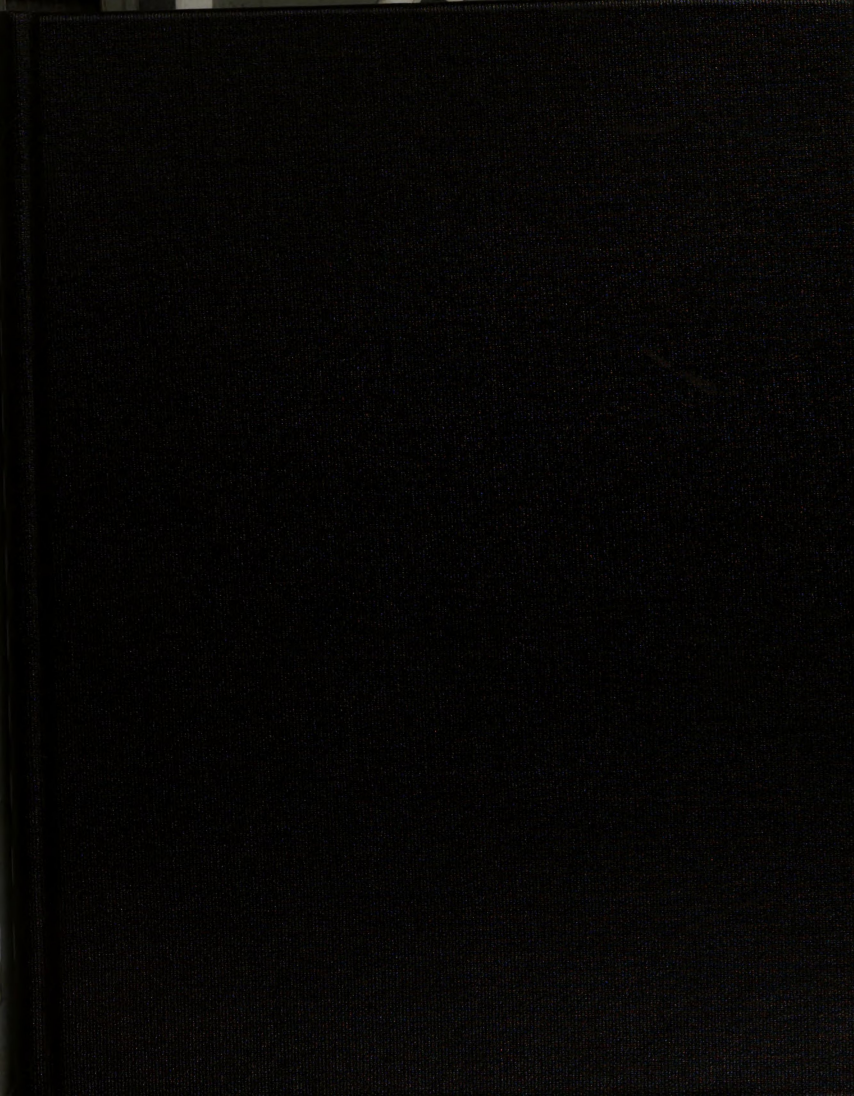




2
2002

This is to certify that the
dissertation entitled
**ANALYZING ECONOMIC MULTIPLIERS FOR THE
WOOD PRODUCTS INDUSTRIES**

presented by

PAUL R. BECKLEY

has been accepted towards fulfillment
of the requirements for

Doctoral degree in **Philosophy**

Larry A. Leekers
Major professor

Date March 18, 2002

LIBRARY
Michigan State
University

PLACE IN RETURN BOX to remove this checkout from your record.
TO AVOID FINES return on or before date due.
MAY BE RECALLED with earlier due date if requested.

DATE DUE	DATE DUE	DATE DUE
000 51 4 2003 000 51 4 2003		
JAN 0 8 2007 000 07 0 6		

**ANALYZING ECONOMIC MULTIPLIERS FOR THE
WOOD PRODUCTS INDUSTRIES**

By

Paul R. Beckley

A DISSERTATION

**Submitted to
Michigan State University
In partial fulfillment of the requirements
for the degree of**

DOCTOR OF PHILOSOPHY

Department of Forestry

2002

ABSTRACT

ANALYZING ECONOMIC MULTIPLIERS FOR THE

WOOD PRODUCTS INDUSTRIES

By

Paul R. Beckley

With anticipated increased demand for economic impact analysis to support public policy and investment decisions at all levels, steps to improve the efficiency of the economic impact analysis process must be examined. Presently, the most common method of impact analysis involves the use of ready-made interactive, input-output models. These models along with the necessary data can be very expensive and require substantial skills to use. The primary purpose of this study is to investigate the feasibility of providing more readily available, situation-specific, economic multipliers, to assist in the economic impact analysis process. The study focuses on the wood products industries and also examines the response of multipliers to changes in the size of economic impact areas.

Regression analysis techniques are used to determine the relationship between several types of multipliers and reasonably available explanatory variables including human population, population density, number of economic sectors, total industry output, personal income, and the physical size of the impact area.

Results of the study demonstrate that regression models, based on readily available information, might not be the best approach to predict economic multipliers for the wood

products industries. These models vary greatly in their ability to predict multipliers. In some specific situations the resulting models might be useful but at the same time risky. Additional study would be needed to analyze the relative merits of alternative prediction methods including costs and the extent and quality of results (e.g. statistical confidence) needed to support decisions. Study results also highlight the effect of impact area size on multipliers – as economic impact areas become relatively large, the growth of multipliers rapidly culminates. Therefore, analysts should use caution when increasing the size of impact areas to capture additional economic effects, especially in heavily populated areas with large economies.

ACKNOWLEDGMENTS

I would especially like to thank my Ph.D. advisor, Dr. Larry Leefers, for his support and advice. It was his effort and encouragement that led me back into the Ph.D. program at Michigan State University. I am especially thankful for his patience and constructive criticism in the completion of this dissertation.

In addition I would like to thank the other members of my graduate committee. They are Dr. Daniel J. Stynes, Dr. Dennis B. Propst and Dr. Karen Potter-Witter. Their technical leadership was extremely helpful.

Special thanks to my employer, USDA Forest Service, which provided valuable time and resources necessary to complete this research and dissertation. Especially helpful was the technical and moral support provided by my co-worker and good friend, Michael Niccolucci of the Forest Service Inventory and Monitoring Institute. Assistance from Dr. Susan Winter and Dr. Greg Alward of the Inventory and Monitoring Institute was also extremely valuable as was the continuing wisdom of Dr. Michael Vasievich of the Forest Service, Washington Office, Ecosystem Management group. Thanks is also due to Shirley Hendrick, of Columbia Falls, MT who did much data entry and writing-editing for this project.

Above all I would like to thank my wife PJ for all the sacrifices she has made while I pursued this degree. Her patience has been admirable. I hope I can repay her someday.

TABLE OF CONTENTS

LIST OF TABLES	ix
LIST OF FIGURES	x
CHAPTER ONE: INTRODUCTION	1
Introduction	1
Problem Statement	1
Research Objectives	6
Organization of Dissertation	9
CHAPTER TWO: LITERATURE REVIEW	10
Introduction	10
Economic Impact Analysis	11
Multiplier Analysis	13
Economic base models	14
Input-output analysis	16
Common types of multipliers	20
IMPLAN –An Economic Impact Analysis System	25
Multipliers generated by IMPLAN	28
Types of multipliers	28
Economic variables	28
Other sources of multipliers	29
The Impact Area	30
Multiplier Studies Closely Related to this Study	34
CHAPTER THREE: RESEARCH METHODS	37
Introduction	37
Economic Impact Model Selection	37
Selecting the Study Area(s)	39
Economic Sectors Analyzed	42
Impact Model Construction	48
Trade flows	50
Types of multipliers to be analyzed – types of effects	53
Types of multipliers analyzed – economic variables	54
Response Variables Analyzed	55
Explanatory Variables Analyzed	57
Other explanatory variables considered	63
General Form of model	65
Statistical Analysis Methods	66

CHAPTER FOUR: ANALYSIS AND RESULTS	70
Introduction	70
Study Objective 1: <u>Describe variations in multipliers</u>	70
Number of Observations	71
Measures of Central Tendency	72
Measures of Variability	75
Study Objective 2: <u>Identify the key factors that explain the variation of economic multipliers.</u>	77
Correlation Analysis	78
Measures of Normality	83
Study Objective 3: <u>Determination of the feasibility of using selected, readily available, explanatory variables to estimate economic multipliers</u>	84
Introduction	85
Procedures and Results	85
Regression Analysis	86
Best Subsets Process	86
Developing the Regression Equation	89
Improving the Regression Equation	93
Regional Differences/Adjustments	94
Elimination of Unusual Observations	99
Non-linear Alternatives	99
Additional Explanatory Variables	101
Multicollinearity and Heteroscedasticity	102
The Final Model	105
Examining Differences Between Geographic Impact Levels	109

CHAPTER FIVE: CONCLUSIONS	112
Introduction	112
Research Summary	113
Discussion of Results	115
Study Objective 1: <u>Describe variations in multipliers</u>	116
Study Objective 2: <u>Identify the key factors that explain the variation of economic multipliers</u>	116
Study Objective 3: <u>Determination of the feasibility of using selected, readily available, explanatory variables to estimate economic multipliers</u>	117
Policy Implications	122
Recommendations for Future Research	123
REFERENCES	126

APPENDICES

APPENDIX A - Validation of Linear Models	133
APPENDIX B - Best Subsets Regression Analysis Results	137
APPENDIX C - Regression Equation Information – Preliminary Model	149
APPENDIX D - Analysis of Variance Results	155
APPENDIX E – Study Area Map	167
APPENDIX F - States in USDA Forest Service Regions	168
APPENDIX G - Final Regression Analysis Results	169
APPENDIX H - Descriptive Statistics for Impact Area Levels	175

LIST OF TABLES

TABLE 1. Forest Products IMPLAN Sectors	44
TABLE 2. Multiplier Response Variables	45
TABLE 3 Wood Products Industry Production Functions	47
TABLE 4 Regional Purchase Coefficients	51
TABLE 5. Explanatory Variables	57
TABLE 6. Explanatory Variable Codes	57
TABLE 7. Number of Observations	71
TABLE 8. Measures of Central Tendency	73
TABLE 9. Measures of Variability	74
TABLE 10. Standard Deviation and Coefficient of Variation	76
TABLE 11. Correlation Coefficients and P-Values	79
TABLE 12. Correlation Coefficients Between Explanatory Variables	82
TABLE 13. Best-subsets Regression Analysis Results Summary	88
TABLE 14. Regression Equations – Preliminary Model Coefficients	90
TABLE 15. Effects of Regional Stratification on R^2 Values	95
TABLE 16. Multicollinearity Analysis	104
TABLE 17. Final Model Coefficients	105
TABLE 18. Mean and Maximum Absolute Percent Error	108

LIST OF FIGURES

FIGURE 1. Distribution of Multiplier Means (LE2)

77

CHAPTER ONE: INTRODUCTION

Introduction

Planning and quantifying regional economic growth or determining economic effects of public agency land management decisions has been a concern for quite some time, of local, state and national planners (McKusick 1978). The primary purpose of this study is to evaluate the feasibility of providing more convenient input-output multipliers to assist in conducting economic impact analysis for economic planning purposes.

Problem Statement

With anticipated increased demand for economic impact analysis there is also increased demand to simplify the economic impact analysis process. This pressure to simplify and streamline has been hastened by the development of new high-speed computers and software, which have been of great assistance to those already skilled in economic impact analysis. However, it has also encouraged analysts unskilled in economic analysis methods to try their hand at impact analysis and also created a mindset that economic analysis is now automated and can be done by most anyone. This has often resulted in analysts shortcutting and oversimplifying the impact analysis process with little awareness of the detrimental consequences.

The increased demand for more and better economic impact analysis exists not only at the federal government level but also the state and local government level. In addition, the private sector, which is applying for government assistance or providing the

government with analysis results, is in need of better and simpler analysis procedures and a better understanding of the analysis process itself.

The wood products industry, on which this study focuses, is an industry that is significantly affected by public agency decisions, especially those decisions by the USDA Forest Service. Consequently, changes in this industry resulting from government decisions can have significant effects on the social and economic environments of affected regions.

The purpose of economic impact analysis, at least in the case of this study, is to determine the nature and magnitude of the economic effects of the activities of the wood products industries. These activities are frequently affected by decisions of federal and state land management agencies that provide raw material to the industry. The economic impact analysis process commonly specializes in answering the main questions that have been found to be of concern to surrounding communities – how many jobs will be gained/lost and how will income be affected by changes in the wood products industries? And, more broadly, how will these potential changes eventually affect the economic structure and health of the region?

In the past, it was relatively difficult to conduct impact analysis to any degree of accuracy, and the process required extensive knowledge in the concepts of regional economics. It also required access to models and data that were few and expensive. These systems normally required the collection of primary data – taking new surveys.

Eventually secondary data based systems for use on desktop computers were developed which made analysis much faster, easier, and cheaper. These systems include IMPLAN, REMI, and RIMS II (Rickman 1995). Although these systems are a vast improvement over what was previously available, they can still be expensive and require skills not available in many smaller regions or organizations. An alternative approach has been to use ready-made multipliers (U.S. Department of Commerce 1986). However, these multipliers are normally heavily aggregated and geographically broad. Consequently, these ready-made multipliers could be significantly different for a smaller region, for a specific economic sector, or a small aggregation of sectors. In summary, there is little help at the present time, in the form of ready-made impact information or alternatives to input-output type model building for site-specific conditions. New models must be developed which require skills frequently not available at the local level.

In addition to the above problem there has been the more specific and continual problem of actually defining the region of economic impact. This is frequently affected by local officials who each want to know how their jurisdiction will be affected by proposed activities. Frequently, the analyst does not have enough information to properly determine an area appropriate to help solve the problem at hand. These conditions have lead to incorrectly designated geographic impact areas, inaccurate analyses, and consequently incorrect results. Analysts designing the impact study also might be unfamiliar with the behavior of economic outcomes or multipliers that are influenced by changes in local or regional assumptions or conditions.

An example of the above problem is where there is always pressure, from certain interest groups, or at least the tendency to expand the impact area, to capture every last possible economic effect. Of course the price the analyst or decision maker must pay for this is that the effects become diluted as the area increases. For example, are we willing to double the size of the impact area to capture a 2 percent increase in effects at the expense of diluting the locational specificity on the first 98 percent? This mistake is continually repeated because the behavior of economic multipliers is not very well understood by many analysts. This is especially critical in the wood products industry because of the complexity of the industry and its associated trade flows.

Until the past few years, the smallest unit of geography in available impact models, such as IMPLAN and RIMSII, has been the county. Consequently, most impact areas have been designated as individual counties or combinations of counties. In addition to the tendency of making impact areas too large there is the opposite and equally important error of making them too small.

The same political forces mentioned above that resulted in areas that might have been too large, could also result in impact areas that are too small. This occurs as a result of political leaders wanting to know what is happening in just their county. This creates a situation where a significant amount of the effects associated with changes in that county or county group might occur outside that area and be missed by the analyst. Part of the problem is that there is a tendency to use readily available multipliers or simple rules of thumb in designing impact areas. This occurs when the analyst does not have enough

information to design an impact area that fits the local situation. For example if the analyst does not know where the raw materials (e.g. logs) originate, where the workforce lives or spends its money, or what the byproducts of the industry might be, they might be tempted to use a ready-made region such as a county, or a Bureau of Economic Analysis (BEA) economic area or region. Occasionally, the analyst might not even know what the question is that needs to be answered. If there are not time or other resources to build a model and estimate economic multipliers and other useful information there might be a temptation to use a ready-made multiplier. These ready-made multipliers and regions will be discussed in more detail in another section of this report.

Recently, the capability of designing models at the zip code level, has been available and an irresistible temptation for some analysts. This could compound the problem of designing impact areas that are too small. However, this study will not test models at that smaller geographic level because data quality at the present is unacceptable for research purposes.

In summary, although vast improvements have been made in the process of economic impact analysis, it can still be expensive. Economic multipliers used to estimate economic effects can be generated in haste without adequate skills or information. The implications of these problems remain unknown to many analysts. The feasibility of ready-made multipliers at the local level needs to be investigated. The same type of problem exists with identifying the appropriate economic impact areas. Where areas that are too large have been designed, site-specific effects can be smoothed over and not

identified. The specific effects might be localized within a large impact area and the advantages associated with a smaller impact area might be lost. With areas that are too small a significant amount of the effects might occur in adjacent areas and be missed by the analyst or decision maker. The objectives of this study and how some of the above problems will be addressed are included in the next section.

Research Objectives

The goal of this study is to investigate the behavior of certain economic multipliers of the wood products industry to determine the feasibility of providing more readily available, situation-specific, economic multipliers to assist in the economic impact analysis process. In conjunction with this, the study results should provide information that analysts can use for testing the effectiveness of designing economic impact areas or how multipliers respond to changes in the size of the impact area.

The above goals will be addressed by creating economic response surfaces (geographic regions) and examining how selected factors influence them. Regression analysis will be the primary analysis tool to address these goals.

Currently, most timber industry analysts and economic development teams do not have access to such information. As previously discussed, existing multipliers are either very general or they need to be determined through an extensive computer analysis by someone with specialized training. Economic impact areas are frequently selected with

inadequate information or skill with little awareness of the implications of incorrectly chosen areas.

The following specific questions need to be addressed in order to provide a better understanding of multipliers useful to the wood products industry and to conduct this study. These questions serve as a basis for the specific objectives.

- What is the variation in economic multipliers for wood products sectors across impact areas of differing sizes and characteristics, and what are the variables that might be responsible for these differences?
- Is there a detectible pattern for the variation? Do some variables cause multipliers to increase while others cause multipliers to decrease or have no effect at all?
- How can these findings be used? Can they be used to develop more site-specific readily usable multipliers? Can they be used to establish or test the appropriateness of selected impact areas?
- How can analysts working for the government, private contractors, or the wood products industry apply the appropriate multipliers for a given region?

Consequently, an approach is needed as explained above, that can simplify the problem of generating multipliers for a given region and problem and take some of the uncertainty out of estimating economic impacts. Because of the nature of the wood products industry, with clear and significant backward linkages from the major producers to the incidental producers, multipliers for the few most significant actors in the wood products

industry should be able to fairly represent the wood products industry as a whole and address the questions being asked.

National Forests, administered by the USDA Forest Service, serve as the nuclei of the economic impact regions in this study. They will also be the primary beneficiaries of the results of this study.

There are three research objectives for this study:

1. Describe the range of variation in economic multipliers for representative sectors in the wood products industry across geographic regions of varying size and characteristics.
2. Identify the key factors that explain the variation of economic multipliers from Objective 1, which might serve as variables in a multiplier prediction model.
3. Determine the feasibility of using selected explanatory variables from Objective 2 to estimate economic multipliers for the wood products industry.

Data used for explanatory variables must be readily available to provide an advantage over constructing input-output or economic base models. The study will not develop a ready-made set of multipliers or impact regions to be immediately applied. Rather it will focus on the feasibility of doing so.

To achieve these objectives, the study has developed a set of various economic multipliers representing selected wood products industry sectors and investigated the role

of selected explanatory variables that were used to predict those multipliers. The multipliers, serving as response or dependent variables in the regression model, were developed using the IMPLAN system. These multipliers served as the baseline, or were deemed to be the true multipliers for purposes of this study. Regression techniques were used to analyze the relationships between the multipliers and the explanatory variables.

Organization of Dissertation

This dissertation report begins with Chapter 1, which describes the problem and its importance, along with specific questions the study is designed to answer and the resulting study objectives. Can multipliers be developed to be useful for analysts that do not require the time or resources to calculate multipliers or identify regions on their own? Can economic multipliers be used to assist in determining the major extent of economic impact areas? Chapter 2 presents a review of available literature dealing with the subjects of economic impact analysis, input-output analysis, economic multipliers, and regional delineation. Past use of these procedures for evaluating management decisions for the wood products industry also are discussed. Chapter 3 describes the methods used to collect and analyze data in this study. Chapter 4 discusses the analysis and study results. Chapter 5 discusses the conclusion of the study and describes policy implications and suggestions for further research.

CHAPTER TWO: LITERATURE REVIEW

Introduction

This chapter reviews basic literature on economic impact analysis, multiplier analysis, and input-output analysis to provide a conceptual framework for the study. The chapter is divided into five sections. The first section briefly discusses economic impact analysis – its uses and relationship to recent regulatory requirements. The second section reviews the fundamentals of multiplier analysis – how multipliers are calculated and used. The third section presents the fundamentals of the IMPLAN economic impact analysis system, its history, products, and present capabilities in relation to calculating multipliers. The fourth section discusses how economic impact areas are defined. The fifth section discusses past studies.

The literature review begins with an extensive search of available databases. The most significant are those available through the Rocky Mountain Research Station of the USDA Forest Service. The Station uses a variety of specialized databases both on CDs and through commercial online vendors, through several universities, with access to over 400 bibliographic databases. The databases that were most productive were TreeCD, which has worldwide coverage of forestry literature – 1939 to present; Agricola, which has worldwide coverage of agricultural/rangelands literature – 1970 to present and Sociological Abstracts, which has worldwide coverage of literature on theoretical and applied sociology/behavioral sciences (including economics) – 1963 to present. The

Internet was also searched for applicable literature. The search results were reviewed and the literature that appeared to be most useful for this study was acquired and reviewed.

Economic Impact Analysis

Input-output analysis is one of many tools available for economic impact analysis although, for various reasons, the input-output approach is the most popular (Richardson 1979, 1985). Its popularity will continue to increase because the development of computers and software has made impact analysis much easier and more informative than it has been in the past.. Input-output analysis will be discussed in more detail later in this chapter.

With the assistance of computers and the availability of secondary data based models, the potential and demand for higher quality economic impact analysis is on the rise. This is a result of the increasing intensity of economic and community planning along with a myriad of new governmental regulations and policies requiring economic impact analysis prior to implementing new activities or proposals.

Moreover, economic impact analysis will be necessary to satisfy the requirements of the most recent National Forest Management Act Regulations (USDA 2000). The purpose of the regulations includes statements such as:

“Sustainability, composed of interdependent ecological, social, and economic elements, embodies the principles of multiple-use and sustained yield without impairment to the productivity of the land. Sustainability means meeting needs of

the present generation without compromising the ability of future generations to meet their needs (Section 219.1(b)(3))”.

The harvest of wood products plays a role in the sustainability goals. Section 219.21 elaborates on the importance of social and economic information to provide assessment and economic impact information to aid in the planning and decision making process.

The recently released Interior Columbia Basin Final Environmental Impact Statement, which could substantially change management direction on FS and BLM lands for a large geographic region, proposes agency decisions that “support economic and/social needs of people, cultures, and communities, and provide sustainable and predictable levels of products and services, from lands administered by the Forest Service or the BLM” (USDA, USDI 2000). Economic impact analysis using input-output analysis was used to determine many of the economic effects of this new policy proposal (Quigley et al. 1996). Economic impact analysis is needed to assess the economic effects as the new policy is implemented which will include various new roles (e.g., ecosystem restoration) for the harvest of wood products.

A variety of social and economic disciplines call for economic impact analysis including distributive justice (Phelps 1989, Wagner et al. 1992), social impact assessment (Burdge 1994, 1999, 1991, Bryan 1996), community planning (Rasker 1994, Hart 1999, Bauen 1996), rural development (Holland et al. 1995, 1997, Vasievich 1999, Fossum 1993) and environmental economics (Ekins 1992, Power 1988, 1996).

Multiplier Analysis

Economic multipliers are most frequently generated from input-output analysis.

Economic multipliers are a way of measuring the economic interdependence of a region.

Multiplier analysis can be an important procedure in assessing the economic effects of changes in the wood products industry on the economy of a region. Multipliers are used to estimate the direct, indirect and total economic impacts resulting from a change in “final demand”. The idea is to multiply the multiplier by some economic measures in order to get the total impacts (Aldwell 1984).

Economic multipliers generated by regional input-output models can be useful in providing detailed information about an industry’s effect on the regional economy. By computing various kinds of multipliers, one can measure the impact of an industry on economic variables such as employment, income, and the activity level of all industries in the region (Burford et al. 1981).

Economic multipliers vary according to the economic variables to which they apply (e.g. total industry output, income, employment, etc.) and how they are defined (e.g. Type I, Type II, Type SAM). Type I multipliers address direct and indirect effects. Type II and Type SAM multipliers address total effects. These types of multipliers can be referred to as “ratio multipliers” because they are the ratio of total economic effects in relation to direct economic effects. Multiplier variables and types of multipliers will be explained in more detail later in this chapter under “Input-output Analysis”.

Multiplier Analysis

Economic multipliers are most frequently generated from input-output analysis.

Economic multipliers are a way of measuring the economic interdependence of a region.

Multiplier analysis can be an important procedure in assessing the economic effects of changes in the wood products industry on the economy of a region. Multipliers are used to estimate the direct, indirect and total economic impacts resulting from a change in “final demand”. The idea is to multiply the multiplier by some economic measures in order to get the total impacts (Aldwell 1984).

Economic multipliers generated by regional input-output models can be useful in providing detailed information about an industry’s effect on the regional economy. By computing various kinds of multipliers, one can measure the impact of an industry on economic variables such as employment, income, and the activity level of all industries in the region (Burford et al. 1981).

Economic multipliers vary according to the economic variables to which they apply (e.g. total industry output, income, employment, etc.) and how they are defined (e.g. Type I, Type II, Type SAM). Type I multipliers address direct and indirect effects. Type II and Type SAM multipliers address total effects. These types of multipliers can be referred to as “ratio multipliers” because they are the ratio of total economic effects in relation to direct economic effects. Multiplier variables and types of multipliers will be explained in more detail later in this chapter under “Input-output Analysis”.

Various combinations of multipliers and types are possible. For example, it is possible to derive a Type I income multiplier, a Type II value-added multiplier, and a type SAM employment multiplier. Keynesian-type multipliers, can be derived from the IMPLAN output and an independent estimate of the total amount of new dollars spent in a region as a result of some action (Stynes 1999). These multipliers are commonly related to some physical unit of change in production. An example of this would be the number of jobs lost per million board feet reduction in timber harvest or million AUMs of grazing (Kolison et al. 1992). Response coefficients are not addressed in this study. They need to be calculated separately and locally because of the need to include physical output.

Among the most commonly used methods for estimating economic multipliers and ratios are 1) economic base models and 2) input-output models (Richardson 1985). Of the two methods input-output analysis has received the most use in recent years (Hastings 1993). This is probably due to the availability of new software and high-speed personal computers that make computations simple compared to the past as well as providing more detailed sectoral information from secondary data.

Economic base models - The assumption for economic base models is that all economic activity within a region must be classified as basic or non-basic. Basic sectors export their output to markets outside the region therefore bringing new money to the region. The derivative sectors arise from serving markets inside the region or it can be said that the non-basic sectors are attributable to the basic sectors. Therefore the

economic base multiplier (i.e. income) is expressed as the ratio of total income to the basic income.

A frequent approach for identifying the basic component of an economy is through the use of location quotients. The location quotient assumes that any production in excess of local consumption is exported and that percent is applied to the economic variable of interest to determine the basic component of the economy (Pleeter 1980, Krikelas 1992).

One of the more consistent lines of criticism has been that base studies perform well in description but poorly in prediction. Base/non-base classification is a useful way to characterize or describe a regional economy. But it is not certain that a single base/non-base ratio can yield a useful estimate of the multiplicative effects associated with a change in export activity (Bills and Zygodlo 1978).

Several other shortcomings of economic base models include (1) failure of the model to reckon with supply inelasticities, (2) drift over time of estimated parameter values due to evolving local economies, (3) focus on exports to the exclusion of other autonomous sources of demand and, (4) weaknesses that any Keynesian consumption function exhibits. In addition, estimated economic base multipliers exhibit wide variability (Frey 1989).

Economic base models will not be discussed further because this approach cannot be used to address the objectives of this study.

Input–output analysis – Francois Quesnay first described inter-industry relationships in 1758. However, the empirical application had to wait until the 20th century when Wassily Leontief developed the concept of multipliers from input-output tables. He received a Nobel Prize in 1973 for his work.

Input-output analysis is frequently chosen for regional analysis because it provides several types of information. It is an excellent descriptive tool showing in detail the structure of an existing regional economy. It provides important information on individual industrial sector size, and its behavior and interaction with the rest of the economy. It shows the relative importance of sectors in terms of their sales, wages and employment. It also provides a way to predict how the economy will respond to exogenous changes or changes that are planned. Therefore, it is useful in prescriptive exercises where various actions are being considered and the relative merits are to be determined based on alternative outcomes (Hastings 1993).

Input-output analysis is a means of examining relationships within an economy both between business and between businesses and final consumers. It captures all monetary market transactions for consumption in a given period of time. The resulting mathematical formulae allow one to examine the effects of a change in one or several economic activities on an entire economy.

A primary input-output study is based on data collected directly from industries. An example is the United States' Benchmark Study on Input-Output Accounts (U.S.

Department of Commerce 2001). Secondary input-output studies rely on data collected from other sources to construct the accounts. The inter-industry transaction information usually comes from some other primary study. IMPLAN (Minnesota IMPLAN Group 1999) is an example of a secondary input-output modeling system.

Trade flows are also part of the descriptive input-output model. They describe the movement of goods and services within a region and from and to the outside world (regional imports and exports).

By adding, what is referred to by economists, as social accounting data an analyst can examine non-industrial transactions such as payment of taxes by businesses and households. Social accounting data includes tax collection by governments, and payments to households and business. Input-output accounting describes the flow of commodities from producers to intermediate and final consumers. Social Accounting Matrices (SAMs) show the flow of money between all institutions (Maki 1997).

The regional economic accounts are used to construct local level multipliers. Multipliers describe the response of the economy to a stimulus (a change in demand or production). Purchases for final use (final demand) drive an input–output model. That is, industries producing goods and services for consumption purchase goods and services from other producers, and these other producers in turn purchase goods and services (indirect purchases). These indirect purchases (or indirect effects) continue until leakages from the region stop the cycle.

The indirect effects and the effects of increased household spending (induced effects) can be mathematically derived as sets of multipliers. The derivation is called the Leontief inverse (Miernyk 1965). The resulting sets of multipliers describe the change of output for each industry caused by a one-dollar change in final demand for any given industry.

There are a large number of references on input-output analysis covering an array of inter-related matrices, matrix algebra, and trade flow mechanics (Miernyk 1965, Minnesota IMPLAN Group 2000, Miller 1985, Otto and Johnson 1993, USDA 1978, Rose et al. 1989, Fjeldsted 1990, and Hewings 1985). Technical details, to the extent necessary, will be discussed as part of the study methodology. However, this study focuses primarily on the behavior of economic multipliers, which are described in the following section.

Although input/output analysis has gained in popularity as the tool of choice for economic impact analysis it is not without critics or shortcomings. A few of the concerns associated with input/output analysis include

- Models are especially sensitive to trade flow assumptions.
- Multipliers are based on technology matrices that are often quite old.
- New sectors, not presently existing, have to be added to the model by the user as increases in demand for certain commodities or industries reach threshold levels. Otherwise these inputs are imported and outputs are exported resulting

in incorrect impact analysis solutions. Information is usually not available as to what the thresholds are or what sectors might emerge.

- Production functions are assumed to have constant return to scale, which means they are considered linear. If additional output is required, all inputs increase proportionately.
- Data for some individual sectors, especially those in agriculture, have a substantial chance to be in error as the data is highly disaggregated based on questionable variables.
- There are no supply constraints. An industry is treated as having unlimited access to raw materials and its output is limited only by the demand for its products.
- Another assumption is that there is a fixed commodity input structure which means that price changes do not cause a firm to buy substitute goods which is what would happen in real situations.
- The model also assumes there is homogeneous sector output or the proportions of all the commodities produced by that industry, remain the same, regardless of total output. An industry will not increase the output of one product without proportionately increasing the output of all its other products.
- The industry technology assumption comes into play when data is collected on an industry-by-commodity basis and then converted to industry-by-industry matrices. It assumes that an industry uses the same technology to

produce all its products. In other words, an industry has a primary or main product and all other products are byproducts of the primary product.

Common Types of Multipliers

Type I - Type I multipliers give the direct and indirect effect only – that is, the original expenditures resulting from the impacts plus the indirect effects of industries buying from industries. An example of a direct effect would be a sawmill employing workers. An indirect effect would be the local auto repair shop hiring another mechanic who will work on the sawmill's vehicles. Whether or not an effect is direct or indirect depends on the starting point of the analysis. For example if the starting point of the analysis was the mechanics shop then the hiring of the mechanic would be a direct effect and the hiring of an extra worker at the parts store (where the mechanics shop buys its parts) would be an indirect effect. However, the starting point for this analysis was the wood products sectors chosen for analysis. Household expenditure effects – i.e. induced effects are not estimated with the Type I multiplier. The Type I multiplier can be defined as the ratio of the direct and indirect changes in the variable of interest to direct changes in that variable.

Type II – Type II multipliers account for the direct, indirect, and induced effects where the induced effects are based on income. Examples of direct and indirect effects are given above. Induced effects result from expenditures of the household sector on final consumption. When the mill worker, or mechanic, and parts store worker spend their income it generates new employment and income. These rounds of spending

continue until all the effects of the original spending eventually leak out of the selected impact area.

The relationship between personal consumption expenditures (PCE) (induced effects) and income is based on resident-only income from the Social Accounting Matrix (SAM) accounts. The assumption is that there is a linear relationship between local income and local expenditures. The Type II multiplier can be defined as the ratio of the direct, indirect and induced changes in the selected variable, to direct changes in that variable.

Type II multipliers are created by incorporating household expenditures into the rounds of inter-industry purchases. In this way, money flowing to the labor force is recycled through the economy and is not lost as a leakage as it would be with the Type I multiplier.

There are three different Type II formulations: “Standard” (which is not available in IMPLAN software), “SAM Based” and “Specified Disposable Income”. The difference in all three methodologies lies in how the denominator representing total household expenditures is derived for conversion of household expenditures (the PCE vector) to coefficients. Each dollar going to household income from either employee compensation or proprietor income is spent through the PCE coefficients in each of the rounds of direct and indirect effects. The household income row and PCE column coefficients can be incorporated in the Leontief matrix the same as any other industry. Common to all three methodologies is that:

- 1) The PCE final demand (the household expenditures column) represents expenditures for goods and services made by residents of the region being modeled.
- 2) The household row represents income earned by labor (usually employee compensation plus self-employment income) paid by regional industries and institutions for household production – i.e., labor.

Therefore, the household expenditures column is residence based and the household income row is by place of work.

The “standard” formulation for Type II multipliers can be found in the input-output textbooks (Miller and Blair 1985, Hoover et al. 1984). The denominator, by which the PCE column coefficients are derived, is the sum of payments made by industry and institutions for labor – i.e., employee compensation and proprietor income. The PCE column represents expenditures by regional households (from income from all sources) and that income is work place based. This creates problems in places that have large retirement industries, and consequently, large amounts of retirement income (Olson, 1997). This can coincide with high natural amenity areas that have significant wood products industries – the subject of this study. The amount of PCE reflects both retirement and labor income but the denominator only includes labor income. Resulting coefficients and therefore multipliers, can be overstated unless a certain number of retirees are assumed to come with each new job. Similar problems exist if a region does not incorporate the labor force and has a significant number of commuters – i.e., changes

in regional household spending does not directly reflect changes in work-based labor income, which are respent through PCE. These are important implications when dealing with regions such as those occupied by National Forests or the timber industry where aesthetic values are generally above average and consequently the retirement component of the population is disproportionately high.

The above situation is important in respect to the proposed study because frequently workers will live in high amenity rural areas but commute into urban areas to work. As the impact areas are enlarged, as will be the case in this study, we should see this situation reflected in the resulting induced (Type II) multipliers. Workers will tend to spend much of their money in the larger area. This suggests that the design of impact areas should take this situation into consideration and maybe include those areas where much of the induced spending occurs. It also suggests that analyzing the multipliers might help in the design of impact areas – part of the subject of this study.

IMPLAN models have an advantage over many earlier input-output models because data are available for a complete set of “Social Accounting Matrices” (SAM) (Alward et al. 1996). The PCE expenditure vector in the “standard” input-output account only includes expenditures for goods and services provided by industry sectors (local or imported). To derive PCE coefficients the PCE coefficient column generated by the SAM is used. The coefficients can now represent the disposal of each dollar of regional household income regardless of the source. The Type II multiplier in this formulation, therefore, includes spending and re-spending of each dollar of labor income through the PCE coefficients –

transfer payments are no longer a part of the impacts. However, significant commuting into the region may cause overestimation of multipliers, as the model will assume their labor income will be spent locally instead of in the region from which the commuters came.

When using the “Specified Disposable Income Ratio” method, IMPLAN borrows the methodology from the BEA’s RIMS II project, which essentially is a SAM based multiplier without the SAM data (Minnesota IMPLAN Group 1999). First the PCE vector is normalized by summing the PCE values and dividing by the total (the BEA uses the national PCE vector while IMPLAN models are regionalized). The resulting PCE coefficients sum to one. Each dollar of labor income can now be spent through this vector, but first we must account for the fact that each dollar in income is not spent solely on PCE. Some of it is used to pay taxes and some goes to savings. This is simulated by applying a disposable-to-total income ratio to the PCE coefficients. Nationally, this ratio is 0.85. The PCE coefficients are reduced so that they now sum to 0.85 – i.e., every dollar of labor income generates 85 cents of PCE. In IMPLAN, the user can edit the disposable income ratio. This is useful for studying any policies affecting household disposable income rates (taxes, savings). Type II multipliers simulating loss of income through commuters from outside can be generated by proportionately reducing the disposable income rate (Minnesota IMPLAN Group 1997).

Type SAM – Type SAM multipliers are the direct, indirect, and induced effects where the induced effect is based on information in the social accounting matrix. This

relationship accounts for social security and income tax leakage, institution savings, and commuting. It also accounts for inter-institutional transfers. The SAM multiplier was the indirect multiplier used for this study because it has the capability of accounting for all expenditures regardless of source. This is important because there are significant transactions between government institutions and the wood products industry and their omission could result in inaccurate multipliers.

IMPLAN – An Economic Impact Analysis System.

This study used the IMPLAN economic impact analysis system to generate impact models and resulting economic multipliers. IMPLAN (IMpact Analysis for PLANing), was originally developed by the USDA Forest Service in cooperation with the Federal Emergency Management Agency and the USDI Bureau of Land Management to assist the Forest Service in land and resource management planning (Siverts et al. 1985). It is one of the most widely used economic impact tools for wood products industry applications (Flick and Teeter 1988, Aruna et al. 1997, Zeng and Harou 1988, Lord and Strauss 1993). This study used the IMPLAN Pro, Version 2.0 (hereafter referred to as IMPLAN) economic impact system to generate impact models and resulting economic multipliers. The multipliers generated by IMPLAN were used as the response (dependent) variables in the regression models for this study. IMPLAN also served as a source of data used as a few of the explanatory (independent) variables.

The IMPLAN system has been in use since 1979 and has evolved from a mainframe, non-interactive application that ran in “batch” mode to a menu-driven microcomputer program that is completely interactive (Minnesota IMPLAN Group 1999). The Minnesota IMPLAN Group (MIG) began work on IMPLAN databases in 1987 at the University of Minnesota. In 1993, MIG, Inc. (MIG) was formed to privatize the development of IMPLAN data and software. Version 1 of the Windows software was developed by MIG and released in June of 1996.

The IMPLAN database, created by MIG, consists of two major parts:

- 1) National-level technology matrices;
- 2) Estimates of regional data for institutional demand and transfers, value-added, industry output and employment for each county in the U.S. as well as state and national totals.

The IMPLAN data and accounts closely follow the accounting conventions used in the “Input-Output Study of the U.S. Economy “ by the Bureau of Economic Analysis (1980).

There are two components to the IMPLAN system, the software and the database.

The software performs the necessary calculations using the impact area or region data to create the models. It also provides an interface for the user to change the region’s economic description, create impact scenarios and introduce changes to the local model.

The databases provide all the information needed to create regional IMPLAN models.

The IMPLAN system can be used to analyze a wide variety of issues including those associated with Federal land management.

IMPLAN's regional social accounting system allows a user to:

- 1) Develop a set of balanced economic/social accounts.
- 2) Develop multiplier tables.
- 3) Change any component of the system (i.e. production functions, trade flows or database).
- 4) Create custom impact analysis by entering final demand changes.
- 5) Obtain any report in the system and examine the models assumptions and calculations.
- 6) Define regions appropriate to address problems being analyzed.
- 7) Select and focus on just those sectors of interest.

Because the purpose of this study was primarily to determine the feasibility of predicting multipliers, given certain explanatory variables, the discussion of the IMPLAN model was limited to what is needed to understand the study. This included, the various types of multipliers, and the economic variables generated by IMPLAN that the multipliers apply to.

Multipliers generated by IMPLAN.

Types of multipliers

IMPLAN generates Type I, II, and SAM multipliers. If it is desirable for an analysis to include the induced effects generated by indirect spending, the type of induced multiplier to construct (e.g. Type II or SAM), must be selected. IMPLAN allows the user to select what institutions to include in the calculation of the SAM multiplier.

Economic Variables

IMPLAN calculates economic ratios and multipliers for various impact measures that include total industry output (sales); the components of value added such as labor income, employee compensation, proprietor's income, other property income and indirect business taxes; and employment.

Total Industry Output – Total industry output (TIO) is the value of production by industry(s) for a given time period. For IMPLAN, TIO is annual calendar year production. Output can be measured either by the total value of purchases by intermediate and final consumers, or by intermediate outlays plus value added. Output can also be thought of as the value of sales plus or minus the changes in inventories. The TIO multiplier is the change in output resulting from an increase of \$1.00 in final demand.

Value Added – Income multipliers (or for any of the value-added components) are derived from the relationship between income and output. Study area data has the total industry output and total income for each sector. From this, income per dollar of output

can be calculated. An industry multiplier is split into the direct and indirect effects and then multiplied by the income per dollar of output ratio to get the income direct and indirect effects.

Employment – The employment multiplier is created in the same manner as the income multiplier, but using output per worker ratios instead of output per dollar of income. First the employment per dollar of output is calculated, and then the direct and indirect effects are estimated (Type I multipliers). The level of employment per million dollars of output is then multiplied by the output multiplier. Employment data in the IMPLAN model is from local sources.

It is important to know that employment, as reported by IMPLAN, is a single number of jobs. It includes both full time and part-time workers – not full-time equivalents.

Other sources of multipliers.

In the mid-1970's the Bureau of Economic Analysis (BEA) completed development of a method for estimating regional Input-Output multipliers known as RIMS (Regional Industrial Multiplier System). More recently BEA completed an enhancement of RIMS known as RIMS II (Regional Input-Output Modeling System). A product of the RIMS II system is a published set of multiplier tables by industry aggregation for each state in the U.S. They are "total" multipliers for output, earnings and employment. These multipliers are sometimes useful when resources are not available to calculate more

specific multipliers. However, inaccurate results can arise from these multipliers when the situation at hand is different from the average. Also, these multipliers are not useful to this study because they do not present geographic variation about an area of potential change in production and the industry-sectoring scheme for the wood products industry is not nearly as useful as in IMPLAN (U.S. Department of Commerce 1986).

The Impact Area

One of the first tasks in economic impact analysis is identifying the impact area or describing the spatial context of an impact analysis. In this study it was essential to get a wide variety of impact regions for analysis so identifying the most appropriate areas for impact analysis was only a product of chance rather than purposive or biased. The procedure used for this particular study will be explained in the next chapter. When normally selecting the appropriate impact area the analyst will usually define a county or groups of counties (and/or states) for conducting the economic impact analysis. Also, parts of counties (zip code areas) can now be included with other areas or analyzed separately (Olfert et al. 1994, Goldman et al. 1997)

The following factors are presented as guides in selecting the appropriate counties for the analysis. There are no hard-and-fast rules that can be mechanically applied to determine the area to use. However, using information contained in the impact system database, the user's own knowledge of the area, and the question being evaluated, appropriate areas can be developed.

The impact area should be defined as (1) a functional economic unit of a size appropriate to the policy issue and (2) an area that includes most of economic factors that are most directly affected by the policy (USDA 1988).

In terms of size the area should be large enough to:

- 1) Include all relevant activities (forward and backward linkages) related to the question, and
- 2) Serve as a functional economic area.

But the area should be small enough to:

- 1) Be geographically oriented to the question, and
- 2) Represent the individuals, institutions and industries most affected by the proposed action.

Factors to consider include:

- 1) Problem definition. - What question is the study expected to address? Region definitions depend on the issues under study. The boundaries depend on the purpose of the analysis. However, there is a tendency where regional boundaries are drawn for narrow purposes to end up with regions that geographers call “lifeless”. A lifeless region offers little reason for its definition beyond the narrow purposes for which it is defined (Robison 1997).
- 2) Trade patterns - What are the principal trading patterns of key industries or institutions? This includes governments and households. Where do local residents spend their money? Where does the labor force live? Where are the

travel corridors? What is the location of supporting industries and services (USDA 1988)?

- 3) Targeted audience - Is the analysis being conducted for a particular group or audience as opposed to being prepared for general public knowledge? If a particular group is targeted, a study area that represents their interests should be used. For example a Board of Commissioners might only be interested in economic effects in their particular county.
- 4) Spatial grouping. - Generally the impact area is made up of contiguous counties. However, for some purposes leaving gaps or holes may be appropriate. If impacts are desired both at a local level and a broader regional or state level it may be necessary to develop several models with increasingly larger areas included. This is similar to what was done in this study.
- 5) Leakages. - Generally, the smaller the area, the less diverse the economy and the more spending leaks from the area. It might be useful to increase the size of the impact area to reduce the leakage. This will also increase the income and employment impact.
- 6) Functional economic areas. - This type of area is considered to be an ideal strictly from an economic perspective. It contains a resident labor force for local industries, a source for consumer purchases, and essential support

services. However, it may not be the ideal for a particular impact analysis if some other (e.g., political) consideration takes precedence (USDA 1992).

Predefined Study Areas - The U.S. Bureau of Census and the Department of Commerce have predefined study areas such as the Census Commuting Area or the Metropolitan Statistical Areas (MSA). Both of these are based on counties. The MSA's capture metropolitan regions quite well, whereas the commuting areas tend to be large (USDA 1999).

Tolbert and Kizer (1990) have defined 382 labor market areas (875 sub-market areas) for the U.S., based on county-level journey to work data from the 1990 Census. There are also Bureau of Economic Analysis economic regions and areas based on counties but each is centered on a Census defined Metropolitan area. The BEA areas are defined to have a minimum population of 100,000.

The USDA Forest Service has identified a network of impact areas for the purpose of accomplishment reporting for several forest resources. An example of this is that done for several National Forests in Minnesota, Wisconsin and Michigan (Retzlaff et al. 2000).

At the present, a doubly constrained gravity model is being developed by the Minnesota IMPLAN Group and the USDA Forest Service Inventory and Monitoring Institute to estimate trade flows for over 500 commodities between all counties in the U.S. The IMPLAN software and national database of counties will be used to create the attracting

masses (supply and demand). Impedances will be derived from the Oakridge National Laboratory's Transportation Network model between centroids for all U.S. counties to represent distances between the masses. The resulting trade flow estimates between counties will allow for the replacement of the current IMPLAN methodology of econometrically derived Regional Purchase Coefficients with "observed" county level RPCs. This will be a substantial improvement in the identification of impact areas and determination of trade flows (Olson 2000).

Multiplier Studies Closely Related to this Study.

One of the objectives of the literature review was to be reasonably assured that the proposed study was not a duplication of an already completed study as well as acquire ideas useful to this study. The search indicated that there were similar studies but each of these studies had some major differences from this study. Several studies investigated the possibility of estimating multipliers through regression models but did not apply it to the wood products industry. Several other studies analyzed wood products industry multipliers but did not propose alternatives to using existing input/output models. These studies are briefly discussed in the following section.

Burford and Katz (1981) felt that acceptable input-output models could be built by using census data to determine the regional purchase coefficients and then using these in conjunction with a technical coefficient matrix could estimate the matrix of regional coefficients. These results could then be used as explanatory variables in a regression model to predict multipliers. Their study showed merit but is not presently useful

because the effort required to develop the explanatory variables is probably more costly than using a secondary data model such as IMPLAN.

A study by Zheng and Harou (1988) successfully estimated multipliers in the wood products industry by using the “internal purchase ratio” and the “intra-regional sale ratio” as explanatory variables in a regression model. Again, these variables are not readily available outside the input/output model and their acquisition and use would require special skills. This technique is similar to what is used by the U.S. Department of Commerce, Bureau of Economic Analysis (1986) to set up their modeling systems (RIMS and RIMSII) that was then used to develop their regional multiplier tables. The above approach arose from the work done by Drake (1976) whose idea was to estimate multipliers without the creation of an entire regional input-output model.

Regression analysis was used by Mulligan and Gibson (1984) to estimate economic base multipliers for small communities. Their objectives were similar to this study in that the authors were looking for readily available explanatory variables but still ended up including components of the input-output model itself.

A study by Wen-Huei Chang (2001) describes and explains how multipliers used to analyze recreation and tourism impacts vary across a wide range of conditions.

Conditions included population, population density, area, and geographic location. These variables were also included in this study.

In summary, there have been a number of studies to simplify the estimation of economic multipliers. Generally, they have depended on regression models that required data that was not readily available. Moreover, a number of studies used input-output model data

as explanatory variables for estimating multipliers. Since the completion of those studies secondary data models have become relatively available, fast and inexpensive.

Moreover, contemporary models do not require the skills that input-output analysis methods previously required. This study does not replicate those studies because this study proposes using information that is readily available.

One of the purposes of the literature review was to identify explanatory variables that would be useful in explaining changes in economic multipliers. It was found that variables used in past studies were either not generally available or had already been considered for inclusion in this study prior to the review of the existing literature. However, it did reinforce the notion that the selected variables were worthy of analysis. Discussion of the variables used in this study begins in Chapter 3.

CHAPTER THREE: RESEARCH METHODS

Introduction

The research objectives of this study are to describe the range of variation in economic multipliers in the wood products industries, to identify the key factors that explain the variation, and to determine the feasibility of using these key factors in building a model that could be used as an inexpensive tool to predict economic multipliers. To achieve these objectives 650 economic impact regions were selected that varied in size and economic characteristics. Input-output models, using IMPLAN, were developed for determining multipliers for each of these regions for selected sectors of the wood products industries. Variations in multipliers were examined by comparing characteristics of the regions selected including industry production functions. Regression models were developed to identify regional characteristics that explain the variations in multipliers and evaluating the potential for generating multipliers from readily available explanatory variables.

Economic Impact Model Selection

The 1997 IMPLAN database, which is the most recent year available, was used to build all the regional models. For purposes of this study IMPLAN multipliers are treated as representing the actual interactions within a local impact area. However, the multipliers are derived from the national-level technology matrix, which serves as a basis for each sector's production functions and resulting multipliers. Regional data is applied to

national matrices (absorption and byproducts) to create a set of regional accounts. The value added and final demand components of the transactions table are from regional data. Thus, the assumptions used to calculate the multipliers may influence the direction and magnitude of local interactions if the national production functions do not resemble the local production functions (local interactions).

In an attempt to determine if there was the potential to have major differences in production functions from region to region, the average firm size was examined. It was assumed that significantly different sizes of operations would lead to different economies of scale and possibly different production functions from region to region. This was done by determining the number of firms and total jobs for SIC 24, Lumber and Wood Products sector from County Business Patterns. From this the average firm size in terms of employment for different regions was determined. It was found that there was very little difference in firm size from region to region. The average firm size for the United States for 1997 was 21.3 workers. For regions of the United States the average was as follows: western, 21.8, northcentral, 23.8, southeast, 21.9, and northeast, 16.1. These differences are not large enough to suspect they would cause major differences in production functions.

The IMPLAN data files are available at the county, state and national levels. Models can also be developed at the zip code level (Olson 1997). For this study only county-level databases were used and combined as described as follows to form the appropriate impact regions. The current version of IMPLAN (IMPLAN Pro 2.0), was used in this study to

generate economic multipliers. The only other feasible available system to use was RIMSII. However, it was not used because of the initial investment required and lack of technical support for using the model. The IMPLAN software program and data is presently accessible from the USDA Forest Service along with technical support.

Selecting the Study Area(s)

Because the primary use of the results of this study will be for economic effects analysis for timber management programs for units of the USDA Forest Service, the impact areas were centered around National Forests. Impact regions were designed for each National Forest. Although all National Forests do not have a timber program, all were included because they at least had the potential to have a timber sale program (Appendix E).

The National Grasslands were not included because they generally had no potential for timber sale programs. Alaska and Puerto Rico were not included in the study because conditions in those states are vastly different than other states and it was felt that any attempt to develop meaningful predictors for Alaska and Puerto Rico based heavily on mainland data would be unsuccessful. This resulted in 155 sets of impact areas being designed.

For each National Forest a set of up to five impact areas, were developed where data was available. The reason five areas were developed for each National Forest was to assure an adequate sample size and to also assure an adequate range in the size of impact areas – in terms of the explanatory variables to be used in the analysis.

The first impact area (Level A) consisted of the most populous of all individual counties that contain any lands within the designated National Forest. This was the smallest of the set of 5 impact areas. Although this region might exclude much of the NF land in some cases it could capture a significant amount of the economic activity.

The second impact area (Level B) included all the counties having lands within the designated National Forest. This model might show that only including the geographic source of raw materials could exclude much of the economic activity resulting from the wood products industry. This was usually the next to smallest impact region and always included the county comprising Level A.

The third impact area (Level C) included all the counties in Level B along with all the counties adjacent to the counties in Level B. This was an attempt to view the effectiveness of a mechanical construct that attempts to capture most of the economic activity without knowing where the impact variables (e.g. where do people work, spend their money etc.) actually are. It also provides additional diversity in impact area size.

Level D is the impact region that is used by the Forest Service for the TSPIRS (Timber Sale Program Information Reporting System) report system. TSPIRS was developed in response to Congressional direction contained in the Conference Committee Report on the 1985 Interior Appropriations Bill as a result of concern over “below cost” timber sales. One part of the annual report displays the economic impact of the Forest Service timber program, for the year of the report, in terms of employment and income. Of the

five levels of impact area used in this study this is thought to be the most appropriate impact area and is constructed pursuant to criteria in the Forest Service Handbook 1909.17 and Retzlaff et al. (2000). However, this level covers only approximately half the National Forests in the study because of missing data.

Level E consists of all the counties in the Bureau of Economic Analysis (BEA) Economic Area (s) that the National Forests predominantly fall within (BEA 1995, Johnson 1995). This was usually the largest of the impact areas and often occurred in several states.

IMPLAN models were developed for each of the impact areas described above (650 models). Some Forests had only four impact areas because of the lack of information needed to design the Level D impact area. The study objectives require that a wide variety of conditions among impact areas exist in order to evaluate the effects of those conditions on the resulting economic multipliers. It would also provide a realistic array of impact areas. This was evaluated by the first objective that describes the range of variation in the multipliers and the second objective that looks at the range of conditions within the impact areas that could cause the variation.

Another approach considered, to design economic impact areas, was an expanding core in the form of concentric circles or layers of counties. For example another layer of adjacent counties would be added to Level C. The feasibility of this approach was tested and found to be somewhat possible in the east and south where counties are relatively small. However, in the remainder of the country the resulting impact areas were often

enormous and frequently crossed several states. A modification of this approach would be to add counties one at a time, each time creating a new model. However, there is no clear criteria as to what sequence the counties would be added and it could create a substantial amount of extra work with no obvious benefit. Also, the study objectives should still be achievable without this extra model-building as there remains a large range of situations within the existing design.

One of the reasons for including a wide variety of regions is to get an ample range in the value of multipliers to see how the variation is related to the various characteristics of regions. This is an important assumption in linear regression that will be addressed later in this report (Appendix A). The list of impact areas is available from the author upon request.

Economic sectors analyzed

The purpose of the study was to analyze the behavior of the economic multipliers of the wood products industry in response to a number of variables. Several assumptions were made in identifying what economic sectors are normally considered part of the “wood products industry”. This is necessary to determine which sectors this study should focus on.

The main assumption is that the primary use of multipliers is to measure the marginal effects resulting from management decisions that might affect the “wood products

industry” and consequently the surrounding economy. In the Forest Service or other federal agencies this would normally be a decision that would affect the timber supply to local processors. Therefore we must ask the question – will there be any direct effects on a particular sector from a change in timber supply? For example, should the home construction industry, which uses wood, be considered part of the wood products industry? In this study it would not be included because it is not directly affected by a change in the local timber supply. The home construction industry would exist even if the timber industry were not present. The presence of the industry is a function of local demand for housing, not the supply of raw materials. Beckley (1998) showed that a significant part of the lumber used in housing, even in those regions where the lumber supply exceeds demand, is imported from outside the region. There is a substantial amount of cross hauling. The same can be said for other wood-using sectors such as millwork, wood kitchen cabinets, wood containers, and furniture. Their existence is not a function of favorable wood prices because of the presence of a local sawmill. Using the above logic, Table 1 shows the IMPLAN sectors that could potentially be considered part of the wood products industry for the production and analysis of economic multipliers in this study. These sectors need to be located reasonably close to raw materials because of high transportation costs.

The next step is to decide which, if not all, of these sector’s multipliers should be analyzed. The first, alternative might be to form an aggregate of all the wood products industries. Although aggregating speeds up the model development process, and reduces the size of reports, it can introduce substantial error due to the loss of data detail. Errors

Table 1 – Forest Products IMPLAN Sectors	
IMPLAN Sector	IMPLAN Sector Name (1987 SIC Codes)^a
22	Forest Products (Commodity – no SIC code)
24	Forestry Products (0810, 0830, 0970)
26	Ag, Forestry, and Fishery Services (710, 720, 750, 760, 0254, 0850, 0920)
133	Logging Camps and Contractors (2410)
134	Sawmills and Planing Mills (2421)
135	Hardwood Dimension and Flooring Mills (2426)
139	Veneer and Plywood (2435, 2436)
146	Reconstituted Wood Products (2493)
161	Pulp Mills (2610)
162	Paper Mills (2620)
163	Paperboard Mills (2630)

^a Detailed descriptions of each of these sectors can be found in the Standard Industrial Classification Manual – 1987. Executive Office of the President, Office of Management and Budget.

are introduced from production functions, output per worker averages, and other value added ratios. Aggregating the region's industry sectors before generating multipliers has the effect of taking several individual industries and combining them to form a totally new industry. Dramatic errors can happen when multipliers are derived from the production functions of aggregated industries. The production function of the new aggregated industry becomes the weighted average of the individual production

Table 2 – Multiplier Response Variables

Economic Sector (IMPLAN Code) ^a	Multiplier Class	Type Multiplier	Variable Code
Logging Camps and Contractors (133)	Employment	Type I	LE1
		Type SAM	LE2
	Personal Income	Type I	LP1
		Type SAM	LP2
Sawmills and Planing Mills (134)	Employment	Type I	SE1
		Type SAM	SE2
	Personal Income	Type I	SP1
		Type SAM	SP2
Veneer and Plywood (139)	Employment	Type I	VE1
		Type SAM	VE2
	Personal Income	Type I	VP1
		Type SAM	VP2

^a First Digit ---- *Economic Sector*

Logging Camps and Contractors = L

Sawmills and Planing Mills = S

Veneer and Plywood = V

Second Digit---- *Multiplier Class*

Employment Multiplier = E

Personal Income Multiplier = I

Third Digit ---- *Multiplier Type*

Type I Multiplier = 1

Type SAM Multiplier = 2

functions. Industries with the greatest outputs have the greatest influence on the aggregated industry, but the new industry's production function may not truly represent an industry being impacted. This generates an aggregation-induced error (Olson 1995, 1995a). Test runs on several counties showed that there is a significant variation in Type SAM employment multipliers among the 11 wood products sectors. Some are

considerably above and some considerably below the Type SAM aggregate multiplier. Another problem with an aggregated multiplier is that it would be difficult to extrapolate to other regions because it is doubtful that the sector proportions would be the same or that the region would even have the same combinations of sectors.

For the above reasons no attempt was made to aggregate sectors in calculating multipliers. In addition, only the most significant sectors in the wood products industry (Table 2) were analyzed.

There are several reasons for examining a subset of the sectors. First, there are 11 sectors identified, which could make the study too burdensome. Many of the sectors are already backward linked. For example the sawmills and planing mills (sector 134) sector has backward linkages to four other wood products sectors. This was determined by examining the production function in the absorption matrix generated by IMPLAN (Table 3). For example the logging sector (133) buys 32 percent of its inputs from the forestry products sector (24) and the sawmills sector (134) buys 28 percent of its input from the logging sector. All wood products sectors buy more from other wood products sectors than all other sectors combined. These transactions will all be accounted for as indirect effects. Employee compensation varies from 16 percent for the logging sector (the least labor intensive) to 25 percent for the veneer sector (the most labor intensive).

Table 3 - Wood Products Industry Production Functions^a			
Industry	Logging	Sawmills	Veneer
Commodity Demand:			
Forest Products Sectors:			
Forestry Products (24)	.32	.09	.09
Logging (133)	.08	.28	.16
Sawmills (134)	.00	.08	.01
Veneer (139)	.00	.00	.08
Total Forest Industry	.40	.45	.34
Other Commodity Demand	.22	.25	.28
Total Commodity Demand	.62	.70	.62
Value Added:			
Employee Comp	.16	.18	.25
Proprietary Income	.02	.02	.04
Other Property Income	.19	.08	.08
Indirect Business Taxes	.01	.00	.01
Total Value Added	.38	.30	.38
Total	1.00	1.00	1.00

^a Data in Table 3 are Gross Absorption Coefficients from the IMPLAN Industry Balance Sheet from report united states.iap. 1997 data.

The study will focus on the sectors that are normally significant and higher on the processing ladder.

The study will include Sector 133, Logging Camps and Contractors; Sector 134, Sawmills and Planing mills; and Sector 139, veneer and plywood.

These sectors are described as follows:¹

Logging Camps and Contractors (133) – Establishments primarily engaged in cutting timber and in producing rough, round, hewn, or riven primary forest or wood raw materials, or in producing wood chips in the field.

Sawmills and Planing Mills (134) – Establishments primarily engaged in sawing rough lumber and timber from logs and bolts, or resawing cants and flitches into lumber, including box lumber and softwood cut stock; planing mills combined with sawmills; and separately operated planing mills which are engaged primarily in producing surfaced lumber and standard workings or patterns of lumber.....

Veneer and Plywood (139) – Establishments primarily engaged in producing commercial veneer and those primarily engaged in manufacturing commercial plywood or pre-finished plywood. This includes non-wood backed or faced veneer and plywood, from veneer produced in the same establishment or from purchased veneer.

Sector 146, Reconstituted Wood Products; Sector 161, Pulp Mills; Sector 162, Paper Mills; and Sector 163, Paperboard Mills will not be included in the analysis because they occur so infrequently it would be difficult to statistically evaluate. Those sectors that are minor and well represented as indirect effects will not be included in the analysis (Sector 22, Forest Products; Sector 24, Forestry Products; and Sector 26, Agricultural, Forestry and Fishery Services). Finally, if a sector occurs very infrequently and is also probably

¹ Descriptions of wood products sectors are from the Standard Industrial Classification Manual – 1987. Executive Office of the President, Office of Management and Budget.

independent of local raw material supply, it is excluded (Sector 135, Hardwood Dimension and Flooring Mills).

Impact Model Construction

There are two different impact model components that are normally constructed by IMPLAN for each of the impact areas selected. The *descriptive* component describes the transfers of money between industries and institutions for the impact area being analyzed. It contains the social accounts and the input-output accounts. The social accounts include variables such as employment, income, and total industry output for the impact area being analyzed. The input-output accounts show the flow of funds between economic sectors (intermediate and final demand).

The *predictive* component is the set of input-output multipliers, which “predict” total regional activity based on a change in consumption or demand. This includes the calculation of the economic multipliers. For this study, all that was needed from the *predictive* component were the economic multipliers, which were used as the only response variables in regression analysis. The *descriptive* component must be generated before IMPLAN generates the *predictive* component.

The following sections discuss several decisions and assumptions that were made which affect the multipliers that were generated and analyzed.

Trade Flows

When constructing the social accounts, certain assumptions must be made with respect to trade flows (imports/exports/transshipments). This is an important assumption in that it determines how much is spent within the impact region, how much is spent outside the impact region from activity originating within the impact region, and consequently how large the resulting multipliers will be. In this study, trade flows will be estimated using the *Regional Purchase Coefficient* (RPC) approach. RPCs are derived with an econometric equation that predicts local purchases based on the region's characteristics. The RPC is the percent of local demand satisfied by local production. For example, an RPC of 0.75 for a given commodity means that for each \$1 of local need, 75% will be purchased from local producers. This method is based on the characteristics of the region and describes the actual trade flows for the region mathematically. IMPLAN software generates RPCs automatically with a set of econometrically based equations. There is a different equation for each commodity with variables filled by study area data. The RPCs are limited by the supply/demand-pooling ratio (explained below). The ratio of locally purchased to imported goods is perhaps the most significant factor affecting subsequent multipliers (MIG 2000). The greater the quantity of goods purchased locally, the more local economic activity will be stimulated and hence the larger the resulting multipliers. IMPLAN allows RPCs to be edited if the user feels they have better information. This has not been done in this study as there is no known better secondary data.

Table 4 shows an example of RPCs for four different states in different regions of the United States. For example in Georgia economic impact areas, 76 percent of what the sawmills sector buys from the logging sector is bought locally. Table 4 also shows there is little difference from region to region with exception of the sawmill sector in Michigan, which is significantly smaller. It also shows that the logging and sawmill sectors do not buy anything locally. This is consistent with Table 3, which shows that the production functions for both logging and sawmills do not include the veneer sector. The production functions and regional purchase coefficients must both be examined when

Table 4 - Regional Purchase Coefficients ^a			
State/Industry	Logging (133)	Sawmills (134)	Veneer (139)
Georgia			
133	.96	.96	.97
134	.76	.76	.76
139	.00	.00	.81
Michigan			
133	.93	.93	.93
134	.43	.43	.43
139	.00	.00	.60
Montana			
133	.93	.93	.93
134	.78	.78	.78
139	.00	.00	.76
Oregon			
133	.98	.98	.98
134	.78	.78	.78
139	.00	.00	.83

^a Data in Table 4 are RPCs from the IMPLAN Industry Balance Sheet from report [united states.iap](#). 1997 data.

considering the economic effects of changes in the wood products industry. The production function for an industry could show a particular industry to provide significant inputs. However, if the RPC is small the effects on the impact area will be small.

An alternative to using RPCs to calculate trade flows is the *Supply/Demand Pooling* approach, which is offered by IMPLAN. Supply/demand pooling assumes that local demand will get as much locally as possible; all local need that can possibly be met by local producers will be purchased locally. Since this minimizes imports it will maximize local economic activity and the resulting multiplier. The percent of local usage is based on physical capacity for the region. The total commodity supply is divided by the demand. If the resulting ratio is 0.8 then 80% of local needs will be met by local demand. If supply is greater than demand, 100% of that demand will be met by local production and the remainder is exported. Because this scenario is not realistic it will not be used in this study as the primary trade flow assumption. However, as mentioned above, it will serve as the upper limit for the RPC ratio.

The third potential assumption for trade flows is the *Location Quotient* approach. This approach is also offered by IMPLAN and based on commodity output. The location quotient equation is a fixed equation. It compares the ratios of local production to national production ratios or other base regions as desired. This implies that the base region is self-sufficient. If commodity production for a region approaches the similar proportion as the base region, the RPC approaches 1. For example if the veneer and plywood sector is 2.3% of the U.S. economy anything in excess of this in the region

being analyzed would be assumed to be exported. This again has the same weaknesses as the supply/demand pooling approach, in that it ignores cross hauling/transshipments, and consequently was not used in this study.

When dealing with multi-state regions some additional assumptions were made with respect to trade flows. IMPLAN offers three alternatives. The first is *Maximum RPC*, which says that the combined region's RPCs will be at least equal to the maximum of the individual RPCs. The second option, *First RPC* is an arbitrary system previously used by the US Postal Service that simply takes the first state from a multi-state list. The last alternative is *Average RPC*, which is a weighted average based on output for all the states in the selected region. This is the assumption that was used in this study because it is the only one that resembles reality and is the IMPLAN recommended default.

As previously mentioned in this report the Forest Service in conjunction with the Minnesota IMPLAN Group is developing a gravity model that will determine trade flows. However, this will not be available as part of the IMPLAN model until 2002 at the earliest.

Types of multipliers to be analyzed – types of effects.

One study objective focuses on the behavior of economic multipliers of a region as various selected economic characteristics are changed (e.g. population, size, geographic region etc.). Of the multipliers available, only the Type I, Type II and the Type SAM are developed by IMPLAN. They are also the most useful current method of measuring

economic impacts. The Type SAM multiplier option in IMPLAN allows the user to select what institutions to include or not include. The Type I multiplier, which is used to determine direct and indirect effects, is a prerequisite and is generated regardless of whether the Type II or Type SAM multiplier is selected.

The Type I multiplier was one of the two multipliers selected as a dependent variable in regression analysis and analyzed in this study. It was anticipated that it might display different behavior than the multipliers that include induced effects. The Type I multiplier is usually a good indicator of the development of the industry being analyzed; that is the more diverse the industry is the higher will be the Type I multiplier.

The Type SAM multiplier was chosen as the other dependent variable to be analyzed. It was selected over the Type II multiplier because it gives a more realistic view of the economy in that it can account for transactions for all institutions, including households and government transactions.

This study was limited to analyzing the Type I multiplier and the Type SAM multiplier. This is consistent with the objectives of the study as these are the multipliers most commonly used by Forest Service and other analysts at the present.

Types of multipliers to be analyzed – economic variables.

Multipliers, including the Type I and Type SAM multipliers, are commonly expressed in terms of several economic variables. These include employment, output (sales), and

value added. Value added includes employee compensation, proprietor income (self-employed income), other property income and indirect business taxes. All of these multiplier variables can be generated by IMPLAN. However, in the interest of keeping the study focused, only employment and personal income (determined by adding employee compensation and proprietor income) multipliers were analyzed. This was determined based on the demand by analysts for employment and income multipliers and the lack of demand for value added multipliers. Value added data is rarely encountered in economic impact studies and not well understood by the concerned public or many analysts determining economic impacts.

Response variables analyzed

Table 2 shows the structure of the multipliers used in the study. The first and most basic statement is that the multipliers will eventually be used as the response (dependent) variables for regression analysis in this study. The explanatory (independent) variables will be discussed later in this section.

Three economic sectors from the IMPLAN economic impact model were chosen for the study. Rationale for this selection was presented earlier in this chapter under section “Economic Sectors Analyzed” (pg 42). For each of these three sectors both the Type I and Type SAM multipliers were calculated for both the employment multiplier and the personal income multiplier. This means that up to a total of 12 multipliers were generated for each impact model. Less than 12 were generated if some of the wood products sectors were not present. Up to five impact area models were generated for each

National Forest therefore up to 60 multipliers could be developed for an individual National Forest impact area. However, every multiplier will fall into one of the 12 classes listed below. Also, each of the multipliers listed below will be analyzed separately throughout this study.

Each of the 12 different classes of multipliers was given a descriptive variable code (Table 2, pg 45). This code will be used for descriptive purposes throughout this study for the simple reason of saving space. The code is somewhat easy to follow. For example the first multiplier shown in Table 2, LE1 is a multiplier for, Logging Camps and Contractors sector for Employment, Type 1.

Explanatory variables to be analyzed

Introduction

One of the objectives of the study was to identify key factors that explain the variation of economic multipliers. The explanatory variables chosen for analysis were a result of the literature review, discussion with experts, and reviewing the industry production functions to get a better understanding of the industry structures. Another objective was to determine the feasibility of using these variables individually or in combination to predict the economic multipliers (response variables). The impact areas analyzed have been described by a number of characteristics that helped in analyzing and classifying multipliers. These explanatory variables served as the initial independent variables in the study and were analyzed using regression techniques to determine if they are useful in predicting multipliers (Table 5).

Table 5 - Explanatory Variables		
Variable	Indicator Used	Source
Human Population	Total population	1997 Population from U.S. Census estimate.
	Population Density	
Size of Economy	Total Industry Output	IMPLAN
	Total Personal Income	
	Number of Economic Sectors	
Geographic Size	Square miles	U.S. Census

Each of the variables in Table 5, were included with other information for each economic impact area analyzed, in an Excel spreadsheet. They were then transferred to MINITAB for statistical analysis. The independent or explanatory variables (Table 5) were analyzed, through regression analysis, in relation to the economic multipliers described earlier. The economic multipliers were modeled as the dependent variables.

Explanatory variables analyzed

Table 6 Explanatory Variable Codes	
Explanatory Variables (Base year-1997)	Variable Code
Population (millions) (+) ^a	Pop
Population Density (+)	PopDen
Number of Economic Sectors (+)	Sect
Total Industry Output (billions\$) (+)	TIO
Personal Income (billions\$) (+)	PerInc
Economic Impact Area (million acres) (+)	EIA
Geographic Region	Reg

^a Sign in brackets indicates the anticipated direction of change of multiplier. For example (+) means that the multiplier will be expected to increase as the value of the explanatory variable increases.

Table 6 shows the variables that have been chosen to analyze as potential explanatory variables, the unit of measure, the code used to identify the variable in this report, and the expected direction of change. Although much of the data was taken from the IMPLAN databases, it is available from alternate sources as described below. The use of IMPLAN data was mostly a matter of expedience. A brief discussion of each of the variables follows.

Population - The expectation is that as human populations increase, so will economic multipliers. This would be especially true for the Type SAM multiplier, which includes induced effects. As the population increases the economy becomes more diverse and spending leakage is reduced. As more local money is spent within the economic impact region the multipliers should increase. However, this assumption is not as obvious with the Type I multiplier. Does the wood products industry become more developed in larger impact areas therefore have larger multipliers? The analysis should answer this question.

The human population for each county was determined from the County Population Estimates Program of the U.S. Census Bureau for the year 1997. The program estimates populations along with demographic components of change for each year between decennial censuses. The year 1997 was used because it is the latest year reflected in the most recent IMPLAN model. All impact areas were either individual counties or in most cases aggregates of counties where populations were simply summed.

Population Density - As the human population becomes more dense (more people per unit of land area) how do the economic multipliers change? Prior research investigating the behavior of multipliers for the tourism industry, showed Type II multipliers increase with increasing population density (Chang 2001). It was assumed the same direction of change will occur with the wood products industry.

Human population density, in terms of people per square mile, was determined by dividing the total population, explained above for 1997, for each impact area, by the area (square miles) in each impact area. The area (square miles) was determined as explained below under *Economic Impact Area*. The U.S. Census Bureau computation of population density was not used because the latest computation was from 1990, which is not consistent with the base year 1997 used for other variables.

Number of Economic Sectors - It is assumed that the larger the number of economic sectors in the region being analyzed, the larger will be the Type SAM multiplier for that region. As sectors are added, the economy becomes more diverse and spending leakage is reduced. As more local money is spent within the economic impact region the multipliers should increase. This assumption was based on previous research (Zheng et al. 1988). However, this assumption is not as obvious with the Type I multiplier. It is not intuitive or apparent that the wood products industry becomes more developed in impact areas with more sectors. The analysis addresses this issue.

The number of economic sectors used as an explanatory variable is based on the number of sectors in the IMPLAN economic impact model. The IMPLAN basic impact area data report shows the total number of sectors for each model. A non-IMPLAN substitute for this variable could be the number of sectors at the 3-4-digit level for Bureau of Economic Analysis data. Another source would be the ES202 state and county employment and income data collected under the guidance of the Bureau of Labor Statistics, U.S. Department of Labor.

Total Industry Output - The total industry output is the total value of production by industry for a given time period – in this case, a year, for the selected impact area. It is the total value of purchases by intermediate and final consumers including exports.

The assumption is that as an economy grows, it becomes more diverse and there is less spending leakage from the economic impact area. For this reason, at least the Type SAM multiplier should increase as the total industry output increases. . This assumption was based on previous research (Zheng et al. 1988). Whether or not the Type I multiplier increases (normally a function of the complexity of the wood products industry) will need to be answered by the analysis.

The total industry output information is taken from IMPLAN Report #SA050, Output, Value Added and Employment. An alternative source of this information would be the Bureau of Economic Analysis Annual Survey of Manufactures (USDC 2001).

Personal Income - Personal income in this case, is the sum of employee compensation and proprietor income, for the selected impact area. It is assumed that as personal income increases so will population and the diversity of the economy. This will result in less spending leakage to other economies and result in larger multipliers – at least the Type SAM multiplier. It is uncertain as to whether the Type I multiplier will increase. This information is taken from IMPLAN Report #SA050, Output, Value Added and Employment. Other Property Income was not included because it does not have a direct relationship or influence on economic multipliers. A non-IMPLAN source of this data would be the REIS CA5 tables (USDC 2001).

Economic Impact Area – The physical size of the economic impact areas (square miles) was determined for each county and aggregated as necessary for each multi-county impact area. As the size of the impact area increases it is assumed that the size and diversity of the economy will increase resulting in larger multipliers – at least for the Type SAM multiplier. The direction of change for the Type I multiplier is uncertain without further analysis. The size of the impact areas was taken from the U.S. Census Bureau, Table 1. Land Area, Population, and Density for States and Counties from 1990 Summary Tape File 1C.

Geographic Variables

To address one of the study objectives (Objective 3), it was necessary to look at possible regional differences to determine if separate equations for different regions would

improve the ability to predict multipliers. Two sets of regions were identified as potentially useful. These delineations are not used as explanatory variables in the study. They are potential criteria to be used to develop sub-populations through analysis of variance to determine if there are differences between geographic regions and used as dummy variables in the model building process. They are as follows:

USDA Forest Service Regions – Impact areas and resulting multipliers were determined for all Forest Service Regions with the exception of the Alaska Region (R-10). All multipliers have been attributed by FS Region and can be reviewed for certain regional differences. These will be discussed in later sections. Forest Service Regions were chosen because they are well distributed throughout the U.S., and the benefit of this study accrues primarily to that agency (Appendix E and F).

Bureau of Economic Analysis Economic Regions - The entire U.S. is divided into Economic Regions by the Bureau of Economic Analysis . Regions are then subdivided into economic areas. All multipliers for all levels are attributed by the primary BEA Economic Regions within the economic impact area for which the multiplier applies. Also, the level five impact area (Chapter 3) is based on BEA Economic areas. Economic areas are generally centered around metropolitan areas (USDC 1995).

It was expected that the complexity of the wood products industry could change throughout the U.S. with some areas more developed than others. This could result in

better predictive models. Analysis of variance techniques were used to address this possibility. However, there were no specific expectations as to where multipliers of all types might be higher or lower or where the effects of explanatory variables differ.

Other explanatory variables considered

Another variable considered, in addition to those in Table 1, was the size of the timber industry in the impact area – both relative and absolute. However, this was dropped as it was not possible to collect this information in a reasonable time at a reasonable cost, and it would not serve as a reasonably available explanatory variable to use in lieu of impact model building. The objectives of the study require that data must be reasonably available, outside the IMPLAN system. Although, IMPLAN data was used for some of the explanatory variables, it was done for expediency reasons – the data could have been acquired from an additional source, but at additional expense to the study.

The number of sectors in the wood products industry (as opposed to all economic sectors), was also considered as an explanatory variable, and is closely related to the variable discussed above. There were several problems with this idea. The first problem was that it is difficult and time-consuming to determine how many wood products sectors there actually are in a particular impact area without using IMPLAN. The objectives of this study require alternative and reasonably available sources of data. The second problem is the degree of connectedness among wood products sectors. Based on reviewing the production functions of each wood products sector in an impact region, the

amount of economic transactions among wood products sectors is often less than half of the of the total value of purchases.

In addition to the number of economic sectors as an indicator of economic diversity, the use of the Shannon-Weaver index (USDA Forest Service 2001) as an indicator, was also considered. However, this idea was dropped because the indices are only available for individual counties and cannot be aggregated. The calculation of multi-county indices is complicated and requires the analysis of the IMPLAN multi-county data matrix. There are plans to have this automated by IMPLAN in the near future.

Another independent variable considered was the acres of timberland in an impact area and the timberland acres as a percent of total impact area acres. This variable was not used because of the unavailability of timberland data. However, this data will be available from the USDA Forest Service in the near future.

Interaction variables (cross-product terms) were also examined. For example, population was multiplied by the number of economic sectors and then evaluated as an independent variable. The new interaction terms were evaluated with and without the original terms.

Where original terms remained the model produced slightly higher adjusted R^2 (regression coefficient) values but introduced significant additional multicollinearity indicated by variance inflation factor (VIF) values greater than 5.0. The VIF is used to detect whether one explanatory variable has a strong linear association with the remaining explanatory variables (the presence of multicollinearity among the explanatory

variables). VIF measures how much the variance of an estimated regression coefficient increases if your explanatory variables are correlated. The presence of this condition resulted in eventually eliminating the interaction variables in favor of more significant individual variables. Where original terms were excluded, adjusted R^2 values were slightly lower than they were when in the original model where the original terms were separate and the interaction term was excluded.

General Form of Model

Given the data described above, it was possible to represent the models in a general form. The general linear form of a multiple regression model is:

$$multipliers_i = \beta_0 + \sum_j \beta_j * X_j + \varepsilon_i$$

where $multipliers_i$ is the value of the economic multiplier i , (Table 2), β_0 is the

intercept, $\sum_j \beta_j * X_j$ is the sum of all coefficients multiplied by the respective variable

attribute for each attribute j , (Table 5), and ε_i is the error for each multiplier i . The

focus of the study is determining the feasibility of predicting $multiplier s_i$.

Although the functional form $y = x$ (linear) was selected for the preliminary models other forms were considered. These included $\ln y = \ln x$ (double log), $\ln y = x$ (dependent) single log) and $y = \ln x$ (independent) single log.

Statistical Analysis Methods

The statistical analysis done for this study was based on techniques from several sources. They include Koutsoyiannis (1979), Gujarati (1995), Freedman et al (1991), Snedecor et al (1980), Mendenhall et al. (1994), and the MINITAB users guide (1999). The statistical analysis methods will be described as they are applied for each of the study objectives. A summary of analysis procedures leading up to statistical analysis will precede the discussion of the statistical analysis used in this study.

Study Objective 1: Describe Variations in Multipliers.

The IMPLAN economic impact system was used to generate the economic multipliers to be analyzed as dependent variables in regression analysis. Type I and Type SAM multipliers were generated to reflect indirect, induced, and total economic effects. Type I and Type SAM multipliers were generated to reflect effects in terms of employment and personal income. The multipliers analyzed were limited to the three predominant sectors of the wood products industries: logging camps and contractors, sawmills and planing mills, and veneer and plywood. Trade flows were estimated using the Regional Purchase Coefficient option in IMPLAN.

The study areas, to be used as samples, were based on National Forests with each National Forest serving as a center for up to five, different sized, impact areas. This resulted in 155 sets of impact areas for 650 separate impact areas to be analyzed. Impact areas varied from single counties to aggregates of more than 20 counties. Sets of data,

including economic multipliers, were developed for each of the 650 impact areas.

Geographic regions were established based on USDA Forest Service regions and impact areas grouped accordingly for additional comparison.

Descriptive statistics generated consisted of means, trimmed means, medians, ranges, maximums, minimums, and standard deviations of multipliers. These were computed for each type of multiplier for each sector selected for all geographic regions (listed in Table 2).

Study Objective 2: Identify the key factors that explain the variation of economic multipliers.

The second study objective is to identify the key factors that explain the variations in economic multipliers under varying conditions. The explanatory variables selected as possible factors that might help to explain the variation in multipliers were: population (human), population density, total number of economic sectors, total industry output, personal income, and the physical size of the economic impact area. Other explanatory variables considered but not used were: total acres of timberland within the impact area, percent of total acres in an impact area consisting of timberlands, the size of the timber industry, and cross-product terms. Non-linearity was tested using natural log transformations. Dummy variables were used to investigate possible regional differences. Building separate models for separate geographic regions was also investigated.

Correlation coefficients were computed between all response and explanatory variables to determine the direction and strength of relationships among these variables. This process also helped identify those variables that had the best chance for inclusion in the model developed for Objective 3. At the same time the relationship between explanatory variables and other explanatory variables was examined through correlation analysis to determine the potential for multicollinearity.

Study Objective 3: Determine the feasibility of using selected, readily available, explanatory variables to estimate economic multipliers.

Linear regression techniques were used to quantitatively evaluate the relationship between the explanatory variables and the response variables (multipliers). MINITAB Best Subsets process was used to identify the best combinations of variables to be used in potential regression models to predict multipliers. R^2 , adjusted R^2 , and Mallows' C_p values were generated by MINITAB to measure the quality of the models. Regression models were examined for the presence of heteroscedasticity by visual analysis of the residual plot graphs.

After the preliminary regression models were established for each multiplier group additional procedures were used in an attempt to improve the models. Analysis of variance procedures were used to determine if there were differences in multipliers between geographic regions (based on USDA Forest Service regions) that might improve the models. Regions or the best combinations of regions were represented as dummy variables in the regression analysis process. Additional variables such as the size of the

wood products industry and the percent of the region bearing harvestable timber crops were analyzed as explanatory variables. Natural log transformations were used to generate non-linear explanatory variables in an attempt to improve the predictive strength of the models. Multicollinearity was measured using variance inflation factors (VIF). If A VIF greater than 5.0 indicated that multicollinearity was a problem. Models were then adjusted to reduce multicollinearity.

The analysis produced 12 final models designed to predict two types of multipliers, for two categories of economic effects, for each of three wood products industries sectors.

CHAPTER FOUR: ANALYSIS AND RESULTS

Introduction

Results of the study are presented for each of the three objectives of the study. The objectives are: describe variations in the multipliers of the wood products industries, identify the key factors that explain the variations in multipliers, and determine the feasibility of using selected explanatory variables to predict multipliers. The following discussion applies to all National Forest impact regions in the U.S. unless otherwise indicated.

Study Objective 1: Describe Variations in Multipliers.

This section describes the response variables (multipliers) in terms of generally used descriptive statistics discussed in Chapter 3, Methods. These parameters were generated with the use of the MINITAB statistical software program. This information is generally referred to as “descriptive statistics” as opposed to “inferential statistics” which will be discussed under Study Objective 3, later in this chapter. The review of the data characteristics is important in order to evaluate potential use of the data for statistical inferences.

Descriptive statistics “consists of procedures used to summarize and describe the important characteristics of a set of measurements” (Mendenhall et al. 1994). The following section describes each selected characteristic for each multiplier group and each of the explanatory variables.

The relationship between the multipliers and the individual explanatory variables is reported and discussed under Objectives 2 and 3. This was done by correlation analysis and also regressing each explanatory variable on each individual multiplier in the Best Subsets procedure.

Number of Observations

As can be seen in Table 7, the number of observations (impact models) varies among the three economic sectors studied. A total of 674 impact areas were modeled in IMPLAN.

Twenty-four of these impact areas had none of the 3 wood products economic sectors therefore were not represented in the models. Almost all impact areas that had Logging Camps and Logging Contractors (650 areas) also had Sawmills and Planing mills (647 areas). In contrast, only 384 impact areas were represented by the Veneer and Plywood sector. This is economically logical as the economic threshold to enter this sector is much greater than the other sectors and the scale of operation must be much larger to continue operations. There is no such thing as a small veneer and plywood mill.

Table 7 Number of Observations

Multiplier Variable Code	Number of Observations
LE1	650
LE2	650
LP1	650
LP2	650
SE1	647
SE2	647
SP1	647
SP2	647
VE1	384
VE2	384
VP1	384
VP2	384
Pop	674
PopDen	674
Sect	674
TIO	674
PerInc	674
EIA	674

Measures of Central Tendency

Means - MINITAB calculates the arithmetic mean, or average. The mean is a commonly used measure of the center of a population or sample of numbers (Table 8). An explanation of the multipliers in Table 8 can best be given by an example. The mean multiplier for LE1 (logging, employment, Type I), is 1.39. This means that for each direct job created or lost from change in output another .39 indirect jobs will be created or lost. For example, an increase in raw material availability from a local National Forest might have the potential to create another 100 jobs in the logging industry. It would also create another 39 indirect jobs in those industries that the logging industry buys from. The multiplier for LE2 (logging, employment, Type SAM) creates an additional 32 jobs ($1.71 - 1.39$) from induced spending (household spending of wages paid by firms causing direct and indirect effects). Total direct, indirect, and induced effects would result in 171 additional jobs in the selected economic impact area.

The first procedure performed after the means were calculated was to review the relationship between the Type 1 and Type SAM multipliers. The mean Type SAM Multiplier is and should be in all cases larger as it includes the induced as well as direct and indirect effects. Also, an important procedure was to review all the resultant multipliers to see if they appeared to be reasonable for the given circumstances. This was done by ordering the multipliers from top to bottom and vice versa to search for extreme values. All multipliers of all types were found to be above 1.0 and below 4.0. Anything below 1.0 would have

indicated a technical error.

Anything above 4.0 would be possible but highly unlikely.

Multiplier means of all types are consistently larger in the sawmill and planing mills sector than the other sectors (Table 8). They are lowest in the logging sector.

Employment multipliers are consistently higher than

personal income multipliers. This might suggest that income per job for the direct and induced effects is lower on the average than the wood products sectors.

Table 8 – Measures of Central Tendency

Multiplier Variable Code ^a	Means	Trimmed Means	Medians
LE1	1.39	1.39	1.40
LE2	1.71	1.71	1.71
LP1	1.35	1.34	1.32
LP2	1.60	1.59	1.57
SE1	1.76	1.77	1.80
SE2	2.24	2.25	2.26
SP1	1.68	1.67	1.65
SP2	2.00	1.99	1.96
VE1	1.53	1.53	1.54
VE2	1.99	1.99	2.00
VP1	1.46	1.46	1.46
VP2	1.75	1.75	1.75

^a See Table 2 (pg 45) for variable code descriptions.

Trimmed Mean - A 5% trimmed mean is calculated by MINITAB, which removes the smallest 5% and the largest 5% of the values, and then averages the remaining values as it does for the arithmetic mean. The value of this calculation is to determine if there is a

Table 9 - Measures of Variability				
Multiplier Variable Code	Minimum Value	Means	Maximum Value	Range (difference)
LE1	1.09	1.39	1.59	0.50
LE2	1.18	1.71	2.13	0.95
LP1	1.10	1.35	1.87	0.77
LP2	1.18	1.60	2.26	1.08
SE1	1.35	1.76	2.05	0.70
SE2	1.58	2.24	2.91	1.33
SP1	1.23	1.68	2.99	1.76
SP2	1.38	2.00	3.51	2.13
VE1	1.24	1.53	1.69	0.45
VE2	1.47	1.99	2.44	0.97
VP1	1.15	1.46	1.90	0.75
VP2	1.33	1.75	2.29	0.96

significant effect on the arithmetic mean from the presence of unusual values/outliers.

Table 8 indicates the possibility that there is little difference between the 2 sets of means therefore suggesting that unusual values to one side of the mean or the other, might not have very much effect on the arithmetic mean.

Median – The median is the middle measurement in a set of measurements. Half the observations are less than or equal to it.

The median can be very similar or exactly the same as the arithmetic mean or it can be very different depending how individual observations (multipliers, etc.) are distributed about the mean. It is a rough indicator of whether or not a population is “normally distributed”. A comparison of the arithmetic means and medians in Table 8 shows very little difference. Graphs, generated by MINITAB, showed the distribution of median values about their means to be similar to the distribution of the value of the means.

Measures of Variability

Range – The range is the difference between the largest and smallest data value. It is one of several measures of the distribution of observations (multipliers etc.) about the population means.

The minimum and maximum value and range for each class of multiplier shows a fairly broad range of values (Table 9). However, by itself, range is not a good predictor of a normal distribution about the mean because it depends on only 2 values, which might be outliers and not very representative of the population. A normal distribution is essential in evaluating the predictive qualities of the data. The sawmill and planing mills type SAM multiplier shows the widest range and this seems reasonable because the industry exists over a wide range of conditions (population, number of sectors, etc.). Generally, the minimum values occurred in counties with small economies and the maximum values occurred in medium to large counties. The absolute smallest multiplier (1.09) is for the Level 1 area for the William F. Bankhead National Forest in a rural area of Alabama. The absolute largest multiplier (3.51) is for the Level 3 impact area for the Appalachian

National in the heavily populated Florida panhandle. The area has both a large population and well developed wood products industry.

Table 10 - Standard Deviation and Coefficient of Variation			
Multiplier Variable Code	Means	Standard Deviation	Coefficient of Variation
LE1	1.39	0.10	.07
LE2	1.71	0.17	.10
LP1	1.35	0.13	.10
LP2	1.60	0.18	.11
SE1	1.76	0.14	.08
SE2	2.24	0.26	.12
SP1	1.68	0.18	.11
SP2	2.00	0.24	.12
VE1	1.53	0.10	.07
VE2	1.99	0.21	.11
VP1	1.46	0.10	.07
VP2	1.75	0.15	.09

Standard Deviation - The standard deviation is another method, which provides a measure of how spread out the data are (dispersion). While the "range" explained above was based on only the 2 most dispersed observations, the standard deviation considers all observations and their relationship to the mean (Table 9).

Comparing the ranges from above to the standard deviations for each multiplier category shows many similarities. The major similarity is that the Type SAM multipliers show more dispersion than the Type 1 multipliers under both income and employment measures and whether expressed as an absolute (standard deviation) or as a percent (coefficient of variation) . This is probably due to the inclusion of induced effects that

vary more over a wide range of conditions. Also, sawmills and planing mills show more dispersion than the other two sectors therefore being consistent with the behavior of the means.

Figure 1 – Distribution of Multiplier Means

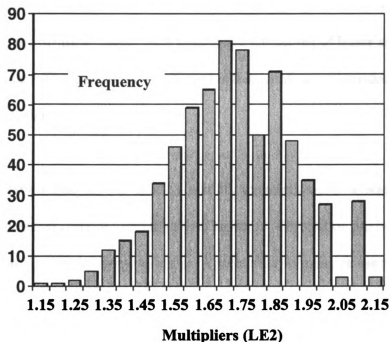


Figure 1 is an example of the normal shaped distribution of the means and the small degree of deviation.

Study Objective 2: Identify the key factors that explain the variation of economic multipliers.

While Study Objective 1 identified differences in multipliers throughout the impact areas established for the study, Study Objective 2 evaluates the relationship of the multipliers

to a variety of explanatory variables (Table 5). The reason for this analysis was to determine their potential as explanatory variables in the final regression model.

Correlation Analysis – Closely related to but conceptually different from regression analysis (to be discussed under Study Objective 3) is correlation analysis where the primary objective is to measure the degree of linear association between two variables. The correlation coefficient is the quantitative measure of the linear association (Mendenhall et al. 1994).

Correlation analysis was used to quantify the relationship between the explanatory variables and response variables (multipliers). Correlation analysis was also used to compare the explanatory variables to each other to measure the potential for multicollinearity. The presence of multicollinearity was measured by variance inflation factors (VIF).

Correlation analysis was used (MINITAB) to calculate the coefficient of correlation for pairs of variables. This was done to determine the degree of the relationships between the explanatory variables and the response variables. The correlation coefficient assumes a value between -1 and $+1$. If one variable tends to increase as the other decreases, the correlation coefficient is negative. Conversely, if the two variables tend to increase or

Table 11 – Correlation Coefficients and P-Values

Variables	Pop	PopDen	Sect	TIO	PerInc	EIA
LE1	0.193 ^a 0.000 ^b	0.020 0.612	0.517 0.000	0.193 0.000	0.193 0.000	0.463 0.000
LE2	0.172 0.000	-0.031 0.437	0.509 0.000	0.172 0.000	0.173 0.000	0.423 0.000
LP1	0.323 0.000	0.206 0.000	0.510 0.000	0.322 0.000	0.317 0.000	0.360 0.000
LP2	0.405 0.000	0.220 0.000	0.643 0.000	0.402 0.000	0.397 0.000	0.422 0.000
SE1	-0.214 0.000	-0.312 0.000	0.118 0.003	-0.212 0.000	-0.213 0.000	0.273 0.000
SE2	-0.109 0.006	-0.227 0.000	0.251 0.000	-0.105 0.008	-0.106 0.007	0.305 0.000
SP1	0.022 0.583	-0.128 0.001	0.204 0.000	0.025 0.519	0.019 0.625	0.224 0.000
SP2	0.143 0.000	-0.071 0.072	0.392 0.000	0.145 0.000	0.138 0.000	0.319 0.000
VE1	-0.309 0.000	-0.318 0.000	0.021 0.688	-0.312 0.000	-0.317 0.000	0.032 0.537
VE2	-0.205 0.000	-0.234 0.000	0.122 0.017	-0.204 0.000	-0.207 0.000	0.211 0.017
VP1	0.079 0.124	-0.052 0.312	0.296 0.000	0.075 0.145	0.067 0.193	0.102 0.046
VP2	0.236 0.000	0.024 0.644	0.535 0.000	0.233 0.000	0.226 0.000	0.248 0.000

^a Correlation coefficients are presented in the top row of data for each variable.

^b P-values are presented in the bottom row for each variable.

decrease together, the correlation coefficient is positive. In Table 10 the top number in each set of rows is the correlation coefficient. The bottom number is the p-value for individual hypothesis tests. P-values are often used in hypothesis tests where you either accept or reject a null hypothesis. The p-value represents the probability of rejecting the null hypothesis when it is true. The smaller the p-value, the smaller the probability that

you would be making a mistake by rejecting the null hypothesis. The null hypothesis in this case is that the correlation coefficients in Table 10 are zero. A cutoff value often used is 0.05, that is, reject the null hypothesis when the p-value is less than 0.05. Technically, the p-value is defined as the lowest significance level at which a null hypothesis can be rejected.

Table 11 shows that in most cases there is little chance that you would be making a mistake by rejecting the null hypothesis. However, in some cases (i.e. VE1/IS = 0.688) the p-value is quite large suggesting there is a high probability that rejecting the null hypotheses would be a mistake. Some p-values are unusually high for population density. High p-values are associated with low correlation coefficients – this is to be expected. Other than that there is no obvious pattern.

The strongest positive linear relationships, between the response variables (multipliers) and the explanatory variables (Table 5) indicated by the correlation coefficients, was between the multipliers and the number of economic sectors. This would be expected as the more diverse the economy the more spending, direct, indirect, and induced, will probably occur in the impact area. Also, consistently positive was the relationship between the multipliers and the geographic size of the area. It would be expected that the larger the area the more economic sectors it would have – on the average. This would be consistent with the findings that the more economic sectors in an impact area, the larger the multipliers. Negative correlations were dominant for population density suggesting that areas with higher population density have less developed wood products industries

although the industry was present throughout a wide variety of population densities. This would also seem logical as many wood harvesting and processing centers are in rural areas – but areas that are apparently economically diverse. Overall, negative correlations were most frequent in the sawmills and planing mills sector followed by the veneer and plywood sector. The logging camps and contractors sector was positively correlated in most all cases. The reasons for this are not apparent. The veneer and plywood sector, employment variable, for both types of multipliers, are mostly negative. In the same group but only for personal income, the correlations coefficients are mostly positive. The same pattern occurs for the sawmills and planing mills sectors. Why employment is positively correlated and personal income negatively correlated is not clear but it is fairly consistent. This might suggest that incomes on the average tend to be lower in wood products industry dominated regions. Other than the observations above there is not a significant pattern of highs and lows, positives and negatives. Why negative values exist for some variables and not others is not apparent. The lowest correlation coefficients are for industry output. This is probably due to the fact that industry output is not a good indicator of income or employment when examined in the aggregate. A region could have a large value for industry output but low income, employment, population and economic diversity resulting in low multipliers – or just the opposite could occur.

Although the relationships of individual variables might show small correlation coefficients, their predictive powers could be significant when combined with other explanatory variables in multiple regression analysis. In summary, Table 11 shows that

there are definitely relationships between the variables selected and further analysis is justified.

The next step was to look at the relationships among the explanatory variables to investigate the potential for multicollinearity. Multicollinearity exists where explanatory variables overlap. For example the relationship between an explanatory variable such as population and a particular type of

economic multiplier may be very similar to another explanatory variable such as total industry output or the number of economic sectors. The implication of this condition generally means that the confidence in the models ability to predict multipliers will be less. Sometimes one variable can predict as well as two or more variables and the fewer variables necessary the less the cost of collecting and managing data.

Correlation coefficients were estimated for each pair explanatory variables (Table 12). It can be seen that there are several high correlation coefficients therefore potential multicollinearity. This is most noticeable between the number of economic sectors , total industry output , and personal income. The only negative correlation is between

Table 12 – Correlation Coefficients Between Explanatory Variables

	Pop	PopDen	Sect	TIO	PerInc
PopDen	0.376 ^a				
	0.000 ^b				
Sect	0.628	0.301			
	0.000	0.000			
TIO	0.992	0.396	0.620		
	0.000	0.000	0.000		
PerInc	0.994	0.387	0.615	0.998	
	0.000	0.000	0.000	0.000	
EIA	0.365	-0.157	0.491	0.346	0.347
	0.000	0.000	0.000	0.000	0.000

^a Correlation coefficients are presented in the top row of data for each variable.

^b P-values are presented in the bottom row of data for each variable.

population density and geographic area and this value is low. This would seem logical because areas that are physically large, normally have less dense populations. P-values were low indicating a high level of confidence in the coefficient values. This multicollinearity problem will be dealt with during the regression analysis process. All of the variables in Table 11 were carried into regression analysis.

In summary, it is expected that the number of economic sectors and population density will serve as the best explanatory variables in predicting multipliers of all types. Population, industry output, and total income should increase the strength of the model but may represent the same conditions and cause excessive multicollinearity.

Measures of Normality - In order for standard statistical assumptions to be valid in making statistical inferences, the response (dependent) variables need to be “normally” distributed about their means. There are several methods in assessing the normal distribution. They include comparing the means to the medians, determining the range of the observations, and standard deviations. However, there are several other more accurate methods of determining the degree of normality. The first method is by simply reviewing histograms for each variable to see how closely they resemble a bell curve. This was done for all response variables, and the bell shape pattern of the histogram existed in all cases. Figure 1 is an example of the normal distribution of multiplier values.

The second method is more complex and systematic. The method used was the Ryan-Joiner test, which is a correlation-based test (available in MINITAB), similar to the commonly used Shapiro-Wilke test (not available in MINITAB). The graphical output is a plot of normal probabilities versus the data. The input data are plotted as the x-values. Then the probability of occurrence is calculated, assuming a normal distribution, and then calculated probabilities are plotted as y-values. The grid on the graph resembles the grids found on normal probability paper, with a log scale for the probabilities. A least-squares line is fit to the plotted points and drawn on the plot for reference. The line forms an estimate of the cumulative distribution function for the population from which data are drawn.

In this study, the data departed from the fitted line most evidently in the extremes, or distribution tails. In most cases there was a slight tendency for the data to be lighter in the tails than a normal distribution because the smallest points are below the lines and the largest points are just above the lines. The test indicated that all multiplier classes showed a high degree of normality with all R^2 values above 0.96 and several above 0.99 and most p-values less than 0.00100. In summary, the data distributions were found to be highly normal and fit for developing statistical inferences using regression analysis without any need for scale transformation (log scale, square root etc) or further analysis.

Study Objective 3: Determine the feasibility of using selected, readily available, explanatory variables to estimate economic multipliers.

Introduction

One of the primary objectives of this study was to investigate the feasibility of estimating economic multipliers using variables such as population and industry output (Table 5).

This would allow the analyst to enter a few variables into an equation and generate various types of multipliers to estimate the economic impacts of a particular decision for a particular area.

This analysis used “inferential” statistical methods as opposed to “descriptive” methods as previously described in study objectives 1 and 2. However, descriptive statistics have been included in the report to give a better understanding of the characteristics of the variables. “Inferential statistics” consists of procedures used to make inferences about population characteristics (Mendenhall et al. 1994). In this particular case we would “infer” a particular multiplier from characteristics of the population being studied.

There are certain assumptions that must hold when constructing linear least squares models (1995). Each of these ten assumptions was addressed in respect to this study to validate the models constructed (Appendix A).

Procedures and results

Linear regression analysis was used to determine the relationship between the explanatory (independent) variables and the response (dependent) variables. The following sections will describe the processes used to eventually develop predictive/probabilistic models through the use of regression and other associated techniques.

Regression Analysis

Regression analysis is concerned with the study of the dependence of one variable (in this case – economic multipliers), on one or more other variables, the explanatory variables, (in this case population, number of economic sectors, etc.) with a view to estimating and or predicting the mean value of the former in terms of the known or fixed values of the latter (Gujarati 1995). The end product of regression analysis is a multivariate equation useful in predicting the dependent variable along with the associated probability distribution. This is different than correlation discussed earlier where the end product was a measure of the strength of the association between the dependent variable and an explanatory variable with no attempt at prediction.

The “least squares” linear regression method was used (as opposed to “logistic regression”) because of the continuous nature of the dependent variable (multipliers).

The least squares method estimates parameters in the model so that the fit of the model is optimized by minimizing the sum of squared errors. This means that the linear regression line is positioned so the sum of the distances between the actual measurement and the regression line are equal above and below the regression line.

Best Subsets Process - The first step in the regression analysis process was to evaluate the explanatory variables using the MINITAB Best Subsets process. Best Subsets regression generates regression models and then selects the two models giving the largest adjusted R^2 . MINITAB displays information on these models, again examines all two-predictor models, and again selects the two models with the largest adjusted R^2 .

The process continues so all possible combinations are evaluated. In this analysis each dependent variable (multiplier) is evaluated (regressed) with respect to every combination of explanatory variables – individually and in combinations. The Best Subsets regression procedure can be used to select a group of likely models for further analysis. The general method is to select the smallest subset that fulfills certain statistical criteria. The reason that you would use a subset of variables rather than a full set is because the subset model may actually estimate the regression coefficients and predict future responses (multipliers) with smaller variance than the full model using all the explanatory variables. Also, the more variables used, the more expensive the process of generating, analyzing and managing data.

The statistics R^2 , adjusted R^2 , Mallow's C_p , and s (square root of MSE) were calculated by the Best Subsets procedure and were used for comparison criteria.

Normally, analysts might only consider subsets that provide the largest R^2 value.

However R^2 always increases with the size of the subset therefore R^2 is most useful when comparing models of the same size. This study uses adjusted R^2 and Mallow's C_p to compare models with different numbers of explanatory variables. The adjusted R^2 is the regular R^2 adjusted by the degrees of freedom. Mallow's C_p , which is used to determine goodness of fit of the model, is the ratio of the standard error of the model divided by the number of independent (explanatory) variables in the equation. The lower the C_p the better the model. Generally, it is good to have a C_p value lower than the number of response variables in the model (Draper et al. 1981).

In general this study favored models where C_p is small and close to or less than the number of explanatory variables in the model. In the case of this study it should be between 1 and 6 for the range of models being evaluated. A small value of C_p indicates that the model is relatively precise (has small variance) in estimating the true regression coefficients and predicting future responses.

The output for the Best Subsets process is quite lengthy (12 tables) and is presented as Appendix B. The analysis of each dependent variable (multiplier) generates a separate report showing the results of the regression analysis in respect to the explanatory variable(s) in all possible combinations. Using adjusted R^2 and Mallows' C_p as criteria the best model was picked for each multiplier. Table 13 presents a summary of each of the best models.

Table 13 – Best Subsets Regression Analysis Results Summary											
Multiplier Models	Number of Variables	R^2	R^2_{adj}	C_p	S	Pop	Pop Den	Sec	TIO	Per Inc	EIA
LE1	4	37.6	37.3	4.9	0.07702	X		X		X	X
LE2	6	36.6	36.0	7.0	0.13587	X	X	X	X	X	X
LP1	3	28.9	28.5	3.5	0.10629		X	X			X
LP2	3	43.6	43.3	3.3	0.13769		X	X			X
SE1	5	27.5	26.9	5.3	0.12202	X	X	X		X	X
SE2	5	27.6	27.0	5.1	0.21814	X	X	X		X	X
SP1	5	10.7	10.0	6.5	0.17518		X	X	X	X	X
SP2	5	21.9	21.3	6.8	0.21701		X	X	X	X	X
VE1	4	25.3	24.6	3.5	0.08366		X	X	X	X	
VE2	4	20.0	19.2	4.1	0.18610	X	X	X			X
VP1	6	16.2	14.8	7.0	0.09601	X	X	X	X	X	X
VP2	5	34.2	33.3	5.6	0.11940	X	X	X	X	X	

The Best Subsets process does not generate a regression equation, but yields an R^2 and adjusted R^2 value and shows which variables have the best opportunities to produce the most useful and efficient equations. Table 13 shows a range of adjusted R^2 values from 43.3 for LP2 to only 10.0 for SP1. With the exception of LP1 and LP2, all multiplier groups have better adjusted R^2 values using four or more explanatory variables. The number of economic sectors appears in all multiplier groups. Population density is only absent from LE1. At the other extreme is total industry output that is only present in 6 multiplier groups. This is consistent with what was found in correlation analysis. Table 13 also shows that most of the models have a Cp value slightly greater than the number of variables.

Developing the Regression Equation – The next step in the model building process was to develop the preliminary regression equation for each of the Best Subsets of variables in Table 12. Because of the length of the MINITAB output for each regression equation and associated data only a summary is presented in Table 14. The detailed tables are shown in Appendix C.

All type SAM multipliers were, as they should be, larger than Type 1 values.

Table 14 - Regression Equations -- Preliminary Model Coefficients

Multiplier Model	Intercept (Beta)	Explanatory Variables					
		Pop (mil)	Pop Dens	Sect	TIO (mil\$)	PerInc (bil\$)	Area (mil ac)
LE1	1.23	-.0512		.000517		.000002	.00129
LE2	1.42	-.1000	-.000089	.001010	-.137	.000006	.00165
LP1	1.18		.000092	.000474			.00103
LP2	1.29		.000095	.000961			.00136
SE1	1.64	-.0788	-.000210	.000508		.000002	.00139
SE2	1.93	-.1640	-.000311	.001210		.000005	.00223
SP1	1.55		-.000181	.000472	.248	-.000006	.00089
SP2	1.69		-.000238	.001070	.305	-.000007	.00123
VE1	1.44		-.000140	.000397	.104	-.000003	
VE2	1.75	-.0300	-.000202	.000913			.00081
VP1	1.33	.0374	-.000103	.000484	.157	-.000005	.00044
VP2	1.46	.0322	-.000123	.000986	.167	-.000005	

The population coefficient was negative in most of the models it was included. This means that as the population increased, multipliers decreased. This result seems reasonable in that most timber industry activities are in less populated areas. Exceptions were VP1 and VP2 where multipliers increased slightly with increasing population. However, it is expected that veneer and plywood plants be in larger areas than the other sectors.

The population density coefficient was negative in most of the models it was included (included in 11 of 12 models). This means that areas with higher population densities will have lower multipliers. This seems reasonable in that most wood products operations are in rural regions. This is consistent with what was found in correlation analysis.

The coefficient for the number of economic sectors occurs in all models and is positive in all cases. This means that the more total economic sectors a region has the higher will be the economic multiplier

The coefficient for total industry output occurs in only 6 models. Five of the six models are positive meaning that as total industry output becomes larger the multipliers increase. The primary reason for its exclusion in half the models is that it is highly collinear with population and income and its exclusion lowers the VIF of the model.

The regression coefficient for total personal income occurs in 9 of 12 models. Four are positive and five are negative. There is no apparent reason why some are positive and some are negative in light of the fact that closely related variables and are substantially positive.

The regression coefficients for the physical size of the analysis area occur in 10 of 12 models and in all cases vary positively with the multipliers. As the physical size of the analysis area increases so does the multiplier. This would be expected - the economy would, on the average, get larger and more diversified (more sectors) as the size of the area increased.

The Type SAM multiplier is larger than the Type 1 multiplier in all cases. This is because the SAM multiplier includes induced effects as well as the indirect effects.

Statistical parameters that explain statistical significance were calculated for each term in the regression equation. Because of the large volume of data, the statistical parameters for the model coefficients shown in Table 14, are included in Appendix C. Each of these parameters is briefly discussed as follows:

P-Value – The p-value was used to determine which of the effects in the model are statistically significant. The p-value was calculated for each of the regression equation coefficients. The traditional procedure was to:

- Identify the p-value for the effect or coefficient you want to evaluate.
- Compare this p-value to your selected significance level. In the case of this study the significance level is .05.
 - o If the p-value was less than or equal to .05, conclude that the effect is significant – therefore there is a significant linear effect for the coefficient of interest (e.g. population).
 - o If the p-value was not less than .05, there is no significant effect, that is the multiplier does not change with changes in the coefficient of interest (e.g. personal income).

VIF - VIF (variance inflation factor) measures how much the variance of an estimated regression coefficient increases if your explanatory variables are correlated. A VIF in excess of the number of independent variables, as a general rule, is deemed unacceptable. For this study a VIF greater than 5.0, a common level used in practice, is

deemed to be excessive. How excessive VIF values were dealt with is explained in Chapter 3.

Because of the large number of samples taken, p-values were almost always less than 0.05 and therefore deemed significant at the 5.0 percent level for the preliminary regression equations in Table 14. P-values can be seen in Appendix C. In brief, coefficient values are likely to be what is shown in Table 14. For this reason only the few exceptions where p-values exceeded 0.05 will be briefly discussed. For Type 1 multipliers only VP1 for the area variable exceeded 0.05. This is consistent with the result of it also being the only negative “area” variable coefficient for the preliminary models. The area variable coefficient is positive for all the other models. Other coefficients that are non-significant are VP2 - population coefficient with a p-value of 0.110; LE2 – output coefficient with a p-value of 0.056; VE2 – area coefficient with a p-value of 0.131.

Improving the Regression Equation – After the preliminary equations/models shown in Table 14 were established additional variations were examined to improve their predictive strength. Several factors were taken into consideration. These include possible regional differences, eliminating some of the observations, testing for natural log transformations, and reducing multicollinearity. These will be discussed separately in the following sections.

Regional Differences/Adjustments – In an attempt to improve the models, possible regional differences were examined. Analysis of variance (ANOVA) was used to explore the possibility of subdividing the observations into geographic regions to determine if there are real differences between regions.

ANOVA is similar to regression in that it is used to investigate and model the relationship between a response variable (multipliers) and one or more explanatory variables. However, analysis of variance differs from regression in that the independent variables are qualitative (categorical) (Minitab 2000).

To look at possible regional differences in multiplier observations, one-way analysis of variance was used. One-way analysis of variance tests the equality of population means when classification is by one variable (such as Forest Service (FS) Regions). The classification variable, or factor, usually has three or more levels, where the level represents the treatment applied. In this case there were seven FS Regions. FS Regions were set up as dummy variables for regression analysis. The regions were based on USDA Forest Service Regions, because they are discretely distributed throughout the continental U.S. The resulting geographic regions can be seen in Appendix D. The states in each FS Region are in Appendix E and F.

Multiple comparisons of means allowed for the examination of which means are different and to estimate by how much they are different. These differences can be examined in

Appendix D of this report. Because of the need to examine all pairwise comparisons of means, Fisher's least significant difference (LSD) method was used.

All multipliers benefited from regional differentiation. All multiplier groups were separated into either two or three sub-groups. Although differences in means were quite small the standard deviations were also quite small allowing for a separation of regions.

Table 15 shows adjusted R^2 values before and after regional differentiation. Absolute and percent changes were substantial in some cases. In other cases changes were minimal but in all cases changes were positive. There is no clear pattern as to what caused some areas to benefit more than others other than the wood products industry in some regions must increase with the size (population, income etc) of the region on a relatively predictable basis.

Multiplier means are different enough from region to region that stratification reduces the variation enough, with the result being better regression results.

Some of the more notable differences between regions disclosed through analysis of variance are as follows.

Table 15 – Effects of Regional Stratification on R- Square Value (adjusted).

Multiplier Group	Before	After	%Change
LE1	37.3	50.5	35.4%
LE2	36.0	56.2	56.1%
LP1	28.5	47.2	65.6%
LP2	43.3	55.2	27.5%
SE1	26.9	30.5	13.4%
SE2	27.0	34.3	27.0%
SP1	10.0	28.4	184.0%
SP2	21.3	35.9	68.5%
VE1	24.6	34.9	41.9%
VE2	19.2	35.2	83.3%
VP1	14.8	17.4	17.6%
VP2	33.3	35.2	5.7%

Employment multipliers for the *logging camps and contractors* sector (LE1, LE2) are significantly lower in the east than in the west. Type SAM multipliers follow the same pattern as the Type I multipliers therefore attributing the regional differences to the Type I multiplier. This could be a result of a lesser-developed logging industry in the east. Also, with counties being smaller in the east so would the impact areas likely be smaller. This study has already shown that areas that are geographically smaller will likely have smaller multipliers. The personal income multipliers for the logging sector (LP1, LP2) follow a similar pattern but with one major difference – the northwest (FS Regions 1 and 6) are lower than the east. The Type SAM multipliers again follow the same pattern as the Type 1 multiplier. This could be due to low income to output ratios in the northwest from advanced technology and the effect of low local wage structure in parts of the northwest. Logging industry multipliers in the southwest remain high in all cases. This could be due to the extremely large impact area size in the southwest.

As previously mentioned, multipliers for the *sawmills and planing mills* sector, as a group, are higher than multipliers for the logging, and veneer and plywood sectors. The multiplier patterns for sawmills and planing mills are somewhat similar to those for the logging sector. Employment multipliers for sawmills and planing mills (SE1, SE2) for the northeast are lower than the other regions. This is probably due to the same reasons logging sector multipliers are lower. The southwest and California regions are also quite low. Again, Type SAM multipliers follow the same pattern as Type I multipliers. This suggests there is little difference in the induced component of the multipliers. Personal

income multipliers (SP1, SP2) do not follow the same pattern as the employment multipliers but loosely follow the same pattern as the personal income multipliers for the logging industry. An exception is that the Rocky Mountain region has relatively high personal income multipliers while having below average employment multipliers. This difference might be the result of high income to output ratios for the sawmills and planing mills sector.

Employment multipliers for the *veneer and plywood* sectors (VE1, VE2) are lowest in the arid regions of the southwest and the Rocky Mountains (FS Regions 2,3,4). The industry is small and not well developed or diverse, due to the scarcity of raw materials (large diameter logs). The large impact areas do not compensate for the lack of industry development. Employment multipliers are highest in the northeast and south where the industry is well developed due to the proximity to raw materials. Personal income multipliers (VP1, VP2) have similar patterns to the employment multipliers. They are highest in the east and south and lowest in the central Rocky Mountains and Great Basin. Patterns for the Type SAM multipliers, for both employment and personal income of all types, from region to region are similar to the Type I multipliers therefore being consistent with the other wood products sectors studied.

In summary, the differences in the total (SAM) multipliers, throughout all three wood products sectors, are primarily a result of differences in the direct (Type 1) multipliers. This is evident because the Type SAM multipliers follow the same pattern from region to region as the Type 1 multiplier. This suggests there is little difference in the induced

multiplier component of the Type SAM multipliers of all types from region to region. It also appears that differences in impact area sizes affect the size of multipliers from region to region. Some regions have much larger impact areas (larger counties) on the average and consequently higher multipliers.

Another approach in lieu of using dummy variables to investigate regional differences originated from the review of the U.S. map in Appendix E. One of the most obvious features is that the geographic size of counties in the west is much greater than the east. This would cause impact areas in the west to be larger than the east and might be the cause of differences in economic multipliers. This difference was investigated by dividing the sample data (multipliers) into two separate populations – east and west. The east consisted of Forest Service Regions 8 and 9. The west consisted of Forest Service Regions 1,2,3,4,5, and 6. Separate models were developed for the east and west for each multiplier group – 12 models for the east, 12 for the west. These models were then compared to the results in Table 15. In almost all cases the dummy variable approach to regional differentiation produced better models. Adjusted R^2 values were higher for all but one model. P-values were equally significant. Another benefit of the dummy variable approach is that there is just one set of 12 models instead of 24 models. For the above reasons the idea of separate sets of regional models to account for different regions was not pursued any further.

Elimination of Unusual Observations - Another strategy considered to improve the model, involving geographic adjustments, was to look at unusual observations (outliers) that might not contribute to the model and could be eliminated without endangering the objectives of the study. A visual review of scatter plots of regression residuals was done for all models. A formal method (e.g. Studentized Residual) was deemed unnecessary as the outliers appeared to be obvious in the scatter plots. The review showed extreme values for the impact areas in Southern California. The explanatory variables such as population, total industry output, and total personal income were significantly higher than other areas. To determine the effects of this situation, the impact areas in Southern California were eliminated from the study population and the regression analysis was repeated. The new adjusted R^2 values showed only a minimal improvement – in most cases less than 1%. It was assumed that although these values were extreme there were too few of them to make a difference. Given the small effects of excluding the outliers, the Southern California impact areas were kept in the study.

Non-linear Alternatives – Further efforts to improve the model included natural log transformation to test the possibility of non-linearity. Independent (explanatory) and dependent (response) variables were tested individually and in unison. Of special interest were the population, industry output, and personal income variables. It was originally assumed that these explanatory variables were positively correlated with the multipliers. The results of the correlation analysis (Table 11) show differently with either low or negative correlation. Because the number of economic sectors showed strong positive correlation with the multiplier values it was considered a possibility that the other

economic magnitude variables might show the same results if transformed to natural log form. This transformation did result in most of the signs changing from negative to positive. However, improvement (less than 2% adjusted R^2) resulted in only a few of the models. Most models showed no improvement or a reduction in adjusted R^2 values. The most significant problem was a substantial increase in p-values. Almost all models resulted with non-significant natural log variables. This further substantiated the position that the proper functional form is linear.

Potential non-linearity was also tested by graphical analysis of residuals. These tests also indicated the best overall fit and functional form was linear.

Interaction effects (cross-product terms) were also tested. For example, population was multiplied by the number of economic sectors and then evaluated as an explanatory variable. This was done for all combinations of two explanatory variables. Three-variable combinations were not attempted because of the inability to explain the interactions. The new interaction terms were evaluated with and without the original terms. Where original terms remained, the model produced slightly higher R^2 values but introduced significant additional multicollinearity indicated by variance inflation factor (VIF) values greater than 5.0. The presence of this condition resulted in eventually eliminating the interaction variables in favor of more significant individual variables. When interaction terms and original terms were retained in the model the same situation again occurred. Given the presence of multicollinearity and the limited increase in

significance led to eventually eliminating the interaction variables in favor of more significant individual variables.

Additional Explanatory Variables - After completion of the preliminary models there was evidence that additional explanatory variables might improve the model. These ideas came from further literature review and consultation with experts. These variables were 1) total acres of timberland within the impact area and 2) percent of total acres in an impact area consisting of timberlands.

It was hypothesized that the more heavily timbered an area was the more developed the timber industry would be and consequently have higher economic multipliers - at least the Type I multiplier. Because of the labor involved in populating a database with timberland areas for over 700 impact areas, the analysis was done on a sample basis (242 impact areas). The USDA Forest Service Forest Inventory and Assessment (FIA) Data Base was the source of the area of timberlands. The FIA program is an inventory of vegetation on all lands in the U.S. Total area (acres) of the economic impact areas has already been discussed in a previous section of this report.

The acres of timberland and the percent of acres of timberland were analyzed using the same methods used to analyze the other explanatory variables. Best Subsets regression was used to test the new variables by themselves and in conjunction with all combinations of the original explanatory variables. There was no clear evidence, in the aggregate, that the additional variables significantly improved the existing models.

Adjusted R^2 values were only slightly improved (less than 2%) for some of the multipliers and actually decreased for others. Even where improvement occurred, it was done at the expense of increased VIFs - substantial in some cases. For this reason the timberland acres and percentages were not used as explanatory variables in the final models.

The National Resources Inventory (NRI) coordinated by the Natural Resources Conservation Service (NRCS) was evaluated as a source of timberland data. However, the survey does not include Federal lands therefore it is not useful for this study.

Another approach to improve the regression models was to treat the different impact area levels (Level A – E) as dummy variables. The results were mixed. In several of the models the adjusted R^2 increased by as much as 8 %. Some increased by no more than 2%. Multicollinearity was not a problem as VIF values were well below 5.0. However, p-values were very large for the Level regression coefficients with most all well exceeding the .05 significance level. Therefore the final models in this study have not been adjusted to include this idea.

Multicollinearity and Heteroscedasticity - Correlation analysis, indicated a great degree of correlation among several of the explanatory variables. These included total personal income, total industry output, and total economic sectors (Table 12). The regression analysis process in MINITAB also generates a variance inflation factor (VIF), which is used to detect correlation amongst variables (i.e. multicollinearity). VIF measures how

much the variance of an estimated regression coefficient increases if your explanatory variables are correlated. A VIF of 1.0 indicates no correlation; greater than one, otherwise. It has been suggested that when the VIF is greater than 5 to 10, the regression coefficients standard errors are inflated (Gujarati 1995). If this occurs recommendations to mitigate these effects include: collecting additional data, deleting explanatory variables that are causing the multicollinearity, using different explanatory variables, or an alternative to least squares regression (Montgomery et al. 1982). The most expedient alternative, as a first step, is to reduce the VIF to acceptable levels by systematically excluding those explanatory variables that are the greatest cause of multicollinearity (Table 16).

The VIF factors for each explanatory variable for each final model are included in Appendix G of this report. The next step in the analysis was to eliminate the explanatory variable with the highest VIF and then re-determine the regression equation to see if VIF values are all below 5.0 without a substantial reduction in the adjusted R^2 value. Table 16 is a summary of the results of adjusting the model to reduce VIF values to less than 5.0. Detailed data is presented in Appendix G.

Table 16 – Multicollinearity Analysis

Multiplier Group	Explanatory Variables Eliminated	Explanatory Variables Retained	Adjusted R^2 Before	Adjusted R^2 After
LE1	PerInc	PerInc,Sect,Area, RegDum	50.5%	50.5%
LE2	PerInc	PopDen,Sect,TIO Area,RegDum	56.2%	56.1%
LP1	None	PopDen,Sect, Area,RegDum	47.2%	47.2%
LP2	None	PopDen,Sect Area,RegDum	55.2%	55.2%
SE1	PerInc	Pop,PopDen,Sect, Area,RegDum	30.5%	29.8%
SE2	PerInc	Pop,PopDen,Sect, Area,RegDum	34.3%	33.4%
SP1	PerInc	PopDen,Sect,TIO Area,RegDum	28.4%	28.4%
SP2	PerInc	PopDen,Sect,TIO Area,RegDum	35.9%	35.8%
VE1	TIO	PopDen,Sect, PerInc,RegDum	34.9%	34.2%
VE2	TIO	Pop,PopDen,Sect, RegDum	35.2%	35.2%
VP1	TIO, PerInc	PerInc,PopDen, Sect,Area, RegDum	17.4%	15.0%
VP2	TIO, PerInc	Pop,PopDen,Sect, RegDum	35.2%	33.9%

Only in two cases (VP1, VP2) did more than one variable need to be eliminated to reduce the VIF value to less than 5.0. Eliminating total personal income was the most common remedy. Regression coefficients either did not change at all or only declined slightly in all cases. Therefore, the resulting models are significantly free of multicollinearity.

Also of concern is the possible presence of heteroscedasticity. Heteroscedasticity generally means that the residual values (differences between the regression line and the

actual measured values) are not evenly distributed over the range of values of the independent variables. This was addressed by a graphical analysis of the residuals for each model. In most cases heteroscedasticity was not a problem. An exception existed at the extremes of the regression model where extreme values existed (e.g. large population in Los Angeles County). Other statistical parameters, remained relatively constant (Appendix G).

The Final Model – Table 17 shows a summary of the final models. Details of the final regression analysis are included in the report as Appendix G. All SAM multipliers, as they should, remained larger than their Type I counterpart.

The population coefficient remained negative in all models. The population density

Table 17 – Final Model Coefficients

Multiplier Model	Explanatory Variables							
	Intercept-Beta	Pop (mil)	PopDen	Sectors	TIO (bil\$)	PerInc (bil\$)	Area (mil ac)	RegDum Variable
LE1	1.25	-.0110		.000603			.000387	-.0855
LE2	1.48		-.000120	.001180	-.3240		-.000380	-.1840
LP1	1.01		.000072	.000476			.000748	.1020
LP2	1.08		.000072	.000964			.001040	.1190
SE1	1.53	-.0228	-.000193	.000473			.001280	.0626
SE2	1.78	-.0459	-.000326	.001270			.001000	.0933
SP1	1.25		.000157	.000588	-.0552		.000310	.1470
SP2	1.42		-.000223	.001100	-.1110		.000846	.1560
VE1	1.32		-.000130	.004000		-.0000		.0683
VE2	1.24	-.0277	-.000220	.001040				.2330
VP1	1.15		-.000076	.000478		-.0000	.000041	.0920
VP2	1.35	-.0075	-.000123	.000995				.0583

coefficient remained negative in 9 of the 11 models it occurred, which is similar to the preliminary model. The coefficient for the number of economic sectors occurred in all 12 models and remained positive in all models. Total industry output coefficients remained in three of the final models and all those coefficients were negative. This reduction was a result of efforts to reduce multicollinearity. The personal income coefficient only remained in 2 of the 12 final models and had very slight negative values. The absence of the personal income coefficient was also a result of reducing multicollinearity. The regression coefficients for the physical size of the area occur in 9 of the 12 models. All but one of the coefficients (LE2) are positive.

The regional variable coefficients R which were based on USDA FS regions, and included in the regression analysis as dummy variables, are included in all models.

Reliability parameters such as standard errors, p-values and variance inflation factors indicate that the R^2 values and adjusted R^2 values are highly reliable (Appendix G). This is mostly due to a large number of samples. Enlarging the survey would add very little to the confidence in the analysis. However, R^2 values and adjusted R^2 values are quite low, indicating a large portion of the variation is unexplained (Table 16). Only three multiplier groups have R^2 and adjusted R^2 values exceeding 50%. Three groups are less than 30% with one of those at 15%.

The Logging Camps and Contractors sector has significantly higher R^2 and adjusted R^2 values than the other sectors (> 50%). This could be because the logging industry is

more consistently found in small rural economies than are the other two sectors, which are usually found throughout a variety of circumstances. SAM multipliers have higher R^2 and adjusted R^2 values than Type 1 multipliers. This is probably because the Type SAM multiplier includes the effects of induced spending (i.e. spending by households) which increases with increases in variables directly related to the size of the economy (i.e. personal income, number of sectors, population). There is no clear pattern to the relationship between employment multipliers and income multipliers. As previously mentioned, high income does not necessarily mean a high number of jobs and just the opposite can also be true.

Statistical parameters that explain statistical significance were calculated for each term in the regression equation. Because of the large volume of data, the statistical parameters for the model coefficients shown in Table 17, are included in Appendix G. Definitions are covered earlier in this chapter. As in the preliminary models, p-values are quite small in relation to the chosen significance level of 0.05. The largest p-value is for VP1 for the area variable coefficient at 0.892. Others are : LE2 – area with p-value of 0.139; SP1 – industry output with p-value of 0.406; SP1 – area with p-value of 0.346; SP2 – industry output with a p-value of 0.184. The conclusion is that in most cases there is high probability that most of the coefficients in the models are correct. The remainder, with large p-values, do not meet the standards for confidence.

Absolute Percent Error

In addition to the measures taken above, the absolute percent errors were calculated to provide further evidence as to the validity of the regression models. The results were

examined in respect to both mean and maximum percent error. This was done by taking the difference between the IMPLAN generated multipliers and comparing them to the multipliers generated by the regression models (Table 18). The mean absolute percent error is the mean of the percent differences for each multiplier group. The maximum absolute percent error is the individual maximum percent difference for each multiplier group.

Table 18 - Mean and Maximum Absolute Percent Error		
Multiplier Group	Mean Absolute Percent Error	Maximum Absolute Percent Error
LE1	3.9%	17.9%
LE2	5.3%	33.1%
LP1	5.0%	29.0%
LP2	5.7%	31.0%
SE1	5.4%	36.0%
SE2	7.5%	48.7%
SP1	6.5%	45.1%
SP2	6.8%	45.5%
VE1	4.1%	19.4%
VE2	7.5%	36.6%
VP1	4.8%	22.9%
VP2	5.0%	28.9%

The results in Table 18 are somewhat similar to the regression results in Table 17 in some respects and very different in others. The mean absolute percent errors are generally small (good) varying from a low of 3.9% for the LE1 Group to a high of 7.5% for both the SE2 and VE2 Groups. Maximum percent errors vary from a low of 17.9% for LE1 to a high of 48.7% for SE2. The sawmill sector has the highest mean absolute percent errors, the highest maximum absolute percent errors, and the lowest R^2 values. While the logging sector had the highest R^2 values it had the highest (poorest) absolute percent

errors. The absolute percent errors are higher for the Type SAM multipliers than for the Type I multipliers. This is just the opposite of the R^2 values.

In general the mean absolute percent errors are quite small. However, most of the maximum percent errors are quite high. This suggests that there is high potential for variation about the mean values and the potential for substantial error in the prediction models.

Examining Differences Between Geographic Impact Levels. Previous sections of this analysis demonstrated the relationship between economic multipliers and an array of explanatory variables. The purpose was to analyze their potential in measuring the values of various types of economic multipliers. The intent was to determine the feasibility of building statistical regression models that might forgo the need to build formal input-output impact models.

The objectives of this study, stated in the beginning of this report, also included explaining the differences in impact area characteristics (multipliers and explanatory variables) as they expanded from impact region Level A to Level E around each National Forest impact centroid and how they might be used in designing future impact areas.

Each National Forest was the geographic focal point for the design of 5 separate economic impact areas. The criteria used to design these areas are described earlier in this report but generally it was expected that they would increase in physical size from

Level A to Level E. Initial review showed that Level D (FS TSPIRS criteria) was an exception in that they were on the average larger than level E but also showed great fluctuation. For these reasons Level D was not included in this part of the study.

The initial idea was to visualize the impact areas growing as a series of concentric circles, which have been historically used to explain the concept of central place theory (Isard 1975). However, political boundaries, which are also data area boundaries, did not in most cases, produce a configuration resembling concentric circles. The addition of a county(s) produced results that were unpredictable. They varied from sparsely populated areas to major metropolitan areas with little relationship to the preceding impact area.

On the average, the expected patterns held true. As the geographic area was enlarged the multipliers became greater. This is consistent with regression analysis findings where the coefficient for the size of the geographic area was positive in 11 of 12 models (Table 17). With the exception of population density, all of the explanatory variables increased on the average along with increases in geographic area. This is consistent with the correlation analysis done to validate potential explanatory variables (Table 11).

Population density would be expected to be higher in Level A than B because A is the most populated county in the Level B aggregate.

The discussion above verifies the results of the regression analysis earlier in this report, at least from the perspective of direction of change. Although the extremes were quite dispersed, the standard deviations indicated the multiplier populations were tightly

distributed about their means. However, the explanatory variables were quite diverse. As areas grew quite large or populated, multipliers grew slowly. This suggests that although the appropriate functional form was determined to be linear, the upper range of the regression line flattens to a non-linear configuration. The implications of this will be discussed further in the next chapter. Appendix H shows the descriptive statistics of the analysis described.

CHAPTER FIVE - CONCLUSIONS

Introduction

An increased demand for economic impact analysis has increased the need to simplify the economic impact analysis process. Much of this demand has resulted from recently developed laws and regulations requiring government agencies to determine the economic impacts of their proposed projects, programs, and policies.

The general goal of this study was to investigate the behavior of certain economic multipliers of the wood products industries and the feasibility of providing more readily available, situation-specific, economic multipliers to assist in the economic impact analysis process.

The three research objectives are to describe the range of variation in economic multipliers for certain sectors of the wood products industries, to identify the key factors that explain the variation of economic multipliers, and to determine the feasibility of using selected explanatory variables to simplify the determination of economic multipliers used for economic analysis. The study was not intended to develop a system to be used immediately by analysts but only determine whether or not such tools might be feasible.

Research Summary

The first step in the research study was a literature review. The literature provided the theoretical foundation for the eventual analysis. The literature also provided many ideas as to what variables could be used to achieve the desired results as well as what analysis techniques might be available and the appropriate functional form to use.

The first operational step in the research process was to determine the scope of the study. It was decided that it would not be feasible to analyze all the sectors of the wood products industry – there were just too many and conclusions could probably be reached with just analyzing the major, most frequently occurring sectors. A variety of explanatory variables were then selected. This selection was based on discussions with subject area experts and a review of the literature. These were variables that were thought to be closely related to the complexity of the wood products industry and the economy as a whole as well as being variables that could be easily measured from readily available data. Additional variables were tested as the study progressed. Other decisions that were made included the kind of multipliers to be tested (Type I, Type SAM, etc.), the economic parameter that the multiplier should represent (income, employment, value added etc.), and how the impact area boundaries were determined and sampled. These processes are described in detail in previous sections of this dissertation.

Twelve separate models were developed representing three wood products sectors, two types of multipliers, and two economic variables. Regional differences were accounted for via dummy variable values within each of the twelve models.

The first step in the statistical analysis was to determine and review the basic descriptive statistics of the multipliers and explanatory variables. Means, medians, standard deviations, distributions, ranges, and measures of normality of economic multipliers were analyzed for all selected wood products sectors. These analyses set the stage and provided guidance for further analysis.

The next step was to do correlation analysis to determine the strength of the relationships between the explanatory variables and the multipliers to validate the selection of explanatory variables. Also correlation analysis was done among the explanatory variables to determine the potential for multicollinearity.

Best-subsets linear regression analysis was then performed to determine the linear relationships between and among all combinations of explanatory and response (multipliers) variables. The end product of this process was a “preliminary” regression equation or model. The models were examined for the presence of heteroscedasticity.

The next steps were done to try to improve the model. An analysis of variance was done to determine whether or not there were regional differences – and there were.

Non-linear assumptions were tested by using logarithm scales on both and each axis. The original correlation analysis and the best-subsets regression analysis indicated a high potential for multicollinearity between explanatory variables. Sets of explanatory variables were reviewed and excluded of unnecessary variables to reduce variance inflation factors (VIF) to acceptable levels.

Several additional potential response variables were tested for their predictive ability.

The most significant of these was the magnitude of the timber supply within the economic impact area represented by the total acres of timberland in the economic impact area and timberland acres as a percent of total acres within the impact area.

An attempt was also made to develop interaction variables (e.g. cross-product terms) by multiplying explanatory variables by each other to create a new and unique explanatory variable.

Another approach to improve the regression models was to treat the different impact area levels (Level A – E) as dummy variables. The results were mixed. Some models were reasonably improved with higher adjusted R^2 s and others were not. However, p-values were very large for the Level regression coefficients so the idea was not carried further.

The last analysis procedure was to examine the statistical differences between different geographic impact levels originating from the same event. This was done to answer the question of what happens to the multiplier as the same impact area is systematically expanded and what are the implications.

Discussion of the results

Twelve different regression equation models were developed to predict economic multipliers for three different wood products industry sectors (Table 17). The final

results and notable findings as the models were developed are discussed in this section.

The development of the least squares regression model was based on the set of assumptions in Appendix A.

Study Objective 1: Describe Variations in Multipliers.

The number of observations (impact areas) in the study for each multiplier group, varied from a low of 384 to a high of 674 observations. Most impact areas had Logging Camps and Contractors (Sector 133) and Sawmills and Planing Mills (Sector 134) but many impact areas did not have Veneer and Plywood (Sector 139). The number of observations was more than adequate for the objectives of the study with standard deviations and standard errors being quite small. A wide variation in multiplier values occurred, which is a prerequisite to regression models.

Study Objective 2: Identify the key factors that explain the variation of economic multipliers.

Normality of the response variable (multipliers) was analyzed by visual observation and the use of the Ryan-Joiner test. Results showed good normal distribution allowing for the necessary statistical inferences.

Correlation analysis results (Tables 11 and 12) indicated a low to moderate degree of correlation among most of the variables. The strongest linear relationships were between the multipliers and the number of economic sectors and the geographic size of the area. Correlation coefficients were highest for the logging sector and lowest for planing mills. Correlation coefficients for personal income multipliers were generally larger than for employment multipliers and larger for Type SAM multipliers than Type 1 multipliers. Correlation coefficients were very high between several of the explanatory variables suggesting potential for multicollinearity. This situation was especially pronounced between population and number of economic sectors, total industry output, and total personal income. However, this was to be expected and the problem was addressed later in the analysis.

Study Objective 3: Determine the feasibility of using selected, readily available, explanatory variables to estimate economic multipliers.

The purpose of regression analysis, was to develop a set of regression equations to serve as predictive models. The first phase of Best-Subsets analysis generated a set of models for each of the twelve multiplier groups based on all the possible combinations of the explanatory variables and their relationship to the multipliers. Non-linear relationships were considered by examining residual plots and logarithmic transformation. The results can be seen in Appendix B. The single strongest variables, those with the highest adjusted R^2 value, are the same as those with the highest correlation coefficients; the number of economic sectors and the geographic size of the area. The single weakest

explanatory variables do not follow a pattern and vary depending on the type of multiplier and industry sector. However, as explanatory variables are aggregated the adjusted R^2 values increase with corresponding decreases in the Cp values. The best of each set is then selected for preliminary regression model development (Table 14). The best model could have as few as three or as many as six explanatory variables. The most common explanatory variable was the number of sectors, which occurs in every model. The least common variable is total industry output, which occurs in only half the models. Models for the SAM multipliers are generally stronger (have higher adjusted R^2) than for the Type I multipliers. This is because leakage from indirect expenditures is less in regions with more people and economic sectors. Models for the Logging Camps and Contractors (sector 133) are the strongest. Models for Sawmills and Planing Mills (sector 134) are the weakest. There is no consistency between the personal income and employment multipliers across the economic sectors analyzed.

Regional differences (Forest Service Regions were used) were analyzed using one-way analysis of variance. The outcome can be seen in Appendix D. Regional groupings were done based on the graphic illustrations in Appendix D and new models were developed.

Substantial improvement was made with all models with some of the R^2 values more than doubling (Table 15). As previously mentioned non-linearity was addressed with no potential improvement apparent.

Multicollinearity, which was suggested as a potential problem in predictive models, during the discussion of correlation analysis, became a reality as a result of regression

analysis. All but two of the multiplier variables needed at least one of the highly collinear explanatory variables removed from the model to reduce VIF values below the acceptable standard of 5.0 (Table 15). Only two of the multiplier variables required two explanatory variables to be removed. The most frequently removed explanatory variable was personal income, which was removed from six of the models. Removal of these variables resulted in very little decrease of R^2 values. The number of sectors was still the most common explanatory variable in the final models, personal income, the least frequent.

The models were examined for the presence of heteroscedasticity. This was done with a visual examination of graphs of the regression residuals. The presence of heteroscedasticity was not to the extent that it would affect the results of the study.

As previously mentioned, interaction variables or cross-product terms were tested. This only resulted in slight increases in adjusted R^2 values but substantially increased multicollinearity (VIF). Therefore, the interaction variables were not carried any further in the analysis.

After the final models in Table 17 were developed, another potentially viable explanatory variable was analyzed. Timber supply in the impact area was thought to possibly have an effect on the magnitude and complexity of the timber industry. This might influence the size of the economic multipliers. This element was represented by two variables – total timberland area and timberland area as a percent of total area in the impact region. The

details of this analysis are included in a previous section of this report. The results were inconclusive because of several factors including the absence of reliable data. In some cases the adjusted R^2 values slightly improved the models, but introduced a substantial amount of additional variation (higher VIF values). In other cases there was little or no improvement. No conclusions can be made on the usefulness of this variable until a full set of reliable data is available.

Among the objectives of the study was the intent to analyze the behavior of multipliers as specific impact areas are systematically increased. The descriptive statistics developed by this analysis are in Appendix H. With the exception of population density, all the explanatory variables increased as the impact area increased in size. As expected, as the geographic area was enlarged the multiplier usually also got larger, but in the extremes of its range, not in a linear fashion. For example the geographic area (mean) between level B and C approximately tripled while the increase in the multiplier is very small. The main point here is that, at least on the average, additional area past a certain point, tends to dilute the effects of the impact analysis. Therefore, indiscriminately using the mechanical rules to build impact areas (e.g. all counties with National Forest land) or aggressively expanding the impact area in the fear of missing something, is not recommended.

The question remains – how feasible is it to construct ready-made, situation specific economic multipliers for the wood products industry. The study shows that predictive models can be developed from readily available data (Table 17). Reliability parameters

such as standard errors, p-values and variance inflation factors indicate that the R^2 values and adjusted R^2 values are highly reliable. However, R^2 values and adjusted R^2 values are quite low indicating a large portion of the variation is unaccounted for (Table 16). Only three multiplier groups have R^2 and adjusted R^2 values exceeding 50%. Three groups are less than 30% with one of these groups at 15%.

The Logging Camps and Contractors sector has significantly higher R^2 and adjusted R^2 values than the other sectors (> 50%). This could be due to the logging industry being more consistently found in small rural economies than are the other two sectors, which can be found throughout a greater range of regions with more diverse characteristics.

SAM multipliers have higher R^2 and adjusted R^2 values than Type 1 multipliers. This is probably because the Type SAM multiplier includes the effects of induced spending (i.e. spending by households) which increases with increases in variables directly related to the size of the economy (i.e. personal income, number of sectors, population). There is no clear pattern to the relationship between employment multipliers and income multipliers.

In addition to the measures taken above to determine the validity of the regression models, the results were also examined in respect to mean and maximum percent error. This was done by taking the difference between the IMPLAN generated multipliers and comparing them to the multipliers generated by the regression models (Table 18). While mean absolute percent errors were small the maximum values were quite large meaning

that for any particular multiplier estimate the multiplier estimated by the regression models could be substantially different than the actual multiplier estimated by IMPLAN. This analysis did not serve to invalidate the regression results.

The information above indicates that only those willing to take high risks might chose to use the models for Logging Camps and Contractors to estimate the multipliers for just that economic sector. However, there is usually very little demand for just a single sector multiplier within the wood products industries. R^2 and adjusted R^2 values for the other sectors are so low that they should not be considered a viable alternative. This leaves the analyst with the undesirable option of using a regression model derived multiplier for one sector (Logging Camps and Contractors) and multipliers from another source (i.e. input-output model) for the remaining sectors.

This study suggests the feasibility of generating economic multipliers from regression models, using readily available explanatory variables, is limited. The risk of error is high and must be weighed against more expensive methods with lower risk. However, there are still potential ways to improve the models (see Recommendations for Future Research).

Policy Implications

Based on this research, the development of predictive regression models to estimate economic multipliers appears to be questionable. Also, is it practical in light of other

alternatives? The primary alternative at the present time is developing site-specific multipliers using model-building software programs such as IMPLAN and RIMSII. Both approaches require substantial initial investments. To develop a set of predictive models requires skills in econometric modeling. Software systems like IMPLAN are already developed. However, systems like IMPLAN have high acquisition costs and also require skilled analysts to develop dependable results.

The use of regression models requires the user to have access to the data representing the model variables. The IMPLAN user already has most of the information they need in the IMPLAN data set. However, it would be possible to build a regression based software model where all the basic information is in the model data set just as it is in IMPLAN. Policy decisions would depend heavily on the total relative costs of the alternative approaches and the quality of the results. This is beyond the scope of this study. The behavior of multipliers resulting from changes in explanatory variables should be considered in writing guidelines for economic impact area design.

Recommendations for Future Research

Economic feasibility and assumption of risk seems to be the primary issue as to whether regression model multipliers should be used in lieu of models such as IMPLAN. To make this decision research needs to be done looking at the total costs in relation to the product needed. For example IMPLAN is expensive but can do a great many things.

However, many of these products of IMPLAN are rarely used. A regression-based model could be designed to produce only those reports that are commonly used and require less skill/expense. If the results from developing regression-based models do not justify the expense then further development is likely doomed unless it makes the model more economically feasible.

A second recommendation is to further explore timber supply as an explanatory variable. Within the next year the USDA Forest Service should have complete coverage of the U.S. for timberlands. The present inventory that is available for use does not include some substantial timbered areas such as Oregon and Washington. Other vegetation inventories do not separate out timberlands. Actual timber harvest levels could also be used as an explanatory variable if this data could be found in a readily usable form.

A third recommendation is to analyze more of the multipliers such as those for value-added. It was not included in this study because of an attempt to keep the number of variables within reason and value-added is not a commonly used measure of economic activity for economic impact analysis.

A fourth recommendation is to evaluate employment as an explanatory variable. It was not included in this study because of the concern about the different definitions of employment and what jobs are included in the readily available data. However, there might be some merit to looking at the definition in terms of total employment as determined by the Bureau of Economic Analysis.

A fifth possibility for future research is to generate sales or output multipliers and use these multipliers as the dependent variables in regression models. Then analyze for correlation and regression coefficients. Employment and income multipliers could then be created from local or regional output to jobs or income ratios.

A final recommendation is to further investigate the non-linear characteristics of multiplier predictions. Part of this study looked at the behavior of multipliers as specific impact areas were expanded. There was evidence to demonstrate that as the size of specific impact areas reached certain levels, multipliers grew slowly. This has important implications where analysts are trying to geographically focus on impacted areas.

REFERENCES

Aldwell, P.H.B. 1984. Variability of regional income and employment multipliers in forest-based industries in New Zealand. (New Zealand Forest Service). First published in *New Zealand Economic Papers* 16 (1982):113-131.

Alward, Gregory S. and Scott Lindall. 1996. Deriving SAM multiplier models using IMPLAN. Presented at the 1996 National IMPLAN Users Conference, Minneapolis, MN, August 15-17, 1996.

Archer, Brian H. 1983. Economic impact: misleading multiplier. *1984 Annals of Tourism Research*:517-18.

Aruna, P.B., Frederick Cubbage, Karen J. Lee, and Clair Redmond. 1997. Regional economic contributions of the forest-based industries in the south. *Forest Products Journal* 47(7/8):35-45.

Bauen, Rebecca, Bryan Baker, and Kirk Johnson. 1996. *Sustainable Community Checklist*. Seattle: Northwest Policy Center, University of Washington.

Beckley, Paul. Telephone conversations with local retail lumber dealers and brokers. Kalispell, MT, 18 March 1998.

Bills, Nelson, and Linda Zygadlo. 1978. Regional development. In *Regional Development And Plan Evaluation – The Use Of Input-Output Analysis*. Economics, Statistics, and Cooperatives Service, Agriculture Handbook No.530. Washington, D.C.

Bryan, Hobson. 1996. The assessment of social impacts. In *Natural Resource Management – The Human Dimension*. Boulder, CO: Westview Press.

Burdge, Rabel J. 1994. *A Conceptual Approach To Social Impact Assessment*. Middleton, WI: Social Ecology Press.

---. 1999. *A Community Guide to Social Impact Assessment*. Middleton, WI: Social Ecology Press.

Burfurd, Roger L., and Joseph L. Katz. 1981. A method for estimation of input-output-type output multipliers when no I-O model exists. *Journal of Regional Science* 21 (2):151.

Chang, Wen-Huei. 2001. *Variations In Multipliers And Related Economic Ratios For Recreation And Tourism Impact Analysis*. Draft dissertation. East Lansing: Michigan State University.

- Drake, Ronald L. 1976. A short-cut approach to estimates of regional input-output multipliers: methodology and evaluation. *International Regional Science Review*, 1:1-17.
- Draper, Norman, and Harry Smith. 1981. *Applied Regression Analysis*, 2d ed., John Wiley and Sons, New York, pp.266-274.
- Ekins, Paul. 1992. *The GAIA Atlas Of Green Economics*. London:Anchor Books – Doubleday.
- Field, Donald R., and William R. Burch, Jr. 1991. *Rural Sociology And The Environment*. Middleton, WI.: Social Ecology Press.
- Fjeldsted, Boyd L. 1990. *Regional Input-Output Multipliers: Calculation, Meaning, Use And Misuse*. Bureau of Business and Economic Research, Graduate School of Business, University of Utah, 50 (10).
- Flathead County, Montana, lumber retailers. 1998. Personal communications with the author.
- Flick, Warren A., and Lawrence D. Teeter. 1988. Multiplier effects of the southern forest industries. *Forest Products Journal* 38 (11/12):69-74.
- Fossum, Harold L. 1993. *Communities In The Lead – The Northwest Rural Development Sourcebook*. Seattle: Northwest Policy Center, University of Washington.
- Freedman, David, Robert Pisani, Roger Purves, and Ani Adhikari. 1991. *Statistics*, 2nd ed. NY: W.W. Norton and Company, Inc.
- Frey, Donald E. 1989. A structural approach to the economic base multiplier. *Land Economics* 65(4):352-8.
- Goldman, George, Anthony Nakazawa, and David Taylor. 1997. Determining economic impacts for a community. *Economic Development Review* 15(1):48-51.
- Gujarati, Damodar N. 1995. *Basic Econometrics*, 3rd ed. NY: McGraw- Hill, Inc.
- Hart, Maureen. 1999. *Guide To Sustainable Community Indicators*. North Andover, MA: Hart Environmental Data.
- Hastings, Steven E., and Sharon M. Brucker. 1993. An introduction to regional input-output analysis. Chapter 1 in *Microcomputer-Based Input-Output Modeling: Applications To Economic Development*. Boulder: Westview Press.
- Hewings, Geoffrey J.D. 1985. *Regional Input-Output Analysis*, Vol. 6. Beverly Hills, CA: Sage Publications.

Holland, David W., Bruce A. Weber, and Edward C. Waters. 1996. Modeling the economic linkage between core and periphery regions: the Portland, Oregon trade area. In *Rural –Urban Interdependence And Natural Resource Policy*. Corvallis: Oregon State University.

Holland, David W., Hans T. Geir, and Ervin G. Schuster. 1997. *Using IMPLAN To Identify Rural Development Opportunities*. U. S. Department of Agriculture, Forest Service, General Technical Report INT-GTR-350. Washington, D.C.

Hoover, Edgar M., and Frank Giarratani. 1984. *An Introduction To Regional Economics*, 3rd ed. NY:Alfred A. Knopf, Inc.

IMPLAN Pro Version 2.0. 1999. Users guide. Minnesota IMPLAN Group, Stillwater, MN. First edition.

Isard, Walter. 1975. *Introduction To Regional Science*. Englewood Cliffs, NJ: Prentice-Hall Inc.

Kolison, Stephen H. Jr., Rodney L. Busby, and James E. Granskog. 1992. Alabama's lumber and wood products exports: status, importance, and competitive edge. In *Proceedings Of The 1992 Southern Forest Economics Workshop On The Economics Of Southern Forest Productivity: Competing In World Markets*. Columbus, GA: Mead Southern Wood Products.

Koutsoyiannis, A. 1979. *Theory Of Econometrics*, 2nd ed. Hong Kong: The McMillan Press Ltd.

Krikelas, Andrew C. 1992. Why regions grow: A review of research on the economic base model. *Economic Review* 77 (4):16-29.

Lord, Bruce E., and Charles H. Strauss. 1993. A review and validation of the IMPLAN model for Pennsylvania's solid hardwood product industries. In *9th Central Hardwood Forest Conference Proceedings*. West Lafayette: Purdue University.

Maki, Wilbur. 1997. Accounting for local economic change in regional input-output modeling. *The Journal Of Regional Analysis And Policy* 27 (2): 95-109.

McKusick, Robert. 1978. Regional development and plan evaluation: the use of input-output analysis. Chapter 1 in *Regional Development and Plan Evaluation*. U.S. Department of Agriculture. Economics, Statistics, and Cooperatives Service. Agriculture Handbook No.530. Washington, D.C.

Mendenhall, William, and Robert J. Beaver. 1994. *Introduction To Probability And Statistics*, 9th ed. Belmont, CA: Duxbury Press.

- Miernyk, William H. 1965. *The Elements Of Input-Output Analysis*. NY: Random House.
- Miller, Ronald E., and Peter D. Blair. 1985. *Input-Output Analysis, Foundations And Extensions*. Englewood Cliffs, NJ: Prentice-Hall.
- MINITAB users guide, Release 13 for Windows, February 2000.
- Minnesota IMPLAN Group. 1997. Type II Multipliers and IMPLAN. In *IMPLAN News* (December). Stillwater, MN
- Montgomery, D.C., and E.A. Peck. 1982. *Introduction To Linear Regression Analysis*. John Wiley and Sons.
- Mulligan, Gordon F., and Lay James Gibson. 1984. Regression estimates of economic base multipliers for small communities. *Economic Geography* 60: 225-237.
- Olfert, M. R., and Jack C. Stabler. 1994. Community level multipliers for rural development initiatives. *Growth And Change* 25:467-86.
- Olson, Doug. 1995. When the parts are greater than the whole. *IMPLAN News* 14:2. Stillwater, MN.
- Olson, Doug. 1997. Type II multipliers and IMPLAN. *IMPLAN News* 20:1. Stillwater, MN.
- . 1995a. When the parts are greater than the whole – aggregation error. *IMPLAN News* 15:6. Stillwater, MN.
- . 1997. Zip code files. *IMPLAN News* 18:2. Stillwater, MN.
- Olson, Doug, and Gregory Alward. 2000. Updating IMPLAN RPCs. In *Proceedings Of The 2000 National IMPLAN User's Conference*. Fort Collins: Colorado State University.
- Otto, Daniel M., and Thomas G. Johnson. 1993. *Microcomputer-Based Input-Output Modeling: Applications To Economic Development*. Boulder, CO: Westview Press, Inc.
- Palmer, C. and L.E. Siverts. 1985. IMPLAN analysis guide. U.S. Department of Agriculture, Forest Service, Systems Application Unit, Land Management Planning, Fort Collins, CO.
- Phelps, Edmund S. 1989. Distributive justice. In *The New Palgrave – Social Economics*. NY:W.W. Norton and Company, Inc.

Pleeter, Saul. 1980. Methodologies of economic impact analysis: an overview. Chapter 1 in *Economic Impact Analysis: Methodology And Applications*. Boston: Martinus Nijhoff Publishing.

Power, Thomas Michael. 1988. *The Economic Pursuit Of Quality*. Armonk, NY: M.E. Sharp, Inc.

---. 1996. *Lost Landscapes And Failed Economies*. Washington, D.C.: Island Press.

Quigley, Thomas M., Richard W. Haynes, and Russell T. Graham. 1996. *Integrated Scientific Assessment For Ecosystem Management In The Interior Columbia Basin And Portions Of The Klamath And Great Basins*. USDA Forest Service and USDI Bureau of Land Management. General Technical Report PNW-GTR-382:35.

Rasker, Ray, Jerry Johnson, and Vicky York. 1994. *Measuring Change In Rural Communities*. The Wilderness Society, Bolle Center for Ecosystem Management. Missoula, MT.

Retzlaff, Mike, Larry Leefers, and Paul Monson. 2000. *Defining Economic Impact Analysis Areas In The Eastern Region*. An unpublished report. U. S. Department of Agriculture, Forest Service. Denver, CO.

Richardson, Harry W. 1979. *Regional Economics*. Urbana: University of Illinois Press.

---. 1985. Input-output and economic base multipliers: Looking backward and forward. *Journal of Regional Science* 25(4): 607-661.

Rickman, Dan S., and R. Keith Schwer. 1995. A comparison of the multipliers of IMPLAN, REMA, and RIMS II: Benchmarking ready-made models for comparison. *The Annals of Regional Science* 29:363-374.

Robison, Hank. 1997. *Defining Regional Boundaries For Forest Service Economic Impact Analysis*. An unpublished report. Moscow, ID: Economic Modeling Specialists, Inc.

Rose, Adam, and William Miernyk. 1989. Input-output analysis: the first fifty years. *Economic Systems Research* 1(2).

Snedecor, George W., and William G. Cochran. 1980. *Statistical Methods*, 7th ed. Ames: The Iowa State University Press.

Stynes, D.J. 1999. *Understanding Multipliers And How To Interpret Them*. (online). Available at <http://www.msu.edu/course/prr/840/econimpact> (July 12, 2000).

Teeter, Lawrence, Gregory S. Alward, and Warren A. Flick. 1989. Interregional impacts of forest-based economic activity. *Forest Science* 35(2):515-531.

U.S. Department of Agriculture. Economics, Statistics, and Cooperatives Service. 1978. *Regional Development And Plan Evaluation – The Use Of Input-Output Analysis*. Agriculture Handbook No.530. Washington, D.C.

U.S. Department of Agriculture. Forest Service. 1988. *Defining The Relevant Impact Area*. Forest Service Handbook 1909.17. Washington, D.C.

---. Forest Service. 1992. *Supplemental Readings On Input-Output And Micro-IMPLAN*. An unpublished report compiled by Ecosystem Management. Fort Collins, CO.

---. Forest Service. 2000. *36 CFR Parts 217 And 219 National Forest System Land And Resource Management Planning: Final Rule*. Federal Register, November 9, 2000: 67514-67581.

U.S. Department of Agriculture, Forest Service, Inventory and Monitoring Institute. 2001. *Database of Economic Diversity Indices for US Areas*. Fort Collins, CO. Available at www.fs.fed.us/institute/economic_center/spatialdata3.html.

U.S. Department of Agriculture, Forest Service, and U.S. Department of Interior, Bureau of Land Management. 1997. *Upper Columbia River Basin Draft Environmental Impact Statement*, Vol. 1. Boise, ID.

---. 2000. *Interior Columbia Basin Final Environmental Impact Statement*. Boise, ID.

U.S. Department of Commerce. Bureau of Economic Analysis. 1986. *Regional Multipliers: A User Handbook For The Regional Input-Output Modeling System (RIMS II)*. Washington, D.C.

---. Bureau of Economic Analysis. 1995. *BEA Economic Area Component County List*. Available at www.bea.doc.gov/bea/regional/docs/econlist.htm.

---. Bureau of Economic Analysis, Inter-industry Division. 2001. *Annual Input-Output Accounts Of The U.S. Economy, 1997*. Washington, D.C.

Vasievich, Mike. 1999. Here comes the neighborhood – a new gold rush and eleven other trends affecting the Midwest. *NC News* (Aug/Sept):1-3. St. Paul, MN:USDA Forest Service, North Central Research Station.

Wagner, John E., Steven C. Deller, and Greg Alward. 1992. Estimating economic impacts using industry and household expenditures. *Journal Of The Community Development Society* 23(2).

Zheng, Chinlong, and Patrice A. Harou. 1988. A method to estimate input-output multipliers for the forestry sector without an I-O table. *Forest Science* 34(4):882-893.

Appendix A – Validation of Linear Models

Assumptions Underlying the Method of Least Squares - There are 10 basic assumptions that must be met to have a valid least squares regression model (Koutsoyiannis 1979). Each is discussed relative to this research.

Assumption 1: Linear regression model. The regression model is linear in the parameters. The parameters, which are economic multipliers and associated explanatory variables are assumed to be linearly related and are analyzed accordingly. In addition, the X and Y variables were put in natural log form and tested for non-linearity.

Assumption 2: X values are fixed in repeated sampling. Values taken by regressor X are considered fixed in repeated samples. More technically, X is assumed to be non-stochastic. The X values in this study are fixed known constants such as population, number of economic sectors etc. The Y values are economic multipliers and are stochastic.

Assumption 3: Zero mean value of disturbance U_i . Given the value of X , the mean, or expected, value of the random disturbance term U_i is zero. Technically the conditional mean value of U_i is zero. The analysis was designed to produce a least squares linear equation therefore equally distributing the observations about the regression line. This, by definition, produces a net zero value for the residuals.

Appendix A (Cont.)

Assumption 4: Homoscedasticity or equal variance of U_i . Given the value of X , the variance of U_i is the same for all observations. That is, the conditional variances of U_i are identical. This was confirmed by examining scatter plots with residual values plotted along the X values. A slight amount of heteroscedasticity was observed in some of the models but it was not significant enough to justify using a procedure such as the Lagrange-Multiplier test. The implications of non-constant variance which leads to non-constant standard errors affects our ability to make precision statements in hypothesis testing.

Assumption 5: No autocorrelation between the disturbances. Given any two X values, X_{ij} ($i \neq j$), the correlation between any two U_i ($i \neq j$) is zero. The assumption of ordinary least squares is that the successive values of the random variable u are temporally independent, that is, that the value which U assumes in any one period is independent from the value which it assumed in any previous period. In brief, this is a problem associated with time series data. This study does not include time series – it is cross-sectional therefore autocorrelation is not a problem.

Assumption 6: Zero covariance between U_i and X_i . The disturbance u and explanatory variable X are uncorrelated. This is similar to homoscedasticity in Assumption 4. As X changes, the magnitude of the error terms should not change. Examination of the arrangement of the error terms around their mean (zero) shows this not to be a problem.

Appendix A (Cont.)

Assumption 7: The number of observations n must be greater than the number of parameters to be estimated. Alternatively, the number of observations n must be greater than the number of explanatory variables. In this study there are at from 3 to 6 explanatory variables plus the Y intercept. There are a minimum of 384 observations thus satisfying this assumption.

Assumption 8: Variability in X values. The X values in a given sample must not all be the same. Technically, variable (X) must be a finite positive number. For example, if there was little variation in population or any of the other explanatory variables, we would not be able to explain much of the variation in the multipliers. There is significant variation in all the explanatory variables in this study.

Assumption 9: The regression model is correctly specified. Alternatively, there is no specification bias or error in the model used in empirical analysis. Specification should be addressed in terms of the variables chosen and the form of the variables chosen. The variables chosen were the result of experience in the economic impact field and and the available literature. It was felt there were no other readily available variables that would strengthen the model. The assumption of linearity was tested by looking at the graphic distribution of the residuals and natural log transformations.

Appendix A (Cont.)

Assumption 10: There is no perfect multicollinearity. That is, there are no perfect linear relationships among the explanatory variables. This was tested with correlation analysis among the explanatory variables. Although some R-values were very high in the draft models, none were close to perfect. All of the highly collinear values were removed in the final model. Remaining variables in the final model showed very low collinearity. The effects of multicollinearity include enlarging the standard error and variance, which is measured by the variance inflation factor (VIF). Variables with high VIF values were usually removed in the regression analysis process.

Appendix B – Best Subsets Regression Analysis Results

B-1. Regression: Best Subsets

LE1 versus Pop, PopDen, Sect, TIO, PerInc, EIA

Response (dependent variable) is LE1

Vars	R-Sq	R-Sq (adj)	Cp	S	Pop	Pop Den	Sect	TIO	Per Inc	EIA
1	26.9	26.8	109.4	0.083181			X			
1	21.5	21.4	165.5	0.086214						X
1	3.7	3.6	349.0	0.095477				X		
1	3.7	3.6	349.2	0.095488					X	
1	3.7	3.6	349.3	0.095491	X					
2	32.8	32.6	50.7	0.079821			X			X
2	29.9	29.7	80.3	0.081510	X		X			
2	29.7	29.5	82.9	0.081654			X	X		
2	29.6	29.4	84.0	0.081716			X		X	
2	29.1	28.8	89.5	0.082026		X	X			
3	36.6	36.3	13.7	0.077606	X		X			X
3	36.1	35.8	18.9	0.077912			X	X		X
3	36	35.7	19.6	0.077954			X		X	X
3	33.2	32.9	48.5	0.079641		X	X			X
3	31	30.7	71.5	0.080960	X	X	X			
4	37.6	37.3	4.9	0.077022	X		X		X	X
4	37.1	36.7	10.7	0.077372	X		X	X		X
4	36.6	36.2	15.7	0.077666	X	X	X			X
4	36.1	35.7	20.5	0.077953			X	X	X	X
4	36.1	35.7	20.9	0.077973		X	X	X		X
5	37.8	37.3	5.0	0.076972	X		X	X	X	X
5	37.7	37.2	6.7	0.077072	X	X	X		X	X
5	37.1	36.6	12.5	0.077421	X	X	X	X		X
5	36.1	35.6	22.5	0.078013		X	X	X	X	X
5	31.8	31.3	66.7	0.080584	X	X	X	X	X	
6	37.8	37.2	7.0	0.077030	X	X	X	X	X	X

649 cases used, 86 cases contain missing values.^b

^a Shaded row was selected as best subset for regression model development.

^b Cases with missing values are those impact areas that did not have the sector being analyzed.

Appendix B (Cont.)

B-2. Regression: Best Subsets

LE2 versus Pop, PopDen, Sect, TIO, PerInc, EIA

Response (dependent variable) is LE2

Vars	R-Sq	R-Sq (adj)	Cp	S	Pop	Pop Den	Sect	TIO	Per Inc	EIA
1	26	25.9	104.6	0.14625		X				
1	17.9	17.8	186.5	0.15402						X
1	3	2.8	337.7	0.16745					X	
1	2.9	2.8	338.0	0.16748				X		
1	2.9	2.8	338.1	0.16748	X					
2	30	29.8	65.5	0.14229			X			X
2	29.8	29.6	67.8	0.14252		X	X			
2	29.7	29.5	68.9	0.14263	X		X			
2	29.5	29.3	71.4	0.14288			X	X		
2	29.3	29.1	73.4	0.14308			X		X	
3	34.5	34.2	22.7	0.13782	X		X			X
3	34	33.7	27.5	0.13833			X	X		X
3	33.8	33.5	29.2	0.13850			X		X	X
3	31.9	31.6	48.4	0.14046	X	X	X			
3	31.8	31.5	49.6	0.14059		X	X			X
4	35.6	35.2	13.3	0.13675	X		X		X	X
4	34.9	34.5	19.9	0.13743	X	X	X			X
4	34.8	34.4	21.3	0.13758	X		X	X		X
4	34.5	34.1	24.8	0.13794		X	X	X		X
4	34.4	33.9	25.9	0.13806		X	X		X	X
5	36.3	35.8	8.7	0.13615	X	X	X		X	X
5	36.1	35.6	10.0	0.13629	X		X	X	X	X
5	35.4	34.9	16.8	0.13701	X	X	X	X		X
5	34.6	34.1	25.4	0.13790		X	X	X	X	X
5	33.3	32.8	38.8	0.13929	X	X	X	X	X	
6	36.6	36.0	7.0	0.13587	X	X	X	X	X	X

649 cases used, 86 cases contain missing values.

Appendix B (Cont.)

B-3. Regression: Best Subsets

LP1 versus Pop, PopDen, Sect, TIO, PerInc, EIA

Response (dependent variable) is LP1

Vars	R-Sq	R-Sq (adj)	Cp	S	Pop	Pop Den	Sect	TIO	Per Inc	EIA
1	26.1	26.0	24.4	0.10815			X			
1	13.0	12.8	143.6	0.11739						X
1	10.4	10.3	166.7	0.11910	X					
1	10.3	10.2	167.5	0.11915				X		
1	10.1	9.9	170.1	0.11935					X	
2	27.8	27.6	11.4	0.10702			X			X
2	26.4	26.2	23.9	0.10804		X	X			
2	26.1	25.9	26.3	0.10823			X	X		
2	26.1	25.9	26.4	0.10823			X		X	
2	26.1	25.9	26.4	0.10823	X		X			
3	28.9	28.5	3.5	0.10629		X	X			X
3	27.8	27.5	13.3	0.10709	X		X			X
3	27.8	27.5	13.3	0.10710			X		X	X
3	27.8	27.5	13.4	0.10710			X	X		X
3	26.4	26.1	25.7	0.10811		X	X		X	
4	29.1	28.7	3.6	0.10621	X	X	X			X
4	29.1	28.6	3.7	0.10622		X	X		X	X
4	29.1	28.6	3.9	0.10624		X	X	X		X
4	28.0	27.5	13.7	0.10704			X	X	X	X
4	28.0	27.5	13.9	0.10706	X		X	X		X
5	29.1	28.6	5.2	0.10626		X	X	X	X	X
5	29.1	28.6	5.2	0.10627	X	X	X	X		X
5	29.1	28.5	5.5	0.10629	X	X	X		X	X
5	28.0	27.5	15.3	0.10710	X		X	X	X	X
5	26.5	25.9	29.1	0.10823	X	X	X	X	X	
6	29.2	28.5	7.0	0.10633	X	X	X	X	X	X

649 cases used, 86 cases contain missing values.

Appendix B (Cont.)

B-4. Regression: Best Subsets

LP2 versus Pop, PopDen, Sect, TIO, PerInc, EIA

Response (dependent variable) is LP2

Vars	R-Sq	R-Sq (adj)	Cp	S	Pop	Pop Den	Sect	TIO	Per Inc	EIA
1	41.4	41.3	23.5	0.14004			X			
1	17.8	17.6	293.6	0.16593						X
1	16.4	16.2	309.7	0.16735	X					
1	16.1	16.0	312.3	0.16758				X		
1	15.8	15.6	316.6	0.16796					X	
2	43.0	42.8	7.6	0.13825			X			X
2	41.5	41.3	24.8	0.14007		X	X			
2	41.4	41.3	25.5	0.14015	X		X			
2	41.4	41.3	25.5	0.14015			X		X	
2	41.4	41.3	25.5	0.14015			X	X		
3	43.6	43.3	3.3	0.13769		X	X			X
3	43.0	42.8	9.4	0.13834	X		X			X
3	43.0	42.7	9.5	0.13835			X		X	X
3	43.0	42.7	9.5	0.13836			X	X		X
3	41.5	41.2	26.7	0.14016			X	X	X	
4	43.7	43.3	3.7	0.13762	X	X	X			X
4	43.7	43.3	3.8	0.13764		X	X		X	X
4	43.7	43.3	4.0	0.13766		X	X	X		X
4	43.1	42.8	10.0	0.13830			X	X	X	X
4	43.1	42.8	10.1	0.13831	X		X	X		X
5	43.7	43.3	5.2	0.13768	X	X	X	X		X
5	43.7	43.3	5.3	0.13769		X	X	X	X	X
5	43.7	43.3	5.5	0.13771	X	X	X		X	X
5	43.2	42.7	11.6	0.13836	X		X	X	X	X
5	41.6	41.1	30.1	0.14033	X	X	X	X	X	
6	43.8	43.2	7.0	0.13777	X	X	X	X	X	X

649 cases used, 86 cases contain missing values.

Appendix B (Cont.)

B-5. Regression: Best Subsets										
SE1 versus Pop, PopDen, Sect, TIO, PerInc, EIA										
Response (dependent variable) is SE1										
Vars	R-Sq	R-Sq (adj)	Cp	S	Pop	Pop Den	Sect	TIO	Per Inc	EIA
1	10.3	10.2	148.6	0.13527		X				
1	7.4	7.3	174.2	0.13745						X
1	4.6	4.4	199.2	0.13953	X					
1	4.6	4.4	199.4	0.13956					X	
1	4.5	4.3	200.1	0.13961				X		
2	18.6	18.4	77.5	0.12897	X					X
2	18.1	17.8	82.4	0.12941					X	X
2	17.9	17.7	83.7	0.12952				X		X
2	15.6	15.3	104.0	0.13133		X	X			
2	15.4	15.2	105.5	0.13147		X				X
3	23.6	23.2	35.8	0.12508	X	X	X			
3	23	22.6	41.3	0.12559	X		X			X
3	22.8	22.5	42.2	0.12567		X	X		X	
3	22.7	22.3	43.7	0.12581		X	X	X		
3	22	21.6	49.9	0.12638			X		X	X
4	26.7	26.2	10.6	0.12262	X	X	X			X
4	25.7	25.2	19.0	0.12341		X	X		X	X
4	25.5	25.1	20.6	0.12357		X	X	X		X
4	24.2	23.7	32.6	0.12469	X	X	X	X		
4	24.2	23.7	32.6	0.12469	X	X	X		X	
5	27.5	26.9	5.3	0.12202	X	X	X		X	X
5	27.4	26.9	5.7	0.12206	X	X	X	X		X
5	25.8	25.2	20.3	0.12345		X	X	X	X	X
5	24.2	23.6	34.3	0.12476	X	X	X	X	X	
5	23.4	22.8	41.0	0.12538	X		X	X	X	X
6	27.5	26.8	7.0	0.12209	X	X	X	X	X	X
646 cases used, 89 cases contain missing values.										

Appendix B (Cont.)

B-6. Regression: Best Subsets

SE2 versus Pop, PopDen, Sect, TIO, PerInc, EIA

Response (dependent variable) is SE2

Vars	R-Sq	R-Sq (adj)	Cp	S	Pop	Pop Den	Sect	TIO	Per Inc	EIA
1	9.3	9.1	158.8	0.24343						X
1	6.3	6.2	184.8	0.24734			X			
1	5.2	5.0	195.2	0.24890		X				
1	1.2	1.0	230.2	0.25405	X					
1	1.1	1.0	230.8	0.25413					X	
2	18.2	18.0	81.7	0.23128	X		X			
2	17.6	17.3	87.7	0.23223			X	X		
2	17.4	17.2	88.8	0.23240			X		X	
2	16.6	16.4	95.9	0.23354		X	X			
2	14.8	14.5	112.5	0.23615	X					X
3	23.9	23.5	34.1	0.22336	X		X			X
3	23.8	23.5	34.2	0.22337	X	X	X			
3	23.0	22.6	41.9	0.22465		X	X		X	
3	22.9	22.6	42.3	0.22471		X	X	X		
3	22.8	22.4	43.7	0.22495			X	X		X
4	26.3	25.8	14.9	0.21998	X	X	X			X
4	25.2	24.7	24.5	0.22159		X	X		X	X
4	25.1	24.7	24.9	0.22167		X	X	X		X
4	24.9	24.4	26.8	0.22198	X	X	X		X	
4	24.8	24.3	27.6	0.22212	X		X		X	X
5	27.6	27.0	5.1	0.21814	X	X	X		X	X
5	27.2	26.6	8.8	0.21877	X	X	X	X		X
5	25.2	24.6	26.5	0.22176		X	X	X	X	X
5	24.9	24.3	28.6	0.22212	X		X	X	X	X
5	24.9	24.3	28.8	0.22215	X	X	X	X	X	
6	27.6	26.9	7.0	0.21830	X	X	X	X	X	X

646 cases used, 89 cases contain missing values.

Appendix B (Cont.)

B-7. Regression: Best Subsets

SP1 versus Pop, PopDen, Sect, TIO, PerInc, EIA

Response (dependent variable) is SP1

Vars	R-Sq	R-Sq (adj)	Cp	S	Pop	Pop Den	Sect	TIO	Per Inc	EIA
1	5.0	4.9	39.6	0.18015						X
1	4.1	4.0	45.8	0.18096			X			
1	1.6	1.5	63.7	0.1833		X				
1	0.1	0.0	75.0	0.18477				X		
1	0.0	0.0	75.2	0.18479	X					
2	8.2	7.9	18.7	0.17723		X	X			
2	6.2	5.9	33.1	0.17916			X			X
2	6.0	5.7	34.2	0.17931	X		X			
2	6.0	5.7	34.6	0.17936			X		X	
2	5.9	5.6	35.1	0.17943		X				X
3	8.9	8.5	15.3	0.17664	X	X	X			
3	8.9	8.4	15.9	0.17673		X	X		X	
3	8.8	8.3	16.7	0.17683		X	X			X
3	8.7	8.3	16.9	0.17685		X	X	X		
3	8.4	8.0	19.1	0.17716	X		X			X
4	9.9	9.3	10.5	0.17585		X	X	X	X	
4	9.8	9.3	11.0	0.17592	X	X	X			X
4	9.7	9.1	11.9	0.17605		X	X		X	X
4	9.5	9.0	13.1	0.17621		X	X	X		X
4	9.5	9.0	13.2	0.17622	X	X	X	X		
5	10.7	10.0	6.5	0.17518		X	X	X	X	X
5	10.5	9.8	8.2	0.17541	X	X	X	X		X
5	10.0	9.3	11.5	0.17586	X	X	X	X	X	
5	10.0	9.3	11.9	0.17591	X	X	X		X	X
5	9.1	8.4	18.1	0.17676	X		X	X	X	X
6	10.9	10.1	7.0	0.17511	X	X	X	X	X	X

646 cases used, 89 cases contain missing values.

Appendix B (Cont.)

B-8. Regression: Best Subsets

SP2 versus Pop, PopDen, Sect, TIO, PerInc, EIA

Response (dependent variable) is SP2

Vars	R-Sq	R-Sq (adj)	Cp	S	Pop	Pop Den	Sect	TIO	Per Inc	EIA
1	15.4	15.3	52.1	0.22514			X			
1	10.2	10.0	95.0	0.23199						X
1	2.1	2.0	161.1	0.24217				X		
1	2.1	1.9	161.5	0.24224	X					
1	1.9	1.8	162.6	0.24240					X	
2	19.5	19.2	20.5	0.21980		X	X			
2	17.6	17.3	36.3	0.22241			X			X
2	17.2	16.9	39.3	0.22290	X		X			
2	17.1	16.9	39.7	0.22297			X		X	
2	17.0	16.7	40.9	0.22316			X	X		
3	20.2	19.8	16.9	0.21903	X	X	X			
3	20.1	19.7	17.4	0.21911		X	X			X
3	20.1	19.7	17.6	0.21914		X	X		X	
3	20.0	19.6	18.5	0.21930		X	X	X		
3	19.7	19.3	20.7	0.21966	X		X			X
4	21.1	20.6	11.0	0.21789	X	X	X			X
4	21.0	20.5	12.1	0.21806		X	X		X	X
4	21.0	20.5	12.2	0.21808		X	X	X	X	
4	20.9	20.4	13.3	0.21826		X	X	X		X
4	20.7	20.2	14.5	0.21847	X	X	X	X		
5	21.9	21.3	6.8	0.21701		X	X	X	X	X
5	21.8	21.1	7.9	0.21719	X	X	X	X		X
5	21.3	20.7	11.7	0.21783	X	X	X		X	X
5	21.1	20.5	13.1	0.21807	X	X	X	X	X	
5	20.3	19.7	19.8	0.21920	X		X	X	X	X
6	22.1	21.4	7.0	0.21687	X	X	X	X	X	X

646 cases used, 89 cases contain missing values.

Appendix B (Cont.)

B-9. Regression: Best Subsets

VE1 versus Pop, PopDen, Sect, TIO, PerInc, EIA

Response (dependent variable) is VE1

Vars	R-Sq	R-Sq (adj)	Cp	S	Pop	Pop Den	Sect	TIO	Per Inc	EIA
1	10.1	9.9	74.4	0.091452		X				
1	10.1	9.8	74.5	0.091457					X	
1	9.8	9.5	76.1	0.091622				X		
1	9.6	9.3	77.1	0.091721	X					
1	0.1	0.0	124.8	0.096400						X
2	18.7	18.2	33.2	0.087095			X		X	
2	18.2	17.8	35.3	0.087320			X	X		
2	17.8	17.4	37.3	0.087535	X		X			
2	14.3	13.8	55.3	0.089416		X			X	
2	14.0	13.6	56.6	0.089553	X	X				
3	24.3	23.7	6.5	0.084112		X	X		X	
3	23.9	23.3	8.7	0.084349	X	X	X			
3	23.7	23.1	9.7	0.084462		X	X	X		
3	19.6	19.0	30.3	0.086689			X		X	X
3	19.1	18.5	33.0	0.086977			X	X		X
4	25.3	24.6	3.5	0.083664		X	X	X	X	
4	24.4	23.6	8.4	0.084203	X	X	X		X	
4	24.3	23.5	8.5	0.084222		X	X		X	X
4	23.9	23.1	10.6	0.084453	X	X	X	X		
4	23.9	23.1	10.7	0.084458	X	X	X			X
5	25.4	24.5	5.0	0.083722	X	X	X	X	X	
5	25.4	24.4	5.4	0.083768		X	X	X	X	X
5	24.4	23.4	10.4	0.084315	X	X	X		X	X
5	23.9	22.9	12.6	0.084563	X	X	X	X		X
5	20.1	19.0	31.9	0.086665	X		X	X	X	X
6	25.4	24.3	7.0	0.083831	X	X	X	X	X	X

383 cases used, 352 cases contain missing values.

Appendix B (Cont.)

B-10. Regression: Best Subsets

VE2 versus Pop, PopDen, Sect, TIO, PerInc, EIA

Response (dependent variable) is VE2

Vars	R-Sq	R-Sq (adj)	Cp	S	Pop	Pop Den	Sect	TIO	Per Inc	EIA
1	5.5	5.2	66.6	0.20149		X				
1	4.3	4.0	72.2	0.20276					X	
1	4.2	4.0	72.5	0.20282	X					
1	4.2	3.9	72.7	0.20287				X		
1	1.5	1.2	85.4	0.20569						X
2	15.3	14.9	22.1	0.19095			X		X	
2	15.3	14.8	22.3	0.19099	X		X			
2	15.2	14.8	22.6	0.19107			X	X		
2	10.1	9.6	46.9	0.19678		X	X			
2	9.6	9.1	49.1	0.19729	X					X
3	19.5	18.9	4.4	0.18641	X	X	X			
3	19.4	18.7	5.2	0.18661		X	X		X	
3	19.1	18.4	6.6	0.18695		X	X	X		
3	17.4	16.8	14.2	0.18882	X		X			X
3	17.2	16.6	15.1	0.18904			X		X	X
4	20.0	19.2	4.1	0.18610	X	X	X			X
4	19.8	18.9	5.3	0.18638		X	X		X	X
4	19.6	18.7	6.0	0.18657		X	X	X	X	
4	19.6	18.7	6.2	0.18662	X	X	X	X		
4	19.5	18.7	6.4	0.18666	X	X	X		X	
5	20.1	19.0	5.8	0.18627	X	X	X	X		X
5	20.0	19.0	5.9	0.18630		X	X	X	X	X
5	20.0	18.9	6.1	0.18634	X	X	X		X	X
5	19.7	18.7	7.4	0.18667	X	X	X	X	X	
5	17.5	16.4	18.0	0.18927	X		X	X	X	X
6	20.2	19.0	7.0	0.18632	X	X	X	X	X	X

383 cases used, 352 cases contain missing values.

Appendix B (Cont.)

B-11. Regression: Best Subsets

VP1 versus Pop, PopDen, Sect, TIO, PerInc, EIA

Response (dependent variable) is VP1

Vars	R-Sq	R-Sq (adj)	Cp	S	Pop	Pop Den	Sect	TIO	Per Inc	EIA
1	8.8	8.5	30.1	0.09949			X			
1	1.0	0.8	64.8	0.10362						X
1	0.6	0.4	66.7	0.10384	X					
1	0.6	0.3	67.0	0.10387				X		
1	0.5	0.2	67.5	0.10393					X	
2	11.4	11.0	20.3	0.09816		X	X			
2	11.4	10.9	20.6	0.09820			X		X	
2	11.1	10.6	21.7	0.09835			X	X		
2	10.9	10.4	22.6	0.09845	X		X			
2	8.9	8.5	31.4	0.09953			X			X
3	12.9	12.2	15.6	0.09747		X	X		X	
3	12.6	12.0	16.8	0.09761		X	X	X		
3	12.6	11.9	17.0	0.09764	X	X	X			
3	12.4	11.7	17.7	0.09773		X	X			X
3	12.4	11.7	18.0	0.09776			X	X	X	
4	14.6	13.7	9.9	0.09663		X	X	X	X	
4	13.6	12.7	14.5	0.09721	X	X	X		X	
4	13.6	12.7	14.6	0.09722	X		X	X	X	
4	13.4	12.5	15.5	0.09733		X	X		X	X
4	13.2	12.3	16.4	0.09745		X	X	X		X
5	15.6	14.5	7.5	0.09619	X	X	X	X	X	
5	15.0	13.9	10.3	0.09655		X	X	X	X	X
5	14.3	13.1	13.6	0.09697	X	X	X		X	X
5	13.6	12.4	16.5	0.09734	X		X	X	X	X
5	13.2	12.0	18.3	0.09756	X	X	X	X		X
6	16.2	14.8	7.0	0.09601	X	X	X	X	X	X

383 cases used, 352 cases contain missing values.

Appendix B (Cont.)

B-12. Regression: Best Subsets

VP2 versus Pop, PopDen, Sect, TIO, PerInc, EIA

Response (dependent variable) is VP2

Vars	R-Sq	R-Sq (adj)	Cp	S	Pop	Pop Den	Sect	TIO	Per Inc	EIA
1	28.7	28.5	29.2	0.12365			X			
1	6.2	5.9	158.0	0.14182						X
1	5.6	5.4	161.2	0.14225	X					
1	5.5	5.2	162.0	0.14235				X		
1	5.2	4.9	163.8	0.14259					X	
2	31.5	31.2	14.7	0.12129		X	X			
2	31.0	30.6	17.8	0.12177			X		X	
2	30.8	30.5	18.8	0.12192			X	X		
2	30.7	30.3	19.6	0.12204	X		X			
2	28.7	28.3	31.2	0.12381			X			X
3	32.8	32.3	9.7	0.12035		X	X		X	
3	32.6	32.1	10.7	0.12050		X	X	X		
3	32.6	32.1	10.7	0.12052	X	X	X			
3	31.9	31.4	14.6	0.12112		X	X			X
3	31.5	30.9	17.2	0.12153	X		X		X	
4	33.7	33.0	6.2	0.11965		X	X	X	X	
4	33.1	32.4	9.9	0.12023	X	X	X		X	
4	32.9	32.2	11.1	0.12042		X	X		X	X
4	32.7	32.0	12.0	0.12056		X	X	X		X
4	32.7	32.0	12.2	0.12059	X	X	X			X
5	34.2	33.3	5.6	0.11940	X	X	X	X	X	
5	33.8	32.9	7.9	0.11975		X	X	X	X	X
5	33.2	32.4	11.1	0.12026	X	X	X		X	X
5	32.7	31.8	14.0	0.12072	X	X	X	X		X
5	32.1	31.2	17.6	0.12129	X		X	X	X	X
6	34.3	33.3	7.0	0.11946	X	X	X	X	X	X

383 cases used, 352 cases contain missing values.

Appendix C – Regression Equation Information – Preliminary Model

C -1. Regression Analysis: LE1 versus Pop, Sect, PerInc, EIA

The regression equation is

$$LE1 = 1.23 - 0.0512 \text{ Pop} + 0.000517 \text{ Sect} + 0.000002 \text{ PerInc} + 0.00129 \text{ EIA}$$

650 cases used 85 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.23082	0.01041	118.22	0.000	
Pop	-0.05119	0.01254	-4.08	0.000	95.4
Sect	0.00051656	0.00004123	12.53	0.000	1.9
PerInc	0.00000158	0.00000048	3.29	0.001	92.8
EIA	0.0012899	0.0001490	8.66	0.000	1.3

S = 0.07720 R-Sq = 37.4% R-Sq(adj) = 37.0%

C -2. Regression Analysis: LE2 versus Pop, PopDen, Sect, TIO, PerInc, EIA

The regression equation is

$$LE2 = 1.42 - 0.100 \text{ Pop} - 0.000089 \text{ PopDen} + 0.00101 \text{ Sect} - 1.37 \text{ TIO} + 0.000006 \text{ PerInc} + 0.00164 \text{ EIA}$$

649 cases used 86 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.42122	0.01848	76.92	0.000	
Pop	-0.10010	0.02218	-4.51	0.000	96.3
PopDen	-0.00008941	0.00004018	-2.23	0.026	1.5
Sect	0.00101245	0.00007515	13.47	0.000	2.0
TIO	-1.3717	0.7161	-1.92	0.056	290.7
PerInc	0.00000605	0.00000176	3.44	0.001	400.9
EIA	0.0016483	0.0002837	5.81	0.000	1.6

S = 0.1359 R-Sq = 36.6% R-Sq(adj) = 36.0%

Appendix C (Cont.)

C-3. Regression Analysis: LP1 versus PopDen, Sect, EIA

The regression equation is

$$LP1 = 1.18 + 0.000092 \text{ PopDen} + 0.000474 \text{ Sect} + 0.00103 \text{ EIA}$$

649 cases used 86 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.18059	0.01288	91.66	0.000	
PopDen	0.00009170	0.00002914	3.15	0.002	1.3
Sect	0.00047360	0.00005229	9.06	0.000	1.6
EIA	0.0010269	0.0002168	4.74	0.000	1.5

S = 0.1063 R-Sq = 28.9% R-Sq(adj) = 28.5%

C-4. Regression Analysis: LP2 versus PopDen, Sect, EIA

The regression equation is

$$LP2 = 1.29 + 0.000095 \text{ PopDen} + 0.000961 \text{ Sect} + 0.00136 \text{ EIA}$$

649 cases used 86 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.28837	0.01669	77.22	0.000	
PopDen	0.00009467	0.00003775	2.51	0.012	1.3
Sect	0.00096086	0.00006774	14.18	0.000	1.6
EIA	0.0013628	0.0002809	4.85	0.000	1.5

S = 0.1377 R-Sq = 43.6% R-Sq(adj) = 43.3%

Appendix C (Cont.)

C-5. Regression Analysis: SE1 versus Pop, PopDen, Sect, PerInc, EIA

The regression equation is

SE1 = 1.64 - 0.0788 Pop -0.000210 PopDen +0.000508 Sect +0.000002 PerInc
+ 0.00139 EIA

646 cases used 89 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.63906	0.01651	99.31	0.000	
Pop	-0.07876	0.01987	-3.96	0.000	95.9
PopDen	-0.00021024	0.00003510	-5.99	0.000	1.4
Sect	0.00050794	0.00006723	7.55	0.000	2.0
PerInc	0.00000207	0.00000076	2.70	0.007	94.0
EIA	0.0013856	0.0002556	5.42	0.000	1.6

S = 0.1220 R-Sq = 27.5% R-Sq(adj) = 26.9%

C-6. Regression Analysis: SE2 versus Pop, PopDen, Sect, PerInc, EIA

The regression equation is

SE2 = 1.93 - 0.164 Pop -0.000311 PopDen + 0.00121 Sect +0.000005 PerInc
+ 0.00223 EIA

646 cases used 89 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.93169	0.02951	65.46	0.000	
Pop	-0.16438	0.03553	-4.63	0.000	95.9
PopDen	-0.00031113	0.00006275	-4.96	0.000	1.4
Sect	0.0012079	0.0001202	10.05	0.000	2.0
PerInc	0.00000471	0.00000137	3.44	0.001	94.0
EIA	0.0022270	0.0004569	4.87	0.000	1.6

S = 0.2181 R-Sq = 27.6% R-Sq(adj) = 27.0%

Appendix C (Cont.)

C -7. Regression Analysis: SP1 versus PopDen, Sect, TIO, PerInc, EIA

The regression equation is

SP1 = 1.55 -0.000181 PopDen +0.000472 Sect + 2.48 TIO -0.000006 PerInc
+0.000889 EIA

646 cases used 89 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.54900	0.02380	65.09	0.000	
PopDen	-0.00018108	0.00005085	-3.56	0.000	1.5
Sect	0.00047173	0.00009601	4.91	0.000	2.0
TIO	2.4813	0.9124	2.72	0.007	283.9
PerInc	-0.00000553	0.00000189	-2.92	0.004	279.4
EIA	0.0008890	0.0003648	2.44	0.015	1.6

S = 0.1752 R-Sq = 10.7% R-Sq(adj) = 10.0%

C -8. Regression Analysis: SP2 versus PopDen, Sect, TIO, PerInc, EIA

The regression equation is

SP2 = 1.69 -0.000238 PopDen + 0.00107 Sect + 3.05 TIO -0.000007 PerInc
+ 0.00123 EIA

646 cases used 89 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.69454	0.02948	57.48	0.000	
PopDen	-0.00023813	0.00006299	-3.78	0.000	1.5
Sect	0.0010699	0.0001189	9.00	0.000	2.0
TIO	3.046	1.130	2.69	0.007	283.9
PerInc	-0.00000681	0.00000234	-2.90	0.004	279.4
EIA	0.0012273	0.0004519	2.72	0.007	1.6

S = 0.2170 R-Sq = 21.9% R-Sq(adj) = 21.3%

Appendix C (Cont.)

C -9. Regression Analysis: VE1 versus PopDen, Sect, TIO, PerInc

The regression equation is

$$VE1 = 1.44 - 0.000140 \text{ PopDen} + 0.000397 \text{ Sect} + 1.04 \text{ TIO} - 0.000003 \text{ PerInc}$$

383 cases used 352 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.44065	0.01686	85.45	0.000	
PopDen	-0.00014037	0.00002452	-5.72	0.000	1.3
Sect	0.00039712	0.00005731	6.93	0.000	1.7
TIO	1.0431	0.4632	2.25	0.025	293.6
PerInc	-0.00000276	0.00000096	-2.88	0.004	288.5

S = 0.08366 R-Sq = 25.3% R-Sq(adj) = 24.6%

C -10. Regression Analysis: VE2 versus Pop, PopDen, Sect, EIA

The regression equation is

$$VE2 = 1.75 - 0.0314 \text{ Pop} - 0.000202 \text{ PopDen} + 0.000912 \text{ Sect} + 0.000812 \text{ EIA}$$

383 cases used 352 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.75302	0.03730	47.00	0.000	
Pop	-0.031417	0.004603	-6.83	0.000	2.0
PopDen	-0.00020185	0.00005785	-3.49	0.001	1.4
Sect	0.0009125	0.0001356	6.73	0.000	2.0
EIA	0.0008123	0.0005372	1.51	0.131	1.6

S = 0.1861 R-Sq = 20.0% R-Sq(adj) = 19.2%

Appendix C (Cont.)

C -11. Regression Analysis: VP1 versus Pop, PopDen, Sect, TIO, PerInc, EIA

The regression equation is

$$VP1 = 1.33 + 0.0374 \text{ Pop} - 0.000103 \text{ PopDen} + 0.000484 \text{ Sect} + 1.57 \text{ TIO} - 0.000005 \text{ PerInc} - 0.000437 \text{ EIA}$$

383 cases used 352 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.33386	0.01941	68.71	0.000	
Pop	0.03741	0.01632	2.29	0.022	94.8
PopDen	-0.00010326	0.00003042	-3.39	0.001	1.5
Sect	0.00048424	0.00007036	6.88	0.000	2.0
TIO	1.5675	0.5358	2.93	0.004	298.3
PerInc	-0.00000486	0.00000133	-3.65	0.000	421.8
EIA	-0.0004367	0.0002786	-1.57	0.118	1.6

S = 0.09601 R-Sq = 16.2% R-Sq(adj) = 14.8%

C -12. Regression Analysis: VP2 versus Pop, PopDen, Sect, TIO, PerInc

The regression equation is

$$VP2 = 1.46 + 0.0322 \text{ Pop} - 0.000123 \text{ PopDen} + 0.000986 \text{ Sect} + 1.67 \text{ TIO} - 0.000005 \text{ PerInc}$$

383 cases used 352 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.45712	0.02406	60.56	0.000	
Pop	0.03225	0.02013	1.60	0.110	93.3
PopDen	-0.00012316	0.00003508	-3.51	0.000	1.3
Sect	0.00098639	0.00008216	12.01	0.000	1.8
TIO	1.6679	0.6655	2.51	0.013	297.6
PerInc	-0.00000497	0.00000166	-3.00	0.003	421.8

S = 0.1194 R-Sq = 34.2% R-Sq(adj) = 33.3%

Appendix D – Analysis of Variance Results

D-1. Analysis of Variance for LE1 ^a					
Source	DF	SS	MS	F	P
GF	7	1.518250	0.21689	30.11	0
Error	6424	6.2528	0.0072		
Total	6496	14.353			

Individual 95% CIs For Mean Based on Pooled StDev					
Level	N	Mean	StDev	-----+-----+-----+-----	
1	75	1.3947	0.0785	(---*---)	
2	59	1.4350	0.0777	(---*---)	
3	39	1.4261	0.0714	(---*---)	
4	81	1.4637	0.0872	(---*---)	
5	63	1.4323	0.0728	(---*---)	
6	103	1.4038	0.1010	(---*---)	
8	154	1.3302	0.0861	(--*--)	
9	76	1.3355	0.0827	(---*---)	

Pooled StDev = 0.0849	1.350	1.400	1.450
-----------------------	-------	-------	-------

Family error rate = 0.507
Individual error rate = 0.0500

Critical value = 1.964

^a The distribution of means and standard deviations in the one-way analysis of variance above shows clearly that the multipliers for Forest Service Regions 1-6 (dummy = 0) are different than those for Regions 8 and 9 (dummy = 1). This presents an opportunity to divide the samples into two groups to try to improve the strength of the model. This was done by inserting dummy variables in the regression equation. This has been labeled the LE1Dum.

Appendix D (Cont.)

D-2. Analysis of Variance for LE2 ^a					
Source	DF	SS	MS	F	P
GF	7	5.7783	0.8	40.94000	0
Error	642	12.9444	0.0		
Total	649	18.7227			

Individual 95% CIs For Mean Based on Pooled StDev					
Level	N	Mean	StDev	-----+-----+-----+-----+-----	
1	75	1.7767	0.1455		(---*---)
2	59	1.7227	0.1028		(---*---)
3	39	1.6637	0.1109		(-----*-----)
4	81	1.7984	0.1338		(---*---)
5	63	1.7993	0.1244		(---*---)
6	103	1.8177	0.1747		(---*---)
8	154	1.5915	0.1515	(--*--)	
9	76	1.6034	0.1307	(---*---)	
				-----+-----+-----+-----+-----	
Pooled StDev = 0.1420				1.600	1.680 1.760 1.840

Fisher's pairwise comparisons

Family error rate = 0.507
Individual error rate = 0.0500

Critical value = 1.964

^a The distribution of means and standard deviations in the one-way analysis of variance above shows clearly that the multipliers for Forest Service Regions 1-6 (dummy = 0) are different than those for Regions 8 and 9 (dummy = 1). Also, Region 3 seems to be separated from other regions. This presents an opportunity to at least divide the samples into 2 groups to try to improve the strength of the model. This was done by inserting dummy variables in regression equation. This has been labeled the LE2Dum.

Appendix D (Cont.)

D-3. Analysis of Variance for LP1 ^a					
Source	DF	SS	MS	F	P
GF	7	3.2130	0.4590	41.84	0.000
Error	642	7.0424	0.0110		
Total	649	10.2554			

Individual 95% CIs For Mean					
Based on Pooled StDev					
Level	N	Mean	StDev	-----+-----+-----+-----+	
1	75	1.2666	0.0746	(--*)	
2	59	1.4279	0.1030		(--*)
3	39	1.5459	0.1457		(---*--)
4	81	1.3954	0.1604		(--*)
5	63	1.3460	0.1009	(--*)	
6	103	1.2745	0.0686	(-*)	
8	154	1.3275	0.0899	(-*)	
9	76	1.3350	0.1034	(--*)	
-----+-----+-----+-----+					
Pooled StDev = 0.1047				1.30	1.40 1.50 1.60

Fisher's pairwise comparisons

Family error rate = 0.507
Individual error rate = 0.0500

Critical value = 1.964

^a The distribution of means and standard deviations in the one-way analysis of variance above shows clearly that the multipliers for Forest Service Regions 1 and 6 (dummy = 0) are different than 2,4,5,8,9 (dummy = 1) which are different than those for Region 3 (dummy = 2). This presents an opportunity to at least divide the samples into 3 separate groups to try to improve the strength of the model. This was done by inserting dummy variables in the regression equation. This has been labeled the LP1Dum.

Appendix D (Cont.)

D-4. Analysis of Variance for LP2^a

Source	DF	SS	MS	F	P
GF	7	5.2676	0.7525	29.45	0
Error	642	16.4065	0.0256		
Total	649	21.6741			

Individual 95% CIs For Mean Based on Pooled St Dev

Level	N	Mean	StDev	-----+-----+-----+-----
1	75	1.4940	0.1148	(--*---)
2	59	1.7104	0.1543	(---*--)
3	39	1.8493	0.1943	(---*---)
4	81	1.6676	0.2325	(--*--)
5	63	1.6232	0.1575	(--*---)
6	103	1.5226	0.1232	(--*-)
8	154	1.5606	0.1473	(-*--)
9	76	1.5836	0.1585	(--*--)

Pooled St Dev = 0.1599 1.56 1.68 1.80

Fisher's pairwise comparisons

Family error rate = 0.507
Individual error rate = 0.0500

Critical value = 1.964

^a The distribution of means and standard deviations in the one-way analysis of variance above shows clearly that the multipliers for Forest Service Regions 1 and 6 (dummy = 0) are different than 2,4,5,8,9 (dummy = 1) which are different than those for Region 3 (dummy = 2). This presents an opportunity to at least divide the samples into 3 separate groups to try to improve the strength of the model. This was done by inserting dummy variables in the regression equation. This was labeled the LP2Dum.

Appendix D (Cont.)

D-5. Analysis of Variance for SE1 ^a					
Source	DF	SS	MS	F	P
GF	7	1.3002	0.1857	10.01	0
Error	639	11.8528	0.0185		
Total	646	13.153			

Individual 95% CIs For Mean Based on Pooled St Dev					
Level	N	Mean	St Dev	-----+-----+-----+-----	
1	73	1.8286	0.0928		(---*---)
2	61	1.7899	0.1176		(---*---)
3	38	1.7157	0.1676	(-----*-----)	
4	76	1.7815	0.1574		(---*---)
5	65	1.7288	0.1731	(---*---)	
6	102	1.7994	0.1084		(---*---)
8	155	1.7667	0.1417	(--*---)	
9	77	1.6733	0.1312	(---*---)	
				-----+-----+-----+-----	
Pooled St Dev = 0.1362				1.680	1.740 1.800

Fisher's pairwise comparisons

Family error rate = 0.507
Individual error rate = 0.0500

Critical value = 1.964

^a The distribution of means and standard deviations in the one-way analysis of variance above shows clearly that the multipliers for Forest Service Regions 3 and 5 and 9 (dummy = 0) are different than the other regions which are fairly similar (dummy = 1). This presents an opportunity to at least divide the samples into 2 geographic groups which might improve the strength of the model. They have been labeled the SE1Dum.

Appendix D (Cont.)

D-6. Analysis of Variance for SE2 ^a					
Source	DF	SS	MS	F	P
GF	7	9.3939	1.342	26.22	0
Error	639	32.7006	0.0512		
Total	646	42.0945			

Individual 95% CIs For Mean					
Based on Pooled StDev					
Level	N	Mean	StDev	-----+-----+-----+-----	
1	73	2.3904	0.1738		(--*--)
2	61	2.1737	0.1707	(---*---)	
3	38	2.1020	0.2168	(---*---)	
4	76	2.3071	0.2998		(---*--)
5	65	2.2656	0.2831		(---*---)
6	102	2.4122	0.2116		(--*--)
8	155	2.1768	0.2251	(-*--)	
9	77	2.0518	0.1934	(---*--)	
-----+-----+-----+-----					
Pooled StDev = 0.2262			2.10	2.25	2.40

Fisher's pairwise comparisons

Family error rate = 0.507
Individual error rate = 0.0500

Critical value = 1.964

^a The distribution of means and standard deviations in the one-way analysis of variance above shows clearly that the multipliers for Forest Service Regions 2,3,8 and 9 (dummy = 0) are different than 1 and 6 (dummy = 1) which are different than 4 and 5 (dummy = 2). This presents an opportunity to at least divide the samples into 3 geographic groups which might improve the strength of the model. They have been labeled the SE2Dum.

Appendix D (Cont.)

D-7. Analysis of Variance for SP1 ^a					
Source	DF	SS	MS	F	P
GF	7	4.8545	0.6935	25.84	0
Error	639	17.1475	0.0268		
Total	646	22.002			

Individual 95% CIs For Mean					
Based on Pooled StDev					
Level	N	Mean	StDev	-----+-----+-----+-----	
1	73	1.6369	0.1111	(--*--)	
2	61	1.9066	0.1814		(---*--)
3	38	1.6899	0.1403	(----*--)	
4	76	1.6841	0.2316	(--*--)	
5	65	1.5807	0.0972	(---*--)	
6	102	1.5892	0.0772	(-*--)	
8	155	1.6936	0.1630	(-*--)	
9	77	1.7105	0.2361	(---*--)	
			-----+-----+-----+-----		
Pooled StDev = 0.1638			1.56	1.68	1.80 1.92

Fisher's pairwise comparisons

Family error rate = 0.507
Individual error rate = 0.0500

Critical value = 1.964

^a The distribution of means and standard deviations in the one-way analysis of variance above shows clearly that the multipliers for Forest Service Regions 5 and 6 (dummy = 0) are different than 1,3,4,8 and 9 (dummy = 1) which are different than Region 2 (dummy = 2). This presents an opportunity to at least divide the samples into 3 geographic groups which might improve the strength of the model. They have been labeled the SP1Dum.

Appendix D (Cont.)

D-8. Analysis of Variance for SP2 ^a					
Source	DF	SS	MS	F	P
GF	7	7.0007	1.0001	20.23	0
Error	639	31.5845	0.0494		
Total	646	38.5852			

Individual 95% CIs For Mean					
Based on Pooled St Dev					
Level	N	Mean	St Dev	-----+-----+-----+-----	
1	73	1.9302	0.1592	(---*--)	
2	61	2.2828	0.2464		(---*---)
3	38	2.0243	0.1831	(----*----)	
4	76	2.0156	0.3167	(--*---)	
5	65	1.9051	0.1390	(---*---)	
6	102	1.8974	0.1279	(-*--)	
8	155	1.9900	0.2262	(--*--)	
9	77	2.0271	0.2945	(--*--)	
-----+-----+-----+-----					
Pooled St Dev = 0.2223				1.95	2.10 2.25

Fisher's pairwise comparisons

Family error rate = 0.507
Individual error rate = 0.0500

Critical value = 1.964

^a The distribution of means and standard deviations in the one-way analysis of variance above shows clearly that the multipliers for Forest Service Regions 1,5 and 6 (dummy = 0) are different than 3,4,8 and 9 (dummy = 1) that are different than Region 2 (dummy = 2). This presents an opportunity to at least divide the samples into 3 geographic groups that might improve the strength of the model. They have been labeled the SP2Dum.

Appendix D (Cont.)

D-9. Analysis of Variance for VE1 ^a					
Source	DF	SS	MS	F	P
GF	7	0.70949	0.10136	13.44	0
Error	376	2.83474	0.00754		
Total	383	3.54423			

Individual 95% CIs For Mean

Based on Pooled StDev

Level	N	Mean	StDev	
1	41	1.5953	0.0483	(---*---)
2	11	1.4811	0.0523	(-----*-----)
3	12	1.4143	0.1032	(-----*-----)
4	17	1.4620	0.1198	(-----*-----)
5	42	1.4805	0.1135	(---*---)
6	96	1.5587	0.0858	(--*--)
8	117	1.5412	0.0889	(-*--)
9	48	1.4911	0.0689	(---*---)

Pooled StDev = 0.0868 1.400 1.470 1.540 1.610

Fisher's pairwise comparisons

Family error rate = 0.505
Individual error rate = 0.0500

Critical value = 1.966

^a The distribution of means and standard deviations in the one-way analysis of variance above shows clearly that the multipliers for Forest Service Regions 2,3,4,5 and 9 (dummy = 0) are different than 1,6, and 8 (dummy = 1). This presents an opportunity to at least divide the samples into 2 geographic groups which might improve the strength of the model. They have been labeled the VE1Dum.

Appendix D (Cont.)

D-10. Analysis of Variance for VE2 ^a					
Source	DF	SS	MS	F	P
GF	7	4.3638	0.6234	19.53	0
Error	376	12.0036	0.0319		
Total	383	16.3674			

Individual 95% CIs For Mean			
Based on Pooled StDev			
Level	N	Mean	StDev
1	41	2.2058	0.1336
2	11	1.9261	0.1212
3	12	1.7311	0.1889
4	17	1.9303	0.1982
5	42	1.9344	0.2151
6	96	2.0681	0.1859
8	117	1.9577	0.1913
9	48	1.8685	0.1217

Pooled StDev = .1787	1.80	2.00	2.20
----------------------	------	------	------

Fisher's pairwise comparisons

Family error rate = 0.505
Individual error rate = 0.0500

Critical value = 1.966

^a The distribution of means and standard deviations in the one-way analysis of variance above shows clearly that the multipliers for Forest Service Regions 3 (dummy = 0) is different than 2,4,5,6,8,and 9 (dummy = 1) which is different that region 1 (dummy = 2). This presents an opportunity to at least divide the samples into 3 geographic groups which might improve the strength of the model. They have been labeled the VE2Dum.

Appendix D (Cont.)

D-11. Analysis of Variance for VP1 ^a					
Source	DF	SS	MS	F	P
GF	7	0.3268	0.0467	4.6	0
Error	376	3.8115	0.0101		
Total	383	4.1383			

Individual 95% CIs For Mean					
Based on Pooled StDev					
Level	N	Mean	StDev	-----+-----+-----+-----	
1	41	1.4218	0.0568	(---*---)	
2	11	1.4468	0.0265	(-----*-----)	
3	12	1.4750	0.1569	(-----*-----)	
4	17	1.3837	0.1315	(-----*-----)	
5	42	1.4397	0.0965	(---*---)	
6	96	1.4723	0.0738	(--*--)	
8	117	1.4777	0.1182	(--*--)	
9	48	1.5064	0.1139	(---*---)	
				-----+-----+-----+-----	
Pooled StDev = 0.1007				1.380	1.440 1.500

Fisher's pairwise comparisons

Family error rate = 0.505
Individual error rate = 0.0500

Critical value = 1.966

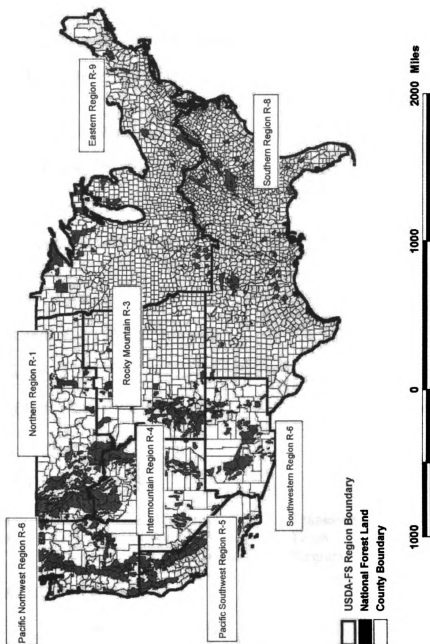
^a The distribution of means and standard deviations in the one-way analysis of variance above shows clearly that the multipliers for Forest Service Region 4 (dummy = 0) is different than 1,2,3,5,6,8,and 9 (dummy = 1). This presents an opportunity to at least divide the samples into 2 geographic groups which might improve the strength of the model. They have been labeled the VP1Dum.

Appendix D (Cont.)

D-12. Analysis of Variance for VP2 ^a					
Source	DF	SS	MS	F	P
GF	7	0.4308	0.0615	2.99	0.005
Error	376	7.7449	0.0206		
Total	383	8.1757			
Individual 95% CIs For Mean					
Based on Pooled StDev					
Level	N	Mean	StDev	-----+-----+-----+-----	
1	41	1.6955	0.0877	(-----*-----)	
2	11	1.7929	0.0313	(-----*-----)	
3	12	1.7747	0.2043	(-----*-----)	
4	17	1.6597	0.1745	(-----*-----)	
5	42	1.7385	0.1363	(-----*-----)	
6	96	1.7615	0.1236	(--*--)	
8	117	1.7483	0.1634	(---*---)	
9	48	1.8009	0.1578	(-----*-----)	
				-----+-----+-----+-----	
Pooled StDev = 0.1435		1.600	1.680	1.760	1.840
Fisher's pairwise comparisons					
Family error rate = 0.505					
Individual error rate = 0.0500					
Critical value = 1.966					

^a The distribution of means and standard deviations in the one-way analysis of variance above shows clearly that the multipliers for Forest Service Regions 1 and 4 (dummy = 0) are different than 2,,5,6,8,and 9 (dummy = 1). This presents an opportunity to at least divide the samples into 2 geographic groups which might improve the strength of the model. They have been labeled the VP2Dum.

APPENDIX E – Study Area Map



Appendix F – States in USDA Forest Service Regions

REGION 1 – NORTHERN REGION

Idaho Montana North Dakota

REGION 2 – ROCKY MOUNTAIN REGION

Colorado South Dakota
Nebraska Wyoming

REGION 3 – SOUTHWESTERN REGION

Arizona New Mexico

REGION 4 – INTERMOUNTAIN REGION

Idaho Utah
Nevada Wyoming

REGION 5 – PACIFIC SOUTHWEST REGION

California

REGION 6 – PACIFIC NORTHWEST REGION

Oregon Washington

REGION 8 – SOUTHERN REGION

Alabama Georgia Mississippi Tennessee
Arkansas Kentucky North Carolina Texas
Florida Louisiana South Carolina Virginia

REGION 9 – EASTERN REGION

Illinois Minnesota Ohio West Virginia
Indiana Missouri Pennsylvania Wisconsin
Michigan New Hampshire Vermont

REGION 10 – ALASKA REGION

Alaska

Appendix G – Final Regression Analysis Results

G-1. Final Regression Analysis: LE1 vs Pop, Sect, EIA, LE1Dum

The regression equation is

$$LE1 = 1.25 - 0.0110 \text{ Pop} + 0.000603 \text{ Sect} + 0.000387 \text{ EIA} - 0.0855 \text{ LE1Dum.}$$

650 cases used 85 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.25381	0.00938	133.72	0.000	
Pop	-0.010950	0.001471	-7.44	0.000	1.7
Sect	0.00060279	0.00003712	16.24	0.000	2.0
EIA	0.0003873	0.0001443	2.68	0.007	1.6
LE1Dum.	-0.085494	0.006205	-13.78	0.000	1.2

S = 0.06843 R-Sq = 50.8% R-Sq(adj) = 50.5%

G-2. Final Regression Analysis: LE2 vs PopDen, Sect, TIO, EIA, LE2Dum

The regression equation is

$$LE2 = 1.48 - 0.000120 \text{ PopDen} + 0.00118 \text{ Sect} - 0.324 \text{ TIO} - 0.000380 \text{ EIA} - 0.184 \text{ LE2Dum}$$

649 cases used 86 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.47588	0.01543	95.62	0.000	
PopDen	-0.00012024	0.00003296	-3.65	0.000	1.4
Sect	0.00118375	0.00006276	18.86	0.000	2.1
TIO	-0.32382	0.04752	-6.81	0.000	1.9
EIA	-0.0003796	0.0002565	-1.48	0.139	1.9
LE2Dum	-0.18431	0.01025	-17.98	0.000	1.2

S = 0.1126 R-Sq = 56.4% R-Sq(adj) = 56.1%

Appendix G (Cont.)

G-3. Final Regression Analysis: LP1 vs PopDen, Sect, EIA, LP1Dum

The regression equation is

$$\text{LP1} = 1.01 + 0.000072 \text{ PopDen} + 0.000476 \text{ Sect} + 0.000748 \text{ EIA} + 0.102 \text{ LP1Dum}.$$

649 cases used 86 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.00563	0.01601	62.83	0.000	
PopDen	0.00007240	0.00002508	2.89	0.004	1.3
Sect	0.00047649	0.00004495	10.60	0.000	1.6
EIA	0.0007485	0.0001873	4.00	0.000	1.5
LP1Dum	0.101804	0.006727	15.13	0.000	1.0

S = 0.09136 R-Sq = 47.5% R-Sq(adj) = 47.2%

G-4. Final Regression Analysis: LP2 vs PopDen, Sect, EIA, LP2Dum

The regression equation is

$$\text{LP2} = 1.08 + 0.000072 \text{ PopDen} + 0.000964 \text{ Sect} + 0.00104 \text{ EIA} + 0.119 \text{ LP2Dum}.$$

649 cases used 86 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.08465	0.02143	50.61	0.000	
PopDen	0.00007220	0.00003358	2.15	0.032	1.3
Sect	0.00096423	0.00006018	16.02	0.000	1.6
EIA	0.0010385	0.0002507	4.14	0.000	1.5
LP2Dum	0.118536	0.009007	13.16	0.000	1.0

S = 0.1223 R-Sq = 55.5% R-Sq(adj) = 55.2%

Appendix G (Cont.)

G-5. Final Regression Analysis: SE1 vs Pop, PopDen, Sect, EIA, SE1Dum

The regression equation is

$$SE1 = 1.53 - 0.0228 \text{ Pop} - 0.000193 \text{ PopDen} + 0.000473 \text{ Sect} \\ + 0.00128 \text{ EIA} + 0.0626 \text{ SE1Dum}.$$

646 cases used 89 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.53447	0.02422	63.35	0.000	
Pop	-0.022765	0.002758	-8.26	0.000	1.9
PopDen	-0.00019295	0.00003420	-5.64	0.000	1.4
Sect	0.00047331	0.00006578	7.20	0.000	2.0
EIA	0.0012831	0.0002498	5.14	0.000	1.6
SE1Dum	0.06261	0.01077	5.82	0.000	1.0

S = 0.1196 R-Sq = 30.3% R-Sq(adj) = 29.8%

G-6. Final Regression Analysis: SE2 vs Pop, PopDen, Sect, EIA, SE2Dum

The regression equation is

$$SE2 = 1.78 - 0.0459 \text{ Pop} - 0.000326 \text{ PopDen} + 0.00127 \text{ Sect} \\ + 0.00100 \text{ EIA} + 0.0933 \text{ SE2Dum}.$$

646 cases used 89 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.77667	0.03345	53.12	0.000	
Pop	-0.045871	0.004737	-9.68	0.000	1.9
PopDen	-0.00032552	0.00005972	-5.45	0.000	1.4
Sect	0.0012687	0.0001150	11.04	0.000	2.0
EIA	0.0010022	0.0004532	2.21	0.027	1.7
SE2Dum	0.09327	0.01080	8.64	0.000	1.1

S = 0.2083 R-Sq = 34.0% R-Sq(adj) = 33.4%

Appendix G (Cont.)

G-7. Final Regression Analysis: SP1 vs PopDen, Sect, TIO, EIA, SP1Dum

The regression equation is

$$\text{MSP1} = 1.25 - 0.000157 \text{ PopDen} + 0.000588 \text{ Sect} - 0.0552 \text{ TIO} + 0.000310 \text{ EIA} + 0.147 \text{ SP1Dum}.$$

646 cases used 89 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.25007	0.03067	40.76	0.000	
PopDen	-0.00015717	0.00004502	-3.49	0.001	1.4
Sect	0.00058816	0.00008583	6.85	0.000	2.0
TIO	-0.05522	0.06642	-0.83	0.406	1.9
EIA	0.0003097	0.0003283	0.94	0.346	1.6
SP1Dum	0.14724	0.01114	13.22	0.000	1.1

S = 0.1563 R-Sq = 29.0% R-Sq(adj) = 28.4%

G-8. Final Regression Analysis: SP2 vs PopDen, Sect, TIO, EIA, SP2Dum

The regression equation is

$$\text{SP2} = 1.42 - 0.000223 \text{ PopDen} + 0.00110 \text{ Sect} - 0.111 \text{ TIO} + 0.000846 \text{ EIA} + 0.156 \text{ SP2Dum}.$$

646 cases used 89 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.41916	0.03420	41.49	0.000	
PopDen	-0.00022334	0.00005646	-3.96	0.000	1.4
Sect	0.0011006	0.0001072	10.26	0.000	2.0
TIO	-0.11053	0.08303	-1.33	0.184	1.9
EIA	0.0008461	0.0004092	2.07	0.039	1.6
SP2Dum	0.15555	0.01249	12.45	0.000	1.0

S = 0.1960 R-Sq = 36.3% R-Sq(adj) = 35.8%

Appendix G (Cont.)

G-9. Final Regression Analysis: VE1 vs PopDen, Sect, PerInc, VE1Dum

The regression equation is

$$\text{VE1} = 1.32 - 0.000130 \text{ PopDen} + 0.000400 \text{ Sect} - 0.000000 \text{ PerInc} + 0.0683 \text{ VE1Dum}.$$

383 cases used 352 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.32105	0.02153	61.35	0.000	
PopDen	-0.00012976	0.00002235	-5.81	0.000	1.2
Sect	0.00040029	0.00005334	7.50	0.000	1.7
PerInc	-0.00000049	0.00000007	-6.72	0.000	1.9
VE1Dum	0.068341	0.008729	7.83	0.000	1.1

S = 0.07813 R-Sq = 34.9% R-Sq(adj) = 34.2%

G-10. Final Regression Analysis: VE2 vs Pop, PopDen, Sect, VE2Dum

The regression equation is

$$\text{VE2} = 1.24 - 0.0277 \text{ Pop} - 0.000220 \text{ PopDen} + 0.00104 \text{ Sect} + 0.233 \text{ VE2Dum}$$

383 cases used 352 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.24418	0.06098	20.40	0.000	
Pop	-0.027703	0.003933	-7.04	0.000	1.8
PopDen	-0.00022042	0.00004743	-4.65	0.000	1.2
Sect	0.0010385	0.0001142	9.10	0.000	1.7
VE2Dum	0.23262	0.02355	9.88	0.000	1.0

S = 0.1664 R-Sq = 36.0% R-Sq(adj) = 35.4%

Appendix G (Cont.)

G-11. Final Regression Analysis: VP1 vs PerInc, PopDen, Sect, EIA, VP1Dum

The regression equation is

VP1 = 1.15 -0.000000 PerInc -0.000076 PopDen +0.000478 Sect
+0.000041 EIA + 0.0920 VP1Dum.

383 cases used 352 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.14964	0.05459	21.06	0.000	
PerInc	-0.00000024	0.00000009	-2.62	0.009	2.1
PopDen	-0.00007592	0.00003016	-2.52	0.012	1.5
Sect	0.00047818	0.00007047	6.79	0.000	2.0
EIA	0.0000413	0.0003033	0.14	0.892	1.9
VP1Dum	0.09198	0.02628	3.50	0.001	1.2

S = 0.09592 R-Sq = 16.1% R-Sq(adj) = 15.0%

G-12. Final Regression Analysis: VP2 vs Pop, PopDen, Sect, VP2Dum

The regression equation is

VP2 = 1.35 - 0.00751 Pop -0.000123 PopDen +0.000995 Sect
+ 0.0583 VP2Dum.

383 cases used 352 cases contain missing values

Predictor	Coef	SE Coef	T	P	VIF
Constant	1.35134	0.03855	35.06	0.000	
Pop	-0.007508	0.002813	-2.67	0.008	1.8
PopDen	-0.00012342	0.00003402	-3.63	0.000	1.2
Sect	0.00099530	0.00008151	12.21	0.000	1.7
VP2Dum	0.05834	0.01726	3.38	0.001	1.0

S = 0.1189 R-Sq = 34.6% R-Sq(adj) = 33.9%

Appendix H – Descriptive Statistics for Impact Area Levels

H-1. Descriptive Statistics for Impact Area Levels - Population

Impact Level ^a	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	321474	79039	1107081	91310	9004	9116506
B	579538	187058	1649802	136073	6639	11446843
C	1688949	689737	3209368	264704	89040	19743435
E	1926399	789207	2883554	237831	92639	17078752

H-2. Descriptive Statistics for Impact Area Levels - Population Density

Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	151.4	47.7	317.9	26.2	1.8	2245.4
B	49.19	24.3	74.67	6.16	1.7	488.6
C	55.4	36.7	62.51	5.16	3.8	384.7
E	75.73	44.9	94.08	7.79	4.2	647.7

H-3. Descriptive Statistics for Impact Area Levels - No. of Economic Sectors

Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	191.68	175	76.76	6.33	95	483
B	236.41	222	82.8	6.83	61	495
C	338.57	335	76.28	6.29	182	505
E	345.37	338	83.06	6.85	183	502

H-4. Descriptive Statistics for Impact Area Levels - Total Industry Output

Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	16297	3340	62339	5142	229	513996
B	27180	7696	87046	7179	175	614238
C	80868	29706	164582	13575	3923	1008888
E	103669	32777	164357	13556	4310	874502

- ^a Impact Level A – Most populous county in National Forest.
 Impact Level B – All counties in National Forest.
 Impact Level C – Level 2 plus all adjacent counties.
 Impact Level E – BEA economic areas.

Appendix H (Cont.)

H-5. Descriptive Statistics for Impact area Levels - Personal Income

Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	7.792E+09	1.65E+09	2.85E+10	2.35E+09	962092	2.34E+11
B	1.35E+10	3.63E+09	4.14E+10	3.41E+09	83153000	2.90E+11
C	3.92E+10	1.46E+10	7.98E+10	6.58E+09	1.66E+09	4.91E+11
E	4.86E+10	1.64E+10	7.83E+10	6.46E+09	1.77E+09	4.25E+11

H-6. Descriptive Statistics for Impact Area Levels - Area (sq mi)

Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	2958	1909	3317	274	251	20062
B	12700	10335	9852	813	428	48674
C	37780	31482	30367	2505	3489	189303
E	33543	32019	22810	1881	3430	87808

H-7. Descriptive Statistics for Impact Area Level Multipliers - LE1

Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	1.3518	1.3579	0.1046	0.0092	1.0868	1.5554
B	1.357	1.3521	0.0941	0.0078	1.1661	1.5775
C	1.4183	1.4186	0.0903	0.0074	1.1946	1.5863
E	1.4318	1.4422	0.0849	0.007	1.2409	1.5592

H-8. Descriptive Statistics for Impact Area Level Multipliers - LE2

Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	1.6253	1.6375	0.1679	0.0148	1.1825	2.0092
B	1.6509	1.6534	0.1614	0.0134	1.2699	2.1341
C	1.7664	1.7632	0.1648	0.0136	1.3524	2.1144
E	1.7782	1.7787	0.1533	0.0127	1.3962	2.0955

Appendix H (Cont.)

H-9. Descriptive Statistics for Impact Area Level Multipliers - LP1						
Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	1.3116	1.2709	0.1301	0.0115	1.0966	1.7477
B	1.3103	1.281	0.1179	0.0098	1.156	1.8096
C	1.3677	1.3524	0.1165	0.0096	1.1905	1.8338
E	1.3943	1.387	0.1202	0.01	1.1858	1.8688

H-10. Descriptive Statistics for Impact Area Level Multipliers – LP2						
Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	1.5296	1.4886	0.1815	0.016	1.1789	2.0935
B	1.5333	1.5052	0.1656	0.0138	1.2778	2.1882
C	1.6465	1.6313	0.1652	0.0136	1.3764	2.2553
E	1.6869	1.6824	0.166	0.0138	1.3897	2.2496

H-11. Descriptive Statistics for Impact Area Level Multipliers – SE1						
Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	1.6651	1.6887	0.1535	0.0137	1.3521	1.9333
B	1.7363	1.7627	0.1262	0.0105	1.3792	2.0281
C	1.8044	1.8173	0.1271	0.0105	1.4443	2.032
E	1.8095	1.8345	0.1342	0.0111	1.448	2.0476

H-12. Descriptive Statistics for Impact Area Level Multipliers – SE2						
Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	2.07	2.1167	0.2539	0.0227	1.5777	2.6335
B	2.1665	2.1753	0.2156	0.018	1.6984	2.7162
C	2.3199	2.3223	0.2406	0.0198	1.8343	2.9118
E	2.3276	2.3392	0.2318	0.0192	1.7588	2.907

Appendix H (Cont.)

H-13. Descriptive Statistics for Impact Area Level Multipliers – SP1

Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	1.5945	1.5527	0.2069	0.0185	1.2314	2.5955
B	1.6648	1.6037	0.2127	0.0178	1.2824	2.9886
C	1.7189	1.696	0.1462	0.0121	1.3076	2.2151
E	1.7297	1.7071	0.1551	0.0128	1.3803	2.0308

H-14. Descriptive Statistics for Impact Area Level Multipliers – SP2

Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	1.8572	1.7952	0.2514	0.0225	1.3758	2.977
B	1.9476	1.8942	0.2669	0.0223	1.5184	3.5145
C	2.0676	2.0455	0.1898	0.0157	1.5774	2.6687
E	2.0925	2.0651	0.2077	0.0172	1.6683	2.4792

H-15. Descriptive Statistics for Impact Area Level Multipliers – VE1

Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	1.456	1.4447	0.1162	0.0181	1.2406	1.6551
B	1.5019	1.5177	0.0853	0.01	1.268	1.6526
C	1.5406	1.5443	0.0907	0.0089	1.3333	1.6892
E	1.5416	1.5482	0.093	0.0092	1.3119	1.6712

H-16 Descriptive Statistics for Impact Area Level Multipliers – VE2

Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	1.8661	1.8253	0.2223	0.0347	1.5506	2.2757
B	1.9099	1.8901	0.1828	0.0215	1.4734	2.2621
C	2.0121	2.0088	0.203	0.0199	1.5587	2.4034
E	2.0278	2.0303	0.2089	0.0206	1.5144	2.4437

Appendix H (Cont.)

H-17. Descriptive Statistics for Impact Area Level Multipliers – VP1

Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	1.3913	1.4042	0.0894	0.014	1.1479	1.5974
B	1.4387	1.4301	0.0974	0.0115	1.2217	1.7591
C	1.4858	1.4814	0.1009	0.0099	1.2079	1.7735
E	1.4874	1.4733	0.1094	0.0108	1.2422	1.9037

H-18. Descriptive Statistics for Impact Area Level Multipliers – VP2

Impact Level	Mean	Median	StDev	SE Mean	Minimum	Maximum
A	1.6327	1.622	0.122	0.0191	1.3272	1.8577
B	1.6812	1.6856	0.1299	0.0153	1.3296	2.0062
C	1.7881	1.7877	0.1421	0.0139	1.3847	2.1842
E	1.8056	1.806	0.1387	0.0137	1.4862	2.2922

MICHIGAN STATE UNIVERSITY LIBRARIES



3 1293 02334 0866