

Some Pitfalls in Large Econometric Models: A Case Study

I. Introduction

In this paper, I am going to discuss some of the difficulties in large-scale econometric modeling. As a stalking-horse, I will use the Regional Energy Activity and Demography (READ) model. This model was based on work by Harris and Hopkins (1972); development continued under the aegis of the Energy Information Administration, Department of Energy, Washington, D.C. By 1978, with development still incomplete, costs were estimated at over a million dollars, and a review was undertaken; this preceded any use of the model by the agency.¹ As a consequence of the review, work on the model was stopped.

* I would like to thank my colleagues David Brillinger and Thomas Rothenberg for many hours of helpful conversations about econometrics. Needless to say, while they share the credit for the sensible arguments in this paper, they bear no responsibility for the other ones. I would also like to thank a very patient and helpful referee.

1. The agency has begun an active program of external review for its data and models. This particular review was requested by Dr. Lincoln Moses, the administrator; it was organized by Dr. George Lady, Office of Analysis Oversight and Access. I was one of the reviewers. The description of READ is based on material developed in the review, including model documentation (Donnelly et al. 1976; Gamson et al. 1977; Havenner and Donnelly 1977; Tannen 1979). Parts of this paper were presented in the review, and an early draft was given at the National Bureau of Standards Workshop on Validation and Assessment of Energy Models, 1979.

This paper points to some of the difficulties in large-scale econometric modeling. The major issues here include the logic of the equations, the nature of the stochastic disturbance terms, the logic of the fitting procedures, and the quality of the underlying data. One specific model is reviewed as an example: the Regional Energy Activity and Demography model developed for the Department of Energy.

The READ model is something of an outlier in the population of econometric models. Hasty generalizations should therefore be avoided. On the other hand, much can be learned from studying such an extreme case, where several issues come into sharp focus. Further more, though an outlier, READ is a recognizable member of its population. In my opinion, even the best econometric models may be found to share some of READ's weaknesses.

This assessment of READ will focus on the following issues: the logic of the equations, the nature of the stochastic disturbance terms, the logic of the fitting procedures, and the quality of the underlying data. The READ model is a large-scale annual econometric model of the United States, with very fine detail. The basic geographical unit is the county, and there are nearly 50 industries in the present version of the model. The object is "regional impact analysis," that is, analysis of the impact of energy policy on regional economic activity. The model has four basic sectors: (1) industrial location (output by industry and county); (2) demography (employment by industry and county, labor force, and population by county); (3) construction (over 100 building types, by industry and county); and (4) government (local, state, and federal). Energy prices appear as explanatory variables in most of the econometric equations. There is also a linear programming transportation submodel, which in effect moves industrial output over a transportation network with about 500 nodes. Shadow prices from this submodel are used as explanatory variables in the industrial location sector.

Macroeconomic variables (GNP and its major components) are treated as exogenous, and can be supplied to READ from any of the standard models of the U.S. economy. Regional energy prices, too, are exogenous and can be supplied to READ from Energy Information Administration models. In effect, READ disaggregates the macro-level forecasts down to the county level. Then, it is possible to aggregate back up to any desired geographical level. The fitting period for the model is 1965-74, and forecasts will be made out to 1990.

II. A Closer Look at the READ Model

This section describes the following components of the READ model: the industrial location sector, the demographic sector, and the transportation submodel.

Industrial Location Sector

"Industry" is a catchall term: READ industry no. 1 is "agricultural production," and READ industry no. 40 is "wholesale and retail trade." The READ industries are defined in terms of the usual Standard Industrial Classification (SIC) codes. For instance, READ

industry no. 1 corresponds to SIC codes 01–02, and READ industry no. 40 corresponds to SIC codes 50–59. The present version of the model has 47 industries, a private household sector, and several government sectors (local, state, and federal).

“Industrial location” is a bit of a misnomer. This sector predicts annual changes in the value of output (e.g., sales) by READ industry and county, using a linear regression equation. A “change” is the difference between the value of output for the current year and for the previous year. Value of output is denominated in 1967 dollars. The independent variables are: value of output in the previous year; capital stock; wage rate in the industry; land prices; energy prices; transportation costs (shadow prices from the transportation submodel described below); tax rates; macroeconomic variables (GNP, PCE, etc.); agglomeration variables,² and environmental variables (weather).

Demographic Sector

This sector has equations to predict employment, unemployment, wages, and population. The first equation in the sector predicts changes in the level of employment in an industry in a county using a linear regression equation. The explanatory variables are: level of employment, change in output, change in capital stocks, level of capital stocks, energy prices, and a technology index (not further defined). The other equations in this sector are similar in conception and will not be discussed here.

Transportation Submodel

In each county, there are two sources of supply for the output of READ's industries, and seven sources of demand: (1) *supply*: imports, industrial output; (2) *demand*: exports, interindustry demand, personal consumption, construction, equipment investment, state and local government, and federal government. Counties are grouped into the 473 Department of Transport regions. Supply and demand are summed algebraically over each region, which becomes either a net supplier or a net demander of each type of output. In effect, a linear program then ships the various industrial outputs over the network, the shipping costs being derived from the Census of Transportation (1967, 1972) and from Interstate Commerce Commission data.

III. The Logic of the Equations

This section presents a series of criticisms of the equations in the READ model. The points are all straightforward but should not be

2. Not further specified in the READ documentation (Donnelly et al. 1976; Gamson et al. 1977; Havenner and Donnelly 1977; Tannen 1979), but see Harris and Hopkins (1972).

brushed aside for that reason, nor should they be dismissed as "academic." Indeed, when faced with a model like READ, the key question to ask is, Why should we believe the predictions? With READ, the answer to this question is almost necessarily an argument *a priori*: The forecasts are very difficult to compare with observations, because READ mainly predicts unobservable quantities (this is discussed in detail in Section VI, below). In an argument *a priori*, only sound reasoning compels conviction. On this count, READ does not do so well.

The first point to discuss is the "cancellation of errors" defense for modeling in fine detail. There is usually little interest in predicting the output of one industry in one county, and much aggregation will be done before forecasts are made. So each equation in READ may be suspect, the argument goes, but the errors tend to cancel out during aggregation.

Cancellation of errors is the grand illusion of statistics. If errors are independent, they do tend to cancel (on a relative basis). If the errors are correlated, they pile up; aggregation, then, can increase the relative error. The errors in READ are likely to be correlated, as shown below. The cancellation argument is therefore quite shaky.

In general, I see no theoretical argument in favor of great detail in regression models, and much can be said against.³ This may seem paradoxical. For instance, a large-scale pilot plant is surely more predictive than a small-scale one. However, a regression model is not a microcosm described in algebraic notation. The prudent view is that at best each equation represents an approximation to some more or less stable statistical relationship among the variables. Clearly, the multiplication of statistical half-truths is a dangerous business.

The next point has to do with the fact that the basic geographical unit in READ is the county. But the county is not a natural unit of economic analysis, since county economies interact in complex ways. For example, take my home Standard Metropolitan Statistical Area (SMSA), San Francisco–Oakland. This comprises five counties: Alameda, Contra Costa, Marin, San Francisco, and San Mateo. "Wholesale and retail trade" is READ industry no. 40. Changes in the level of employment in this industry for each county are postulated by the model to depend only on other variables describing that same county. However, the reality is very different, since people drive heedlessly across county lines to do their shopping.

To make this more definite, suppose the federal government closes an army base in Marin. This exogenous shock will depress wholesale and retail trade in all five counties for some years. In the model, this

3. In practice, complex regression models seem to do neither better nor worse at forecasting than simple ones. For a discussion of predictive accuracy in forecasting, see McNees (1973, 1974, 1975, 1977, 1979); Christ (1975); Ascher (1978); and Zarnowitz (1979).

shock has to be absorbed into the stochastic disturbance terms, because the explanatory variables for wholesale and retail trade in one county do not include army bases in another county. Indeed, there are no variables in the equation to capture the impact of one county on another except the shadow prices from the transportation submodel which appear in the industrial location equation. The estimation strategy of READ precludes other choices, as will be seen in the next section. There are no variables to capture the impact of one industry on another except the "agglomeration variables" in the industrial location equation; and nothing is said in the documentation about the strategy for defining or using these variables.

Next, despite the complexity of the model, many important variables are omitted from the equations. For instance, the prices of an industry's major inputs or outputs are not used as explanatory variables in the equations predicting output or employment. Indeed, prices do not appear in the READ database at all, apart from the ones mentioned in Section II.⁴

In any large econometric model, there will be omitted variables. Sometimes this is because no data are available (in READ, however, this constraint was not binding). Sometimes variables are omitted to conserve degrees of freedom. In any case, the statistical impact of omitted variables is bias in the estimates of the coefficients.

The next point is that READ uses the same linear functional form for all industries. There is a separate equation for each industry, but in each such equation the same coefficients are used for all counties. This is illogical. Industry no. 47, for instance, is "health, legal, and other services." In Alameda County, California, this industry almost comes down to the University of California, Berkeley. In another county, the main component of this industry might be a private university or a hospital. The industry will show different economic behavior in different counties and this should be described by different equations. The same argument applies to virtually all of READ's industries.

Also, there is a serious endpoint problem in using linear equations with coefficients constant across counties when an industry has been dormant in some counties over the fitting period. For example, take the industrial location equation, where the dependent variable is the change in value of output; explanatory variables include GNP and energy prices. The GNP almost has to come into the equation with a positive coefficient, the same for all counties. Suppose the model starts forecasting, in a scenario for economic expansion. The GNP term on the right-hand side of the equation forces positive changes in output. The base is zero, so we get positive output. Dormant industries spring into action all across the United States, driven by the expanding

4. How is the value of output computed? This question will be answered in Section VI.

economy. The equations grow wheat in Boston and strike oil in Berkeley.

Of course, this economic renaissance could be cut off by an increase in energy prices; for many industries, these must come into the equation with negative coefficients. If the scenario makes OPEC too exigent, the industrial location sector could easily predict negative changes in output. For a dormant industry, the base is zero. Are we ready for the concept of negative output?

The magnitude of these effects is hard to judge. However, economic activity is highly concentrated; for instance, 10% of the counties have 80% of the manufacturing. As a result, in any industry, there will be many, many counties with output either zero or slightly positive over the fitting period. In effect, READ is fitting a straight line to a function which is clearly nonlinear, and then using the line in the vicinity of that nonlinearity.⁵

Next, many equations in READ are autoregressive. The employment equation, for instance, has the form

$$L_{ijt+1} - L_{jt} = \beta L_{jt} + \text{other terms.}$$

Here, L_{jt} is the level of employment in the industry in county j and year t . Consider forecasting with such an equation in a scenario where all the explanatory variables (GNP, energy prices, etc.) are constant. Now, L_{jt} should converge to its equilibrium value as time goes on, and it would if the coefficient β were slightly negative. However, the economy was expanding over the fitting period 1965–74, so β is expected to be positive: +0.05 is a typical value in preliminary fits. As a result, L_{jt} can explode geometrically fast to either plus infinity or minus infinity, depending on the coefficients and on the initialization of the other variables. Over a 15-year forecasting period, changes in L_{jt} on the order of 50% may be expected, even with all other variables being held constant. This is quite paradoxical.⁶

To sum up, the equations in READ—and in many other large-scale models—have paradoxical implications, especially when they interact under circumstances different from the ones which obtained during the fitting period. A typical defense is that the equations only represent local approximations. Perhaps they do. However, econometric models

5. The READ group comments that if an industry was dormant in a county over the fitting period, this county is dropped from the data set before fitting the equation. They also write: "... the simulation routine we are developing contains a broad set of initial conditions and checks with regard to applying each equation correctly when forecasting. One of these checks is to ensure that employment changes are not forecast unless industrial production is present. There are additional checks present for alternative situations in which impossible contingencies would otherwise occur." No further details are available. However, since the model is more complex than a set of simultaneous linear equations, the rationale for the statistical procedures used to estimate the coefficients is further weakened.

6. For a discussion of explosive autoregressions, see Anderson (1959).

are often used to extrapolate into the distant future. The READ model, for example, was fitted on data for 1965–74 and extrapolates out to 1990. For this purpose, local approximations are not good enough. Long-term forecasts may reflect all the paradoxical implications built into the equations.⁷

IV. The Logic of the Fitting Procedure

Pooled time-series, cross-sectional regressions are used in READ. To be more definite, consider the employment equation for a READ industry (like wholesale and retail trade). Let L_{jt} be the level of employment in this industry in county j and year t , and V_{jt} the value of output (sales). Let D be the first difference operator: $DX_{jt} = X_{jt+1} - X_{jt}$. The READ employment equation is

$$DL_{jt} = \beta_0 + \beta_1 L_{jt} + \beta_2 DV_{jt} + \beta \cdot U_{jt+1} + \delta_{jt+1}, \quad (1)$$

where β_0 , β_1 , and β_2 are scalar parameters; β is a vector of parameters; U_{jt} is a vector of explanatory variables describing county j in year t ; and δ_{jt} is a stochastic disturbance term. Notice that the parameters do not depend on j or t : They are constant across counties and years, as stated in the previous section.

The READ strategy is to pool the observations from all 3,000+ counties (indexed by j) and all 10 years of the fitting period (indexed by t), fitting equation (1) to the resulting 30,000+ data points. This pooling is what forces the coefficients in the regressions to be constant across counties. And this pooling obliterates all the economic and demographic interrelationships between counties.

In its present form, READ uses ordinary least squares. The estimates β are chosen to minimize:

$$\sum_{jt} (DL_{jt} - \hat{\beta}_0 - \hat{\beta}_1 L_{jt} - \hat{\beta}_2 DV_{jt} - \hat{\beta} \cdot U_{jt+1})^2. \quad (2)$$

The δ 's are correlated with each other and with the explanatory variables due to the county interactions and other specification errors discussed above. So these estimates for the coefficients, and the associated standard errors, are biased. Since equations were developed by retaining variables whose estimated coefficients were significant and had the expected sign, even the choice of explanatory variables is suspect.

The READ group proposed to get around these problems by using a variant of two-stage least squares, along the lines discussed in Brundy and Jorgenson (1971). However, this method cannot be expected to succeed either. Indeed, the two main assumptions of the proposed

7. On this point, see, e.g., pp. 241–59 of Zarnowitz (1972). Also see Christ (1975) and Lucas and Sargent (1979).

fitting method⁸ are as follows: (1) Within a sector, disturbance terms corresponding to different counties are uncorrelated; and (2) across sectors, disturbance terms are uncorrelated. These assumptions are not very plausible. In brief, one main source for the disturbance terms is specification error, which will be correlated with everything.

To make this more vivid, suppose that in year $t + 1$ the Russians have a bad harvest and buy large quantities of American wheat. What are the implications for the employment equation (1) in agriculture? For many agricultural counties, DL_{it} and DV_{it} will be large, but just how large and in what relationship must depend on factors such as the proportion of agricultural production accounted for by wheat, the spare capacity, etc. These factors will vary systematically from one county to another. So even if equation (1) were correct, county by county, forcing the coefficients to be constant across counties would correlate the errors with the explanatory variables and across counties.⁹ Since DV_{it} is the dependent variable in the industrial location equations, Russian grain purchases turn up in the stochastic disturbance terms of that sector. And since the equations are autoregressive, the correlation spreads through the variables over time.

This is only one example. But any exogenous shock will propagate itself through the system in a similar way. Clearly, the assumptions made in the READ model about the stochastic disturbance terms were not designed to stand up to rigorous scrutiny. It may even seem unfair to be examining them so closely. However, such assumptions can have a material impact on the estimates of the coefficients and hence on the forecasts.

Few large econometric models have explicit mechanisms for generating the stochastic disturbance terms. As a result, the assumptions made about these terms are usually ad hoc. Of course, to some extent the assumptions could be tested empirically. However, it is uncommon to see any serious data analysis done on the residuals, which remain the orphans of econometric work. I saw no evidence of serious data analysis in the READ project.

Some readers may find this discussion too metaphysical. After all, whatever the δ 's may be, it is feasible to do the minimization in equation (2) and come up with the estimates $\hat{\beta}$. But what do these estimates tell us? The crucial point is the concentration of economic activity, noted above. In particular, a plot of the data will show a large

8. See Havenner and Donnelly (1977).

9. Note the "if"; the assertion is almost a theorem. Suppose, e.g., that $y_i = a_i x_i + \epsilon_i$, where the ϵ 's are independent among themselves and independent of the (random) x 's, and have mean zero. Suppose too that the (nonrandom) a_i 's show variability in i . Now write $y_i - a_i x_i = \delta_i$. Clearly, $\delta_i = (a_i - \bar{a}) x_i + \epsilon_i$. So δ_i is correlated with x_i . And, if the x 's are correlated among themselves—neighboring counties respond similarly to Russian grain purchases—the δ 's will be correlated among themselves too. Of course, the ordinary least-squares estimate is thrown off by the puzzling term, $\sum_i (a_i - \bar{a}) x_i / \sum_i x_i^2$.

cloud of points near the origin for the small counties with a few points straggling off to represent the big counties.¹⁰ The regression plane will be largely determined by these outliers. What β does, then, is to measure the contrast between the many small counties and the few big ones. This contrast between the two groups of counties does not seem relevant in predicting how the economies of the counties in either group will evolve over time or respond to energy policy.

It is not uncommon to pool temporal and geographical variation. Often, however, this is a mistake: the causes for the two kinds of variation can be quite different. Usually the functional form of the relationships among the variables is unknown and the equations in the model are no more than reasonable guesses. Under such circumstances, it is very difficult to control for confounding factors. The coefficients in the equation can turn out to be measuring the effects of omitted variables rather than the effects of the variables in the equation. This seems to be the case for READ.

V. The Transportation Submodel

The following points provide a brief critique of the transportation submodel:¹¹ (1) The model only covers truck and rail shipments; air-borne and waterborne shipments are neglected and so are pipelines. (2) Department of Transport regions may reflect the nodes in the highway system, but they are not well related to the rail system. (3) All within-zone shipments are eliminated by netting out. This eliminates cross hauling and distorts the optimization. (4) The netting out is quite unrealistic due to the inhomogeneity of READ industries. For instance, READ industry no. 1 makes both apples and oranges. However, some regions will export apples and import oranges, and these two transactions do not cancel—for the usual reasons. (5) Each industry's output is handled separately. So it is impossible to consider joint shipping schemes, where, for example, a truck exports some output from READ industry no. 1 and imports some output from READ industry no. 2 on the return trip. The reason is that the output of each industry is denominated in dollars, and no conversion factors from dollars to physical units (tons, cubic yards) are available. (6) The solution to a transportation problem is usually degenerate: The dual problem will then have multiple solutions, no one of which gives appropriate shadow prices.

To sum up, the submodel does not seem to be a good representation

10. For this purpose, size could be measured by value of output, or labor force, or even the first differences. These measures will all be quite highly correlated.

11. This section paraphrases reports of others, esp. Karen Polenske, presented during the READ review. Two useful references are Aucamp and Steinberg (1978) and Fjeldsted and South (1978).

of the transportation process. Therefore, it is questionable whether shadow prices from this submodel are appropriate explanatory variables.

VI. The Data Problem

There are very few solid, relevant county-level data, and this situation is unlikely to change in the next few years. As a consequence, most of the data fed into the READ regressions are synthetic; the term of art is "allocated." The accuracy of the allocation procedures is usually impossible to assess, because there is usually no standard of comparison. The use of synthetic data is crucial to the READ model. The allocation procedures will be explained in this section, and then the epistemological status of the main variables in the model will be reviewed. The statistical issues created by allocation will be discussed in Section VII.

Probably the single most important variable in the READ model is industrial output, which is measured by "value of shipments" (e.g., sales).¹² However, value-of-shipments data are collected at the county level only by the Census of Business at 5-year intervals (1967, 1972, etc.). The model needs annual data from 1965 to 1974. These data are derived by an allocation procedure which varies a bit from industry to industry. Interestingly, the Census of Business data are not used, even in the years for which they are available.

To be perfectly definite, take READ industry no. 26 (SIC code 35), which manufactures "machinery, except electrical." Focus on Santa Clara County, California, in the year 1972. County-level payroll data by industry are available from the Bureau of Economic Analysis, Department of Commerce; caveats to these data are discussed below. State-level value-of-shipments data for manufacturing industries are available from the Annual Survey of Manufactures, Bureau of the Census. The payroll data are used to allocate the state-level value-of-shipments data to the county level. More precisely, the value of shipments by READ industry no. 26 in Santa Clara County in 1972 is taken to be the California state value of shipments in the industry in 1972 times the fraction: Santa Clara County payroll in the industry in 1972 ÷ California state payroll in the industry in 1972. The key assumption in this procedure is that the ratio of outputs to labor inputs is constant across counties. This assumption is not plausible. Since READ industry no. 26 is quite heterogeneous, the parts of it in different counties are likely to have quite different labor productivities. This point will now be discussed in more detail.

12. An interesting twist is that the production functions behind the READ equations do not have physical inputs; the focus is on value added. But since there are no data on value added, value shipped is used in fitting the equations.

TABLE 1 Labor Productivity (Value of Shipments/Payroll) in Several California Manufacturing Industries in 1972

SIC Code	Name	Productivity
352	Farm and garden machinery	4.4
357	Office and computing machines	2.9
35	Machinery, except electrical	3.1

By a quirk of fate, "machinery, except electrical" includes both farm machinery and computers. Santa Clara County is known locally as "silicon gulch." It specializes in computers, not farm machinery, and these two branches of READ industry no. 26 exhibit very different economic behavior. In 1972, the farm machinery branch had much higher labor productivity (see table 1).

Table 2 shows what happens if a READ-type allocation procedure is used to create value-of-shipments data for READ industry no. 26 in 12 California SMSAs: The San Jose SMSA is Santa Clara County. The allocated data are quite good (on a percentage basis) in Los Angeles. They are not as good in San Jose or Fresno, the errors being about 10% and 25%, respectively. San Jose makes the computers (SIC code 357, labor productivity 2.9). Fresno is in the San Joaquin valley, a rich agricultural area, and specializes in farm machinery (SIC code 352, labor productivity 4.4).

TABLE 2 READ-Type Allocated Value of Shipments for "Manufacturing Machinery, Except Electrical" in 12 California SMSAs in 1972

SMSA	Actual	Allocated	% Error
Los Angeles	2,148	2,115	-2
San Jose	738	805	10
San Francisco	466	455	-28
Anaheim	397	403	-2
San Diego	261	291	11
Riverside	83	79	-5
Fresno	66	49	-26
Oxnard	47	40	-15
Sacramento	28	23	-18
Salinas	18	15	-17
Bakersfield	12	10	-17
Modesto	6.7	7.2	7

SOURCE.—The starting point for table 2 is the 1972 *Census of Manufactures*, vol. 3, *Area Statistics*, part 1, Alabama-Montana. Table 5 in this publication gives state figures for payroll and value of shipments, while table 6 gives the SMSA figures; the first 12 SMSAs in table 6 are used. The census lists the SMSAs in alphabetical order, while table 2 above reorders them by size. The allocation in table 2 is to SMSA rather than county level, because the census does not publish county-level data. Data from ASM and BEA were not used. The actual READ allocations, using ASM and BEA data and going down to the county level, must be worse.

NOTE.—% error = (allocated - actual)/actual. The recipe used for the allocation is SMSA payroll/state payroll $\times$ state value of shipments; units = millions of 1972 dollars.

The choice of READ industry no. 26 was arbitrary, for illustrative purposes only. The allocation procedure is the same for any manufacturing industry. For nonmanufacturing industries, the value of shipments by state is unknown and the payroll data are used to allocate national value-of-shipments data to the county level.¹³ The inhomogeneity of READ industry no. 26 is quite typical, too: The economy is just too complicated to divide up into 50 or 100 homogeneous "industries."

Would more industrial detail help? Probably not. For one thing, three- and even four-digit SIC industries still turn out to be quite heterogeneous. For another thing, the number of firms in each county and industry would get quite small, creating serious instabilities of another kind.

The payroll data (wages and salaries, employment levels) constitute the key data set in the READ model. The payrolls are the basis for allocating value of shipments, investments, even exports and imports. The employment levels by county and industry are crucial to the demographic sector of the model. A close look at the payroll data is therefore in order. Unfortunately, the Bureau of Economic Analysis (BEA) does not provide adequate documentation for its procedures.¹⁴

The BEA claims that about 80% of the payroll data (by dollars) is based on administrative records—state unemployment insurance reports. At the county level, however, this claim is suspect. Large corporations file only one unemployment insurance report for each state in which they operate. This report gives the total payroll and the total employment count for the entire state. It does not give any county-level detail. The state totals must then be allocated to counties and industrial activities within counties, because corporations often have activities covered by several different two-digit SIC codes. The allocation is done either by the state agency or by the BEA, apparently on the basis of the reported county totals for the small firms which operate only in single counties. For example, a chain like Sears could be distributed across counties like all the other corner dry-goods stores.

The Bureau of Economic Analysis admits that 20% of the payroll data is allocated to counties from state or national totals. The sectors of the economy particularly hard-hit by allocation are agriculture, transportation, services, and government. Agricultural payrolls, for instance, are allocated to counties using shares from the 1969 Census of Agriculture. Thus, regional differences in the response of agriculture to

13. The three construction industries (READ nos. 8,9,10; SIC codes 15,16,17) are exceptions to this rule, as is anthracite mining (READ no. 3, SIC code 11).

14. The best documentation is the introduction to U.S. Department of Commerce (1977). Subsequent quotations are from this source.

increasing energy prices are not captured in the READ database. The allocation procedure for transportation should be read in full. I quote only the discussion of railroads:

County estimates of wages in the railroad industry were based on the biennial employment series for Class 1 railroads (line-haul and switching and terminal companies) prepared by the Association of American Railroads (AAR). These AAR data are for selected SMSA counties and account for approximately 73 percent of total railroad employment. For the remaining counties, AAR's residual-counties employment estimate in each State was disaggregated in proportion to the railroad employment reported by counties in the 1970 Census of Population. Estimates for the intervening years were derived by averaging the biennial benchmark data. The resulting employment series was used to allocate the State totals of wages and salaries in the railroad industry. The 1975 AAR data were available for distributing the 1975 State totals.

Private household workers, to take one clear example in the service sector, are distributed to counties according to the 1970 Census of Population. Turning to the public sector, local government payrolls are estimated by linear interpolation between the 1967 and 1972 Census of Governments. State government payrolls are estimated from the 1967 census only, "since the 1972 Census of Governments did not include the information necessary to prepare a new benchmark." Federal government payrolls are allocated from states to counties by year-end employment head count, because federal government agencies do not seem to know their payroll by county.

One thing is clear. The payroll data, which are used to allocate so many other variables, are themselves the result of a complex allocation process performed by the BEA. Most of the variables in the READ model are suspect, having been derived by allocation procedures similar to the ones described above.¹⁵ This also applies to the dependent variables. As a result, it would be almost impossible to judge the accuracy of READ forecasts by comparing them to observations, since most of the equations predict variables which are measured only in census years, if at all.

I am very critical of the READ model because its database is so poor. However, there are some widely used econometric models fitted to data which are not much better.¹⁶ Econometricians seldom pay enough attention to the quality of their data, and this seems to me to have unfortunate consequences for the models.

15. Further examples are discussed in Appendix A.

16. For a discussion of one example, see Freedman, Rothenberg, and Sutch (1980a, 1980b).

VII. Statistics of Allocated Data

Consider the model $y = \beta x + \epsilon$, where ϵ is a stochastic disturbance term independent of x . The parameter β can be estimated, as usual, by the regression of y on x . Suppose, however, that neither y nor x is directly observable and that they are estimated by $\hat{y}$ and $\hat{x}$, respectively. The regression of $\hat{y}$ and $\hat{x}$ may—or may not—give a good estimate of β . This is the “errors in variables” problem, and it is central to READ, because almost all the data are allocated. This problem will be discussed in some detail in Appendix B. The analysis will perhaps be oversimplified, but the results are suggestive. The main conclusion is that synthetic data behave quite differently from real data when regressions are run. In the context of the READ model, this conclusion can be made more specific:

1. With 10 years of data and 3,000+ counties, there appear to be 30,000+ degrees of freedom. When the data are allocated, however, this is a gross exaggeration. For example, suppose all the data are derived from national totals using the same allocator. Then the coefficients from a county-level regression based on the allocated data are the same as the coefficients that would be obtained from a regression using the national totals. In fact, there are only 10 degrees of freedom in the allocated data, and the READ-type estimates for standard errors are too optimistic by a factor of $\sqrt{3,000} \approx 50$.

2. When different allocators are used for different variables in the equation, as is typical in READ, a substantial bias may be introduced in the estimates due to the errors in the allocation. Usually, the magnitude of the bias cannot be estimated from the data.

3. Errors in one variable can bias the estimates for the coefficients of other variables.¹⁷

VIII. Summary and Conclusions

The READ model represents an ambitious effort to model county-level economic activity, with attention to the impact of energy variables. The modelers undertook a difficult assignment and worked hard to complete it. However, the prospects of success were remote, for four reasons:

1. There is no adequate theory to guide the choice of equations. The READ equations seem mechanical and in many cases unrealistic.

2. The county is not a natural unit of economic analysis, because counties interact with each other in complicated ways.

3. There is no rationale for the fitting procedures used in the READ model.

4. At the county level, the database is so shaky.

17. Some calculations are presented in Freedman (1980*b*).

Modelers sometimes argue as follows: "It's an imperfect world, we have to do the best we can. Policymakers want numbers, so we need a model. The equations may not be perfect, but they're the best we can come up with. We know there are problems with the data, but it's the best we have. If you don't like what we're doing, just tell us how to do it better. Besides, if we don't develop the model, some other agency will, and they'll do it worse." The result is a variant of Gresham's law, in which bad analysis drives out good, and a fog of misinformation settles over the policy process.

One lesson of READ is that in an imperfect world, making the best possible model may be a waste of time. After all, if you do not believe the equations or the assumptions about the stochastic disturbance terms or the data, why should you believe the forecasts? When the basic theory is incomplete, or the data sparse, ad hoc analysis by experts may be better than a large-scale econometric model.¹⁸ In some cases, it may be still better to tell the policymaker that his question is unanswerable. This might prompt a search for policies which do not depend on knowing the unknowable.

Appendix A

Some Further Detail on READ Data

Energy prices. Coal prices do not seem to be in the READ database. Electricity and natural gas prices were obtained from the Electric Power Research Institute by utility district, not by county. They were converted to county level by an unspecified procedure. Fuel oil prices were obtained from the Federal Energy Data System (FEDS) database, the original source being Platt's Oilgram. Platt's surveys prices in 46 cities. State prices were constructed for FEDS by averaging Platt's estimates for the survey cities in that state. If there were no such cities, the average price for neighboring states was used. Within each state, the READ database takes the price of fuel oil to be the same for all counties.¹⁹

Investment. New and replacement equipment investment are handled separately. In manufacturing industries, replacement investment is allocated to states using Annual Survey of Manufactures (ASM) total investment shares, and then to counties by payroll shares. In nonmanufacturing industries, re-

18. The available evidence suggests, for instance, that judgmental forecasts of the economy are as good as model-based forecasts (Zarnowitz 1979). As is customary, the tables in that paper are based on medians, which might in principle distort the comparisons due to variations in sample size. Professor Zarnowitz was kind enough to redo the tables, first combining the errors across time within forecasters and then across forecasters. This made very little difference to the results. Of course, the judgmental forecasts by experts are often based partly on the results from the models—and model-based forecasts usually incorporate many judgmental factors—which certainly clouds the issue.

19. For a discussion of FEDS, see Freedman et al. (1980a).

placement investment is allocated directly to counties by payroll shares. As a consequence, in any particular nonmanufacturing industry, either all counties where the industry operates show positive replacement investment or all are negative. New investment is allocated from national totals by construction shares.

Interindustry demand. Annual data do not exist at the county level. The READ model creates such data from the industrial-output-by-county "data" discussed in Section V. Given a figure for the output of an industry in a county, the 1967 Bureau of Labor Statistics national input-output table is used to compute what the demand for intermediate goods "must" have been. This ignores all regional differences in technology. It also ignores all changes since 1967 in relative prices—for instance, the price of labor relative to energy.

Personal consumption expenditures (PCE). Comprising about two-thirds of GNP, PCE is the largest source of demand for READ's commodities. However, PCE is not measured at the county level. The READ model creates the data by a very elaborate procedure whose starting point is the 1972 Census of Retail Trade. This census collected county-level data on sales for 10 different types of retail outlets and SMSA-level data on sales for about 200 merchandise lines. These data are used to construct county shares of PCE by Bureau of Economic Analysis (BEA) consumption category. National figures for PCE are then shared down and converted to county demands for READ's commodities by the 1967 input-output table. However well this worked in 1972, the READ database ignores any regional trends in consumption: The raw data are available only in 1972. In forecasting mode, national PCE is exogenous and is shared down to county level by an econometric equation. The equation, however, is fitted to the "data" just described.

State and local government variables. This sector of the model is quite elaborate, though the documentation is rather sketchy. State and local government activities appear to be endogenous. Variables include expenditures by function, revenues by type, tax rates, and indebtedness. County-level data of this kind are available for local governments from the Census of Governments (1967, 1972). Also, the Bureau of the Census has an annual survey of local governments which provides data for sample counties; however, there is no sensible way of estimating the behavior of individual nonsampled counties, and there are no data on state government expenditures by county, even in census years. County-level data are needed both for the government sector and for the transportation submodel. According to the documentation (Gamson et al. 1977), READ creates the data as follows: "State level controls for state and local government expenditure were obtained for the Morlan model data base at BLS. . . . For the preliminary model, state and local controls will be allocated to counties using [1967 and 1972] Census of Government distributions *for local government only*" (emphasis added).

Federal government. The federal government is almost totally ignored by READ; it comes into the database only through the BEA payroll data. As the READ documentation notes, "Data from the Federal Government has not been consolidated to date."

Appendix B

Statistics of Allocated Data

Consider first a regression with one explanatory variable and no constant term—for instance, of output on investment. (This equation does not appear in the READ model, but it illustrates an important point and keeps the algebra within bounds.) Fix one industry, say, “wholesale and retail trade.” Let y_t = national output figure for this industry in year t , x_t = national investment figure for this industry in year t , p_{jt} = payroll in this industry in county j in year t , and

$$f_{jt} = p_{jt} / \sum_i p_{it}. \quad (\text{B1})$$

Thus, f_{jt} is the fraction of the payroll going to county j in year t , and a READ-like allocation procedure²⁰ is to use

$$\hat{y}_{jt} = f_{jt} y_t \quad (\text{B2})$$

$$\hat{x}_{jt} = f_{jt} x_t \quad (\text{B3})$$

for “data” on output and investment in county j for year t . The true—but unknown—numbers will be denoted y_{jt} and x_{jt} .

The next step is to choose $\hat{\beta}$ to minimize the sum of squares:

$$\sum_{jt} (\hat{y}_{jt} - \hat{\beta} \hat{x}_{jt})^2. \quad (\text{B4})$$

Here, j runs over all 3,000+ counties, and t runs over the 10 years in the fitting period 1965–74. Apparently, there are 30,000+ degrees of freedom in (B4). However, substitution of (B2) and (B3) into (B4) shows that $\hat{\beta}$ minimizes

$$\sum_t w_t (y_t - \hat{\beta} x_t)^2, \quad (\text{B5a})$$

where

$$w_t = \sum_j f_{jt}^2. \quad (\text{B5b})$$

In other words, $\hat{\beta}$ is the regression coefficient of the national totals y_t on x_t . The only effect of allocating data to the county level is the introduction of the weights w_t . In principle, these are data-dependent weights varying with t . They weight more heavily those years with higher economic concentration. In practice, the w_t 's are likely to be essentially constant. In any case, there are only 10 degrees of freedom available for estimating β .

The kind of stochastic model suggested by READ is

$$y_{jt} = \beta x_{jt} + \epsilon_{jt}, \quad (\text{B6})$$

where y_{jt} and x_{jt} are the true (but unknown) county-level output and investment figures and ϵ_{jt} is a stochastic disturbance term assumed to have a mean of zero. As j and t vary, these disturbances are assumed independent. Furthermore, x_{jt} and f_{jt} are going to be treated as if they were constant rather than stochastic. These assumptions simplify the analysis and are similar to those in READ.

Now for a modification. There are big counties and little ones, and presumably the likely size of ϵ_{jt} is bigger in the big counties. It may even be reasonable

20. In READ, “replacement” investment is in fact allocated by payroll shares but new investment is not.

to suppose that the likely size of the disturbance is approximately proportional to the payroll share f_{jt} , so that

$$\text{var } \epsilon_{jt} = \sigma^2 f_{jt}. \quad (\text{B7})$$

Under these conditions, it is reasonable to estimate β from a weighted regression, choosing $\hat{\beta}$ to minimize

$$\sum_{jt} (y_{jt} - \hat{\beta} x_{jt})^2 / f_{jt} = N \sum_{jt} (y_{jt} - \hat{\beta} x_{jt})^2, \quad (\text{B8})$$

where $N = 3,000+$ is the number of counties. This weighted regression on the synthetic county-level data produces exactly the same results as a regression using the observed national totals, with the spurious 30,000+ degrees of freedom deflated to the corrected 10. This conclusion applies to multiple regression when the intercept is forced to zero, the same allocator (f_{jt} in the example) is used for all the variables, and the error SD is also proportional to this allocator.

With a constant term, the algebra is a bit more complicated. The model is

$$y_{jt} = \alpha + \beta x_{jt} + \epsilon_{jt}, \quad (\text{B9})$$

where the ϵ_{jt} are stochastic disturbances with a mean of zero, assumed independent. For ordinary least squares, the estimators $\hat{\alpha}$ and $\hat{\beta}$ are chosen to minimize the sum of squares:

$$\sum_{jt} (y_{jt} - \hat{\alpha} - \hat{\beta} x_{jt})^2. \quad (\text{B10})$$

(As before, y_{jt} and x_{jt} are the true but unknown county-level figures; $\hat{\beta}$ and $\hat{\alpha}$ represent the allocated data.) Solving the usual normal equations gives

$$\hat{\beta} = \left[\sum_{jt} w_{jt} x_{jt} - \frac{T}{N} \bar{x} \bar{y} \right] / \left[\sum_{jt} w_{jt} x_{jt}^2 - \frac{T}{N} \bar{x}^2 \right], \quad (\text{B11})$$

where $T = 10$ is the number of time periods and $N = 3,000+$ is the number of counties. The weight w_{jt} was defined by (B5b) and exceeds $1/N$ by Schwarz's inequality. Finally,

$$\hat{\alpha} = \frac{1}{T} \sum_{jt} y_{jt} - \hat{\beta} \bar{x} \quad \text{and} \quad \bar{y} = \frac{1}{T} \sum_{jt} y_{jt}. \quad (\text{B12})$$

Again, there really are only 10 degrees of freedom.

The model (B9) implies

$$y_{jt} = N(\alpha + \beta x_{jt}) + \epsilon_{jt}, \quad \text{where} \quad \epsilon_{jt} = \sum_i \epsilon_{ijt}. \quad (\text{B13})$$

Taking x_{jt} and w_{jt} to be nonstochastic, the bias in using synthetic county-level data is easy to work out from (B11) and (B13):

$$\text{bias} = E(\hat{\beta}) - \beta = \alpha \frac{\sum_{jt} w_{jt} x_{jt} - (T/N) \bar{x}}{\sum_{jt} w_{jt} x_{jt}^2 - (T/N) \bar{x}^2}. \quad (\text{B14})$$

Since there are only a few large counties, w_{jt} will be appreciably larger than $1/N$, so the bias in (B14) can be considerable. An alternative is to start from (B13), using the observed national data instead of the synthetic county-level data. This eliminates bias, and the reduction in degrees of freedom is more apparent than real.

Of course, the model (B9) can be revised to incorporate the idea that a small county should have a small intercept and small error. Suppose, as is compatible

with the READ allocation scheme, that the intercept is proportional to the payroll share f_{jt} and so is the likely size of the error:

$$y_{jt} = \alpha f_{jt} + \beta x_{jt} + \epsilon_{jt}, \quad \text{var } \epsilon_{jt} \sim \sigma^2 f_{jt}^2. \quad (\text{B15})$$

Then α and β would be chosen to minimize the weighted sum of squares:

$$\sum_{jt} (\hat{y}_{jt} - \alpha f_{jt} - \hat{\beta} x_{jt})^2 / f_{jt}^2 = N \sum_t (y_t - \hat{\alpha} - \hat{\beta} x_t)^2. \quad (\text{B16})$$

Again, the weighted regression is equivalent to discarding the synthetic county-level data and running the regression on the observed national-level aggregates. This conclusion applies to multiple regression when the same allocator is used for all the variables, the intercept is proportional to this allocator, and the error SD is also proportional to the same allocator.

These conclusions do not apply when different allocators are used for different variables. In such cases, however, substantial bias may be anticipated, because we have an "errors in variables" situation. As an illustrative case, consider another regression of output on investment, where output is allocated from state totals, while investment is allocated from national totals. The new point is the interplay between the two different geographical levels.

Following previous notation, let

$$y_{st} = \text{output in the industry in state } s \text{ in year } t. \quad (\text{B17})$$

Abbreviate $s(j)$ for the state containing county j , and let

$$h_{st} = \sum_{it \in s} f_{it}, \quad (\text{B18a})$$

$$g_{jt} = f_{jt} / h_{s(j)t}. \quad (\text{B18b})$$

Thus, h_{st} is the state s share of the national payroll, and g_{jt} is the county j share of the state payroll. In this situation, the county-level output figure may be allocated as

$$\hat{y}_{jt} = g_{jt} y_{s(j)t}. \quad (\text{B19})$$

However, the county-level investment figure is still given by (B3). The model is still (B6); the intercept is forced to zero.

Two cases have to be considered, according as the regression on the county-level "data" is run weighted or unweighted. Take the unweighted case. Then $\hat{\beta}$ is chosen to minimize the sum of squares:

$$\sum_{jt} (\hat{y}_{jt} - \hat{\beta} x_{jt})^2 = \sum_{st} u_{st} (y_{st} - \hat{\beta} h_{st} x_t)^2, \quad (\text{B20a})$$

where

$$u_{st} = \sum_{j \in s} g_{jt}^2. \quad (\text{B20b})$$

The index s on the right-hand side of (B20a) runs over all 50 states, while t runs over the 10 years in the fitting period.

Note that $h_{st} x_t$ on the right-hand side of (B20a) is an allocation of the national investment figure x_t to the state s by payroll share. The weight u_{st} measures economic concentration. If the payroll for the industry in state s in year t is concentrated in one county, then $u_{st} = 1$; otherwise, $u_{st} < 1$. The minimum for u_{st} occurs when the payroll is spread evenly over all counties in the state s . So u_{st} does not seem like a sensible weight. Taking the other case for a moment,

it is possible to weight county j 's contribution to the sum of squares by $1/f_{jt}$; then u_{jt} gets replaced by u_{jt}/h_{jt} , where h_{jt} is the number of counties in state s ; this too seems an unreasonable choice of weights for a regression on national aggregates.

Coming back to the unweighted case, (B20a and B20b) is minimized by taking

$$\hat{\beta} = \frac{\sum_{jt} u_{jt} h_{jt} x_{jt} y_{jt}}{\sum_{jt} u_{jt} h_{jt}^2 x_{jt}^2}. \quad (\text{B21})$$

The model (B6) implies

$$y_{jt} = \beta x_{jt} + \epsilon_{jt}, \quad \text{where} \quad \epsilon_{jt} = \sum_{it} \epsilon_{it}, \quad (\text{B22})$$

Taking payroll and investment to be nonstochastic, the bias in $\hat{\beta}$ can in theory be computed from (B21) and (B22):

$$\frac{E(\hat{\beta})}{\beta} = \frac{\sum_{jt} u_{jt} h_{jt} x_{jt} x_{jt}}{\sum_{jt} u_{jt} h_{jt}^2 x_{jt}^2}. \quad (\text{B23})$$

In practice this expression cannot be estimated, because the state-level investment figure x_{jt} is unknown. The "error in variable" here is the discrepancy between the true state-level figure x_{jt} and the allocated figure $h_{jt}x_{jt}$. The naive "commonsense" approach of putting $h_{jt}x_{jt}$ in for x_{jt} just assumes the problem away; algebraically, this substitution makes the right-hand side of (B23) equal to one. The conclusion is that when different allocators are used for different variables in the equation, as is typical in READ, a bias may be anticipated in the estimates. Though the magnitude of the bias usually cannot be estimated from the data, it may be large.

So far, the discussion has focused on ordinary regression with weights allowed. In principle, instrumental variables can be used to get around the problem of measurement error. For READ, this would involve writing down a proper stochastic model, including stochastic equations which link the allocated values to the unobserved county-level data. The source and nature of the stochastic disturbances would have to be analyzed in some detail. Then, it would be necessary to produce some instrumental variables orthogonal to the errors. The orthogonality would have to be argued *a priori*, on the basis of economic theory, because the errors are unobservable. Omitted variables, county interactions, and other specification errors would all turn up in the disturbance terms, introducing correlations over counties and years as discussed in Section V. The allocation procedure itself spreads national aggregates (and their errors) over counties and even over years, another source of correlations. Consequently, neither lagged variables nor national aggregates are plausible candidates for instruments. Therefore, the instrumental-variable approach seems unlikely to solve the problem. The READ documentation says little of substance about this issue, because the documentation does not address the relationship between the allocated data and the data.

Literature Cited

- Anderson, J. W. 1959. On asymptotic distributions of estimates of parameters of stochastic difference equations. *Annals of Mathematical Statistics* 30:676-87.

- Ascher, W. 1978. *Forecasting: An Appraisal for Policy-Makers and Planners*. Baltimore: Johns Hopkins University Press.
- Aucamp, D. C., and Steinberg, D. I. 1978. On the nonequivalence of shadow prices and dual variables. Technical report. Carbondale: Southern Illinois University.
- Brundy, J., and Jorgenson, D. 1971. Efficient estimation of simultaneous equations by instrumental variables. *Review of Economics and Statistics* 53:207-24.
- Christ, C. 1975. Judging the performance of econometric models of the U.S. economy. *International Economic Review* 16:54-74.
- Donnelly, W. A., et al. 1976. Estimating a comprehensive county-level forecasting model of the United States—READ. Washington, D.C.: Federal Energy Administration.
- Fjeldsted, B. L., and South, J. B. 1978. A note on the multiregional multi-industry forecasting model. Technical report. Salt Lake City: University of Utah.
- Freedman, D. 1980a. An assessment of the READ model. In Saul Gass (ed.), *Validation and Assessment Issues of Energy Models*. Special publication no. 569. Washington, D.C.: National Bureau of Standards.
- Freedman, D. 1980b. Uncertainties in model forecasts. Technical report in progress. Berkeley, Calif.: Lawrence Berkeley Laboratory.
- Freedman, D.; Rothenberg, T.; and Sutch, R. 1980a. An assessment of the Federal Energy Data System. Technical report, LBID no. 202. Berkeley, Calif.: Lawrence Berkeley Laboratory.
- Freedman, D.; Rothenberg, T.; and Sutch, R. 1980b. The demand for energy in the year 1990: an assessment of the Regional Demand Forecasting Model. Technical report, LBID no. 199. Berkeley, Calif.: Lawrence Berkeley Laboratory.
- Gamson, N., et al. 1977. Users' guide to the READ regression file. Publication no. TM/EU/78. Washington, D.C.: Energy Information Agency.
- Harris, C. C., and Hopkins, F. 1972. *Locational Analysis*. Lexington, Mass.: Heath.
- Havener, A., and Donnelly, W. 1977. Estimation from a pooled time-series of cross-sections of simultaneous equations. Washington, D.C.: Federal Energy Administration.
- Lucas, R. E., Jr., and Sargent, T. J. 1979. After Keynesian macroeconomics. Federal Reserve Bank of Minneapolis, *Quarterly Review* 3, no. 2:1-16.
- McNees, S. 1973. The predictive accuracy of econometric forecasts. *New England Economic Review* (September/October), pp. 3-22.
- McNees, S. 1974. How accurate are economic forecasts? *New England Economic Review* (November/December), pp. 2-19.
- McNees, S. 1975. An evaluation of economic forecasts. *New England Economic Review* (November/December), pp. 3-39.
- McNees, S. 1977. An assessment of the Council of Economic Advisers' forecasts of 1977. *New England Economic Review* (September/October), pp. 30-44.
- McNees, S. 1979. The forecasting record for the 1970s. *New England Economic Review* (September/October), pp. 1-21.
- Tannen, M. 1979. Preliminary version of the structure of the demographic, employment, and income sector of the READ model. Washington, D.C.: Energy Information Administration. [The correct year is probably 1978.]
- U.S. Department of Commerce. 1977. *Local Area Personal Income 1970-75*. Vol. 1. Summary. PB 270 880. Washington, D.C.: Department of Commerce.
- Zarnowitz, V. 1979. An analysis of annual and multiperiod quarterly forecasts of aggregate income, output, and the price level. *Journal of Business* 52:1-34.

Bibliography

- Cooper, R. L. 1972. The predictive performance of quarterly econometric models of the United States. In B. G. Hickman (ed.), *Econometric Models of Cyclical Behavior*. New York: Columbia University Press (for the National Bureau of Economic Research).
- Hausman, J. 1974. Full information instrumental variable estimation of simultaneous equation systems. *Annals of Economic and Social Measurement* 3:641-52.
- Hausman, J. 1975. Project independence report: an appraisal of U.S. energy needs up to 1985. *Bell Journal of Economics* 6:517-51.

- Kuh, E., and Neese, J. 1979. Parameter sensitivity and model reliability. Technical report. Cambridge, Mass.: Massachusetts Institute of Technology, Center for Computational Research in Economics and Management Science.
- Kuh, E., and Welsch, R. E. 1980. Econometric models and their assessment for policy. In *Validation and Assessment Issues of Energy Models*, edited by Saul Gass. National Bureau of Standards, special publication no. 569.
- Leontieff, W. 1971. Theoretical assumptions and nonobserved facts. *American Economic Review* 61:1-7.
- Mincer, J. (ed.). 1969. *Economic Forecasts and Expectations*. New York: Columbia University Press (for the National Bureau of Economic Research).
- Morganstern, O. 1963. *On the Accuracy of Economic Observations*. Princeton, N.J.: Princeton University Press.
- Stekler, H. O. 1970. *Economic Forecasting*. New York: Praeger.
- Su, V., and Su, J. 1975. An evaluation of ASA/NBER business outlook survey forecasts. *Explorations in Economic Research* 2:588-618.
- Zarnowitz, V. 1972. *The Business Cycle Today*. New York: Columbia University Press (for the National Bureau of Economic Research).

Copyright of Journal of Business is the property of University of Chicago Press and its content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.