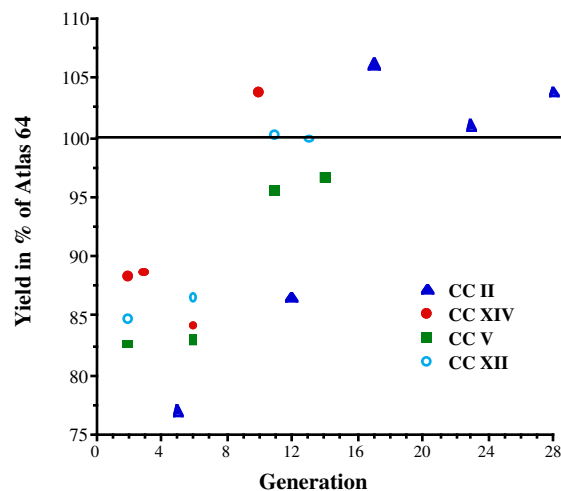


Figure 30.6. The yields of four different bulked composite lines of barley against a check (Atlas 64) as a function of generations of bulk selection. Composite cross II (CC II) results

from the complete intercross of 28 varieties, with bulk selection starting on 378 separate F_1 families. CC V results from 31 varieties, CC XII from 26 varieties, and CC XIV resulted from 9 varieties of California-adapted varieties. While CCs II, V, and XII almost entirely selfed, CC XIV contains a male sterile gene to facilitate outcrossing. Bulk selection was successful in producing yields equal to the value of Atlas 64, but required at least ten or more generations. After Suneson (1956).

A second problem (not suffered by pedigree selection) is that bulk selection does not directly select on yield per se, and high yielding lines may actually be selected against. For example, in certain crops, tall plants seem to have a competitive advance over short plants (Bal et al. 1959 for barley, Jennings and Herrera 1968 for rice, Khalifa and Qualset 1975 for wheat). In many cases, however, short-statured plants have a greater yield potential (Briggle and Vogel 1968, Athwal 1971). Sunsen (1949) examined the competition between four pure lines, finding that the dominant strain out-competes the strain with the best agricultural features (Figure 30.7).

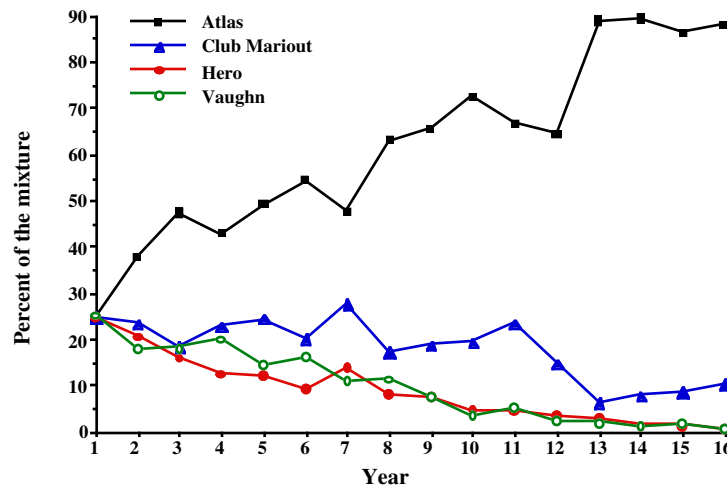


Figure 30.7. Suneson (1949) followed the frequencies of four barley varieties grown (and hence selected) in bulk over a 16 year period (1933 to 1948). As the following figure shows, the Atlas variety wins, but the Vaughn variety is superior from an agricultural standpoint, with a higher yield (107% of Atlas), and earlier heading date, and much better disease resistance.

Despite these potential concerns, bulk selection has shown some success (Figure 30.6). It certainly seems a reasonable approach to screen for lines that can survive extreme conditions for certain growing regions. Ikehashi and Fujimaki (1980) report that 21 of 33 rice cultivars released by the Japanese government in 1977 were developed by bulk selection. Likewise, (as of 1980) lines accounted for 44% of the total rice area in Japan were developed using Bulk selection. Bulk and pedigree methods resulted in the same increase in yield in Soybeans (Raebler and Weber 1953), while Torrie (1958) found that bulk and pedigree methods gave the same yield in 4 of 6 lines, with bulk breeding having a higher yield in the other two tested F_6 lines.

Hamblin and Rowell (1975) tried to quantify the effects of competition on the decision of bulk vs. pedigree selection. If β is the slope of the regression of yield in mixed cultures vs. yield in a pure culture, then if $\beta < -1$, bulk selection can lose high-yielding lines, and

pedigree selection is appropriate. If $\beta > -1$ then bulk selection is reasonable.

Is Early Generation Selection Effective?

Two issues: correlation of genotypes over generations (e.g., does F_3 well-predict F_∞ ?, resulting G $\times$ E issues

A major motivation for early generation selection is to avoid the loss of desirable alleles that occurs in the absence of such selection. As inbreeding progresses, individuals carrying favorable alleles at a large number of loci become increasingly infrequent (Van Der Kley 1955, Snee 1977, Snape and Simpson 1984). To see this, consider the probability that a locus carries a favorable allele. We assume that where alleles differ between the pure lines being crossed, alleles either have an equivalent effect on the character of interest (and hence are ignored) or one is more favorable than the other. Overdominance obviously does not fit into this pattern, but since we are breeding lines to fixation such loci are not considered here (Chapter 31 discusses cross-breeding schemes for such loci in inbred lines). In the F_1 all individuals are identical, being completely heterozygous for all loci at which the lines differ. Because of this, no selection is practiced at this stage unless we are considering a collection of lines from different crosses (in which case the F_1 s differ over crosses). In the F_2 and subsequent generations, the probability that an individual carries at least one favorable allele at a particular locus is the frequency of heterozygotes plus the frequency of favorable homozygotes. In the absence of selection, this is a standard Hardy-Weinburg under inbreeding (LW Table 10.1), with the frequency of the favorable allele being $1/2$. Hence, the probability π that an individual carries at least one favorable allele at a particular locus is

$$\begin{aligned}\pi_t &= \text{Freq}(Aa) + \text{Freq}(AA) = (1 - f_t) \cdot \frac{1}{2} + \left[(1 - f_t) \cdot \frac{1}{4} + f_t \cdot \frac{1}{2} \right] \\ &= \frac{3 - f_t}{4} = \frac{1}{2} \left(1 + \left(\frac{1}{2} \right)^{t+1} \right)\end{aligned}\quad (30.11)$$

where $f_t = 1 - (1/2)^t$ is the inbreeding coefficient and t number of generations of selecting. This ranges from $3/4$ in the F_2 ($t = 0, f = 0$) to $1/2$ in the F_∞ ($f = 1$). If the trait of interest is governed by 10 favorable loci, then (assuming loci are unlinked) the probability that an individual contains at least one favorable allele at each locus is 5.6% in the F_2 , but only 0.098% in the F_∞ . On average, one out of every 18 F_2 individuals contains favorable alleles at all loci, while only one out of every 1,024 F_∞ individuals does. More generally, if loci are unlinked, then the probability that an individual contains at least k out of n loci with favorable alleles follows from the binomial distribution with success parameter π_t and sample size n .

Hence, in the absence of selection, if the trait of interest is controlled by even a modest number of loci that differ between the lines, unless a very large collection of F_∞ lines is maintained, it is highly unlikely that the best possible genotypes will be included in the sample to be yield-tested. Thus, despite all of the concerns addressed about early generation selection, the situation may be more dire without such selection.

The literature reports mixed results on the effectiveness of early-generations selection. Traits with high heritabilities seem to respond better than traits with lower heritabilities (such as yield). For example, McKenzie and Lambert (1961) and Hanson et al. (1979) found that early generation selection was of little value for yield in Barley, while Frey (1954) found that F_2 performance was good indicator of F_3 performance for yield. Jinks and Pooni (1981a) found early generation selection in tobacco (*Nicotiana rustica*) was not useful.

As mentioned, a major concern with early-generation testing is that dramatic environmental differences can occur year-to-year even in the same location, and genotype $\times$ environment interaction can pose serious problems in such cases. For example, Mahumud

and Kramer (1951) found little association for yield between F_3 and their selected F_4 descents in soybeans when different plant spacing and seasons were involved. Conversely, the F_3 and F_4 showed good associations when grown (using remnant seed) in the same location and year. For yield in Dry Beans (*Phaseolus vulgaris*), Hamblin and Evans (1976) found that there was little effect of generation on the correlation, provided both lines were grown at crop densities. Generations using more space-planted densities were much poorer predictors.

Perhaps the most detailed study was that of Whan et al. (1982), who examined the increase in yield in wheat, summarized in Table 30.2. When lines are tested in the same year and site, early generations selection for yield was clearly effective, with mean response of selected lines being 120-130% of unselected lines. When tested in the same year, but different sites, early generations selection was still effective, with responses of 110-120%. When derived lines were tested in different years (but the same sites) as their ancestors, the response was essentially the same as for unselected lines.

Table 30.2. Results from a multi-generational wheat yield study by Whan et al (1982). Starting with 56 F_2 families (the total for the two crosses followed), the authors advanced lines by randomly selected one line per family each generation, with the pedigree data carefully tracked. By retroactively using this pedigree data, the authors could simulate the results for selection response at various generations by asking about the performance of descendant lines from the top 10% of lines. Selection was either simulated using individual value (F_k -derived lines) or lines chosen based on their family value, F_k (based on F_{k+1} family) lines. Using remnant seed allowed lines from different generations to be grown in the same location and same year in some of the experiments.

Type of Selection	Percent of Mean Response		
	F_3 -derived lines	F_4 -derived lines	F_5 -derived lines
Same year, site			
F_2 -derived lines	125	111	113
F_2 (based on F_3 family)	—	119	135
F_3 -derived lines	—	123	134
F_3 (based on F_4 family)	—	—	132
F_4 -derived lines	—	—	128
Same year, different site			
F_2 -derived lines	117	120	126
F_2 (based on mean of F_3 lines)	112	117	123
Same site, different years			
F_2 -derived lines	106	104	101
F_2 (based on mean of F_3 lines)	101	91	111

Table 30.3 summarizes the between-generation correlations in yield for a number of crops (grown in different years except where noted). While correlations were positive and significant, a number of authors have questioned whether they are sufficiently high to justify the considerable increase in effort required for **early generations yield testing (EGYT)**. For example, even with a correlation as high as 0.60, Knott (1972) and Knott and Kumar (1975) questioned whether EGYT was worth the increased effort.

Table 30.3. Between-generation correlations in yield. Where reported, level of significance is indicated by ** = 0.01, *** = 0.001.

Winter wheat (<i>Triticum aestivum</i> em. Thell.)	
F_4 - F_5 correlation of 0.555***	Lupton and Whitehouse (1957)

Spring wheat

F ₃ - F ₅ correlation of 0.847**	Shebeski (1967)
F ₃ - F ₅ correlation of 0.83**	Briggs and Shebeski (1971)

Wheat

F ₃ - F ₄ correlation of 0.59**	De Pauw and Shebeski (1973)
F ₃ - F ₅ correlation of 0.56**	De Pauw and Shebeski (1973)
F ₄ - F ₅ correlation of 0.75**	Busch et al. i (1974)
F ₃ - F ₅ correlations of 0.29** and 0.14**	Knott and Kumar (1975)
F ₂ - F ₅ correlation of 0.81** (grown in same year)	Cregan and Busch (1977)

Barley (*Hordeum vulgare*)

F ₃ - F ₆ correlations of 0.543*** and 0.313**	McKenzie and Lambert (1961)
--	-----------------------------

Soybeans (*Glycine max*)

F ₂ - F ₃ correlation of 0.85	Leffel and Hanson (1961)
---	--------------------------

Another issue, noted by Snee (1984), relating to early generation testing is that for certain small grains (such as barley and wheat), the number of seeds per plant is too small to provide sufficient power for yield-testing of their progeny (especially when considerable G X E exists, which can only be mitigated by extensive replication across environments). Hence, an F₃ test is not that efficient, and instead the F₃ should be used to create a larger F₄ family for yields testing.

Single Seed Descent (SSD) Selection

The motivation behind the **single seed descent (SSD)** is that since in many cases early-generations selection can be misleading, the simplest approach is to inbred the lines in (as much as possible) the absence of selection. One way to do this is to propagate a single seed from each selfed plant. This is done under space planted conditions (to reduce competition) until the F₅ or F₆ generation, at which point the resulting lines are expanded and yield-tested. This approach was initially suggested by Goulden (1939), and later by Grafius (1965), Brim (1966, who called this approach the **modified pedigree method**), and Kaufman (1971, who referred to this as the **random method**). The term single seed descent is due to Empig and Fehr (1971).

One subtle, but important, advantage of SSD is that it allows populations to be grown in environments quite different from those in which the final cultivar will be grown, without being unduly influenced by selection. Hence, lines can be advanced towards complete inbreeding in a greenhouse and/or in winter nurseries experiencing a very different environment than the final cultivar. SSD has been used on a number of crops – most cereals, soybeans, tomato, lettuce, and safflower to name a few (Stoskopf et al. 1993). SSD requires more work than bulk selection, as single seeds must be harvested from each plant, but this is far less work than pedigree selection.

When starting a set of lines for SSD, the F₁ population should be as large as possible. The number of lines limits the possible range of genotypes that will be recovered (we discuss this further below in the context of selection limits). Further, the population size is expected decline over time, as some plants will die out or not seed. Roy (1976) found in wheat that SSD lines are impacted by competition but not to the extent of bulk selection. He suggests growing plants under conditions of relatively uniform plant survival to mitigate the effects of competition on the development of SSD lines.

In the absence of selection, SSD and pedigree inbreeding are expected to generate the same F_∞ distribution of lines. This expectation was indeed observed by Jinks and Pooni

(1984), who found no differences in phenotypic and genotypic distributions of seven quantitative traits in *Nicotiana rustica* lines developed from SSD and pedigree inbreeding.

Modifications of the standard SSD design have been proposed. Casali and Tigchelaar (1975a) suggest pedigree selection in the F_2 and F_3 and then use SSD to advance to the F_6 for yield traits. Compton (1968) suggested that instead of propagating lines from a collection of F_1 individuals from a single cross, that instead a large number of crosses be made, and a single F_1 line is propagated from each cross. This exploits both between-line as well as within-line selection, increasing response.

Doubled Haploid (DH) Selection

In a number of plants, cytological techniques can allow one to directly create completely homozygous line by producing **doubled haploids** (Kermicle 1969, Nitzsche and Wenzel 1977, Choo 1981). As with SSD, the motivation for DH is to sample a collection of completely inbred lines that have undergone as little selection as possible. The advantage of DH over SSD is that a single generation takes a line to complete homozygosity, as opposed to selfing for 5-6 generations (which only takes a line to, respectively, 96.9% and 98.4 % homozygosity). This advantage is partially offset by the additional effort required to haploidize a line. In the absence of selection, SSD, and DH produce an identical distribution of pure lines, except when linkage is present (Jinks and Pooni 1981). Since DH lines typically undergo only a single round of meiosis (with gametes from the F_1 's being haploidized), the resulting DH lines can be strongly influenced by parental linkage phase, even for moderately linked loci. Favorable genes in repulsion in parents tend to remain so, reducing the joint probability of fixation relative to SSD. Conversely, favorable genes in coupling have a higher chance of being jointly fixed than under SSD. Powell et al. (1986) found no evidence of such linkage effects for yield in spring barley, where SSD and DH lines showed the same performance. If linkage is a concern, its effect can be reduced by allowing several generations of meiosis before haploization.

Comparison of the Different Methods

The breeder is faced with several options for advancing generations: active selection (pedigree), passive selection (bulk), or random selection (SSD, DH). The general consensus is that early generation selection (usually referred to as EGT, for **early generation testing**) is certainly appropriate for simply-inherited characters, or for traits with high heritabilities. For low heritability traits, it is not clear if the additional effort required for EGT is worthwhile, and several workers have used computer simulations in an attempt to examine the relative efficiencies of the different methods. Casali and Tigchelaar (1975b) compared SSD, bulk and pedigree selection. At heritabilities above 0.5, pedigree selection was most efficient, while at heritability of 0.25, mass selection was most efficient, while all methods were roughly equivalent at low heritabilities (0.10). SSD was the most efficient approach when multiple characters were under selection. Yonezawa et al. (1987) showed that, relative to pedigree selection, DH had an advantage provided the trait is influenced by a relatively small number of loci, favorable alleles are recessive, and genes are not strong linked. Van Oeveren and Stam (1992) compare selection starting in the F_3 with selection in the F_6 following SSD (a scheme they called **early selection**, or **ES**). The ES F_6 lines were compared with the F_7 from SSD. ES has a significant advantage over SSD when heritability is high. When less than 100 F_2 plants are used, the results from SSD are poor. The number of loci (of equal effect) underlying the trait is also important, with both methods giving a larger response when fewer loci are involved, while the difference between methods diminishes as the number of loci becomes smaller. When more than one cross is involved, SSD outperformed ES, especially under low heritabilities.

While simulation studies can be used to examine various genetic models, they do not fully accommodate important concerns, such as the role of genotype $\times$ environment interaction (e.g., year, local and competition effects). A number of workers have directly compared various methods in the field, and their results (summarized in Table 30.4) find that no single method is clearly superior (or inferior) to the others. Given this, early-generation traits for yield and other low heritability traits are generally not an efficient use of effort, and a passive (Bulk) or random (SSD, DH) selection approach is favored because they do not require expensive yield testing (until later generations) while still allowing for a rapid generation advance.

Table 30.4. Comparison of response under different generation advancement methods. A * indicates experiments where one (or more) methods were found to be superior to others tested. SSD = single-seed-decent selection, PS = pedigree selection (various schemes), BS = bulk selection, EGYT = early-generation yield testing (various schemes), DH = doubled haploid selection, SPS = single plant (mass) selection.

Soybeans (<i>Glycine max</i>)	
No consistent F_6 differences between PS, BS	Raeber and Weber (1953)
No consistent F_6/F_7 differences between EGYT (F_4, F_5), PS, SSD	Luedders et al. (1973)
No consistent F_8 differences between EGYT, PS, SSD	Boerma and Copper (1975)
SSD, BS showed no consistent differences	Empig and Frey (1971)
BS, PS equal in 4 of 6 crosses	Toorie (1958)
SSD, PS equivalent	Byron and Orf (1991)
* BS significantly greater than PS equal in 2 of 6 crosses	Toorie (1958)
* EGYT (F_4) superior to BS, PS	Voigt and Weber (1960)
Wheat (<i>Triticum aestivum</i>)	
*SSD, SPS equivalent; both superior to BS	Srivastave et al. (1989)
BS, SSD methods give identical results	Tee and Qualset (1975)
Barley (<i>Hordeum vulgare</i>)	
BS, PS equally effective	Harlan et al. (1940)
DH, SSD, PS equally effective	Park et al. (1976), Powell et al. (1986)
Chickpea (<i>Cicer arietinum</i>)	
*SPS on seed size gave greater yield than SSD, which gave greater yield than SPS on yield itself	Bisenn et al. (1985)
Greengram (<i>Vigna radiata</i>) CHECK	
*SPS gave greater response than F_3 SSD lines	Dahiya and Singh (1986)
Long Bean (<i>Vigna sesquipedalis</i>)	
PS, SPS, BS gave similar results	Yap et al. (1977)
Tomato	
*PS slightly better than SSD	Peirce (1977)

Early generations testing seems to be most efficient in culling the lowest performing strains. For example, Knott and Kumar (1975) examined yield in both yield-tested, YT (i.e., early generation selection) and SSD lines. Although the mean yield was higher for YT lines, this was entirely due to the exclusion of low-performance strains. The best performing (upper

20%) of SSD lines were at least as good as the YT lines. Similar findings were seen for barley yield by Park et al (1976), who examined pedigree-selected, SSD, and DH lines. Mean grain yield in superior lines was similar in all three methods, while the lower limit of yield was higher in the pedigree-selected lines than in SSD or DH lines.

One reason for little discrimination between the various methods for low heritability traits may be that SSD and DH lines are expected to retain more genetic variance than lines subjected to early selection, which lowers the effective population size. Thus, a slight advantage from early generations selection can be offset by the increase in genetic variance in DH/SSD lines. In barley, Park et al (1976) indeed observed that pedigree-selected lines showed lower genetic variance across lines than did SSD or DH lines. However, Boerma and Copper (1975) found that pedigree-selected soybean lines appeared to retain more genetic variance in yield than did SSD lines.

SAMPLING OF GENOTYPES AND SELECTION LIMITS

One concern with SSD and DH lines, which holds for pedigree selection on low heritability traits as well, is that only a very small subset of the possible genotypes are represented in the final collection of lines to be tested. If the number of final lines is too small, the chances of recovering favorable genotypes are greatly diminished. Likewise, the selection limit imposed by choosing the best genotype in our sample is also greatly affected by the sample size.

Suppose we wish to recover at least r lines with a particular genotype (which occurs with probability q). How many lines must be sampled so have probability p that at least r such lines are included? This is a binomial sampling problem, with success parameter q , sampling from N lines. The solution is given by the smallest N satisfying

$$\sum_{j=r}^N \frac{(N-j)! j!}{N!} q^j (1-q)^{N-j} \geq p \quad (30.12a)$$

This is usually more easily computed by considering the complementary event,

$$1 - \sum_{j=0}^{r-1} \frac{(N-j)! j!}{N!} q^j (1-q)^{N-j} \geq p \quad (30.12b)$$

One quick approximation (which is usually a conservative overestimate) is to take

$$N \geq r \frac{\ln(1-p)}{\ln(1-q)} \quad (30.13)$$

This is simply r times the exact solution for $r = 1$. While better approximations are presented by Sedcole (1977), Equation 30.12b can quickly be solved by trial and error using the N from Equation 30.13 for an initial starting point.

Example 30.7. Suppose we are interested in ultimately recovering an SSD (or DH) line that is homozygous for five particular (unlinked) loci. Here, $q = (1/2)^5$. How many SSD or DH lines must be examined to have a 95% chance of seeing at least 5 such lines? Equation 30.13 gives

$$N \geq 5 \cdot \frac{\ln(1-0.95)}{\ln(1-0.5^5)} = 471.8$$

Trail and error using Equation 30.12b gives $N = 427$.

A related problem is find the sample size required to ensure a certain number of lines that contains at least k loci homozygous for favorable alleles (England 1981, Jansen and Jansen 1990, Jansen 1992). Suppose the two pure lines to be crossed differ at n loci with favorable alleles for the trait of interest. In the inbreeding limit (assuming unlinked loci), the probability of any particular homozygote at a locus is just $1/2$. Thus, the probability q that any single line contains at least k of n of the loci homozygous for the favored allele as

$$\begin{aligned} q &= \sum_{i=k}^n \frac{(n-i)! i!}{n!} (1/2)^i (1 - 1/2)^{n-i} \\ &= (1/2)^n \sum_{i=k}^n \frac{(n-i)! i!}{n!} \end{aligned} \quad (30.14)$$

This value of q is substituted into Equations 30.12a and b. More generally, under SSD (or other selfing schemes), many generations are required for a random loci to reach homozygosity. Recalling LW Table 10.1 for the frequencies of genotypes under inbreeding, the probability of a particular homozygote after t generations of selfing (i.e, the S_{t+2} for a pure line cross, as the S_0 is the F_2) is

$$\pi_{Hom}^{(t)} = (1/4)(1/2)^t + (1/2)(1 - 2^{-t}) = \frac{1}{2} \left(1 - \frac{1}{2^{t+1}} \right) \quad (30.15)$$

Thus, the probability in the F_t (which has undergone $t - 2$ generations of selfing) that an SSD line contains at least k loci homozygous for favorable alleles is

$$\begin{aligned} q(t) &= \sum_{i=k}^n \frac{(n-i)! i!}{n!} \left(\pi_{Hom}^{(t)} \right)^i \left(1 - \pi_{Hom}^{(t)} \right)^{n-i} \\ &= \left(\frac{1}{2} \right)^n \cdot \sum_{i=k}^n \frac{(n-i)! i!}{n!} \left(1 - \frac{1}{2^{t-1}} \right)^i \left(1 + \frac{1}{2^{t-1}} \right)^{n-i} \end{aligned} \quad (30.16)$$

Likewise, recall that we had previously (Equation 30.11) computed the probability that a locus contains a favorable allele. Substituting this result gives the probability an individual contains favorable alleles at k or more loci in the F_{t+2} under SSD (or selfing in the absence of selection) is

$$q = \left(\frac{1}{2} \right)^n \cdot \sum_{i=k}^n \frac{(n-i)! i!}{n!} \left(1 + \frac{1}{2^{t-1}} \right)^i \left(1 - \frac{1}{2^{t-1}} \right)^{n-i} \quad (30.17)$$

Expressions when loci are linked can be found in Jansen and Jansen (1990) and Jansen (1992). As was pointed out by England (1981), the expected number of SSD lines required is usually larger (1.2 to over 3-fold larger) than DH lines when the SSD lines are not completely inbred (for example, using the typical F_6 lines).

Both SSD and DH capture only a very small subset of the possible genotypes that can segregate from an F_1 . For example, if there are 10 favorable loci between the two strains, then $2^{10} = 1024$ different homozygous lines can be produced, while for 25 loci, there are roughly 3.4 million different possible homozygous genotypes. This has obvious consequences for the

selection limit, and a number of authors have investigated the effect of the number of lines on the final limit (Bliss and Gates 1968, Bailey and Comstock 1976, Bailey 1980, England 1981, Schwarzbach 1981).

The most useful starting point is the simulations of Schwarzbach (1981), who obtained the theoretical maximum response to selection on inbred lines when N lines are maintained. Schwarzbach assumed in the F_2 that the individual with the highest percentage of favorable alleles is chosen to form the next generation. In the F_3 again the individual with the highest fraction of favorable alleles is chosen and so on to fixation of a pure line. This admittedly unrealistic scheme gives the maximum possible response under any selection scheme involving pure selfing, and in fact most selection schemes are far, far less efficient. Figure 30.8 summarized Schwarzbach's results. Bailey and Comstock (1976) and Bailey (1980) examine the consequences of linkage. As mentioned, linkage effects are most serious with DH lines, as these are typically formed after only a single round of recombination, as opposed to multiple rounds of recombination during SSD towards the formation of a pure line.

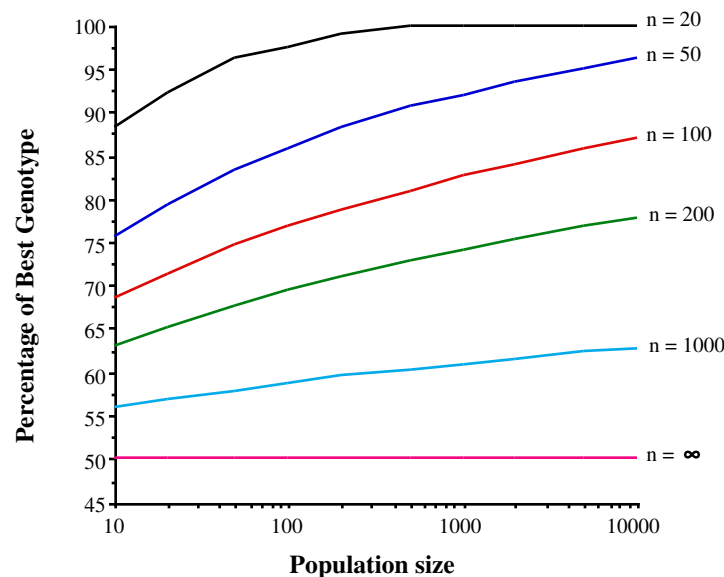


Figure 30.8. The theoretical maximum response to selection on inbred lines given n loci and N lines. From the (more extensive) tables in Schwarzbach (1981).

Bliss and Gates (1968) showed for n loci of equal effect, that the fraction of the maximal possible response given selection among a population of pure lines is

$$\bar{i} \sqrt{\frac{H^2}{n}} \quad (30.18)$$

This result assumes an infinite number of lines and corrects for the strength of selection and number of loci, but ignores the sampling consequences of a finite number of lines. Thus, it also provides an upper bound for the response, but one based on number of loci as opposed to number of lines.

As is the case for mass selection in an outbred population (Chapter 26), there is an optimal selection intensity for maximum response. Under the simulation conditions assumed by Schwarzbach (1981), the optimal intensity was between 12 and 25 percent selection and

the response curve (as a function of intensity) was rather flat in this region. The lowest total response occurred under the highest selection intensity.

CHOICE OF PARENTS

The final issue for breeding an optimal pure line is the choice of parents. In particular, what is the best approach to predict **parental prepotency** — finding the parents whose inbred hybrid offspring will yield the best lines. The principal difficulty is that our interest is in the completely-inbred (F_∞) lines, and these take a number of generations to realize. Further, since our goal is extracting the *best* lines, the mean of the F_∞ lines from a particular cross may be below average, but this cross may also segregate out the most exceptional lines.

Often parental lines are chosen because we wish to combine favorable features of each in the derived lines, such as crossing one parent with strong yield with a second with better disease resistance, or crossing a line with high yield in one location with a line that has high yield in a different location. In such cases, the choice of parents can be relatively obvious, especially if the traits of interest have a simple (few loci) genetic basis. For more genetically complex traits, the choice of parents even in this case is not necessarily obvious.

A more difficult issue is how to choose parents such that the resulting lines show **transgressive segregation** for the trait of interest — with some of the derived hybrid lines exceeding the values of either parent. Such is usually the case for yield, where we wish to extract lines with higher yield than either parent. The idea is that each line contains favorable alleles that the other lacks. Recombination generates lines containing more favorable alleles than either parent, and hence a more extreme character value.

Ideally, one could use just phenotypic information from parental lines (avoiding the cost and time for line crosses), and we examine such strategies first. The other approach is to actually cross at least certain lines and use this information to increase our predictable ability. One obvious approach is to pick those lines with the best mean values over the first few generations of inbreeding. While a number of workers suggested that using a bulk (measuring all descendent lines in a single batch) was sufficient for predicting the crosses that will generate a high proportion of high-yield lines (Harlan et al 1940, Immer 1941, and Smith and Lambert 1968 for barley; Harrington 1940, Lupton 1961, Busch et al. 1974, and Cregan and Busch 1977 for wheat; Leffel and Hanson 1961 for soybeans), others questioned the predictive value of these bulks (Fowler and Heyne 1955 for wheat; Weiss et al. 1947 and Kalton 1948 for soybeans; Atkins and Murphy 1949 for oats; Grafius et al. 1952 for barley). Two other strategies (diallel analysis and generations-means analysis) using information from crosses have been proposed for predicting the best parents and we examine these in detail below.

Using Only Parental Phenotypic Information

As mentioned, when we wish to combine desirable features, such as performance in particular traits or over particular environments, the choice of parental lines can be rather obvious. At a minimum, it is often clear that certain lines should be avoided as parents, and this limits the combinations we may wish to further test. Conversely, there is no general consensus on the best approach for choosing parents to achieve transgenic segregation. Langham (1961) proposed the **high-low method**, suggesting that transgressive segregation would generate lines with larger yields (on average) from such crosses as opposed to high-high or low-low cross. The support for this view is mixed. In maize, Johnson and Hayes (1940) found that high $\times$ low and high $\times$ high crosses gave the same proportion of high-yielding hybrids (with the most extreme lines coming from the H $\times$ L crosses), while Green (1948) found that high $\times$ high crosses were best, as did Lonnquist (1968). In wheat, Busch et al. (1974) found that

while the highest performing line came from a high $\times$ low cross, the high $\times$ high cross gave the highest frequency of superior lines. On theoretical grounds, Bailey and Comstock (1976) and Bailey (1980) suggest that when one of the parental lines is significantly superior to the other, then any given derived line is less likely to contain more favorable alleles than the original superior parent. Example 30.8 illustrates their logic.

Example 30.8. Suppose the two strains being crossed differ at 60 unlinked loci of identical effects. By first computing the probability of fixation π for a single locus and then using this as the success parameter in a binomial distribution with parameters $n = 60$ and π , Bailey and Comstock (1976) found the probability of more than k favorable alleles in a line were as follows:

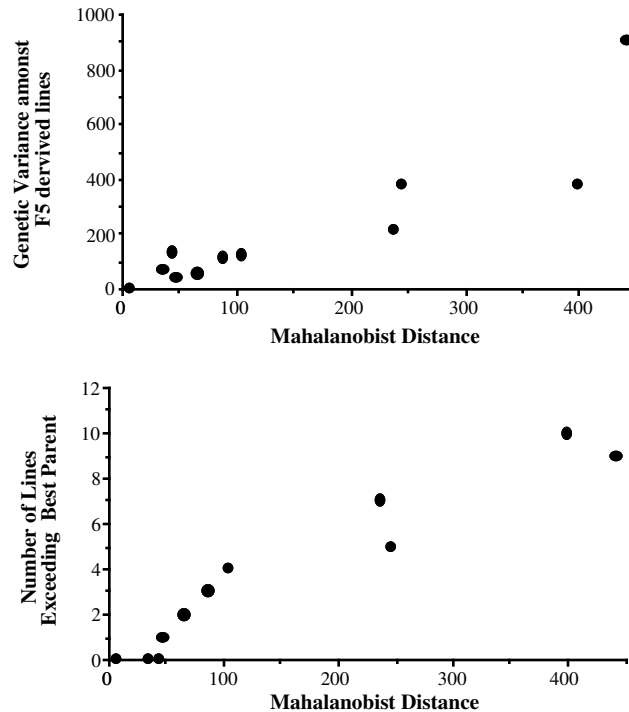
k	No selection	10% selection intensity	
		$h^2 = 0.25$	$h^2 = 0.5$
30	0.4487	0.8660	0.9534
33	0.1831	0.6293	0.8157
36	0.4627	0.3248	0.5443
39	0.0067	0.1057	0.2455
42	0.0005	0.0196	0.0660
45	—	0.0019	0.0094
48	—	—	0.0006

Thus, if the two strains each have 30 favorable loci, then there is a significant probability (87% and 95% under $h^2 = 0.25$ and $h^2 = 0.5$, respectively) that a resulting pure line derived from this cross exceeds either parent. Conversely, suppose one of the strains contains 42 favorable alleles. The resulting probabilities of a line exceeding the best parent are 2% (for $h^2 = 0.25$) and 6.6% (for $h^2 = 0.5$). The average number of favorable alleles is the average of both parents. If one parent far exceeds this average, only a small fraction of the resulting F_∞ lines will equal or exceed this parent.

In an attempt to provide a more quantitative approach to choosing parents solely on their phenotypes, Bhatt (1970, 1973) suggested a multivariate method that does not require crossing (but does require lines are replicated so that genetic and environmental variances can be estimated). Bhatt's idea is to cross those parents that are genetically most distinct, and these extreme parents are chosen as follows: First, we measure a number of potentially informative characters for the trait of interest (for example, for yield, related components such as tiller number, seed weight, and number are also measured). The idea here is that a particular component may have a much higher heritability than the full trait of interest and hence respond better to selection. Next, one computes the genetic variances and covariances for these traits in a sample of pure lines, and the **Mahalanobi's D^2 statistic**, a multivariate measure of distance (see Appendix 5), is then used to pick those parents that are the farthest from each other (the most distant in the geometry defined by the genetic variance-covariance structure).

Example 30.9. Bhatt (1973) measured yield per acre and several related yield components (days to ear emergence, height, number of tillers, yield per plant, and kernel weight) in 40 genotypes of bread wheat (*Triticum aestivum* L.) from different regions of Australia. For a

data set consisting of 11 crosses, the correlation between genetic distance between parents (measured by Mahalanobi's D^2 statistic) and the genetic variation seen among 27 measured derived F_5 lines (from each cross) was highly significant (Spearman's rank correlation of 0.864). Likewise, the four lines with the largest D^2 values had the most derived lines exceeding their parents, and again the correlation between number of transgenic lines and D^2 was striking. Both data sets are plotted below.



Using Information from Crosses: Diallels and DHs

While one expects that information from crossing parental lines is far more predictive than simply using the phenotypic values of the parents, this is not always the case. For example, while the **hybrid potential** of a particular cross can be immediately assessed from the F_1 , our concern is the **inbreeding potential** of the resulting lines (the F_∞), and these may be poorly correlated with the F_1 .

The most straightforward solution is to use doubled-haploid lines when these can be generated. DH lines immediately represent the completely inbred offspring from a hybrid. There is the minor complication that linkage can result in a different distribution of lines than F_∞ lines produced by inbreeding. However, in most cases, a set of DH lines from a hybrid provides a quick assessment of parental potential. Reinbergs et al. (1976) used 100 DH lines per cross, but by random elimination of the data, they found that 20 lines gave essentially the same amount of signal in predicting the most superior crosses.

An alternative suggestion is to use a diallel analysis (all pairwise crosses of parents, see LW Chapter 20). Using this approach, Whitehouse et al. (1958) and Lupton (1961) recommend a strategy that amounts to crossing the two parents with the largest general combining abilities (GCAs), as (everything else being equal) these contain the most additive factors. Major limitations of a diallel design is that it requires $n(n-1)/2$ crosses for n tested parents,

and just how efficient diallels are at predicting F_∞ lines remains unclear. Another difficulty is that when epistasis is present, estimates of GCAs are biased (Matzinger and Kempthorne 1956, Baker 1978).

Example 30.10. Thurling and Ratinam (1987) compared three methods of parent selection for yield improvement in the cowpea (*Vigna unguiculata*) in western Australia. These were: selection based simply on parental yield, selection based on the estimated GCA of the parents (requiring parental and F_1 data), and selection based on the F_2 lines. A total of 10 potential parents were tested, resulting in a total of 45 crosses. Seven of the 45 F_2 lines had a higher yield than the best of the 10 original parents, while 12 had yields below the worst parent. The F_2 yield from the cross between the lines (G and C) with the two largest GCAs was exceeded by 10 of the other F_2 lines. Indeed, C and G had a rather significant negative interaction, with 6 of the other 8 F_2 's involving C having higher values than the $C \times G$ cross (only one of the other F_2 's involving G had a higher value). Hence, GCA was a poor predictor of F_2 yield. When F_3 bulks (selfed collections of the F_2 s planted in bulk) were examined, yield was higher for bulks whose parents were chosen by their F_2 values than by their GCA values (10.75 gr/plot vs. 8.78 gr/plot), but the difference was not significant. Both the GCA and F_2 selection schemes gave larger F_3 bulk yields than the cross chosen simply from parental yield (6.44 gr/plot).

Using Information from Crosses: The Jinks-Pooni Method

A different approach to estimating the expected distribution of the F_∞ is based on joint-scaling tests (also called generations-means analysis). The basis of such tests (reviewed in LW Chapter 9; also see our Chapter 17) are to use a series of crosses to estimate both composite genetic effects underlying the difference between two lines (such as additive and additive $\times$ additive contributions) as well as the genetic variances generated in a cross of such lines. The initial suggestion was by Jinks and Perkins (1972) for predicting the expected range of the distribution, with a much more complete treatment suggested by Jinks and Pooni (1976).

Let α_1^c and α_2^c denote the composite additive and additive $\times$ additive effects, defined with respect to the parental lines being crossed, see LW Chapter 9. (A brief notational aside. We use the notation in LW Chapter 9 (the Edinburgh school, e.g., Falconer 1989), while most of the literature discussing this test uses the Birmingham school notation (e. g., Mather and Jinks 1982), where $[d]$ and $[i]$ denote the additive and additive $\times$ additive effects, and $D/2$ the additive genetic variance.)

If no genotype-environment interaction is present and epistasis is negligible, then Jinks and Pooni (1976) show that the probability that any given F_∞ line exceeds either parent (i.e., shows transgressive segregation) is given by the two-tailed unit normal probability value corresponding to $\alpha_1^c/(\sqrt{2}\sigma_A)$. In particular, if the mean of $P_1 > P_2$, then

$$\Pr(F_\infty > P_1) = \Pr(F_\infty < P_2) = \Pr\left(U > \frac{\alpha_1^c}{\sqrt{2}\sigma_A}\right) \quad (30.19)$$

where U is a unit normal. This assumes an approximate normal distribution of the F_∞ lines and that each is measured with sufficient accuracy that we can ignore any sampling error (i.e., σ_e^2 is small). In such cases, the between-line variance is $2\sigma_A^2$ (assuming σ_D^2 and σ_{ADI} can be ignored, as is the case for a pure-line cross). Under strict additivity, equal percentages of lines are expected to exceed the best parents and under perform the worst parent. If there is heterosis ($F_1 > P_1, P_2$) then

$$\Pr(F_\infty > F_1) = \Pr\left(U > \frac{\delta_1^c}{\sqrt{2}\sigma_A}\right) \quad (30.20)$$

where δ_i^c is the dominance composite effect ($[h]$ in the Birmingham notation). Finally, if there is epistasis, this introduces skew in the distribution of F_∞ lines (Snape and Riggs 1975), but an approximate solution (assuming only additive $\times$ additive epistasis) is given by

$$\Pr(F_\infty > P_1) = \Pr\left(U > \frac{\alpha_1^c + \alpha_2^c}{\sqrt{2\sigma_A^2 + 4\sigma_{AA}^2}}\right) \quad (30.21a)$$

and

$$\Pr(F_\infty < P_2) = \Pr\left(U < \frac{-\alpha_1^c + \alpha_2^c}{\sqrt{2\sigma_A^2 + 4\sigma_{AA}^2}}\right) \quad (30.21b)$$

Thus, with epistasis, these probabilities differ. Although these expressions formally involve σ_{AA}^2 , in principle they are applied using only the estimate for σ_A^2 . Complications and extensions examined by Pooni et al. (1977), Pooni and Jinks (1978, 1979, 1981) and Jinks and Pooni (1982).

To apply these expressions, a joint-scaling (= generations mean) test is applied to both estimate the composite effects (α_i^c , etc.) and test the model fit (Example 30.11). This is done using a set of crosses (all grown in the same environment, i.e., the same location and year) and then using least-squares to fit successfully more complete models (additive, additive + dominance, additive + dominance + additive $\times$ additive epistasis, etc.) until the addition of a new parameter does not significantly improve model fit (see LW Example 9.2). The basic set of crosses used are both parents, F_1 , F_2 , and both backcrosses ($B_1 = F_1 \times P_1$ and $B_2 = F_1 \times P_2$).

A more difficult issue is to obtain a clean estimate of σ_A^2 . While this can be done using the F_2 and backcross generations alone, the **TTC design** using the three F_2 **triple test crosses** ($L_1 = P_1 \times F_2$, $L_2 = P_2 \times F_2$, and $L_3 = F_1 \times F_2$) provides a cleaner estimate of σ_A^2 (Keasery and Jinks 1968, Jinks and Perkins 1970, Pooni et al 1978, Pooni and Jinks 1979). Jinks and Pooni (1980) also suggest that a quick and dirty approach is to use F_3 s to estimate the additive variance by twice the variance of F_3 family means. This includes a contribution from the dominance variance ($\sigma_D^2/8$), but the bias is often expected to be small. Indeed, Tapsell and Thomas (1981) found similar results for yield in Barley when estimating the additive variance by the TTC and F_3 designs. It must be stressed that all generations in the design need be grown in the same year. Caligari et al (1985) found (for Barley) that a TTC analysis performed in one year was a very poor predictor of another year's performance.

Example 30.11. Snape (1982) examined yield and four related components in two crosses of wheat. Through the use of stored seed, both parents, the first three filial generations, both backcrosses, and the three triple test crosses involving the F_2 were all grown at the same time. The data for yield and spikelet number for one of the crosses (Hobbit sib $\times$ Sava) is given below.

Yield (gm/plant)									
P_1	P_2	F_1	F_2	F_3	B_1	B_2	L_1	L_2	L_3
20.26	19.60	28.82	27.30	24.75	24.82	27.56	27.70	28.87	27.94
Spikelet number/ear									
P_1	P_2	F_1	F_2	F_3	B_1	B_2	L_1	L_2	L_3
20.54	21.34	23.22	21.63	22.32	22.06	22.40	22.41	22.19	22.52

Joint-scaling tests (LW Chapter 9) were applied, and the additive + dominance model provided an adequate fit for yield, while no model (including those with epistasis) provided a suitable

fit for spikelet number (the residual sums of squares exceeded the critical χ^2 value). The estimates for composite effects of interest were:

Yield (gm/plant)		Spikelet number/ear		
α_1^c	δ_1^c	α_1^c	δ_1^c	α_2^c
0.67	13.98	0.38	3.07	0.14

An ANOVA using the triple test data was used to estimate twice the additive variance, giving values of 22.79 and 3.485 for yield and spikelet number (respectively). Thus, for yield the expected probability of an F_∞ exceeding the best parent is

$$\Pr(F_\infty > P_1) = \Pr(U > 0.67/\sqrt{22.79}) = \Pr(U > 0.140) = 0.44$$

Likewise, there is a 44.0% chance of being below the inferior parent. A total of 116 lines were scored, and 64 exceeded the best parent while 40 were under the poorer parent. This difference is not significantly different from the exceeded number of 51 for each class. Similarly, the probability that a random F_∞ line from this cross exceeds the F_1 is

$$\Pr(F_\infty > F_1) = \Pr(U > 13.98/\sqrt{22.79}) = \Pr(U > 2.93) = 0.002$$

Thus, to have a 95% probability of obtaining an F_∞ line that outperforms the F_1 , recall Equation 30.13, we need to score at least $\ln(1-0.95)/\ln(1-0.002) = 1496$ lines!

Turning to spikelet number, we must account for epistasis. The expected frequencies of transgressive segregants thus become

$$\Pr(F_\infty > P_1) = \Pr\left(U > \frac{0.38 + 0.14}{\sqrt{3.458}}\right) = \Pr(U > 0.280) = 0.39$$

and

$$\Pr(F_\infty < P_2) = \Pr\left(U < \frac{-0.38 + 0.14}{\sqrt{3.458}}\right) = \Pr(U < -0.129) = 0.45$$

Again testing 116 lines, we would expect 45 lines larger than P_1 and 53 smaller than P_2 . What was observed were 71 larger and 28 smaller, a significant difference (at the 5% level). Snape suggests that this discrepancy is not unexpected, given that no composite model from the joint scaling test fully fitted the data.

Number of Crosses vs. Number of Lines per Cross

There are two areas of risk when trying to choose elite genotypes: not choosing the best crosses and not choosing the best lines within a cross. What is the optimal strategy for balancing these? If a total of N lines are to be used, the breeder must weight the number of crosses m against the number of lines n from each cross (with $N = mn$). In the absence of any trait information, how should the breeder choose m and n so as to minimize the risk that no elite genotypes exist in the lines being examined? This question has been examined by Yonezama and Yamagata (1978), Weber (1979), and (most generally) Wricke and Weber (1986). Wricke and Weber (1986) show that, in the absence of any phenotypic information on the crosses, the risk is minimized by maximizing the number of crosses (taking $m = N$ and hence $n = 1$). There is also a tradeoff between the within and between cross variance. Widecrosses are expected to have somewhat intermediate performance but a high within-cross segregation variance, while crosses between closely-related lines may have higher mean cross values

but also smaller within-cross variance, reflecting fewer genetic differences than between the parents of wide-species crosses.

RECURRENT SELECTION

Silvela and Diez-Barra (1985) compared recurrent selection to self-pollinating schemes, finding a clear superiority of RS over selfers when negative disequilibrium is present (favorable alleles in repulsion) and/or epistasis.

Literature Cited

- Allard, R. W. 1960. *Principles of plant breeding*, Wiley, New York. [30]
- Anderson, J. A. D., and H. W. Howard. 1981. Effectiveness of selection in early stages of potato breeding programmes. *Potato Res.* 24: 289–299. [30]
- Athwal, D. S. 1971. Semidraft rice and wheat in global food needs. *Quart. Rev. Biol.* 46: 1–34. [30]
- Atkins, R. E. 1964. Visual selection for grain yield in barley, *Crop Science* 4: 494–496. [30]
- Atkins, R. E., and H. C. Murphy, 1949. Evaluation of yield potentialities of oat crosses from bulk hybrid tests. *Agron. J.* 41: 41–45. [30]
- Bailey, T. B. 1980. Selection limits in self-fertilizing populations following the cross of homozygous lines. In E. Pollak, O. Kempthorne, and T. B. Bailey, Jr. (eds.), *Proceedings of the international conference on quantitative genetics*, pp. 399–412. Iowa State Univ. Press, Ames. [30]
- Bailey, T. B., and R. E. Comstock. 1976. Linkage and the synthesis of better genotypes in self-fertilizing species. *Crop Science* 16: 363–370. [30]
- Baker, R. J. 1968. Extent of intermating in self-pollinated species necessary to counteract the effects of genetic drift. *Crop Sci.* 8: 547–550. [30]
- Baker, R. J. 1978. Issues in diallel analysis. *Crop Sci.* 18: 533–536. [30]
- Bal, B. S., C. A. Suneson, and R. T. Ramage. 1959. Genetic shift during 30 generations of natural selection in barley. *Agron. J.* 51: 555–557. [30]
- Bartlett, M. S. 1978. Nearest-neighbour models in the analysis of field experiments (with Discussion). *J. Royal Stat. Soc. Series B* 40: 147–174. [30]
- Besag, J., and R. Kempton. 1986. Statistical analysis of field experiments using neighbouring plots. *Biometrics* 42: 231–251. [30]
- Bhatt, G. M. 1970. Multivariate analysis approach to selection of parents for hybridization aiming at yields improvement in self-pollinated crops. *Aust. J. Agric. Res.* 21: 1–7. [30]
- Bhatt, G. M. 1973. Comparison of various methods of selecting parents for hybridization in common bread wheat (*Triticum aestivum* L.). *Aust. J. Agric. Res.* 24: 457–464. [30]
- Bisen, M. S., S. P. Singh, and S. K. Rao. 1985. Effectiveness of selection methods in chickpea *Cicer arietinum* under different environments. *Theor. Appl. Gene.* 70: 661–666. [30]
- Bliss, F. A., and C. E. Gates. 1968. Directional selection in simulated populations of self-pollinated plants. *Aus. J. Biol. Sci.* 21: 705–719. [30]
- Branch, W. D., J. S. Kirby, J. C. Wynne, C. C. Holbrook, and W. F. Anderson. 1991. Sequential vs. pedigree selection method for yield and leafspot resistance in peanut. *Crop Science* 31: 274–276. [30]
- Briggle, L. W., and O. E. Vogel. 1968. Breeding short stature, disease resistant wheats in United States. *Euphytica* Suppl 1: 107–130. [30]
- Boerma, H. R., and R. L. Cooper. 1975. Comparison of three selection procedures for yield in soybeans. *Cop Science* 15: 225–229. [30]
- Bonnet, O. T., and W. B. Bever. 1947. Head-hill method of planting head selections of small grains. *J. Amer. Soc. Agron.* 39: 442–445. [30]
- Bos, I. 1977. More arguments against intermating F_2 plants of a self-fertilizing crop. *Euphytica* 26: 33–46. [30]
- Bradu, D., and K. R. Gabriel. 1978. The biplot as a diagnostic tool for models of two-way tables. *Technometrics* 20: 47–68. [30]
- Briggs, K. G., and L. H. Shebeski. 1970. Visual selection for yielding ability of F_3 lines in a hard red spring wheat breeding program. *Cop Science* 105: 400–402. [30]

- Briggs, K. G., and L. H. Shebeski. 1971. Early generation selection for yield and breadmaking quality of hard red spring wheat. *Euphytica* 20: 53–463. [30]
- Busch, R. H., J. C. Janke, and R. C. Frohberg. 1974. Evaluation of crosses among high and low yielding parents of spring wheat (*Triticum aestivum* L.) and bulk prediction of line performance. *Crop Science* 14: 47–50. [30]
- Byron, D. F., and J. H. Orf. 1991. Comparison of three selection procedures for development of early-maturing soybean lines. *Crop Science* 31: 656–660. [30]
- Caligari, P. D. S., W. Powell, and J. L. Jinks. 1985. The use of doubled haploids in barley breeding. 2. An assessment of univariate cross prediction methods. *Heredity* 54: 353–358. [30]
- Carmer, S. G. 1976. Optimal significance levels for application of the least significant difference in crop performance trails. *Crop Sci.* 16: 95–99. [30]
- Carmer, S. G., and M. R. Swanson. 1971. Detection of differences between means: a Monte Carlo study of five pairwise multiple comparison procedures. *Agron. J.* 63: 940–945. [30]
- Carmer, S. G., and M. R. Swanson. 1973. An evaluation of ten pairwise multiple comparison procedures by Monte Carlo methods. *J. Amer. Stat. Assoc.* 68: 66–74. [30]
- Casali, V. W. D., and E. C. Tigchelaar. 1975a. Breeding progress in tomatoe with pedigree selection and single seed descent. *J. Amer. Soc. Hort. Sci.* 100: 362–364. [30]
- Casali, V. W. D., and E. C. Tigchelaar. 1975b. Computer simulation studies comparing pedigree, bulk, and single seed descent in self pollinated populations. *J. Amer. Soc. Hort. Sci.* 100: 3642–367. [30]
- Choo, T. M. 1981. Doubled haploids for studying the inheritance of quantitative characters. *Genetics* 99: 525–540. [30]
- Compton, W. A. 1968. Recurrent selection in self-pollinated crops without extensive crossing. *Crop Science* 8: 773. [10, 11]
- Cooper, M., and I. H. DeLacy. 1994. Relationships among analytical methods used to study genotypic variation and genotype-by-environment interaction in plant breeding multi-environment experiments. *Theor. Appl. Genet.* 88: 561–572. [30]
- Cornelius, P. 1993. Statistical tests and retention of terms in the additive main effects interaaction model for cultivar trails. *Crop Sci.* 33: 1186–1193. [30]
- Cornelius, P., M. Seyedsadr, and J. Crossa. 1992. Using the shifted multiplicative model to search for “separability” in crop cultivar trails. *Theor. Appl. Genet.* 84: 161–172. [30]
- Cregan, P. B., and R. H. Busch. 1977. Early generation bulk hybrid yield testing of adapated hard red spring wheat crosses. *Crop Science* 17: 887–891. [30]
- Crossa, J. 1990. Statistical analysis of multilocation trails. *Advances in Agronomy* 44: 55–85. [30]
- Crossa, J., P. N. Fox, W. H. Pfeiffer, S. Rajaram, and H. G. Guach, Jr. 1991. AMMI adjustment for statistical analysis of an international wheat yeild trail. *Theor. Appl. Genet.* 81: 27–27. [30]
- Crossa, J., H. G. Guach, Jr., and R. W. Zobel. 1990. Additive main effects and multiplicative interaction analysis of two international maize cultivar trails. *Crop Sci.* 30: 493–500. [30]
- Dahiya, B. N., and V. P. Singh. 1986. Comparison of single seed descent, selective intermating and mass selection for seed size in greengram (*Vigna radiata* (L.) Wilczek). *Theor. Appl. Genet.* 72: 678–681. [30]
- Dendrious, A. D. 1931. The question of shape in field experimentation. *Intl. Assn.. Plant Breeders Bull.* 4: 106–113. [30]
- De Pauw, R. M., and L. H. Shebeski. 1973. An evaluation of an early generation yield testing procedure in *Triticum aestivum*. *Can. J. Plant Sci.* 53: 465–470. [30]
- Empig, L. T., and W. R. Fehr. 1971. Evaluation of methods for generation advance in bulk hybrid soybean populations. *Crop Sci.* 11: 51–54. [30]
- England, F. 1977. Response to family selection based on replicated trails. *J. Agri./Sci.* 88: 127–134. [8, 9]

- England, F. J. W. 1981. Population size and the chance of success in a barley breeding programme based on the use of doubled haploids (DH) or single-seed descent (SSD). In *Barley Genetics IV. Proceedings of the Fourth International Barley Genetics Symposium*, pp. 176–178. Edinburgh. [30]
- Falconer, D. S. 1989. *Introduction to quantitative genetics*. 3rd Ed. Longman Sci. and Tech., Harlow, UK. [30]
- Fasoulas, A.. 1973. A new approach to breeding superior yielding varieties. Dept. Genetics and Plant Breeding, Aristotelian Univ. of Thessalonika (Greece) Publ. No. 3. 114pp. [30]
- Finney, D. J. 1958. Plant selection for yield improvement. *Euphytica* 7: 83–106. [30]
- Fowler, W. L., and E. G. Heyne. 1955. Evaluation of bulk hybrid tests for predicting performance of pure line selection in hard red winter wheat. *Agron. J.* 47: 430–434. [30]
- Frey, K. J. 1954. The use of F_2 lines in predicting the performance of F_3 selection in two barley crosses. *Agron. J.* 46: 541–544. [30]
- Gabriel, K. R. 1971. Biplot display of multivariate matrices with applications to principal component analysis. *Biometrika* 58: 453–467. [30]
- Gabriel, K. R. 1978. Least squares approximation of matrices by additive and multiplicative models. *J. R. Stat. Soc. Ser. B* 40: 186–196. [30]
- Gardner 1961 **strated mass selection ref – add.** [30]
- Garner, H. V., and J. W. Weil. 1939. Standard errors of field plots at Rothamsted and outside centers. *Emp. J. Exp. Agr.* 7: 369–379. [30]
- Gauch, H. G., Jr. 1988. Model selection and validation for yield trials with interaction. *Biometrics* 44: 705–715. [30]
- Gauch, H. G., Jr. 1992. *Statistical analysis of regional yield trials: AMMI analysis of factorial designs*. Elsevier, Amsterdam, the Netherlands. [30]
- Gauch, H. G., Jr., and R. W. Zobel. 1988. Predictive and postdictive success of statistical analysis of yield trials. *Theor. Appl. Genet* 76: 1–10. [30]
- Gauch, H. G., Jr., and R. W. Zobel. 1989. Accuracy and selection success in yield trial analysis. *Theor. Appl. Genet* 77: 473–481. [30]
- Gauch, H. G., Jr., and R. W. Zobel. 1990. Inputting missing yield trial data. *Theor. Appl. Genet* 79: 753–761. [30]
- Gauch, H. G., Jr., and R. W. Zobel. 1996. Optimal replication in selection experiments. *Crop Sci.* 36: 838–843. [30]
- Gibbons, J. D., I. Okin, and M. Sobel. 1997. *Selecting and ordering populations: A new statistical methodology*. Wiley, New York. [30]
- Goodman, L. A., and S. J. Haberman. 1990. The analysis of non-additivity in two-way analysis of variance. *J. Am. Stat. Assoc.* 85: 139–145. [30]
- Gollob, H. F. 1968. A statistical model which combines features of factor analysis and analysis of variance techniques. *Psychometrika* 33: 73–115. [30]
- Gomez, K. A., and A. A. Gomez. 1984. *Statistical procedures for agricultural research*. John Wiley and Sons, New York. [30] **GET!!**
- Goulden, C. H. 1939. Problems in plant selection, in R. C. Burnett (ed.) *Proceedings of the 7th International Genetics Congress.*, Edinburgh, pp. 132–133. Cambridge University Press, Cambridge. [30]
- Grafius, J. E. 1965. Short cuts in plant breeding. *Crop Science* 5: 377. [30]
- Grafius, J. E., W. L. Nelson, and V. A. Dirks. 1952. The heritability of yield in barley as measured by early generation bulk progenies. *Agron. J.* 44: 253–257. [30]
- Green, J. M. 1948. The inheritance of combining ability in maize hybrids. *J. Amer. Soc. Agron.* 40: 58–63. [30]

- Hallauer, A. R., and J. B. Miranda. 1981. *Quantitative genetics in maize breeding*. Iowa State Univ. Press, Ames. [30]
- Hamblin, J., and J. G. Rowell. 1975. Breeding implications of the relationship between competitive ability and pure culture yield in self-pollinated grain crops. *Euphytica* 24: 221–228. [30]
- Hamblin, J., and A. M. Evans. 1976. The estimation of cross yield using early generation and parentla yields in dry beans (*Phaseolus vulgaris* L.). *Euphytica* 25: 515–520. [30]
- Hanson, P. R., G. Jenkins, and B. Westcott. 1979. Early generation selection in a cross of spring barley. *Z. Pflanzenzüchtg.* 83: 64–80. [30]
- Hanson, W. D., R. C. Leffel, and H. W. Johnson. 1962. Visual discrimination for yield among soybean populations. *Crop Sci.* 9: 93–96. [30]
- Harlan, H. V., and M. L. Martini. 1929. A composite hybrid mixture. *J. Amer. Soc. Agron.* 21: 487–490. [30]
- Harlan, H. V., and M. L. Martini. 1938. The effect of natural selection in a mixture of barley varieties. *J. Agr. Res.* 57: 189–199. [30]
- Harlan, H. V., M. L. Martini, and H. Stevens. 1940. A study of methods of barley breeding. USDA tech. Bull. No. 720. [30]
- Harrington, J. B. 1940. Yielding capacity of wheat crosses as indicated by bulk hybrid tests. *Canadian J. Res. C* 18: 578–583. [30]
- Hill, R. R., Jr., and J. L. Rosenberger. 1985. Methods for combining data from germplasm evaluation trails. *Crop Sci.* 25: 467–470. [30]
- Ikehashi, H., and H. Fujimaki. 1980. Modified bulk population method for rice breeding. In, *Innovative approaches to rice breeding*, pp.163–182. International Rice Research Institute, Los Banos, Philippines. [30]
- Immer, F. R. 1941. Relation between yielding ability and homozygosis in barley crosses. *J. Am. Soc. Agron.* 33: 200–206. [30]
- Ivers, D. R., and W. R. Fehr. 1978. Evaluation of the pure-family method for cultivar development. *Crop Science* 18: 541–544. [30]
- Jansen, R. C. 1992. On the selection for specific genes in doubled haploids. *Heredity* 69: 92–95. [30]
- Jansen, R. C., and J. J. Jansen. 1990. On the selection for specific genes by single seed descent. *Euphytica* 51: 131–140. [30]
- Jennings, P. R., and R. M. Herrera. 1968. Studies on competition in rice. II. Competition in segregating populations. *Evolution* 22: 332–336. [30]
- Jinks, J. L., and J. M. Perkins. 1970. A general method for detection of additive, dominance and epistatic components of variance. *Heredity* 25: 419–429. [30]
- Jinks, J. L., and J. M. Perkins. 1972. Predicting the range of inbred lines. *Heredity* 28: 399–403. [30]
- Jinks, J. L., and H. S. Pooni. 1976. Predicting the properties of recombinant inbred lines derived by single seed descent. *Heredity* 36: 253–266. [30]
- Jinks, J. L., and H. S. Pooni. 1980. Comparing predictions of mean performance and environmental sensitivity of recombinant inbred lines based upon F_3 and triple test cross families. *Heredity* 45: 305–312. [30]
- Jinks, J. L., and H. S. Pooni. 1981a. Comparative results of selection in the early and late stages of an inbreeding programme. *Heredity* 46: 1–7. [30]
- Jinks, J. L., and H. S. Pooni. 1981b. Properties of pure-breeding lines produced by dihaploidy, single seed descent and pedigree breeding. *Heredity* 46: 391–395. [30]
- Jinks, J. L., and H. S. Pooni. 1982. Predicting the properties of pure lines extractable from a cross in the presence of linkage. *Heredity* 49: 265–270. [30]

- Jinks, J. L., and H. S. Pooni. 1984. Comparison of inbred lines produced by single seed descent and pedigree inbreeding. *Heredity* 53: 299–308. [30]
- Johnson, J.J., J. R. Alldredge, S. E. Ullrich, and O. Dangi. 1992. Replacement of replications with additional locations for grain sorghum cultivar evaluation. *Crop Sci.* 32: 43–46. [30]
- Kalton, R. R. 1948. Breeding behavior at successive generations following hybridization in soybeans. *Iowa State Coll. Agr. Expt. Station Research Bull.* 358. pp. 671–732. [30]
- Kearsey, M. J., and J. L. Jinks. 1968. A general method for the detection of additive, dominance, and epistatic components of variance. I. Theory. *Heredity* 23: 403–409. [30]
- Kempton, R. A. 1984. The use of biplots in interpreting variety by environment interactions. *J. Agric. Sci.* 103: 123–135. [30]
- Kermicle, J. L. 1969. Androgenesis conditioned by a mutation in maize. *Science* 166: 1422–1424. [30]
- Khalifa, M. A., and C. O. Qualset. 1975. Inter-genotypic competition between tall and dwarf wheats: II. in hybrid bulks. *Crop Science* 15: 640–644. [30]
- Knott, D. R. 1972. Effects of selection for F_2 plant yield on subsequent generations in wheat. *Can. J. Plant Sci.* 52: 721–726. [30]
- Knott, D. R., and J. Kumar. 1975. Comparison of early generation yield testing and a single seed descent procedure in wheat breeding. *Crop Sci.* 15: 295–299. [30]
- Kulshrestha, V. P. 1989. A modified pedigree method of selection. *Crop Science* 78: 173–176. [30]
- Kwon, S. H., and J. H. Torrie. 1964. Visual discrimination for yield in two soybean populations. *Crop Sci.* 4: 287–290. [30]
- LeClerg, E. L. 1966. Significance of experimental design in plant breeding. in K. J. Frey (ed) *Plant breeding*, pp. 243–313. Iowa State University press, Ames, Iowa. [30]
- LeClerg, E. L., W. H. Leonard, and A. G. Clark. 1962. *Field Plot Technique*, 2nd ed. Bruggess, Minneapolis. [30]
- Leffel, R. C., and W. D. Hanson. 1961. Early generation testing of diallel crosses of soybeans. *Crop Sci.* 1: 169–174. [30]
- Little, T. M. and F. J. Hillis. 1978. *Agricultural experimentation: Design and analysis*. John Wiley, New York. [30] **GET**
- Lonquist, J. H. 1968. Further evidence on testcross versus line performance in maize. *Crop Science* 8: 50–53. [30]
- Lowe, H. H. 1927. A program for selecting and testing small grains in successive generations following hybridization. *American Society of Agronomy Journal* 19: 705–712. [30]
- Luedders, V. D., L. A. Ducloes, and A. L. Matson. 1973. Bulk, pedigree, and early generation testing breeding methods compared in soybeans. *Crop Sci.* 13: 363–364. [30]
- Lupton, F. G. H. 1961. Studies in the breeding of self-pollinating cereals. 3. Further studies in cross-prediction. *Euphytica* 10: 209–224. [30]
- Lupton, F. G. H., and R. N. H. Whitehouse. 1957. Studies on the breeding of self-pollinating cereals. I. Selection methods in breeding for yield. *Euphytica* 6: 169–184. [30]
- Ma, R. H., and J. B. Harrington. 1948. The standard errors of different designs of field experiments at the University of Saskatchewan. *Sci. Agr.* 28: 461–474. [30]
- MacNeil, M. D., D. D. Kress, A. E. Flower, and R. L. Blackwell. 1984. Effects of mating systems in Japanese quail. 2. genetic parameters, response and correlated response to selection. *Theor. Appl. Genet.* 67: 407–412. [30]
- Mahmud I., and H. H. Kramer. 1951. Segregation for yield, height, and maturity following a soybean cross. *Agron. J.* 43: 605–609. [30]
- Mandel, J. 1971. A new analysis of variance model for non-additive data. *Technometrics* 13: 1–8. [30]

- Marani, A. 1975. Simulation of response of small self-fertilizing populations to selection for quantitative traits: I. effect of number of loci, selection intensity and initial heritability under conditions of no dominance. *Theor. Appl. Genet.* 46: 221–231. [30]
- Mather, K., and J. L. Jinks. 1982. *Biometrical genetics*. 3rd Ed. Chapman and Hall, NY. [30]
- Matzinger, D. F., and O. Kempthorne. 1956. The modified diallel table with partial inbreeding and interactions with environment. *Genetics* 41: 822–833. [30]
- Mayo, O. 1987. *The theory of plant breeding, 2nd Edition*. Clarendon, Oxford. [8A]
- McKenzie, R. I. H., and J. W. Lambert. 1961. A comparison of F_3 lines and their related F_6 lines in two barley crosses. *Crop Sci.* 1: 246–249. [30]
- Mercer, W. B., and A. D. Hall. Experimental error in field trials. *J. Agr. Sci.* 4: 107–127. [30]
- Meredith, W. R., and R. R. Bridge. 1971. Breakup of linkage blocks in cotton, *Gossypium hirsutum* L. *Crop Sci.* 11: 695–698. [30]
- Miller, P. A., and J. O. Rawlings. 1967. Breakup of initial linkage blocks through intermating in a cotton breeding population. *Crop Sci.* 7: 199–204. [30]
- Mitchell, J. W., R. J. Baker, and D. R. Knott. 1982. Evaluation of honeycomb selection for single plant yield in Durum Wheat. *Crop Science* 22: 840–843. [30]
- Nachit, M. N., G. Naachit, H. Keata, H. G. Gauch Jr., and R. W. Zobel. 1992. Use of AMMI and linear regression models to analyze genotype-environment interactions in durum wheat. *Theor. Appl. Genet.* 83: 597–601. [30]
- Nitzsche, W., and G. Wenzel. 1977. *Haploids in plant breeding*. Paul Parey, Hamburg. [30]
- Nordskog, A. W. and J. Hardiman. 1980. Inbreeding depression and natural selection as factors limiting progress from selection in poultry. In A. Roberston, (ed.), *Selection experiments in laboratory and domestic animals*, pp. 91–99. Commonwealth Agricultural Bureaux. [30]
- Nilsson-Ehle, H. 1910. Arbetena med hvete och havre vid Svalöf under år 1909. *Sv. Utsädesf. Tidskr.* 20: 332–353. [30]
- Park, S. J., E. J. Walsh, E. Reinbergs, L.S. P. Song, and K. J. Kasha. 1976. Field performance of doubled haploid barley lines in comparison with lines developed by pedigree and single seed descent methods. *Can. J. Plant Sci.* 56: 467–474. [30]
- Patterson, H. D., V. Silvery, M. Talbot, and S. T. C. Weatherup. 1977. Variability of yields of cereal varieties in U.K. trials. *J. Agric. Sci.* 89: 239–245. [30]
- Pearce, S. C. 1983. *The agricultural field experiment: A statistical examination of theory and practice*. John Wiley, New York. [30] **GET**
- Pearce, S. C., G. M. Clarke, G. V. Dyke, and R. E. Kempson. 1988. *A manual of crop experimentation*. Oxford University Press, Oxford, England. [30] **GET**
- Pederson, D. G. 1974. Arguments against intermating before selection in a self-fertilising species. *Theor. Appl. Genet.* 45: 157–162. [30]
- Petersen, R. G. 1994. *Agricultural field experiments: Design and Analysis*. Marcel Dekker, New York. [30] **GET**
- Piepho, H.-P. 1994. Best linear unbiased prediction (BLUP) for regional yield trials: a comparison to additive main effects and multiplicative interaction (AMMI) analysis. *Theor. Appl. Genet.* 89: 647–654. [30]
- Pierce, L. C. 1977. Impact of single seed descent in selecting for fruit size, earliness, and total yield in tomato. *J. Am. Soc. Hort. Sci.* 102: 520–522. [30]
- Ponni, H. S., and J. L. Jinks. 1978. Predicting the properties of recombinant inbred lines derived by single seed descent for two or more characters simultaneously. *Heredity* 40: 349–361. [30]
- Ponni, H. S., and J. L. Jinks. 1979. Sources and biases of the predictors of the properties of recombinant inbred lines derived by single seed descent. *Heredity* 42: 41–48. [30]

- Ponni, H. S., and J. L. Jinks. 1981. Sources of predictions and their reliability in predicting the properties of recombinant inbred lines which can be obtained from a cross by single seed descent. In *Barley Genetics IV. Proceedings of the Fourth International Barley Genetics Symposium*, pp. 73–78. Edinburgh. [30]
- Ponni, H. S., J. L. Jinks, and M. A. Cornish. 1977. The causes and consequences of non-normality in predicting the properties of recombinant inbred lines. *Heredity* 38: 329–338. [30]
- Ponni, H. S., J. L. Jinks, and N. E. M. Jayasekara. 1978. An investigation of gene action and genotype $\times$ environment interaction in two crosses of *Nicotiana rustica* by triple test cross and inbred line analysis. *Heredity* 41: 83–92. [30]
- Powell, W., P. D. S. Caligari, and W. T. B. Thomas. 1986. Comparison of spring barley lines produced by single seed descent, pedigree inbreeding and doubled haploidy. *Plant Breeding* 97: 138–146. [30]
- Raeber, J. G., and C. R. Weber. 1953. Effectiveness of selection for yield in soybean crosses by bulk and pedigree systems of breeding. *Agronomy Journal* 45: 362–366. [30]
- Rasmusson, D. C., and J. W. Lambert. 1961. Variety $\times$ environment interactions in barley variety tests. *Crop Sci.* 1: 261–262. [30]
- Reinbergs, E., S. J. Park, and L. S. P. Song. 1976. Early identification of superior barley crosses by the doubled haploid technique. *Z. Pflanzenzüchtg.* 76: 215–244. [30]
- Riggs, T. J., P. R. Hanson, and N. D. Start. 1981. Modification of the pedigree method in spring barley breeding by incorporating bulk selection in F_4 . In *Barley Genetics IV. Proceedings of the Fourth International Barley Genetics Symposium*, pp. 138–141. Edinburgh. [30]
- Roy, N. N. 1976. Inter-genotypic plant competition in wheat under single seed descent breeding. *Euphytica* 25: 219–223. [30]
- Schwarzbach, E. 1981. The limits of selection in segregating autogamous populations under ideal and under realistic assumptions. In *Barley Genetics IV. Proceedings of the Fourth International Barley Genetics Symposium*, pp. 154–158. Edinburgh. [30]
- Sedcole, J. R. 1977. Number of plants necessary to recover a trait. *Crop Science* 17: 667–668. [30]
- Shebeski, L. H. 1967. Wheat and breeding, in K. F. Nielsen (ed) *Proceedings of the Canadian centennial wheat symposium*, pp. 249–272. Modern Press, Saskatoon, Sask. [30]
- Silvela, L., and R. Diez-Barra. 1985. Recurrent selection in autogamous species under forced random mating. *Euphytica* 34:817–832. [30]
- Simmonds, N. W. 1979. *Principles of crop improvement*, Longman, London. [30]
- Smith, E. L., and J. W. Lambert. 1968. Evaluation of early generation testing in spring barley. *Crop Sci.* 8: 490–492. [30]
- Snape, J. W. 1982. Predicting the frequencies of transgressive segregants for yield and yield components in wheat. *Theor. Appl. Genet* 62: 127–134. [30]
- Snape, J. W., and E. Simpson. 1984. Early generation selection and rapid generation advancement methods in autogamous crops. In W. Lange, A. C. Zeven, and N. G. Hogenboom . (eds.), *Efficiency in Plant Breeding: Proceedings 10th Congress of the European Association for Research in Plant Breeding, EUCARPIA*, pp. 82–86. Wageningen, Netherlands. [30]
- Snape, J. W., and T. J. Riggs. 1975. Genetical consequences of single seed descent in the breeding of self-pollinating crops. *Heredity* 35: 211–219. [30]
- Sneep, J. 1977. Selection for yield in early generations of self-fertilizing crops. *Euphytica* 26: 27–30. [30]
- Sneep, J. 1984. Is selection on quantitative characters between F_3 lines in small grains feasible?. In W. Lange, A. C. Zeven, and N. G. Hogenboom . (eds.), *Efficiency in Plant Breeding: Proceedings 10th Congress of the European Association for Research in Plant Breeding, EUCARPIA*, pp. 87–89. Wageningen, Netherlands. [30]

- Srivastava, R. B., R. S. Paroda, S. C. Sharma, and Md. Yunus. 1989. Genetic variability and advance under four selection procedures in wheat pedigree breeding programme. *Theor. Appl. Gene.* 77: 516–520. [30]
- Stam, P. 1977. Selection response under random mating and under selfing in the progeny of a cross of homozygous parents. *Euphytica* 26: 169–184. [30]
- St.-Pierre, C. A., H. R. Klinck, and F. M. Gauthier. 1967. Early generation selection under different environments as it influences adaptation of barley. *Can. J. Plant Sci.* 47: 507–517. [30]
- Stoskopf, N. C., D. T. Tomes, and B. R. Christie. 1993. *Plant breeding: Theory and practice*. Westview Press, Boulder. [30]
- Stroup, W. W., and D. K. Mulitze. 1991. Nearest neighbor adjusted best linear unbiased predictors. *Am. Stat.* 45: 194–200. [30]
- Suneson, C. A. 1949. Survival of four barley varieties in a mixture. *Agronomy J.* 41: 459–461. [30]
- Suneson, C. A. 1956. An evolutionary plant breeding method. *Agronomy J.* 48: 188–191. [30]
- Talbot, M. 1984. Yield variability of crop varieties in the U. K. *J. Agric. Sci.* 102: 315–321. [30]
- Tapsell, C. R., and W. T. B. Thomas. 1981. Estimating the genetical components for cross-prediction of yield and its components in barley. In *Barley Genetics IV. Proceedings of the Fourth International Barley Genetics Symposium*, pp. 79–83. Edinburgh. [30]
- Taylor, F. W. 1909. The size of experimental plots for field crops. *Proc. Amer. Soc. Agron.* 1: 56–58. [30]
- Tee, T. S., and C. O. Qualset. 1975. Bulk populations in wheat breeding: comparison of single-seed descent and random bulk methods. *Euphytica* 24: 393–405. [30]
- Thurling, N., and M. Ratnam. 1987. Evaluation of parent selection methods for yield improvement of cowpea (*Vigna unguiculata* (L.) Walp.). *Euphytica* 36: 913–926. [30]
- Torrie, J. H. 1958. A comparison of the pedigree and bulk method of breeding soybeans. *Agronomy Journal* 50: 198–200. [30]
- Van der Kely, F. K. 1955. The efficiency of some selection methods in populations resulting from crosses between self-fertilizing plants *Euphytica* 4: 58–66. [30]
- Van Oeveren, A. J., and P. Stam. 1992. Comparative simulation studies on the effects of selection for quantitative traits in autogamous crops: early selection versus single seed descent. *Heredity* 69: 342–351. [30]
- Valentine, J. 1979. The effect of competition and method of sowing on the efficiency of single plant selection for grain yield, yield components and other characters in spring Barley. *Z. Pflanzenzüchtg.* 83: 193–204. [30]
- Voigt, R. L., and C. R. Weber. 1960. Effectiveness of selection methods for yield in soybean crosses. *Agron. J.* 52: 527–530. [30]
- Weber, W. E. 1979. Number and size of cross progenies from a constant total number of plants manageable in a breeding program. *Euphytica* 28: 453–456. [30]
- Weber, W. E. 1984. Selection in early generations. In W. Lange, A. C. Zeven, and N. G. Hogenboom (eds.), *Efficiency in Plant Breeding: Proceedings 10th Congress of the European Association for Research in Plant Breeding, EUCARPIA*, pp. 72–81. Wageningen, Netherlands. [30]
- Weiss, M. G., C. R. Weber, and R. R. Kalton. 1947. Early generation testing in soybeans. *J. Am. Soc. Agron* 39: 791–811. [30]
- Whan, B. R., R. Knight, and A. J. Rathjen. 1982. Response to selection for grain yield and harvest index in F_2 , F_3 and F_4 derived lines of two wheat crosses. *Euphytica* 31: 139–150. [30]
- Whitehouse, R. N.H., J. B. Thompson, and M. A. M. Do. Valle Ribeiro. 1958. Studies on the breeding of self-pollinated cereals. 2. The use of a diallel cross analysis in yield prediction. [30]
- Wilkinson, G. N., S. R. Eckert, T. W. Hancock, and O. Mayo. 1983. Nearest-neighbour (NN) analysis of field experiments (with Discussion). *J. Royal Stat. Soc. Series B* 45: 151–211. [30]

- Williams, W. 1964. *Genetical principles and plant breeding*. Blackwell, Oxford. [30]
- Woods, T. B., and F. J. Stratton. 1910. Interpretation of experimental results. *J. Agr. Sci.* 3: 417–559. [30]
- Wricke, G., and W. E. Weber. 1986. *Quantitative genetics and selection in plant breeding*. Walter de Gruyter and Co., NY. [8, 9]
- Yap, T. C., C. T. Poh, and C. Mak. 1977. Comparative studies on selection methods in long beans. *In Animal Breeding Papers, Third Int. Congr. SABRAO*, Canberra, CSIRO. pp 2c(i) 23-25. [30]
- Yates, F. 1935. Complex experiments. *Suppl. J. Royal Stat. Soc.* 2: 181–247. [30]
- Yonezama, K. and H. Yamagata. 1978. On the number and size of cross combinations in a breeding programme of self-fertilizing crops. *Euphytica* 27: 113-116. [30]
- Yonezama, K., T. Normura, and Y. Sasaki. 1987. Conditions favouring doubled haploid breeding over conventional breeding of self-fertilizing crops. *Euphytica* 36: 441–453. [30]
- Zobel, R. W., M. J. Wright, and H. G. Gauch, Jr. 1988. Statistical analysis of a yield trial. *Agron. J.* 80: 388–393. [30]