

Loglinear Analysis: A New Tool for Educational Researchers

J. Dale Burnett
queen's university

Many situations in educational research result in the collection of data that may be conveniently represented by means of a multi-dimensional contingency table. Recent years have seen the emergence of a number of procedures that examine such structures holistically permitting the examination of various types of higher-order interaction. One such approach is that of loglinear analysis. The present article attempts to set out the basic concepts underlying this approach and to then provide a number of examples from a variety of educational research settings that illustrate the procedure.

Dans un grand nombre de recherches en éducation, les données recueillies peuvent commodément être représentées sous la forme de tableaux de contingence multidimensionnels. Au cours des dernières années, plusieurs méthodes permettant une analyse holistique de tels tableaux et l'étude de divers types d'interaction ont été développées. L'une de ces méthodes est celle de l'analyse linéaire par transformation logarithmique (loglinear analysis). Le présent article tente de faire ressortir les concepts fondamentaux sur lesquels il s'appuie et illustre la méthode en présentant de nombreux exemples tirés d'une variété de situations de recherche en éducation.

This article offers an abbreviated description of loglinear analysis and provides some examples illustrating its use. Although the purpose is straightforward, the journey is less so. Thus a number of issues, of both a general and a specific nature, are encountered and briefly highlighted.

Signpost questions include: what is the nature of the data that lends itself to a loglinear approach? What types of research questions are answerable using this technique? How does one specify a loglinear model? What are some of the special features of this approach? This latter question is approached through a series of three examples. Each example is used to illustrate a specific point.

NATURE OF THE DATA

There are three topics within this section that deserve mention: the data type, the sampling procedure and the method of data organization.

Let us deal with the method of data organization first. The key concept is that of cross-tabulation or cross-classification of frequencies. The simplest possible situation is a 2×2 table. One such example is taken from a compilation of various B.Ed. enrollment statistics (Burnett and Smith, 1977). Although there are simpler ways of examining this table than by

using a loglinear approach, it provides a useful introduction to the technique. The extension to an $I \times J$ table is easily made.

More interesting situations arise when we increase the number of dimensions beyond two. At this stage chi-square approaches begin to run into difficulty. Loglinear analysis on the other hand is ideally suited to this type of situation. An investigation by Fennema and Sherman (1977) of mathematics achievement included a $4 \times 2 \times 4$ table which will be reanalyzed using a loglinear approach. A final example, using a four-dimensional table ($3 \times 4 \times 2 \times 2$), is provided by Statistics Canada (1971) in a survey on student television viewing.

A slight detour is appropriate at this point. In preparing this article, a number of journals and books were examined for complex tables. There were fewer tables than expected. Editorial policies seem to be leaning toward the reporting of summary statistics, thereby shortening the article but concomitantly making it more difficult for the interested reader to re-examine the data for additional insights. Second, it was very difficult to locate tables with more than two dimensions. This may be naturally explained as reflecting, at least until recently, the absence of analytic procedures for examining multi-dimensional structures. Thus authors tend to collapse their tables in a manner consistent with the questions they wish to answer and with analytic techniques that permit answers to these questions. It is not always clear where the prime motivating force lies.

Even within a specific tabular structure there are nuances. Thus in some situations one may wish to distinguish between explanatory and response variables. Such a distinction usually leads the investigator into the realm of logit analysis. Baker (1981) provides a useful comparison of loglinear and logit-linear approaches.

Our second topic, within this theme concerning the nature of the data, is that of appropriate sampling procedures. These are usually given relatively light consideration in most educational research settings. Simple random sampling, usually assumed by the statistical test, is rarely carried out by the investigator, and the effects of the discrepancy remain unknown.

Fienberg (1980) in his excellent introduction to loglinear analysis provides a clear summary of the three sampling models that permit subsequent analysis via loglinear approaches. The critical question in loglinear analysis is how to obtain estimates, in this case of cell frequencies, that satisfy certain well-accepted statistical criteria such as the principle of least squares or maximum likelihood. It turns out that the loglinear procedures to be described below all provide maximum likelihood estimates when any of these three sampling procedures are used.

Thus loglinear approaches are suitable for analyzing tables where (1) the time frame for collecting data is fixed beforehand and then each subject observed during this time is classified according to the established tabular structure as falling into one of the pre-set cells, (2) the sample size

is fixed at N and then the first N subjects observed are classified, or (3) for each category of one variable (usually set up as the row variable), a set number of subjects are observed and classified according to the remaining variables. Variations from these procedures, although common in educational research settings because of various practical considerations, and although they may still give rise to various tabular structures, lead one into a no-man's-land. It may be safe, or it may contain unsuspected danger.

Our third topic concerns the basic type of data that is examined by loglinear analysis. This data consists of frequency counts or its simple derivatives: proportions or percentages. The variables used to provide the tabular structure are usually dichotomous (sex), non-ordered polytomous (school subjects) or ordered polytomous (low, middle, high SES). Although one often sees situations where a continuous variable is "simplified" into three (or five) levels, presumably to facilitate tabular displays, such approaches are not recommended here. There is an attendant loss of information accompanying such a non-linear transformation, which increases the error of measurement in a discipline where this is already a critical problem.

TYPES OF QUESTION

The types of models (to be explicated in a moment) appear to be very similar to ANOVA models. However, there is one fundamental difference between ANOVA models and loglinear models. ANOVA models are predicated on explanations based on a partitioning of variance. Loglinear models are based on the probabilities associated with the various cells in the table in question. The central issue in loglinear analysis is that of structure. To what extent does knowledge of one set of marginal totals provide probabilistic knowledge of the remaining tabular entries?

The basic goal in a loglinear approach is to specify a model that predicts cell frequencies (proportions, probabilities). The overall sequence may be described in five steps.

- (1) *Propose* a model.
- (2) *Compute* (using a computer program) expected cell values assuming the model is appropriate.
- (3) *Compare* observed cells with expected cells (chi-square or likelihood ratio chi-square goodness of fit statistic).
- (4) *Accept or reject* model.
- (5) (a) If accept, *interpret* results. (b) If reject, try another model.

A brief step ahead to discuss the goodness of fit tests will then free us to focus on the notation and the principles for constructing loglinear models. For the moment assume that we have computed a set of expected cell values (e.g. frequencies) which we wish to compare with our observed frequencies. One familiar statistic is

$$\chi^2 = \sum \frac{(O - E)^2}{E}$$

Another slightly less familiar statistic is G^2 , the likelihood ratio chi-square statistic,

$$G^2 = LR\chi^2 = 2 \sum O \cdot \ln \left(\frac{O}{E} \right)$$

Both of these statistics are distributed asymptotically as χ^2 when the sample size is sufficiently large. Colgan and Smith (1978) provide further information on appropriate sample sizes. The degrees of freedom is given by $df = (\text{no. of cells}) - (\text{no. of parameters in model})$.

LOGLINEAR MODELS

The notation, and the general procedures, are best introduced by means of a simple 2×2 case such as that illustrated by Table 1.

Let either 1 or I represent the row variable (sex), and 2 or J represent the column variable (year of enrollment). Let x_{ij} represent the observed frequency for cell i, j , m_{ij} represent the expected frequency for cell (i, j) and p_{ij} represent the cell probability (i.e. probability that a subject selected according to the sampling scheme will fall in the (i, j) cell).

One problem familiar to all educational researchers is the inconsistency of notational convention among authors. The articles and books listed in the references of this paper provide a good example of this problem. This article adopts the conventions used by Fienberg (1980) but the notation used by Baker (1981) prompts a brief comment on subscripts. If values across a specific dimension are summed, then a + symbol is used in place of the subscript. When the values are averaged, instead of summed, then the $\cdot$ is used. Thus

$$x_{i+} = \sum_{j=1}^J x_{ij}$$

$$x_{i\cdot} = \frac{1}{J} \sum_{j=1}^J x_{ij}$$

A second notational convention revolves around the distinction between population parameters and sample estimates. Social scientists tend to prefer using Greek letters for the former and Latin letters for the latter; statisticians prefer using unadorned lower-case letters for parameters and "hatted" letters for their estimates. Once again, in a manner consistent with Fienberg, the "hat" approach is used, although there is the minor inconsistency of letting a bare-headed x_{ij} represent the observed frequency. Thus x represents the raw data, the only known values (in algebra it is the opposite, let x represent the unknown – what a tangled web we weave).

Returning to our 2×2 case (e.g. Table 1), under a model (i.e., assumption) of independence, the probability of a subject falling into cell (i, j) is defined to be the product of the corresponding marginal probabilities:

$$p_{ij} = p_{i+} p_{+j}$$

Also,

$$m_{ij} = N p_{ij} = N p_{i+} p_{+j}$$

Taking logarithms,

$$\log m_{ij} = \log N + \log p_{i+} + \log p_{+j}$$

This is an example of a loglinear model, involving population parameters. It is a linear (no higher order terms) model involving the simple sum of a number of logarithmic terms. In words, the above model states that the logarithm of the expected frequency for each cell is the sum of three terms: a constant term based on the sample size, a term related to the marginal probabilities for row i , and a term related to the marginal probabilities for column j .

In order to further simplify the notation, and using the same rationale found in the treatment of linear models for ANOVA, (e.g. Glass and Stanley, p. 342) this latter equation is often written as

$$\log m_{ij} = u + u_{1(i)} + u_{2(j)}$$

where u is a constant; $u_{1(i)}$ signifies a deviation term involving the first variable – hence the 1, which has i levels or values; and $u_{2(j)}$ represents a deviation term involving the second variable which has j values.

Recall that this model corresponds to the situation where variable 1 and variable 2 (i.e. row and column variables) are assumed independent. However, there are other possible models that may also provide an adequate fit to the observed data in a given contingency table. For example:

$$\log m_{ij} = u \quad (\text{constant})$$

$$\log m_{ij} = u + u_{1(i)}$$

$$\log m_{ij} = u + u_{2(j)}$$

$$\log m_{ij} = u + u_{1(i)} + u_{2(j)} \quad (\text{independence})$$

$$\log m_{ij} = u + u_{1(i)} + u_{2(j)} + u_{1(i)u_{2(j)}} \quad (\text{full model})$$

The last model is called the full (or saturated) model. Since in the general case there are IJ cells in the matrix, and $1 + (I - 1) + (J - 1) + (I - 1)(J - 1) = IJ$ parameters, the degrees of freedom for this model are 0; the full model always provides a perfect prediction.

Moving from the specification of the model, which is in terms of

TABLE 1
B.Ed. Enrolments: Sex by Year

<i>Sex</i>	1974-75		1975-76	
	<i>f</i>	%	<i>f</i>	%
Female	404	59	442	55
Male	278	41	364	45

SOURCE: Burnett, J. D. and Smith, H. A. The B.Ed. Program: 1974-76 A Tabular Description. Research and Information Report, Faculty of Education, Queen's University, 1977 (p. 9).

population parameters, to estimation of these parameters involves the development of specific numerical procedures derived from theoretical statistical principles such as that of maximum likelihood. For the simple 2×2 case we have:

$$\hat{m}_{ij} = N \hat{p}_{ij} = N \hat{p}_{i+} \hat{p}_{+j} = \frac{1}{N} x_{i+} x_{+j}$$

and

$$\log \hat{m}_{ij} = \log x_{i+} + \log x_{+j} - \log N$$

Estimation procedures for more complex models often involve iterative procedures amenable to computer processing (cf. Fienberg, p. 37).

Example 1

This example illustrates the general procedure of model testing. For the 2×2 case (Table 1) there are three possible models of interest. It is a simple matter to test all three models and to summarize the results (Table 2) using a relatively standard format. One last comment on notation deserves mention. When presenting such tables, a shorthand method of specifying the loglinear model is used. An example should make this clear. Thus the model

$$\log m_{ij} = u + u_{1(i)} + u_{2(j)}$$

is represented by [1] [2], where the numbers in square brackets identify the factors (variables) in the model. One question remains to be answered from this table. Which model should we select? Interestingly, the question does not always result in a clear answer. (Shades of regression analysis, factor analysis!). This point will be pursued further when we examine some of the more complex tables. The usual approach, at least for two-dimensional tables, is to set up a formal goodness-of-fit null hypo-

TABLE 2
*Fitted Loglinear Models to
Table 1 Data*

<i>Model</i>	<i>LRχ^2</i>	<i>df</i>	<i>p</i>
[1]	13.26	2	0.001
[2]	30.97	2	0
[1] [2]	2.92	1	0.09

NOTE: variable 1 represents sex; variable 2 represents year.

thesis and tentatively accept any model that satisfies this hypothesis at a pre-determined level of significance (say = 0.01). Thus we keep any model that we are not 99% confident should be rejected. Using this approach with our data, only the third model satisfies the criterion. Thus we would keep the [1] [2] model, which corresponds to a model of independence between the two variables.

Before engaging in a final interpretation based on the particular model that was selected, one should examine the expected cell frequencies, and the corresponding residuals, for this model. Table 3 displays this information. An intuitive view of the expected values and the residuals indicates a reasonably good fit. The standardized residuals have a unit normal distribution under the model's assumptions and thus should be examined for extreme values (e.g., magnitude greater than 1.96) and for patterns of + and - signs which may provide additional information regarding areas where the model may be weak. In this simple case the absolute values of the standardized residuals are all less than 1.00.

Finally, having completed this analysis, what may we conclude? An examination of B.Ed. enrollment statistics for 1974-75 and 1975-76 indicates that the frequency pattern may be adequately replicated using a model that involves both main effect terms (sex, year) but which does not require an interaction term. Thus there are non-random differences in the frequencies between the two years, and between females and males, but these differences are independent of one other. In simple terms, enrollments increased and the number of females was greater than the number of males. Independence is equivalent to saying that the proportion of females remained constant (.59 versus .55).

The observant reader will have noticed a feature of the preceding analysis which sometimes intrudes into situations involving statistics. Strictly speaking, we are dealing with a population (in this case the B.Ed. classes of 74-76 at Queen's University) rather than a sample. Thus all differences in Table 1 represent real differences - they are not due to sampling fluctuations. The actual proportion of females decreased from .59 to .55 during the two-year period in question. Two comments seem

TABLE 3
Comparison of Actual and Expected Values of Table 1 Data Using [1] [2] Loglinear Model

<i>Actual</i>			<i>Expected</i>		<i>Residual</i>		<i>Standardized Residual</i>	
	1974-75	1975-76						
<i>F</i>	404	442	388	458	16	-16	.83	-.76
<i>M</i>	278	364	294	348	-16	16	-.95	.87

TABLE 4
Students in "On Grade" Math Classes (Percentage)

Grade	Sex	School 1	School 2	School 3	School 4
9	F	81	71	76	88
	M	82	69	74	77
10	F	61	51	62	40
	M	59	50	65	36
11	F	36	24	30	22
	M	43	32	34	23
12	F	12	4	11	3
	M	20	15	23	5

SOURCE: Fennema, E. and Sherman, J. Sex-related differences in mathematics achievement, spatial visualization and affective factors. *American Educational Research Journal*, 1977, 14 (1), 51-71 (p. 56).

appropriate. First, one can consider the Queen’s B.Ed. students to be a sample from the larger population of “potential” B.Ed. students. In this case the sampling is not random but highly complex, and hence we adopt a random sampling model as a simplifying assumption. Second, in an attempt to distinguish between minor and substantive differences – a question that involves understanding the situation in question – one may temporarily adopt a hypothetical “suppose my data were a random sample from some larger population” to see if the observed differences are at least statistically significant. This information in turn may help the investigator decide if the differences are practically significant.

Example 2

The Fennema and Sherman (1977) data on mathematics achievement (Table 4) provides an excellent vehicle for describing one basic restriction

TABLE 5
*Selected Fitted Loglinear Models to Table 4 Data (Students
 in "On Grade" Math Classes)*

<i>Model</i>	<i>LRχ^2</i>	<i>df</i>	<i>p</i>
[1][2][3]	45.91	24	0.00
[1][2][3][12][13][23]	1.84	9	.99
backward selection			
[1][2][3] [12][13]	2.83	12	.99
[12] [23]	31.15	18	0.03
[13][23]	14.23	12	0.29
[1] [3] [13]	16.80	16	0.40
[1][2][3] [12][13]	2.83	12	0.99

NOTE: variable 1 represents grade; variable 2 represents sex; variable 3 represents school.

of loglinear models as well as illustrating one of the recommended procedures for selecting a model. The basic restriction is known as the hierarchy principle, which states that higher-order terms may be included only if the related lower-order terms are included. Thus requesting the term [123] also implies including [1], [2], [3], [12], [13] and [23]. There are both technical and interpretive reasons for this principle. The technical reasons involve problems of estimation using iterative programming procedures. Fienberg (p. 43) prefers to emphasize the interpretive reason, pointing out that if the interaction terms are significant, it is of little interest that the main effects may be zero.

Another issue confronting the individual conducting a loglinear analysis involves the final selection of a model. There is no commonly agreed upon procedure for comparing models (a situation similar to stepwise regression), although the procedure summarized in Table 5 is often used when the data becomes more complex, involving four or more dimensions. It is more easily shown with three dimensions. Since [1][2][3] is not an adequate model and since the model containing all two-way interactions provides a statistically significant fit, one is interested in models between these two. One may either proceed by beginning with the main effects model and adding interaction terms, one at a time, to see which is the best; or one may start with the full two-way interactions model and delete terms one at a time. Using the backward selection procedure the [1][2][3][12][13] model was identified.

This example also illustrates the holistic, simultaneous nature of a loglinear approach. This is the most noteworthy feature of the procedure and may be contrasted with a more common approach which would examine the three-dimensional matrix as a set of three independent

TABLE 6
*Comparison of Actual and Expected Values of Table 4 Data Using
[1] [2] [3] [12] [13]*

Actual				Expected				Standardized Residuals			
81	71	76	88	83	72	77	84	-0.26	-0.07	-0.08	0.40
82	69	74	77	80	68	73	81	0.26	0.07	0.08	-0.40
61	51	62	40	61	51	64	38	0.06	0.00	-0.26	0.27
59	50	65	36	59	50	63	38	-0.06	0.00	0.26	-0.27
36	24	30	22	36	26	29	21	-0.04	-0.34	0.11	0.30
43	32	34	23	43	30	35	24	0.04	0.31	-0.11	-0.27
12	4	11	3	10	6	11	3	0.52	-0.86	0.01	0.26
20	15	23	5	22	13	23	5	-0.36	0.59	-0.01	-0.18

TABLE 7
*Students in "On Grade" Math
Classes (Collapse of Table 4
Across Schools) (Percentage)*

Grade	Sex	
	F	M
9	79	75.5
10	53.5	52
11	28	33
12	7.5	15.8

two-dimensional tables with corresponding chi-square tests of independence. Finally, this example may be used to illustrate the manner in which loglinear analysis facilitates the “collapsing” of tables. The operative principle (Fienberg, p. 51) is that it is permissible to collapse a table if the variable collapsed over is independent of at least one of the other two variables. Clearly, any collapsing results in a loss of detail, and should be approached with caution. However, the above principle ensures that the interaction between the remaining two terms will be unaffected by collapsing under this condition. And sometimes the resulting table permits one to see this interaction pattern more clearly.

Consideration of Table 5 indicates that Table 4 may be collapsed over either school or sex. Since specific schools are usually of little interest, at least in research where one is looking for general patterns, one may decide to collapse over schools. The result is shown in Table 7. It is an easy matter to construct a graph corresponding to Table 7 and to note the

disordinal interaction between grade level and sex. It is also interesting to compare the results of Table 5 with the results described in the Fennema article of conducting an ANOVA on math achievement errors. One obtains exactly the same patterns, three main effects plus two interaction terms. It is intriguing to speculate on the degree to which such findings would replicate, permitting researchers to focus on enrollment patterns without having to administer a battery of additional tests.

Example 3

This example (Table 8) of a four-dimensional table ($3 \times 4 \times 2 \times 2$) is taken from a Statistics Canada publication. It is difficult to find such examples in the literature, in part because the subsequent analysis is unlikely to reflect totally such a high-order structure. Another reason is that such studies require relatively large sample sizes in order to provide some stability to the individual cell frequencies.

Reynolds (1977), in a section before he treats loglinear approaches, discusses the problem of skewed marginal distributions and suggests that one way to reduce the effect of low-frequency components is to rescale the values along the relevant dimension from frequencies to percentages. Clearly such a non-linear adjustment may result in an alteration to the underlying tabular structure. It seems prudent in such situations to conduct analyses on both sets of data (frequencies and percentages) and compare the results in order to determine the effect of using the percentages. This was done for the television viewing data of Table 8. The fitted loglinear models are shown in Table 9 (frequencies) and Table 10 (percentages). The results are quite striking. The percentages result in a much simpler loglinear model involving only two 2-way interaction terms. Since the frequencies for all four cities were adjusted to values totalling 100, this had the effect of removing the cities term from the model. The two interaction terms are $[13]$, grade by TV viewing, and $[34]$, SES by TV viewing. Returning to the original table, we notice the nature of the interaction. The higher grade students watch less television than do the lower grade students; the high SES group tend to be in the low TV viewing group whereas the low SES group is more evenly divided between high and low TV viewing. However, when one analyzes the raw frequencies, which vary considerably across the four cities, a very complex structure results. The only three-way interaction NOT in the model is $[234]$. It is difficult to adequately express in words the final loglinear model $[1][2][3][4][12][13][14][23][24][34][123][124][134]$ except to say that it is very complex. At least one should be warned to avoid overly simplistic generalizations about this data. Only two standardized residuals exceeded 1.96 for the percentage case and only one exceeded this value for the frequencies data.

TABLE 8
Television Viewing by Socio-Economic Status, Grade and City, 1971 (Percentage)

	Socio-Economic Status							
	Vancouver		Winnipeg		Toronto		Ottawa	
	High SES	Low SES	High SES	Low SES	High SES	Low SES	High SES	Low SES
<i>Grade 7</i>								
Low TV	60 (26)	51 (34)	59 (78)	50 (95)	67 (39)	48 (91)	55 (109)	21 (22)
High TV	40 (17)	49 (33)	41 (54)	50 (95)	33 (19)	52 (99)	45 (90)	79 (84)
<i>Grade 9</i>								
Low TV	62 (21)	56 (216)	78 (166)	52 (131)	78 (124)	44 (58)	75 (200)	51 (122)
High TV	38 (13)	44 (169)	22 (47)	48 (119)	22 (35)	56 (75)	25 (66)	49 (115)
<i>Grade 11</i>								
Low TV	86 (91)	74 (193)	81 (138)	68 (77)	80 (121)	75 (111)	83 (184)	73 (171)
High TV	14 (15)	26 (63)	19 (32)	32 (37)	20 (30)	25 (38)	17 (38)	27 (63)

SOURCE: Statistics Canada Catalogue 81-561. Socio-Cultural Characteristics of Elementary and Secondary Students 1971, Table 3, p. 75.
NOTE: *N* in parentheses.

TABLE 10
Selected Fitted Loglinear Models to Table 8 Data (Student Television Viewing) (Percentages)

<i>Model</i>		<i>LRχ^2</i>	<i>df</i>	<i>p</i>
[1][2][3][4]		247.63	40	0.00
/	[12][13][14][23][24][34]	41.97	23	0.01
/	/	13.02	6	0.04
forward selection				
[1][2][3][4]	[12]	247.63	34	0.00
/	[13]	125.01	38	0.00
/	[14]	247.63	38	0.00
/	[23]	242.22	37	0.00
/	[24]	247.63	37	0.00
/	[34]	174.40	39	0.00
/	[13]	125.01	38	0.00
/	[13][34]	51.77	37	0.05
backward selection				
[1][2][3][4]	[13][14]	51.77	40	0.10

SOURCE: variable 1 represents grade; variable 2 represents city; variable 3 represents TV viewing; variable 4 represents SES.

CONCLUSION

This article has attempted to provide an introduction to loglinear analysis. It has illustrated a number of major issues surrounding the technique through a series of three examples. The reader may sense that some issues, such as stepwise model selection, require further refinement. Other issues, such as the testing of cell outliers (cells whose standardized residual exceed 1.96) by deleting them from the table and reanalyzing, and the effect of small and large sample sizes, have not been discussed. Further information on these issues may be found in an article by Colgan and Smith (1978). Fienberg (1980) provides an excellent introduction to loglinear analysis, as do Everitt (1977) and Upton (1978), and Bishop, Fienberg and Holland (1975) have produced a definitive handbook on loglinear approaches. For the technically inclined, an extensive bibliography on the analysis of contingency tables for the period up to 1974 has been prepared by Killion and Zahn (1976).

Also, there is no implication that one should use a particular computer program. The analyses for this article were conducted using an APL program called CONTABLE which is available in many APL libraries. Bock's MULTIQUAL program is well known although it utilizes a fairly sophisticated approach. Another computer package which imbeds log-linear modeling within it is the GLIM (Generalised Linear Interactive Modeling) system developed at Rothamsted Experimental Station in England. The BMD-P series also includes a program based on a loglinear model. Interested readers should contact their computing centre or statistical advisor for the programs available at their location.

Finally, the article does not suggest that all 3-dimensional, or higher, tables be analyzed via loglinear models. Taylor and Chappell (1980) suggest a partitioning of variance approach. Bishop, Fienberg and Holland devote an entire chapter of their book to other approaches such as information-theory, partitioning chi-square and logit analysis. Baker (1981) also provides further information on logistic models.

The article was written with the goal of increasing the probability that educational researchers will include loglinear analysis in their toolbox of potential data analytic techniques. This goal could also be approached by incorporating loglinear analysis into graduate level educational research courses. With the increased use of such techniques an understanding of the basic principles of a loglinear approach will become increasingly necessary for the intelligent reading and interpretation of the research literature.

REFERENCES

- Baker, F. B. Log-linear, logit-linear models: A didactic. *Journal of Educational Statistics*, 1981, 6(1), 75-102.

- Baker, R. J., & Nelder, J. A. *The GLIM System, Release 3, generalised linear interactive modelling*. Oxford, England: The Numerical Algorithms Group, 1978.
- Bishop, Y. M. M., Fienberg, S. E., & Holland, P. W. *Discrete multivariate analysis: Theory and practice*. Cambridge, Mass.: MIT Press, 1975.
- Bock, R. D. *Multivariate statistical methods in behavioral research*. New York: McGraw-Hill, 1975. (Chapter 8).
- Bock, R. D., & Yates, G. *MULTIQUAL: Log-linear analysis of nominal or ordinal qualitative data by the method of maximum likelihood*. Chicago: National Educational Resources, Inc., 1973.
- Burnett, J. D., & Smith, H. A. The B.Ed. Program: 1974-76 A Tabular Description. Research and Information Report, Faculty of Education, Queen's University, 1977.
- Colgan, P. W., & Smith, J. T. Multidimensional contingency table analysis. In P. W. Colgan (Ed.), *Quantitative ethology*. New York: John Wiley, 1978.
- Dixon, W. J., & Brown, M. B. *BMD/P-79: Biomedical computer programs P-series*. Berkeley: University of California Press, 1979.
- Everitt, B. S. *The analysis of contingency tables*. London: Chapman and Hall, 1977.
- Fennema, E., & Sherman, J. Sex-related differences in mathematics achievement, spatial visualization and affective factors. *American Educational Research Journal*, 1977, 14(1) 51-71.
- Fienberg, S.E. *The Analysis of cross-classified categorical data* (2nd ed.). Cambridge, Mass.: MIT Press, 1980.
- Glass, G. V., & Stanley, J. C. *Statistical methods in education and psychology*. Englewood Cliffs, N.J.: Prentice-Hall, 1970.
- Goodman, L. A. *Analyzing qualitative/categorical data*. Cambridge, Mass.: Abt Books, 1978.
- Haberman, S. *Analysis of qualitative data*. New York: John Wiley, 1978.
- Killion, R. A., & Zahn, D. A. A bibliography of contingency table literature: 1900 to 1974. *International Statistical Review*, 1976, 44(1), 71-112.
- Plackett, R. L. *The analysis of categorical data*. London: Griffin, 1974.
- Reynolds, H. T. *Analysis of nominal data*. Beverly Hills: SAGE Pub., 1977.
- Smith, J. E. K. Analysis of qualitative data. *Annual Review of Psychology*, 1976, 27, 487-499.
- Statistics Canada Catalogue 81-561. Socio-Cultural Characteristics of Elementary and Secondary Students, 1971.
- Taylor, K. W., & Chappell, N. L. Multivariate analysis of qualitative data. *The Canadian Review of Sociology and Anthropology*, 1980, 17(2) 93-108.
- Upton, G. J. G. *The analysis of cross-tabulated data*. Chichester: John Wiley, 1978.

J. Dale Burnett is an Associate Professor in Educational Psychology, Faculty of Education, Queen's University, Kingston, Ontario K7L 3N6.