

20000814146

EFFICIENT IDENTIFICATION OF
IMPORTANT FACTORS IN LARGE
SCALE SIMULATIONS

by

Carl A. Mauro

Applied Research in Statistics - Mathematics - Operations Research

ADA 173 497

**EFFICIENT IDENTIFICATION OF
IMPORTANT FACTORS IN LARGE
SCALE SIMULATIONS**

by

Carl A. Mauro

TECHNICAL REPORT NO. 123-2

October 1986

**This study was supported by the U.S. Navy under
Contract No. N00014-84-C-0573**

**Reproduction in whole or in part is permitted
for any purpose of the United States Government**

Approved for public release; distribution unlimited

TABLE OF CONTENTS

	<u>Page</u>
ABSTRACT	1
1. INTRODUCTION	1
2. A SCREENING MODEL	3
3. SUPERSATURATED SCREENING DESIGNS	4
3.1 Random Balance Designs	4
3.2 Systematic Supersaturated Designs.	8
3.3 Group Screening Designs.	9
3.4 Modified Group Screening Designs	11
3.5 T-Optimal Designs.	12
3.6 R-Optimal Designs.	13
3.7 Search Designs	16
4. SUMMARY DISCUSSION	17
REFERENCES	18

ABSTRACT

Large, complex computer simulation models can require prohibitively costly and time-consuming experimental programs to study their behavior. Therefore we may want to concentrate the analysis on the set of "most important" factors (i.e., input variables). Factor screening experiments, which attempt to identify the more important variables, can be extremely useful in the study of such models. The number of computer runs available for screening, however, is usually severely limited. In fact, the number of factors often exceeds the number of available runs. In this paper we present a survey of supersaturated designs for use in factor screening experiments. The designs considered are: random balance, systematic supersaturated, group screening, modified group screening, T-optimal, R-optimal, and search designs. We discuss in general terms the basic technique, advantages, and disadvantages of each procedure surveyed.

1. INTRODUCTION

Large-scale computer simulation models, because of their size and running time, can require prohibitively costly and time-consuming experimental programs to study their behavior. Often it is anticipated, however, that only relatively few factors (i.e., input variables) will have major effects. Therefore, one may want to conduct an efficient preliminary experiment to determine the subset of "most important" factors. Once the most important factors have been identified, subsequent experimentation can focus on these particular factors, thereby eliminating experimentation with relatively unimportant factors which can needlessly consume resources.

In some situations, of course, prior information concerning

the experimental factors is available and can be used to identify those factors which are most likely to be of importance.

Although the use of prior information can be beneficial to the screening process, in this paper we assume that there is no prior information available. This would not exclude, for example, situations in which prior information is available on the importance of some of the factors, and it is desired to examine the remaining factors in a screening experiment.

The function of a factor screening experiment is to sort the factors into two groups. One group consists of the important factors which are judged worthwhile to investigate further, while the other consists of the remaining factors classified as unimportant. In general, in screening experiments we want (a) to detect as many important factors as possible, (b) to declare important as few unimportant factors as possible, and (c) to expend as few computer runs as possible. Since these are conflicting objectives, one must generally trade off how many runs a screening method requires against how accurately it classifies factors.

The screening problem can occur in two general situations. These are the unsaturated/saturated and the supersaturated situations. In the unsaturated/saturated situation, one can afford to invest more runs than there are factors. In the supersaturated situation, the number of factors equals or exceeds the number of runs available for screening. Although screening can be done more effectively in the unsaturated/saturated situation, the supersaturated situation is a common and practical situation in the analysis of large-scale simulation models. Here, once again, design economy is the primary consideration.

Supersaturated design procedures are not customarily discussed in textbooks on experimental design, and there are few

examples of such experiments in the statistical or simulation literature. This paper presents a survey of supersaturated designs for use in factor screening, with application to large-scale simulation models. The designs considered are: (i) random balance designs, (ii) systematic supersaturated designs, (iii) group screening designs, (iv) modified group screening designs, (v) T-optimal designs, (vi) R-optimal designs, and (vii) search designs. Our intent is to provide a broad overview of supersaturated screening methods and to discuss in general terms the basic technique, analysis procedures, advantages, and disadvantages of each method surveyed. Appropriate references are provided if further information is desired.

2. A SCREENING MODEL

For the purpose of detecting the factors which have major effects, it generally suffices to assume the first-order model

$$y_i = \beta_0 + \sum_{j=1}^K \beta_j x_{ij} + \epsilon_i$$

where y_i is the value of the response (i.e., output variable) in the i th simulation run, K is the total number of factors, each of which is at two normalized levels (coded ± 1), x_{ij} is the level of the j th factor during the i th simulation run, β_0 is a constant component common to all observations, β_j ($j \geq 1$) is the (linear) effect of the j th factor, and ϵ_i is a random error component with mean 0 and unknown variance σ^2 .

A common interpretation of this model is that it represents a first-order Taylor series approximation to the true relationship between the output y_i and the normalized input

variables $x_1, x_2, \dots, x_K$. Moreover, the coefficients $\beta_1, \beta_2, \dots, \beta_K$ can be related to the sensitivity of the output variable y_1 to changes made in the input variables, at least in the vicinity of their nominal values. Ordinarily this model would be used over a relatively small region of the factor space.

3. SUPERSATURATED SCREENING DESIGNS

Screening designs can be classified as either "fixed" or "sequential" designs. In a fixed, or nonsequential, design, the factors are screened based on a given set of observations (i.e., computer runs). In a sequential design, results from a first-stage design are used to provide information on how to set up the design used in the next stage, and so on. All of the designs considered in this paper are fixed, with the exception of group screening designs which are sequential.

3.1 Random Balance Designs

Random balance (RB) designs, introduced by Budne (1959a, 1959b) and Satterthwaite (1959), were discussed at length by Anscombe (1959) and Youden et al. (1959). See also Dempster (1960), Mauro and Smith (1984), and Mauro and Burns (1984).

In a two-level (± 1) RB design, each column of the design matrix consists of $N/2$ +1's and $N/2$ -1's where N , an even number, denotes the total number of runs to be made. The +1's and -1's in each column are assigned randomly, making all possible combinations of $N/2$ +1's and $N/2$ -1's equally likely, with each column receiving an independent randomization.

The principal advantage of the RB method is its flexibility. The number of runs N can be selected independently of the number

of factors K. No mathematical restriction or relationship, except that N be an even number, need exist between N and K. A second advantage is that RB designs are very easy to prepare for any combination of N and K. This latter advantage can be an important consideration when K is large.

The major disadvantage of RB designs is that confounding is random. Anscombe (1959) has written:

The fact that the degree of nonorthogonality or unbalance is random can be made the basis for an objection to the whole notion of random balance designs. Such designs may work well on the average, but should I trust to one on this occasion?

Indeed, the lack of control over the confounding in RB designs has been a controversial aspect since such experimentation was first proposed.

Another disadvantage of RB designs is that there is no generally accepted or established method of analysis for these designs. In fact, the problem of analysis is characteristic of supersaturated designs with irregular confounding patterns and is not peculiar to the RB method. The simplest analysis approach is to consider each factor separately, ignoring all other factors, and apply some standard analysis technique such as an F-test. More sophisticated analysis methods which can be used include variable selection procedures such as least squares stepwise and stagewise regression (see, for example, Draper and Smith 1981).

The simple least squares estimator of β_j obtained by ignoring all other factors is given by

$$\hat{\beta}_j = (\bar{y}_{+j} - \bar{y}_{-j})/2$$

where $\bar{y}_{+j}(\bar{y}_{-j})$ is the average value of the response over the N/2

runs at the +1(-1) level of the j th factor. Let $\underline{y}$ denote the $N \times 1$ vector $(y_1, y_2, \dots, y_N)'$ of responses and $\underline{x}_j$ denote the $N \times 1$ vector $(x_{1j}, x_{2j}, \dots, x_{Nj})'$. Further, let $\underline{X} = [\underline{1}, \underline{x}_1, \underline{x}_2, \dots, \underline{x}_K]$, where $\underline{1}$ designates an $N \times 1$ vector of +1's. In an RB design, the matrix $\underline{X}$ is, by construction, stochastic (except, of course, for the initial column of +1's). Assuming that $\underline{X}$ and the ϵ_i are independent, it is easily shown that conditional on $\underline{X}$,

$$E(\hat{\beta}_j | \underline{X}) = \beta_j + \left(\sum_{i \neq j} \beta_i \underline{x}_i' \underline{x}_j \right) / N \quad (1)$$

and

$$V(\hat{\beta}_j | \underline{X}) = \sigma^2 / N. \quad (2)$$

The conditional mean square error (MSE) of $\hat{\beta}_j$ is then

$$MSE(\hat{\beta}_j | \underline{X}) = \sigma^2 / N + \left(\sum_{i \neq j} \beta_i \underline{x}_i' \underline{x}_j \right)^2 / N^2.$$

Unconditionally, it can be shown that

$$E(\hat{\beta}_j) = \beta_j \quad (3)$$

and

$$V(\hat{\beta}_j) = \tau_j^2 / (N-1) + \sigma^2 / N \quad (4)$$

where $\tau_j^2 = \sum_{m \neq j} \beta_m^2$.

As Box (1959) pointed out, equations (1) and (2) refer to the behavior of the estimates for repetitions of a particular RB design. Equations (3) and (4), on the other hand, refer to the

behavior of the estimates if we average over the random choice of RB designs. Box noted that although $\hat{\beta}_j$ is unconditionally unbiased, the effect of the conditional bias term in (1) is transferred to the unconditional variance of $\hat{\beta}_j$ which now contains terms from every other factor present. Conditionally, therefore, one pays the price of having a biased estimator; unconditionally, one pays the price of having a potentially inflated variance. From either point of view, RB sampling would seem inefficient for detecting all but the very large effects. In a study of the RB method, Mauro and Burns (1984) derived formulas for determining (unconditional) power probabilities of detecting factor effects when separate F-tests are used as the method of analysis.

Some further results regarding the use of separate F-tests are that for $i \neq j$

$$\text{cov}(\hat{\beta}_i, \hat{\beta}_j) = \beta_i \beta_j / (N-1)$$

and
$$\text{corr}(\hat{\beta}_i, \hat{\beta}_j) = \beta_i \beta_j / (\tau_i^2 + \sigma^2)(\tau_j^2 + \sigma^2). \quad (5)$$

The correlation expressed in equation (5) is a measure of the confounding between $\hat{\beta}_i$ and $\hat{\beta}_j$. It is interesting to note that an increase in N does not reduce the confounding in an RB design where simple least squares is used as the estimation method. As indicated by (5), the confounding between $\hat{\beta}_i$ and $\hat{\beta}_j$ is primarily a function of σ^2 and the magnitudes of the effects in the model.

It should also be noted that more sophisticated analysis techniques such as variable selection procedures are not immune to the adverse effects of the irregular confounding patterns characteristic of supersaturated designs. It is well known that multicollinearity in the predictor variables can cause serious computational and statistical difficulties. See, for example,

Belsley, Kuh, and Welsch (1980) and Silvey (1969). Furthermore, sequential selection procedures have the added problem that it is difficult to control the true overall significance level (i.e., the probability of declaring important a negligible factor) of such procedures.

3.2 Systematic Supersaturated Designs

Because of the random confounding that occurs in RB designs, Booth and Cox (1962) introduced balanced (i.e., an equal number of +1's and -1's in each design column) two-level designs which systematically attempt to minimize confounding. Since not all design columns can be orthogonal when $N \leq K$, Booth and Cox constructed designs, which they termed systematic supersaturated (SS) designs, that minimize $\max_{i \neq j} |c_{ij}|$ where $c_{ij} = \sum x_i x_j$. For two or more designs with the same minimax value, the preferred SS design is the one which minimizes the number of pairs of columns attaining the minimax value.

The cosine of the angle θ_{ij} , $0 \leq \theta_{ij} \leq \pi$, between any two column vectors, $\underline{x}_i$ and $\underline{x}_j$, is defined by

$$\cos \theta_{ij} = c_{ij} / (c_{ii} c_{jj})^{1/2} = c_{ij} / N.$$

The vectors $\underline{x}_i$ and $\underline{x}_j$ are orthogonal if $c_{ij} = 0$ or, in other words, if the cosine of the angle between them is zero. The absolute value of the inner product, therefore, is a measure of the orthogonality of any two design columns. In a certain sense, then, SS designs are constructed as nearly orthogonal as possible.

Booth and Cox tabulated their SS designs for the following seven combinations of (N, K) : (12,16), (12,20), (12,24), (18,24), (18,30), (18,36), and (24,30). As they pointed out, designs for

intermediate values of K can be formed by dropping the final columns from the next largest SS design. The designs indicated above were obtained with the aid of an iterative computer search procedure, since it was impracticable to enumerate all possible designs and select the best.

The principal advantage of SS designs is that they attempt to minimize the confounding which inevitably occurs when $N \leq K$. The principal disadvantage is that these designs are not readily available for combinations of N and K other than those already tabulated. Furthermore, the time and expense needed to write and to run a computer program for the generation of these designs may be prohibitive, especially if K is large. Also, as in RB designs, there remains the difficulty that the analysis of SS designs is complicated by the confounding of factor effects.

A quick comparison of SS with RB designs can be made based on an analysis of the variance of the inner product of two columns chosen at random from the design. For RB designs, for example, the variance of $x_i x_j$ for any i and j is $N^2/(N-1)$. Booth and Cox made such a comparison for the seven SS designs they derived and observed that SS designs are substantially better than RB designs when $N > K/2$. As would be expected, SS designs lose their advantage when N is small relative to K.

3.3 Group Screening Designs

In a group screening (GS) design we partition the individual factors into groups of suitable sizes and then test these groups by considering each as a single factor. The level of a "group-factor" is defined by assigning the same level, either +1 or -1, to each component factor. Because the number of group-factors is generally much smaller than the original number of factors, we can usually study the group-factors in a standard

orthogonal design such as a Plackett-Burman (PB) design. PB designs are two-level (+1) orthogonal designs for studying up to $K=4m-1$ factors in $N=4m$ runs. PB designs were tabulated by Plackett and Burman (1946) for $N \leq 100$ and are minimal in the number of runs required to achieve orthogonality. When N is a power of two, PB designs are the same as resolution III 2^{k-p} fractional factorial designs (see Box and Hunter 1961).

The grouping and testing process can be repeated for any number of stages. In each stage, however, we repartition into smaller groups only those groups determined to be significant in the previous stage. Further, we hold at a constant level any factor not included in a subsequent stage so not to bias any of the later-stage group-factor estimates. In the final stage of screening, we test factors individually (i.e., the group size is unity).

GS designs were introduced initially by Watson (1961), who considered screening in two stages. The GS method was then generalized to more than two stages by both Li (1962) and Patel (1962). For an excellent overview of GS designs, the reader may consult Kleijnen (1975).

There are two major advantages to GS designs. The first of these is that we can to a certain extent control the confounding pattern, since factors within a group are completely confounded and factors in different groups are not confounded. Secondly, the grouping process reduces the dimensionality of the model and enables the use of orthogonal main effect designs, such as PB designs, to test the significance of group-factors. Moreover, such designs can be analyzed by the usual analysis of variance procedures for factorial experiments.

There are two major disadvantages of the GS method. The

first of these is that the total number of runs required by a GS procedure is not fixed, since the number of group-factors carried over from stage to stage (when one progresses beyond the first stage) is random. Thus, in a GS strategy, one generally does not know prior to experimentation the exact total number of runs that will be expended. A second major disadvantage is that the optimal choice of group sizes and significance levels used in the various stages of group screening requires prior information on certain properties of the underlying model, such as the proportion of important effects. The problem of selecting optimal two-stage group screening designs has been discussed by Mauro (1984) and Patel and Ottieno (1984).

Another consideration in the use of the GS approach is that important effects may cancel within a group. As a simple example, consider two factors which have effects that are negatives or near negatives of each other. If these two factors are the only important factors in a group, their effects will essentially cancel and their combined effect may be masked by experimental error. Cancellation of effects cannot occur, of course, if factor levels are assigned a priori so that all effects are in the same direction. See Mauro (1984), Mauro and Burns (1984), and Mauro and Smith (1982) for a more detailed analysis of the effects of cancellation on the performance of two-stage GS designs.

3.4 Modified Group Screening Designs

Because an analyst might be reluctant to use a screening strategy in which the total number of runs cannot be predetermined, Mauro and Burns (1984) suggested a modified GS procedure in two stages where the total number of runs can be fixed prior to experimentation. Consider a two-stage GS strategy where the analyst decides beforehand not only on the number of

group-factors which will be tested in the first stage but also on the number of group-factors, say m , which will be carried over to the second stage. After the first-stage experiment, the m group-factors with the largest estimated effects are determined and their component factors tested individually in a second-stage experiment. In this strategy, the number of first- and second-stage runs are both predetermined.

The advantages of modified GS designs are the same as for regular GS designs. The same disadvantages also apply except of course that the total number of runs is fixed in modified GS and is random in regular GS. An additional disadvantage of modified GS is that the prespecification of m , the number of group-factors to be carried over to the second stage, may be inadequate in order to reasonably insure that all of the apparently significant groups reach the second stage.

Further research and practical experience on modified GS designs is needed. However, preliminary indications are that the performance of this strategy is comparable to that of regular two-stage GS if the proportion of important effects is not too large.

3.5 T-Optimal Designs

In a two-level (± 1) design, the inner product between any two design columns is a measure of their orthogonality. SS designs, which were discussed in Section 3.2, are constructed with the objective of minimizing the maximum absolute inner product between any two distinct design columns. One can define other criteria, however, for measuring the optimality of a supersaturated design.

We define a design to be T-optimal in a given class of

designs if it minimizes, over all designs in that class, the trace of $(\underline{X}'\underline{X})^2$, where $\underline{X}$ is as defined previously. Equivalently, a design is T-optimal if it minimizes the sum of squared inner products of all pairs of columns in $\underline{X}$. Thus, the columns of $\underline{X}$ are, in a certain sense, as nearly orthogonal as possible.

The principal disadvantage of supersaturated T-optimal designs is that rules for their general construction have not been developed, nor have any such designs been tabulated within the class of two-level (± 1) designs. However, in studies where the constant term β_0 is known or an advance estimate is available, rules for the construction of supersaturated T-optimal designs can essentially be found in Morris and Mitchell (1983), who derived such designs in the process of obtaining their trace-L optimal designs for detecting two-factor interactions. A second disadvantage is that, as in RB and SS designs, the analysis of T-optimal designs is made difficult by the confounding of factor effects.

3.6 R-Optimal Designs

In matrix terms, the screening model introduced in Section 2 can be expressed as

$$\underline{y} = \underline{X}\underline{\beta} + \underline{\epsilon}$$

where $\underline{y}$ is an $N \times 1$ vector of responses, $\underline{\beta} = (\beta_0, \beta_1, \dots, \beta_K)'$ is a $(K+1) \times 1$ vector of parameters, $\underline{\epsilon}$ is an $N \times 1$ vector of random error terms with mean 0 and variance σ^2 , and $\underline{X}$ is an $N \times (K+1)$ matrix of coefficients of the parameters $\beta_0, \beta_1, \dots, \beta_K$. We shall assume that $N \leq K$ and that $\underline{X}$ is of rank N . For simplicity in the following discussion, we shall also assume that $\sigma^2 = 0$, so that $\underline{y} = \underline{X}\underline{\beta}$.

Consider, then, the supersaturated system of linear equations $y = X\beta$. Since $N \leq K$, this system is underdetermined and therefore possesses infinitely many solutions. It can be shown, however, that the solution which has minimum length is given by

$$\hat{\beta}_m = X'(XX')^{-1}y. \quad (6)$$

Because $\hat{\beta}_m$ is the minimum length solution, it insures, in some sense, that factors are treated equally in the estimation process. For example, if two or more factors are completely confounded, the effect of each factor will be estimated by the average of their combined effects.

Substituting $y = X\beta$ into (6), we obtain, in terms of the true β ,

$$\hat{\beta}_m = R\beta$$

where $R = X'(XX')^{-1}X$. We observe that R is a $(K+1) \times (K+1)$ symmetric idempotent matrix of rank N and is the projection matrix operator onto the space spanned by the rows of X . In addition, since R is a projection matrix, we note that $0 \leq r_{ii} \leq 1$, where r_{ii} designates the i th diagonal element of R . Furthermore, the sum of all diagonal elements of R equals N , since the trace of R equals its rank.

The notion of R -optimal designs was introduced by Mitchell, Hunter, and Showers (1980), who considered a Bayesian analysis of the supersaturated screening problem. Assuming that $\sigma^2 = 0$ and no constant term β_0 is present and making certain prior assumptions regarding $\beta_1, \beta_2, \dots, \beta_K$, $\hat{\beta}_m$ was obtained as the Bayes estimate of β . It may be noted that the coefficient matrix X , in this case, does not include a column of +1's corresponding to the constant term β_0 .

The variance-covariance matrix of the posterior distribution of β , as derived by Mitchell, Hunter, and Showers, was found to be proportional to the matrix $I - R$. Consideration of this result lead to the following design criteria. A design is said to be R-optimal in a given class of designs if it minimizes, over all designs in that class, the maximum diagonal element of R . This amounts to making the diagonal elements of R , hence also $I - R$, as nearly equal as possible (since the trace of R is fixed). Thus, an R-optimal design would, in some sense, provide equal information on all the β_i 's. It follows that if a design exists such that all r_{ii} are equal, that design is R-optimal.

In a related Bayesian treatment of this problem, Anscombe (1963) imposed certain restrictions on the coefficient matrix X in order to expedite calculation of the posterior distribution. One of these was that the rows of X are orthogonal. It is easy to show that if X is a row-orthogonal matrix, the diagonal elements of R are all equal, and consequently the associated design is R-optimal. This indicates that, for certain values of N and K , a supersaturated R-optimal design can be readily obtained by transposing a $(K+1) \times N$ ($N \leq K$) column-orthogonal matrix (where a row of +1's is reserved as coefficients of the constant term β_0).

The principal advantage to R-optimal designs is that, under certain restrictions, estimation of β is isotropic, i.e., the posterior variances of the β_i are all equal. The principal disadvantage of such designs is that their general performance characteristics in factor screening experiments have not been fully evaluated.

3.7 Search Designs

The factor screening problem can also be formulated as a "search" problem, following Srivastava (1975). Suppose it is known that at most k of the K factor effects are non-zero and the remaining $(K-k)$ effects are zero. The goal of the search problem is to determine the k non-zero effects and estimate them.

Srivastava developed design criteria for obtaining search designs in the case where no experimental error is present (i.e., $\sigma^2=0$). We can write the coefficient matrix $\underline{X}$ as $\underline{X}=(\underline{1} \ \underline{D})$, where $\underline{1}$ is an $N \times 1$ vector of +1's and $\underline{D}$ is an $N \times K$ design matrix. In order to ensure that the k non-zero effects can be uniquely determined from the $K!/k!(K-k)!$ possible subsets, every subset of $2k$ columns of $\underline{D}$ together with the column vector $\underline{1}$ must have a combined rank of $2k+1$. This implies, of course, that $N \geq 2k+1$.

Smaller sized search designs are possible, however, under the further restriction that the response vector $\underline{y}$ does not lie in the intersection of two competing subspaces. An equivalent restriction is that no linear relationships exist among the β_i 's. In this case, $\underline{X}$ may be chosen so that no two $(k+1)$ -dimensional column subspaces are identical, where each subspace includes the column vector $\underline{1}$. This implies, of course, that $N \geq k+1$.

The principal advantage of search designs is that theoretically the non-negligible effects should be identified and estimated with reasonable power, since this is an inherent condition of the construction of such designs.

There are three major disadvantages of search designs. The first of these is that construction of two-level (± 1) search designs is extremely difficult, particularly for large-scale simulation studies. A second disadvantage is that it is assumed that the maximum number of non-negligible effects, k , is known. It is unclear, however, what impact misspecification of k will

have on the search procedure. Finally, the analysis of search designs, which is based on subset regressions, may require a prohibitive number of computations even for moderate values of N , K , and k .

4. SUMMARY DISCUSSION

In this paper we have presented a description and comparative discussion of eight different types of supersaturated designs which have been suggested for use in factor screening experiments, with application to the study of large-scale computer simulation models. Because of the lack of comparative performance data, there are currently no definitive guidelines for the selection and use of supersaturated screening methods. Nevertheless, of those methods surveyed, the group screening method has been generally recommended, and we would concur with this recommendation except when the number of runs relative to the number of factors is severely limited. In such a case, Mauro and Burns (1984) found that the performance of group screening can be extremely poor, even for detecting the large effects. In such situations, then, alternative design strategies, such as systematic supersaturated designs, should be considered. From a practical point of view, although the screening plans considered in this paper are appealing, further theoretical development of these and other methods is needed, particularly in relation to the study of computer simulations per se.

REFERENCES

- Anscombe, F. J. (1959). Quick analysis methods for random balance screening experiments. Technometrics 1, 195-209.
- Anscombe, F. J. (1963). Bayesian inference concerning many parameters with reference to supersaturated designs. Bulletin of the International Statistical Institute 40, 721-733.
- Belsley, D. A., Kuh, E., and Welsch (1980). Regression Diagnostics: Identifying Influential Data and Sources of Collinearity. John Wiley & Sons, Inc., New York.
- Booth, K. H. V. and Cox, D. R. (1962). Some systematic supersaturated designs. Technometrics 4, 489-495.
- Box, G. E. P. (1959). Discussion of the papers of Messrs. Satterthwaite and Budne. Technometrics 1, 174-180.
- Box, G. E. P. and Hunter J. S. (1961). The 2^{k-p} fractional factorial designs, part I and part II. Technometrics 3, 311-351 and 449-458.
- Budne, T. A. (1959a). The application of random balance designs. Technometrics 1, 139-155.
- Budne, T. A. (1959b). Random balance: Part I -- The missing statistical link in fact finding techniques. Industrial Quality Control 15, 5-10.
- Dempster, A. P. (1960). Random allocation designs I: On general classes of estimation methods. Annals of Mathematical

Statistics 31, 885-905.

Draper, N. R. and Smith, H. (1981). Applied Regression Analysis,
Second Edition. John Wiley & Sons, Inc., New York.

Kleijnen, J. P. C. (1975). Statistical Techniques in
Simulation - Part II. Marcel Dekker, Inc., New York.

Li, C. H. (1962). A sequential method for screening experimental
variables. Journal of the American Statistical Association
57, 455-477.

Mauro, C. A. (1984). On the performance of two-stage group
screening. Technometrics 26, 255-264.

Mauro, C. A. and Burns, K. C. (1984). A comparison of random
balance and two-stage group screening designs: A case
study. Communications in Statistics A13, 2625-2647.

Mauro, C. A. and Smith, D. E. (1982). The performance of
two-stage group screening in factor screening experiments.
Technometrics 24, 325-330.

Mauro, C. A. and Smith, D. E. (1984). Factor screening in
simulation: Evaluation of two strategies based on random
balance sampling. Management Science 30, 209-221.

Mitchell, T. S., Hunter, W. G., and Showers, D. K. (1980).
Design of experiments for studying computer codes.
Unpublished manuscript, University of Wisconsin.

Morris, M. D. and Mitchell, T. J. (1983). Two-level multifactor
designs for detecting the presence of interactions.
Technometrics 25, 345-355.

Patel, M. S. (1962). Group screening with more than two stages. Technometrics 4, 209-217.

Patel, M. S. and Ottieno, J. A. M. (1984). Optimum two stage group-screening designs. Communications in Statistics A13, 2649-2663.

Plackett, R. L. and Burman, J. P. (1946). The design of optimum multifactor experiments. Biometrika 33, 305-325.

Satterthwaite, F. W. (1959). Random balance experimentation. Technometrics 1, 111-137.

Silvey, S. D. (1969). Multicollinearity and imprecise estimation. Journal of Royal Statistical Society, Series B, 31, 539-552.

Srivastava, J. N. (1975). Designs for searching non-negligible effects. In: A Survey of Statistical Design and Linear Models (J. N. Srivastava, ed.). North-Holland Publishing Company, New York, 507-519.

Watson, G. S. (1961). A study of the group screening method. Technometrics 3, 371-388.

Youden, W. J., Kempthorne, O., Tukey, J. W., Box, G. E. P., and Hunter, J. S. (1959). Discussions of the papers of Messrs. Satterthwaite and Budne. Technometrics 1, 157-193.

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM
1. REPORT NUMBER 123-2	2. GOVT ACCESSION NO.	3. RECIPIENT'S CATALOG NUMBER
4. TITLE (and Subtitle) EFFICIENT IDENTIFICATION OF IMPORTANT FACTORS IN LARGE SCALE SIMULATIONS		5. TYPE OF REPORT & PERIOD COVERED Technical Report
		6. PERFORMING ORG. REPORT NUMBER
7. AUTHOR(s) Carl A. Mauro		8. CONTRACT OR GRANT NUMBER(s) N00014-84-C-0573
9. PERFORMING ORGANIZATION NAME AND ADDRESS Desmatics, Inc., P.O. Box 618 State College, PA 16804		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS NR 042-529
11. CONTROLLING OFFICE NAME AND ADDRESS Office of Naval Research Arlington, VA 22217		12. REPORT DATE October 1986
		13. NUMBER OF PAGES 20
14. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office)		15. SECURITY CLASS. (of this report) Unclassified
		15a. DECLASSIFICATION/DOWNGRADING SCHEDULE
16. DISTRIBUTION STATEMENT (of this Report) Distribution of this report is unlimited.		
17. DISTRIBUTION STATEMENT (of the abstract entered in Block 20, if different from Report)		
18. SUPPLEMENTARY NOTES		
19. KEY WORDS (Continue on reverse side if necessary and identify by block number) Simulation Factor Screening Supersaturated Designs Important Factors		
20. ABSTRACT (Continue on reverse side if necessary and identify by block number) Large, complex computer simulation models can require prohibitively costly and time-consuming experimental programs to study their behavior. Therefore we may want to concentrate the analysis on the set of "most important" factors (i.e., input variables). Factor screening experiments, which attempt to identify the more important variables, can be extremely useful in the study of such models. The number of computer runs available for screening, however, is usually severely limited. In fact, the number of factors often exceeds the number of available runs. In this paper we present a survey of supersaturated designs for use in		

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE(When Data Entered)

factor screening experiments. The designs considered are: random balance, systematic supersaturated, group screening, modified group screening, T-optimal, R-optimal, and search designs. We discuss in general terms the basic technique, advantages, and disadvantages of each procedure surveyed.

END

DATE
FILMED

12-86