

TABLE OF CONTENTS

<u>Title</u>	<u>Page</u>
DECOMPOSING THE COEFFICIENT OF DETERMINATION IN MULTIPLE REGRESSION	1
Lee M. Wolfle, Virginia Polytechnic Institute and State University	
CAUSAL MODELS WITH UNMEASURED VARIABLES: AN INTRODUCTION TO LISREL	9
Lee M. Wolfle, Virginia Polytechnic Institute and State University	
TIME TO SECURE EMPLOYMENT ANALYZED BY SETWISE REGRESSION ANALYSIS	55
Jole A. Williams, The Grafton State School and John D. Williams, University of North Dakota	

DECOMPOSING THE COEFFICIENT OF DETERMINATION IN MULTIPLE REGRESSION

LEE M. WOLFLE

In multiple regression, the coefficient of determination (or R-square) has a very useful interpretation. The statistic is the ratio of the variation that is explained by the regression equation to the total variation of the dependent variable. For example, a coefficient of determination equal to .45 indicates that the independent variables can explain 45 percent of the variation in the dependent variable.

It follows immediately that a person's next instinct is to want to allocate among the several independent variables the explained variation in the dependent variable. For example, many people would like to say that if the regression of a dependent variable on three independent variables explains 45 percent of the variance, that (say) 25 percent was due to the first independent variable, 15 percent to the second, and 5 percent to the third. While this interpretation is tempting, it should be avoided. The reason for avoiding it is that there is no unique way of decomposing the explained variance, and if there is no unique way of doing so, then there is no meaningful way of doing so.

For example, consider the regression of one dependent variable, Y, on two independent variables:

$$Y = a + bX + cZ + u,$$

where Y is the dependent variable, X and Z are independent variables, "a" is the intercept, "b" and "c" are partial regression coefficients, and "u" is the disturbance or error term. From this regression, one would obtain the coefficient of determination, $R^2_{y.xz}$, which is merely a convenient notation for R-square of Y regressed on X and Z. By the method of part correlations, it may be shown that:

$$R^2_{y.xz} = R^2_{y.x} + R^2_{y(z.x)}$$

where $R^2_{y.x}$ is the squared zero-order correlation of Y and X, and $R^2_{y(z.x)}$ is the squared part correlation of Y with Z residualized for X. (This equation is analogous to formula 5.10 in Kerlinger and Pedhazur, 1973.)

But it is also true that:

$$R^2_{y.xz} = R^2_{y.z} + R^2_{y(x.z)},$$

and in general $R^2_{y.x}$ does not equal $R^2_{y(x.z)}$. If these two quantities are not equal, by which quantity therefore does one measure the unique contribution of X to the explained variation in Y? Because there are two answers, two different answers, there is no unique solution.

These quantities may appear mysterious in symbolic form, but they are familiar quantities, which appear in the SPSS regression output. They appear in the summary table in a column of numbers entitled, "R-Square Change." People often want to interpret these quantities as measuring the amount of variance explained by each independent variable. This temptation should be avoided.

Suppose, for example, that one regresses educational attainment on two independent variables, father's education and father's occupational

status. Suppose the coefficient of determination for this equation is equal to .31. If father's education was added to the regression as the first independent variable, then one would learn from the SPSS summary table that the R-Square Change for father's education was .27 and for father's occupation was .04. If, however, father's occupation were to have been listed first, then one would learn that the R-Square Change for father's occupation was .23 and for father's education was .08. Question: Does father's education explain 27 percent of variation in son's education, or does it explain 8 percent? The answer is, "Yes, it does." That is, without a unique way to decompose the explained variance, there is no unique answer to the question.

Let us try another approach. It is well known (e.g., Kerlinger and Pedhazur, 1973, formula 4.17) that the coefficient of determination may be decomposed into the sum of the products of the zero-order correlations and their associated beta-weights. One may, therefore, be tempted to interpret the product of say R_{yx} times $B_{yx.z}$ (where R_{yx} is the zero-order correlation of Y and X, and $B_{yx.z}$ is the beta-weight of Y regressed on X controlling for Z) as the amount of variance in Y explained by X. The problem with this approach, however, is that the zero-order correlation and the beta-weight are not constrained to have the same sign. In such cases, one would have to interpret the product as being a negative component to the explained variance, which is a very troublesome concept.

For example, consider the regression of son's occupational status on his educational attainment, his father's education, and his father's

occupational status. The correlations of son's occupation with these three independent variables are in one sample, respectively, .47, .38, and .73. The corresponding beta-weights are, respectively, .21, -.11, and .69. The coefficient of determination may be decomposed:

$$R^2 = .56 = (.21)(.47) + (-.11)(.38) + (.69)(.73) \text{ or}$$

$$R^2 = .56 = .10 - .04 + .50.$$

It is with no relish whatsoever that one should interpret the amount of variance explained by father's education as being minus 4 percent. That is equivalent to saying that the addition of father's education to the regression equation takes away four percent of the variance in son's occupational status. Not only does that interpretation not make any sense in a substantive way, it is mathematically impossible.

Finally, consider the decomposition:

$$R^2_{Y.XZ} = B^2_{Y.X} + B^2_{Y.Z} + 2(B_{Y.X})(B_{Y.Z})(R_{XZ}),$$

In which $B^2_{Y.X}$ is the square of the beta-weight of Y regressed on X controlling for Z (and analogously for $B^2_{Y.Z}$), $B_{Y.X}$ is the beta-weight of Y regressed on X, and R_{XZ} is the zero-order correlation between X and Z. This decomposition seemingly contains a portion (the squared beta-weight) that can be uniquely attributed to the independent variable, but the decomposition also contains an explicit term (or more than one term if there are more than two independent variables) representing the contribution to the explained variance in Y that is shared by both independent variables. What this decomposition indicates is that the explained variance in Y cannot be decomposed into unique separate components due to each independent variable (unless $R_{XZ} = 0$, a very rare occurrence).

If the coefficient of determination cannot be uniquely partitioned into amounts of variance explained by each independent variable, how then does one measure the contribution of each independent variable to the dependent variable? The solution would appear to be to use either:

- (1) the metric partial regression coefficients; or
- (2) the standardized partial regression coefficients,
or beta-weights.

Notice that neither of these are interpretable as components of explained variance. Metric partial regression coefficients are to be interpreted as the amount Y changes for a one-unit increase in one independent variable while the other independent variables are held constant. Standardized partial regression coefficients are interpretable as the number of standard deviations Y changes for a one standard deviation increase in one independent variable while the other independent variables are held constant.

The standardized regression coefficients have the advantage of being standardized. That is, the size of the metric regression coefficients depend upon the metric in which the independent variables have been measured. If one of the independent variables is income, for example, the metric regression coefficients will be different if income is measured in increments of thousand dollars versus increments of single dollars. In any event, these coefficients will be different from those of another independent variable measured in, say, years of schooling. Standardized coefficients get around this problem by measuring all the variables in standard deviation units. Thus, the standardized

coefficients are comparable among independent variables. A beta-weight of .5 for one independent variable means that Y changes twice as much as it does when another variable, which has a beta-weight of .25, changes.

Seemingly, therefore, the standardized regression coefficient is better than the metric coefficient. But wait! The standardized coefficient has to be standardized in terms of something, and that "something" turns out to be a quantity which is not invariant across either samples or populations. I am referring to the ratio of the standard deviations of the independent to the dependent variables. That is,

$$\text{BETA} = B (S_x/S_y),$$

where BETA is a beta-weight and B is the corresponding metric regression coefficient. If, for example, one is interested in comparing the effects of one independent variable on a dependent variable, and wants to compare the size of this effect across two populations (e.g., freshman versus sophomores, men versus women, blacks versus whites, etc.) then the beta-weights can change as a function of a change in the ratio of standard deviations; even while the structural coefficients, the metric coefficients, remain constant across populations.

Therefore, in reporting regression results one should always report both the standardized and the metric coefficients. The former are useful in comparing the relative effects of independent variables within a sample or population, while the latter are useful for comparing the relative effects of independent variables across samples or populations.

In any event, the amount of variance explained by each independent variable is not a quantity that can be uniquely estimated; the use of such estimates is to be discouraged.

(Kerlinger and Pedhazur, 1973, pp. 297-305, discuss a method called "Commonality Analysis," which can sometimes be used to estimate amounts of variance explained by each independent variable. In essence, the approach measures the portion of explained variance for a single independent variable as that portion unexplained by all of the other independent variables. The method results in measures of unique contributions and common contributions. In most real-life cases, the common portions far outweigh the unique portions. Another problem is the proliferation of higher-order commonalities. With five independent variables, commonality analysis produces five unique components, and 26 common components. In my view, metric and standardized regression coefficients are to be preferred in reporting the results of regression analyses.)

REFERENCES

Kerlinger, Fred N., and Pedhazur, Elazar J. Multiple Regression In Behavioral Research. New York: Holt, Rinehart and Winston, 1973.

CAUSAL MODELS WITH UNMEASURED VARIABLES AN INTRODUCTION TO LISREL

LEE M. WOLFLE

ABSTRACT

Recent methodological advances (Jöreskog and Sörbom, 1978) now allow the estimation of causal models which incorporate structural relationships among latent variables, and confirmatory factor procedures to estimate latent variables and the measurement properties of their manifest indicators. The examination of structural effects among theoretical variables not directly measurable offers great promise for developing and testing theory in educational research. This paper presents an introduction to such models.

(This paper won the outstanding award at the 1981 annual meeting of the American Educational Research Association Division D. The paper was presented at the MULTIPLE LINEAR REGRESSION SPECIAL INTEREST GROUP MEETING and, according to policy, is being published without review.)

CAUSAL MODELS WITH UNMEASURED VARIABLES:

AN INTRODUCTION TO LISREL ¹

Whenever one uses ordinary least squares (OLS) regression, one is making an implicit assumption about measurement error. As Blalock (1964, p. 49) noted, one assumes that "there may be errors of measurement with respect to the dependent variable Y, but that all of the independent variables have been measured without error." Such an assumption is obviously unrealistic for most social data. Until recently, one had three alternative methods for estimating such regression models. By far the most common was to naively assume that the variables were measured without error, and wistfully hope the resulting estimates were robust. A second alternative was to correct correlation coefficients for attenuation, and use the corrected estimates as inputs to the regression analysis. The procedure, however, required a priori knowledge of the reliability coefficients for the variables; furthermore, one had to assume the reliabilities were invariant from one application to another. These restrictions have severely limited the use of regression analyses based on correlations corrected for attenuation. Yet a third alternative was to measure implied coefficients between latent variables for which one had multiple manifest indicators. Siegel and Hodge (1968), for example, explicated several such models in their paper directed to sociologists; furthermore, they noted that correlations corrected for attenuation were merely special cases of their multiple indicator models. (One should note that the sociologists' preoccupation

with path analysis and causal models as approaches to measurement error were anticipated by Sewall Wright nearly sixty years ago. For a summary and appreciation of Wright's seminal work in structural modeling, see Goldberger [1972].)

The problem with the third alternative, as noted by Hauser and Goldberger (1971) and Long (1976), is its casual approach toward statistical estimation and hypothesis testing. The problem results from overidentified models, which yield multiple estimates of the associations among latent variables. In response, some authors have chosen to ignore one or more of the identifying equations (e.g., Blalock, 1970; Land, 1970); others have averaged the estimates from the several equations (e.g., Hauser, 1970). A better alternative would be to obtain estimates of the overidentified parameters by maximum likelihood estimation (MLE). These procedures grew out of the work of Lawley (1943), but the immense computational load required for their iterative estimation prevented their application in practice. Thus, the application of more adequate statistical procedures languished until Jöreskog (1966, 1967, 1969) discovered an efficient MLE computational procedure, soon to be followed by a computer program for confirmatory factor analysis (Jöreskog, Grunvæus, and van Thillo, 1970). The resulting variances and covariances of the latent factors could be used to estimate the parameters of a structural model assumed to exist among the factors, and Jöreskog and Sörbom (1978) have provided a program which incorporates maximum likelihood estimation procedures for both the confirmatory factor analysis measurement model, and the linear structural model among the factors. This program is called LISREL, an acronym for linear structural relationships, and possesses the

potential for revolutionizing the way social scientists test hypothesized relationships among theoretical, unmeasured latent variables. Yet I hasten to caution both users of LISREL, and those who would uncritically accept LISREL solutions; strong models, which yield useful information, result from researchers who know their subject matter, and who postulate disprovable hypotheses (Platt, 1964). As Cooley (1978) reminded us, sometimes a simple contingency table analysis provides better causal estimates than more complex estimation procedures, when the researcher using the contingency table uses the right variables, and thereby controls for alternative explanations for observed associations. LISREL is, therefore, like most estimation procedures, an aid to thinking about one's subject matter. It is not a substitute. Readers (and editors!) should not be lulled into the belief that LISREL solutions are ipso facto good research. Duncan (1975) drew the useful distinction between the easy part of causal modeling (the estimation of causal parameters, and their algebraic manipulation) and the hard part (knowing one's subject matter, and having a stylish appreciation of alternative explanations). Solving the easy problems should not mislead anyone to believe that the hard problems have been resolved. This paper will become difficult enough both symbolically and algebraically, but does not even approach the complexity involved in doing good research.

The purpose of this paper is to provide a nonmathematical introduction to LISREL. Those interested in the basic papers that develop the mathematics of LISREL may read Jöreskog and Sörbom (1979). Those interested in reading well-founded applications of LISREL (or earlier versions of the program) may read Mason, et al. (1976) and Mare and Mason (1980),

who analyzed the errors involved in children's reports of parental characteristics; Bielby, et al. (1977), who analyzed a LISREL model of status attainment and compared the resulting estimates with OLS estimates; and Werts, et al. (1977), who develop a simplex model of academic achievement over time. For those interested in alternative introductions to the analysis of covariance structures, see Burt (1973) or Long (1976).

SPECIFICATION OF A LISREL MODEL

In presenting any causal model, it is convenient to use a path diagram. One such model is shown in Figure 1. In Figure 1, variables enclosed in ellipses are latent, unobserved variables; variables enclosed in rectangles are manifest, observed variables; the unenclosed variables represent errors of measurement and residual errors of prediction. A one-way arrow represents a hypothesized causal effect; the arrow points toward the affected variable. A curved, double-headed arrow represents a correlation to which no causal interpretation is attached. Dashed lines represent associations added to the model in subsequent analyses; these will be explained below.

Before proceeding to discuss the model in Figure 1, let me explain why the discussion is couched in matrix notation represented by Greek letters. These merely follow the tradition established by Jöreskog; thus, learning the notation once pays off when reading any of the basic papers on confirmatory factor analysis. Moreover, setting up the LISREL computer program involves specifying the type and size of several matrices, and involves specifying whether the elements within these matrices are fixed at prespecified values, or are free parameters to be estimated by

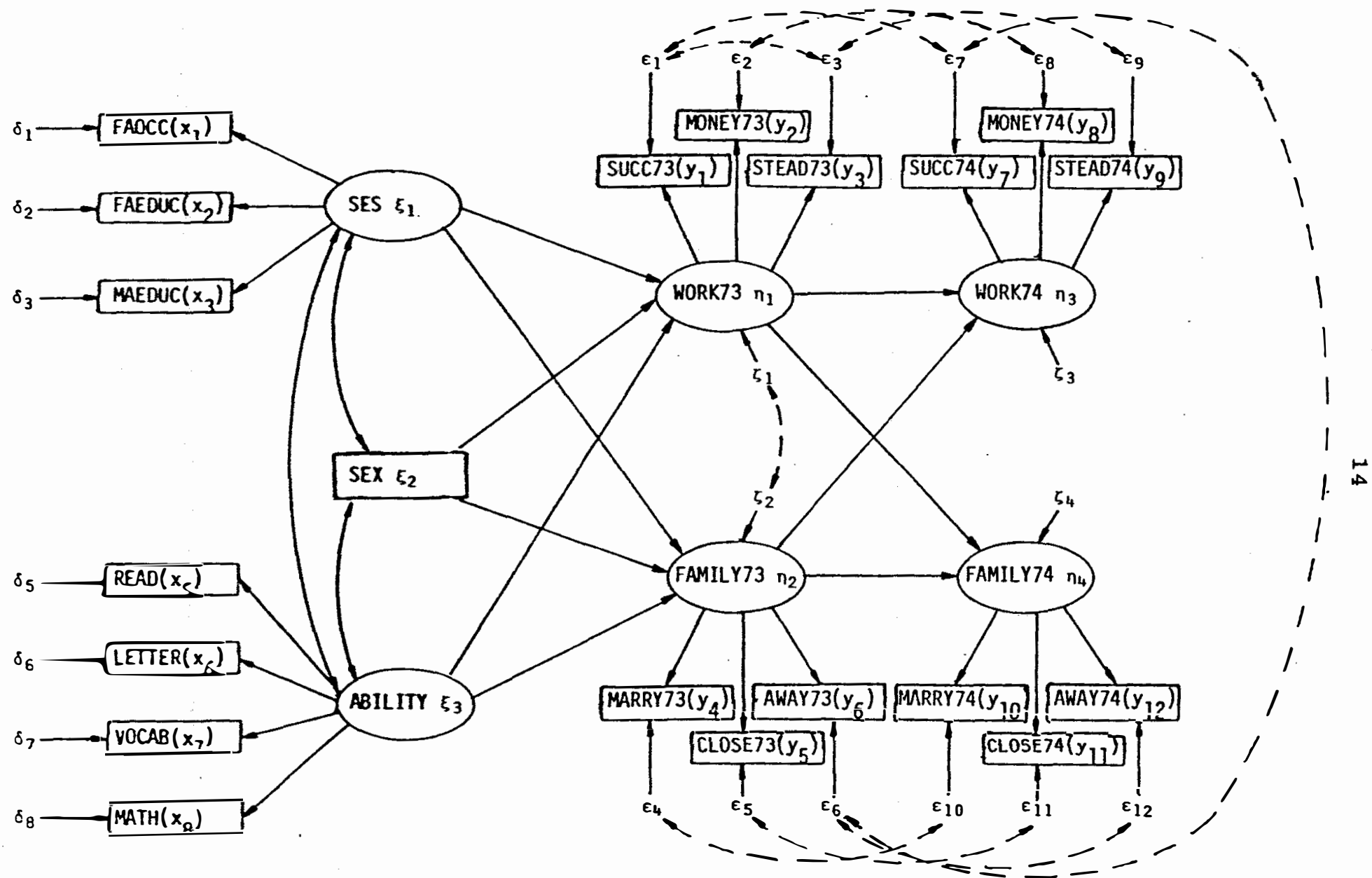


Figure 1. LISREL Model of the Development and Stability of Attitudes Toward Work and Family; White High School Graduates, 1972 (N = 751).

the program. The program's notation for these matrices is representative of their Greek names; for example, the Λ_x (lambda-x) matrix, which relates factors to manifest variables, is referred to as LX in the LISREL program. Thus, a little effort invested in learning the notation system will aid in reading the LISREL literature, and in using the LISREL program.

There are eight matrices in the LISREL model, and a listing of them will help to understand the discussion to follow. One basic distinction to keep in mind is between exogenous variables, whose causes are unanalyzed in the model, and endogenous variables, whose causes are included in the model. Another distinction is made between those matrices that define the measurement portion of the model, and those matrices that define the structural portion of the model. The measurement model is defined by two regression matrices, and two variance-covariance matrices among errors of measurement:

- 1) Λ_y (lambda-y, or LY) is a p by m regression matrix, which relates the m endogenous factors to each of the p endogenous manifest variables;
- 2) Λ_x (lambda-x, or LX) is a q by n regression matrix, which relates the n exogenous factors to each of the q exogenous manifest variables;
- 3) Θ_e (theta-epsilon, or TE) is a symmetrical p by p variance-covariance matrix among the errors of measurement for the p endogenous manifest variables; when one assumes these errors of measurement are uncorrelated, Θ_e may be specified to be a diagonal matrix pf order p;

- 4) $\underline{\theta}_\delta$ (theta-delta, or TD) is a symmetrical q by q variance-covariance matrix among the errors of measurement for the exogenous manifest variables; when one assumes these errors are uncorrelated, $\underline{\theta}_\delta$ may be specified to be a diagonal matrix of order q .

The structural model is defined by two additional regression matrices, and two variance-covariance matrices, one among the exogenous factors, one among the residual errors of prediction:

- 5) $\underline{\gamma}$ (gamma, or GA) is an m by n regression matrix, which relates the n exogenous factors to the m endogenous factors;
- 6) $\underline{\beta}$ (beta, or BE) is an m by m matrix which contains ones on the diagonal, and on the off-diagonal contains negative regression coefficients which relate the m endogenous factors to each other;
- 7) $\underline{\phi}$ (phi, or PH) is an n by n symmetrical variance-covariance matrix among the n exogenous factors; and
- 8) $\underline{\psi}$ (psi, or PS) is an m by m symmetrical variance-covariance matrix among the m residual errors of prediction for the m endogenous factors.

In addition, η_i (eta) denotes the $i = 1 \dots m$ endogenous factors, each of which has a ζ_i (zeta) residual error of prediction; ξ_j (xi or ksi) denotes the $j = 1 \dots n$ exogenous factors. Note that the ξ_j exogenous factors may be standardized by inputting correlations and specifying the diagonal of $\underline{\phi}$ to be ones. The LISREL model does not, however, permit a priori specification of the variance-covariance matrix among the η_i .

endogenous factors. This is not a shortcoming of the LISREL model, but is inherent to the model's mathematical logic. Thus, LISREL estimates are unstandardized even when analyzing a correlation matrix; standardized solutions are available in LISREL by adjusting the unstandardized estimates by estimated standard deviations for the ξ and η factors. The exact specification of these solutions is given in Jöreskog and Sörbom (1978, p. 60).

THE STRUCTURAL MODEL

With these matrices in mind, we may again consider the model depicted in Figure 1, and then express the relationships in matrix terms. Figure 1 represents a longitudinal model of the development and stability of attitudes toward work and family. Respondents in the National Longitudinal Study of the High School Class of 1972 (see Levinsohn, et al., 1978) were asked "How important is each of the following to you in your life?"

- A. Being successful in my line of work,
- B. Finding the right person to marry and having a happy family,
- C. Having lots of money,
- E. Being able to find steady work,
- H. Living close to parents and relatives,
- I. Getting away from this area of the country.

Items A, C, and E were used as indices of attitudes toward work, while items B, H, and (the additive inverse of) I were used as indices of attitudes toward the family. (Items omitted from this list dealt with

measures of community orientation.) The respondents were asked to indicate whether these statements were either not important, somewhat important, or very important. The items were thus scaled or rescaled such that higher numbers reflected more importance; but item I was scaled to its additive inverse. There is little reason to be concerned about the ordinal measurement of the manifest indicators; they are assumed to be imperfect indicators of underlying interval-level scales. It is granted without question that the ordinal measures are not exactly isomorphic with the underlying latent variable, assumed to be normally distributed, and therefore efficiency is lost to the extent that there is a lack of correspondence. This may be thought of as a form of measurement error, and in no way suggests that interval-level statistics are inappropriately applied to these data (see Borgatta and Bohrnstedt, 1980).²

These attitudes were measured twice (in this application in 1973 and 1974), and the model specifies at time 1 that latent attitudes about work and family are caused by the respondent's ability, sex, and the socioeconomic status of their family of origin. At time 2, work and family attitudes are specified to be caused by the previous expression of these attitudes. Thus, socioeconomic status, sex, and ability are assumed to cause work and family attitudes at time 2 only to the extent that these exogenous variables affect the development of the same attitudes at time 1. Furthermore, no causal nexus is assumed to exist between work and family attitudes at either time 1 or time 2. The causal model among the latent variables is therefore hierarchical, but not fully recursive; in another context Wolfle (1980a) has called such models block-recursive.

The specification of this causal model is arguable. For example, how realistic is the postulation that the exogenous variables have no direct influence on work and family attitudes at time 2? In this case, the postulation is easily testable. One may simply respecify the model with the appropriate paths included, and determine if the parameter estimates are equal to values within appropriate ranges expected by chance if the parameters were in fact zero. In another instance, why specify that work at time 2 depends upon family attitudes at time 1, and conversely that family at time 2 depends upon work at time 1? Why not instead specify that work and family at time 2 are reciprocal causes of each other? In this case, there is no statistical test to fall back upon. To permit both sets of effects creates an underidentified model, which has no unique solution. Thus, one must choose between models on the basis of one's knowledge about the subject matter, or the analytic purpose of the model (see Wolfle, 1980a, pp. 203-204). To the extent that the reader finds these choices implausible, thus will the results be implausible (regardless of the method of estimation). In my view, the advantage of path diagrams and analyses of structural models is that their presentation requires a degree of explicitness which allows readers to decide for themselves how plausible or implausible are the models.

In essence, all we have considered so far in this discussion of the structural model are the two matrices of regression effects, gamma and beta. There are two further matrices to consider, phi and psi. Phi is the symmetrical variance-covariance matrix among the three exogenous factors; this matrix will have five unknown parameters to be estimated --

three covariances, and the variances of the socioeconomic and ability factors. Here, sex is shown to be a manifest variable, hence its variance is known a priori. Psi is the variance-covariance matrix among the four ζ_i ($i = 1 \dots 4$). In an initial estimation, this matrix is assumed to be a diagonal matrix, implying that none of the residual terms are correlated. Later applications will relax this assumption.

These relationships may be expressed algebraically:

$$\eta_1 = \gamma_{11}\xi_1 + \gamma_{12}\xi_2 + \gamma_{13}\xi_3 + \zeta_1$$

$$\eta_2 = \gamma_{21}\xi_1 + \gamma_{22}\xi_2 + \gamma_{23}\xi_3 + \zeta_2$$

$$\eta_3 = \beta_{31}\eta_1 + \beta_{32}\eta_2 + \zeta_3$$

$$\eta_4 = \beta_{41}\eta_1 + \beta_{42}\eta_2 + \zeta_4$$

Rearranging the equations so that all of the eta variables are to the left of the equality, allows the expression of this set of equations in matrix notation:

$$\underline{\beta} \underline{\eta} = \underline{\Gamma} \underline{\xi} + \underline{\zeta}$$

where $\underline{\eta}' = (\eta_1, \eta_2, \eta_3, \eta_4)$, $\underline{\zeta}' = (\zeta_1, \zeta_2, \zeta_3, \zeta_4)$, and $\underline{\xi}' = (\xi_1, \xi_2, \xi_3)$.

$\underline{\beta}$ is a 4 x 4 matrix such that:

$$\underline{\beta} = \begin{bmatrix} 1.0 & 0 & 0 & 0 \\ 0 & 1.0 & 0 & 0 \\ -\beta_{31} & -\beta_{32} & 1.0 & 0 \\ -\beta_{41} & -\beta_{42} & 0 & 1.0 \end{bmatrix}$$

$\underline{\Gamma}$ is a 4 x 3 matrix such that:

$$\Gamma = \begin{bmatrix} \gamma_{11} & \gamma_{12} & \gamma_{13} \\ \gamma_{21} & \gamma_{22} & \gamma_{23} \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix}$$

Furthermore, let $\Psi = \text{diag}(\psi_1, \psi_2, \psi_3, \psi_4)$, and:

$$\Phi = \begin{bmatrix} \sigma_{11} & & \\ \sigma_{21} & \sigma_{22} & \\ \sigma_{31} & \sigma_{32} & \sigma_{33} \end{bmatrix},$$

in which σ_{jj} denotes the variance of ξ_j , and σ_{jk} ($j \neq k$) denotes the covariance between ξ_j and ξ_k .³

Estimation of the free parameters of these matrices in LISREL requires that the user specify for each element in each matrix whether the element is free (a value to be estimated by the program), or fixed. If the latter, the user has to specify the fixed value. For example, the diagonal of β is to be fixed at unity, but other elements of β and some in Γ are to be fixed at zero. Furthermore, σ_{22} is to be fixed equal to the variance of the manifest variable, sex.

THE MEASUREMENT MODEL

In the structural portion of the model, the vectors η and ξ are not observed, but $y' = (y_1 \dots y_{12})$ and $x' = (x_1 \dots x_8)$ are observed, such that:

$$y = \Lambda_y \eta + \varepsilon,$$

and

$$\underline{x} = \underline{\Lambda}_x \underline{\xi} + \underline{\delta} ,$$

in which $\underline{\epsilon}$ and $\underline{\delta}$ are vectors of errors of measurement in $\underline{y}$ and $\underline{x}$, respectively. These errors of measurement represent both specific and random components of variation (Alwin and Jackson, 1979). The vectors $\underline{\epsilon}$ and $\underline{\delta}$ are assumed to be uncorrelated with $\underline{\eta}$, $\underline{\xi}$, and $\underline{\zeta}$, but may have elements internal to each which are correlated to other elements of the same vector. (LISREL does not permit elements in $\underline{\epsilon}$ to be correlated with elements in $\underline{\delta}$, but this is easily overcome by specifying all of the variables to be endogenous.) The matrices $\underline{\Lambda}_y$ (12 x 4) and $\underline{\Lambda}_x$ (8 x 3) are regression matrices.

In this analysis, socioeconomic status of the respondent's family of origin is indexed by father's occupational status, father's educational attainment, and mother's educational attainment. These manifest indicators were taken from the NLS data file, V2468, V1627, and V1628, respectively (see Levinsohn, et al., 1978); these measures were composite variables of the father's Duncan (1961) socioeconomic index, and father's and mother's education. This model specifies that these three variables have a common cause, assumed to be familial socioeconomic status. The model is not deterministic, however, and assumes that the underlying factor does not completely determine the observed variation in the three manifest variables; this specification is represented by the three δ_i ($i = 1, 2, 3$). These error terms have been loosely called errors of measurement, but in fact contain both random and specific errors of prediction from the underlying factor.

Alwin and Jackson (1979) have discussed the several measurement models that are applicable to such models. There is therefore no need to go into such a discussion in detail. A brief review, however, will sensitize us to assumptions being made about the models of measurement. Consider the three equations from socioeconomic status:

$$x_1 = \lambda_{11}\xi_1 + e_1 + u_1$$

$$x_2 = \lambda_{21}\xi_1 + e_2 + u_2$$

$$x_3 = \lambda_{31}\xi_1 + e_3 + u_3 ,$$

in which the x_i are the three manifest variables specified above; λ_{ij} are regression coefficients; ξ_1 is a latent factor of socioeconomic status; e_i are random components of error; and u_i are components of error specific to each manifest variable. In practice, the e_i and u_i are inseparable; all one is able to estimate are the $\delta_i = e_i + u_i$.

A metric for ξ_1 is imposed by setting $\lambda_{11} = 1.0$; as a result the λ_{ij} have arbitrary levels, but regardless of which one of the λ_{ij} are specified to be unity, the ratio of the λ_{ij} to each other will be constant. If the equations above are squared, and expectations taken, one obtains:

$$\sigma_{x_1}^2 = \sigma_{\xi_1}^2 + \sigma_{e_1}^2 + \sigma_{u_1}^2$$

$$\sigma_{x_2}^2 = \lambda_{21}^2 \sigma_{\xi_1}^2 + \sigma_{e_2}^2 + \sigma_{u_2}^2$$

$$\sigma_{x_3}^2 = \lambda_{31}^2 \sigma_{\xi_1}^2 + \sigma_{e_3}^2 + \sigma_{u_3}^2 .$$

This is called a common-factor measures model; for each manifest variable there exists a component of error that is due to random error, and a

component specific to each manifest variable. LISREL assumes as the most general case the common-factor model, which is not a classical true-score model. Therefore, if one wants to interpret the coefficients in terms of classical true-score theory, one has to make restrictive assumptions about the nature of the error term. Jöreskog (1971) used the common-factor model to classical true-score theory by defining the congeneric measures model. In such a model, each of the manifest variables x_i is assumed to have a separate true score, ξ_{1i} , which are perfectly correlated with each other, with the implication that a random variable, ξ_1 , exists such that all of the ξ_{1i} are linearly related to it:

$$\xi_{1i} = \mu_i + \lambda_i \xi_1$$

By making the further assumption that the errors across measures are entirely random (i.e., $u_i = 0$), and because $x_i = \xi_{1i} + e_i$, it follows that:

$$x_i = \mu_i + \lambda_i \xi_1 + e_i$$

If these three equations are squared, and expectations taken, one obtains:

$$\sigma_{x_1}^2 = \sigma_{\xi_1}^2 + \sigma_{e_1}^2$$

$$\sigma_{x_2}^2 = \lambda_{21}^2 \sigma_{\xi_1}^2 + \sigma_{e_2}^2$$

$$\sigma_{x_3}^2 = \lambda_{31}^2 \sigma_{\xi_1}^2 + \sigma_{e_3}^2$$

Alwin and Jackson (1979) suggest that it is difficult to imagine a latent factor score that all variation between its manifest indicators is due to the factor alone. The congeneric measures model is thus an ideal to be sought after (Bohrnstedt and Borgatta, 1980). As we will see below,

the model being discussed in this paper exhibits some congeneric constructs, but others definitely fall short of the ideal.

If one is willing to make a further restriction in the measurement model, namely that $\lambda_{11} = \lambda_{21} = \lambda_{31}$, the model becomes:

$$\sigma_{X_1}^2 = \sigma_{\xi_1}^2 + \sigma_{e_1}^2$$

$$\sigma_{X_2}^2 = \sigma_{\xi_1}^2 + \sigma_{e_2}^2$$

$$\sigma_{X_3}^2 = \sigma_{\xi_1}^2 + \sigma_{e_3}^2 ,$$

which is called a tau-equivalent measures model. The LISREL program allows users to specify, in addition to fixed and free parameters, parameters that are equivalent. Thus, the tau-equivalent model may be estimated by specifying $\lambda_{11} = \lambda_{21} = \lambda_{31}$, on the assumption that $u_i = 0$, which implies that $\sigma_{\delta_i \delta_j} = 0$ for all $i \neq j$.

Finally, by specifying not only that $u_i = 0$, and that $\lambda_{11} = \lambda_{21} = \lambda_{31}$, but in addition that $\sigma_{e_1}^2 = \sigma_{e_2}^2 = \sigma_{e_3}^2$, one obtains the parallel measures model:

$$\sigma_{X_1}^2 = \sigma_{\xi_1}^2 + \sigma_e^2$$

$$\sigma_{X_2}^2 = \sigma_{\xi_1}^2 + \sigma_e^2$$

$$\sigma_{X_3}^2 = \sigma_{\xi_1}^2 + \sigma_e^2 .$$

This is the most restrictive, least seldom seen in application, of the measurement models.

Of these models, this analysis assumes that father's occupational status, father's education, and mother's education are congeneric measures

of socioeconomic status. This assumption implies that u_i ($i = 1, 2, 3$) = 0; that is, the only error of prediction is random. A test of this assumption is possible to the extent that the overidentifying restrictions in the measurement model may be relaxed by permitting covariance parameters among measurement error terms to be free parameters in the model.

A second exogenous variable, sex, is included, not as a latent factor, but as a single manifest variable. This variable is a composite measure of sex (V1626) taken from the NLS data file (Levinsohn, et al., 1978). It is incorporated into the LISREL model by specifying:

$$x_4 = \xi_2 ,$$

that is, $\lambda_{42} = 1.0$, and $\delta_4 = 0$.

Several ability subtests were administered to the NLS respondents during their senior year of high school. Four of these, reading (V618), letter groups (V619), vocabulary (V614), and math (V620), are assumed to be congeneric measures of a latent ability factor.

The measures of the endogenous latent factors, work and family, have been described above. Here we adopt a common-factor measures model, because it seems unlikely that the $u_i = 0$, since errors of measurement of a particular attitude at time 1 are likely to be correlated with the corresponding errors at time 2.

The measurement portion of the model is defined by 20 equations, relating each of the 20 manifest variables to its respective latent factor. These have been omitted here; the measurement model is more easily represented in matrix notation. For the exogenous variables, the appropriate equation is:

$$\underline{x} = \underline{\Lambda_x} \underline{\xi} + \underline{\delta}$$

in which $\underline{x}' = (x_1, \dots, x_8)$, $\underline{\xi}' = (\xi_1, \xi_2, \xi_3)$, $\underline{\delta}' = (\delta_1, \dots, \delta_8)$,

and

$$\underline{\Lambda_x} = \begin{bmatrix} 1.0 & 0 & 0 \\ \lambda_{21} & 0 & 0 \\ \lambda_{31} & 0 & 0 \\ 0 & 1.0 & 0 \\ 0 & 0 & 1.0 \\ 0 & 0 & \lambda_{63} \\ 0 & 0 & \lambda_{73} \\ 0 & 0 & \lambda_{83} \end{bmatrix}.$$

Furthermore, let $\underline{\Theta}_{\delta}$ be the symmetrical variance-covariance matrix among the δ_i ($i = 1 \dots 8$). Initially, this matrix is assumed to be diagonal (all off-diagonal elements are fixed at zero) with the further specification that $\Theta_{\delta_{44}} = 0$.

For the endogenous variables, the measurement model is defined by:

$$\underline{y} = \underline{\Lambda_y} \underline{\eta} + \underline{\epsilon}$$

in which $\underline{y}' = (y_1, \dots, y_{12})$, $\underline{\eta}' = (\eta_1, \eta_2, \eta_3, \eta_4)$, $\underline{\epsilon}' = (\epsilon_1, \dots, \epsilon_{12})$,

and

$$\Lambda_y = \begin{bmatrix} 1.0 & 0 & 0 & 0 \\ \lambda_{21} & 0 & 0 & 0 \\ \lambda_{31} & 0 & 0 & 0 \\ 0 & 1.0 & 0 & 0 \\ 0 & \lambda_{52} & 0 & 0 \\ 0 & \lambda_{62} & 0 & 0 \\ 0 & 0 & 1.0 & 0 \\ 0 & 0 & \lambda_{83} & 0 \\ 0 & 0 & \lambda_{93} & 0 \\ 0 & 0 & 0 & 1.0 \\ 0 & 0 & 0 & \lambda_{11,4} \\ 0 & 0 & 0 & \lambda_{12,4} \end{bmatrix}$$

Furthermore, let Θ_{ϵ} be the symmetrical variance-covariance matrix among the ϵ_i ($i = 1 \dots 12$). Initially, this matrix is assumed to be diagonal.

This completes the definition of the initial LISREL model. Below we will consider certain alterations of the model, some to be made on practical grounds, some to be made on the basis of more theoretically oriented concerns.

IDENTIFICATION OF THE MODEL

Before the model is estimated, the identification problem must be examined. In the first case, an equation such as:

$$\tilde{x} = \Lambda_{\tilde{x}} \tilde{\epsilon} + \tilde{\delta}$$

cannot be directly estimated, since everything on the right side of the

equation is unobserved. Thus, the covariance structure of the observed variables is analyzed in terms of the postulated causal structure. If the equation above is multiplied by $\underline{x}'$, and expectations taken, one obtains the variance-covariance matrix of $\underline{x}$:

$$\begin{aligned}\Sigma_{\underline{x}} &= E(\underline{x} \underline{x}') = E[(\Lambda_{\underline{x}} \underline{\xi} + \underline{\delta})(\Lambda_{\underline{x}} \underline{\xi} + \underline{\delta})'] \\ &= \Lambda_{\underline{x}} \Phi \Lambda_{\underline{x}}' + \Theta_{\underline{\delta}}.\end{aligned}$$

In a more complicated way, the variance-covariance matrix of $\underline{y}$, and the covariance matrix, $\Sigma_{\underline{xy}}$, may be derived (see, for example, Long, 1976).

Analysis of the model in terms of covariance structures does not, however, resolve the identification problem. The model will be under-identified unless certain constraints are made. Specifically, there must be no more free parameters than there are elements in the lower triangle of the observed variance-covariance matrix, including the diagonal.

If there are t free parameters in the model, then:

$$t \leq \frac{1}{2}(p + q)(p + q + 1).$$

In this model, $p = 12$ and $q = 8$; thus t must be less than or equal to 210.

As specified above for the structural portion of the model, there are 4 free parameters in $\underline{\beta}$, 6 in $\underline{\Gamma}$, 5 in $\underline{\Phi}$, and 4 in $\underline{\Psi}$. In the measurement portion of the model, there are 8 free parameters in $\Lambda_{\underline{y}}$, 5 in $\Lambda_{\underline{x}}$, 12 in $\Theta_{\underline{\epsilon}}$, and 7 in $\Theta_{\underline{\delta}}$. Thus, $t = 51$, which is clearly less than 210.

Unfortunately, this condition is necessary, but not sufficient, for identification.

In addition, each and every equation in the model must be identified. In the measurement error portion of the model, this may be accomplished by setting one element in each column of the two Λ matrices to some fixed

value, usually 1.0.⁴ This becomes the reference indicator, and serves to provide a metric for the underlying factor. One may alternatively fix the variance of the underlying factor; this is commonly set to unity when analyzing standardized-solution models. Notice, however, that this is possible only for the exogenous factors in the Φ matrix. The variances of the η factors may not be fixed. Thus, reference indicators must be specified for the endogenous factors even when analyzing a correlation matrix. In each column of Λ_x and Λ_y there must also be one fewer zeros than there are n and m factors respectively. Examination of Λ_x and Λ_y reveals that there is a 1.0 in every column, and at least 2 zeros in each column of Λ_x and 3 zeros in each column of Λ_y . Satisfaction of these rules is still only a necessary condition for identification.

The structural portion of the model must also be internally identified. This means that excess degrees of freedom in the measurement model may not be used to identify an underidentified structural model. In this case there are $(\frac{1}{2})(m + n)(m + n + 1) = 28$ variances and covariances in the structural model, and 19 free parameters; thus the structural model satisfies the counting rule for identification. Unfortunately, it is often difficult to determine if every parameter is estimable, and there exists no general set of rules that applies to every model. Rules have been developed for certain types of models (see references in Jöreskog and Sörbom, 1978, p. 10), but the solutions are often tedious.

Users of LISREL should be aware that the program does provide solutions for underidentified models. Underidentification simply means that there is no unique solution to the model; the LISREL program simply stops when it finds one of the solutions. Use different starting values,

and one would obtain another solution. This foible of LISREL has caught at least one author unawares, who published his substantive interpretation of an underidentified model. Fortunately, LISREL does provide a check for identification. The iterative subprogram that calculates the standard errors of the estimates begins by calculating the information matrix for all the independent unknown parameters.

If this matrix is positive definite it is almost certain that the model is identified. On the other hand, if the information matrix is singular, the model is not identified and the following message will be printed

THE n-TH FREE PARAMETER MAY NOT BE IDENTIFIED .

This is a strong indication that the n-th free parameter in the pattern vector is not identified. The n-th parameter is usually the last parameter in a group of parameters connected in an indeterminacy (Jöreskog and Sörbom, 1978, p. 11).

As a result, it is recommended that one always request on the program's output parameter card that the standard errors be calculated and printed. Publishing these standard errors is in good taste, and should also be encouraged.

ESTIMATION AND TESTING

The independent parameters in the model are estimated by the method of maximum likelihood estimation as described by Gruvaeus and Jöreskog (1970). This method assumes the distribution of $(\underline{y}', \underline{x}')$ is multivariate normal; the loss of efficiency by violations of this assumption has not

yet been established. The estimation problem is essentially that of fitting the variance-covariance matrix imposed by the model to the sample variance-covariance matrix. A fitting function (see Jöreskog and Sörbom, 1978, p. 13) is minimized by using first and second order derivatives, and converges rapidly from most arbitrary starting points. If there are several local minima of the fitting function, however, there is no guarantee that the method will converge to the absolute minimum. Users are well advised to reestimate models with different starting values in case the program has converged to a local minimum of the fitting function.

When the maximum likelihood estimates of the parameters have been obtained, a χ^2 -measure (as distinguished from χ^2 , which is the sampling distribution to which χ^2 is compared) of overall goodness-of-fit of the model is calculated. This statistic may be regarded as a test of the specific model against the most general alternative that the estimated variance-covariance matrix is any positive definite matrix. The likelihood ratio χ^2 statistic is $(N/2) F_0$, where F_0 is the minimum value of the fitting function, and N is the sample size. In large samples this statistic is distributed as χ^2 with degrees of freedom:

$$df = (\frac{1}{2})(p + q)(p + q + 1) - t ,$$

where t is the number of independent parameters. This test tells if the model fits or does not fit, but if the latter cannot tell where the model does not fit.

The likelihood ratio χ^2 statistic is sensitive to even small deviations from perfect fit. Particularly when the sample size is large,

it is common to find large, or even very large, values of χ^2 relative to the degrees of freedom. When the sample size is large, one common rule of thumb is to adopt as appropriate a model in which χ^2/df is less than 5.

Another alternative is to fit a more restrictive model against a less restrictive model, and to compare the resulting two χ^2 measures. The difference between the two likelihood ratio χ^2 measures is distributed approximately as χ^2 with degrees of freedom equal to the difference in the number of independent parameters in the two models. If there is a large drop in χ^2 from one model to the other relative to the difference in degrees of freedom, then the changes made in the second model represent a real improvement in fit. These procedures will be illustrated below; Bentler and Bonett (1980) have investigated significance testing in models such as these.

Sometimes changes in models can be suggested on the basis of logic. For example, the model under consideration in this paper postulates no direct effects from the three exogenous factors to work and family attitudes at time 2. If the initial model does not fit, perhaps it would be wise to relax this restrictive assumption, and see if the new model produces a real improvement in the fit. In the measurement portion of this model, the initial model as specified assumes that the errors of measurement for work and family attitudes are uncorrelated from time 1 to time 2. However, it is often the case that such errors of measurement are correlated. Accordingly, a new model may be estimated with these covariances included as free parameters. The difference in χ^2 measures would tell us whether the errors of measurement were in fact correlated.

At other times, changes in the model may be suggested by an inspection of the first-order derivatives of the fitting function with respect to the fixed parameters. This table is available in LISREL, and it is a good practice to request its printing. Sörbom (1975) has published a useful example describing how the procedure works in practice. One may want to relax the restriction in the model which is least probably zero. The table of first-order derivatives suggests which fixed parameter if set free will give the largest decrease in the fitting function. If doing so does not violate the logic of the causal structure, one should find the fixed parameter whose absolute value of its first-order derivative is the greatest. A new model is then fitted with the restriction removed, and the χ^2 measures from the two models compared to see if the less restrictive model provides a significant improvement in fit. In any event, it is wise to keep in mind Jöreskog's advice, "Ultimately the criteria for goodness of the model depends on the usefulness of it and the results it produces (Jöreskog, 1969, p. 201)."

THE DATA

Data for the model depicted in Figure 1 were obtained from the National Longitudinal Study of the High School Class of 1972 (Levinsohn, et al., 1978). The NLS was designed to provide data on the development of the educational, vocational, and personal aspects of the lives of adolescents as they make the transition from high school to the adult world. The population analyzed here includes only those white respondents who completed the study's ability tests in their senior year of high school. There were 12,844 such respondents. From this group,

a random subsample was selected, and correlations calculated for all those with listwise present data. The analysis to be described below was therefore based on 751 white respondents. The correlations and standard deviations for the 20 variables are shown in the appendix.⁵

FITTING THE MODEL

Because LISREL is an iterative program, it must be provided with starting values for all of the elements in every matrix not defined to be either an identity matrix or a zero matrix. The default value assumed by the program is zero, so only those starting values for non-zero fixed parameters, and free parameters need be specified. These starting values may be chosen arbitrarily, but must in the first iteration produce an estimated variance-covariance matrix which is positive definite. If this does not hold, the program terminates abnormally.

Selecting appropriate starting values is not always easy. Because the correlation coefficients of one variable with two other variables restrict the range of the correlation coefficient of the second and third variables, it is possible to generate estimated correlations or covariances in the first iteration which are internally inconsistent. Indeed, with the model at hand several different sets of starting values were required before the program would run. It is a good idea to create simple summated factors of the variables in the analysis, and use OLS regression to estimate appropriate starting values of LISREL. Not only does this help avoid the problem of producing a nonpositive matrix, but it reduces the computer time of LISREL by providing it

starting values more likely to be close to those in the final solution. (By the way, a mistake I find myself repeating is to forget that elements in β are reversed in sign.)

We will now turn to the presentation of results for the model described above. Summary goodness-of-fit statistics will be presented for a number of specifications. When a model is found deemed to be best, its parameter estimates will be discussed in substantive terms.

The summary measures of goodness-of-fit are shown in Table 1. Model A of Table 1 assumes errors of measurement are entirely random, and imposes a causal structure among the latent factors as shown in Figure 1. The likelihood ratio χ^2 value for this model is 715.29 with 159 degrees of freedom. At first blush, this value suggests the model does not fit the data very well, but there is more wrong with this model than a mere lack of fit. An examination of the LISREL estimates (not shown here) reveals that the residual variance for family at time 2 is a negative number. This occurrence, all too frequent, is known as the Heywood Case (apocryphally after the Rev. Christopher Heywood, who kept coming up with a negative number of angels who could stand on the head of a pin). Variance estimates zero or greater are logically permissible; negative values are not. It is entirely possible, however, as in this case, for maximum likelihood estimates to converge at a value less than zero. The standard fix for the Heywood Case is to constrain the offending variance to zero, or a small positive number. In this instance, ψ_{44} was set to zero, and the model reestimated. Substantively, this means that family attitude at time 2 is assumed to be perfectly predicted by some linear combination of family and work attitudes at

Table 1. Goodness-of-Fit Statistics for Models of Development and Stability of Attitudes toward Work and Family

Model	χ^2	d.f.	Prob.	$\Delta\chi^2$	d.f.	Prob.
A. Random errors	715.29	159	0.0			
B. Random errors; $\psi_{44} = 0$	744.81	160	0.0			
C. Random errors; $\psi_{44} = 0$; ψ_{21} free	733.97	159	0.0	10.84	1	.001
D. Random errors; $\psi_{44} = 0$; ψ_{21} , γ_{33} free	731.02	158	0.0	2.95	1	.086
E. Random errors; $\psi_{44} = 0$; ψ_{21} , ψ_{43} , γ free	728.00	152	0.0	5.97	7	.543
F. Covariance between work and family errors time 1 and time 2; $\psi_{44} = 0$; ψ_{21} free	333.85	153	0.0	400.12	6	0.0
G. Model F, with $\theta_{\epsilon_{31}}$ free	321.57	152	0.0	12.28	1	.001
H. Model G, with $\theta_{\epsilon_{76}}$ free	315.13	151	0.0	6.44	1	.011
I. Model H, with $\theta_{\epsilon_{41}}$ free	313.04	150	0.0	2.09	1	.148

time 1. The new result is shown in Model B of Table 1. By setting $\psi_{44} = 0$, one degree of freedom has been gained, and the likelihood ratio χ^2 is 744.81 with 160 degrees of freedom. All of the LISREL estimates in the new model are logically permissible, but the χ^2 value suggests a poorly fitting model. How should the model be changed to see if a better fitting model is possible? One reasonable strategy, suggested by Kenny (1979, p. 161), is to fit a just-identified structural model. Any lack of fit of the model to the correlations could therefore be attributed to specification errors in the measurement model. This was done in stages.

An examination of the first-order derivatives among the structural matrices produced by Model B suggested that the covariation between ζ_1 and ζ_2 (i.e., ψ_{21}) was the value most likely not to be zero.⁶ In substantive terms, this indicates that the three exogenous factors have not explained all of the covariation between family attitudes and work attitudes at time 1; it is likely that the multiple partial correlation between WORK73 and FAMILY73 is nonzero. Permitting this parameter to be freely estimated in a new model, Model C, provides a significant improvement in the fit of the model; the χ^2 value is 733.97 with 159 degrees of freedom. The difference between the χ^2 measures for Models B and C is 10.84, which is distributed as χ^2 with one degree of freedom; the value is significant.

An examination of the table of first-order derivatives from Model C revealed that, among the fixed structural coefficients, γ_{33} was the new value most likely not to be zero. This is the direct effect of ability on work at time 2. Model D was estimated with this value set free, but the improvement of fit between Model C and Model D was not significant.

Model E was estimated (as suggested by Kenny, 1979) with the structural portion of the model completely identified. These results reveal two items of interest. First, the χ^2 value of 728.00 suggests a severe lack of fit, but because the structural model is completely identified (save for $\psi_{44} = 0$) we now know that the lack of fit must be in the measurement portion of the model. Second, comparing the χ^2 measures of Model E with Model C reveals that the seven additional parameters set free in Model E do not yield a significant improvement in fit. That is, the seven additional parameters differ from zero only as a matter of chance. Thus, our initial postulation that the background factors do not directly influence work and family attitudes at time 2 is confirmed. Furthermore, the covariance between work and family attitudes at time 2 is almost completely explained by the previous expression of these attitudes. The only change we have made in the structural portion of the model is to allow the residuals of work and family at time 1 to covary; this merely represents the inability of the three exogenous factors to completely explain their covariation.

Model C represents the best fitting of the structural models, but it must be admitted that it does not fit very well. The model assumes that errors of measurement are random, but the model's lack of fit suggests that the assumption is not tenable. The model contains the same variables measured at two different times. It is well known that when the same measuring instrument is used at two or more occasions, there is a tendency for the errors in each variable to covary over time because of memory or other re-test effects. Accordingly, Model F was estimated with the errors for each of the three work manifest variables allowed to covary

with the equivalent errors for the questions repeated at time 2; the equivalent parameters were also set free for family attitudes. Comparison of the fit of this model will be made to that of Model C, since Model C was the accepted structural model. The difference in χ^2 measures was 400.12 with 6 degrees of freedom. This is a major improvement in fit, and suggests that we were well advised to consider the possibility of correlated errors over time.

It is still possible, of course, for there to be other sources of covariation among the errors of measurement. For example, respondents who express a desire for being successful in work may systematically overidentify that response with the desire to have money. Thus, the errors of measurement for these two questions will covary to the extent that the underlying factor is unable to explain the covariation between the manifest indicators. The choice of example was prophetic. An examination of the first-order derivatives of the fixed parameters in Model F revealed $\theta_{\epsilon_{31}}$ was the value most likely not to be zero; this is the covariation between the questions at time 1 about being successful in work and being able to find steady work. Model G was estimated with this new parameter set free. The difference in χ^2 measures between Model F and Model G was 12.28 with one degree of freedom -- a significant improvement in fit.

The process continued. The first-order derivatives for Model G were examined with the result that $\theta_{\epsilon_{76}}$ was discovered to be the value most likely not to be zero. This is the covariation between the errors for the question about moving away from one's family, asked at time 1, and the question about being successful in one's work, asked at time 2.

One can see how this may be plausibly explained; respondents who systematically overstated or understated their desire to move away at one time may be those who overstate or understate their desire for success at another time. If I had been unable to construct a plausible explanation, I would not have estimated Model H, which allowed this parameter to be free. This decision obviously went beyond the statistical information available (see Blalock, 1972, p. 448), and reflects my belief that no model may be said to fit, which produces uninterpretable results.

Model H, when compared to Model G, suggests a significant improvement in fit. An examination of the first-order derivatives from Model H revealed $\theta_{\epsilon_{41}}$ was the value most likely not to be zero. Model I was therefore estimated with this parameter set free, but the improvement of fit over that of Model H was not significant. Thus, Model H is accepted as the best fitting model to explain both the structural and measurement properties of a model of the development and stability of attitudes toward work and family.

EMPIRICAL RESULTS

This section presents the results from the measurement and structural analysis of Model H. The measurement properties will be discussed first. These results are shown in Table 2, and mirror the LISREL solutions reproduced in the appendix, which also contains the LISREL program set up used to generate the parameter estimates. Table 2 contains the parameter estimates for the true score variances, error variances, and the slopes of the manifest variables as regressed on the latent factors.

Table 2. Model H Measurement Parameter Estimates

Latent Factor	Manifest Variable**	True Score Variance	Error Variance	Slope	Reliability
Work 73	V685	.018	.170	1.00*	.096
	V687		.286	1.458	.118
	V689		.244	2.546	.323
Family 73	V686	.019	.160	1.00*	.106
	V692		.355	1.477	.105
	V693		.329	.679	.026
Work 74	V1573	.037	.164	1.00*	.184
	V1575		.308	1.446	.201
	V1577		.143	2.285	.575
Family 74	V1574	.035	.161	1.00*	.179
	V1580		.368	1.228	.125
	V1581		.382	.271	.007
Socioeconomic Status	V2468	204.782	340.066	1.00*	.376
	V1627		.172	.083	.891
	V1628		.662	.041	.342
Ability	V618	59.244	27.346	1.00*	.684
	V619		44.007	.705	.401
	V614		42.170	.891	.527
	V620		37.382	.897	.560

* Fixed parameter.

** Variable labels from Levensohn, et al. (1978).

The last column shows the estimated reliability coefficients. These may be estimated by:

$$\lambda_{ij}^2 (\sigma_j^2 / \sigma_i^2)$$

(Wolfe, 1980b), where λ_{ij}^2 is the estimated slope of the i -th manifest variable regressed on the j -th latent factor, σ_j^2 is the estimated variance of the j -th latent factor, and σ_i^2 is the sample variance of the i -th manifest variable; or they may be estimated by:

$$(\lambda_{ij}^2 \sigma_j^2) / (\lambda_{ij}^2 \sigma_j^2 + \theta_i^2)$$

(Jöreskog, 1971), where λ_{ij}^2 and σ_j^2 are as previously defined, and θ_i^2 is the estimated error variance (from either the θ_e or θ_δ matrix) of the i -th manifest indicator. These two expressions are algebraically equivalent, and in either case measure an indicator's reliability as the proportion of its expected variance due to the variance of its corresponding latent factor, or true score.

The estimated reliability coefficients for work and family attitudes are very low. We have already seen that the measurement errors contained unique components that were correlated from one administration of the survey to the next. We now see that there are also rather large random errors associated with these variables. There is, of course, some variance explained among the manifest indicators by the latent true scores, and to that extent there do seem to be underlying factors which measure one's general attitudes toward work and family. Insofar as measuring these constructs is concerned, however, these manifest variables do not do a very adequate job. Most of the variance in the manifest indicators is explained by errors of measurement, both unique and random.

Among the exogenous factors, we have already confirmed that father's occupation and education, and mother's education, are congeneric measures of socioeconomic status; that is, we found evidence to suggest their errors of measurement were probably just random. The same may also be said of the four indicators of ability. Examining the estimated reliability coefficients, one sees that father's education is by far the most reliable indicator of socioeconomic status; this result supports a similar analysis reported by Wolfle and Robertshaw (1981). Among the indicators of ability, the reading and math subtests are more reliable than the letter-group and vocabulary subtests; this result also is reflected in other analyses incorporating these variables (Wolfle and Lichtman, 1981; Wolfle and Robertshaw, 1981).

We turn now to the discussion of the structural portion of Model H. These results are shown in Table 3. When the development of attitudes toward work are considered, it may be seen that the effect of socioeconomic status is not significant. The variable, sex, was coded 1 = male and 2 = female; therefore the negative coefficient indicates that women consider work values less important than do men, ceteris paribus. The effect of ability is negative; that is, the greater one's score on the ability factor, the less importance is attached to work values. Attitudes toward family values are also negatively influenced by ability and socioeconomic status. Women, however, place more importance on family values than do men. These exogenous variables explain about one-fourth of the variance in work and family attitudes, but do not contribute to the further measurement of these variables once previous measures of the same attitudes have been controlled for. Both work attitudes and family attitudes

Table 3. Model H Structural Parameter Estimates

Dependent Factors	Independent Factors					Proportion Explained Variance
	Socioeconomic Status	Sex	Ability	Work 73	Family 73	
Standardized Coefficients						
Work 73	-.052	-.350	-.276			
Family 73	-.249	.241	-.255			
Work 74	- - -	- - -	- - -	.780	.046	
Family 74	- - -	- - -	- - -	.053	.978	
Unstandardized Coefficients*						
Work 73	-.000 (.001)	-.094 (.020)	-.005 (.001)			.22
Family 73	-.002 (.001)	.066 (.022)	-.005 (.001)			.26
Work 74	- - -	- - -	- - -	1.124 (.213)	.065 (.105)	.65
Family 74	- - -	- - -	- - -	.074 (.159)	1.329 (.255)	1.00

* Standard errors shown in parentheses.

are relatively stable variables (that is, the latent factors, not their manifest indicators); the correlation between the work factors implied by this analysis is .799, and the correlation between family factors over time is .999 -- a very stable construct, indeed.

CONCLUSION

Previous introductory discussions of LISREL (e.g., Jöreskog, 1974; Long, 1976; Jöreskog and Sörbom, 1978) have produced a number of books covering a range of possible applications, but none have gone into the development and estimation of a single model to the extent presented here. The lengthy specification of a single model may have seemed tedious to some; I hope it proves useful to others.

I would like to close with some reflections on the utility of LISREL. Kerlinger (1977) has said that we are in the midst of a revolution in research thinking, and cited LISREL by way of example. More recently Cooley (1978) reminded us that knowing one's subject matter is very much more important than using any particular set of mathematical tools. Using LISREL is an aid to thinking about and analyzing complex causal models, but like many new analytic techniques runs the risk of being faddishly adopted to implausible applications. Using LISREL presumes one has multiple indicators of underlying variables which are of theoretical interest. This means that one knows beforehand which indices accurately measure variables of theoretical interest, or that one has obtained alternative measures of underlying traits. This may mean conducting panel surveys as reported in this paper, or using alternative means of collecting data at one point in time (e.g., mailed questionnaires, telephone surveys, and personal interviews). In any event, the collection

of such data is likely to be expensive. In other words, like Leontief's input-output charts of the economy, LISREL may not simply be used in place of OLS regression or exploratory factor analysis, but requires changes in the data collection process, and the way researchers think about their analyses. Such changes will not come easily or cheaply, and we are likely to see implausible and incorrect analyses using LISREL along the way. Yet in the long-run, I share Kerlinger's (1977) optimism. LISREL provides a synthesis of analytic procedures formerly considered disjointly, and in applicable situations will allow the analysis of theoretical models not previously estimable. The accumulation of knowledge of educational relationships will benefit accordingly.

FOOTNOTES

¹ This is a revised version of a paper originally presented at the annual meetings of the American Educational Research Association, Los Angeles, 1981. I would like to thank Alan Acock, Jeffrey Danes, and anonymous referees for their comments on an earlier draft of this paper. I am also indebted to Dianne Robertshaw and Barbara Patteson, who will be found between the lines.

² Note, however, that these models assume a linear relationship between the ordinal indicators and the (assumed to be) interval scaled latent variable. Little is known about how robust this assumption of linearity is.

³ Normally, one could designate the elements of ϕ as ϕ_{11} , etc., but since variances are usually designated by lower-case sigma, I have done so here.

⁴ Which parameter is set to unity will affect the absolute values of the other free parameters for the same latent variable, but not their relative magnitudes. For example, if $\lambda_1 = 1.0$ in a two-variable model, and λ_2 was found to be .8 upon estimation, then if λ_2 were set equal to unity, λ_1 would be found equal to 1.25. Thus, $\lambda_1/\lambda_2 = \lambda_1/\lambda_2$ regardless of which parameter was set equal to unity.

⁵ The appendix has not been published with the body of this manuscript, but may be obtained upon request by writing to the author: Lee M. Wolfle / College of Education / University City Office Building / Virginia Tech / Blacksburg, Virginia 24061 .

⁶ Jöreskog has recently discovered that the magnitude of the first-order derivatives depend on the magnitude of the data and the parameter values. As a result, the first-order derivatives are not standardized, and are not strictly speaking comparable. Unfortunately, there now exists no readily available alternative to the first-order derivatives. However, when the next version of LISREL is released,

the program will provide a matrix of the ratio between the squared first-order derivative and the second-order derivative, which provides the best index of the fixed parameter which when relaxed will maximally improve the fit. In any event, relaxing parameters should be done only when it makes substantive sense to do so. See Karl G. Jöreskog, "Basic issues in the application of LISREL," DATA, 1981, 1 (June): 1-6.

⁷ This is not strictly true, since one may specify each latent variable to be perfectly measured by a single manifest variable. But such models are just special cases of the general LISREL model, which unlike other methods is unique in its ability to estimate both measurement error parameters and structural parameters simultaneously.

REFERENCES

- Alwin, D. F., and Jackson, D. J. Measurement models for response errors in surveys: Issues and applications. In K. F. Schuessler (ed.), Sociological Methodology 1980. San Francisco: Jossey-Bass, 1979.
- Bentler, P. M., and Bonett, D. G. Significance tests and goodness of fit in the analysis of covariance structures. Psychological Bulletin, 1980, 88, 588-606.
- Bielby, W. T., Hauser, R. M., and Featherman, D. L. Response errors of black and nonblack males in models of the intergenerational transmission of socioeconomic status. American Journal of Sociology, 1977, 82, 1242-1288.
- Blalock, H. M., Jr. Causal Inferences in Nonexperimental Research. Chapel Hill: University of North Carolina Press, 1964.
- Blalock, H. M., Jr. Estimating measurement error using multiple indicators and several points in time. American Sociological Review, 1970, 35, 101-111.
- Blalock, H. M., Jr. Social Statistics, 2nd edition. New York: McGraw-Hill, 1972.
- Bohrnstedt, G. W., and Borgatta, E. F. Forward to special issue on measurement. Sociological Methods & Research, 1980, 9, 139-146.
- Borgatta, E. F., and Bohrnstedt, G. W. Level of measurement: Once over again. Sociological Methods & Research, 1980, 9, 147-160.

- Burt, R. S. Confirmatory factor-analytic structures and the theory construction process. Sociological Methods & Research, 1973, 2, 131-190.
- Cooley, W. W. Explanatory observational studies. Educational Researcher, 1978, 7(9), 9-15.
- Duncan, O. D. A socioeconomic index for all occupations. In A. J. Reiss, Jr. (ed.), Occupations and Social Status. New York: The Free Press of Glencoe, 1961.
- Duncan, O. D. Introduction to Structural Equation Models. New York: Academic Press, 1975.
- Goldberger, A. S. Structural equation methods in the social sciences. Econometrica, 1972, 40, 979-1001.
- Gruvaeus, G. T., and Jöreskog, K. G. A computer program for minimizing a function of several variables. Research Bulletin 70-14. Princeton, N.J.: Educational Testing Service, 1970.
- Hauser, R. M. Educational stratification in the United States. Sociological Inquiry, 1970, 40, 102-129.
- Hauser, R. M., and Goldberger, A. S. The treatment of unobservable variables in path analysis. In H. L. Costner (ed.), Sociological Methodology 1971. San Francisco: Jossey-Bass, 1971.
- Jöreskog, K. G. UMFLA: A computer program for unrestricted maximum likelihood factor analysis. Research Memorandum 66-20. Princeton, N.J.: Educational Testing Service, 1966.

- Jöreskog, K. G. Some contributions to maximum likelihood factor analysis. Psychometrika, 1967, 31, 165-178.
- Jöreskog, K. G. A general approach to confirmatory maximum likelihood factor analysis. Psychometrika, 1969, 34, 183-202.
- Jöreskog, K. G. Statistical analysis of sets of congeneric tests. Psychometrika, 1971, 36, 109-133.
- Jöreskog, K. G. Analyzing psychological data by structural analysis of covariance matrices. In D. H. Krantz, et al. (eds.), Contemporary Developments in Mathematical Psychology, Volume II. San Francisco: W. H. Freeman, 1974.
- Jöreskog, K. G., Gruvaeus, G. T., and van Thillo, M. ACOVIS: A General Computer Program for Analysis of Covariance Structures. Princeton, N.J.: Educational Testing Service, 1970.
- Jöreskog, K. G., and Sörbom, D. LISREL IV: Analysis of Linear Structural Relationships by the Method of Maximum Likelihood. Chicago: National Educational Resources, 1978.
- Jöreskog, K. G., and Sörbom, D. Advances in Factor Analysis and Structural Equation Models. Cambridge: Abt Books, 1979.
- Kenny, D. A. Correlation and Causality. New York: John Wiley, 1979.
- Kerlinger, F. N. The influence of research on educational practice. Educational Researcher, 1977, 6(8), 5-12.
- Land, K. C. On the estimation of path coefficients for unmeasured variables from correlations among observed variables. Social Forces, 1970, 48, 506-511.

- Lawley, D. N. The application of the maximum likelihood method to factor analysis. British Journal of Psychology, 1943, 33, 172-175.
- Levinsohn, J. R., Henderson, L. B., Roccobono, J. A., and Moore, R. P. National Longitudinal Study: Base Year, First, Second and Third Follow-Up Data File Users Manual, Volumes I and II. Washington, D.C.: National Center for Education Statistics, 1978.
- Long, J. S. Estimation and hypothesis testing in linear models containing measurement error: A review of Joreskog's model for the analysis of covariance structures. Sociological Methods & Research, 1976, 5, 157-206.
- Mare, R. D., and Mason, W. M. Children's reports of parental socioeconomic status: A multiple group measurement model. Sociological Methods & Research, 1980, 9, 178-198.
- Mason, W. M., Hauser, R. M., Kerckhoff, A. C., Poss, S. S., and Manton, K. Models of response error in student reports of parental socioeconomic characteristics. In W. H. Sewell, R. M. Hauser, and D. L. Featherman (eds.), Schooling and Achievement in American Society. New York: Academic Press.
- Platt, J. R. Strong inference. Science, 1964, 146, 347-353.
- Siegel, P. M., and Hodge, R. W. A causal approach to the study of measurement error. In H. M. Blalock, Jr. and A. B. Blalock (eds.), Methodology in Social Research. New York: McGraw-Hill, 1968.
- Sörbom, D. Detection of correlated errors in longitudinal data. British Journal of Mathematical & Statistical Psychology, 1975, 28, 138-151.

- Werts, C. E., Linn, R. L., and Jöreskog, K. G. A simplex model for analyzing academic growth. Educational and Psychological Measurement, 1977, 37, 745-757.
- Wolfle, L. M. Strategies of path analysis. American Educational Research Journal, 1980a, 17, 183-209.
- Wolfle, L. M. Unmeasured variables in path analysis: Addendum. Multiple Linear Regression Viewpoints, 1980b, 10(4), 61-63.
- Wolfle, L. M., and Lichtman, M. Educational attainment among whites, blacks, and Mexican-Americans. Paper read at American Educational Research Association's Annual Meeting, Los Angeles, 1981.
- Wolfle, L. M., and Robertshaw, D. Latent causal structures in panel surveys. Paper read at American Educational Association's Annual Meeting, Los Angeles, 1981.

TIME TO SECURE EMPLOYMENT ANALYZED BY SETWISE REGRESSION ANALYSIS

JOLE A. WILLIAMS

JOHN D. WILLIAMS

Abstract - Graduates of Trinity (Texas) University were surveyed regarding the length of time necessary for them to secure their first post-baccalaureate employment. Five different sets of predictor variables were used: biographical variables (age and sex); ethnic background and socio-economic status; academic major and perception of major and Trinity University; employment seeking behaviors (salary, number of resumes, number of interviews, and number of offers) and geographic and size of employer preferences. The employment seeking behaviors set was clearly the most important set of predictors. Those who took longer to find employment generally worked harder, sending out more resumes, having more interviews, but receiving lower salaries.

During World War II and for many years thereafter, the need was not great for placement assistance for graduating college students. However, with the onset of economic recession in the late 1960's and early 1970's, desirable employment was more difficult to obtain. Another change occurred affecting the number of available jobs usually filled by the previously small number of college graduates. With the availability of many forms of financial aid, the economic and sociologic level of the members of the student bodies changed considerably. People who formerly could not have afforded to attend college were now enrolled, were graduating and entering the already shrinking job market (Wolfbein, 1970).

In addition, the discontinuance of the draft also had an affect on the number of students seeking employment after obtaining a baccalaureate degree.

With the state of the national economy and the abundance of college graduates, employers became more selective. Therefore, recruiters of college graduates were faced with the task of finding the best qualified job applicants from a labor pool that was almost unlimited from their viewpoint.

A few favored career fields were still experiencing immediate employment but they were the exceptions. Some majors, particularly in the liberal arts, but also including fields such as biology and psychology, are experiencing difficulty in obtaining employment in their own field, and after some failures in obtaining suitable employment in their chosen career field, opt for unrelated employment.

The purpose of this study was to determine if there is a significant relationship between length of time it takes college graduates to obtain employment and job related variables. The group surveyed in this study was the Trinity (Texas) University graduates who had received bachelor's degrees in December 1978, May 1979, and August 1979. A questionnaire was sent to each of the Trinity University graduates; 370 useable questionnaires were returned (65 percent).

The Criterion: Time to Secure Employment

The questionnaire required the respondent to choose a predetermined category in regard to length of time to secure employment. In turn the time variable was transformed into a continuous variable as follows:

<u>Value</u>	<u>Time to Secure Employment</u>
0	Employed before graduation
1	1 week after graduation
3	2-4 weeks after graduation
6	5-7 weeks after graduation
10	8-11 weeks after graduation
20	3-6 months after graduation
39	6-12 months after graduation
52	Over one year after graduation
Not included in analysis	Attending graduate school
65	Unemployed

Several points should be made about the constructed criterion. For values up to $Y = 39$, the criterion represents weeks after graduation that employment was secured. Those who secured employment over a year after graduation should cause little distortion, since the maximum time since graduation would have been less than 14 months. Graduate students were excluded from the analysis for the simple reason that no acceptable value could be applied uniformly across persons and programs. Some people may go on to graduate school because they see little likelihood of securing employment utilizing the skills they learned as an undergraduate. At the other extreme, some students may be quite confident that suitable employment awaits them; they can afford the luxury of acquiring additional skills so that they may become even more marketable. One might suspect that most students who went on to graduate school may have had motives somewhere between these two extremes.

While 54 of the 367 graduates (14.71%) attended graduate school, such attendance was not uniform across disciplines. Among smaller departments (less than 15 graduates), half or more attended graduate school from the

following: environmental studies, art history, economics, social sciences, French, and chemistry. Among larger departments, 52.94% of political science majors, 33.33% of psychology majors and 27.27% of biology majors attended graduate school. At the other end of the spectrum, no graduates in either sociology or elementary education attended graduate school, while 3.03% of graduates in engineering related majors, 10.53% of business majors and 11.11% of business administration majors attended graduate school immediately upon their receiving their bachelor degrees.

Inclusion of those unemployed in the analysis was done for reasons of best exploiting the data. It was felt that including them in the analysis with the highest possible value for the criterion ($Y = 65$) was better than excluding them in finding relationships among the data. To take an example of a particular department, seven of the fifteen graduates in psychology were employed no later than four weeks after graduation; one person was employed after one year; five went on to graduate school and two were unemployed. If these last seven graduates had been left out of the analysis, it would appear that all but one graduate was hired immediately upon graduation. While the value 65 does not represent 65 weeks, it does assign a reasonable number to be used in a regression analysis.

Predictor Variables Used in the Analysis

The following variables were used in the analysis:

Biographical Variables

X_1 = age;

X_2 = sex; 1 if male, 0 if female;

Ethnic Background and Socio-Economic Status (SES)

X_3 = White = 1, 0 otherwise;

X_4 = Black = 1, 0 otherwise;

X_5 = American-Indian, 0 otherwise;

(The zero coded variable for X_3 , X_4 and X_5 is Mexican-American)

X_6 = SES;

Academic Major

X_7 = major, 1 if engineering related, 0 otherwise

X_8 = value of major;

X_9 = value of Trinity degree;

Employment Seeking Behaviors

X_{10} = salary;

X_{11} = number of resumes;

X_{12} = number of interviews;

X_{13} = number of job offers;

Geographic and Size of Employer Preference

X_{14} = geographic preference, 1 = yes, 0 = no; and

X_{15} = size of employer preference, 1 = yes, 0 = no.

Before an analysis was completed, the 54 graduates who attended graduate school full time were eliminated; also, any graduate who was missing data on any of the predictor variables was not included. A setwise regression (Williams and Lindem, 1971, a, b) was used to analyze the data. Results are shown in Table 1.

Table 1

Setwise Regression Analysis Using all Five Sets of
Variables and Time to Secure Employment
(N = 269)

Step	Set Eliminated	r	R for Set Eliminated	R
1	None			.40847 ^a
2	Biographical Variables		.06984	.40746 ^a
	X ₁ : age	-.042		
	X ₂ : age	-.056		
3	Ethnic Background and SES		.03254	.40521 ^a
	X ₃ : White	-.002		
	X ₄ : Black	.009		
	X ₅ : American-Indian	.030		
	X ₆ : SES	-.001		
4	Geographic and Size Preference		.03572	.39594 ^a
	X ₁₄ : Geographic Pre- ference	.036		
	X ₁₅ : Size of Employer Preference	.014		
5	Academic Major		.17079 ^b	.38611 ^a
	X ₇ : Major	-.112		
	X ₈ : Value of major	.135 ^b		
	X ₉ : Value of Trinity degree	.115 ^b		
6	Employment Seeking Behaviors		.38611 ^a	
	X ₁₀ : Salary	-.212 ^a		
	X ₁₁ : Number of resumes	.255 ^a		
	X ₁₂ : Number of Interviews	.248 ^a		
	X ₁₃ : Number of job offers	-.028		

a_p < .01

b_p < .05

First, if the zero-order correlations are inspected, age and sex show no significant relationship, suggesting that at least these forms of discrimination don't occur in regard to length of time to obtain employment. However, discrimination may exist even in regard to time; while it is not reported in Table 1, sex correlates .28 with initial salary; even in the presence of the remaining variables, the partial relationship of sex to salary remains significant, and favors men. It might well be that women graduates "settle" on a less attractive employment situation in approximately the same time frame that men acquire somewhat more attractive (in terms of salary at least) employment.

Ethnicity seems to be unrelated to length of time to secure employment; geographic and size of employer preferences also appear to be only slightly related to length of time to secure employment.

The academic major variables are significant as a set; engineering related majors are becoming employed sooner. On the other hand, those who see their majors or their degree from Trinity as being valuable take longer to become employed.

The most important set is the employment seeking behaviors set. The interpretation of that set is not particularly encouraging to those who are having difficulty finding employment. Those who take longer to find employment tend to send more resumes, have more interviews, and then get a smaller salary. But then, who said life was fair?

Although they are not shown here, other predictor variables were thought relevant and used in an initial analysis. Such variables, discarded because they showed little relationship to the time criterion included: relatedness of major to field of employment, grade point average, marital status, perceived number of jobs available in the major, and necessity of seeking employment.

References

Williams, J.D. and Lindem, A.C. Setwise regression analysis; A stepwise procedure for sets of variables. Educational and Psychological Measurement, 1971a, 31, 747-748.

Williams, J.D and Lindem, A.C. Setwise regression analysis - A new data analytic tool. Multiple Linear Regression Viewpoints, 1971b, 2, 25-27.

Wolfbein, S.L. Seven signs for the seventies. Journal of College Placement, 1970, 31, 22-27.