



National Library of Canada

Cataloguing Branch
Canadian Theses Division

Ottawa, Canada
K1A 0N4

Bibliothèque nationale du Canada

Direction du catalogage
Division des thèses canadiennes

NOTICE

The quality of this microfiche is heavily dependent upon the quality of the original thesis submitted for microfilming. Every effort has been made to ensure the highest quality of reproduction possible.

If pages are missing, contact the university which granted the degree.

Some pages may have indistinct print especially if the original pages were typed with a poor typewriter ribbon or if the university sent us a poor photocopy.

Previously copyrighted materials (journal articles, published tests, etc.) are not filmed.

Reproduction in full or in part of this film is governed by the Canadian Copyright Act, R.S.C. 1970, c. C-30. Please read the authorization forms which accompany this thesis.

**THIS DISSERTATION
HAS BEEN MICROFILMED
EXACTLY AS RECEIVED**

AVIS

La qualité de cette microfiche dépend grandement de la qualité de la thèse soumise au microfilmage. Nous avons tout fait pour assurer une qualité supérieure de reproduction.

S'il manque des pages, veuillez communiquer avec l'université qui a conféré le grade.

La qualité d'impression de certaines pages peut laisser à désirer, surtout si les pages originales ont été dactylographiées à l'aide d'un ruban usé ou si l'université nous a fait parvenir une photocopie de mauvaise qualité.

Les documents qui font déjà l'objet d'un droit d'auteur (articles de revue, examens publiés, etc.) ne sont pas microfilmés.

La reproduction, même partielle, de ce microfilm est soumise à la Loi canadienne sur le droit d'auteur, SRC 1970, c. C-30. Veuillez prendre connaissance des formules d'autorisation qui accompagnent cette thèse.

**LA THÈSE A ÉTÉ
MICROFILMÉE TELLE QUE
NOUS L'AVONS REÇUE**

15

A Study of a Technique
for
Validating Learning Hierarchies

by

Wai-Man Chan

Thesis presented to the School of Graduate
Studies of the University of Ottawa in
partial fulfillment of the requirements
for the degree of Master of Arts
(Education)

Ottawa, Ontario, 1979

(c) Wai-Man Chan, 1979

© W.-M. Chan, Ottawa, Canada, 1979

ACKNOWLEDGEMENT

In preparing this thesis, valuable advice and guidance were given to the researcher by Dr. Marvin Boss and Dr. Jean-Paul Dionne. Dr. Boss acted as thesis committee chairman.

The researcher is indebted to the Ontario Ministry of Colleges and Universities for its financial support in terms of an Ontario Graduate Scholarship.

TABLE OF CONTENTS

Chapter	page
I. Critical Review of the Literature	1
Learning hierarchies and their validation	1
The White & Clark Model	7
The MAX technique	11
Owston's design and findings on MAX	18
Generating random samples	19
Estimating risk of decision errors	22
Specifying population parameters	24
Generating datasets to examine the risks of decision errors	27
Setting criteria on risks of committing decision errors	28
Owston's findings on the risks of committing decision errors	29
Weaknesses of the MAX technique	31
The research problem	37
II. Design and Experimentation	39
Two procedures for implementing the MAX technique	39
Design for collecting data	44
Criteria for analysing the data	49
III. Presentation and discussions of results	53
Results obtained from using procedure A	54
Comparisons of Owston's results with the results obtained by using procedure A	67
Results obtained from using procedure B	74
Comparisons of results obtained from procedures A and B	83
Summary and Conclusions	91
Bibliography	100
Appendices	102
Abstract	127

LIST OF TABLES

Table	page
I. Probabilities of members of the samples being classified in particular cells.....	10
II. Parameters Estimation Using the Method of Scoring..	17
III. Schematic Representation of Owston's design for testing the MAX technique, over different decision criteria, sample sizes, measurement errors, skill difficulty levels and validity levels.....	25
IV. Error categories and error types identified when the MAX technique was pilot tested using Owston's computer program.....	32
V. Error types that are checked for by using procedure B, procedure A, or both procedures A and B.....	41
VI. Six sets of P's corresponding to 3 levels of skill difficulty and 2 levels of validity of learning hierarchy.....	47
VII. Schematic Representation of the design for applying procedures A and B to test the MAX technique over different decision criteria, sample sizes, measurement errors, skill difficulty levels and validity levels.....	50
VIII. Proportions of the various types of errors detected by using procedure A, for the different levels of validity, skill difficulty, probability of guessing and sample size.....	55
IX. Proportions of correct decisions and incorrect decisions at 95% and 90% confidence intervals, made by using procedure A, over the various combinations of validity, skill difficulty, probability of guessing and sample size levels.....	64

- X. Comparisons of the proportions of "correct decisions", at decision criteria of 95% and 90% confidence intervals, among results obtained from three sources (Owston's results; Results from procedure A; Results from procedure A with adjustments to include some "No decisions" cases as "correct decisions" cases), over the various combinations of validity (valid & invalid), skill difficulty, probability of guessing, and sample size levels..... 68
- XI. Proportions related to the various types of errors, detected by using procedure B, over the various combinations of validity, skill difficulty, probability of guessing and sample size levels..... 76
- XII. Proportions of correct and incorrect decisions, made by using procedure B at different confidence intervals, over the various combinations of validity, skill difficulty, probability of guessing and sample size levels..... 82
- XIII. Comparisons of the proportions of "No decisions" and the proportions of correct decisions, at 95% and 90% confidence interval, for procedures A and B, over the various combinations of validity, skill difficulty, probability of guessing and sample size levels..... 85
- XIV. Comparisons of the adjusted (to include some of the "No decisions" cases as "correct decisions" cases) proportions of "correct decisions" at 95% and 90% confidence intervals, for procedures A and B, over the different combinations of validity, skill difficulty, probability of guessing and sample size levels..... 86
- XV. Adjusted proportions of "correct decisions", at 85%, 80%, 75% and 70% confidence intervals, for procedure B, over the different combinations of validity, skill difficulty, probability of guessing and sample size levels..... 89

LIST OF FIGURES

Figure	page
1. Flowchart for Random Generation of Subjects' Responses to Two Questions per Skill for a Sample of Size N	20
2. A 3x3 pattern in which the error type "f=0" would occur	34
3. Hypothetical observed frequencies used in illustrating the method of calculating initial parameter estimates	105a

INTRODUCTION

One of the techniques for testing the validity of a learning hierarchy was carefully examined. This validating technique, called MAX, was developed by Owston. The maximum likelihood method for parameter estimation was used in the MAX technique.

With computer simulated data, the performance of the MAX technique was compared (by Owston) with that of other validating techniques. It was claimed that the performance of MAX was consistently better.

However, a number of weaknesses were identified in the original procedure for applying the MAX technique. These weaknesses are described near the end of the first chapter, after a brief review of the development of the MAX technique has been made.

Because of the weaknesses in the original procedure, modifications to this procedure appear necessary. Two modified procedures were developed and are discussed in Chapter II. These two procedures are basically the same as the original procedure, except that, the consequence of

ignoring the weaknesses is considered. The design for data collection is similar to that used in the original study. Computer simulated data are also used in this study. The results obtained from using the two modified procedures are presented and discussed in the last chapter. Some technical details involved in this study are shown in appendix A to E.

Chapter I

CRITICAL REVIEW OF THE LITERATURE

The concept of learning hierarchies has been used in different areas in the field of education. Among the techniques available to validate the logically derived learning hierarchies, Owston (1978) demonstrated that his MAX technique performed better than other techniques. However, there are some weaknesses in his application of the technique that may have biased the research findings.

In this chapter, after a brief review of the concepts relating to learning hierarchies and their validation methodologies, descriptions are made of the MAX technique and the general design of Owston's study. Then, weaknesses of the technique are discussed. In the last section of the chapter the research problem is discussed.

1.1 LEARNING HIERARCHIES AND THEIR VALIDATION

A learning hierarchy is a network of well-defined intellectual skills organized into levels, with the terminal skill(s) at the top level and subordinate skill(s) at the lower level(s); any subordinate skill is hypothesized to be

a prerequisite to the directly connected skill(s) in the higher level(s). The simplest (Durell, 1974, p. 4) learning hierarchy consists of two levels, with one skill in each level. A complicated hierarchy may be composed of numerous connected pairs (Owston, 1978, p. 20 and White, 1974, p. 61) of skills, each pair being equivalent to the simplest hierarchy.

The concept of learning hierarchies as networks of prerequisite relationships of intellectual skills was first introduced by Gagne (Gagne, 1962 and Gagne & White, 1974, p. 19). Uses of the concept are reported (Durell, 1974, p. 1 and Owston, 1978, p. 8-10) in the designing of individualized instruction, mastery learning sequences, computer-based instructional systems, and in subject areas (Gagne & White, 1974, p. 21) such as mathematics and science, where quantitative skills are essential.

To generate a learning hierarchy, a curriculum designer begins with the terminal skill(s) to be learned and asks the question (Gagne, 1962, p. 358) "What would the individual(s) have to be able to do in order that he(they) can attain successful performance on this task (skill expressed in behavioral terms)?" The answer to the question leads to the identification of immediately prerequisite skill(s). By reapplying the question to each of those newly identified

skills not yet possessed by the individual(s), subordinate skill(s) of lower level(s) are derived. The questioning process is repeated on the derived skills until (Cotton et. al., 1977, p. 190) the last identified skill can no longer be decomposed into simpler prerequisite skills which are not already possessed by the prospective learner(s).

For the simplest learning hierarchy, the process of its generation is as follows: On applying the question to the terminal skill, the curriculum designer can identify only one prerequisite skill the individual(s) has(have) to learn. On reapplying the question to the prerequisite skill, he finds that the prospective learner(s) does (do) not need to learn other subordinate skills for the attainment of the prerequisite skill. The process of questioning will stop, with two skills being in the learning hierarchy.

There is no certainty that hierarchies, derived by using the above procedures, are valid (Durell, 1974, p. 2). Two approaches are available to test the validity. The first approach is based upon the assumption that in training a learner on a certain skill, his training on an immediately subordinate skill will lead to positive transfer (Gagne 1965, p. 233 and Gagne, 1968) to the learning of the present skill. The second approach is based on the performance patterns of the learners. It involves utilizing the

performance pattern to test the hypothesis that, except by chance, learners cannot be successful in answering questions on a certain skill if they are not successful in answering questions on the prerequisite skill(s).

The first approach requires (White, 1973, p. 373) setting up a large number of experimental groups and instructional interventions even for a fairly simple hierarchy to be validated. For example, a hierarchy consisting of three skills may require setting up $3!$ (or 6) experimental groups; a hierarchy of k skills may require $k!$ groups. Hence, this approach is not commonly used. The second approach is less rigorous than the first one in ensuring positive transfer of learned capacities to a skill from its immediate prerequisite skill(s). However, due to the practical importance of requiring less time and instructional intervention, the second approach is widely used (Durell, 1974, p. 2 and White, 1973, p. 373).

Two classes (Owston, 1978, p. 29) of techniques for validating learning hierarchies may be identified in the second approach. They are deterministic techniques and probabilistic techniques. Deterministic techniques (White, 1974, p. 65) have one or both of the following types of short-comings. These shortcomings include the ignoring of errors of measurements in analysing test results of the

learners, and judging the criterion of skill possession in terms of the number of test items answered correctly.

White & Clark devised a probabilistic technique (White, 1974, p. 64 and White & Clark, 1973) which is free from the above short-comings. This technique is useful for validating two-skill learning hierarchies (the White & Clark Model of learning hierarchies; to be discussed in the next section). However, the White & Clark technique has been criticized (Owston, 1978, p. 62) as being weak in the estimation of parameters contained in the White & Clark Model. Information available for parameter estimation has not been fully utilized in this technique.

Owston (1978, p. 62-79), using some of the assumptions of the White & Clark Model, developed a probabilistic validating technique, called the MAX technique. This technique involves using the maximum likelihood method for estimating population parameters.

No matter whether it is a deterministic or a probabilistic technique, there may always be decision errors (Owston, 1978, p. 29-30) involved in using the technique. For a hierarchy that is valid, the application of a technique may lead to an indication of invalidity. For an invalid hierarchy, the result from applying the technique may cause the researcher to think that the hierarchy is valid.

Owston considered that the accuracy (Owston, 1978, p. 107) of a technique may be judged by the extent of its simultaneously satisfying predetermined probabilities for each type of decision error. The emphasis on simultaneity may be explained with an extreme example; a technique that indicates correctly all the hierarchies that are valid and indicates incorrectly all the hierarchies that are invalid would be useless as would one that performs in an opposite manner.

In trying to establish the accuracy of a technique or to compare two or more different techniques, one more complication exists; in real life, one never knows whether a derived hierarchy is valid or invalid, for if the underlying validity were known, no testing would be required (Durell, 1974, p. 3). Computer simulation provides a means of generating simulated learner performance patterns related to hierarchies of "known" validity. With computer simulation, the researcher can determine the validity of hierarchies by specifying population parameters. Also, he can easily vary the size of samples, size of measurement errors, difficulty levels, validity of the hierarchy, etc..

Using computer simulated data, Owston (1978) studied the influences of measurement errors, skill difficulty levels and sample sizes on the accuracies of the MAX

technique, the White & Clark technique, and two (Owston, 1978, p. 81-82) deterministic techniques, the Walbesser Ratios and the Difference Ratio. The learning hierarchies Owston generated were the simplest (Owston, 1978, p. 20 and White, 1974, p. 65) ones with only two skills involved. Based on his findings, he concluded (Owston, 1978, p. 141) that the MAX technique was more accurate than the other techniques over different skill difficulty levels, measurement errors and sample sizes commonly encountered in research.

To have a better understanding of the MAX technique, it is important to know the White & Clark Model upon which the MAX technique is based. It is also important to have an understanding of Owston's experimental design used in testing the accuracy of MAX. In the next few sections, the White & Clark Model and the MAX technique will be described, followed by a critical review of Owston's study and his technique.

1.2 THE WHITE & CLARK MODEL

White & Clark developed a probabilistic model (White & Clark, 1973 and Owston, 1978, p. 53-62) to test the validity of a hierarchical connection between a pair of skills, called skill I and skill II, with skill I being the subordinate skill. Two test items were used to test each skill. Measurement errors were considered.

In a population of students four mutually exclusive groups can be formed. The four groups, respectively, contain students who possess, neither skill I nor skill II, only skill I, only skill II, both skill I & skill II. The corresponding proportions of students within the population are denoted by P_0 , P_I , P_{II} and P_B .

For skill I, the two test items are assumed to be very similar to, but independent of, each other. This means that, for any student, the chance of correctly answering the first item is the same as the chance of correctly answering the second item; however, the actual outcomes (correct or incorrect) of answering the two items, are independent of each other. This is also true for the two items measuring skill II.

Based on the above assumption, probabilities of correctly answering the test items may be associated with each student in accordance with his state of skill possession. For a student who possesses skill I (he is either from the P_I or P_B group), his probability, of correctly answering each of the two independent skill I test items, is indicated by θ_a . For a student who does not possess skill I (he is either from the P_0 or P_{II} group), the corresponding probability of correctly answering each of the two skill I test items is indicated by θ_b . Similarly, two

probabilities are related to skill II test items. For a student who possesses skill II (he is either from the P_{II} or P_B group), his probability, of correctly answering each of the two independent skill II test items, is denoted by θ_c . For a student who does not possess skill II (he is either from the P_O or P_I group), the corresponding probability of correctly answering the two items is denoted by θ_d .

Note that θ_b and θ_d are related to the measurement errors of correct guessing on skill I and skill II test items respectively. Furthermore, note that $(1-\theta_a)$ and $(1-\theta_c)$ are related to the measurement errors of blundering (failing to answering correctly when possessing the skill), respectively, on skill I and skill II test items.

With the four population proportions, P_O , P_I , P_{II} and P_B , and the four probabilities of obtaining a correct answer, θ_a , θ_b , θ_c and θ_d , the White & Clark Model can be formed. Table I shows the model which contains nine cells and the corresponding theoretical probability functions associated with each cell.

Corresponding to the theoretical model shown in Table I, a 3x3 table of observed response patterns can be obtained from a sample of students who respond to the skill I and skill II test items.

TABLE I

Probabilities of Members of the Sample
being Classified in Particular Cells *

Skill II

Questions correct	0	1	2
2	$p_{20} = p_1$ cell 1 $= p_{00}^2 (1-\theta_d)^2$ $+ p_{10}^2 (1-\theta_d)^2$ $+ p_{11}^2 (1-\theta_c)^2$ $+ p_{20}^2 (1-\theta_c)^2$	$p_{21} = p_2$ cell 2 $= 2p_{00}^2 \theta_d (1-\theta_d)$ $+ 2p_{10}^2 \theta_d (1-\theta_d)$ $+ 2p_{11}^2 \theta_c (1-\theta_c)$ $+ 2p_{20}^2 \theta_c (1-\theta_c)$	$p_{22} = p_3$ cell 3 $= p_{00}^2 \theta_d^2$ $+ p_{10}^2 \theta_d^2$ $+ p_{11}^2 \theta_c^2$ $+ p_{20}^2 \theta_c^2$
1	$p_{10} = p_4$ cell 4 $= 2p_{00} \theta_b (1-\theta_b) (1-\theta_d)^2$ $+ 2p_{10} \theta_a (1-\theta_a) (1-\theta_d)^2$ $+ 2p_{11} \theta_b (1-\theta_b) (1-\theta_c)^2$ $+ 2p_{20} \theta_a (1-\theta_a) (1-\theta_c)^2$	$p_{11} = p_5$ cell 5 $= 4p_{00} \theta_b (1-\theta_b) \theta_d (1-\theta_d)$ $+ 4p_{10} \theta_a (1-\theta_a) \theta_d (1-\theta_d)$ $+ 4p_{11} \theta_b (1-\theta_b) \theta_c (1-\theta_c)$ $+ 4p_{20} \theta_a (1-\theta_a) \theta_c (1-\theta_c)$	$p_{12} = p_6$ cell 6 $= 2p_{00} \theta_b (1-\theta_b) \theta_d^2$ $+ 2p_{10} \theta_a (1-\theta_a) \theta_d^2$ $+ 2p_{11} \theta_b (1-\theta_b) \theta_c^2$ $+ 2p_{20} \theta_a (1-\theta_a) \theta_c^2$
0	$p_{00} = p_7$ cell 7 $= p_{00} (1-\theta_b)^2 (1-\theta_d)^2$ $+ p_{10} (1-\theta_a)^2 (1-\theta_d)^2$ $+ p_{11} (1-\theta_b)^2 (1-\theta_c)^2$ $+ p_{20} (1-\theta_a)^2 (1-\theta_c)^2$	$p_{01} = p_8$ cell 8 $= 2p_{00} (1-\theta_b)^2 \theta_d (1-\theta_d)$ $+ 2p_{10} (1-\theta_a)^2 \theta_d (1-\theta_d)$ $+ 2p_{11} (1-\theta_b)^2 \theta_c (1-\theta_c)$ $+ 2p_{20} (1-\theta_a)^2 \theta_c (1-\theta_c)$	$p_{02} = p_9$ cell 9 $= p_{00} (1-\theta_b)^2 \theta_d^2$ $+ p_{10} (1-\theta_a)^2 \theta_d^2$ $+ p_{11} (1-\theta_b)^2 \theta_c^2$ $+ p_{20} (1-\theta_a)^2 \theta_c^2$
	f	e	d

* Reproduced from White & Clark (1973, p. 77);

a, b, c, d, e, f are marginal frequencies.

Using the theoretical model and the observed 3x3 table together, it is possible to derive useful relationships for inferring validity of the learning hierarchy characterized by the population.

The MAX technique, to be described next, is one of the validating techniques based on the White & Clark Model of two-skills learning hierarchy.

1.3 THE MAX TECHNIQUE

The MAX technique (Owston, 1978, p. 62-79) consists in producing an interval estimate of the parameter P_{II} , which is the proportion of the population that possess the superordinate skill without possessing the subordinate skill. For a valid hierarchy P_{II} is equal to zero; for an invalid hierarchy P_{II} is larger than zero (but less than 1). The estimation is obtained through the application of the method of maximum likelihood. If the confidence interval includes zero, it is concluded that the hierarchy is valid, otherwise it is concluded that the hierarchy is invalid.

In the White & Clark Model, of the eight parameters ($P_O, P_I, P_{II}, P_B, \theta_a, \theta_b, \theta_c, \theta_d$), only seven have to be estimated because P_B can be expressed as $1 - (P_O + P_I + P_{II})$. For convenience in deriving formulae to show the estimation process, new symbols are used such that $P_O = \theta_1, P_I = \theta_2, P_{II} = \theta_3, \theta_a = \theta_4, \theta_b = \theta_5, \theta_c = \theta_6, \theta_d = \theta_7$.

It is assumed in Rao's approach, as presented by Owston, that in a random sample of n subjects, the observed frequencies in the nine cells in Table I happen to be $(f_1, f_2, \dots, f_8, f_9)$. The probability of occurrence of such an observed pattern can be expressed as a multinomial distribution function called the likelihood function (L).

Symbolically,

$$L = \frac{n!}{\prod_{i=1}^9 f_i!} \prod_{i=1}^9 p_i^{f_i} \quad (1)$$

where $0 \leq f_i \leq n$, $\sum_{i=1}^9 f_i = n$, and p_i is a function of the seven parameters, θ_1 to θ_7 .

The seven parameter estimates that maximize L are called the maximum likelihood estimates (m.l.e.) of the parameters.

To simplify the derivation, the natural logarithm $\text{Log } L$ is used; since it is a monotonic (Owston, 1978, p. 66 and Allendoerfer & Oakley, 1963, p. 209) increasing function of L , the estimates that maximize L will also maximize $\text{Log } L$. To find the m.l.e. from $\text{Log } L$, a scoring method can be used. Corresponding to the seven parameters to be estimated, there are seven efficient scores, the j^{th} efficient score (Rao, 1968, p. 302-305) is defined as the partial derivative of $\text{Log } L$ with respect to θ_j , that is

$$S_j = \frac{\partial \text{Log } L}{\partial \theta_j} \quad (2)$$

where
$$\frac{\partial \text{Log } L}{\partial \theta_j} = \sum_{i=1}^9 \frac{f_i \partial p_i}{p_i \partial \theta_j} \quad j = 1, 2, \dots, 7 \quad (3)$$

Theoretically, by setting the seven efficient scores in equation (2) to zeros, and solving for θ_1 to θ_7 , the solution (Rao, 1968, p. 302) thus obtained will be the maximum likelihood estimator. In practice, for estimates that are functions of several parameters, it is too complicated to proceed with that approach.

By using the scoring procedures, an initial set (Rao, 1968, p. 302) of parameter estimates $\hat{\theta}^0 = (\hat{\theta}_1^0, \hat{\theta}_2^0, \dots, \hat{\theta}_7^0)$ is used as a first approximation of the solution. The efficient scores in equation (2) are expanded about the initial set via the Taylor series. Only the first order terms in $\delta \theta_j^0$ are retained, where $\delta \theta_j^0 = \theta_j - \hat{\theta}_j^0$ is the deviation of the j^{th} initial parameter estimate from the j^{th} maximum likelihood estimate ($\delta \theta_j^0$ can be interpreted as a correction of the initial approximation). Following Taylor's expansion, the efficient scores are equated to zero, the following equation is obtained:

$$0 = S_j \Big|_a + \delta \theta_1^0 \frac{\partial^2 \text{Log } L}{\partial \theta_1 \partial \theta_j} \Big|_a + \dots + \delta \theta_6^0 \frac{\partial^2 \text{Log } L}{\partial \theta_6 \partial \theta_j} \Big|_a + \delta \theta_7^0 \frac{\partial^2 \text{Log } L}{\partial \theta_7 \partial \theta_j} \Big|_a + \dots \quad (4)$$

where $j=1,7$ and the symbol $\Big|_a$ implies evaluating the immediately preceding function at the point "a"; in this case "a" is the point $(\hat{\theta}_1^0, \hat{\theta}_2^0, \dots, \hat{\theta}_7^0)$.

For large samples (Rao, 1968, p. 302, and Owston, 1978, p. 68), the information for score j , I_{jk} , can be defined as

$$-\frac{\partial^2 \text{Log } L}{\partial \theta_j \partial \theta_k} \quad \text{where} \quad \begin{array}{l} k = 1, 2, \dots, 7 \\ j = 1, 2, \dots, 7 \end{array} \quad (5)$$

$$\text{that is } I_{jk} = -\frac{\partial^2 \text{Log } L}{\partial \theta_j \partial \theta_k} = n \sum_{i=1}^g \frac{1}{P_i} \frac{\partial P_i}{\partial \theta_j} \frac{\partial P_i}{\partial \theta_k} \quad (6)$$

Using a small circle on the top of the symbols to indicate initial values, equation (4) can be rewritten as

$$0 = \overset{\circ}{S}_j - \overset{\circ}{I}_{jk} \times \overset{\circ}{\delta\theta}_k + \dots \quad (7)$$

where $j, k=1, 2, 3, \dots, 7$. By rearranging the terms in (7),

$$\overset{\circ}{\delta\theta}_k = \overset{\circ}{I}_{jk}^{-1} \times \overset{\circ}{S}_{jk} \quad (8)$$

In matrix form, equation (8) becomes

$$\begin{bmatrix} \theta_1 - \hat{\theta}_1 \\ \theta_2 - \hat{\theta}_2 \\ \vdots \\ \theta_7 - \hat{\theta}_7 \end{bmatrix} = \begin{bmatrix} \overset{\circ}{\delta\theta}_1 \\ \overset{\circ}{\delta\theta}_2 \\ \vdots \\ \overset{\circ}{\delta\theta}_7 \end{bmatrix} = \begin{bmatrix} \overset{\circ}{I}_{11} & \overset{\circ}{I}_{12} & \dots & \overset{\circ}{I}_{17} \\ \overset{\circ}{I}_{21} & \overset{\circ}{I}_{22} & \dots & \overset{\circ}{I}_{27} \\ \vdots & \vdots & \ddots & \vdots \\ \overset{\circ}{I}_{71} & \overset{\circ}{I}_{72} & \dots & \overset{\circ}{I}_{77} \end{bmatrix}^{-1} \begin{bmatrix} \overset{\circ}{S}_1 \\ \overset{\circ}{S}_2 \\ \vdots \\ \overset{\circ}{S}_7 \end{bmatrix} \quad (9)$$

The left hand side of equation (9) represents a column vector of additive corrections to the initial set of parameter estimates. This column vector is found by calculating the elements of the 7x7 information matrix using equation (6), evaluated at initial parameter estimate values, then obtaining the inverse of the 7x7 matrix and

multiplying the inverted matrix by a column vector of efficient scores calculated from equation (2).

With the initial set of parameter estimates and the derived vector of additive corrections, a new set of parameter estimates closer to the m.l.e. can be found by adding the initial set and the vector of corrections. Let the new set of estimates be $\hat{\theta}^1$, then from (9)

$$\begin{bmatrix} \hat{\theta}_1^1 \\ \hat{\theta}_2^1 \\ \vdots \\ \hat{\theta}_7^1 \end{bmatrix} = \begin{bmatrix} \hat{\theta}_1^0 \\ \hat{\theta}_2^0 \\ \vdots \\ \hat{\theta}_7^0 \end{bmatrix} + \begin{bmatrix} \delta\hat{\theta}_1 \\ \delta\hat{\theta}_2 \\ \vdots \\ \delta\hat{\theta}_7 \end{bmatrix} \quad (10)$$

Repeating the process with the new approximations, i.e. substituting $\hat{\theta}_j^1$ values into equation (3), (6) and (9), a new set of additive corrections can be obtained. The process is repeated (Rao, 1968, p. 305 and Owston, 1978, p. 70) until the newly derived additive corrections are sufficiently small, for example, less than 0.0001. The final set of parameter estimates, with stable values, are good approximations of the maximum likelihood estimates.

After the m.l.e. are found, their variances (Elandt-Johnson, 1971, p. 306) can also be obtained. The variance of the j^{th} parameter estimate, $\text{var}\hat{\theta}_j$, is approximately equal to the value of the j^{th} diagonal element in the inverted information matrix with elements calculated at values of the

maximum likelihood estimates. Since the parameter estimate of interest in the MAX technique is $\hat{P}_{II} = \hat{\theta}_3$, its variance, $\text{var}\hat{\theta}_3$, is approximated by the (3,3)th element of the 7x7 inverted information matrix just mentioned.

The MAX technique is based on the fact that maximum likelihood estimates are approximately normally distributed (Owston, 1978, p. 78). A confidence interval can be built around $\hat{\theta}_3$ at a level of confidence of $(1-\alpha)$, such that,

$$\hat{\theta}_3 - z_{\alpha/2} \sqrt{\text{var}\hat{\theta}_3} < \theta_3 < \hat{\theta}_3 + z_{\alpha/2} \sqrt{\text{var}\hat{\theta}_3} \quad (11)$$

$$\text{or } \hat{P}_{II} - z_{\alpha/2} \sqrt{\text{var}\hat{P}_{II}} < P_{II} < \hat{P}_{II} + z_{\alpha/2} \sqrt{\text{var}\hat{P}_{II}} \quad (12)$$

where $z_{\alpha/2}$ is the standard normal deviate (Miller & Freund, 1965, p. 146-148) with the fraction of the area under the normal curve and to the right being equal to $\alpha/2$.

To illustrate the process of building a confidence interval around $\hat{P}_{II}$ and the subsequent decision on the validity of the hierarchy, part of the example given in Owston's thesis is used here. In Table II, the initial set of parameters is identified by the iteration number 0. This initial set is obtained from the 3x3 table of observed frequencies by using the marginal frequency method (devised by White & Clark; see Appendix B). The transition from one iteration to another is attained by the method of scoring.

Table II
Parameter Estimation Using
the Method of Scoring *

Iteration Number	$\hat{P}_O$	$\hat{P}_I$	$\hat{P}_{II}$	$\hat{P}_B$	$\hat{\theta}_a$	$\hat{\theta}_b$	$\hat{\theta}_c$	$\hat{\theta}_d$
0	.0886	.3318	.0000	.5796	.9731	.0000	1.0000	.1929
1	.0959	.2351	.0055	.6635	.9799	.0694	.9389	.1085
2	.1006	.2351	.0068	.6584	.9829	.0910	.9432	.0988
3	.0996	.2308	.0070	.6624	.9824	.0866	.9406	.0938
4	.1001	.2316	.0070	.6611	.9827	.0886	.9414	.0954
5	.0996	.2312	.0070	.6617	.9826	.0878	.9410	.0947
6	.1000	.2314	.0070	.6615	.9827	.0881	.9412	.0949
7	.1000	.2314	.0070	.6616	.9827	.0880	.9411	.0949
8	.1000	.2314	.0070	.6615	.9827	.0881	.9412	.0949
9	.1000	.2314	.0070	.6615	.9827	.0881	.9412	.0949
10	.1000	.2314	.0070	.6615	.9827	.0881	.9412	.0949

* Reproduced from Owston, p. 77.

After the 8th iteration, the new additive corrections are so small (less than .0001 here), that the 8th set can be assumed to be the set of maximum likelihood estimates. $\hat{P}_{II}$ has the value of .0070, and from the inverse of the information matrix calculated from the m.l.e. (maximum likelihood estimates), the estimated variance of $\hat{P}_{II}$ has a value of 7.451×10^{-5} .

For a 95% confidence interval around $\hat{P}_{II}$, $z_{\alpha/2}$ is about 1.96. Substituting $z_{\alpha/2}$, $\hat{P}_{II}$ and var $\hat{P}_{II}$ values into equation (12), the following interval is obtained,

$$-0.00944 < P_{II} < 0.02394 .$$

That is, the population value of P_{II} is estimated to lie between (-0.00944) and (0.02394). Since zero is included in this interval, it is concluded that, the hierarchy is valid, with skill I being a prerequisite to skill II.

1.4 OWSTON'S DESIGN AND FINDINGS ON MAX

The theoretical basis of the MAX technique has been discussed. An example was presented to illustrate how the MAX technique functions.

In this section, Owston's design to test the MAX technique is described. Also, his findings are discussed.

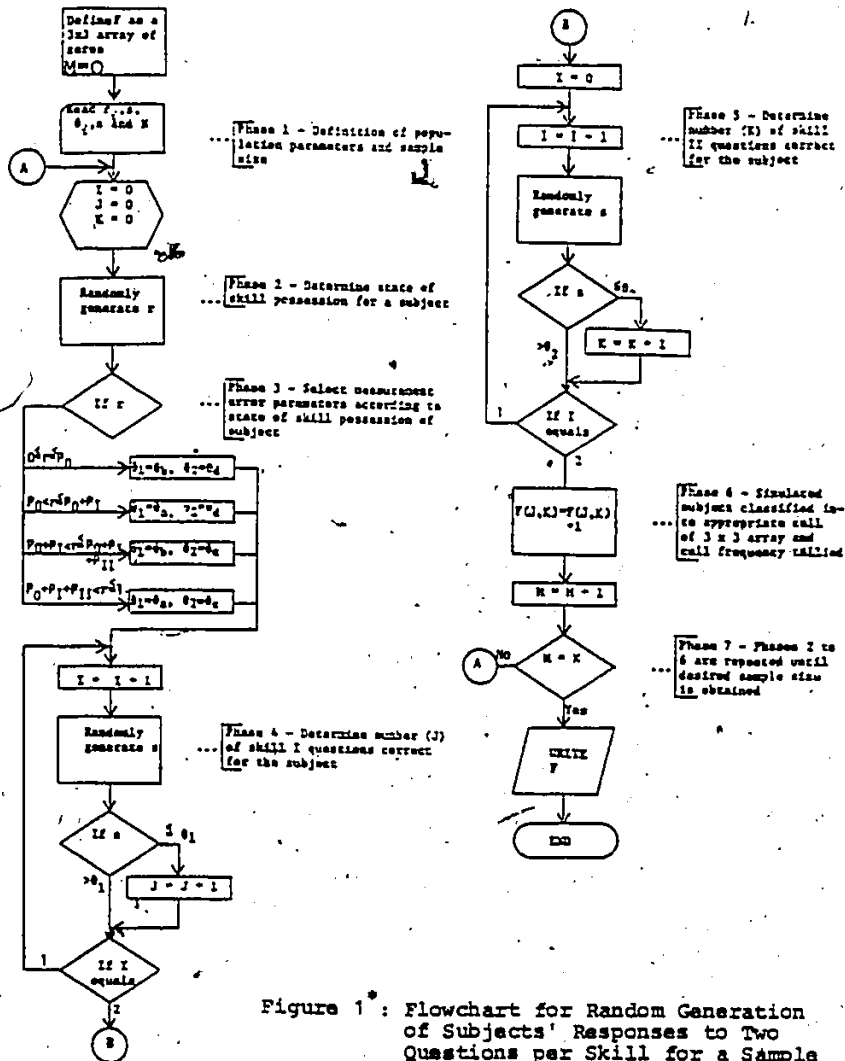
1.4.1 Generating random samples

In real life one never knows whether a learning hierarchy is valid or invalid. This is a complication in testing the performance of a validating technique. In view of this, the testing problem can be approached via computer simulation.

Figure 1 shows the computer flowchart, used by Owston, for random generation of simulated subjects' responses to the test items related to the two-skill hierarchy described in the White & Clark Model.

The most important symbols, used in the flowchart, are those related to the White & Clark Model. Recall that, according to this model, four proportions exist in a population of subjects. The four proportions, P_0 , P_I , P_{II} and P_B , correspond to four groups of subjects who possess "neither skill I nor skill II", only skill I, only skill II, and both skill I and skill II, respectively. Also recall that in this model 2 test items (or questions) per skill are used to test skill possession. Measurement error of correct guessing on skill I and skill II test items are indicated by θ_b and θ_d respectively. The measurement errors of blundering on skill I and skill II test items are denoted by $(1-\theta_a)$ and $(1-\theta_c)$ respectively. Note that θ_a is the probability of correctly answering skill I test items, for

Figure 1



* Reproduced from Owston, p. 87.

those subjects possessing skill I. θ_c is the probability of correctly answering skill II test items, for those subjects possessing skill II.

Notice that there are seven phases indicated in the flowchart. Beginning with phase 1, the four P's and four θ 's, just described, are assigned numerical values. Owston specified different sets of P's to represent valid and invalid hierarchies at different levels of skill difficulty for skill I and skill II. He also specified different sets of θ 's to represent various levels of measurement errors. (The details of specification of P's and θ 's will be discussed later in this section). Now, assume that four P and four θ values are specified to describe a population of "subjects" related to a hierarchy of certain level of skill difficulty and of a certain level of measurement errors. Then, a simulated sample of size N can be generated via the algorithm schematically represented in phase 2 to 7 of the flowchart.

Basically, it is through computer generation of a random number that Owston simulated a "subject" and assigned the "subject" to one of the four groups described in the White & Clark Model. Then, according to which group the "subject" was being assigned, the probabilities of correctly answering skill I and skill II test items, respectively were

determined. Furthermore, the "actual" responses of the "subject" to the test items were also simulated via a random number generation algorithm represented in phases 4 and 5.

The simulated response pattern of the first "subject" was then recorded in a 3x3 table. The process, in phase 2 to phase 6, was repeated N-1 times. By the end of this iterative process, the 3x3 table would contain the response patterns, of the N "subjects" in the sample, to the two test items per skill.

Then, Owston applied the MAX technique to the 3x3 pattern of simulated responses. The result from applying MAX would be a validity indication. Whether the indication is correct or incorrect depends on the validity of the hierarchy characterized by the "population of subjects".

1.4.2 Estimating risk of decision errors

Actually, using a set of population parameters (P's and θ 's) a large number of samples, of differing sizes, can be generated. Owston selected four levels of sample sizes to be 25, 40, 100 and 400. At each level of sample size, 1000 samples were generated. The MAX technique was applied, to make validity decisions, on each sample generated. Two levels of decision criteria, the 95% confidence interval and the 90% confidence interval were used. The correct

decisions, at each decision criterion, were tallied. So, for a given set of population parameters, at certain level of sample size, the results, of generating 1000 samples and applying the MAX technique 1000 times, would be two tallied totals. The two tallied totals correspond to the number of correct decisions made at 95% confidence interval and 90% confidence interval, respectively, on the 1000 samples generated. Owston determined the total number of incorrect decisions, at each decision criterion, simply by subtracting each tallied total from 1000.

Recall that there are two types of decision errors that can be made. One type is related to incorrectly indicating valid hierarchies as "invalid", the "I-V" type. Another type is related to indicating invalid hierarchies as "valid", the "V-I" type. Owston estimated the risk of committing decision errors by expressing the total number of incorrect decisions as a fraction of 1000.

Hence, for a valid hierarchy, the risk of committing decision errors (due to the application of MAX) would be

$$P(I-V) = \frac{\text{total no. of incorrect decisions when hierarchy is valid}}{1000}$$

For an invalid hierarchy, the corresponding risk of committing decision errors would be

$$P(V-I) = \frac{\text{total no. of incorrect decision when hierarchy is invalid}}{1000}$$

1.4.3 Specifying population parameters

So far, in this section, discussions have been centered around one set of population parameters that characterizes a learning hierarchy with a certain level of skill difficulty and with a certain level of measurement errors. From this set of parameters, 1000 samples at certain sample sizes were hypothetically generated. A procedure for estimating the risk of committing decision errors was also outlined. However, the details of specification of different sets of parameters were not discussed. In the next few paragraphs, Owston's procedure for specifying 48 sets of population parameters is presented.

In Table III a schematical representation of Owston's design for testing the MAX technique is shown. Notice that the 48 sets of population parameter values are formed by combining the different levels of validity, skill difficulty, and measurement errors.

Table III

Schematic Representation of Govton's design for testing the MAX technique, over different decision criteria, sample sizes, measurement errors, skill difficulty levels and validity levels.

48 sets of population parameter values were formed by the following combinations:	4 levels of skill difficulty (skill I, skill II) or (θ_1, θ_{II})		4 levels of measurement errors of guessing (θ_1, θ_2)	Based on each set of parameter values, random samples were generated. 1000 samples were generated at each of the following 4 levels of sample size:	For each sample generated, MAX technique was applied to construct a confidence interval (C.I.) around $\hat{\theta}_{II}$, then it was applied to make a validity decision. 2 levels of decision criteria were used for each sample. The 2 levels were:
	1 levels of validity				
Valid		(0.95, 0.90)	0.05	25	95% C.I.
Invalid ($\delta=0.00$)		(0.70, 0.50)	0.10	100	
Invalid ($\delta=0.25$)		(0.80, 0.25)	0.20	100	
Invalid ($\delta=0.25$)		(0.10, 0.05)	0.25	800	

First, Owston determined four levels of skill difficulty D_I for skill I as 0.95, 0.70, 0.40 and 0.10 . He argued that the performance of a validating technique might be related to the skill difficulty levels.

Then, he used his definition of a valid hierarchy to find the corresponding skill difficulty levels D_{II} for skill II as 0.90, 0.50, 0.25 and 0.05 . By Owston's definition of a valid hierarchy, skill II is more difficult than skill I, $P_{II}=0$, and the two skills are correlated with a phi correlation coefficient, ϕ , of approximately 0.68 . With this definition, once D_I is determined, D_{II} and the four P values can also be determined. (Appendix A shows the calculations involved in finding P values).

Corresponding to the four levels of skill difficulty pairs, (D_I, D_{II}) , four sets of P's were determined. These four sets of P's were related to valid learning hierarchies.

Owston determined four other sets of P's to characterize invalid hierarchies with skill I and skill II being independent. He did this by using the same four levels of skill difficulties, (D_I, D_{II}) , already determined, and by setting $\phi=0$.

Using the same four levels of skill difficulty, he also specified four sets of P's to characterize invalid

hierarchies with skill I and skill II being moderately correlated, $\phi = 0.25$. .

After specifying 12 sets of P values, Owston determined four sets of measurement error probability values. For each of these four sets, he assumed that the probabilities of measurement error of blundering were constant and equal to 0.05 . That is, $(1-\theta_a) = (1-\theta_c) = 0.05$; so, $\theta_a = 0.95$ and $\theta_c = 0.95$. He also assumed that the probabilities of guessing on skill I and skill II test items were equal. That is $\theta_b = \theta_d$. However, he allowed the probabilities of guessing to vary at four levels, with each level corresponding to one set of θ values. The four levels of guessing errors were 0.05, 0.10, 0.20 and 0.25 . In his discussions, .25 and .20 levels were related to the guessing probabilities of multiple choice types of items with four and five alternatives respectively. The two levels, .10 and .05, were related to guessing probabilities of short-answer types of items.

By combining 12 sets of P's with four sets of θ 's, Owston obtained 48 sets of population parameters for random generation of simulated subjects' responses.

1.4.4 Generating datasets to examine the risks of decision errors

Recall that four levels of sample sizes were used for

each set of population parameters. At each level of sample size, 1000 random samples were generated. The MAX technique was applied to make 1000 validity decisions with 95% confidence intervals and simultaneously 1000 validity decisions with 90% confidence intervals.

Also recall the outlined procedure for estimating the risks, $P(I-V)$ and $P(V-I)$, of committing decision errors.

By combining 4 levels of sample sizes with the 48 sets of population parameters, 192 datasets were obtained for estimating the risks of decision errors under different conditions. Of these 192 datasets, 64 datasets were used to estimate $P(I-V)$, 64 to estimate $P(V-I)$ when $\phi=0$, another 64 to estimate $P(V-I)$ when $\phi=0.25$.

1.4.5 Setting criteria on risks of committing decision errors

Of the two types of risks, $P(I-V)$ and $P(V-I)$, Owston considered $P(V-I)$ as more serious because it was related to indicating an invalid hierarchy as "valid". So, he set up a more stringent criterion on $P(V-I)$. Over the 128 datasets involving invalid hierarchies, Owston set a criterion of $P(V-I) \leq 0.05$; over the 64 datasets involving valid hierarchies, he set a criterion of $P(I-V) \leq 0.20$. Then, using 95% and 90% confidence intervals, Owston determined how often these criteria were met simultaneously on

populations that differed only in whether the hierarchy was valid or invalid.

In the next subsection, Owston's findings on the risks of decision errors, over the various pairs of datasets, will be presented.

1.4.6 Owston's findings on the risks of committing decision errors

Owston found that the risks of committing decision errors were related to the validity of hierarchy, decision criteria, sample sizes, probabilities of measurement errors, and the skill difficulty levels.

Recall that 95% and 90% confidence intervals were used as decision criteria. P(I-V) risks were found to be lower with 95% confidence interval than with 90% confidence interval. However, P(V-I) risks were higher with 95% confidence interval than with 90% confidence interval.

To compare the relative usefulness of 95% and 90% confidence intervals, Owston utilized the results obtained on the 64 datasets related to valid hierarchies of different levels of skill difficulty, and he also utilized the corresponding results on the 64 datasets related to invalid hierarchies with $\phi=0.25$ (He did not use the results on the 64 datasets related to $\phi=0.00$, because he argued that

datasets with $\phi=0.25$ would provide more stringent comparisons).

With valid hierarchies, the $P(I-V)$ values for 95% and 90% confidence intervals were less than the 0.20 criterion over all the 64 datasets. However, with invalid hierarchies ($\phi=0.25$) the $P(V-I)$ criterion of less than 0.05 was reached in 12 datasets using 95% confidence interval, and in 14 datasets using 90% confidence interval. In other words, the 90% confidence interval was found to be relatively more useful than the 90% confidence interval. Based on this finding, Owston suggested the use of shorter confidence intervals in future research.

Besides validity and decision criteria, sample sizes also played an important role in affecting the size of the error risks. In general, with all other things being equal, an increase in sample size corresponded to a decrease in the risk of committing decision errors.

Also, fewer decision errors were made with a decrease in the probability of measurement errors.

The only other variable is skill difficulty. With fixed levels on measurement errors, sample size and validity there was a general pattern on the size of risks over the four levels of skill difficulty. For both 95% and 90%

confidence intervals, the risks of $P(V-I)$ and $P(I-V)$ were generally higher for the two extreme difficulty levels of (.95, .90) and (.10, .05) than for the two intermediate levels of (.70, .50) and (.40, .25). Owston considered that this pattern might be due to the higher differences of $D_I - D_{II}$ at the two intermediate levels. For example, with the two intermediate levels, $D_I - D_{II}$ were .20 and .15 respectively. However, at the two extreme levels, the $D_I - D_{II}$ values were both .05.

1.5 WEAKNESSES OF THE MAX TECHNIQUE

Owston's study of the MAX technique appears to be theoretically and practically logical. However, in a pilot study conducted to test the MAX technique, using Owston's computer program, various types of errors were identified. All these types of errors occurred when using the subroutine program in which Owston implemented the MAX technique.

In Table IV, 11 types of identified errors are shown.

These error types were grouped into three error categories, I, II, and III. Category I error types, if not detected and by-passed, would cause computer interruptions to occur. It is assumed that Owston had somehow taken care of these errors. Error types in category II and category III would not cause any interruption to the running of the computer.

Table IV: Error categories and error types identified when the
 All technique was pilot tested using Guston's computer program

Error Category	Category Description	Error type		Error Type Description
		No.	Label	
I	Error types in this category cause computer program interruptions. It is essential to check these types of errors.	1	"f=0"	Marginal frequency $f=0$. Leads to "division by zero".
		2	"p=0"	probability values $P_1, P_2, \dots, P_9$ shown in Table I, approach zero.
		3	"Extremes"	Intermediate estimates take on extreme values.
		4	"SQRT5"	At the end of 5 iterations, variance of $\hat{P}_{II}$ may be a negative value. This causes computation of square root of a negative value.
		5	"SQRT2"	It is the same as "SQRT5" except that convergence is checked after each iteration instead of at the end of the 5th iteration.
II	Error types in this category are related to non-convergence of parameter estimates.	1	"NCONV5"	"Non convergence" of estimates is detected by an "inadequate" check at the end of 5 iterations.
		2	"PSCONV5"	"Inadequate" check indicates "convergence" but "proper" check indicates non-convergence.
		3	"NCONV10"	Non-convergence of estimates at the end of 10 iterations.
		4	"INTVAL"	Intermediate estimates take on unreasonable values.
III	This category of errors are associated with improperly convergent estimates	1	"-P _{II} "	Convergent $\hat{P}_{II}$ value happens to be a large negative value.
		2	"PIIFIN"	The final convergent estimates, other than $\hat{P}_{II}$, are outside the normal (0,1) range.

"inadequate" check for convergence is $\|\hat{H}\| \leq .001$, where $\hat{H}$ is the vector of corrections.

"proper" check is $\|\hat{H}\| \leq .001$

program. It is very likely that these error types had not received Owston's attention because in his computer program, no error checks were available to monitor these errors.

In error category I, the error type " $f=0$ " would occur if a randomly generated response pattern had a form similar to that shown in Figure 2, where the marginal frequency f is equal to zero. The reason for the occurrence of this error type in the MAX technique is that Owston used the marginal frequency approach devised by White & Clark (see Appendix B for details) to obtain an initial set of parameter estimates. With this approach the value of f is used as denominator in some computational formulae. When f is equal to zero, computer interruption would take place. The identification of this error type leads the researcher to consider that one weakness of the MAX technique is that it may not always be possible to obtain a feasible set of initial parameter estimates from the 3×3 patterns of observed responses.

This error type is expected to occur more often over datasets associated with small sample sizes and easy skills. Under the above conditions, the cells corresponding to either zero skill I or zero skill II questions correctly answered would more likely be empty. Hence it is also more likely that the marginal frequency value f is equal to zero.

		No. of skill II questions correctly answered			
		0	1	2	
No. of skill I questions correctly answered	2	0	3	15	a
	1	0	3	2	b
	0	0	1	1	c
		f	e	d	

Figure 2: A 3x3 pattern in which the error type "f=0" would occur

The error type "p=0" occurred when the probability function(s) (such as $p_1, p_2, \dots, p_9$), shown in Table I, had value(s) very close to zero. Computer interruption may take place due to "memory overflow" as a result of dividing some expressions by these function values. This error type is expected to occur more often in datasets with low measurement errors. Assume that the measurement errors were so low that there were no measurement errors at all (i.e. $\theta_b = \theta_d = 0$, and $\theta_a = \theta_c = 1$); then, in the White & Clark Model (shown in Table I) only the four corner cells (cells 1, 3, 7, 9) may contain non-zero probability functions. The other five probability functions (p_2, p_4, p_5, p_6, p_8) would be

zero. Furthermore, of the four possible non-zero probability functions, (p_1 , p_3 , p_7 and p_9), p_9 would be zero over the datasets associated with a valid learning hierarchy because the population value of P_{II} is then equal to zero; however, for the datasets associated with invalid hierarchies, p_9 is different from zero.

The error type "extremes" causes computer "memory overflow" due to abnormally large estimates occasionally obtained during the parameter estimation process of the MAX technique. Since the iterative parameter process involves complicated computations, such as the inversion of a 7×7 matrix at each iteration, this error type could be due to some computer rounding errors in the lengthy and complicated computations.

The two error types "SQRT5" and "SQRTP2" are both associated with computing the square root of the variance of $\hat{P}_{II}$ when this variance happened to be a negative value. For "SQRT5", convergence of estimates was checked at the end of the 5th iteration of the parameter estimation process; if convergence was detected, then the square root of the variance of $\hat{P}_{II}$ would be computed. "SQRTP2" is the same as "SQRT5" except that convergence of estimates was checked after each iteration rather than at the end of the 5th iteration.

The error types in category II are related to the problem of non-convergence of parameter estimates. These errors seemed to occur more often than the errors in category I and III.

Owston implemented a one-sided check for convergence of parameter estimates at the end of the 5th iteration of the parameter estimation process. This check, although inadequate, might still be used in detecting a large portion of the non-convergence cases. However, no non-convergence cases were reported by Owston. One possible explanation for this is that a simple but serious program logic error was committed by Owston (see Appendix E for details). The logical error, in effect, might lead Owston to consider a "non-convergent" case as "convergent".

By implementing a proper two-sided (using absolute value) check for convergence immediately after Owston's one-sided check, it was found that there were indeed some non-convergent cases that could not be detected by the one-sided check.

Since only 5 iterations were used by Owston, non-convergence of estimates appeared very likely. When ten iterations were used, there appeared to be a sizeable decrease in non-convergence cases; however, in some samples, the estimates never converged. It was identified that, in

these samples, the intermediate estimates usually had unreasonable values. Supposedly, the estimates should have values in the (0,1) range. Some of the intermediate estimates had values outside the (-1,2) range.

The error types in category III are associated with improperly converging estimates. It was found that some of the converging P_{II} estimates had negative and relatively large values (e.g. less than -.05). It was also identified that other estimates (there were six other estimates besides $\hat{P}_{II}$) might also be improperly converging. Since the MAX technique is mainly based on the P_{II} estimate, it appeared reasonable to emphasize only the proper convergence of $\hat{P}_{II}$; however, it may also be reasonable to consider the proper convergence of the other estimates at the same time because the estimates are interrelated in the process of estimation.

1.6 THE RESEARCH PROBLEM

According to Owston, the MAX technique appears to be a potentially very useful device for validating learning hierarchies. However, different types of errors existed in the computer program used by Owston to implement the MAX technique. This means that Owston's findings on MAX might be biased due to his failure to account for these errors.

It is therefore appropriate to raise the following questions about the MAX technique:

1. If all the errors were accounted for, would the MAX technique be useful for validating learning hierarchies?
2. If Owston's findings were seriously biased, is it possible to improve the procedure for validity decision making using the MAX technique?

Chapter II
DESIGN AND EXPERIMENTATION

At the end of the previous chapter, two questions were raised. The first question is related to the extent of error bias in Owston's data when using his computer program to implement the MAX technique. The second question is associated with the possible improvements on the implementation procedure of the MAX technique.

Corresponding to the two questions, two procedures, A and B, were developed. These two procedures are modified versions of Owston's procedure for implementing the MAX technique.

In this chapter, there are three major sections. The first section contains a description of procedures A and B. In the second section the design for collecting data by means of these two procedures is discussed. Suitable criteria to be used for analysing the data to be collected are reported in the final section.

2.1 TWO PROCEDURES FOR IMPLEMENTING THE MAX TECHNIQUE

Various types of error exist in Owston's computer program. Procedures A and B contain error checks that can be used to count the frequencies of occurrence of the different types of errors. There are similarities and differences between procedure A and procedure B. First, similarities are described.

Both procedures, A and B, are modified versions of Owston's procedure for implementing the MAX technique. These two procedures contain essential error checks that are mainly used to avoid computer program interruptions. Table V shows the various types of errors identified, together with the indications of in which procedure(s) each error type is being checked. For example, in error category I, errors related to " $f=0$ ", " $p=0$ " and "extremes" cases are checked in both procedures. In fact, procedures A and B are designed in such a parallel manner that the three types of errors, just mentioned, are monitored by error checks which are common to the two procedures. Note that the two other error-types, in category I, are basically the same. They are related to the problem of computing the square root of a negative number. Two error checks are used, one for procedure A and the other one for procedure B, to check for "negative-square-root" cases.

Table VI: Error types that are checked for by using
Procedure 1, Procedure 2, or both procedures 1 and 2

Error Category	Category Description	Error type		Procedure(s) used to check the error type.	Error Type Description
		No.	Label		
I	Error types in this category cause computer program interruptions. It is essential to check these types of errors.	1	"f=0"	A, B	Marginal frequency $f=0$. Leads to "division by zero".
		2	"p=0"	A, B	probability values $P_1, P_2, \dots, P_N$, shown in Table I, approach zero.
		3	"Extremes"	A, B	Intermediate estimates take on extreme values.
		4	"SORT5"	A	At the end of 5 iterations, variance of $\hat{P}_{II}$ may be a negative value. This causes computation of square root of a negative value.
		5	"SORTP2"	B	It is the same as "SORT5" except that convergence is checked after each iteration instead of at the end of the 5th iteration.
II	Error types in this category are related to non-convergence of parameter estimates.	1	"SCORV5"	A	"Non convergence" of estimates is detected by an "inadequate" check at the end of 5 iterations.
		2	"PSCORV5"	A	"Inadequate" check indicates "convergence" but "proper" check indicates non-convergence.
		3	"SCORV10"	B	Non-convergence of estimates at the end of 10 iterations.
		4	"INTVAL"	B	Intermediate estimates take on unreasonable values.
III	This category of errors are associated with improperly converged estimates	1	"-P _{II} "	B	Converged $\hat{P}_{II}$ value happens to be a large negative value.
		2	"FINRAB"	B	The final converged estimates, other than $\hat{P}_{II}$, are outside the normal (0,1) range.

* "inadequate" check for convergence is $|\hat{W}| < .001$, where $\hat{W}$ is the vector of additive scores.
 ~ "proper" check is $|\hat{W}| < .001$

Since procedure A is used to collect data for determining the extent of bias in Owston's data, procedure A is set up to be as similar as possible to Owston's procedure. In procedure A, 5 iterations are used for the parameter estimation process. A check for convergence of estimates is performed at the end of the 5th iteration. In contrast, procedure B is used to provide possible improvements. In procedure B, a maximum of 10 iterations is allowed for the process of parameter estimation. Convergence check is made after each iteration so as to detect convergence of estimates which may occur between the first and the tenth iteration.

Another difference between the two procedures is the nature of the checks for convergence. Owston used an inadequate(one-sided; see Table V) check for convergence. This inadequate check is included in procedure A; it is followed by a proper(two-sided) check. The latter check detects cases that are non-convergent but that would be indicated as "convergent" by the inadequate check. In procedure B, only one proper check for convergence is utilized.

For the last error-type involving non-convergence of estimates, a conservative range check of $(-1,2)$ is used to detect unreasonable intermediate parameter estimates outside this range. This range check is only used in procedure B.

Two other error-checks unique to procedure B, are utilized to detect improperly converging estimates. The first error-type in category III is related to a negative converging $\hat{P}_{II}$ value. A conservative check, $\hat{P}_{II} < -.05$, is used to determine if the value of $\hat{P}_{II}$ is negative and relatively large. The final error-check, unique to procedure B, is used to detect improperly converging estimates other than $\hat{P}_{II}$. A liberal range of $(-.05, 1.05)$ is used to check for estimates (besides $\hat{P}_{II}$) having values outside this range.

Thus far, the similarities and the major differences between the two procedures, A and B, have been described. At the same time, all the important error checks used in these two procedures have been discussed. These error checks were implemented in a modified version of Owston's computer program. A computer flowchart showing the various error checks used in procedures A and B, appears in Appendix C.

Before proceeding on to the next section, it may be appropriate to mention another difference between procedures A and B. In procedure A, two levels are used in deriving confidence intervals around $\hat{P}_{II}$. The two levels are the 95% and 90% confidence intervals. In procedure B, besides these two levels, four other levels are used. They are the

85%, 80%, 75% and 70% confidence intervals. These additional levels are included because of Owston's suggestion that the MAX technique may perform better with shorter confidence intervals.

2.2 DESIGN FOR COLLECTING DATA

With minor modifications, Owston's design for studying the MAX technique was used. Modifications are required to accommodate the two procedures, A and B. At the same time, through modifications, the scope of the current study is reduced to a manageable size.

The random sample generation procedure for the current study is the same as Owston's. Owston used four levels of sample sizes of 25, 40, 100 and 400. The sample size of 25 appears to be inappropriate for the MAX technique because of the theoretical assumption of large sample size in the parameter estimation process. Also, it is suggested that the various types of identified errors occur frequently at sample sizes of 25. Consequently only three levels of sample sizes were used in this study. These levels are 40, 100 and 400.

In specifying population parameters, Owston used three levels of validity. The levels are valid, invalid with $\phi=0$, and invalid with $\phi=.25$. Owston later argued that two

levels would be sufficient. The level, invalid with $\phi=0$, is somewhat redundant. So, in this study, two levels of validity were adopted. The level invalid with $\phi=0$ was eliminated.

Besides validity of the hierarchy, one other factor, used by Owston in parameter specification, is skill difficulty. From Owston's four levels of skill difficulty pairs, (D_I, D_{II}) , where D_I and D_{II} are skill difficulties of skill I and skill II respectively, three levels were chosen. These three levels are $(.95, .90)$, $(.70, .50)$ and $(.40, .25)$. The level not adopted was related to extremely difficulty skills with (D_I, D_{II}) being $(.10, .05)$; the reason for not adopting this level is that it might not be frequently encountered in practice.

In the White & Clark Model of learning hierarchies, there are four population parameters related to four mutually exclusive groups of subjects with different status of skill possession. These four parameters P_O , P_I , P_{II} and P_B correspond to the four proportions of subjects who possess "neither skill", only skill I, only skill II, and "both skill I and skill II" respectively.

When three levels of skill difficulty, (D_I, D_{II}) , are combined with two levels of validity, six sets of (P_O, P_I, P_{II}, P_B) values can be determined by following the procedure

used by Owston. The six sets of P values are shown in Table VI.

Besides the four P parameters, there are four other parameters, in the White & Clark Model, related to the probabilities of correctly answering skill I and skill II test items. Owston allowed the probabilities of correctly answering the items through guessing to vary at four levels. These four levels were .25, .20, .10 and .05. In Owston's discussions, .25 and .20 levels would be related to the guessing probabilities of multiple choice types of items with four and five alternatives respectively. The two levels, .10 and .05, would be related to guessing probabilities of short-answer types of items. Three of the four levels of probabilities of guessing were used in this study. These three levels are .25, .20, .05.

By combining six sets of P values and 3 levels of probabilities of guessing, 18 sets of population parameters were formed. Using each of the 18 sets of population parameters random samples of three different sample sizes were generated. The three sample sizes were 40, 100 and 400. At each level of sample size, for each set of population parameters, 1000 random samples were generated. Each "1000 samples" formed a dataset for testing the MAX technique via procedures A and B.

Table VI: Six sets of P's corresponding to 3 levels of skill ,
difficulty and 2 levels of validity of learning hierarchy

(Diff. skill I, Diff. skill II) (D _I , D _{II})	Valid	Invalid
(.95, .90)	(P ₀ , P _I , P _{II} , P _B) (0.05, 0.05, 0.00, 0.90)	(P ₀ , P _I , P _{II} , P _B) (0.02, 0.08, 0.03, 0.87)
(.70, .50)	(0.30, 0.20, 0.00, 0.50)	(0.21, 0.29, 0.09, 0.41)
(.40, .25)	(0.60, 0.15, 0.00, 0.25)	(0.50, 0.25, 0.10, 0.15)

*Each set of P's contains four values representing four mutually exclusive proportions of subjects with different status of skill possession.

By combining 18 sets of parameters with 3 sample sizes, 54 sets of 1000 random samples were drawn, 27 of these datasets were related to valid hierarchies at 3 different levels of skill difficulty, probability of guessing and sample size. The other 27 datasets were related to invalid hierarchies with the corresponding levels of skill difficulty, probability of guessing, and sample size.

In Owston's study, for each dataset, his procedure was applied to make 1000 validity decisions using 95% confidence intervals and to make another 1000 validity decisions using 90% confidence intervals. However, due to the existence of the various types of errors, some modifications were made in this part of his design. First, assume that in each dataset x samples are detected by procedure A to be associated with the different error-types; correspondingly, y samples are detected by procedure B. Note that x and y are variables taking on different values over the different datasets. Then, the number of samples, in each dataset, that are used to make decisions in procedure A, would be "1000- x ". Correspondingly, the number of samples, used to make decisions, in procedure B, would be "1000- y ".

For procedure A, two confidence intervals (C.I.), 95% and 90%, were used. "1000- x " decisions were therefore made at 95% confidence interval, and another "1000- x " decisions at 90% confidence intervals.

For procedure B, six levels of confidence intervals, 95%, 90%, 85%, 80%, 75% and 70%, were used. "1000-y" validity decisions were made at each of the six levels of confidence intervals.

As a summary of this section, a schematical representation of the design for the current study is shown in Table VII. The technical details of this design can be found in Appendix D, where the related computer program statements are listed.

2.3 CRITERIA FOR ANALYSING THE DATA

Validity decisions may be correct or incorrect. However, before applying procedures A and B to make validity decisions, the presence of errors was checked.

Samples, detected in procedure A as containing error, were not used for making validity decisions. Then, procedure A was, in effect, applied to the remaining samples to make decisions at two levels of confidence intervals. At each level of confidence interval, the results were two proportions(out of 1000), one proportion being related to correct decisions, and another proportion being related to incorrect decisions. Hence, in a dataset of 1000 random samples to test the MAX technique for procedure A, there was a constant proportion, $P(ND)$, of samples related to "no-

Table VII
Schematic representation of the design
for applying procedures A and B to test the HIT
technique over different decision criteria, sample
sizes, measurement errors, skill difficulty levels
and validity levels

16 sets of population parameters are formed by the following combinations:			Based on each set or parameter values, random samples are generated. 1000 samples are generated at each of the following 3 levels of sample size:	Two procedures, A and B, are applied to the 1000 samples generated, one sample at a time. A and B have more different error checks, so the two procedures may have different error counts. Furthermore A and B do not have same levels of decision criteria. The set-ups for A and B are as follows:			
2 levels of validity	3 levels of skill difficulty (skill I, skill II) OR. (D ₁ , D ₂)	3 levels of measurement errors of guessing		procedure	# of samples with error(s)	# of samples used for decision making	levels of decision criteria used on each sample involved in decision making
valid	(0.95, 0.90)	0.05	40	A	x	1000-x	95% C.I.
							90% C.I.
invalid	(0.70, 0.50)	0.20	100	B	y	1000-y	95% C.I.
							90% C.I.
invalid ($\rho=0.25$)	(0.40, 0.25)	0.25	400				95% C.I.
							90% C.I.

* C.I. indicates confidence interval

decision" because of errors. The remaining proportion, $1-P(\text{ND})$, was divided into a proportion of correct decisions, $P(\text{COR})$, and a proportion of incorrect decisions, $P(\text{INC})$. $P(\text{COR})$ and $P(\text{INC})$ could vary from the 95% C.I. to the 90% C.I. level, but the sum of $P(\text{COR})$ and $P(\text{INC})$ was a constant for that particular dataset.

Similarly, procedure B was applied in a like fashion, except that with procedure B, four additional levels of decision criteria were used in constructing confidence intervals.

Owston used the proportion of incorrect decisions as an estimate of the risk of committing decision errors. In view of the additional proportion of "no-decision" cases, it appeared to be more useful to adopt the proportion of correct decisions as a measure of the performance of the MAX technique.

Corresponding to Owston's criteria of requiring the proportion of incorrect decisions to be less than or equal to .20 for valid hierarchies, and less than or equal to .05 for invalid hierarchies, proportion of correct decisions, $P(\text{COR})$ was considered as satisfactory if it was greater than .80 for valid hierarchies, and .95 for invalid hierarchies.

The (.80, .95) criteria, just mentioned, will be

utilized, in the next chapter, as a basis for comparing the current results with those found by Owston.

Chapter III

PRESENTATION AND DISCUSSIONS OF RESULTS

At the end of the first chapter, two research questions were raised. The first question is related to determining the extent to which the various types of identified errors affect the performance of the MAX technique for validating learning hierarchies. The second question is about the possibilities of improving the procedure for implementing the MAX technique.

Corresponding to these two questions, two procedures, A and B, were designed. Procedure A is identical to Owston's procedure for implementing the MAX technique, except that, in procedure A checks were provided to detect the various types of identified errors. Procedure B was designed to improve the performance of the MAX technique; additional error checks and modified decision criteria were utilized in procedure B.

Procedures A and B were applied to 27 datasets characterising valid learning hierarchies (for convenience, these datasets are called valid datasets) and to 27 corresponding datasets characterising invalid

hierarchies (for convenience, these datasets are called invalid datasets). Corresponding to each type of identified errors relative frequencies of occurrence were calculated when applying procedures A and B. Also, proportions of correct and of incorrect decisions at each level of decision criteria were calculated.

In this chapter, results obtained from using procedure A and B, over the 27 valid and the 27 invalid datasets, are presented and discussed. For the purpose of comparison, pertinent results from Owston's study are also included.

3.1 RESULTS OBTAINED FROM USING PROCEDURE A

The results, obtained from using procedure A, were divided into two parts. The first part is related to the various error types. The second part is related to the proportions of correct and incorrect validity decisions. Since procedure A was designed mainly for investigating the extent of influences due to the identified errors, the emphasis is placed on the first part of the results. In the next section, the proportions of "correct decisions" obtained by using procedure A are compared with the corresponding findings in Owston's study.

In Table VIII, the frequencies of occurrence of the various types of errors in each different dataset (1000

samples per dataset) are shown as proportions. For each dataset, the sum of the proportions due to the six types of errors ("f=0", $p=0$,...etc.) is denoted as the proportion of "no decisions". In general, the proportions of "no decisions" are fairly high; over the valid datasets, the range is from .043 to .948 and over the invalid datasets, the range is from .141 to .970. For each dataset, a high proportion of "no decisions" implies that there would be a small number of samples for which validity decisions (correct or incorrect decisions) could be made. Also, it implies that some proportions related to the various types (6 types) of identified errors are high. Recall that these six types of errors belong to two error categories. The first four of the six error types belong to error category I which is related to errors that may cause computer interruptions. The last two error types, "non-convergence" and "pseudo-convergence", belong to error category II which is related to the problems of non-convergence of parameter estimates.

It may be observed, in Table VIII, that the proportions related to the error types in category I are generally low as compared with the proportions related to the error types in category II. This implies that, in general, computer interruptions would not occur very frequently. Furthermore, it implies that, using procedure A, non-

51

convergence of parameter estimates occurs frequently. Since, in procedure A only five iterations were allowed for the parameter estimation process, it is not surprising to observe high proportions of category II error types.

It may be interesting to note, that for each error type the patterns of the proportions associated with the valid datasets may be very different from the corresponding patterns of the proportions associated with the invalid datasets. For example, for the error type related to very small probability values ($p=0$), the proportions in some valid datasets (e.g. dataset # 16,17,18,25,26,27) were in general much higher than in the invalid counterparts. The explanation of this phenomenon is postponed until each error type is studied in more detail. Another example can be used to illustrate a somewhat reversed pattern in which the proportions related to invalid datasets are higher than those from the valid counterparts. For the error type, "non-convergence" of estimates, it is generally true that the proportions in the invalid datasets are relatively higher than the corresponding proportions in the valid datasets. The contrast is especially clear in the same datasets (# 16,17,18,25,26,27) as when the error type " $p=0$ " was found to occur more often over the valid hierarchies.

In the subsequent paragraphs, each error type with the

two error categories will be studied more carefully. The incidence of each type of error will be compared for valid or invalid hierarchies. Frequent references will be made to Table VIII where the error proportions are shown. For the detailed descriptions of each error type, the reader is referred to sections 1.5 and 2.1 .

The error type "marginal frequency equals zero" may occur in the computation of initial parameter estimates. One of the marginal frequencies, f , used in calculating the initial estimates may happen to be zero. If that were the case, computer interruption would occur due to an undefined computational situation of dividing some values by zero. An error check was used to detect this error type and is used to by-pass computer interruptions.

As shown in Table VIII, this error type was detected in datasets associated only with easy skills learning hierarchies and small sample sizes (mainly sample sizes of 40). Probability of guessing appears to be related to the size of the proportions due to this type of error. As the probability of guessing decreases, the magnitudes of the error proportions also decrease. These patterns were expected to occur. Reasons for such expectations were given in the part of section 1.5 where this error type was discussed in detail. Also, it may be noticed that the

proportions associated with this error type are relatively small. Over the valid datasets, the range is from .000 to .081. Over the invalid datasets, the range is from .000 to .099.

The error type "probability values tend to be zero" may occur during the parameter estimation process when the probability functions in the White & Clark Model (shown in Table I) are used to compute the first approximation of the parameter estimates. If one or more of these probability functions approach zero, computer interruption may result from a near zero denominator in some expressions. In Table VIII, proportions due to this type of error are found to be generally higher in the valid datasets than in the invalid datasets. The range, over the valid datasets, is from .000 to .170. The range, over the invalid datasets, is from .000 to .047. It is particularly interesting to note that the major differences in proportions, between valid and invalid datasets, due to this error type, occurred in the datasets related to moderate and difficult skills hierarchies with a probability of guessing level of .05.

One possible explanation for the above phenomenon is that the probability functions (shown in Table I) are different over the valid and invalid datasets. For example, assume that there were no measurement errors (such that the

probability of guessing is zero, etc.), one of the (Probability functions(p_{02} ; shown in Table I) is in fact equal to zero in the valid datasets; however this probability function is different from zero in the invalid datasets.

The error type "extremely large estimates" may occur during the parameter estimation process. Due to some unknown reasons (most likely due to the complicated matrix inversion process involved in the MAX technique), the intermediate parameter estimates may(although rarely) be abnormally large (larger than ten to the power 6). An error check was used to detect this error type so as to avoid computer interruptions due to "computer memory overflow". As shown in Table VIII, the proportions associated with this error type are of small magnitude. Over the valid datasets, the range of proportions, due to this error type, is from .000 to .008 . Over the invalid datasets, the range is from .000 to .009 .

The error type "negative variance of $\hat{P}_{II}$ " could occur when the confidence interval around $\hat{P}_{II}$ is constructed. Basically, this error type could be due to rounding errors in the complicated computational processes involved in the MAX technique. However, it may be observed, in Table VIII, that this type of error affected the valid datasets slightly

more often than the invalid datasets. Over the valid datasets, the proportions over this error type range from .000 to .007; while over the invalid datasets, the range is from .000 to .003. Particularly interesting to note is that this error type was never detected in the invalid datasets associated with moderate and difficult skills hierarchies. The real reasons behind this phenomenon are uncertain. Additional data may be useful for further understanding. However, one possible explanation is that the variance estimate of P_{II} is directly related to the estimated value of P_{II} which may be negative due to the normal approximation of the distribution of the parameters around zero. For the valid datasets, the parameter P_{II} was set at zero; as for the invalid datasets, P_{II} was larger than zero. In the valid datasets, it is more likely to have negative estimates of P_{II} ; this may lead to the more frequent occurrence of negative variance estimate of P_{II} in the valid datasets.

So far, each of the error types, in error category I, which can be detected when using procedure A, has been discussed. Also mention was made of the "non-convergence" error type in error category II. The other error type, "pseudo convergence of estimate", in category II, may be studied together with the "non-convergence" type because these two error types are both related to the problem of

non-convergence of parameter estimate. The reason for designing a separate error check was to investigate the deficiencies of Owston's check (one-sided check on additive corrections; the proper check should be based on absolute values of the additive corrections for testing the convergence condition, e.g. " $| \text{additive corrections} | < .001$ ").

In general, the error types in category II occur more often over the invalid datasets. A possible explanation of this phenomenon is that the initial parameter estimate values of P_{II} (the proportion of subjects possessing skill II without possessing skill I) were set to be zero for both valid and invalid datasets. For the valid datasets, the population parameter values of P_{II} were in fact zero; however, for the invalid datasets, the corresponding values were greater than zero; for the easy, moderate and difficult skills hierarchies, the P_{II} values for invalid datasets were, respectively .03, .09 and .10. It would be reasonable to expect that, for the invalid datasets, convergence of the parameter estimate of P_{II} from zero to the population values (.03, .09 and .10) was more difficult than the convergence of P_{II} estimate value from zero to zero in the valid datasets.

Besides validity of the hierarchy, other factors such as sample sizes, probability of guessing levels and skill

difficulty levels also appeared to affect the error types in category II. For datasets based on large sample sizes of 400, the error proportions were generally lower than the corresponding proportions related to sample sizes of 100 and 40. The lower error proportions also tend to occur with the low measurement error for guessing. Furthermore, the error proportions, in the datasets associated with moderate skills learning hierarchies, were generally lower than those from the difficult and easy skill learning hierarchies. This phenomenon may be due to the fact that the difference, $D_I - D_{II}$ (where D_I, D_{II} are the difficulty levels of skill I and skill II respectively), is larger over the moderate skill learning hierarchies. For the moderate skills learning hierarchies, $D_I - D_{II}$ is 0.20; for the difficult skills hierarchies, $D_I - D_{II}$ is 0.15; while, for the easy skills, $D_I - D_{II}$ is 0.05.

So far, only the error proportions have been discussed. The rest of this section is used to present and discuss the proportions of correct and incorrect decisions.

As shown in Table IX, over the valid datasets, the proportions of correct decisions tend to be smaller at the narrower 90% confidence interval. However, over the invalid datasets, the proportions of correct decisions tend to be larger at the 90% confidence interval. According to the

Procedure A (Part II)

Probability of Difficulty Levels (p, D, II)	Sample Size N	Data for VMD Hierarchies				Data for IRMAID Hierarchies			
		95% Confidence Interval		90% Confidence Interval		95% Confidence Interval		90% Confidence Interval	
		Correct	Incorrect	Correct	Incorrect	Correct	Incorrect	Correct	Incorrect
Easy (.45-.90)	400	.794	.206	.000	.000	.205	.001	.732	.001
	25	.911	.089	.000	.000	.000	.000	.000	.000
	40	.948	.051	.001	.001	.001	.001	.001	.001
	400	.694	.306	.000	.000	.304	.002	.702	.021
Moderate (.70-.90)	400	.806	.194	.000	.000	.134	.000	.001	.000
	25	.918	.082	.002	.000	.000	.002	.000	.000
	40	.918	.082	.002	.000	.000	.002	.000	.000
	400	.215	.785	.000	.000	.782	.003	.558	.217
Difficult (.40-.25)	400	.605	.395	.000	.000	.395	.000	.816	.000
	25	.784	.215	.001	.001	.214	.002	.001	.000
	40	.784	.215	.001	.001	.214	.002	.001	.000
	400	.308	.692	.002	.002	.682	.010	.338	.309
Very Difficult (.40-.25)	400	.579	.421	.000	.000	.421	.000	.688	.003
	25	.731	.269	.000	.000	.269	.000	.000	.000
	40	.731	.269	.000	.000	.269	.000	.000	.000
	400	.043	.957	.004	.004	.944	.013	.279	.572
Very Very Difficult (.40-.25)	400	.351	.649	.000	.000	.647	.002	.616	.013
	25	.631	.369	.000	.000	.369	.000	.011	.001
	40	.631	.369	.000	.000	.369	.000	.011	.001
	400	.170	.830	.000	.000	.830	.000	.141	.859
Very Very Very Difficult (.40-.25)	400	.103	.897	.000	.000	.896	.001	.482	.237
	25	.317	.683	.000	.000	.683	.000	.497	.006
	40	.317	.683	.000	.000	.683	.000	.497	.006
	400	.002	.998	.003	.003	.190	.008	.580	.093
Very Very Very Very Difficult (.40-.25)	400	.006	.994	.000	.000	.194	.000	.820	.000
	25	.837	.163	.000	.000	.163	.000	.000	.000
	40	.837	.163	.000	.000	.163	.000	.000	.000
	400	.223	.777	.005	.005	.761	.016	.484	.282
Very Very Very Very Very Difficult (.40-.25)	400	.552	.448	.000	.000	.446	.002	.740	.004
	25	.765	.235	.000	.000	.234	.001	.000	.000
	40	.765	.235	.000	.000	.234	.001	.000	.000
	400	.135	.865	.000	.000	.864	.001	.239	.761
Very Very Very Very Very Very Difficult (.40-.25)	400	.137	.863	.000	.000	.863	.000	.564	.225
	25	.671	.329	.000	.000	.329	.000	.000	.000
	40	.671	.329	.000	.000	.329	.000	.000	.000
	400	.000	.000	.000	.000	.000	.000	.000	.000

Table II: Proportions of correct decisions and incorrect decisions at 95% and 90% confidence intervals, made by using Procedure A, over the various combinations of validity, difficulty, probability of guessing and sample size level.

Notes: The proportions of "No Decisions" are the same as the corresponding entries in Table VIII.

nature of the decision process involved in the MAX technique, the above phenomenon is logical. The narrower 90% confidence interval (as compared with the 95% confidence interval) is less likely to contain zero. Hence, the probability of indicating any hierarchy as a valid hierarchy decreases with the decrease in the width of the confidence interval. Consequently, for a valid dataset, the proportion of correct decisions tends to decrease with the decrease in the width of confidence interval. In contrast, for an invalid dataset, the proportion of correct decisions tends to increase with the decrease in the width of the confidence interval; because it is less likely to contain zero.

It may also be observed that the proportions of incorrect decisions are comparatively smaller over the valid datasets than over the invalid datasets. For example, at the 95% confidence interval, the range of incorrect decisions, over the valid datasets, is from .000 to .005. Over the invalid datasets the range is from .000 to .371.

This phenomenon might be related to the values of the P_{II} estimates upon which confidence intervals were built and validity decisions were made. Supposedly, the convergent P_{II} estimate values would be very close to the population values. For the valid datasets, the population values of P_{II} are all equal to zero. For the invalid datasets, the

population values of P_{II} depend on the skill difficulty levels . For the datasets associated with the easy, moderate and difficult skills learning hierarchies, the population P_{II} values are .03, .09 and .10 respectively. These values are distinctly different from zero.

For the invalid datasets, if the estimated values of P_{II} were very close to the corresponding population values, the confidence intervals built around P_{II} estimates should not likely contain zero. Hence, there should be little chance of incorrectly indicating an invalid hierarchy as valid. However, the high proportions of incorrect decisions, over the invalid datasets, indicate otherwise.

Recall that, in the iterative parameter estimation process used in the MAX technique, an initial set of parameter estimates is required. The White & Clark marginal frequency method for obtaining parameter estimates (from subjects' response patterns to two skill testing items per skill) was used. In this method the initial estimate values for P_{II} are set to zero for all the valid and invalid datasets. Also, recall that in discussing the problem of non-convergence of parameter estimates, it was mentioned that it might not be easy to obtain the convergence of $\hat{P}_{II}$ in the invalid datasets.

Another possibility for the occurrence of high proportions of incorrect decisions, over the invalid datasets, is that the variance estimates of P_{II} used to construct confidence intervals were so high that the chances of including zero in the confidence intervals were also high. This also might lead to the high proportions of incorrect decisions over the invalid datasets.

3.2 COMPARISONS OF OWSTON'S RESULTS WITH THE RESULTS OBTAINED BY USING PROCEDURE A

At the end of chapter II, criteria were established for comparing results. For the valid datasets, the proportion of "correct decisions" was required to be larger than .80 before the MAX technique was considered as satisfactory. For the invalid datasets, the proportion of "correct decisions" was required to be larger than .95. Also, these criteria were required to be met simultaneously in the valid datasets and in the corresponding (having the same level of skill difficulty, probability of guessing and sample size) invalid datasets.

Based on these criteria, (.80, .95), Owston's results (retrieved from Appendix of Owston's thesis), and the proportions of correct decisions obtained from using procedure A (retrieved from Table IX) were compared. The data used for comparisons are shown in Table X.

Skill Difficulty Levels (σ_1, σ_2)	Probability of Success (θ)	Sample Size (N)	Proportions of Correct Decisions									
			90%				90%				Confidence Interval	
			Dustin's Results		Adjusted Results from Procedure A		Dustin's Results		Adjusted Results from Procedure A		Results from Procedure A	
			Valid	Invalid	Valid	Invalid	Valid	Invalid	Valid	Invalid	Valid	Invalid
Easy (.75, .90)	.25	400	0.934	0.907	0.704	0.001	0.204	0.253	0.923	0.875	0.205	0.700
		100	0.920	0.893	0.089	0.000	0.097	0.223	0.918	0.708	0.089	0.723
		40	0.921	0.871	0.031	0.086	0.079	0.746	0.941	0.633	0.089	0.949
		10	0.943	0.889	0.301	0.031	0.104	0.722	0.859	0.444	0.304	0.747
Moderate (.70, .80)	.20	400	0.944	0.764	0.134	0.001	0.138	0.904	0.934	0.732	0.138	0.904
		100	0.943	0.727	0.090	0.000	0.115	0.733	0.938	0.546	0.080	0.932
		40	0.982	0.817	0.185	0.000	0.138	0.733	0.938	0.546	0.080	0.932
		10	0.992	0.817	0.185	0.000	0.138	0.733	0.938	0.546	0.080	0.932
Difficult (.40, .50)	.05	400	0.970	0.892	0.215	0.018	0.283	0.821	0.943	0.703	0.212	0.871
		100	0.946	0.829	0.421	0.003	0.421	0.891	0.934	0.853	0.402	0.870
		40	0.932	0.819	0.249	0.002	0.274	0.851	0.924	0.879	0.269	0.851
		10	0.993	0.898	0.933	0.572	0.753	0.851	0.985	0.899	0.944	0.942
	.00	400	0.901	0.895	0.449	0.013	0.450	0.830	0.978	0.903	0.449	0.865
		100	0.900	0.895	0.449	0.013	0.450	0.830	0.978	0.903	0.449	0.865
		40	0.900	0.895	0.449	0.013	0.450	0.830	0.978	0.903	0.449	0.865
		10	0.900	0.895	0.449	0.013	0.450	0.830	0.978	0.903	0.449	0.865
	.00	400	0.900	0.895	0.449	0.013	0.450	0.830	0.978	0.903	0.449	0.865
		100	0.900	0.895	0.449	0.013	0.450	0.830	0.978	0.903	0.449	0.865
		40	0.900	0.895	0.449	0.013	0.450	0.830	0.978	0.903	0.449	0.865
		10	0.900	0.895	0.449	0.013	0.450	0.830	0.978	0.903	0.449	0.865

Table 21: Comparison of the proportions of "correct decisions" at decision criteria of .90 and .95 confidence intervals, among results obtained from three separate (independent) results from Procedure A (results from Procedure A with adjustments to include some "no decisions" cases) and "correct decisions" cases, over the various combinations of validity (valid & invalid), skill difficulty, probability of guessing, and sample size levels.

* Indication that the two criteria, requiring the proportions of "correct decisions" to be larger than .90 and .95, were not simultaneously satisfied.

Indication that data were missing from certain results.

It may be noticed that some extra data grouped under the headings of "Adjusted Results from Procedure A" are included. The reasons for these adjustments will be explained in the next few paragraphs.

As discussed in the last section, different error types may have different effects on the valid and invalid datasets. Thus, it appears possible to utilize this property of the error types to aid in making validity decisions.

As an extreme case, suppose that a certain error type occurred only in the valid datasets and never occurred in the invalid datasets. Further, suppose that there were sufficient evidence to believe that such a pattern was reasonable. If, based on this information, the researcher treated the occurrences of this type of error as an indication of a valid hierarchy, he would increase the proportions of correct decisions over the valid datasets by the corresponding proportions due to this type of error. At the same time, the proportions of correct decisions over the invalid datasets would not be affected by such adjustments over the valid datasets. However, the proportions of incorrect decisions over the invalid datasets would be increased by the proportions related to this error type.

Similarly, an error type that was detected mainly in the invalid datasets, might be treated as an indication of an invalid hierarchy. Consequently, the proportions of correct decisions over the invalid datasets would be increased without affecting the proportions of correct decisions over the valid datasets. Nevertheless, the proportions of incorrect decisions over the valid datasets would be increased.

Assume that, there is another error type which was detected as often among the valid datasets as among the invalid datasets. Suppose this error type had to be classified as an indication of a valid or an invalid hierarchy. The logical action is to treat this type of error as an indication of an invalid hierarchy because, from a conservative viewpoint, it is a less serious decision error to consider a valid learning hierarchy as invalid than to consider an invalid learning hierarchy as valid.

If the information on the relative incidences of proportions of the various types of errors were not fully utilized, there would be sizeable proportions of "no decisions" due to the various error types. However, the information might be profitably used to eliminate the incidences of no decisions.

Recall that two error types, "negative variance of $\hat{P}_{II}$ " and "probability values tend to be zero", discussed in the last section, were detected more often over the valid datasets. There was evidence for believing that such patterns were reasonable. If each of these two types of errors is treated as an indication of a valid learning hierarchy, the proportion of "correct decisions" over each of the valid datasets can therefore be adjusted upward to include the proportions corresponding to the above two error types. This results in little change in the error rates for invalid hierarchies.

The error types, "marginal frequency equals zero" and "extremely large parameter estimates" appear to be detected as often among the valid datasets as among the invalid datasets. As discussed earlier, based on a conservative viewpoint (considering that it is a less serious error to indicate a valid hierarchy as invalid than to indicate an invalid hierarchy as valid), each of these two error types is therefore treated as an indication of an invalid hierarchy. The proportion of "correct decisions" over each of the invalid datasets is increased to include the proportions due to these two error types. These changes over the invalid datasets result in only slight changes in the error rates for the valid hierarchies. For each of the valid datasets, the proportion of "incorrect decisions" is

increased slightly by the inclusion of the proportions due to these two error types.

The other two error types, "non-convergence of estimates" and "pseudo-convergence of estimates", were generally detected more often over the invalid datasets. So, each of these two error types is treated as an indication of an invalid hierarchy. However, it may be observed that, in so doing, the proportion of correct decisions in each of the invalid datasets is increased sizeably (in general). The reason is that in each dataset, the proportion due to each of these two error types is fairly large. At the same time, the proportion of incorrect decisions in each of the valid datasets is adjusted upward to include the proportions due to these two error types.

The adjusting process, just discussed, seems to be reasonable, because the researcher is willing to increase the proportions of correct decisions over one type of learning hierarchies (e.g. invalid hierarchies) at the expense of increasing the proportions of incorrect decisions over the other type of learning hierarchies (e.g. valid hierarchies).

The adjusted results, as discussed above are shown in Table X. However, it may be appropriate to stress here that the adjusted results cannot be considered (in the strict

13-
case) as directly obtained from applying the MAX technique. The reason is that only some of the samples (not all) were involved in the complete process of parameter estimation, confidence interval construction and decision making. However, from a practical point of view, by treating each error type globally as an indication of a valid or an invalid hierarchy, only two categories of proportions, "correct" and "incorrect" were possible.

In fact, the adjusted results may be considered as obtained from a "hybrid" procedure for validating learning hierarchy. The reason for calling it a "hybrid" procedure is that this procedure is based partly on the MAX technique and partly on the proportions of errors detected in applying the MAX technique. Strictly speaking, only the unadjusted proportions of correct decisions may be used for comparing with the corresponding results obtained by Owston.

Using data obtained earlier (see Table X) Owston would claim that the (.80, .95) criteria, mentioned at the beginning of this section, were met 8 times at each of the 95% and 90% confidence intervals. However, from the unadjusted results obtained from using procedure A, the (.80, .95) criteria were not met for any of the same 27 valid and 27 invalid datasets. With the adjusted results, the (.80, .95) criteria were satisfied twice at each of the

95% and 90% confidence intervals. The (.80, .95) criteria were met in the two pairs (each pair contains one valid and one invalid dataset) of datasets related to moderate and difficult skills with probability of guessing of .05 and sample sizes of 400.

From the above comparisons, it can be concluded that Owston's results are different from the results obtained by using procedure A (which is identical to Owston's procedure for data collection with the exception that procedure A contains error checks) before and after the adjustments were made. This conclusion leads the researcher to question the findings obtained by Owston over the datasets utilized in the current study.

In the next section, the results obtained from using procedure B will be presented and discussed.

3.3 RESULTS OBTAINED FROM USING PROCEDURE B

As mentioned before, procedure B was designed as an improved version of Owston's procedure for validating learning hierarchies. In procedure B, some additional error checks were used. Also, there were four extra levels of decision criteria of .85%, .80%, .75% and 70% confidence intervals. Ten iterations rather than five were used in the parameter estimation process.

The results obtained from using procedure B are divided into two parts. The first part (shown in Table XI) is related to the various types of identified errors. There are altogether eight error types (classified into three error categories). Detailed descriptions of each type of error may be found in sections 1.5 and 2.1. The second part (shown in Table XII) of the results is related to the proportions of correct and incorrect decisions at each of the 95%, 90%, 85%, 80%, 75% and 70% confidence intervals.

In studying the first part of the results, the effects (on the various error proportions) of validity, skill difficulty, probability of guessing and sample sizes will be considered.

In Table XI, the proportions related to the various error types are shown. There are altogether three categories of error types detected in procedure B. The category I error types include "marginal frequency equals zero", "probability values tend to be zero", "extremely large estimates", and "negative variance of $\hat{P}_{II}$ ". It may be recognized that these error types were studied in the discussions of the results obtained from procedure A. In fact, the first three of the above four error types are identical to the corresponding ones detected by using procedure A because, for those error types, the checks used were common to both procedure A and procedure B.

		Procedure B (Part II)									
		Data sets for UNID hierarchies					Data sets for INWAD hierarchies				
		Experiment 1					Experiment 2				
Probability of error	Sample size N	Category I					Category II				
		Proportion of no decisions	Normal frequency = 0.70	Probability values tend to zero = 0.0	Extremely large estimates = 0.00	Negative variance of P = 0.00	Non-convergence of estimates = 0.00	Unreasonably large estimates = 0.00	Unreasonably large estimates = 0.00	Unreasonably large estimates = 0.00	Unreasonably large estimates = 0.00
High difficulty (0.1, 0.2, 0.3, 0.4, 0.5, 0.6, 0.7, 0.8, 0.9)	0.05	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
	0.10	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
	0.20	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
	0.30	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
Moderate difficulty (0.40, 0.50)	0.05	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
	0.10	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
	0.20	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
	0.30	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
Difficult difficulty (0.60, 0.70)	0.05	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
	0.10	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
	0.20	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
	0.30	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00

Table II: Proportions related to the various types of errors, detected by using procedure B, over the various combinations of validity, skill difficulty, probability of guessing and sample size.

Note: 1. For each dataset, proportion of "no decisions" is the sum of the proportions related to the various types of errors.
2. This table is to be used together with Table III.

The error type "negative variance of $\hat{P}_{II}$ " was defined differently in procedure B (as compared with procedure A). In procedure A, this error type was checked, after the estimate of P_{II} was obtained, at the end of the fifth iteration. However, in procedure B, check for convergence was made after each iteration. This error type was checked whenever the convergence of the estimate of P_{II} was obtained. Since ten iterations rather than five were allowed in procedure B, this error type occurred somewhere in between the first and the tenth iteration rather than at the end of the fifth iteration. As shown in Table XI, this error type was detected mainly in the valid datasets. It was explained in section 3.1 (where error types detected in procedure A were discussed), that the more frequent occurrences of this error type ("negative variance of $\hat{P}_{II}$ ") could be due to the fact that, in the valid datasets, the P_{II} estimate values are more likely to assume small negative values.

In error category II, the two error types detected in procedure B, are "non-convergence of estimates" and "unreasonably large estimates". Both error types are related to the problems of convergence of parameter estimates. The first type was detected at the end of the tenth iteration of the parameter estimation process. The second type was detected during the parameter estimation

process by means of a range check $(-1, 2)$ which was applied to each of the intermediate estimates.

In Table XI, it is shown that the "non-convergence of estimates" error type was detected more frequently over the invalid datasets. The explanations given to "non-convergence" cases in section 3.1 can also be used here to explain this pattern. However, it may be noted that the magnitudes of the proportions due to this error type (non-convergence of estimates) are smaller in the results obtained from using procedure B as compared with the corresponding results obtained from using procedure A. One explanation for this phenomenon is that there were more iterations allowed in procedure B (10 rather than 5). Another explanation is that the range check for "unreasonably large estimates" was utilized in procedure B to detect the cases which would potentially lead to "non-convergence" of parameter estimates. If this range check were removed, the proportions of "non-convergence" of estimates would be increased.

Over the valid datasets, the proportions, related to this error type, range from .000 to .094. The lower error proportions occurred in datasets associated with sample sizes of 400. The higher error proportions were related to small sample sizes of 40 and easy skill hierarchies. Over

the invalid datasets, the proportions range from .004 to .214. The higher proportions also tend to occur with small sample sizes of 40 and 100. The lower proportions were associated with sample sizes of 400 and small probability of guessing level of .05. The above patterns of error proportions were somewhat similar to those found in procedure A, except that, the magnitudes of the proportions of "non-convergence" in the results from using procedure A were of larger magnitudes.

Also, in Table XI, it is shown that the "unreasonably large estimates" error type was generally more often detected over the invalid datasets.

Over the valid datasets, the lower error proportions, for this error type, appear to be associated with sample sizes of 400, low probability of guessing of .05 and with moderate skill learning hierarchies. The high error proportions appear to be due to small sample sizes of 40. High proportions also seem to be related to easy skills hierarchies.

Over the invalid datasets, similar patterns may be observed, except that the proportions are generally higher than the valid counterparts.

In error category III, two error types are identified,

namely, "estimate $\hat{P}_{II}$ being a large negative value" and "other estimates being outside the (0,1) range". Both error types are related to improperly convergent parameter estimates.)

The first error type, in category III, was rarely detected in the invalid datasets. It is reasonable to expect that the error proportions of this type of error, were higher in the valid datasets because the population values of P_{II} over the valid datasets are all equal to zero. It may be observed, in Table XI, over the proportions in the valid datasets, that this error type was not detected in large sample sizes of 400 and was never detected in datasets with probability of guessing of .05 . It is generally expected that with large sample sizes and low measurement error of probability of guessing, the parameter estimation would be more accurate. According to the nature of this error type (detected when P_{II} is less than $-.05$), it is reasonable to find that this error type never occurred in samples sizes of 400 and guessing probability of .05 .

The second error type, in category III, is related to improperly convergent estimates other than P_{II} . (There were six other parameter estimates). The population parameter values of the other estimates range from .02 to .95 . The check utilized to detect this error type was a very liberal

range check of (-.05 to 1.05). Since the results on this error type were related to 6 parameter estimates, it is difficult to provide clear cut explanations on the observed patterns of proportions. More data may be useful for better understanding of the patterns. However, it may be noted that this error type was detected more frequently in the invalid datasets. Over the invalid datasets, the proportions range from .057 to .320 and over the valid datasets, the proportions range from .000 to .136 .

So far, only the first part of the results obtained from using procedure B has been presented and discussed. The second part of the results are related to correct and incorrect decisions at the various decision criteria. Brief discussions are given in the next paragraph.

As shown in Table XII, over the valid datasets, the proportions of correct decisions decrease as the width of the confidence interval decreases. However, over the invalid datasets, the proportions of correct decisions increase as the width of the confidence interval decreases. Also, the proportions of incorrect decisions are generally higher over the invalid datasets than over the valid datasets. The above phenomena were discussed when the results obtained from procedure A were studied.

Procedure B (Part II)

Difficulty Interval (in %)	Proportion of Successes	Sample Size N	Intervals for UNID Hierarchies										Intervals for INVAID Hierarchies									
			95% Conf. Int.		90% Conf. Int.		85% Conf. Int.		80% Conf. Int.		75% Conf. Int.		95% Conf. Int.		90% Conf. Int.		85% Conf. Int.		80% Conf. Int.		75% Conf. Int.	
			C	I	C	I	C	I	C	I	C	I	C	I	C	I	C	I	C	I	C	I
Easy (1-25%)	.25	100	95	5	90	10	85	15	80	20	75	25	95	5	90	10	85	15	80	20	75	25
	.20	100	95	5	90	10	85	15	80	20	75	25	95	5	90	10	85	15	80	20	75	25
	.15	100	95	5	90	10	85	15	80	20	75	25	95	5	90	10	85	15	80	20	75	25
	.10	100	95	5	90	10	85	15	80	20	75	25	95	5	90	10	85	15	80	20	75	25
Moderate (26-50%)	.50	100	95	5	90	10	85	15	80	20	75	25	95	5	90	10	85	15	80	20	75	25
	.45	100	95	5	90	10	85	15	80	20	75	25	95	5	90	10	85	15	80	20	75	25
	.40	100	95	5	90	10	85	15	80	20	75	25	95	5	90	10	85	15	80	20	75	25
	.35	100	95	5	90	10	85	15	80	20	75	25	95	5	90	10	85	15	80	20	75	25
Difficult (51-75%)	.75	100	95	5	90	10	85	15	80	20	75	25	95	5	90	10	85	15	80	20	75	25
	.70	100	95	5	90	10	85	15	80	20	75	25	95	5	90	10	85	15	80	20	75	25
	.65	100	95	5	90	10	85	15	80	20	75	25	95	5	90	10	85	15	80	20	75	25
	.60	100	95	5	90	10	85	15	80	20	75	25	95	5	90	10	85	15	80	20	75	25

Table XII: Proportions of correct and incorrect decisions, made by using procedure B at different confidence intervals, over the various combinations of validity, skill difficulty, and difficulty of the task. (Note: The numbers in the parentheses indicate the proportion of correct decisions.)

Note: 1. "C" indicates proportion of correct decisions.
2. "I" indicates proportion of incorrect decisions.

In the next section, results obtained from the two modified versions (A and B) of Owston's procedure for validating learning hierarchies will be compared.

3.4 COMPARISONS OF RESULTS OBTAINED FROM PROCEDURES A AND B

There were only two levels of decision criteria used in procedure A. These decision criteria are 95% and 90% confidence intervals. In this section, the proportions of correct decisions obtained in both procedures (A and B) at the 95% and 90% confidence intervals are compared. The comparison criteria of (.80, .95), used in section 3.2, will also be used here. In addition to the proportion of correct decisions, comparisons will be made on the proportion of "no decisions" due to the error types detected in each of the two procedures.

As additional information, the proportions of correct decisions (obtained in procedure B; shown in Table XII) were adjusted by utilizing the error proportions. The adjustment process will be mentioned in this section. The adjusted proportions of correct decisions obtained from each of the two procedures (A and B) were also compared by means of the (.80, .95) criteria, at each of the 95% and 90% confidence intervals.

At the end of this section, the adjusted proportions of "correct decisions", obtained in procedure B, at each of the 85%, 80%, 75% and 70% confidence intervals, will be shown. The incidences of meeting the (.80, .95) criteria will be briefly discussed.

As shown in Table XIII, the proportions of "no decisions" over the 27 valid and 27 invalid datasets, are generally lower in the results obtained from using procedure B. This implies that, by using procedure B, there were generally fewer incidences of "no decisions" cases. One main reason for this phenomenon is that ten rather than five iterations were allowed in procedure B. The proportion of "non-convergence", for each dataset, was thus reduced in procedure B. Consequently, the proportion of "no decisions" was also smaller.

Also, in Table XIII, it may be observed that at each of the 95% and 90% confidence intervals, the proportions of correct decisions are relatively higher in the results obtained from using procedure B. However, it is indicated in the data in Table XIII that the (.80, .95) criteria were not satisfied.

In Table XIV, comparisons between procedure A and B are made by utilizing the adjusted results. The adjusted

Skill Difficulty Level (1 = Easy, 4 = Difficult)	Probability of Success	Sample Size	Proportions of "No Decisions"		Proportions of Correct Decisions									
					95% Confidence Interval					90% Confidence Interval				
					Procedure A		Procedure B		Invalid Ns	Procedure A		Procedure B		Invalid Ns
					Valid Ds	Invalid Ns	Valid Ds	Invalid Ns		Valid Ds	Invalid Ns	Valid Ds	Invalid Ns	
Easy (.95-.90)	.25	100	0.794	0.252	0.206	0.001	0.615	0.004	0.205	0.038	0.413	0.034	0.013	0.034
		100	0.911	0.089	0.089	0.000	0.236	0.000	0.089	0.000	0.236	0.013	0.013	0.013
		40	0.848	0.152	0.351	0.004	0.115	0.024	0.351	0.008	0.115	0.024	0.024	0.024
		100	0.894	0.106	0.308	0.021	0.249	0.017	0.304	0.066	0.246	0.077	0.077	0.077
Moderate (.70-.50)	.20	100	0.866	0.134	0.134	0.001	0.303	0.011	0.134	0.001	0.303	0.015	0.015	0.015
		40	0.818	0.182	0.847	0.008	0.150	0.025	0.840	0.013	0.151	0.022	0.022	0.022
		100	0.815	0.185	0.051	0.002	0.285	0.003	0.051	0.026	0.285	0.030	0.030	0.030
		100	0.805	0.195	0.409	0.006	0.395	0.012	0.395	0.017	0.395	0.023	0.023	0.023
Difficult (.40-.25)	.05	100	0.804	0.196	0.200	0.008	0.200	0.008	0.200	0.008	0.200	0.008	0.008	0.008
		40	0.779	0.221	0.356	0.003	0.444	0.003	0.356	0.013	0.444	0.019	0.019	0.019
		100	0.771	0.229	0.598	0.000	0.402	0.001	0.598	0.000	0.402	0.001	0.001	0.001
		40	0.643	0.357	0.005	0.000	0.991	0.001	0.005	0.000	0.991	0.001	0.001	0.001
	.20	100	0.831	0.169	0.158	0.005	0.842	0.013	0.158	0.009	0.842	0.013	0.013	0.013
		40	0.817	0.183	0.205	0.001	0.795	0.008	0.205	0.009	0.795	0.008	0.008	0.008
		100	0.817	0.183	0.205	0.001	0.830	0.008	0.205	0.009	0.830	0.008	0.008	0.008
		40	0.817	0.183	0.205	0.001	0.830	0.008	0.205	0.009	0.830	0.008	0.008	0.008
	.05	100	0.817	0.183	0.205	0.001	0.830	0.008	0.205	0.009	0.830	0.008	0.008	0.008
		40	0.817	0.183	0.205	0.001	0.830	0.008	0.205	0.009	0.830	0.008	0.008	0.008
		100	0.817	0.183	0.205	0.001	0.830	0.008	0.205	0.009	0.830	0.008	0.008	0.008
		40	0.817	0.183	0.205	0.001	0.830	0.008	0.205	0.009	0.830	0.008	0.008	0.008

Table III: Comparisons of the proportions of "no decisions" and the proportions of correct decisions at 95% and 90% confidence intervals, for procedures A and B, over the various combinations of validity, skill difficulty, probability, sample size, and sample size levels.

Note: "Ds" indicates decision; for example "Valid Ds" indicates data that characterize valid hierarchical.

Adjusted Proportions of "Correct Decisions"														
Skill Difficulty Level (D, P, H)	Probability of Attending (A)	Sample Size (N)	75% Confidence Interval						90% Confidence Interval					
			Procedure A			Procedure B			Procedure A			Procedure B		
			valid D _a	invalid D _b	valid D _c	invalid D _d	valid D _e	invalid D _f	valid D _g	invalid D _h	valid D _i	invalid D _j	valid D _k	invalid D _l
Easy (.95, .90)	.25	400	0.204	0.733	0.615	0.547	0.205	0.780	0.613	0.577	0.205	0.780	0.613	0.577
		100	0.097	0.923	0.230	0.874	0.097	0.923	0.230	0.874	0.097	0.923	0.230	0.874
		40	0.077	0.946	0.143	0.857	0.077	0.946	0.143	0.857	0.077	0.946	0.143	0.857
		100	0.308	0.772	0.749	0.510	0.304	0.767	0.746	0.570	0.304	0.767	0.746	0.570
		40	0.138	0.904	0.307	0.693	0.138	0.904	0.307	0.693	0.138	0.904	0.307	0.693
Moderate (.70, .50)	.25	400	0.112	0.887	0.487	0.506	0.115	0.883	0.487	0.506	0.115	0.883	0.487	0.506
		100	0.709	0.774	0.953	0.704	0.786	0.883	0.950	0.834	0.786	0.883	0.950	0.834
		40	0.400	0.820	0.601	0.772	0.400	0.831	0.601	0.772	0.400	0.831	0.601	0.772
		100	0.293	0.871	0.722	0.584	0.282	0.878	0.722	0.584	0.282	0.878	0.722	0.584
		40	0.298	0.847	0.920	0.638	0.282	0.878	0.920	0.638	0.282	0.878	0.920	0.638
Difficult (.40, .25)	.25	400	0.421	0.691	0.448	0.612	0.421	0.701	0.448	0.612	0.421	0.701	0.448	0.612
		100	0.274	0.851	0.430	0.848	0.274	0.851	0.430	0.848	0.274	0.851	0.430	0.848
		40	0.953	0.051	0.991	0.009	0.944	0.056	0.942	0.056	0.944	0.056	0.942	0.056
		100	0.650	0.630	0.044	0.573	0.640	0.645	0.044	0.645	0.640	0.645	0.044	0.645
		40	0.379	0.807	0.558	0.802	0.379	0.810	0.558	0.802	0.379	0.810	0.558	0.802
	.05	400	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
		100	0.957	0.739	0.984	0.764	0.954	0.800	0.903	0.910	0.954	0.800	0.903	0.910
		40	0.810	0.644	0.848	0.648	0.810	0.610	0.810	0.610	0.810	0.610	0.810	0.610
		100	0.194	0.820	0.180	0.753	0.194	0.753	0.194	0.753	0.194	0.753	0.194	0.753
		40	0.143	0.923	0.229	0.874	0.143	0.923	0.229	0.874	0.143	0.923	0.229	0.874
	.25	400	0.772	0.764	0.933	0.742	0.741	0.808	0.921	0.881	0.721	0.881	0.921	0.881
		100	0.449	0.744	0.635	0.688	0.446	0.744	0.635	0.688	0.446	0.744	0.635	0.688
		40	0.240	0.884	0.175	0.892	0.231	0.884	0.175	0.892	0.231	0.884	0.175	0.892
		100	0.680	1.000	1.000	1.000	0.680	1.000	1.000	1.000	0.680	1.000	1.000	1.000
		40	0.945	0.804	0.984	0.705	0.945	0.705	0.984	0.705	0.945	0.705	0.984	0.705
	.05	400	0.772	0.711	0.863	0.727	0.772	0.727	0.863	0.727	0.772	0.727	0.863	0.727

Table 11: Comparisons of the adjusted (to include some of the "no decisions" cases as "correct decisions") cases) proportions of "correct decisions" at 75% & 90% confidence intervals, for procedures A and B, over the different combinations of validity, skill difficulty, probability of guessing, and sample size levels.

Note: "0.05" indicates "detectable" indicates that the two criteria, requiring the proportions of "correct decisions" to be larger than .40 for valid alternatives and to be larger than .25 for invalid alternatives, were simultaneously satisfied.

results related to procedure A were retrieved from Table X. The adjusted results from procedure B were obtained by the process described in the next paragraph.

Recall that in discussing the various types of errors detected by using procedure B, it was suggested (directly or indirectly) that some error types mainly occurred over the valid datasets while the rest occurred chiefly over the invalid datasets. The three error types, namely, "probability values tend to be zero", "negative variance of $\hat{P}_{II}$ ", and "convergent $\hat{P}_{II}$ estimate being a large negative number", were more often detected over the valid datasets. Each of these error types was treated as an indication of a valid learning hierarchy. For each valid dataset, the proportion of correct decisions at each of the decision criteria was adjusted upward to include the proportions due to these three error types. There was no change in the proportions of "correct decisions" over the invalid datasets. The other five error types, namely, "marginal frequency equals zero", ... etc., were detected more often over the invalid datasets. Each of these five error types was considered globally as an indication of an invalid hierarchy. For each invalid dataset, the proportion of correct decisions at each of the decision criteria was adjusted upward to include the proportions due to these five error types. The adjusted proportions of "correct

decisions" in procedure B were shown in Table XIV and Table XV.

It may be observed, in Table XIV, that at 95% confidence interval, with the adjusted results the (.80, .95) criteria were satisfied twice in the results obtained from each of the two procedures. However, at the 90% confidence interval, these criteria were satisfied three times in the adjusted results obtained from procedure B (as can be observed in the last two columns of Table XIV).

In Table XV, it is shown that, as the width of the confidence interval decreases, the (.80, .95) criteria tended to be more often met. At each of the 75% and 70% confidence intervals, in the adjusted results obtained from procedure B, the (.80, .95) criteria were satisfied eight times. Of these eight times, six occurred over sample sizes of 400 and two occurred over sample sizes of 100. With respect to measurement error of guessing, of these eight times, five occurred over the low guessing level of .05, two over the level of .20 and one over the level of .25. As far as skill difficulty level is concerned, of these 8 times, four occurred in datasets associated with moderate skills learning hierarchies, three in difficult skills learning hierarchies and one in the easy skills learning hierarchies.

Skill Difficulty Levels (D, D, I)	Probability of Success	Sample Size N	Adjusted Proportions of "Correct Decisions" for procedure B											
			80% Confidence Interval				75% Confidence Interval				70% Confidence Interval			
			Valid Data	Invalid Data	Valid Data	Invalid Data	Valid Data	Invalid Data	Valid Data	Invalid Data	Valid Data	Invalid Data	Valid Data	Invalid Data
Easy (.95+.90)	.25	400	0.409	0.406	0.404	0.454	0.393	0.712	0.234	0.804	0.575	0.740	0.220	0.808
		100	0.238	0.077	0.237	0.080	0.142	0.004	0.142	0.004	0.142	0.004	0.142	0.004
		40	0.142	0.004	0.142	0.004	0.004	0.004	0.004	0.004	0.004	0.004	0.004	0.004
		400	0.743	0.656	0.734	0.696	0.712	0.746	0.300	0.863	0.299	0.868	0.299	0.868
		100	0.307	0.053	0.305	0.057	0.183	0.036	0.183	0.036	0.183	0.036	0.183	0.036
Moderate (.70+.50)	.20	400	0.187	0.095	0.183	0.093	0.183	0.093	0.183	0.093	0.183	0.093	0.183	0.093
		100	0.095	0.004	0.093	0.004	0.093	0.004	0.093	0.004	0.093	0.004	0.093	0.004
		40	0.004	0.004	0.004	0.004	0.004	0.004	0.004	0.004	0.004	0.004	0.004	0.004
		400	0.401	0.270	0.374	0.801	0.370	0.801	0.370	0.801	0.370	0.801	0.370	0.801
		100	0.374	0.024	0.370	0.020	0.370	0.020	0.370	0.020	0.370	0.020	0.370	0.020
Difficult (.40+.25)	.05	400	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077
		100	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077
		40	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077	0.077
		400	0.443	0.463	0.428	0.493	0.421	0.863	0.421	0.863	0.418	0.870	0.418	0.870
		100	0.427	0.047	0.426	0.047	0.426	0.047	0.426	0.047	0.426	0.047	0.426	0.047
Very Difficult (.20+.05)	.05	400	0.040	0.007	0.041	0.007	0.041	0.007	0.041	0.007	0.041	0.007	0.041	0.007
		100	0.040	0.007	0.041	0.007	0.041	0.007	0.041	0.007	0.041	0.007	0.041	0.007
		40	0.040	0.007	0.041	0.007	0.041	0.007	0.041	0.007	0.041	0.007	0.041	0.007
		400	0.038	0.007	0.038	0.007	0.038	0.007	0.038	0.007	0.038	0.007	0.038	0.007
		100	0.038	0.007	0.038	0.007	0.038	0.007	0.038	0.007	0.038	0.007	0.038	0.007

Table IV: Adjusted proportions of "correct decisions", at 80%, 75%, 70% confidence intervals, for procedure B, over the different combinations of validity, skill difficulty, probability of guessing and sample size levels. Indicates that the two criteria, requiring the proportions of "correct decisions" to be larger than .80 for valid hierarchy and to be larger than .95 for invalid hierarchy, were simultaneously satisfied.

Recall that, using the results shown in Table X, Owston claimed that the (.80, .95) criteria were satisfied eight times, at each of the 95% and 90% confidence intervals. By comparing Table X with Table XV, it may be observed that the (.80, .95) criteria were met in the same eight pairs of datasets (one pair contains one valid and the corresponding invalid datasets). This implies that, based on the (.80, .95) criteria, the adjusted results obtained from using procedure B at each of the 75% and 70% confidence intervals, were as good as the results obtained by Owston at each of the 95% and 90% confidence intervals.

So far, the results obtained in this study together with the pertinent data obtained by Owston have been presented and discussed. In the next few pages, a brief summary and some general conclusions are made.

SUMMARY AND CONCLUSIONS

A learning hierarchy is a network of intellectual skills arranged in such a manner that the lower-order skills are hypothesized to be the prerequisites of the higher-order skills. White & Clark proposed a theoretical model to describe a two-skill learning hierarchy. In this model, two questions per skill were used to test skill possession; also, measurement errors were considered.

Owston developed a technique, called MAX, for validating two-skill learning hierarchies. The MAX technique is essentially a statistical procedure for estimating one of the parameters, P_{II} , of the White & Clark Model. This parameter is zero for valid hierarchies and larger than zero (but less than 1) for invalid hierarchies.

An iterative maximum likelihood method of estimation was utilized in the MAX technique. From an observed 3x3 pattern of subjects' responses to two questions per skill, a first approximation of the maximum likelihood estimates was obtained via the marginal frequency approach devised by White & Clark. The value of the P_{II} estimate, contained in the final set of estimates, was used for making a decision on the validity of the related hierarchy. To do this, a

confidence interval was constructed around the final P_{II} estimate. If this confidence interval included a zero value, it was concluded that the hierarchy is valid; otherwise it was concluded that the hierarchy is invalid. (note that this conclusion or decision may be correct or incorrect depending on the actual validity of the hierarchy).

Owston then implemented the MAX technique and a data generating procedure in a computer program to test the performance of the MAX technique.

For generating simulated data, Owston specified 48 sets of population parameters. These ~~48~~ sets were formed by combining the various levels of validity of hierarchy, skill difficulty and measurement error. From each of these parameter sets, 1000 random samples were drawn at each of four levels of sample sizes of 25, 40, 100 and 400.

Next, Owston set up criteria for evaluating the usefulness of the MAX technique. He required that the proportion of correct decisions should be higher than .80 for 1000 samples drawn from valid hierarchies and .95 for samples drawn from invalid hierarchies. These criteria, (.80, .95), were applied to analyze the usefulness of the MAX technique over the valid and invalid hierarchies associated with different levels of skill difficulty,

measurement error of guessing and sample size.

However, in Owston's computer program, various errors were identified in the part of the program where the MAX technique was implemented. These identified errors were grouped into three categories.

Category I error types were associated with problems that caused computer interruptions. Category II error types were related to the problems of non-convergence of parameter estimates in the iterative parameter estimation process. Category III error types were related to improperly convergent estimates.

The identification of the various error types leads the researcher to question the results obtained by Owston.

Two questions about the MAX technique were then raised. The first question was about the usefulness of the MAX technique when the identified errors were considered. The second question was about the possibilities of improving the implementation procedure of the MAX technique.

Corresponding to these two questions, two procedures, A and B, were developed. Procedure A was identical to Owston's procedure for implementing the MAX technique except that error checks were installed for counting the frequency of occurrence of the various identified error types. Two levels of decision criteria, the 95% and 90% confidence

intervals, were used in procedure A. Procedure B was designed to improve upon procedure A; ten rather than five iterations were used in procedure B; also four extra levels of decision criteria, the 85%, 80%, 75% and 70% confidence intervals were used in procedure B.

The design for collecting data, in this study, was similar to Owston's design. From Owston's 48 sets of population parameters, 18 sets were chosen. These 18 sets were formed by combining three levels of skill difficulty and three levels of measurement error of guessing, at each of the two levels of validity of learning hierarchy (valid and invalid). The three levels of skill difficulty were easy (0.95, 0.90), moderate (0.70, 0.50) and difficult (0.40, 0.25); the three levels of measurement errors (guessing probabilities) were .05, .20 and .25.

From four levels of sample sizes used by Owston, three levels, namely, 40, 100 and 400, were selected. 1000 random samples were generated from each set of parameters at each level of sample size. Each 1000 random sample thus generated was called a dataset. Hence, by combining 18 sets of population parameters with three levels of sample sizes, 27 datasets related to valid hierarchies (called valid datasets) and 27 corresponding datasets related to invalid learning hierarchies (called invalid datasets) were formed.

Owston's criteria, (.80, .95), for judging the usefulness of the MAX technique, were also utilized for each pair of datasets (one valid and one corresponding invalid datasets).

One unique feature in the current design is that besides calculating the proportions of correct and incorrect decisions when testing the MAX technique, the proportions due to each type of identified error were also calculated. The sum of the error proportions was called the proportion of "no decisions".

Results obtained from using procedure A were then presented. It was found that the proportions of "no decisions" for each pair of datasets were generally high. By analysing the proportions of "no decisions" according to the error types in the error categories, it was found that error types related to non-convergence of parameter estimates were generally detected more often than other error types. Furthermore, it was observed that non-convergence of estimates occurred more frequently in the invalid datasets.

The (.80, .95) criteria on the proportions of correct decisions were not met in any pair of the 27 pairs of datasets. Based on his results, Owston found that the (.80, .95) criteria were met eight times over these dataset pairs.

In analysing the error proportions, it was identified that some error types occurred more often over the valid datasets, while some error types occurred more often over the invalid datasets. These differences among the error types were utilized to aid in validity decision making. An extra set of results was obtained via an adjusting process, in which each error type was either classified as an indication of a valid hierarchy or an invalid hierarchy.

The adjusted results obtained from using procedure A were also compared with Owston's results. It was found that, with the adjusted results, the (.80, .95) criteria were met in two pairs of datasets at each of the 95% and 90% confidence intervals.

Since Owston's results were different from results obtained from using procedure A, before and after adjustments, it was concluded that Owston's results, over the 27 pairs of datasets used in this study, were questionable.

Then, the results obtained from using procedure B was analysed. The proportions of "no decisions" (due to various error types), detected in procedure B, over each pair of datasets, were compared with the corresponding proportions obtained in procedure A. It was found that lower proportions of "no decisions" were detected in procedure B.

One explanation for this phenomenon was that more iterations were allowed in procedure B.

At each of the 95% and 90% confidence intervals, the proportions of correct decisions, obtained in procedures A and B, were also compared. It was noted that, generally, over each dataset, the proportion of correct decisions obtained in procedure B was higher than those obtained in procedure A. This indicates that procedure B was in fact an improved version of the procedure for implementing the MAX technique.

With the error proportions being utilized, the results obtained from procedure B were adjusted. These adjusted results were compared with the corresponding adjusted results obtained from procedure A. At the 90% confidence interval, the (.80, .95) criteria were met, in procedure B results, in three pairs of datasets.

The adjusted results obtained from procedure B, at 85%, 80%, 75% and 70% confidence intervals were also presented and discussed. The (.80, .95) criteria tended to be more often met at the narrower confidence intervals. At each of the 75% and 70% confidence intervals, the (.80, .95) criteria were met eight times over 8 pairs of datasets. These 8 pairs of datasets were the same ones in which Owston found that the (.80, .95) criteria were met at each of the 95% and 90% confidence intervals.

It was observed that, of these 8 pairs of datasets, 6 pairs were associated with sample sizes of 400, 2 pairs were associated with sizes of 100. Based on these observations, it is recommended that in utilizing the MAX technique for validating learning hierarchies, large sample sizes should be used; preferably, the sizes of 400. Generally sample sizes of 40 or less should not be used.

Besides sample size, skill difficulty and measurement error of guessing are two factors that also influence the performance of the MAX technique. It was found that, at 75% and 70% confidence intervals, the (.80, .95) criteria were met in four dataset pairs associated with moderate skill hierarchies, in three dataset pairs associated with difficult skill hierarchies, and only in one dataset pair associated with easy skill learning hierarchies. For easy skills, such as those found in mastery tests, the MAX technique is generally not recommended, except when the probability of guessing is about .05 and when large sample sizes are used.

With respect to the measurement error of guessing, it was noted that at each of the 75% and 70% confidence intervals, the (.80, .95) criteria were satisfied five times when the guessing probabilities were very small; the (.80, .95) criteria were satisfied twice in dataset pairs.

associated with measurement error of guessing of .20, and only once with guessing probability of .25 . This implies that the use of the technique is questionable when using multiple choice test questions with five or four options.

Overall speaking, MAX can be a useful technique for validating learning hierarchies if narrower confidence intervals (such as 75% or 70%) were used, and if errors identified in the technique were utilized to aid in the decision making process.

As has been discussed, non-convergence of parameter estimates is one of the most frequently detected error type in the MAX technique. It is suggested that future researchers may find it interesting to resolve the problems associated with non-convergence of estimates in the MAX technique. Perhaps, if some other methods were used to obtain an initial set of parameter estimates from a 3x3 pattern of subject responses, the frequency of occurrence of the various error types might be drastically reduced. If that were the case, then the proportions of correct decisions and consequently the performance of the MAX technique might be improved.

A probabilistic technique for validating learning hierarchies

BIBLIOGRAPHY

- Allendoerfer, C.B., & Oakley, C.O. Principles of Mathematics (2nd ed.). McGraw-Hill Book Company, Inc., 1963.
- Cotton, J.W., Gallagher, J.P., & Marshall, S.P. The Identification and Decomposition of Hierarchical Tasks. American Educational Research Journal, Summer 1977, 14 (3), 189-212.
- Durell, A.B. A computer Simulation Study of Measures for Validating Learning Hierarchies. Paper presented at the Annual Meeting of the American Educational Research Association. Chicago, April 1974.
- Elandt-Johnson, R.C. Probability Models and Statistical Methods In Statistics. New York: Wiley, 1971.
- Gagne, R.M. The acquisition of knowledge. Psychological Review, 1962, 69 (4), 355-365.
- Gagne, R.M. The Conditions of Learning. New York: Holt, Rinehart, & Winston, 1965.
- Gagne, R.M. Learning Hierarchies. Educational Psychologist, November 1968, 6 (1), 1-9.
- Gagne, R.M., & White, R.T. Past and Future Research on Learning Hierarchies. Educational Psychologist, 1974, 11 (1), 19-28.
- Keith, V. Design and Analysis in Experimentation. Ottawa: University of Ottawa Press, 1972.
- Miller, I., & Freund, J.E. Probability and Statistics for Engineers. Englewood Cliffs, New Jersey: Prentice-Hall, Inc., 1965.
- Owston, R.D. A Comparison of Learning Hierarchy Validation Techniques using Simulated Data. Unpublished doctoral dissertation, University of Ottawa, 1978.
- Rao, C.R. Linear Statistical Inference and its Applications. New York: Wiley, 1968.

White, R.T., & Clark, R.M. A test of inclusion which allows for errors of measurement. Psychometrika, March 1973, 38 (1), 77-86.

White, R.T. Indexes Used in Testing the Validity of Learning Hierarchies. J. of Research in Science Teaching, 1974, 11 (1), 61-66.

APPENDICES

Appendix A

Determination of population parameters from difficulty levels of skills I and II

The difficulty of skill I, D_I , and the difficulty of skill II, D_{II} , can be expressed in terms of the four population parameters, P_O , P_I , P_{II} and P_B . P_O is the proportion of subjects possessing neither skill I nor skill II, P_I is the proportion possessing only skill I, P_{II} is the proportion possessing only skill II, P_B is the proportion possessing both skill I and skill II.

$$\begin{aligned} D_I &= P_I + P_B \\ &= \text{proportion possessing skill I,} \end{aligned}$$

$$\begin{aligned} D_{II} &= P_{II} + P_B \\ &= \text{proportion possessing skill II.} \end{aligned}$$

$$\bar{D}_I = 1 - D_I \quad , \text{ and}$$

$$\bar{D}_{II} = 1 - D_{II} .$$

Beginning with a choice of D_I , for example, to be 0.95, D_{II} can be determined via an assumption that in a valid hierarchy, the phi correlation coefficient, ϕ , is approximately 0.68, and that $P_{II}=0$. Solving a formula

$$\phi^2 = \frac{\bar{D}_I D_{II}}{\bar{D}_{II} D_I} \quad \text{for } D_{II}, \quad D_{II} = 0.898 \approx 0.90 .$$

The four P's can hence be found using the following formulae:

$$P_O = \bar{D}_I = 0.05 ,$$

$$P_I = D_I - D_{II} = 0.05 ,$$

$$P_{II} = 0 ,$$

$$P_B = D_{II} = 0.90 .$$

Appendix A

The P's for the corresponding invalid population (e.g. with $\phi = 0.25$) can be obtained by letting $P_{II} = x$, and $\phi = 0.25$, where $x = \bar{D}_I D_{II} - \phi \sqrt{D_I \bar{D}_I D_{II} \bar{D}_{II}}$. Solving for x using the previous D_I , D_{II} values, $x = 0.287 \pm .03$.

The other P values are found using formulae

$$P_O = \bar{D}_I - x = .02$$

$$P_I = D_I - D_{II} + x = .08$$

$$P_B = D_{II} - x = .87$$

At the skill difficulties of (.95, .90), there are four pairs of populations at the four levels of errors of guessing; however the P values for valid and invalid populations have definite patterns as follow*:

	P_O	P_I	P_{II}	P_B
valid	0.05	0.05	0.00	0.90
invalid	0.02	0.08	0.03	0.87

* Same results as obtained by Owston (p.96, set#1 and #9); the formulae derivations are based on Keith, p.67.

Appendix B

Obtaining an initial set of parameter estimates

Owston used the marginal frequencies method to obtain an initial set of parameter estimates required for the process of scoring. He followed the procedures devised by White & Clark. White & Clark set $\hat{P}_{II}=0$, $\hat{\theta}_b=0$ and $\hat{\theta}_c=1$. That is, the proportion possessing skill II without possessing skill I is set to zero, the probability of guessing skill I items is set to zero, and for those who possess skill II, their chance of answering skill II test items correctly are assumed to be 1. Then White & Clark derived formulae to express the other parameter estimates in terms of the marginal frequencies. With an observed 3x3 table as shown in Figure B (same 3x3 table used by Owston for illustrating the MAX technique), the following expressions can be formed. When the marginal frequencies, (a,b,....,f), are substituted into the expressions, numerical values of parameter estimates can be found.

$$\text{So, } \hat{\theta}_a = \frac{2a}{2a+b} = 0.9731 ,$$

$$\hat{\theta}_d = \frac{e}{e+2f} = 0.1929 ,$$

$$\hat{P}_B = 1 - \frac{(e+2f)^2}{4fN} = 0.5796 ,$$

$$\hat{P}_O = 1 - (P_B + P_I) = 0.0886 .$$

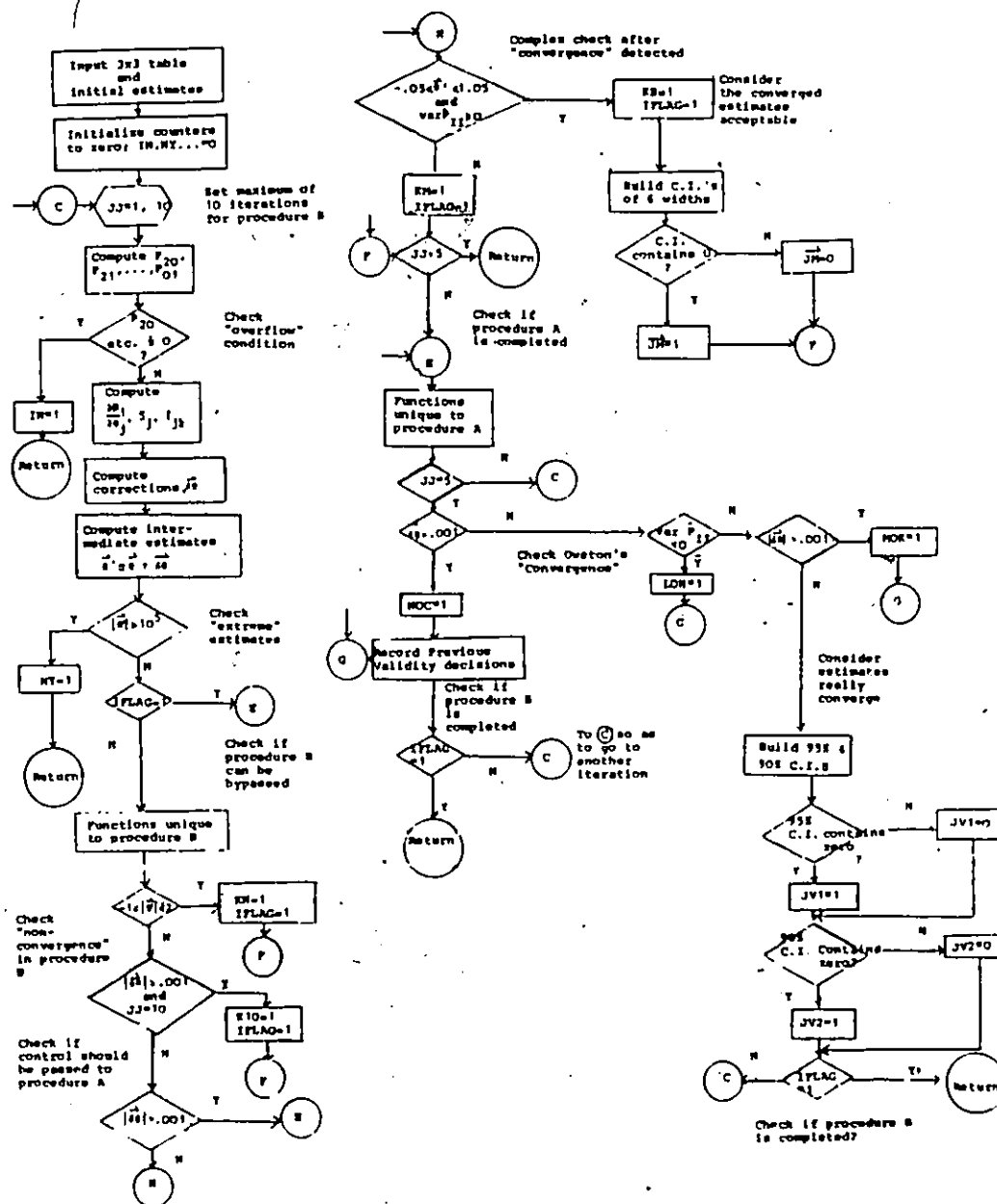
Appendix B

		No. of skill II items answered correctly			
		0	1	2	
No. of skill I items answered correctly	2	31	19	95	a=145
	1	4	0	4	b= 8
	0	11	3	1	c= 15
		f=	e=	d=	N=168
		46	22	100	

Figure 3: Hypothetical observed frequencies used in illustrating the method of calculating initial parameter estimates

Appendix C

Flowchart to show the various checks implemented in procedure A & B



POOR COPY.

Appendix D

```

WRITE(9,43)
WRITE(6,44)
WRITE(9,44)
WRITE(6,45)
WRITE(9,45)
WRITE(6,46)
WRITE(9,46)
WRITE(6,47)
WRITE(9,47)
WRITE(6,48)
WRITE(9,48)
WRITE(6,49)
WRITE(9,49)
43 FORMAT(' ',5X,'A',9X,'MARGINAL COLUMN TOTAL, F, EQUALS TO ZERO, CAU
CSING UNDEFINED SITUATION')
44 FORMAT(' ',5X,'B',9X,'INITIAL PARAMETER ESTIMATE(S) BEING OUTSIDE
CTHE RANGE (0,1)')
45 FORMAT(' ',5X,'C',9X,'MAX TECH. INVOLVES EST. (S) P20,P21,...VALUES
C, ONE OF MORE SUCH VALUE FOUND TO BE ZERO, CAUSING UNDEF. SITUATIONS
C')
46 FORMAT(' ',5X,'D',9X,'DURING ITERATIVE PROCESS, EXTREME ESTIMATES O
UTSIDE (-1,2) RANGE OCCURRED; SUBSEQUENT CONVG. NOT LIKELY')
47 FORMAT(' ',5X,'E',9X,'AT END OF 10 ITERATIONS CONVERGENCE WAS NOT
CATTAINED')
48 FORMAT(' ',5X,'F',9X,'CONV. WITHIN 10 ITER. BUT ESTIMATES OFF (-0.
COS, 1.05) RANGE OR ERROR VARIANCE FOR P2 < 0.')
49 FORMAT(' ',5X,'+',9X,'CONV. WITHIN 10 ITER. AND MAX TECH. CAN BE
CWARNINGFULLY APPLIED TO MAKE VALIDITY DECISIONS')
C** FIRST LOOP TO VARY SAMPLE SIZE AT 4 LEVELS
N=M(3)
C** SET CUMULATIVE COUNTERS TO ZERO
C** FOR EACH SAMPLE SIZE, AFTER GENERATING 1000 SAMPLES ON VALID HIERARCHIES
C** ANOTHER 1000 SAMPLES WILL BE GENERATED ON INVALID HIERARCHIES
C** SECOND LOOP FOR VALID AND INVALID HIERARCHIES
DO 30 J=1,2
WRITE(6,51)
WRITE(9,51)
WRITE(6,36)
WRITE(9,36)
WRITE(6,37)
WRITE(9,37)
WRITE(6,38)
WRITE(9,38)
WRITE(6,39)
WRITE(9,39)
51 FORMAT('1', 'LISTING OF INFORMATION ABOUT EACH SAMPLE GENERATED')
36 FORMAT('0', '1', 2X, 'EP', 1X, 'SAH', 1X, 'AP', 2X, 'AE', 2X, 'AD', 2X, 'BF', 2
CX, 'BE', 2X, 'BD', 2X, 'CF', 2X, 'CE', 2X, 'CD', 2X, 'INIT. PO', 6X, 'P1', 6X, 'P2
C', 6X, 'PB', 6X, 'TA', 6X, 'TB', 6X, 'TC', 6X, 'TD', 11X, 'HALT-ITER')
37 FORMAT(' ', '2', 9X, 'FINAL P20', 10X, 'P21', 10X, 'P22', 10X, 'P10', 10X, 'P
C11', 10X, 'P12', 10X, 'P00', 10X, 'P01', 10X, 'P02', 7X)
38 FORMAT(' ', '3', 9X, 'FINAL H11(P0)', 7X, 'H22(P1)', 8X, 'H33(P2)', 8X, 'H4
C4(TA)', 8X, 'H55(TB)', 8X, 'H66(TC)', 8X, 'H77(TD)', 6X, 'DETERMINANT')
39 FORMAT(' ', '4', 9X, 'FINAL', 2X, 'P0', 11X, 'P1', 11X, 'P2', 11X, 'P3', 11X,
C'TA', 11X, 'TB', 11X, 'TC', 11X, 'TD')
KSS=0
IEE(1)=0
IEE(2)=0
K10K=0
C CUMULATIVE COUNTERS FOR JC1, JC2, JP1, JP2 AND FREQ. OF PSEUDO CONVERGENCE
NJC1=0
NJC2=0
NJP1=0
NJP2=0
HJC1=0
HJC2=0

```

Appendix D

```

HJP1=0
HJP2=0
HOK=0
DO 1 I=1,10
  ICC(I)=0
  INN(I)=0
  KHH(I)=0
  KNN(I)=0
1 CONTINUE
  KHH(11)=0
  KHH(12)=0
C** LOOP FOR INITIALIZING THE COUNTERS TO ZERO AND INDICATOR TO BLANKS
DO 40 K=1,6
  HJP(K)=0.
  HJP(K)=0.
  HJ(K)=0
40 HJ(K)=0
C** THIRPD LOOP TO GENERATE 1000 RANDOM SAMPLES OF SIZE N, VALIDITY J
  IND(1)=AST(1)
  IND(2)=AST(2)
  HH=0
  HJ1=0
  HJ2=0
  HOC=0
  HLOH=0
  HJK1=0
  HJK2=0
  HJH1=0
  HJH2=0
  HJ1=0
  HJ2=0
  HJK1=0
  HJK2=0
  HJH1=0
  HJH2=0
  HNY=0
  HJY1=0
  HJY2=0
  HJY1=0
  HJY2=0
  JY1=0
  JY2=0
C LOOP TO GENERATE 1000 SAMPLES
DO 50 K=1,1000
  JJ=0
  HOC=0
  LOH=0
C JC1,JC2 USED FOR CHECKING PSEUDO CONVERGENCE, ALSO HOK
C JP1,JP2 USED FOR (P20,P21... CASES)
  JC1=0
  JC2=0
  HOK=0
  JP1=0
  JP2=0
  JY1=0
  JY2=0
  NY=0
C** RESTORE THE PARAMETER VALUES FOR EACH SAMPLE
  PO=PP0(J)
  P1=PP1(J)
  P2=PP2(J)
  PB=PPB(J)
  TA=TTA
  TB=TTB
  TC=TTC
  TD=TTD

```

Appendix D

```

      IF(K.LE.25) GOTO 17
      GOTO 18
17  HH=HH+1
      IF(HH.EQ.11) GOTO 82
      GOTO 83
82  WRITE(6,51)
      WRITE(9,51)
      WRITE(6,36)
      WRITE(9,36)
      WRITE(6,37)
      WRITE(9,37)
      WRITE(6,38)
      WRITE(9,38)
      WRITE(6,39)
      WRITE(9,39)
      HH=1
18  IF(K.GT.25) GOTO 19
      GOTO 83
19  HH=HH+1
      IF(HH.EQ.11) GOTO 98
      GOTO 83
98  WRITE(9,51)
      WRITE(9,36)
      WRITE(9,37)
      WRITE(9,38)
      WRITE(9,39)
      HH=1
C** CALL GEN SUBROUTINE TO CREATE A 3X3 TABLE OF OBSERVED FREQUENCIES
C** CALL WC SUBROUTINE TO ARRIVE AT AN INITIAL SET OF STATISTICS TO BE
C** USED BY THE MAX SUBROUTINE
83  CALL GEN(P0,P1,P2,PB,TA,TB,TC,TD,IS,IC,IF)
      IF(IE(2).EQ.1) GOTO 2
      PI(1)=P0
      PI(2)=P1
      PI(3)=P2
      PI(4)=PB
      PI(5)=TA
      PI(6)=TB
      PI(7)=TC
      PI(8)=TD
      IF(IC(9).EQ.1) GOTO 3
      GOTO 4
2   DO 5 I=1,2
      IEE(I)=IEE(I)+IE(I)
5   CONTINUE
      IF(K.LE.25) GOTO 20
      GOTO 21
20  WRITE(6,55) IND(J),ZRP(1),K,AP,AZ,AD,BF,BE,BD,CF,CZ,CD,JJ
55  FORMAT('0','1',1X,A1,1X,A1,I4,9(1X,I3),84X,I2)
21  WRITE(9,55) IND(J),ZRP(1),K,AP,AZ,AD,BF,BE,BD,CF,CZ,CD,JJ
      GOTO 50
3   DO 7 I=1,9
      ICC(I)=ICC(I)+IC(I)
7   CONTINUE
4   D=0.0
C** CALL MAX SUBROUTINE, PASSING THE CONFIDENCE INTERVAL WIDTHS AND
C** RETURN THE CORRESPONDING DECISION VALUES(1 INDICATES VALID, 0 INDICATES
C** INVALID)
      CALL MAX(P0,P1,P2,PB,TA,TB,TC,TD,ZH,JH,IN,KH,KN,K10,K,JJ,KS,PP,HH,
      CD,J,PI,JV1,JV2,JK1,JK2,JH1,JH2,LOH,NOC,NOK,NY)
      IF(NY.EQ.1) GOTO 75
      GOTO 76
75  HNY=HNY+NY
      JY1=JV1
      JY2=JV2
      NJY1=KNY1+JY1

```

Appendix D

```

NJJ2=NJJ2+JY2
NJJ1=NJJ1+1-JY1
NJJ2=NJJ2+1-JY2
76 IF(IN(10).EQ.1)GOTO 8
   IF(KN(9).EQ.1)GOTO 9
   IF(KN(10).EQ.1)GOTO 11
   IF(K10.EQ.1)GOTO 12
   IF(KS.EQ.1) KSS=KSS+KS
   GOTO 62
8 DO 13 I=1,10
  INN(I)=INN(I)+IN(I)
13 CONTINUE
  JP1=JV1
  JP2=JV2
  NJP1=NJP1+JP1
  NJP2=NJP2+JP2
  NJP1=NJP1+1-JP1
  NJP2=NJP2+1-JP2
  GOTO 62
9 DO 14 I=1,9
  KNN(I)=KNN(I)+KN(I)
14 CONTINUE
  GOTO 62
11 DO 16 I=1,12
  KHH(I)=KHH(I)+KH(I)
16 CONTINUE
  GOTO 62
12 K10K=K10K+K10
C** DEPENDENT ON VALIDITY OF HIERARCHIES, STANDARDS OF CORRECT DECISION ARE
C** DIFFERENT
62 NJ1=NJ1+JV1
   NJ2=NJ2+JV2
   NJ1=NJ1+1-JV1
   NJ2=NJ2+1-JV2
   NCCC=NCCC+NOC
   NLOH=NLOH+LOH
   NMOK=NMOK+MOK
   IF(NOC.EQ.1)GOTO 84
   GOTO 85
84 NJK1=NJK1+JK1
   NJK2=NJK2+JK2
   NJK1=NJK1+1-JK1
   NJK2=NJK2+1-JK2
85 IF(LOH.EQ.1)GOTO 86
   GOTO 87
86 NJH1=NJH1+JH1
   NJH2=NJH2+JH2
   NJH1=NJH1+1-JH1
   NJH2=NJH2+1-JH2
87 IF(MOK.EQ.1)GOTO 91
   GOTO 92
91 JC1=JV1
   JC2=JV2
   NJC1=NJC1+JC1
   NJC2=NJC2+JC2
   NJC1=NJC1+1-JC1
   NJC2=NJC2+1-JC2
92 IF(IN(10).EQ.1)GOTO 50
   IF(NY.EQ.1)GOTO 50
   IF(KN(9).EQ.1)GOTO 50
   IF(KN(10).EQ.1)GOTO 50
   IF(K10.EQ.1)GOTO 50
   IF(J.EQ.2)GOTO 90
C** LOOP TO SUM UP NUMBER OF INCORRECT DECISIONS FOR VALID HIERARCHIES
DO 60 L=1,6
  NJ(L)=NJ(L)+1-JH(L)

```

Appendix D

```

      NJ(L)=NJ(L)+JN(L)
60 CONTINUE
      GO TO 50
C** LOOP TO SUM UP NUMBER OF INCORRECT DECISIONS FOR INVALID HIERARCHIES
90 DO 80 L=1,6
      NJ(L)=NJ(L)+JN(L)
      NJ(L)=NJ(L)+1-JN(L)
80 CONTINUE
C** TRIED LOOP COMPLETED
50 CONTINUE
      K=K-1
C** FOR VALID HIERARCHIES
      IF (J.EQ.2) GO TO 100
      NTP=IEE(2)+KNN(9)+KNN(10)+K10K+INN(10)+KNY
      DO 78 L=1,6
      NJP(L)=100*FLOAT(NJ(L))/FLOAT(KSS)
      NJP(L)=100*FLOAT(NJ(L))/FLOAT(KSS)
78 CONTINUE
      WRITE(6,15) NP,D1,D2,TTA,TTB,TTC,TTD
      WRITE(9,15) NP,D1,D2,TTA,TTB,TTC,TTD
      WRITE(6,25) (PP0(I),PP1(I),PP2(I),PPB(I),I=1,2)
      WRITE(9,25) (PP0(I),PP1(I),PP2(I),PPB(I),I=1,2)
      WRITE(6,41) M(3)
      WRITE(9,41) M(3)
      WRITE(6,61)
      WRITE(9,61)
61 POPHAT('-', 'SUMMARY OF RESULTS FOR POPULATION CHARACTERISING VALID
  C HIERARCHIES')
      WRITE(6,67) K
      WRITE(9,67) K
67 POPHAT('0', 'NO. OF SAMPLES GENERATED', 41X, I4, //, 1X, 'SAMPLES REJECTED
  CD DUE TO ERRORS OF THE FOLLOWING TYPES', //, 6X, 'ERROR TYPE', 5X, 'BRIEF
  CF DESCRIPTIONS', 21X, 'FREQ')
      WRITE(6,68) IEE(2), ICC(9), INN(10)
      WRITE(9,68) IEE(2), ICC(9), INN(10)
68 FORMAT(' ', 10X, 'A', 7X, 'F=0', 38X, I3, //, 11X, 'B', 7X, 'INT. EST. NOT IN
  C(0,1) (NOT ADDED)', 21X, I3, //, 11X, 'C', 7X, 'ONE OF P20,P21... EQ.ZERO',
  C16X, I3)
      WRITE(6,69) KNN(9), K10K, KNN(10)
      WRITE(9,69) KNN(9), K10K, KNN(10)
69 FORMAT(' ', 10X, 'D', 7X, 'CPF (-1.0,2.0) RANGE IN ITERATION', 8X, I3, //,
  C11X, 'E', 7X, 'NO CONVERGENCE WHEN 10 ITER. COMPLETED', 3X, I3, //, 11X,
  C'F', 7X, 'OFF (-0.05,1.05) OR ERP. VAR. POP P2 < 0.', 1X, I3)
      WRITE(6,108) KNN(9), KNN(11), KNN(12)
      WRITE(9,108) KNN(9), KNN(11), KNN(12)
108 FORMAT(' ', 18X, 'SUPP. INFO. ON TYPE F PPPOP', //, 19X, 'VAR. P2<0.', 4
  CX, I3, //, 19X, 'P2<-0.05 : ', 5X, I3, //, 19X, 'P2>1.05 : ', 6X, I3)
      WRITE(6,101) KNY
      WRITE(9,101) KNY
101 POPHAT(' ', 10X, 'G', 7X, 'EXTREME PARAM. ESTIMATES', //, 8X, I3)
      WRITE(6,70) NTE, KSS
      WRITE(9,70) NTE, KSS
70 POPHAT(' ', 'TOTAL NO. OF SAMPLES REJECTED', 36X, I4, //, 1X, 'NO. OF S
  AMPLES USED TO INFER POPULATION VALIDITY', 17X, I4)
      WRITE(6,71)
      WRITE(9,71)
71 POPHAT('-', 'DECISION CRITERIA (CONF. INT.)', 1X, 5X, '95%', 12X, '90%',
  C12X, '85%', 12X, '80%', 12X, '75%', 12X, '70%')
      WRITE(6,72) (NJ(L), NJP(L), L=1,6)
      WRITE(9,72) (NJ(L), NJP(L), L=1,6)
72 FORMAT(' ', 'CORRECT DECISIONS (AND S)', 6X, 6(1X, I4, ' (', P6.2, 'X)'))
      WRITE(6,73) (NJ(L), NJP(L), L=1,6)
      WRITE(9,73) (NJ(L), NJP(L), L=1,6)
73 FORMAT(' ', 'INCORRECT DECISIONS (AND S)', 4X, 6(1X, I4, ' (', P6.2, 'X)')
      WRITE(6,74) (KSS, L=1,6)

```

Appendix D

```

WRITE(9,74) (KSS,L=1,6)
74 FORMAT(' ', 'TOTAL NO. OF DECISIONS(AND X)', 1X, 6(1X, I4, ' (100.00%)'
C))
GO TO 88
100 NTE=IEZ(2)+KNN(9)+KMM(10)+K10K+INN(10)+NNY
DO 79 L=1,6
NJP(L)=100*FLOAT(NJ(L))/FLOAT(KSS)
NJP(L)=100*FLOAT(MJ(L))/FLOAT(KSS)
79 CONTINUE
WRITE(6,15) NP,D1,D2,TTA,TTB,TTT,TTD
WRITE(9,15) NP,D1,D2,TTA,TTB,TTT,TTD
WRITE(6,25) (PP0(I),PP1(I),PP2(I),PPB(I),I=1,2)
WRITE(9,25) (PP0(I),PP1(I),PP2(I),PPB(I),I=1,2)
WRITE(6,41) M(3)
WRITE(9,41) M(3)
WRITE(6,64)
WRITE(9,64)
64 FORMAT(' ', 'SUMMARY OF RESULTS FOR POPULATION CHARACTERISING INVAL
CID HIEBACHIES')
WRITE(6,67) K
WRITE(9,67) K
WRITE(6,68) IEZ(2),ICC(9),INN(10)
WRITE(9,68) IEZ(2),ICC(9),INN(10)
WRITE(6,69) KNN(9),K10K,KMM(10)
WRITE(9,69) KNN(9),K10K,KMM(10)
WRITE(6,108) KMM(9),KMM(11),KMM(12)
WRITE(9,108) KMM(9),KMM(11),KMM(12)
WRITE(6,101) NNY
WRITE(9,101) NNY
WRITE(6,70) NTE,KSS
WRITE(9,70) NTE,KSS
WRITE(6,71)
WRITE(9,71)
WRITE(6,72) (NJ(L),NJP(L),L=1,6)
WRITE(9,72) (NJ(L),NJP(L),L=1,6)
WRITE(6,73) (MJ(L),NJP(L),L=1,6)
WRITE(9,73) (MJ(L),NJP(L),L=1,6)
WRITE(6,74) (KSS,L=1,6)
WRITE(9,74) (KSS,L=1,6)
C** SECOND LOOP COMPLETED
88 WRITE(6,15) NP,D1,D2,TTA,TTB,TTT,TTD
WRITE(9,15) NP,D1,D2,TTA,TTB,TTT,TTD
WRITE(6,25) (PP0(I),PP1(I),PP2(I),PPB(I),I=1,2)
WRITE(9,25) (PP0(I),PP1(I),PP2(I),PPB(I),I=1,2)
WRITE(6,41) M(3)
WRITE(9,41) M(3)
IF(J.EQ.2) WRITE(6,64)
IF(J.EQ.2) WRITE(9,64)
IF(J.EQ.1) WRITE(6,61)
IF(J.EQ.1) WRITE(9,61)
WRITE(6,77) K
WRITE(9,77) K
77 FORMAT('0', 'NO. OF SAMPLES GENERATED', 20X, I4)
NI1=K-NJ1
NI2=K-NJ2
WRITE(6,89) NJ1,NJ2,NI1,NI2
WRITE(9,89) NJ1,NJ2,NI1,NI2
89 FORMAT('0', 'IND. VALID', 10X, 2(I4, 6X), //, 1X, 'IND. INVAL', 10X, 2(I4, 6X)
C))
WRITE(6,95) NOCC,NJK1,NJK2,NJK1,NJK2
WRITE(9,95) NOCC,NJK1,NJK2,NJK1,NJK2
95 FORMAT('0', 'NO. OF NON-CONVERGENT CASES', //, 10X, I4, //, 1X, 'IND. VALID
C', 10X, 2(I4, 6X), //, 1X, 'IND. INVAL', 10X, 2(I4, 6X))
WRITE(6,96) NLOH,NJH1,NJH2,NJH1,NJH2
WRITE(9,96) NLOH,NJH1,NJH2,NJH1,NJH2
96 FORMAT('0', '-VE VARIANCE P2', //, 10X, I4, //, 1X, 'IND. VALID', 10X, 2(I4, 6X)

```

Appendix D

```

CX),/,1X,'IND. INVALID',10X,2(I4,6X))
WRITE(6,63) HMOX,NJC1,NJC2,NJC1,NJC2
WRITE(9,63) HMOX,NJC1,NJC2,NJC1,NJC2
63 FORMAT('0','NO. OF PSEUDO CONVERGENT CASES',/,10X,I4,/,1X,'END. VA
CLID',10X,2(I4,6X),/,1X,'IND. INVALID',8X,2(I4,6X))
WRITE(6,66) INH(10),NJP1,NJP2,NJP1,NJP2
WRITE(9,66) INH(10),NJP1,NJP2,NJP1,NJP2
66 POPHAT('0','NO. OF CASES P20 OF P21 ETC. =0 ',/,10X,I4,/,1X,'IND.
CVALID',10X,2(I4,6X),/,1X,'IND. INVALID',8X,2(I4,6X))
WRITE(6,81) NNY,NJY1,NJY2,NJY1,NJY2
WRITE(9,81) NNY,NJY1,NJY2,NJY1,NJY2
81 FORMAT('0','NO. OF EXTREME PARAM. EST. ',/,10X,I4,/,1X,'IND. VALID
C',10X,2(I4,6X),/,1X,'IND. INVALID',10X,2(I4,6X))
WRITE(6,65) (IEE(2),L=1,3)
WRITE(9,65) (IEE(2),L=1,3)
65 FORMAT('0','NO. OF P=0 CASES',/,10X,I4,/,1X,'IND. INVALID',8X,2(I
C4,6X))
NJ1T=NJ1-NJK1-NJH1-NJC1-NJP1-NJY1
NJ2T=NJ2-NJK2-NJH2-NJC2-NJP2-NJY2
NJ1T=K11-NJK1-NJH1-NJC1-NJP1-IEE(2)-NJY1
NJ2T=K12-NJK2-NJH2-NJC2-NJP2-IEE(2)-NJY2
WRITE(6,97) NJ1T,NJ2T,NJ1T,NJ2T
WRITE(9,97) NJ1T,NJ2T,NJ1T,NJ2T
97 FORMAT('0','CAN ONLY CLAIM',2X,'IND. VALID',5X,2(I4,6X),/,1X,'CAN
ONLY CLAIM',2X,'IND. INVALID',5X,2(I4,6X))
30 CONTINUE
C** FIRST LOOP COMPLETED
WRITE(6,31)
WRITE(9,31)
31 FORMAT('1',10(' '))
STOP
END
C*****
SUBROUTINE GEN(P0,P1,P2,PB,TA,TB,TC,TD,IS,IC,IE)
C** SUBROUTINE GEN TO GENERATE RANDOM SAMPLES ACCORDING TO SPECIFIED
C***** POPULATION PARAMETERS
REAL*8 P0,P1,P2,PB,TA,TB,TC,TD
COMMON AP,AE,AD,BF,BE,BD,CF,CE,CD,N
INTEGER AP,AE,AD,BF,BE,BD,CF,CE,CD
C** ALLOCATE MEMORY SPACE FOR 3X3 TABLE, 4 QUESTIONS AT 2 SKILL LEVELS.
C** PROBABILITIES OF CORRECTNESS AND TALLIES OF SCORES OF THE QUESTIONS
DIMENSION M(3,3),A(4),J(4),IC(10),IE(2)
C** SET P=0 & P=0 COUNTERS TO ZERO
IE(1)=0
IE(2)=0
C** INITIALIZE COUNTERS ON OFF RANGE P & THETA VALUES
DO 55 I=1,10
IC(I)=0
55 CONTINUE
C** SET THE 3X3 TABLE INITIALLY WITH ZEROS
DO 300 I=1,3
DO 300 K=1,3
M(K,I)=0
300 CONTINUE
C** SET UP CRITERIA FOR COMPARING WITH NUMBER TO BE GENERATED
RA=P0+P1
RB=P0+P1+P2
C** DETERMINE THE TRUE STATE OF SKILL POSSESSION FOR EACH SUBJECT
305 DO 310 KK=1,N
CALL RANDU(IS,IS,R)
IF (R.LE.P0) GO TO 320
IF (R.LE.RA) GO TO 330
IF (R.LE.RB) GO TO 340
C** THIS SUBJECT POSSESSES BOTH SKILLS
C** HIS CHANCES OF ANSWERING CORRECTLY ARE AS FOLLOW
A(1)=TA

```

Appendix D

```

      A(2)=TA
      A(3)=TC
      A(4)=TC
      GO TO 350
C** THIS SUBJECT POSSESSSES NEITHER OF THE SKILLS
320 A(1)=TB
      A(2)=TR
      A(3)=TD
      A(4)=TD
      GO TO 350
C** THIS SUBJECT POSSESSSES OF SKILL 1 ONLY
330 A(1)=TA
      A(2)=TA
      A(3)=TD
      A(4)=TD
      GO TO 350
C** THIS SUBJECT HAS ONLY SKILL 2
340 A(1)=TB
      A(2)=TR
      A(3)=TC
      A(4)=TC
C*****
C** GENERATE RANDOM NO. TO DETERMINE IF QUESTIONS ARE CORRECTLY ANSWERED
C*****
350 DO 360 I=1,4
      CALL RANDM (IS,IS,R)
C** WHEN A(I) IS GT OR EQ P ANSWER ASSUMED CORRECT
      IF (R-A(I)) 370,370,380
370 J(I)=1
      GO TO 360
380 J(I)=0
360 CONTINUE
C** ADD 1 TO THE SUM OF SCORES TO CHANGE ORIGIN FROM 0 TO 1
      K=J(1)+J(2)+1
      L=J(3)+J(4)+1
C** ENTER THE SCORES INTO THE CELLS OF THE TABLE
      H(K,L)=H(K,L)+1
310 CONTINUE
C*****
C** RENAMN THE CELLS ACCORDING TO MARGINAL FREQUENCIES SYMBOLS
C*****
      AF=H(3,1)
      AE=H(3,2)
      AD=H(3,3)
      BF=H(2,1)
      BE=H(2,2)
      BD=H(2,3)
      CF=H(1,1)
      CE=H(1,2)
      CD=H(1,3)
C** MARGINAL FREQUENCIES OF 3X3 CONTINGENCY TABLE FROM SUBROUTINE GEN
      AA=AF+AE+AD
      B=BF+BE+BD
      C=CF+CE+CD
      D=AD+BD+CD
      E=AE+BE+CE
      F=AF+BF+CF
C** CHECK IF MARGINAL COL TOTAL F=0 OR NOT
      IF (F.EQ.0.) IE(1)=1
C** CHECK IF MARGINAL COLUMN TOTAL F=0 OR NOT
      IF (F.EQ.0.) IE(2)=1
      IF (IE(2).EQ.1) GOTO 50
C** PARAMETER ASSUMPTIONS OF WC MODEL --P(II),THETA(B),THETA(C)
      P2=0.
      TB=0.
      TC=1.

```

Appendix D

```

C** THREE PARAMETERS ASSUMED FIVE OTHER ONES CALCULATED AS FOLLOW
TA=(2*AA)/(2*AA+B)
TD=E/(E+2*P)
PB=1.-((E+2*P)**2)/(4*P*N)
P1=((2*AA+B)**2)/(4*AA*N)-PB
P0=1.-(P1+PB)
IF((P0.LT.0.).OR.(P0.GT.1)) IC(1)=1
IF((P1.LT.0.).OR.(P1.GT.1)) IC(2)=1
IF((P2.LT.0.).OR.(P2.GT.1)) IC(3)=1
IF((PB.LT.0.).OR.(PB.GT.1)) IC(4)=1
IF((TA.LT.0.).OR.(TA.GT.1)) IC(5)=1
IF((TB.LT.0.).OR.(TB.GT.1)) IC(6)=1
IF((TC.LT.0.).OR.(TC.GT.1)) IC(7)=1
IF((TD.LT.0.).OR.(TD.GT.1)) IC(8)=1
DO 60 I=1,8
IF(IC(I).EQ.1) IC(9)=1
60 CONTINUE
50 RETURN
END
C. SUBROUTINE MAX
SUBROUTINE MAX(P0,P1,P2,PB,TA,TB,TC,TD,ZN,JH,IN,KH,KH,K10,K1,JJ,KS
C,PP,HH,D,J1,PI,JV1,JV2,JK1,JK2,JH1,JH2,LOH,NOC,MCK,NY)
IMPLICIT REAL*8 (A-Z)
INTEGER JH,IN,KH,K1,KH,K10,KS,IPLAG,IND,AST(8),NY
REAL*4 ZN
INTEGER JV1,JV2,NOC,LOH,JK1,JK2,JH1,JH2,MCK,LY
DATA AST/'V','I','D','E','P','+', '-', 'S','C'/
DOUBLE PRECISION DABS,DSQRT
DIMENSION S(7),ST(7,1),P(7,1),C(7,1),A(7),PI(8),PP(7)
DIMENSION H(7,7),L(7),M(7),PP(9),HH(7),PZ2(10),PF(9)
DIMENSION ZN(6),JH(6),IN(10),KH(12),KH(10),Z(6)
INTEGER I,J,K,JJ,N,J1
EQUIVALENCE (S(1),ST(1,1))
INTEGER AP,AE,AD,BF,BE,BD,CF,CE,CD
COMMON AP,AE,AD,BF,BE,BD,CF,CE,CD,N
C** K10 TO INDICATE NO CONVERGENCE AT END OF 10 ITERATIONS
C** KS=1 TO INDICATE SUCCESSFUL CONVERGENCE
K10=0
KS=0
C** VECTOR IN STORES COUNTS OF ERROR OF TYPE P==0, IN(10) A FLAG
C** KH STORES FREQ. OF P & THETA VALUES OFF (-0.5,1.5) RANGE DURING ITERATIONS
C** KH(9) BEING INDICATOR TO CHECK FOR EXIT FROM ITERATIONS
C** KH STORES COUNTS OF OFF RANGE (-.5,1.05) FOR P & THETA VALUES, AND
C** ALSO TEST SS ERROR FOR P2 BEING <0
DO 5 I=1,10
PZ2(I)=0.
IN(I)=0
KH(I)=0
KH(I)=0
5 CONTINUE
KH(11)=0
KH(12)=0
DO 1 I=1,7
HH(I)=0.
1 CONTINUE
IPLAG=0
NOC=0
LOH=0
C** FIRST OF ALL, SET INDICATOR TO ASSUME INVALIDITY
DO 209 K=1,6
JH(K)=0
209 CONTINUE
C JJ=NO. OF ITERATIONS
DO 199 JJ=1,10
TAA=1.-TA
TBB=1.-TB

```

Appendix D

```
TCC=1.-TC
TDD=1.-TD
TAS=TA**2
TBS=TB**2
TCS=TC**2
TDS=TD**2
TAAS=TA**2
TBBS=TB**2
TCCS=TCC**2
TDDS=TDD**2
```

C
C
C
C

PROBABILITIES OF A RANDOMLY SELECTED SUBJECT BEING CLASSIFIED
INTO EACH OF 3X3 TABLE

```
P20=P0*TBS*TDDS+P1*TAS*TDDS+P2*TBS*TCCS+PB*TAS*TCCS
P21=2*(P0*TBS*TD*TDD+P1*TAS*TD*TDD+P2*TBS*TC*TCC+
12*PB*TAS*TC*TCC
P22=P0*TBS*TDS+P1*TAS*TDS+P2*TBS*TCS+PB*TAS*TCS
P10=2*(P0*TB*TBB*TDDS+P1*TA*TAA*TDDS+P2*TB*TBB*TCCS+
1PB*TA*TAA*TCCS)
P11=4*(P0*TB*TBB*TD*TDD+P1*TA*TAA*TD*TDD+P2*TB*TBB*TC*TCC
1+PB*TA*TAA*TC*TCC)
P12=2*(P0*TB*TBB*TDS+TA*P1*TAA*TDS+P2*TB*TBB*TCS+PB*TA*TAA*TCS)
P00=P0*TBBS*TDDS+P1*TAAS*TDDS+P2*TBBS*TCCS+PB*TAAS*TCCS
P01=2*(P0*TBBS*TD*TDD+P1*TAAS*TD*TDD+P2*TBBS*TC*TCC+PB*TAAS*TC*TC
1C)
P02=P0*TBBS*TDS+P1*TAAS*TDS+P2*TBBS*TCS+PB*TAAS*TCS
PP(1)=P20
PP(2)=P21
PP(3)=P22
PP(4)=P10
PP(5)=P11
PP(6)=P12
PP(7)=P00
PP(8)=P01
PP(9)=P02
PZ2(JJ)=P02
```

C

```
DO 143 I=1,9
PP(I)=DABS(PP(I))
143 CONTINUE
DO 145 I=1,9
IF(PP(I).LT.1.00D-10) IN(I)=1
145 CONTINUE
DO 10 I=1,9
IF(IN(I).EQ.1) IN(10)=1
10 CONTINUE
IF(IN(10).EQ.1) GOTO 61
GOTO 62
61 WRITE(9,56) AST(J1),AST(8),K1,AP,AZ,AD,BF,BE,BD,CF,CE,CD,(PI(I),I=1
C,8),JJ
WRITE(9,57) (PP(I),I=1,9)
WRITE(9,58) (HH(I),I=1,7),D
WRITE(9,59) P0,P1,P2,PB,TA,TB,TC,TD
WRITE(9,32) (PZ2(LY),LY=1,10)
32 FORMAT(' ',5',10X,5(D15.7),/,1X,11X,5(D15.7))
IF(K1.LE.25) GOTO 16
GOTO 11
16 WRITE(6,56) AST(J1),AST(8),K1,AP,AZ,AD,BF,BE,BD,CF,CE,CD,(PI(I),I=1
C,8),JJ
56 FORMAT('0',1',1X,A1,1X,A1,I4,9(1X,I3),1X,8(1X,F7.4),19X,I2)
WRITE(6,57) (PP(I),I=1,9)
57 FORMAT(' ',2',10X,9(D12.4,1X))
WRITE(6,58) (HH(I),I=1,7),D
58 FORMAT(' ',3',10X,7(D15.7),D12.4)
WRITE(6,59) P0,P1,P2,PB,TA,TB,TC,TD
```

Appendix D

```

59 FORMAT(' ',4*,10X,8(D12.4,1X))
WRITE(6,32) (PZ2(LY),LY=1,10)
11 GOTO 239
C
62 PARTIAL DERIVATIVES
P20P0=TBS*TD*DDS-TAS*TCCS
P20P1=TAS*TD*DDS-TAS*TCCS
P20P2=TBS*TCCS-TAS*TCCS
P20TA=2.*(P1*TA*TD*DDS+PB*TA*TCCS)
P20TB=2.*(P0*TB*TD*DDS+P2*TB*TCCS)
P20TC=2.*(-P2*TB*TC+P2*TB*PB*TAS+TC*PB*TAS)
P20TD=2.*(-P0*TB*TD*P0*TB*P1*TAS+TD*P1*TAS)
P21P0=2.*(TBS*TD*TDD-TAS*TC*TCC)
P21P1=2.*(TAS*TD*TDD-TAS*TC*TCC)
P21P2=2.*(TBS*TC*TCC-TAS*TC*TCC)
P21TA=4.*(P1*TA*TD*TDD+PB*TA*TC*TCC)
P21TB=4.*(P0*TB*TD*TDD+P2*TB*TC*TCC)
P21TC=2.*P2*TB*TC+2.*PB*TAS-4.*PB*TAS*TC
P21TD=2.*P0*TB*TC+2.*P1*TAS-4.*P1*TD*TAS
P22P0=TBS*TDS-TAS*TCS
P22P1=TAS*TDS-TAS*TCS
P22P2=TBS*TCS-TAS*TCS
P22TA=2.*(P1*TD*PB*TCS)*TA
P22TB=2.*(P0*TD*P2*TC)*TB
P22TC=2.*(P2*TB*PB*TAS)*TC
P22TD=2.*(P0*TB*P1*TAS)*TD
P10P0=2.*(TB*TBB*TD*DDS-TA*TAA*TCCS)
P10P1=2.*TA*TAA*(TD*DDS-TCCS)
P10P2=2.*(TB*TBB-TA*TAA)*TCCS
P10TA=(2.*P1-4.*P1*TA)*TD*DDS+(2.*PB-4.*PB*TA)*TCCS
P10TB=(2.*P0-4.*P0*TB)*TD*DDS+(2.*P2-4.*P2*TB)*TCCS
P10TC=2.*(TBB*P2*TB+TAA*TA*PB)*(-2.*+2.*TC)
P10TD=2.*(P0*TB*TBB+P1*TA*TAA)*(-2.*+2.*TD)
P11P0=4.*(TB*TBB*TD*TDD-TA*TAA*TC*TCC)
P11P1=4.*TA*TAA*(TD*TDD-TC*TCC)
P11P2=4.*TC*TCC*(TB*TBB-TA*TAA)
P11TA=(4.*P1-8.*P1*TA)*TD*TDD+(4.*PB-8.*PB*TA)*TC*TCC
P11TB=(4.*P0-8.*P0*TB)*TD*TDD+(4.*P2-8.*P2*TB)*TC*TCC
P11TC=4.*(1.-2.*TC)*(P2*TBB*TB+PB*TA*TAA)
P11TD=4.*(1.-2.*TD)*(P0*TB*TBB+P1*TA*TAA)
P00P1=TAAS*TD*DDS-TAAS*TCCS
P00P0=TBS*TD*DDS-TAAS*TCCS
P00P2=TCCS*(TBB-TAAS)
P00TA=(-2.*+2.*TA)*(P1*TD*PB*TCCS)
P00TB=(-2.*+2.*TB)*(P0*TD*P2*TCCS)
P00TC=2.*(-1.*TC)*(P2*TBB*PB*TAAS)
P00TD=P0*(-2.*+2.*TD)*TBB*P1*(-2.*+2.*TD)*TAAS
P01P0=2.*(TD*TDD*TBB*TC*TCC*TAAS)
P01P1=2.*TAAS*(TD*TDD-TC*TCC)
P01P2=2.*TC*TCC*(TBB-TAAS)
P01TA=(-4.*P1+4.*P1*TA)*TD*TDD+(-4.*PB+4.*PB*TA)*TC*TCC
P01TB=(-4.*P0+4.*P0*TB)*TD*TDD+(-4.*P2+4.*P2*TB)*TC*TCC
P01TC=2.*(1.-2.*TC)*(P2*TBB*PB*TAAS)
P01TD=2.*(1.-2.*TD)*(P0*TBB*P1*TAAS)
P12P0=2.*(TB*TBB*TD*DDS-TA*TAA*TC)
P12P1=2.*TA*TAA*(TD*DDS-TCCS)
P12P2=2.*(TB*TBB*TD*DDS-TA*TAA*TC)
P12TA=(2.*P1-4.*P1*TA)*TD*DDS+(2.*PB-4.*PB*TA)*TCCS
P12TB=(2.*P0-4.*P0*TB)*TD*DDS+(2.*P2-4.*P2*TB)*TCCS
P12TC=4.*(P2*TB*TBB+PB*TA*TAA)*TC
P12TD=4.*TD*(P0*TB*TBB+P1*TA*TAA)
P02P0=TBS*TD*DDS-TAAS*TCCS
P02P1=TAAS*TD*DDS-TAAS*TCCS
P02P2=TBB*TD*DDS-TAAS*TCCS
P02TA=(-2.*P1+2.*P1*TA)*TD*DDS+(-2.*PB+2.*PB*TA)*TCCS
P02TB=(-2.*P0+2.*P0*TB)*TD*DDS+(-2.*P2+2.*P2*TB)*TCCS
P02TC=2.*TC*(P2*TBB*PB*TAAS)

```

Appendix D

P02TD=2.*TD*(P0*TBBS+P1*TAAS)
 COMPUTING EFFICIENCY SCORES
 C S(1)=AF*P20P0/P20+AE*P21P0/P21+AD*P22P0/P22+BF*P10P0/P10+BE*P11P0
 1/P11+BD*P12P0/P12+CF*P00P0/P00+CE*P01P0/P01+CD*P02P0/P02
 S(2)=AF*P20P1/P20+AE*P21P1/P21+AD*P22P1/P22+BF*P10P1/P10+BE*P11P1/
 1P11+BD*P12P1/P12+CF*P00P1/P00+CE*P01P1/P01+CD*P02P1/P02
 S(3)=AF*P20P2/P20+AE*P21P2/P21+AD*P22P2/P22+BF*P10P2/P10+BE*P11P2/
 1P11+BD*P12P2/P12+CF*P00P2/P00+CE*P01P2/P01+CD*P02P2/P02
 S(4)=AF*P20TA/P20+AF*P21TA/P21+AD*P22TA/P22+BF*P10TA/P10+BE*P11TA/
 1P11+BD*P12TA/P12+CF*P00TA/P00+CE*P01TA/P01+CD*P02TA/P02
 S(5)=AF*P20TB/P20+AE*P21TB/P21+AD*P22TB/P22+BF*P10TB/P10+BE*P11TB/
 1P11+BD*P12TB/P12+CF*P00TB/P00+CE*P01TB/P01+CD*P02TB/P02
 S(6)=AF*P20TC/P20+AF*P21TC/P21+AD*P22TC/P22+BF*P10TC/P10+BE*P11TC
 1/P11+BD*P12TC/P12+CF*P00TC/P00+CE*P01TC/P01+CD*P02TC/P02
 S(7)=AF*P20TD/P20+AE*P21TD/P21+AD*P22TD/P22+BF*P10TD/P10+BE*P11TD
 1/P11+BD*P12TD/P12+CF*P00TD/P00+CE*P01TD/P01+CD*P02TD/P02
 C INFORMATION MATRIX ELEMENTS
 H(1,1)=P20P0**2/P20+P21P0**2/P21+P22P0**2/P22+P10P0**2/P10+P11P0**
 12/P11+P12P0**2/P12+P00P0**2/P00+P01P0**2/P01+P02P0**2/P02
 H(1,2)=P20P0*P20P1/P20+P21P0*P21P1/P21+P22P0*P22P1/P22+P10P0*P10P1
 1/P10+P12P0*P12P1/P12+P00P0*P00P1/P00+P01P0*P01P1/P01+P02P0*P02P1/
 2P02+P11P0*P11P1/P11
 H(1,3)=P20P0*P20P2/P20+P21P0*P21P2/P21+P22P0*P22P2/P22+P10P0*P10P2
 1/P10+P11P0*P11P2/P11+P12P0*P12P2/P12+P00P0*P00P2/P00+P01P0*P01P2/P
 201+P02P0*P02P2/P02
 H(1,4)=P20P0*P20TA/P20+P21P0*P21TA/P21+P22P0*P22TA/P22+P10P0*P10TA
 1/P10+P11P0*P11TA/P11+P12P0*P12TA/P12+P00P0*P00TA/P00+P01P0*P01TA/
 2P01+P02P0*P02TA/P02
 H(1,5)=P20P0*P20TB/P20+P21P0*P21TB/P21+P10P0*P10TB/P10+P11P0*P11TB
 1/P11+P12P0*P12TB/P12+P00P0*P00TB/P00+P01P0*P01TB/P01+P02P0*P02TB/
 2P02+P22P0*P22TB/P22
 H(1,6)=P20P0*P20TC/P20+P21P0*P21TC/P21+P10P0*P10TC/P10+P11P0*P11TC
 1/P11+P12P0*P12TC/P12+P00P0*P00TC/P00+P01P0*P01TC/P01+P02P0*P02TC/
 2P02+P22P0*P22TC/P22
 H(1,7)=P20P0*P20TD/P20+P21P0*P21TD/P21+P22P0*P22TD/P22+P10P0*P10TD
 1/P10+P12P0*P12TD/P12+P00P0*P00TD/P00+P01P0*P01TD/P01+P02P0*P02TD/
 2P02+P11P0*P11TD/P11
 H(2,2)=P20P1**2/P20+P21P1**2/P21+P22P1**2/P22+P10P1**2/P10+P11P1**2/
 12/P11+P12P1**2/P12+P00P1**2/P00+P01P1**2/P01+P02P1**2/P02
 H(2,3)=P20P1*P20P2/P20+P21P1*P21P2/P21+P22P1*P22P2/P22+P10P1*P10P2
 1/P10+P11P1*P11P2/P11+P12P1*P12P2/P12+P00P1*P00P2/P00+P01P1*P01P2/P
 201+P02P1*P02P2/P02
 H(2,4)=P20P1*P20TA/P20+P21P1*P21TA/P21+P22P1*P22TA/P22+P10P1*P10TA
 1/P10+P11P1*P11TA/P11+P12P1*P12TA/P12+P00P1*P00TA/P00+P01P1*P01TA/
 2P01+P02P1*P02TA/P02
 H(2,5)=P20P1*P20TB/P20+P21P1*P21TB/P21+P22P1*P22TB/P22+P10P1*P10TB
 1/P10+P11P1*P11TB/P11+P12P1*P12TB/P12+P00P1*P00TB/P00+P01P1*P01TB/P
 201+P02P1*P02TB/P02
 H(2,6)=P20P1*P20TC/P20+P21P1*P21TC/P21+P22P1*P22TC/P22+P10P1*P10TC
 1/P10+P11P1*P11TC/P11+P12P1*P12TC/P12+P00P1*P00TC/P00+P01P1*P01TC/
 2P01+P02P1*P02TC/P02
 H(2,7)=P20P1*P20TD/P20+P21P1*P21TD/P21+P22P1*P22TD/P22+P10P1*P10TD
 1/P10+P11P1*P11TD/P11+P12P1*P12TD/P12+P00P1*P00TD/P00+P01P1*P01TD/
 2P01+P02P1*P02TD/P02
 H(3,3)=P20P2**2/P20+P21P2**2/P21+P22P2**2/P22+P10P2**2/P10+P11P2**
 12/P11+P12P2**2/P12+P00P2**2/P00+P01P2**2/P01+P02P2**2/P02
 H(3,4)=P20P2*P20TA/P20+P21P2*P21TA/P21+P22P2*P22TA/P22+P10P2*P10TA
 1/P10+P11P2*P11TA/P11+P12P2*P12TA/P12+P00P2*P00TA/P00+P01P2*P01TA/
 2P01+P02P2*P02TA/P02
 H(3,5)=P20P2*P20TB/P20+P21P2*P21TB/P21+P22P2*P22TB/P22+P10P2*P10TB
 1/P10+P11P2*P11TB/P11+P12P2*P12TB/P12+P00P2*P00TB/P00+P01P2*P01TB
 2/P01+P02P2*P02TB/P02
 H(3,6)=P20P2*P20TC/P20+P21P2*P21TC/P21+P22P2*P22TC/P22+P10P2*P10TC
 1/P10+P11P2*P11TC/P11+P12P2*P12TC/P12+P00P2*P00TC/P00+P01P2*P01TC
 2/P01+P02P2*P02TC/P02
 H(3,7)=P20P2*P20TD/P20+P21P2*P21TD/P21+P22P2*P22TD/P22+P10P2*P10TD

POOR COPY

Appendix D

```

1/P10+P11P2*P11TD/P11+P12P2*P12TD/P12+P00P2*P00TD/P00+P01P2*P01TD
2/P01+P02P2*P02TD/P02
H(4,4)=P20TA**2/P20+P21TA**2/P21+P22TA**2/P22+P10TA**2/P10+P11TA**
12/P11+P12TA**2/P12+P00TA**2/P00+P01TA**2/P01+P02TA**2/P02
H(4,5)=P20TA*P20TB/P20+P21TA*P21TB/P21+P22TA*P22TB/P22+P10TA*P10TB
1/P10+P11TA*P11TB/P11+P12TA*P12TB/P12+P00TA*P00TB/P00+P01TA*P01TB
2/P01+P02TA*P02TB/P02
H(4,6)=P20TA*P20TC/P20+P21TA*P21TC/P21+P22TA*P22TC/P22+P10TA*P10TC
1/P10+P11TA*P11TC/P11+P12TA*P12TC/P12+P00TA*P00TC/P00+P01TA*P01TC
2/P01+P02TA*P02TC/P02
H(4,7)=P20TA*P20TD/P20+P21TA*P21TD/P21+P22TA*P22TD/P22+P10TA*P10TD
1/P10+P11TA*P11TD/P11+P12TA*P12TD/P12+P00TA*P00TD/P00+P01TA*P01TD
2/P01+P02TA*P02TD/P02
H(5,5)=P20TB**2/P20+P21TB**2/P21+P22TB**2/P22+P10TB**2/P10+P11TB**
12/P11+P12TB**2/P12+P00TB**2/P00+P01TB**2/P01+P02TB**2/P02
H(5,6)=P20TB*P20TC/P20+P21TB*P21TC/P21+P22TB*P22TC/P22+P10TB*P10TC
1/P10+P11TB*P11TC/P11+P12TB*P12TC/P12+P00TB*P00TC/P00+P01TB*P01TC/
2/P01+P02TB*P02TC/P02
H(5,7)=P20TB*P20TD/P20+P21TB*P21TD/P21+P22TB*P22TD/P22+P10TB*P10TD
1/P10+P11TB*P11TD/P11+P12TB*P12TD/P12+P00TB*P00TD/P00+P01TB*P01TD
2/P01+P02TB*P02TD/P02
H(6,6)=P20TC**2/P20+P21TC**2/P21+P22TC**2/P22+P10TC**2/P10+P11TC**
12/P11+P12TC**2/P12+P00TC**2/P00+P01TC**2/P01+P02TC**2/P02
H(6,7)=P20TC*P20TD/P20+P21TC*P21TD/P21+P22TC*P22TD/P22+P10TC*P10TD
1/P10+P11TC*P11TD/P11+P12TC*P12TD/P12+P00TC*P00TD/P00+P01TC*P01TD/
2/P01+P02TC*P02TD/P02
H(7,7)=P20TD**2/P20+P21TD**2/P21+P22TD**2/P22+P10TD**2/P10+P11TD**
12/P11+P12TD**2/P12+P00TD**2/P00+P01TD**2/P01+P02TD**2/P02
DO 99 J=1,5
K=J+1
DO 99 I=K,7
99 H(I,J)=H(J,I)
H(7,6)=H(6,7)
DO 80 I=1,7
DO 80 J=1,7
80 H(I,J)=H(I,J)*K
CALL HINV(H,7,D,L,H)
DO 52 I=1,7
HH(I)=H(I,I)
52 CONTINUE
DO 500 I=1,7
C(I,1)=0.
500 CONTINUE
DO 510 I=1,7
DO 510 J=1,7
C(I,1)=C(I,1)+H(I,J)*ST(J,1)
510 CONTINUE
P(1,1)=P0
P(2,1)=P1
P(3,1)=P2
P(4,1)=TA
P(5,1)=TB
P(6,1)=TC
P(7,1)=TD
DO 101 J=1,7
P(J,1)=P(J,1)+C(J,1)
101 CONTINUE
P0=P(1,1)
P1=P(2,1)
P2=P(3,1)
PB=1-(P(1,1)+P(2,1)+P(3,1))
TA=P(4,1)
TB=P(5,1)
TC=P(6,1)
TD=P(7,1)
NY=0

```

POOR COPY

Appendix D

```
DO 39 J=1,7
PR(J)=DABS(P(J,1))
IF(PR(J).GT.1.0D5) NY=1
39 CONTINUE
IF(NY.EQ.1) GOTO 40
GOTO 42
40 WRITE(9,69) AST(J1), K1, AF, AE, AD, BF, BE, BD, CF, CE, CD, (PI(I), I=1,8), JJ
69 FORMAT('0', '1', 1X, A1, 1X, ' ', I4, 9(1X, I3), 1X, 8(1X, F7.4), 19X, I2)
WRITE(9,57) (PP(I), I=1,9)
WRITE(9,58) (HH(I), I=1,7), D
WRITE(9,59) PO, P1, P2, PB, TA, TB, TC, TD
WRITE(9,32) (PZ2(LY), LY=1,10)
IF(K1.LE.25) GOTO 48
GOTO 47
48 WRITE(6,69) AST(J1), K1, AF, AE, AD, BF, BE, BD, CF, CE, CD, (PI(I), I=1,8), JJ
WRITE(6,57) (PP(I), I=1,9)
WRITE(6,58) (HH(I), I=1,7), D
WRITE(6,59) PO, P1, P2, PB, TA, TB, TC, TD
WRITE(6,32) (PZ2(LY), LY=1,10)
47 IK(10)=0
KN(9)=0
KN(10)=0
K10=0
KS=0
NOC=0
LOH=0
NOK=0
GOTO 239
42 IF(IPLAG.EQ.1) GOTO 45
IF((PO.LT.-1.0).OR.(PO.GT.2.0)) KN(1)=1
IF((P1.LT.-1.0).OR.(P1.GT.2.0)) KN(2)=1
IF((P2.LT.-1.0).OR.(P2.GT.2.0)) KN(3)=1
IF((PB.LT.-1.0).OR.(PB.GT.2.0)) KN(4)=1
IF((TA.LT.-1.0).OR.(TA.GT.2.0)) KN(5)=1
IF((TB.LT.-1.0).OR.(TB.GT.2.0)) KN(6)=1
IF((TC.LT.-1.0).OR.(TC.GT.2.0)) KN(7)=1
IF((TD.LT.-1.0).OR.(TD.GT.2.0)) KN(8)=1
DO 20 I=1,8
IF(KN(I).EQ.1) KN(9)=1
20 CONTINUE
IF(KN(9).EQ.1) GOTO 49
GOTO 50
49 IPLAG=1
IF(KN(9).EQ.1) IND=AST(3)
IF(KN(10).EQ.1) IND=AST(5)
IF(K10.EQ.1) IND=AST(4)
IF(KS.EQ.1) IND=AST(6)
WRITE(9,56) AST(J1), IND, K1, AF, AE, AD, BF, BE, BD, CF, CE, CD, (PI(I), I=1,8)
C, JJ
WRITE(9,57) (PP(I), I=1,9)
WRITE(9,58) (HH(I), I=1,7), D
WRITE(9,59) PO, P1, P2, PB, TA, TB, TC, TD
WRITE(9,32) (PZ2(LY), LY=1,10)
IF(K1.LE.25) GOTO 26
GOTO 29
26 WRITE(6,56) AST(J1), IND, K1, AF, AE, AD, BF, BE, BD, CF, CE, CD, (PI(I), I=1,8)
C, JJ
WRITE(6,57) (PP(I), I=1,9)
WRITE(6,58) (HH(I), I=1,7), D
WRITE(6,59) PO, P1, P2, PB, TA, TB, TC, TD
WRITE(6,32) (PZ2(LY), LY=1,10)
29 IF(JJ.GT.5) GOTO 239
GOTO 45
50 DO 198 K=1,7
A(K)=DABS(C(K,1))
IF((A(K).GT..001).AND.(JJ.EQ.10)) GOTO 25
```

POOR COPY

Appendix D

```
IF (A(K).GT..001) GO TO 45
198 CONTINUE
GOTO 506
25 K10=1
GOTO 49
506 IF((P0.LT.-0.05).OR.(P0.GT.1.05)) KH(1)=1
IF((P1.LT.-0.05).OR.(P1.GT.1.05)) KH(2)=1
IF((P2.LT.-0.05).OR.(P2.GT.1.05)) KH(3)=1
IF((P3.LT.-0.05).OR.(P3.GT.1.05)) KH(4)=1
IF((P4.LT.-0.05).OR.(P4.GT.1.05)) KH(5)=1
IF((P5.LT.-0.05).OR.(P5.GT.1.05)) KH(6)=1
IF((P6.LT.-0.05).OR.(P6.GT.1.05)) KH(7)=1
IF((P7.LT.-0.05).OR.(P7.GT.1.05)) KH(8)=1
IF(H(3,3).LT.0.) KH(9)=1
IF(P2.LT.-0.05) KH(11)=1
IF(P2.GT.1.05) KH(12)=1
DO 46 I=1,9
IF(KH(I).EQ.1) KH(10)=1
46 CONTINUE
IF(KH(10).EQ.1) GOTO 49
KS=1
SQ=DSQRT(H(3,3))
DO 208 K=1,6
Z(K)=ZN(K)*SQ
IF(P2.LT.Z(K)) JH(K)=1
208 CONTINUE
GOTO 49
45 IF(JJ.EQ.5) GOTO 12
GOTO 199
12 WRITE(9,56) AST(J1),AST(7),K1,AF,AE,AD,BF,BE,BD,CF,CE,CD,(PI(I),I=1
C,8),JJ
WRITE(9,57) (PP(I),I=1,9)
WRITE(9,58) (HH(I),I=1,7),D
WRITE(9,59) P0,P1,P2,P3,P4,P5,P6,P7,TA,TB,TC,TD
WRITE(9,87) (C(K,1),K=1,7)
IF(K1.LE.25) GOTO 60
GOTO 63
60 WRITE(6,56) AST(J1),AST(7),K1,AF,AE,AD,BF,BE,BD,CF,CE,CD,(PI(I),I=1
C,8),JJ
WRITE(6,57) (PP(I),I=1,9)
WRITE(6,58) (HH(I),I=1,7),D
WRITE(6,59) P0,P1,P2,P3,P4,P5,P6,P7,TA,TB,TC,TD
WRITE(6,87) (C(K,1),K=1,7)
87 FORMAT(' ',5',10X,3(D12.4,1X),13X,4(D12.4,1X))
63 DO 13 K=1,7
IF(C(K,1).GT..001) GOTO 14
13 CONTINUE
GOTO 15
14 NOK=1
GOTO 17
15 IF(H(3,3).LT.0.) GOTO 28
DO 27 K=1,7
A(K)=DABS(C(K,1))
IF(A(K).GT..001) NOK=1
27 CONTINUE
SQ=DSQRT(H(3,3))
Z1=1.96*SQ
Z2=1.65*SQ
IF(P2-Z1) 18,18,19
18 JV1=1
GOTO 21
19 JV1=0
21 IF(P2-Z2) 22,22,23
22 JV2=1
GOTO 17
23 JV2=0
```

POOR COPY

Appendix D

```

GOTO 17
28 LOH=1
17 CONTINUE
IF(NOC.EQ.1) GOTO 30
GOTO 41
30 JK1=JV1
JF2=JV2
GOTO 44
41 IF(LOH.EQ.1) GOTO 43
GOTO 44
43 JH1=JV1
JH2=JV2
44 IF(IFLAG.EQ.1) GOTO 239
199 CONTINUE
239 RETURN
END

```

C		MINV	10
C		MINV	20
C	SUBROUTINE MINV	MINV	30
C		MINV	40
C	PURPOSE	MINV	50
C	INVERT A MATRIX	MINV	60
C		MINV	70
C	USAGE	MINV	80
C	CALL MINV(A,N,D,L,M)	MINV	90
C		MINV	100
C	DESCRIPTION OF PARAMETERS:	MINV	110
C	A - INPUT MATRIX, DESTROYED IN COMPUTATION AND REPLACED BY	MINV	120
C	RESULTANT INVERSE.	MINV	130
C	N - ORDER OF MATRIX A	MINV	140
C	D - RESULTANT DETERMINANT	MINV	150
C	L - WORK VECTOR OF LENGTH N	MINV	160
C	M - WORK VECTOR OF LENGTH N	MINV	170
C		MINV	180
C	REMARKS	MINV	190
C	MATRIX A MUST BE A GENERAL MATRIX	MINV	200
C		MINV	210
C	SUBROUTINES AND FUNCTION SUBPROGRAMS REQUIRED	MINV	220
C	NONE	MINV	230
C		MINV	240
C	METHOD	MINV	250
C	THE STANDARD GAUSS-JORDAN METHOD IS USED. THE DETERMINANT	MINV	260
C	IS ALSO CALCULATED. A DETERMINANT OF ZERO INDICATES THAT	MINV	270
C	THE MATRIX IS SINGULAR.	MINV	280
C		MINV	290
C		MINV	300
C		MINV	310
C	SUBROUTINE MINV(A,N,D,L,M)	MINV	320
C	DIMENSION A(1),L(1),M(1)	MINV	330
C		MINV	340
C		MINV	350
C		MINV	360
C	IF A DOUBLE PRECISION VERSION OF THIS ROUTINE IS DESIRED, THE	MINV	370
C	C IN COLUMN 1 SHOULD BE REMOVED FROM THE DOUBLE PRECISION	MINV	380
C	STATEMENT WHICH FOLLOWS.	MINV	390
C		MINV	400
C	DOUBLE PRECISION A,D,BIGA,HOLD	MINV	410
C		MINV	420
C	THE C MUST ALSO BE REMOVED FROM DOUBLE PRECISION STATEMENTS	MINV	430
C	APPEARING IN OTHER ROUTINES USED IN CONJUNCTION WITH THIS	MINV	440
C	ROUTINE.	MINV	450
C		MINV	460
C	THE DOUBLE PRECISION VERSION OF THIS SUBROUTINE MUST ALSO	MINV	470
C	CONTAIN DOUBLE PRECISION PORTMAN FUNCTIONS. ABS IN STATEMENT	MINV	480
C	10 MUST BE CHANGED TO DABS.	MINV	490
C		MINV	500

POOR COPY

Appendix D

		MINV 510
		MINV 520
		MINV 530
	SEARCH FOR LARGEST ELEMENT	MINV 540
		MINV 550
	D=1.0	MINV 560
	NK=-N	MINV 570
	DO 80 K=1,N	MINV 580
	NK=NK+N	MINV 590
	L(K)=K	MINV 600
	M(K)=K	MINV 610
	KK=NK+K	MINV 620
	BIGA=A(KK)	MINV 630
	DO 20 J=K,N	MINV 640
	IZ=N*(J-1)	MINV 650
	DO 20 I=K,N	MINV 660
	IJ=IZ+I	MINV 670
	10 IF (DABS(BIGA)-DABS(A(IJ))) 15,20,20	MINV 680
	15 BIGA=A(IJ)	MINV 690
	L(K)=I	MINV 700
	M(K)=J	MINV 710
	20 CONTINUE	MINV 720
		MINV 730
	INTERCHANGE ROWS	MINV 740
		MINV 750
	J=L(K)	MINV 760
	IF (J-K) 35,35,25	MINV 770
	25 KI=K-N	MINV 780
	DO 30 I=1,N	MINV 790
	KI=KI+N	MINV 800
	HOLD=-A(KI)	MINV 810
	JI=KI-K+J	MINV 820
	A(KI)=A(JI)	MINV 830
	30 A(JI)=HOLD	MINV 840
		MINV 850
	INTERCHANGE COLUMNS	MINV 860
		MINV 870
	35 I=M(K)	MINV 880
	IF (I-K) 45,45,38	MINV 890
	38 JP=N*(I-1)	MINV 900
	DO 40 J=1,N	MINV 910
	JK=NK+J	MINV 920
	JT=JP+J	MINV 930
	HOLD=-A(JK)	MINV 940
	A(JK)=A(JT)	MINV 950
	40 A(JT)=HOLD	MINV 960
		MINV 970
	DIVIDE COLUMN BY MINUS PIVOT (VALUE OF PIVOT ELEMENT IS CONTAINED IN BIGA)	MINV 980
		MINV 990
	45 IF(BIGA) 48,46,48	MINV1000
	46 D=0.0	MINV1010
	RETURN	MINV1020
	48 DO 55 I=1,N	MINV1030
	IF (I-K) 50,55,50	MINV1040
	50 IK=NK+I	MINV1060
	A(IK)=A(IK)/(-BIGA)	MINV1070
	55 CONTINUE	MINV1080
		MINV1090
	REDUCE MATRIX	MINV1100
		MINV1110
	DO 65 I=1,N	MINV1120
	IK=NK+I	MINV1130
	HOLD=A(IK)	MINV1140
	IJ=I-N	MINV1150
	DO 65 J=1,N	MINV1160

POOR COPY.

Appendix D

```

      IJ=IJ+N
      IF (I-K) 60,65,60
60    IF (J-K) 62,65,62
62    KJ=IJ-I+K
      A(IJ)=BOLD=A(KJ)+A(IJ)
65    CONTINUE

C      DIVIDE POW BY PIVOT
C
C      KJ=K-N
      DO 75 J=1,N
      KJ=KJ+N
      IF (J-K) 70,75,70
70    A(KJ)=A(KJ)/BIGA
75    CONTINUE

C      PRODUCT OF PIVOTS
C
C      D=D*BIGA

C      REPLACE PIVOT BY RECIPROCAL
C
      A(KK)=1.0/BIGA
80    CONTINUE

C      FINAL ROW AND COLUMN INTERCHANGE
C
      K=N
100   K=(K-1)
      IF (K) 150,150,105
105   I=L(K)
      IF (I-K) 120,120,108
108   JQ=N*(K-1)
      JR=N*(I-1)
      DO 110 J=1,N
      JK=JQ+J
      HOLD=A(JK)
      JI=JR+J
      A(JK)=-A(JI)
110   A(JI)=HOLD
120   J=N(K)
      IF (J-K) 100,100,125
125   KI=K-N
      DO 130 I=1,N
      KI=KI+K
      HOLD=A(KI)
      JI=KI-K+J
      A(KI)=-A(JI)
130   A(JI)=HOLD
      GO TO 100
150   RETURN
      FND
      SUBROUTINE RANDU (IX,IY,YPL)
      IY=IX*65539
      IF (IY) 5,6,6
      IY=IY+2147483647+1
      YPL=IY
      YPL=YPL*.4656613E-9
      RETURN
      END
```

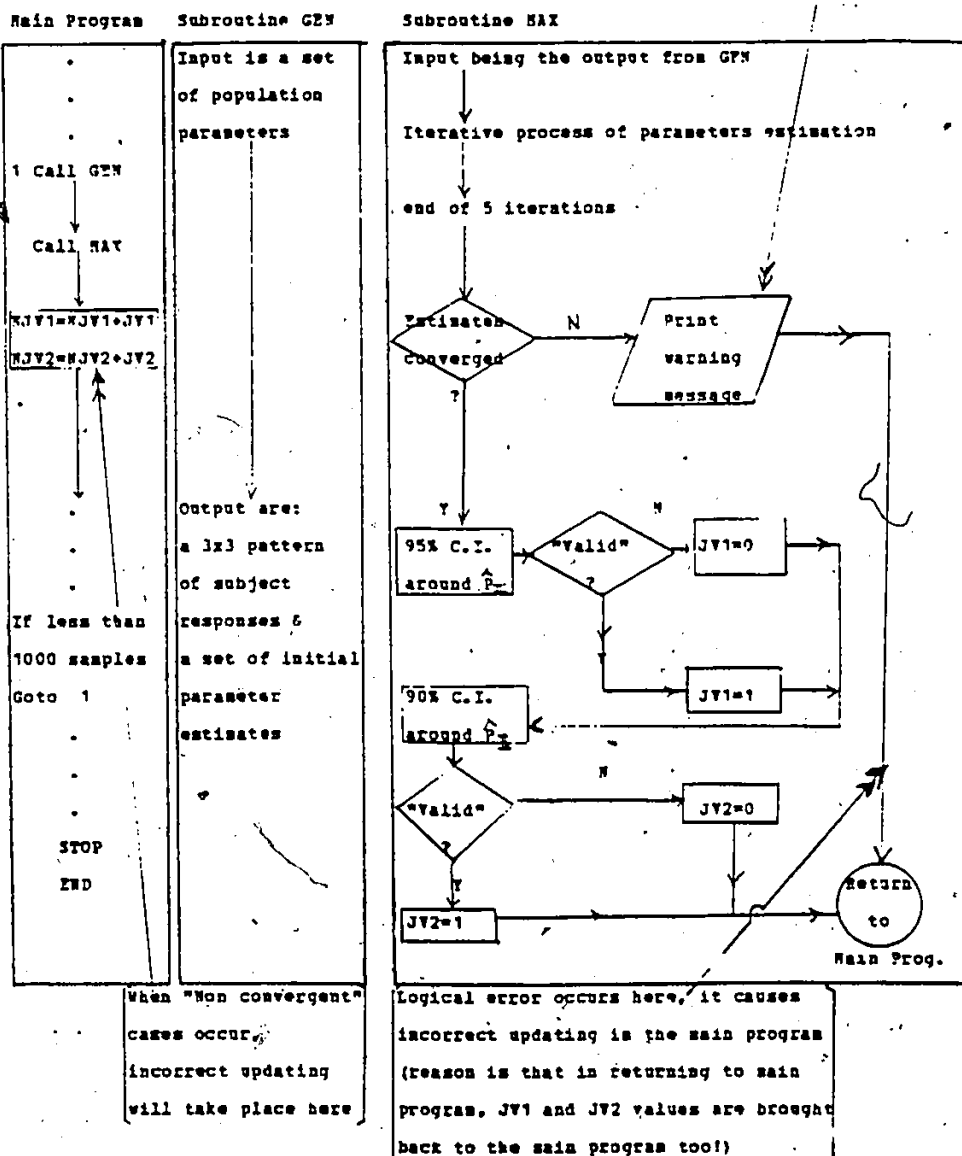
MINV1170
MINV1180
MINV1190
MINV1200
MINV1210
MINV1220
MINV1230
MINV1240
MINV1250
MINV1260
MINV1270
MINV1280
MINV1290
MINV1300
MINV1310
MINV1320
MINV1330
MINV1340
MINV1350
MINV1360
MINV1370
MINV1380
MINV1390
MINV1400
MINV1410
MINV1420
MINV1430
MINV1440
MINV1450
MINV1460
MINV1470
MINV1480
MINV1490
MINV1500
MINV1510
MINV1520
MINV1530
MINV1540
MINV1550
MINV1560
MINV1570
MINV1580
MINV1590
MINV1600
MINV1610
MINV1620
MINV1630
MINV1640
MINV1650
MINV1660
MINV1670
MINV1680

POOR COPY

Appendix E

Flowchart to show a faulty logic in Owston's Fortran IV computer program

one possibility is that Owston had removed this block in data collection



JV1, JV2 are variables to record decisions made by HAX-1 and HAX-2 on each sample.

NJV1, NJV2 are cumulative counters on "valid" indications for HAX-1 and HAX-2 respectively.

ABSTRACT

of

A Study of a Technique for
Validating Learning Hierarchies

A probabilistic technique for validating learning hierarchies was studied using a subset of an original set of hypothetical populations. This technique, called MAX, was developed by Owston. It was based upon the White & Clark Model of learning hierarchies and also was based on the maximum likelihood method of parameter estimation.

Various error types were identified in Owston's computer program in which he implemented the MAX technique. These error types were grouped into three categories. The first category is related to problems causing computer interruptions. The second category is associated with non-convergence of estimates. The third category is related to improperly converging estimates.

Two modified procedures, A and B, were developed. Procedure A was identical to the original procedure, except that the identified errors were considered. Procedure B was designed to improve upon procedure A. Ten rather than five

iterations were used, in procedure B, for the parameter estimation process; also, extra levels of confidence intervals were used.

The results obtained from using procedure A were compared with the corresponding results obtained by Owston. Differences in results were found, which cause one to question the conclusions made by Owston regarding the use of the MAX technique.

Results obtained from using procedure A were then compared with those obtained in procedure B. It was found that B was a better procedure for implementing the MAX technique.