

**Review Article**

Advances in Statistical Methods for Genetic Improvement of Livestock: A Review

C.V. Singh

G.B. Pant University of Agriculture & Technology, Pantnagar-263145, District Udham Singh Nagar (Uttarakhand), Department of Animal Genetics and Breeding College of Veterinary and Animal Science, India

ABSTRACT

Developments in statistics and computing as well as their application to genetic improvement of livestock gained momentum over the last 30 years. This paper reviews and consolidates the statistical methodology used in animal breeding. This paper will prove useful as a reference source for animal breeders, quantitative geneticists, and statisticians working in these areas. The estimates of genetic and phenotypic parameters viz. heritability, genetic and phenotypic correlation are used to determine the method of selection, the intensity of selection for different traits of interest, and prediction of selection response. The unbiased property of ANOVA estimators demands no distributional assumptions of the random effects and the residual error terms in a model but all sampling variance results have been developed based on assuming normality. The parameters are estimated by maximizing the logarithm of the likelihood function. The estimates of predictors of the random effects are expected to be more efficient. The drawbacks of ML are first, that it is downwardly biased because the loss of degrees of freedom due to estimating fixed effects is not taken into account. The estimates of predictors of the random effects are expected to be more efficient. The drawbacks of ML are first, that it is downwardly biased because the loss of degrees of freedom due to estimating fixed effects is not taken into account. Maximum likelihood (ML) restricted maximum likelihood and minimum norm quadratic unbiased estimations (MINQUE) are all preferred to ANOVA because they have built-in properties.

Corresponding Author: C.V. Singh < cvsingh2010@gmail.com >

Cite this Article: Singh, C.V. (2023). Advances in Statistical Methods for Genetic Improvement of Livestock: A Review. *Global Journal of Animal Scientific Research*, 11(1), 64-88.

Retrieved from <http://www.gjasr.com/index.php/GJASR/article/view/152>

Article History: Received: 2022.12.27 Accepted: 2023.02.13

Copyright © 2023 C.V. Singh



This work is licensed under a [Creative Commons Attribution-Noncommercial-No Derivatives 4.0 International License](https://creativecommons.org/licenses/by-nc-nd/4.0/).

MINQUE may considerably be better than the analysis of variance procedures. DFREML was the first public package to implement the derivative-free REML, and it became the standard in the field to which every other program is compared. Its unique feature is the likelihood ratio test for testing the significance of variance component estimates. The use of ML and REML in animal breeding has brought about a change in the random effects fitted in the infinitesimal additive genetic model. In traditional ANOVA and related methods, (co) variance is described in terms of random effect due to single parent (e.g., sire model) or both parents (sire dam model), uniquely partitioning the total sum of the squared deviations of the observations from the grand mean into the sum of squares contributed by each factor in the design. However, over the last decade, considerable research effort has concentrated on the development of specialized and efficient algorithms. This has been closely linked to advances in the genetic evaluation of animals by Best Linear Unbiased Prediction (BLUP). However, ML and REML allow the random effect of models to be expressed in terms of the genetic merit or breeding value of animals. These models are called individual animal models (IAM) and incorporate information on the relationship between all animals. Animal Model (AM) has influenced the use of the mixed model methodology in the statistical analysis of animal breeding data considerably. The AM includes a random effect for the additive genetic merit of each animal, both for animals with records and animals which are parents only, incorporating all known relationship information in the analysis.

Keywords: BLUP, DFREML, REML, Animal Model and MINQUE

INTRODUCTION

The choice of a criterion of selection depends on the heritability estimates, availability of the required information, and the nature of the traits under consideration. If the economic traits are to be included in a breeding program, accurate estimates of breeding values will be needed to optimize the selection program. This requires knowledge of variance and covariance components (Raheja *et al.*, 2000). Traditionally, variance and covariance components were estimated by ANOVA and regression methods.

Analysis of variance estimators such as Henderson's methods 1, 2, and 3 are appropriate in the data where some individuals lack records on some traits as a result of selection on one and other traits. The main assumption of random sampling underlying standard ANOVA-type procedures does not hold. Therefore, the estimates of variance and covariance obtained from these methods are expected to bias by selection (Robertson, 1977; Meyer and Thompson, 1984). The method of least squares (LS) analysis of variance based on paternal half-sib correlation has

widely been used in India for estimating the variance components for animal breeding data.

In contrast, with the analysis of variance estimators, maximum likelihood estimators seem to be free of some forms of selection bias (Schaeffer and Soong, 1979). Minimum Norm Quadratic Unbiased Estimation (Rao, 1971) and restricted maximum likelihood (Paterson and Thompson, 1971) can be used to account for all relationship (Henderson, 1985) which result in an estimate that is less biased by selection and more precise than estimates obtained by traditional methods (Keele and Harvey, 1989). An important advantage of REML, utilizing relationship is that the assumptions that animals (sires) are unrelated and are non-inbred need not be made. Estimating covariance matrices based on mixed model methodology has the basic advantage of using identical models in the prediction of genetic merit and variance component estimation for all possible models.

Reliable estimates of variance and covariance components are also needed for obtaining accurate estimates of genetic and phenotypic parameters. The estimates of genetic and phenotypic parameters viz. heritability, genetic and phenotypic correlation are used to determine the method of selection, the intensity of selection for different traits of interest, and prediction of selection response. Parameter estimates from a sample of data may vary depending on the kind of analysis.

Most of the reported heritability are based on the ratio of variance components estimated mainly by Henderson's method 3. In India, estimates of genetic and phenotypic parameters are also based on least squares analysis of variance and the scientific reports on the use of the restricted maximum likelihood method are scanty (Raheja *et al.*, 2001). The available literature on the various aspects of the present study has been presented under the following heads.

Statistical Methods of Estimating Genetic and Phenotypic Parameters

Fisher's (1918) paper on the theory of quantitative genetics was an important contribution to the development of variance component theory. He made inceptive use of the term "variance and analysis of variance". In his book "statistical methods for research workers", he developed the analysis of variance method from the sum of squares. Eisenhart (1947) made the first precise distinction between the "fixed and random" model (Eisenhart I and II) and before that the name "mixed model" had not been suggested before 1947. Eisenhart's model III described mixed models. In the absence of methods for cross-classified, mixed model data with missing sub-classes; Henderson (1953) derived 2 methods, known as a method I for a random model and method 3 for a mixed model. He also presented method 2, another simple method for the mixed model, but with severe restrictions on the model. Henderson's (1953) methods of estimation from unbalanced data are known as Henderson Model I, II, and III. In the model I all effects are random using quadratic

forms or canonical forms i.e., sum of squares which is analogous to the sum of squares of balanced data. Model II is a fixed model, which assumes no interactions of sub-classes of fixed effects within random effects etc. It is an adaptation of model I and model III is the method of fitting constants for a mixed linear model. Data arising in animal genetics are usually not balanced but methods analogous to the ANOVA have been developed for unbalanced data. In particular, Henderson' (1953) method 3 of fitting constants has found extensive use.

The approach replaces the sum of squares in the balanced ANOVA with quadratic forms involving least-squares solutions of effects for which variances are to be estimated. Its widespread application was greatly aided by the availability of general least squares computer programs tailored towards applications commonly arising in animal breeding.

Traditionally the phenotypic covariance between relatives has been estimated using analysis of variance (ANOVA) or analogous procedures. In general, these require that individuals can be assigned to groups with the same degree of relationship for all members. Family structures considered most often are, for instance, paternal half-sib group or parents and their offspring. Using ANOVA, the covariance among members of a family or groups of relatives is usually determined as the variance components between groups.

In fields such as animal breeding, this evolution has come about because ANOVA and related methods are based on several assumptions that are commonly violated in typical animal breeding data sets. These assumptions are (1) that the data are balanced, i.e., there is equal number of individuals in each subclass (2) that the data are a random sample from an unselected population (3) that the data structure conforms to certain standards or stereotypical designs, e.g., paternal half-sibs or parent(s) offspring and therefore, the only type of relatedness is exploited in the analysis (Shaw, 1987; Mayer, 1989a; Searle, 1989). However, animal breeding data are typically unbalanced, being from selection experiments or livestock improvement schemes in which animals are continuously culled for poor performance and are related in a variety of ways. Hence estimates from ANOVA and related types of analysis are biased (Shaw, 1987; Meyer, 1989a).

From 1956 to 1968 several workers developed formulae for sampling variances of ANOVA estimators and Henderson methods estimators. The unbiased property of ANOVA estimators demands no distributional assumptions of the random effects and the residual error terms in a model but all sampling variance results have been developed based on assuming normality. The ANOVA method for balanced data was well known and convenient estimation methods for nested classifications with unequal numbers had been used. The ANOVA estimators from balanced data are minimum variance unbiased on assuming normality of the random effects and error terms. Later it was shown that even without normality, ANOVA estimators are

minimum variance, quadratic and unbiased. Despite the attractiveness of their properties, ANOVA estimators suffer from one major drawback, i.e. negative estimates of variance components. Searle (1989) and Searle *et al.* (1992) have made a good review of the development of components during this period. The ANOVA methods and Henderson's methods do not have general analytical properties that can be used to determine the relative optimality of any one application of the general ANOVA method over another and they also lack distributional properties.

Later with the development of Rao's MINQUE, Hartley and Rao's ML methods, and Patterson and Thompson's REML, there has been much interest in employing better methods than the ANOVA-type estimators. These were all quadratic, translation invariants, unbiased estimators with no known optimum properties concerning sampling variance. Over the last decade, statistical methods employed to estimate (co) variance components for continuous traits in most fields, such as animal breeding and population biology, have generally evolved from analysis of variance (ANOVA) and related types of analysis (e.g. the general linear model) to maximum likelihood (ML) and related methods (Shaw, 1987; Meyer and Hill, 1992), with the increasing interest and necessity for dealing with multi-trait problems it has become imperative that more attention was paid to estimation of environmental and genetic covariance matrices for multiple traits. An increase in the power of computers and the development of specialized algorithms have aided this evolution (Meyer, 1989a, b; Klassen and Smith, 1990).

In contrast to ANOVA, evidence has been accumulating which indicates that ML and REML may have considerable power to eliminate selection bias, consequently, Henderson (1984) attempted to derive feasible computational strategies for these methods applied to the multiple trait problems. This has resulted in the derivation of a method that involves computing quadratics in $\hat{\mu}^A$ and $\hat{c}^A$, BLUP solutions coming from the mixed model equations. In light of these shortcomings of ANOVA methods, alternative methods like maximum likelihood (ML), restricted maximum likelihood (REML), and minimum norm quadratic unbiased estimation (MINQUE), etc. were considered for the estimation of variance components. Interest in ML and related methods have risen because they have based on sufficiently consistent, asymptotically normal, and efficient statistics (Harville, 1977; Kennedy, 1981). Furthermore, constraints on parameters are imposed in ML to exclude out-of-bounds (Harville, 1977; Shaw, 1987; Searle, 1989; Meyer, 1990). However, out-of-bounds, estimates raise doubts about the validity of the model fitted (Shaw, 1987; Searle, 1989; Klassen and Smith, 1990).

Least Squares Method

The method envisages developing estimators by mathematically minimizing the error variance. Thus, we get estimates which are having minimum variance are the most efficient. They are also generally unbiased. Another advantage of least squares is that there is no requirement of knowing the distribution of the observations or variables. The estimation from non-orthogonal data by least squares is satisfactory for fixed effects. However, for estimating or predicting the random effects, more efficient methods have been developed.

SML Package

During the 1960s, this program by Harvey was a forerunner of present evaluation and estimation programs. It computed simple statistics, solutions, and variance components and tested hypotheses to mixed models with diagonal variance, and covariance matrices. LSML, which was extensively used and cited until a few years ago, used a dense matrix inversion with absorption (Gaussian-elimination) of one effect and variance component estimation by Henderson-3 (Henderson, 1984).

Maximum Likelihood (ML)

Fisher (1925) derived the method of maximum likelihood. Its general application to the estimation of variance components took around 40 years since its derivation. Herbach (1959) derived explicit maximum likelihood estimators for certain balanced data models and took account of the necessity that such estimators must be non-negative (because the method of ML prescribed maximization over the parameters space and variance components are non-negative).

The likelihood function is the likelihood (or chance) of simultaneous occurrence of observations and is generally the product of the density distribution function of the observations (variables). The parameters are estimated by maximizing the logarithm of the likelihood function. The estimates of predictors of the random effects are expected to be more efficient. The drawbacks of ML are first, that it is downwardly biased because the loss of degrees of freedom due to estimating fixed effects is not taken into account. This loss can be severe if there is a large number of fixed effects in the model (Harville, 1977; Kennedy, 1981; Meyer, 1990). Secondly, ML may be further biased because data are required to be normally distributed (Shaw, 1987; Harville, 1977). However, Harville (1977) inferred without proof that ML and related methods may be appropriate even when the distribution of data is not specified. Proof, that in certain instances ML and related procedures are not biased by asymmetrical distribution, was provided by the simulation studies of Banks *et al.* (1985) and Westfall (1987).

Minimum Norm Quadratic Unbiased Estimation (MINQUE)

Townsend (1968), Harville (1969), and Townsend and Searle (1971) began the attempts on finding the minimum variance quadratic unbiased estimators of variance components. This was followed by a series of papers on minimum variance estimation and minimum norm quadratic unbiased estimation. Without making distributional assumptions about the data, it is based on seeking quadratic forms ($Y'AY$) to estimate variance components in such a way that (i) A is symmetric, (ii) $Ax = 0$ to have $Y'AY$ free of the fixed effects (iii) $Y'AY$ is unbiased, and (iv) minimizing a Euclidean norm which under normality equates to a minimum variance property. Estimators obtained by MINQUE are functions of a priori values used in place of variance components in the estimation procedure itself. This may lead to different sets of estimators from the same data and the same model with different priori values. MINQUE estimators are unbiased. MINQUE equations are linear and can be solved without interaction for a given set of a priori values. Searle *et al.* (1992) thought that MINQUE is not a practical method of estimating variance components.

Restricted Maximum Likelihood (REML)

Bias in ML caused by failure to take account of the loss of degrees of freedom due to fitting fixed effects has been corrected for unbalanced data by the modified method called restricted maximum likelihood (REML) developed by Patterson and Thompson (1971). Thompson (1962) first developed REML by introducing the idea of maximizing that part of the likelihood, which is invariant to the location parameters of the model, i.e. to the fixed effects. Thus, the basic idea of REML estimation is that of estimating variance components based on residuals calculated after fitting by ordinary least squares first for the fixed part of the model. It may also be viewed as maximizing a marginal likelihood function (Searle *et al.*, 1992). REML has emerged as the method of choice in estimating covariance matrices in animal breeding. Its widespread use has only become possible with the increasing computer power on one hand and the availability of free software packages on the other. REML is equivalent to performing ML on data that have been standardized to have a mean of zero (Meyer, 1989a; James, 1991). These modifications lead to estimates being identical to those of ANOVA if data are balanced and if out-of-bounds estimates are not excluded. REML shares all the desirable properties of ML (Shaw, 1987). Estimating covariance matrices based on mixed model methodology has the basic advantage of using identical models in the prediction of genetic merit and variance components estimation for all possible models, i.e. also in multi-trait cases.

REML is more difficult computationally than ML, as it requires including the fixed effects (Kennedy, 1981; Mayer, 1989a and 1993), except for a simple and balanced

design. REML estimates require the numerical solution of a constrained non-linear optimization problem. Analytical solutions are impossible and iterative procedures must be used (Harville, 1977; Kennedy, 1981; Klassen and Smith, 1990, Meyer, 1993).

Iterative methods that are both first and second derivatives have been found to cover the quickest (Meyer, 1989a, 1990; Klassen and Smith, 1990). However, these are difficult to calculate in the highly unbalanced data typically found in animal breeding (Klassen and Smith, 1990). Therefore, it is more common to use numerical techniques e.g. the simplex method or statistical approximations of either first or second derivatives, i.e. Quasi-Newton-Raphson method (Graser *et al.*, 1987; Klassen and Smith, 1990).

Maximum likelihood (ML) restricted maximum likelihood and minimum norm quadratic unbiased estimations (MINQUE) are all preferred to ANOVA because they have built-in properties. MINQUE may considerably be better than the analysis of variance procedures. However, ANOVA procedures require fewer computations than MINQUE (Keele and Harvey, 1989). ML and REML have been the methods of choice over MINQUE because in MINQUE we get different values of estimated variance with different sets of pre-assigned values for the same data. REML estimates will be more precise than MINQUE (taking all prior values for components of variance and covariance set to zero except for the environmental variance). However, REML is less biased by selection than MINQUE (taking all prior values for components of variance and covariance set to 1) Sorensen and Kennedy, 1984. MINQUE1 offers a computational advance over REML using the iterative MINQUE algorithm because MINQUE1 requires the same computations as one iterate of iterative MINQUE. REML by the expectation-maximization algorithm or the derivative-free algorithm requires fewer computations per iteration than MINQUE1. MINQUE 0 would be a useful method for estimating components of variance and covariance for characters with low h^2 such as reproductive traits in livestock (Keele and Harvey, 1989). The favored characteristics of REML over ML are (i) with balanced data, REML equations reduced to the same equations as are used in ANOVA estimation, and (ii) further REML takes account of the degree of freedom used for estimating fixed effects. For example, in a simple sample of x in

$N(0, \sigma^2_e)$, the REML estimate of S^2 is $\sum (X_i - \bar{X})^2 / (n-1)$, whereas the ML estimator is $\sum (X_i - \bar{X})^2 / n$. In this simple case, REML is unbiased, but that is not

the general rule, because nothing is unbiased after iteration, neither in ML nor REML. The advantage of ML over REML is that ML procedure includes providing ML estimation of fixed effects. REML method provides no such estimator.

Both ML and REML are based on maximizing the likelihood of error contrasts (Paterson and Thompson, 1971; Smith and Graser, 1986; James, 1991) of note is that, as for other methods, the accuracy and precision of REML estimates increase with sample size regardless of the design and the criterion for choosing estimates (Shaw, 1987). However, even with the best method of analysis, variance components are only as good as the data on which they are based (Meyer, 1989a, and 1990).

REML algorithms

Methods that use both first and second derivatives, i.e. geometrically speaking information on slope and curvature of the function, have been found to cover the quickest a special case is Fisher's method of scoring which requires expected rather than actual values of the second derivatives, which are of often easier to calculate. This has been used by Patterson and Thompson (1971) and, for animal breeding application, in REML algorithms described by Thompson (1973), Schaeffer *et al.*, (1978), Meyer (1983, 1985), and Cude (1986), REML algorithms are often formulated in terms of Henderson's (1973) mixed model equations (MME).

Canonical Transformation (CT)

A re-parameterization of the model or a transformation applied to the MME can often reduce computational requirements substantially. For multivariate analysis, a transformation of the data with corresponding effects is available for special cases. When all traits are recorded for individuals at the same or strictly corresponding time(s), design matrices x and z are equal for all traits, consider p traits for a model with one random factor, and let T and E (of size $p \times p$) denote the covariance matrices of random effects and residuals, respectively. As outlined by Hayes and Hill (1980), a canonical decomposition of T and E then yields a transformation to new traits, so-called canonical variables, which are both genetically and phenotypically uncorrelated.

This reduces the multi-trait analysis to a series of corresponding univariate analyses. REML analysis exploiting the canonical transformation has been described for the method of scoring (Meyer, 1985); the EM-algorithms with tridiagonalization (e.g. Taylor *et al.*, 1985) or diagonalization (Van Randen and Freeman 1986; Lin, 1987) of the coefficient matrix and the derivative-free approach (Mayer, 1988c, Thompson and Juga, 1988). For the latter algorithm, also proved useful for models involving more than one random effect by reducing the number of non-zero off-diagonal elements in the coefficient's matrix and thus the computational effort to evaluate $\log L$.

A special algorithm for diagonal E , i.e., zero residual co-variances, has been described by Schaeffer *et al.* (1978). This is the case when different traits are

measured on different sets of animals for an example growth rate of male and female calves in beef cattle. In special cases, this constellation can also be achieved through a transformation of the data. For instance, when traits are recorded sequentially, i.e. there are missing records but each animal with the same traits t has all records 1 to t , the inverse of the Cholesky, decomposition of E provides a transformation to traits with uncorrelated residuals (Schaeffer, 1986).

Expectation maximization (EM) algorithm

Even for simple models, the calculation of expected derivatives has proven to be computationally highly demanding if not prohibitive. Hence, to date, most REML applications are based on the so-called expectation maximization (EM) algorithm (Dempster *et al.*, 1977) which, implicitly, requires first derivatives of the likelihood to be evaluated. The resulting estimators then have the form of quadratics in the vector of random effects solutions obtained by BLUP for the assumed values of variances to be estimated, which are equated to their expectations.

For the mixed linear model, REML estimates of variance components using the EM algorithm can be obtained as:

$$\hat{\sigma}_{\mu}^2 = \hat{\mu}' A^{-1} \hat{\mu} / [q - \lambda \text{tr}(A^{-1}C)]$$

$$\hat{\sigma}_e^2 = \hat{e}' \hat{e} / [NDF - \lambda \text{tr}(A^{-1}C)]$$

$$\hat{\sigma}_{\mu}^2 = [y'y - y'Xb - y'Zu] / (N - r(X))$$

With $\hat{e} = y - X'\hat{b} - Z'\hat{u} = M(y - Z'\hat{\mu})$. Alternative forms are (Harville, 1977).

$$\hat{\sigma}_{\mu}^2 = [\hat{\mu}' A^{-1} \hat{\mu} + \sigma_e^2 \text{tr}(A^{-1}C)] / q$$

The EM algorithm is slow to converge, especially for low heritabilities. However, modifications have been suggested that improve its performance without increasing the computational burden considerably. These included a re-parameterization estimating variances of family means rather than variances between families, as described by Thompson and Meyer (1986) and Harville and Callanan (1988) who referred to it as 'linearization'. For a multivariate analysis, Robinson (1988) used a so-called secant algorithm that approximates the matrix of second derivatives iteratively Meyer (1986) combined an EM step to estimate the residual components with a method of scoring step to estimate co-variances for random effects. Schaeffer's (1979), common intercept approach attempted to reduce the number of interactions required by predicting changes in estimates.

Though computationally less demanding than the method of scoring, the EM algorithm, in general, requires the direct inverse of matrix of size equal to the number of levels of all random effects in each round of interaction. This imposes

severe restrictions on the kind of analysis feasible, especially for multivariate analysis. Simplifications have been described for special cases.

The familiar mixed model equations (MME) are used widely for variance component estimation with REML algorithm by nearly all animal breeders. As noted by Patterson and Thompson (1971), REML estimators can be expressed in terms of latent roots of matrices in the MME. For a model with one random effect, Dempster *et al.* (1984) described the use of a singular value decomposition of the MME to estimate variances via the EM algorithm. With the strategy, the major computational burden is the calculation of eigenvalues and eigenvectors. However, this is required only once per analysis, and subsequent iterations to solve the above-written equations are fast. Along the same lines, Smith and Graser (1986) advocated the use of a Householder transformation to reduce the coefficient matrix in the MME to tridiagonal rather than diagonal form, which is computationally less demanding. Thompson and Meyer (1988) extended this approach to Reduced Animal Model (RAM).

For models with two random effects, partitioned matrix results have been utilized to reduce the size of the matrix to be inverted in each round of iteration to the number of levels of one of the random effects (Meyer, 1987). Another approach has been to employ a nested two step procedure, combining an EM-step to estimate two variance components, as described above, with a direct search for the maximum likelihood function for the third component (Smith and Graser, 1986).

Derivative-free REML (DF-REML)

The minimum or maximum of a function can be found without knowing its derivatives using numerical techniques called derivative-free (DF) algorithms (Meyer 1989a, 1990). The use of a derivative-free approach for REML estimation of variance components has been considered first by Graser *et al.* (1987) for an AM with animals as the only random effect and univariate analysis. For this case, the REML log-likelihood is:

$$\text{Log } L = -\frac{1}{2} [\text{const} + q \log \sigma^2_{\mu} + \text{NDF} \log \sigma^2_{\epsilon} + y'Py + \log |B| + \log |A|]$$

with $P = S - SZCZ'$'s and $B = C^{-1}$. They showed that the sum of squares of residuals, $y'Py$, and log determinant of the coefficient matrix, $\log |B|$, can be evaluated simultaneously by augmenting B by the vector of right-hand sides and the total sum of squares, $y'y$, and absorbing all rows and columns into the latter. In addition, the residual variance can be estimated directly as $y'Py/\text{NDF}$ so that $\text{Log } L$ can be maximized concerning one parameter only, the variance ratio λ , estimating $\hat{\sigma}^2_{\mu}$ subsequently as $\hat{\lambda} \hat{\sigma}^2_{\epsilon}$. This has been extended to models including additional

random effects, such as an environmental effect due to litter or a maternal genetic effect, and to multivariate analyses (Meyer, 1988 b,c).

The DF approach using the simplex procedure of Nelder and Mead (1965) has proved to be easy and robust (against starting values far from the estimate) for use with animal breeding data (Meyer, 1989b, 1991). This procedure allows the simplex method to re-scale itself automatically in each iteration, changing shape and size according to the landscape of the surface searched. This adaptability is achieved by the combination of reflection, expansion, and contraction (Meyer, 1990). Furthermore, it is highly flexible; accommodating a wide range of models of analysis, in particular facilitating analyses under the animal models, differing in random effects fitted in assumptions about covariance between them, and is of interest for the analysis of animal breeding data (Meyer, 1989a, 1993).

DFREML Package

DFREML, written by Meyer (1988a), was the first public package to implement the derivative-free REML (Smith and Graser, 1986). Extensively cited, it became the standard in the field to which every other program is compared. Its unique feature is the likelihood ratio test for testing the significance of variance component estimates. The documentation is extensive, and it is the only one that has descriptions of all subroutines in the package. DFREML appears very clean of errors. DFRML supports only 10 classes of models, although the important models are included. Also learning curve is high (Misztal, 1994a). The most expensive procedure in a REML program is the computation of the determinant or trace. DFREML computes the determinant by Gaussian elimination and other packages by matrix factorization.

Derivative-free v/s Expectation Maximization

DF and EM denote slightly different concepts (Misztal, 1992). DF denotes all maximization methods that rely only on the value of the maximized function whereas EM is one of the methods using both the values of the function plus its first derivative. To make a more valid comparison, EM will be extended to EM type, which includes all methods that use the first derivative, including accelerated EM. The value of maximum in DF has at most only half as many significant digits as the maximized function. Poor numerical conditioning could lead to false maxima, especially for multiple traits, and when the correlations between variance components are high. One could implement EM-type algorithms by computing the first derivative numerically, disregarding accuracy, using many components +1 function evaluations for each step of EM type. In EM-type, judicious implementation of the sparse inversion, where only selected inverse elements are computed, runs at most only twice as slow as a single evaluation of the likelihood function under DF. Under careful implementation, EM-type algorithms should be

faster and more reliable than DF algorithms, although DF is simpler to implement for general models. Each algorithm operates in quadratic time about the number of equations and none is suitable for application to very large data sets (Misztal, 1992).

Derivative and Derivative-Free Maximization in REML

The speed and accuracy of the REML variance component estimation are dependent on the maximization strategy. The popular derivative-free (DF) maximization used in DFREML is very slow in multiple traits as thousands of rounds of iteration may be needed to obtain convergence. Assuming that the logarithm of the restricted likelihood function (L) is approximately quadratic, Misztal (1994a) calculated that the number of steps to achieve convergence in t -traits is $-T_2$ with good DF algorithms, but it does not depend on t in better derivative (D) algorithms. Better DF algorithms included Powell and Tosenbrock and better D algorithms included quasi-Newton or Broyden (as nonlinear solving). The latter includes the accelerated EM algorithm. These costs are underestimated because the L function is approximately quadratic only close to the maximum. Derivative-free multivariate REML algorithms are computationally expensive, especially, if the likelihood function contains many parameters to be estimated. For example, Misztal (1994b) showed as the number of traits increases the derivative-free methods become less efficient than methods using first derivatives, i.e., procedures based on the EM algorithm. One round of a derivative-free method involves computing the determinant of the coefficient matrix of the mixed model equations. In the EM algorithm, elements in the sparse inverse of this matrix are needed. Worse convergence properties of DF can be seen intuitively by noting that D can sense a desirable direction (gradient) in one round, while DF has to do approximately t^2 rounds to probe all the dimensions. The combined costs of factorization/inversion and maximization are needed for at least t^3 and t^5 numerical preparations for better D and DF algorithms, respectively. Programs using D algorithms are not common because DF algorithms are easier to implement and inversion before FSPAK became available, was very expensive.

Results of analysis with many traits and the general model are not likely to be accurate. First, the accuracy of the factorization /determinant decreases as the traits are becoming more linearly dependent and the MME matrix is larger (Misztal, 1994a). Second, the maximization method may fail. "Faster" DF (such as Powell or Tosenbrock as opposed to simplex) or D (Newton-Raphson or quasi-Newton as opposed to fixed point = EM) algorithms may converge slower far away from the maximum, and many need a fall back to slower algorithms in early rounds to avoid divergence. Also, DF's solutions are less accurate than D because finding a maximum, where the maximized function is flat by definition, is less accurate than

finding a zero of a derivative, which is not flat. Together, a general-model REML and particularly DF may be too expensive and inaccurate with more than 2-4 traits. With many traits, the only feasible procedure at this time is canonical transformation (CT) where the computing cost increases only linearly with the number of traits but only certain models are supported. This procedure is also insensitive to high correlations between the traits, but it is restricted to analyses with the same model for each trait, all traits recorded, and one random effect. CT is supported by DFREML.

The CPU time taken is relatively small for canonical transformation, increasing steeply for DF, D being in between. If a single trait REML took 1 minute of computing time, a 2-trait REML would take at least 1 hr in DF, 8 minutes in D, and 2 minutes in CT. For 5 traits, these times would be 2 days, 2 hrs., and 5 minutes, respectively, and for 15 traits 527 days, 2 days, and 15 minutes, respectively. If the memory required were 2 Mbytes in single traits for D of DF it would be 8 Mbytes in 2 traits, 50 Mbytes in 5 traits, and 450 Mbytes in 15 traits (Misztal, 1994a).

Average Information Restricted Maximum Likelihood (AI-REML)

The use of numerical approximations based on Expectation- Maximization algorithms, which calculate expected second derivatives, has proved to be highly demanding of computation, if not prohibitively demanding, and slow to converge, especially for traits of low heritability (Meyer, 1989a, 1990). However, Johnson and Thompson (1995) have recently developed univariate IAM using the average of observed and expected information. Sparse matrix techniques are employed to derive the coefficient matrix required to calculate the first derivative of the likelihood. The matrix of second derivatives is called the observed information matrix. The second derivatives of the likelihood are calculated by averaging their observed and expected values. This leads to an algorithm called Average Information REML (AI-REML), which is a compromise between the Newton-Raphson and the Fisher scoring algorithms (Madsen *et al.*, 1994; Johnson and Thompson, 1995). The expectation of this matrix is the Fisher information matrix. Both the observed and the expected information matrices involve terms that are difficult to compute. The development of AI-REML algorithms followed the observation that the average of observed and expected information matrices as the second derivative is considerably easier to compute than either of the components due to the cancellation of some terms (Madsen *et al.*, 1994; Johnson and Thompson, 1995). AI-REML was extended to a multivariate analysis by Madsen *et al.* (1994). These extensions to multivariate algorithms were from similar mathematical backgrounds.

DF-REML versus AI-REML

Criticisms of the use of DF algorithms have been based on their slow convergence in multivariate analyses and poor numerical properties, and this has led to an interest in other methods (Madsen *et al.*, 1994). Particular criticisms of the simplex procedure are that its performance becomes less successful as the dimension of search increases with increasing numbers of traits and random effects in the model (Meyer, 1989b, 1991). In test runs, the AI-REML algorithm gave almost identical parameter estimates as DF-REML but with a considerable reduction in computer time. However, AI-REML method converged faster (341.2 s CPU time) than the DF-REML (1837.6s CPU time) on a Cray C92A computer. The AI-REML algorithm converged in 13 or fewer rounds for all the analyses conducted, while DF-REML used 472 rounds for the bivariate analyses (3, 4, and 5 traits DF-REML analyses were not run due to constraints on CPU time) (Madsen *et al.*, 1994). Even though the AI-REML algorithm requires more computations per round of iteration, the savings in computer time were considerable. In the bivariate analyses, the AI-REML algorithm used less than 20% of the computer time used by the DF-REML algorithm. Furthermore, algorithms using the second derivate of the likelihood provide estimates of sampling variance of parameters as a by-product. In contrast, for DF-based algorithms, the additional computational effort is required to estimate sampling error; the extra effort can be quite considerable if a large number of parameters are estimated (Mayer and Hill, 1992). In conclusion, AI-REML algorithms are currently more efficient in the use of computer time and resources than those based on DF multivariate analysis of two or more traits with one random effect. However, for univariate analysis, DF-REML performs similarly to AI-REML. Despite savings in computing time and resources, multivariate IAM AI-REML is at an early stage of development compared with the DF-REML methods of Meyer (1989a) as described by Misztal (1994a). Further work is needed in developing AI-REML to the current standard of DF-REML if AI-REML is to be a worthwhile competitor.

Comparing parameter estimates from ANOVA and related types with those from REML

Few comparisons have been reported of parameter estimates based on REML fitting an IAM and traditional sire models based on ANOVA and its related types. Chauhan (1991) studied a comparison of estimates of heritability of milk yield of Murrah buffaloes from restricted maximum likelihood and Henderson's methods III and revealed that because of its desirable properties, the REML estimator was considered to be more appropriate than that of Henderson's method III. Raheja (1992) made a comparative study of variance, and covariance components of economic traits between different lactations estimated from single and multi traits

procedures and recommended restricted maximum likelihood or maximum likelihood multi traits procedures because they make use of all data available and account for bias in later lactation records due to selection on dairy performance. Raheja *et al.*, (2000) reported a comparison of four methods of variance components estimation for heritability of economic traits in Murrah buffaloes and revealed that the likelihood method always yields positive estimates of heritability and therefore should be the method of choice over the fitting constants procedure (HM3) for imbalance data in mixed models. Jain and Sadana (2000) studied sire evaluation using animal models and conventional methods in Murrah buffaloes and revealed that deciding on the method of sire evaluation to be used for selecting sires with high breeding values criteria of the rank correlation could be misleading and comparison of selected sires likely to give a variable picture. The best linear unbiased prediction method under multi traits animals incorporating first lactation milk yield with first service period as co-variable and age at first calving in the model was found to be more efficient and accurate for sire selection in Murrah buffaloes.

Some considerations in practice of various methods of estimation.

The more refined methods have definite utility. It should be examined whether the type of data we have is suitable for analysis under a method: (i) while using iterative methods, one should be sure that the iteration will converge, and (ii) the estimates should be under global maximum and not under a local maximum of the likelihood function and time and cost of computing.

Animal Model

The use of ML and REML in animal breeding has brought about a change in the random effects fitted in the infinitesimal additive genetic model (Henderson, 1988; Foulley, 1990). In traditional ANOVA and related methods, (co) variance is described in terms of random effect due to single parent (e.g. sire model) or both parents (sire dam model), uniquely partitioning the total sum of the squared deviations of the observations from the grand mean into the sum of squares contributed by each factor in the design (Harville, 1977; Shaw, 1987). However, over the last decade, considerable research effort has concentrated on the development of specialized and efficient algorithms. This has been closely linked to advances in the genetic evaluation of animals by Best Linear Unbiased Prediction (BLUP). However, ML and REML allow the random effect of models to be expressed in terms of the genetic merit or breeding value of animals. These models are called individual animal models (IAM) and incorporate information on the relationship between all animals (Meyer, 1989b, 1991). Animal Model (AM) has influenced the use of the mixed model methodology in the statistical analysis of animal breeding data considerably. The AM includes a random effect for the

additive genetic merit of each animal, both for animals with records and animals which are parents only, incorporating all known relationship information in the analysis. This requires the inverse of numerator relationship matrix A . Hence it gained practical importance with the availability of a procedure to obtain A^{-1} directly from a list of pedigrees (Henderson, 1976; Quaas, 1976) which made the AM computationally feasible for large data sets. Kennedy *et al.* (1988) discussed the genetic properties of animal models, outlining how the AM can account for the change in genetic means and variances. Thus, the AM allows an optimal analysis of data involving multiple generations arising, for instance, from selection experiments (Sorensen and Kennedy, 1986; Kennedy, 1988).

In terms of variance component estimation, the AM had changed thinking from the interpretation of covariance between relatives to a linear model framework where we determine variances directly by fitting corresponding random effects in the model of the analysis. Covariances between random effects for relatives are now taken into account by specifying the variance matrix of random effects accordingly. With the AM, the additive genetic variance is estimated as the variance of animals' additive genetic merit instead of, for example, four times the variance between sires or twice the covariance between parents and offspring. The basic assumptions of the individual animal model (IAM) are : (i) σ^2_e is the same for all observations, (ii) dominance genetic effects are not important and are part of σ^2_p (iii) covariance between animal genetic effect and other random effects in the model are zero (iv) the relative values of variances must be known, and (v) additive genetic effect can include individual without any observations, and in that case corresponding design matrix contain zero columns for those individuals. It is intuitively obvious that an IAM is more correct for animal breeding data since it exploits all known relationships and can therefore account for changes in genetic variance due to both inbreeding and the established linkage disequilibrium (Kennedy and Sorensen, 1988; Henderson, 1990a). Furthermore, the use of an IAM allows more random effects to be fitted, such as maternal and dominance effect, which are known to bias some genetic estimates (Barlow, 1978; Falconer, 1989; Meyer, 1989a; Webb and Bampton, 1990) The other advantages of the animal model are : (i) if data had been collected over many years then the possibility could arise that an individual female animal could appear as one of the measured individuals, but also as the dam of one or more other female animals. Thus, these equations combine information on an animal itself and its progeny, (ii) in an animal model genetic merit of the female to which sires were mated is also considered whereas the same is ignored while evaluating sires solely on their female progeny, and (iii) if only selected animals were allowed to reproduce then biases due to selection can be avoided by use of

numerator relationship matrix. The major disadvantage of an animal model is the larger order of the equations that need to be solved.

Henderson and Quaas (1976) presented a comparable methodology for a multiple-trait animal model. Sorensen and Kennedy (1984) reported that multiple-trait animal models improve the accuracy of prediction, when observations which formed the bases for selection are available, by utilizing information on correlated traits and by accounting for biases that occur with a single trait model. Hudson and Schaeffer (1984), based on the simulation study, reported that the sire model (SM), material grandsire model (MGSM), or SM with regression on dams' predicted milk yield was 43-47% less accurate in evaluating sires than the animal's models for a random mating but selected population, and 65-67% less accurate under assortative mating and selection. Wiggans and Misztal (1987) opined that the main advantage of an animal model over a sire model is that all additive genetic relationships among animals contributes to an animal's evaluation, which improves the accuracy of evaluation and avoids bias due to non-random mating and female selection. They mentioned the disadvantage that many more equations must be solved, and convergence may be slow because animal equations have off-diagonal elements contributed by the relationship matrix. Meyer and Burnside (1988) mentioned that sire model ignores both the dams of the cow (sire's mate) and relationship between females, and therefore, sire proofs may be biased due to non-random mating or selection of cows. On the other hand, animals model evaluates both the sires and cows simultaneously, animals without records (like sires in dairy cattle) are evaluated from the information on their relatives' records. The animal model takes into account all the relationship, adjusts for the non-random mating, account for selection bias, and adds to the accuracy by taking information from the correlated traits. Canon and Chesnais (1989) enumerated following advantages of animal model: (i) permits the use of all additive genetic relationships among animals as 'a priori' information in animals' evaluation, (ii) the predicted genetic merit of sires is free from bias due to non-random mating since the genetic merit of dams of their progenies is taken into account, and (iii) the need for grouping is decreased to account for genetic trends.

Under Indian conditions, the selection of dairy bulls using conventional methods (such as the contemporary comparison of sire evaluation) has long been carried out. The relationships between the individuals of the population are not taken into account in these methods of sire evaluation on which observations have been made, Sun et al, (2010). Therefore, for accurate estimation of breeding values (BVs), different advanced linear models are needed to be which lay more emphasis on the relationship between the individuals of the population. Both sire and animal models consider the relationship and inbreeding coefficient of the individuals of the population, as well as the model, including the numerator relationship matrix

(NRM). In most countries, the sire model is currently used for genetic evaluation of fertility traits, Interbull (2009). In the sire model, it is assumed that the mates are of equal merit which could result in biased estimates of BVs by Mrode (2005). And, in terms of stability and accuracy of the EBV, the animal model had a superior ability to predict breeding value Sun *et al*, (2009). The sire model has the advantage of less computational demand and might have good predictive properties under the conditions no genetic relationship exists between the sire and dam, that is there are no genetic relationships exist between dams, and thus mating is random. However, the assumed conditions necessary for accurate and unbiased EBV using a sire model are frequently violated in current dairy populations. If mates are non-randomly chosen in some manner, and if the model does not account for mating schemes, sire evaluation may be affected adversely and could be biased Schaeffer (1983).

Comparison criteria of models; Information criteria of Akaike (AIC) and Bayesian (BIC) information criteria tests were used in the comparison of the models. In both tests, the most accurate model will be the one that has the highest negative AIC and BIC values. According to these two tests, we will select the model which fits better to data structures. The values of the Akaike information criteria and Bayesian information criteria will be obtained as follows Lukač *et al.*, (2017).

$$AIC = -2 \log (ML_k) + 2p_k$$

$$BIC = -2 \log (ML_k) + p_k \log (n)$$

Where,

ML_k= Maximum Log Likelihood for model k;

p_k = number parameter for model k;

n = number of observations in model k;

CONCLUSION

Sires were ranked according to their breeding values for both models, which indicated that all sires would not rank the same for first lactation and lifetime traits. The top-ranking sires between the SM and AM were inconsistent for some traits. There were changes in the rank of the first top 10 sires of sire evaluation by the Animal and Sire Model. These results indicated that all sires would not rank the same for present data traits. Comparison between the Animal and Sire Models was done by estimating Information criteria of Akaike (AIC) and Bayesian (BIC) information criteria and found that Animal Model was having highest negative value for (AIC) and (BIC), indicating its superiority over than Sire Model for estimating genetic parameters. which suggested that the Animal Model would be the most adequate, while Sire Model has the largest AIC and BIC, which suggests that this model is not an adequate model for the evaluation of genetics parameters and concluded that the Animal Model would be the most adequate model for evaluation of genetics parameters. Using an animal model, particularly with the

multiple traits analysis, for estimating BV's showed higher genetic diversity compared with the sire model which would lead to a rapid genetic gain in the future generations Dangi (2020).

CONFLICT OF INTERESTS

The author declared no competing interests exist.

REFERENCES

- Banks, B.D.; Ma O. I. and Walter, J.P. 1985. Robustness of the restricted maximum likelihood estimator derived under normality as applied to data with showed distributions. *J. Dairy Sci.* 68(7): 1785-1792.
- Barlow, R. 1978. Biological ramification of selection for preweaning growth in cattle. *A review: Anim. Breed, Abstr.* 46 (7): 469-494.
- Canon and Cheshais, J. 1989. Indirect approach to simultaneous sire and cow evaluation by animal model with repeated records and groups for unknown parents. *J. Anim. Bred. Genet.* 106 : 335-344.
- Chauhan, V.P.S. 1991. Comparison of estimates of heritability of milk yield of Murrah buffaloes from restricted maximum likelihood and Henderson's method III. *Indian J. Anim. Sci.* 61(8) : 880-881.
- Cue, R.I. 1986. Variance covariance components estimation for uni-varieties embedded traits. *J. Anim. Bred. Genet.* 103 : 334-341..
- Dangi, Manita.2020. Evaluation of efficiency of sire model and animal model in cross bred cattle using first and lifetime production traits. M.V.Sc,thesis G.B.Pant.University of Agriculture and Tecnology,Pantnagar,U,S,Nagar
- Dempster, A.P.; Lairal, N.M. and Rubin, D.B. 1977. Maximum likelihood from incomplete data via the EM algorithm *J. Royal Statistical Society, Series B* (Methological) 39(1) : 1-38.
- Dempster, A.P.; Seley, M.R.; Patel, C.M. and Roth, A.J. 1984. Statistical and computational aspects of mixed model analysis. *Applied Statistics*, 33(2) : 203-214.
- Eisenhart, C. 1947. The assumptions underlying the analysis of variance. *Biometrics.* 3(1): 1-21.
- Falconer, D.S. 1989. Introduction to quantitative genetics, 3rd edition. Longman, London, U.K. 438 p.
- Fisher, R.A. 1918. The correlation between relations on the supposition of Meradelian in heritance, *Tran. R. Soc. Edinburg.* 52 : 399-433 (cited from compendium of lectures of Nation Training Programme on 'Computer Application in Animal Breeding CAS, division of DCB, NDRI, Kernal during 26 March to 10 April 1997).
- Fisher, R.A. 1925. Statistical methods for research workers. 1st Edn. Oliver and Boyd. Edinburg.
- Foulley, J.L. 1990. Genetic parameters estimation introduction, In Proc. 4th Word Congr. *Genet. Appl. Livest.. Prod.* Vol XIII, pp. 403-406.

- Graser, H.U.; Smith, S.P. and Tier, B. 1987. A derivative free approach for estimating variance components in animal models by restricted maximum likelihood. *J. Anim. Sci.* 64(5): 1362-1370.
- Handerson, C.R. 1984. Estimation of variance and co-variance under multi-trait models. *J. Dairy Sci.* 67 : 1581-1588.
- Harville, D.A. 1977. Maximum likelihood approaches to variance components estimation and to related problems. *J. Anim. Stat. Assoc.* 72(358): 320-340.
- Harville, D.A. and Callanan, T.B. 1988. Computational aspects of likelihood based inference for variance components. In: advances in statistical methods for genetic improvement of livestock. D. Gianola and K. Hammond (editor). Heidelberg, springer Verlag, Berlin, Germany.
- Hays, J.F. and Hill, W.G. 1980. A reparametrization of a genetic selection index to locate its sampling properties. *Biometrics* 36(2) : 237-248.
- Henderson, C.R. 1953. Estimation of variance and covariance components. *Biometrics.* 9: 226-252.
- Henderson, C.R. 1973. Sire evaluation and genetic trends. In *Proc. Anim. Breed Genet. Symposium in honour of Dr. J.L. Lush*, American, Society of Animal Science and American Dairy Science Association, Blackburg, Champaign, Illinois, U.P.A. pp. 10-41.
- Henderson, C.R. 1976. A simple method of for computing the inverse of G numerator relation ship matrix used in prediction of breeding values. *Biometrics.* 32(1): 69-83.
- Henderson, C.R. 1985. MIVQUE and REML estimation of additive and non additive genetic variances. *J. Anim Sci.* 61(1):113-121.
- Henderson, C.R. 1988. Progress in statistical methods applied to quantitative genetics since 1976. In *Proc. 2nd International Conference on Quantative Genetics*. Weir, B.S.; Eisen, E.J. Goodman, M.M. and Namkoong, G (editors). Sinauer, Sunderland Massachusetts, U.S.A. pp. 85-90.
- Henderson, C.R. 1990a. Statistical methods in animal improvement : historical over view. In *Advances in Statistical Methods for Genetic Improvement of Livestock*. Gianola, D. and Hammod, K. (editors). Heidelberg, Springer Verlag, Berlin, Germany, pp. 1-14.
- Henderson, C.R. and Quaas, R.L. 1976. Multiple trait evaluation using relatives records. *J. Dairy Sci.* 43(6) : 1188-1196.
- Herbach, L.H. 1959. Properties of model it type analysis of variance tests. A optimum nature of the F-test for model II in balance case. *Ann. Math Stat.* 30: 939-959.
- Hudson, G.F.S. and Schaeffer, L.R. 1984. Montecarlo comparison of sire evaluation models in population subjected to selection and non random mating. *J. Dairy Sci.* 67(6) : 1264-1272.
- Interbull.2009. National GES information. http://www-interbull.slu.se/national_ges_info2/framesida-ges.htm Accessed January 22, 2009.
- Jain, A. and Sadana, D.K. 2000. Heritability estimates under single and multi trait animal models in Murrah buffaloes. *Asian Aus. J. Anim. Sci.*, 13(5): 575-579.
- James, J.W. 1991. Effect of using an incorrect model in animal model REML estimation of heritability. In : *Proc. Australian Assoc. Anim. Breed. Genet.* 9: 477-479.

- Keele, J.W. and Harvey, W.R. 1989. Estimation of components of variance and covariance by symmetric. Differences squared and minimum norm quadratic unbiased estimation: *A comparison, J. Anim. Sci.* 67(2) : 348-356.
- Kennedy, B.W. 1981. Variance component estimation and prediction of breeding values. *Can. J. Genet, Cytol.* 23(4): 565-578.
- Kennedy, B.W. 1988. Use of mixed model methodology in analysis of designed experiments. In *Advances in Statistical Methods for Genetic Improvement of Livestock* D. Gianola and Hammnond, K. (editors). Heidelberg, Springer verlags Berlin Germany.
- Kennedy, B.W. and Sorensen, D.A. 1988. Properties of mixed model methods. In *Proc. 2nd International Confer. Quantitative Genetics* Weir, B.S. Eisen, E.J. Goodman, M.M. and Namkoong, G. (editors). Sinaur, Sunderland, Massachusetts, USA, pp. 91-100.
- Kennedy, B.W., Schaeffer, L.R. and Sorensen, D.A. 1988. Genetic properties of animal models. *J. Dairy Sci.* 71 Suppl. 2 : 17-26.
- Klassen, D.J. and Smith, S.P. 1990. Animal model estimation using simulated REML. In *Proc. 4th World Congr. Genet. Appl. Livest. Prod.* vol. XIII pp. 472-475.
- Lin, C.Y. 1987. Application of singular value decomposition to restricted maximum likelihood estimation of variance components. *J. Dairy Sci.* 70(12) : 2680-2684.
- Lukač, D., Mišćević, B., Könyves, T., Puvača, N., Džinić, N. & Duragić, O. 2017. Comparison of models for genetic analysis of traits from performance test of gilts using restricted maximum likelihood. *Romanian biotechnological letters*, 22(4): 12802.
- Madsen P, Jensen J, Thompson R (1994) Estimation of (co)variance components by REML in multivariate mixed linear models using average of observed and expected information. *Proc. 5th World Congr Genet Appl Livest Prod*, University of Guelph, Guelph, Canada Vol 22, 19-220.
- Meyer, K. 1983. Maximum likelihood procedures for estimating genetic parameters for later lactations of dairy cattle. *J. Dairy Sci.* 66(9): 1988-1998.
- Meyer, K. 1985. Maximum likelihood estimation of variance components for multivariate mixed model with equal design matrix. *Biometrics*, 41(1): 153-160.
- Meyer, K. 1988a. DFREML a set of programs to estimate variance components under an Individual Animal Model. *J. Dairy Sci.* 71 (supple2): 33-34 (Abstract)
- Meyer, K. 1988b. Estimation of variance components for Individual Animal Models with additional random effects. *J. Dairy Sci.* 71 (supple 1): 263 (Abstract).
- Meyer, K. 1988c. Approximate accuracy of genetic evaluation under an Individual Animal Model. *J. Dairy Sci.* 71 (Suppl): 264 (Abstract).
- Meyer, K. 1989a. Estimation of genetic parameters In : Hill, W.G. and Mickey, T.F.M. (ed.) *Evaluation of Animal Breeding reviews on molecular and quantitative approaches in honour of Alan Robertson*, Chapter 23, CAB. *International, Wallingfort, U.K.* pp. 161-167.
- Meyer, K. 1990. Present status of knowledge about statistical procedures and algorithms available to estimate variance and covariance components. In: *Proc. 4th World Congr. Genet. Appl. Livest. Prod. Vol.. XIII*, pp. 407-418.

- Meyer, K. 1991. Estimating variances and covariances for multivariate animal models by restricted maximum likelihood. *Genet. Sol. Evol.* 23 : 67-68.
- Meyer, K. and Burnside, E.B. 1988. Joint sire and cow evaluation for conformation traits using an individual animal model, *J. Dairy Sci.* 71(4) : 1034-1049.
- Meyer, K. and Hill, W.G. 1992. Approximation of sampling variances and confidence intervals for maximum likelihood estimates of variance components. *J. Anim. Breed. Genet.* 109 : 264-280.
- Meyer, K. and Thompson, R. 1984. Bias in variance and covariance component estimators due to selection on a correlated trait *J. Anim. Breed. Genet.* 101 : 33-5
- Misztal, I, 1994a. Comparison of software packages in animal breeding. In : Proc. 5th World Congr. Genet. Appl. Livest. Prod. during Aug. 7-12, university of Guelph, Guelph, Ontario, Canada vol. 21 pp. 3-10.
- Misztal, I. 1992. Derivative free vs Expectation maximization restricted maximum likelihood algorithms. In Program and Abstracts on the 87th annual meeting of American Dairy Science. Association held at Ohio Univ. during June 21-24, p. 316.
- Misztal, I. 1994b. Comparison of computing properties of derivative and derivative free algorithms in variance components estimation by REML. *J. Anim. Breed. Genet.* (Cited from misztal I, 1994. In : Proc. 5th World Congr. Genet. Appl. Livest. Prod. Vol. 22, pp. 3-10. 55.
- Mrode R.A.2005. Linear Models for the Prediction of Animal Breeding Values. 2nd ed. CAB International, Wallingford, UK.
- Nelder, J.A. and Mead, R. 1965. A simple method for function minimization Computer, J. 7 : 147-151.
- Paterson, H.D. and Thompson, R. 1971. Recovery of inter block information when block sizes are unequal. *Biometrika.* 58(3): 545-554.
- Quaas, R.L. 1976. Computing the diagonal elements and universe of a large numerator relationship matrix. *Biometrics* 32(4): 949-953.
- Raheja, K.L. 1992. Selection free estimates of genetic parameters of production and reproduction traits of first three lactation in Murrah buffaloes. *Indian J. Anim. Sci.* 62(2): 149-154.
- Raheja, K.L.; Vinayak, A.K. and Kalra, S. 2000. Genetic and phenotypic parameters estimated from single and multi-trait restricted maximum likelihood procedure. *Indian J. Anim. Sci.* 70(5) : 497-500.
- Raheja, K.L.; Vinayak, A.K. and Kalra, S. 2001. Comparison of four methods of variance components estimation for heritability of economic traits in Murrah buffalo. *Indian J. Anim. Sci.* 71(2): 146-149.
- Rao, C.R. 1971. Minimum variance quadratic estimation of variance components. *J. of Multivariate Analysis.* 1 : 445-456. CAB
- Robertson, A. 1977. The effect of selection in the estimation of genetic parameters. *J. Anim. Breed. Genet.* 94: 131-135.
- Robinson, J.A.B. 1988. Multiple lactation variance component estimation using restricted maximum likelihood. *Ph.D. Thesis, Cornell Univ., Ithaca, New York, U.S.A.*

- Schaeffer, L. R.(1983). Effectiveness of model for cow evaluation intraherd. *J. Dairy Sci.* 66:874–880
- Schaeffer, L.R. 1979. Estimation of variance and covariance components for average daily gain and back fat thickness in swine. In: Proc. Confer. In honor of C.R. Henderson on variance components in Animal Breeding, pp. 123-137.
- Schaeffer, L.R. 1986. Estimation of variances and covariances within the allowable parameter space *J. Dairy Sci.* 69(1): 187-194.
- Schaeffer, L.R. and Soong, H. 1979. Selection bias and REML variance covariance components estimation. *J. Dairy Sci.* 61(suppl 1): 91 (Abstract).
- Schaeffer, L.R.; Wilton, J.W. and Thompson, R. 1978. Simultaneous estimation of variance and covariance components from multi-trait mixed model equations. *Biometrics*, 34(2): 199-208.
- Searle, S.R. 1989. Variance components. Some history and summary accounts of estimation methods. *J. Anim. Bred. Genet.* 106 : 1-29.
- Searle, S.R.; Casella, G. and Mc Culloch, C.E. 1992. Variance components. John Wiley & Sons Inc, New York.
- Shaw, R.G. 1987. Maximum likelihood approaches applied to quantitative genetics of natural population evaluation, 41: 812-826.
- Smith, S.P. and Graser, H.U. 1986. Estimating variance components in a class of mixed models by Restricted Maximum Likelihood. *J. Dairy Sci.* 69(4): 1156-1165.
- Sorensen, D.A. and Kennedy, B.W. 1984. Estimation of genetic variances from unselected and selected populations. *J. Anim. Sci.* 59(5): 1213-1223.
- Sorensen, D.A. and Kennedy, B.W. 1986. Analysis of selection experiments using mixed model methodology *J. Anim. Sci.*, 63(1): 245-258. 76.
- Sun, C., Madsen, P., Lund, M. S., Zhang, Y., Nielsen, U. S.& Su, G.2010. Improvement in genetic evaluation of female fertility in dairy cattle using multiple-trait models including milk production traits. *J. Anim. Sci.* 88: 871-878.
- Sun, C., Madsen, P., Nielsen, U. S., Zhang, Y., Lund, M. S.& Su, G.2009. Comparison between a sire model and an animal model for genetic evaluation of fertility traits in Danish Holstein population. *J. Dairy Sci.* 92(8): 4063-4071.
- Taylor, J.F.; Bean, B.; Marshall, C.E. and Sullivan, J.J. 1985. Genetic and environmental components of semen production traits of artificial insemination in Holstein bulls. *J. Dairy Sci.* 68(10): 2703-2722.
- Thompson, R. 1973. The estimation of variance and covariance components with an application when records are subject to culling *Biometrics*. 29(3): 527-550.
- Thompson, R. and Atkins, K.D. 1990. Estimation of heritability from selection experiments. In: *Proc. 4th World Congr. Genet. Appl. Livest. Prod.* vol. 13, pp 257-260.
- Thompson, R. and Juga, J. 1988. A derivative-free approach for estimating variance and covariance components in bivariate animal models by Restricted Maximum Likelihood. *Acta Agric. Scand.* (submitted).

- Thompson, R. and Juga, J. 1989. A derivative free approach for estimating variance and co-variance components in bivariate animal models by Restricted Maximum Likelihood. *Acta Agriculture Scandinavica* (Cited from Meyer, 1989a).
- Thompson, R. and Meyer, K. 1986. Estimation of variance components: what is missing in the EM algorithm? *Statistical Computing and Simulation* 24: 215-230 (cited from Meyer, K. 1989 a).
- Thompson, W.A. Jr. 1962. The problem of negative estimates of variance components. *Ann. Math. Stat.* 33 : 273-289.
- Townsend, E.C. and Searle, S.R. 1971. Best quadratic unbiased estimation of variance components from unbalanced data in the 1-way classification. *Biometrics*. 27(3): 643-657.
- Townsend, E.C. 1968. Unbiased estimators of variance components in sample unbalanced designs. *Ph.D. Thesis, Cornell Univ.* Ithaca, New York. CAB.
- Van Raden, P.M. and Freeman, A.E. 1986. Computing restricted maximum likelihood estimates of variance and covariance. *J. Dairy Sci.* 69 (Suppl. 1) : 208 (Abstract).
- Webb, A.J. And Bampton, P.R. 1990. Impact of the new statistical technology on pig improvement. In : New developments in sheep production. Slade, C.F.R. and Lawrence, T.J.L. (editors) occasional publication *British Soc. Anim. Prod.* No. 14, pp. 166-168.
- Westfall, P.H. 1987. A comparison of variance components estimates for arbitrary underling distributions. *J. Anim. Statistical. Assoc.* 82(399): 866-874.
- Wiggans, G.R. and Misztal, I. 1987. Super computer for animal model evaluation of Ayrshire milk yield *J. Dairy Sci.* 70(9) : 1906-1912