

DOKUZ EYLÜL UNIVERSITY
GRADUATE SCHOOL OF NATURAL AND APPLIED SCIENCES

STRUCTURAL EQUATION MODELING (SEM)
AND ITS APPLICATIONS

by
Çiğdem TATAR

January, 2015
İZMİR

STRUCTURAL EQUATION MODELING (SEM) AND ITS APPLICATIONS

**A Thesis Submitted to the
Graduate School of Natural and Applied Sciences of Dokuz Eylül University
In Partial Fulfillment of the Requirements for the Degree of Master of Science
in Statistics**

**by
Çiğdem TATAR**

**January, 2015
İZMİR**

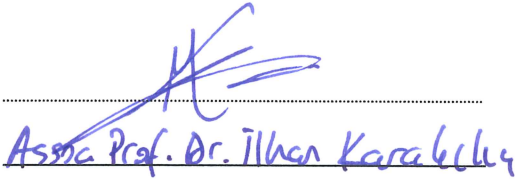
M.Sc THESIS EXAMINATION RESULT FORM

We have read the thesis entitled “STRUCTURAL EQUATION MODELING (SEM) AND ITS APPLICATIONS” completed by ÇİĞDEM TATAR under supervision of ASSOC. PROF. DR. ÖZLEM EGE ORUÇ and we certify that in our opinion it is fully adequate, in scope and in quality, as a thesis for the degree of Master of Science.

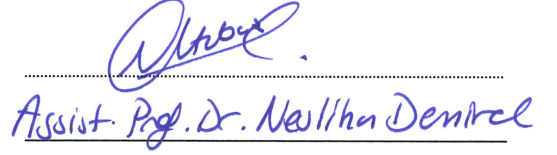


Assoc. Prof. Dr. Özlem EGE ORUÇ

Supervisor



(Jury Member)



(Jury Member)



Prof.Dr. Ayşe OKUR
Director
Graduate School of Natural and Applied Sciences

ACKNOWLEDGMENT

I would like to thank to everybody who has encouraged me on my study. First of all, I wish to express my gratitude to my supervisor Assoc. Prof. Özlem EGE ORUÇ for her guidance, advices, motivation and the support of my thesis.

I also thank all the Dokuz Eylül University academicians for their participation to my research.

I am grateful to Assoc. Prof. Emel KURUOĞLU, Asst. Prof. Süleyman ALPAYKUT, Asst. Prof. Senem VAHAPLAR and Yusuf YENİYAYLA for their contributions and valuable advises.

Finally, I appreciate to my mother and father for their patience trough my study and I would like to thank my dear friend Bora Özkaraoğlu. He has been my biggest source of moral support.

Çiğdem TATAR

STRUCTURAL EQUATION MODELING (SEM) AND ITS APPLICATIONS

ABSTRACT

Nowadays, banks get an opportunity of serving to their customers without the restriction of location and time, thanks to internet. But some customers still do not use it, because of different reasons. In this study, the factors that affect internet banking usage of Dokuz Eylül University's academicians are investigated in AMOS programming. As a consequence of the model analysis, there are significant relationships between the factors "Importance of Internet Banking Needs" and "Benefits of Internet Banking" and the factors " Importance of Internet Banking Needs" and "Communication". Furthermore, according to the structural equations, "Customer Adaptation to Internet Banking " is related to "Communication" and "Convenience".

Keywords: Structural Equation Modeling (SEM), AMOS, model fit indices, internet banking.

YAPISAL EŞİTLİK MODELLEME (YEM) VE UYGULAMALARI

ÖZ

Bankalar son dönemlerde, internet sayesinde pek çok hizmeti zaman ve mekan sınırlaması olmadan müşterilerine sunma imkanına kavuşmuştur. Fakat bazı banka müşterileri farklı nedenlerle internet bankacılığını hala kullanmamaktadır. Bu çalışmada Dokuz Eylül Üniversitesi'nde görevli akademisyenlerin internet bankacılığını kullanmasını etkileyen faktörler, AMOS programında yapısal eşitlik modeli ile araştırılmıştır. Modelin analizi sonucunda; "İnternet Bankacılığının Faydaları" ile "İnternet Bankacılığının Gereksiniminin Önemi" faktörleri arasında ve "İnternet Bankacılığının Gereksiniminin Önemi" ile "İletişim" faktörleri arasında anlamlı bir ilişki olduğu tespit edilmiştir. Ayrıca "İnternet Bankacılığına Müşteri Adaptasyonunun", "İletişim" ve "Uygunluk" ile bağlantılı olduğu kurulan yapısal eşitlik modelleriyle desteklenmektedir.

Anahtar kelimeler: Yapısal Eşitlik Modelleri (YEM), AMOS, model uyum kriterleri, internet bankacılığı.

CONTENTS

	Page
THESIS EXAMINATION RESULT FORM	iii
ACKNOWLEDGMENT.....	iii
ABSTRACT.....	iiiv
ÖZ	v
LIST OF FIGURES	viii
LIST OF TABLES	ix
 CHAPTER ONE- INTRODUCTION	 1
 CHAPTER TWO-STRUCTURAL EQUATION MODELING	 3
2.1 Measurement Model	4
2.2 Structural Model (Path Model).....	5
2.2.1 Path Analysis.....	6
2.3 Assumptions of Structural Equation Modeling	8
2.4 Steps of Structural Equation Modeling.....	9
2.4.1 Model Specification.....	10
2.4.2 Model Identification	10
2.4.3 Model Estimation	11
2.4.4 Model Evaluating	11
2.4.5 Model Modification.....	12
2.5. Software Programs for SEM	12
2.5.1 Analysis of Moment Structure (AMOS) Program.....	12
 CHAPTER THREE- EVALUATION OF THE MODEL.....	 16
3.1 Absolute Fit Indices	17
3.1.1 Standardized Root Mean Square Residual (SRMR)	17
3.1.2 Goodness of Fit Index (GFI)	18
3.1.3 Adjusted Goodness of Fit Index (AGFI).....	18

3.1.4 Root Mean Square Error of Approximation (RMSEA)	19
3.2 Incremental Fit Indices	19
3.2.1 Normed Fit Index (NFI)	20
3.2.2 Non-Normed Fit Index (NNFI) or Tucker and Lewis Index (TLI)	20
3.2.3 Comparative Fit Index (CFI)	21
3.3 Fit Indexes Based On Information Criterion	21
3.3.1 Akaike Information Criterion (AIC)	21
3.3.2 Bayes Information Criterion (BIC)	22
CHAPTER FOUR- APPLICATION	24
4.1 Survey Reliability and Validity Analysis	25
4.2 Factor Analysis	28
4.3 Construct SEM Model Using AMOS	32
CHAPTER FIVE- CONCLUSION	38
REFERENCES	389

LIST OF FIGURES

	Page
Figure 2.1 Path diagram for CFA.....	4
Figure 2.2 Measurement and structural models	5
Figure 2.3 Path diagram for random variables X,Y,Z and W	6
Figure 2.4 Structural equation model path diagram -relationship between textile workers and union sentiment.	7
Figure 2.5 AMOS program menu	13
Figure 2.6 View of the menu bar	14
Figure 4.1 Initial model for the analysis	32
Figure 4.2 Residual covariance matrix	36

LIST OF TABLES

	Page
Table 2.1 Common path diagram symbols.	7
Table 2.2 The grouping of sample size for SEM	9
Table 2.3 Tasks of the buttons in AMOS menu bar.....	14
Table 3.1 Fit indices and their acceptable thresholds	22
Table 4.1 Descriptive Statistics for each survey question score	24
Table 4.2 Descriptive Statistics for all survey question score	25
Table 4.3 The internal consistency analysis.....	26
Table 4.4 Reliability statistics of the survey	27
Table 4.5 ANOVA table and Tukey's test for nonadditivity.....	27
Table 4.6 Hotelling's t-squared test.....	28
Table 4.7 Rotated component matrix for factor analysis	29
Table 4.8 Final survey form after FA.....	29
Table 4.9 Cronbach's Alpha for internal consistency checking	31
Table 4.10 Fit indices for the initial model	33
Table 4.11 Information criteria for the initial model	33
Table 4.12 Fit indices for the last model.....	35
Table 4.13 Information criteria for the last model	35
Table 4.14 Standardized regression weights.....	36
Table 4.15 Results of the hypothesis testing.....	37

CHAPTER ONE

INTRODUCTION

Structural equation modeling (SEM), which is a tool for analyzing multivariate data, has gained popularity across many disciplines in the past two decades. These models go beyond ordinary regression models to incorporate multiple independent and dependent variables as well as hypothetical latent constructs that clusters of observed variables might represent. They also provide a way to test the specified set of relationships among observed and latent variables as a whole, and allow theory testing even when experiments are not possible. SEM has a number of synonyms and special cases in the literature including path analysis, causal modeling, and covariance structure analysis. Statistically, it represents an extension of general linear modeling (GLM) procedures. In simple terms, SEM involves the evaluation of two models: a measurement model and a path model. There are several studies based on SEM in the literature. For instance, Resinger and Mavando (2006) used these techniques to determine the tourist satisfaction. SEM has been recommended for environmental research such as plant conservation by Iriando, Albert and Escurado (2013). Many other uses of SEM have been studied in the literature for example McQuitty (2003), Ustasüleyman and Eyüboğlu (2010), Bollen and Noble (2011) and Demir (2011).

Nowadays, the advancement of internet has provided an opportunity for banking institution in introducing new financial innovations. One of the emerging financial innovations introduced by banking institution is internet banking. Internet banking changed the delivery channels used by the financial services industry. Banks can benefit from much lower operating costs by offering internet banking services. Customers will also benefit from the convenience, speed and round-the-clock availability of internet banking services. However, despite the fact that internet banking provides many advantages there are still a many of customers who does not use such services. Therefore, understanding the reasons for this resistance would be useful for bank managers in formulating strategies aimed at increasing internet banking use. The aim of this study is to investigate the factors influencing the usage of internet banking services for using SEM in AMOS programming.

This study contains five chapters. In chapter I and II, we introduce some common terminologies, general steps and assumptions of SEM. In Chapter III, several popular specialized SEM software programs are briefly discussed with respect to their features and availability. In the last chapter, an application of SEM to the internet banking usage among the academicians at Dokuz Eylül University was conducted. Moreover, particular interpretations are made using the SEM models of the application.

CHAPTER TWO

STRUCTURAL EQUATION MODELING

Structural Equation Modeling (SEM) has gained popularity across many disciplines due perhaps to its generality and flexibility. It is a collection of statistical techniques that allow a set of relationships between one or more independent variables and one or more dependent variables to be examined. Both independent and dependent variables can be either continuous or discrete and can be either factors or measured variables. SEM is a combination of multiple regression analysis, confirmatory factor analysis and path analysis. These analyses make it easy to understand basic SEM models.

The development of SEM methods and software has proceeded rapidly since the 1970s. A historical review of the development of SEM was provided by Bentler (1986), and annotated bibliographies of the technical SEM literature have been provided by Austin & Wolfle (1991) and Austin & Caldero'n (1996). Methodological developments also have been reviewed (Bentler 1980, Bentler & Dudgeon 1996). Software development, with respect to both implementation of methodological advances as well as improvement of user interfaces, has made this sophisticated and powerful method readily accessible to substantive researchers.

SEM enables complex model to be statistically modeled and tested. Additionally it is a method for confirming the theoretical models. On the contrary, basic statistical methods only use a limited number of variables, which are not able to deal with the experienced theories being developed (Schumacker & Lomax, 2010). That is why SEM is reason for preference. In recent years it is generally used in many areas such as sociology, psychology, economics, management, environmental studies, tourism research and marketing.

SEM involves the evaluation of two models: a *measurement* model and a structural model (*path* model). They are described below.

2.1 Measurement Model

Measurement model, which define how implicit variables or theoretical structures rely upon observed variables. Latent variables are variables that are not directly observed but are rather inferred from other variables that are observed. Statistical techniques, such as factor analysis, exploratory or confirmatory, have been widely used to examine the number of latent constructs underlying the observed responses and to evaluate the adequacy of individual items or variables as indicators for the latent constructs they are supposed to measure.

Confirmatory factor analysis (CFA) models, a special case of SEM, are widely used in measurement applications for a variety of purposes. Designs for construct validation and scale refinement, multitrait - multimethod validation, and measurement invariance can be evaluated through specification and testing of CFA models. Let $X_1, X_2, X_3, X_4, X_5, X_6$ are observed variables, F_1, F_2, F_3 are latent variables and $e_1, e_2, e_3, e_4, e_5, e_6$ are the measurement errors. Figure 2.1 shows the path diagram for CFA.

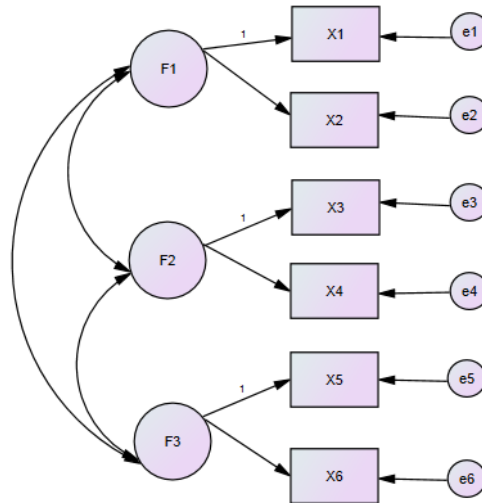


Figure 2.1 Path diagram for CFA

CFA differs from Exploratory Factor Analysis (EFA) in that factor structures are hypothesized a priori and verified empirically rather than derived from the data. EFA

often allows all indicators to load on all factors and does not permit correlated residuals. In contrast, the number of factors in CFA is assumed to be known. In SEM, these factors correspond to the latent constructs represented in the model. CFA allows an indicator to load on multiple factors. It also allows residuals or errors to correlate. Once the measurement model has been specified, structural relations of the latent factors are then modeled essentially the same way as they are in path models. The combination of CFA models with structural path models on the latent constructs represents the general SEM framework in analyzing covariance structures.

2.2 Structural Model (Path Model)

Structural model is an extension of multiple regression in that it involves various multiple regression models or equations that are estimated simultaneously. It can be considered a special case of SEM in which structural relations among observed (vs. latent) variables are modeled.

Let X_1, X_2, X_3, X_4 are observed variables, Y_1, Y_2 are latent variables and e_1, e_2, e_3, e_4 are the measurement errors. A simple structural model and measurement model graphically illustrated in Figure 2.2 (Bayram, 2010).

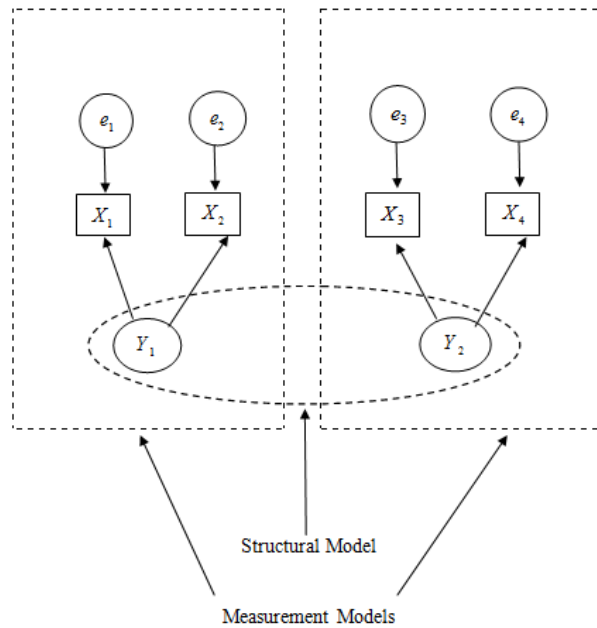


Figure 2.2 Measurement and structural models

In Figure 2.2, there are two measurement models and one structural model.

2.2.1 Path Analysis

Path analysis provides to easily perceive casual links making up the complex systems, to be able to decide under what conditions the variables in causality are one another's cause or effect and to be able to explain this causal connection in mathematical terms has always posed a significant question. In path analysis, one or more multiple regression analyses are performed depending on the variable relationships specified in the path model (Schumacker & Lomax, 2010). A **path diagram** is a pictorial representation of a model. Let X, Y, Z and W be observed variables. Figure 2.3 shows the path diagram for these random variables.

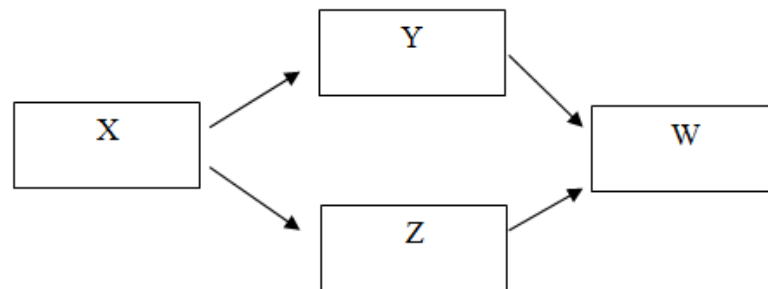


Figure 2.3 Path diagram for random variables X, Y, Z and W

Path diagrams are fundamental to SEM because they allow the researcher to diagram the hypothesized set of relations the model. The diagrams are helpful in clarifying a researcher's ideas about the relations among variables. There is a one to one correspondence between the diagrams and the equations needed for the analysis. Path analysis is not actually a method for discovering causes; it tests theoretical relationships which has been termed causal modeling (Schumacker & Lomax, 2010). Several symbols are used in developing SEM diagrams. The common path diagram symbols are given in Table 2.1.

Table 2.1 Common path diagram symbols.

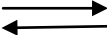
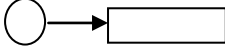
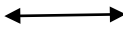
Symbol	Meaning
	Latent Construct, Factor, Unmeasured Variable
	Measured Variable, Observed Variable
	Direct Relationship
	Error Associated With Measured Variable
	Covariance or Correlation

Figure 2.4 shows an example for path analysis. The model consists of five observed variables; the independent variables are the number of years worked in the textile mill and worker age (age); the dependent variables are deference to managers (deference), support for labor activism (support), and sentiment toward unions (sentiment) (McDonald & Clelland ,1984).

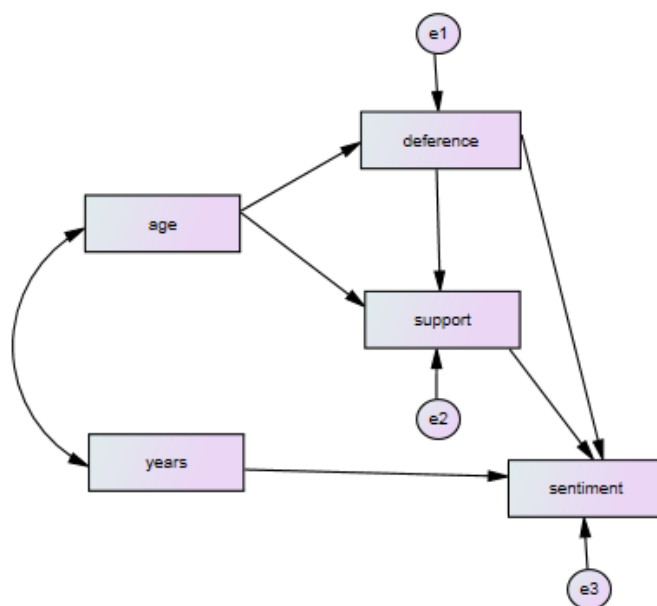


Figure 2.4 Structural equation model path diagram -relationship between textile workers and union sentiment.

The structural equations model for this variables is given below.

$$deference = age + e_1$$

$$support = age + deference + e_2$$

$$sentiment = years + support + deference + e_3$$

2.3 Assumptions of Structural Equation Modeling

As with all statistical methodologies, structural equation modeling requires that certain underlying assumptions be satisfied to ensure accurate inferences. These assumptions pertain to the intersection of the data and the estimation method. This section considers the major assumptions associated with structural equation modeling.

- A basic assumption underlying the standard use of structural equation modeling is that the observations are drawn from a continuous and multivariate normal population. This assumption is particularly important for maximum likelihood (ML) estimation. If the assumption of multivariate normality of the observed data holds, the estimates obtained from maximum likelihood estimation possess certain desirable properties. Specifically, maximum likelihood estimation will yield unbiased and efficient estimates as well as a large sample goodness of-fit test. Departures from normality cause the chi-square test to be larger than it should be and standard errors to be smaller than they should be. If n is very large, do not need to worry about non-normality. Amemiya and Anderson (1988; 1990) show asymptotic robustness of the normality assumption for linear factor analysis model.
- Since SEM is a confirmatory technique, you must specify a full model *a priori* and test that model based on the sample and variables included in experimental measurements.
- Sample size is important for SEM. Because of it has the ability to model complex relationships between multivariate data. A larger sample size is always desired for

SEM. There are some ideas for grouping sample size. These ideas are given in Table 2.2

Table 2.2 The grouping of sample size for SEM

$n < 100$	small
$100 < n < 200$	medium
$n > 200$	large

- SEM techniques only look at first-order (linear) relationships between variables. Linear relationships can be explored by creating bivariate scatter plots for all of variables. Power transformations can be made if a relationship between two variables seems quadratic.
- In SEM, it is assumed that there is no relationship between independent variables. If one of the variables is a perfect linear combination of the other variables, a singular matrix will cause the analysis to crash. Multicollinearity can also be a problem.
- The residuals of the covariances need to be small and centered about zero. Some goodness of fit tests remain robust against highly deviated residuals or non-normal residuals.

2.4 Steps of Structural Equation Modeling

The aims of SEM are to understand the patterns of correlation/covariance among a set of variables and to explain as much of their variance as possible with the model specified. There are five basic steps in SEM analysis which are model specification, model identification, model estimation, model evaluation and model modification. Issues pertaining to each of these steps are discussed below.

2.4.1 Model Specification

The first step of SEM is to specify the model. Model specification refers to the mathematical specification of the relationships that are in the model. Because SEM is a confirmatory technique, it bases on the analysis type that confirms model. Therefore model should be specified correctly. The goal of the researcher is to determine the best possible model that generates the sample covariance matrix (Schumacker & Lomax, 2010). The sample covariance matrix should be the closest matrix to the covariance structure.

2.4.2 Model Identification

The most difficult part of SEM is correctly *identifying* the model. **Identification** involves the study of conditions to obtain a single, unique solution for each and every free parameter specified in the model from the observed data.

Types of model identification

- If a single value cannot be obtained from the observed data for one or more free parameters, then the model is ***underidentified***. In this case the parameters cannot be estimated, and there searcher needs to reduce the number of parameters to be estimated by deleting or fixing some of them.
- If for each free parameter a value can be obtained through one and only one manipulation of the observed data, then the model is ***just identified***. Such a model will fit the data perfectly, and thus is of little use, although it can be used to estimate the values of the coefficients for the paths.
- If a value for one or more free parameters can be obtained in multiple ways from the observed data, then the model is ***overidentified***. Overidentification is the condition in which there are more equations than unknown independent parameters.

There are two basic requirements for the identification of SEM:

- Degrees of freedom of the model should be greater than zero,

$$(df = \frac{p(p+1)}{2} - q, \text{ where } p \text{ is the number of the observed variables and } q \text{ is the number of parameters that will be estimated.})$$
- Every latent variable must be assigned a scale (Kline,2005).

2.4.3 Model Estimation

A properly specified structural equation model often has some fixed parameters and some free parameters to be estimated from the data. The problem of parameter estimation concerns obtaining estimates of the free parameters of the model, subject to the constraints imposed by the fixed parameters of the model. In other words, the goal is to apply a model to the sample covariance matrix $\mathbf{S}$, which is an estimate of the population covariance matrix Σ that is assumed to follow a structural equation model. Estimates are obtained such that the discrepancy between the estimated covariance matrix $\hat{\Sigma} = \Sigma(\hat{\theta})$ and the sample covariance matrix $\mathbf{S}$, is as small as possible. If this difference is equal to zero, then there is a perfect model fit.

There are a lot of estimation methods in SEM. Unweighted Least Squares (ULS), Maximum Likelihood Estimation (MLE), Generalized Least Squares (GLS), Generally Weighted Least Squares (WLS) are some of them. The most common method is MLE which is scale free. Both GLS and MLE have desirable asymptotic properties, which are unbiasedness, minimum variance and large sample (Schumacker & Lomax, 2010). If the assumption of multivariate normal distribution is not met, then the MLE method is not appropriate. In this case, robust ML is the best option, for the usage of AMOS program (Yuan, Bentler & Zhang,2005). Also, WLS is not useful for model that contain many variables ($p > 20$) or the sample size is smaller than 1000. (Iriando, Albert & Escudero, 2003).

2.4.4 Model Evaluating

Once the parameters of the model have been estimated, the model can be tested for global fit to the data. This is essentially a statistical hypothesis-testing problem,

with the null hypothesis being that the model under consideration fits the data. The overall model goodness of fit is reflected by the magnitude of discrepancy between the sample covariance matrix and the covariance matrix implied by the model with the parameter. Most measures of overall model goodness of fit are functionally related to $F_{\min}$. The model test statistic $(n-1)F_{\min}$, where n is the sample size, has a chi-square distribution when the model is correctly specified and can be used to test the null hypothesis that the model fits the data.

2.4.5 Model Modification

If the fit model is not adequate, it has become common practice to modify the model, by deleting parameters that are not significant, and adding parameters that improve the fit. Modifications of the model are often aided by modification indices. Modification index estimates the magnitude of decrease in model chi-square (for 1 degrees of freedom) whereas expected parameter change approximates the expected size of change in the parameter estimate when a certain fixed parameter is freely estimated.

2.5 Software Programs for SEM

Most SEM analyses are conducted using one of the specialized SEM software programs. The commonly used programs for SEM are AMOS (Analysis of Moment Structure), CALIS (Covariance Analysis And Linear Structural Equations), EQS (Equations), LISREL (Linear Structural Relationships), and SEPATH (Structural Equation Modeling And Path Analysis). Special features of AMOS are discussed in this section.

2.5.1 Analysis of Moment Structure (AMOS) Program

AMOS is statistical software and it stands for analysis of a moment structures. It implements the general approach to data analysis known as structural equation modeling (SEM), also known as analysis of covariance structures, or causal

modeling. This approach includes, as special cases, many well-known conventional techniques, including the general linear model and common factor analysis. AMOS is a visual program for structural equation modeling. In AMOS, we can draw models graphically using simple drawing tools. AMOS quickly performs the computations for SEM and displays the results. SPSS Amos is the perfect tool for a variety of purposes .

Amos provides the following methods for estimating structural equation models:

- a. Maximum likelihood
- b. Unweighted least squares
- c. Generalized least squares
- d. Browne's asymptotically distribution-free criterion
- e. Scale-free least squares

With AMOS, you can quickly specify, view, and modify your model graphically using simple drawing tools. Then you can assess your model's fit, make any modifications, and print out a publication-quality graphic of your final model. Simply specify the model graphically AMOS quickly performs the computations and displays the results.

Figure 2.5 shows the AMOS program menu. In the figure, your work area appears. The large area on the right is where you draw path diagrams. The toolbar on the left provides one-click access to the most frequently used buttons. You can use either the toolbar or menu commands for most operations.

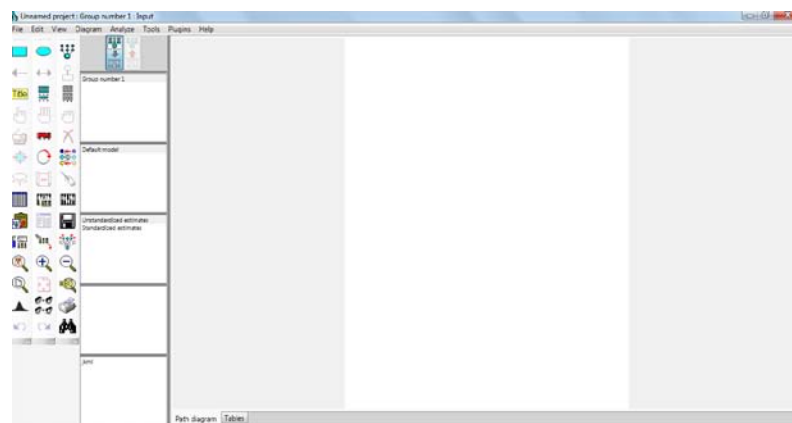


Figure 2.5 AMOS program menu

Figure 2.6 and Table 2.3 demonstrate the view of the menu bar and the tasks of the buttons respectively.



Figure 2.6 View of the menu bar

Table 2.3 Tasks of the buttons in AMOS menu bar

Buttons	Tasks	Buttons	Tasks
	Draws observed variables		Specifies the analyze properties
	Draws latent variables		Calculates the estimates
	Draws latent variables or adds observed variables		Copies the path diagram to the clipboard
	Connects endogenous variable to exogenous variable		Views the output in text form

Table 2.3 Tasks of the buttons in AMOS menu bar (cont.)

	Draws covariances		Saves the current path diagram
	Adds error term to observed variable		Identifies the object properties
	Adds an explanation		Aligns the objects
	Lists the variables in the model		Preserves symmetries
	Lists the variables in the data set		Zooms in a selected area
	Selects one object		Zooms in
	Selects all of the objects		Zooms out
	Unselects the objects		Shows the whole page
	Copies the objects		Reshapes according to the page
	Moves the objects		Zooms in to the path diagram
	Deletes the objects		Does the multiple-group analysis
	Changes the form of the objects		Prints out of the path diagram
	Changes the direction of observed variable		Does the specification research
	Selects the data file and reads it		Changes the observed variable to reverse direction
	Moves the parameter values		Rearranges the arrows in the path diagram

CHAPTER THREE

EVALUATION OF THE MODEL

Fit refers to the ability of a model to reproduce the data. A good-fitting model is one that is reasonably consistent with the data and so does not necessarily require respecification. Also a good-fitting measurement model is required before interpreting the causal paths of the structural model. Fit indices evaluate model fit for the data being examined. They help the researcher determine which proposed model best fit the data by showing how well the parameter estimates account for the observed covariances.

The statistical tests are used for model fit in the analysis of structural equation models. These tests for evaluating structural equation models with unobservable variables are a t-test and a chi-square test. The t-test (the ratio of the parameter estimate to its estimated standard error) indicates whether individual parameter estimates are statistically different from zero. The chi-square statistic compares the "goodness of fit" between the covariance matrix for the observed data and covariance matrix derived from a theoretically specified structure (model).

Most measures of overall model goodness of fit are functionally related to the minimum of the fit function or $F_{\min}$. The model test statistic $(n-1)F_{\min}$, where n is the sample size, has a chi-square distribution.

$$\chi^2 = (n-1)F_{\min} \quad (3.1)$$

When the model is correctly specified and can be used to test the null hypothesis that the model fits the data.

Statistical tests for model fit have the problem that their power varies with the sample size. In reaction to sample size sensitivity problem, a variety of alternative goodness-of-fit indices have been developed to supplement the chi-square statistic. All of these alternative indices attempt to adjust for the effect of sample size, and

many of them also take into account model degrees of freedom, which is a proxy for model size.

Two classes of alternative fit indices, incremental and absolute, have been identified. The most commonly used absolute fit indices are RMSEA (Root-Mean-Square Error Approximation), GFI (Goodness-of-Fit Indices) and AGFI (Adjusted Goodness-of-Fit Indices) (Jöreskog & Sörbom, 1979), SRMR (Standardized Root Mean Square Residual) (Bentler, 1995). Higher values of GFI and AGFI as well as lower values of SRMR and RMSEA indicate better model-data fit. Examples of incremental fit indices include Normed Fit Index (NFI; Bentler & Bonett, 1980), Tucker-Lewis Index (TLI; Tucker & Lewis, 1973), Relative Noncentrality Index (RNI; McDonald & Marsh, 1990), and Comparative Fit Index (CFI; Bentler 1990). Higher values of incremental fit indices indicate larger improvement over the baseline model in fit. Values in the .90s (or more recently $\geq .95$) are generally accepted as indications of good fit.

3.1 Absolute Fit Indices

Absolute fit indices measure the extent to which the specified model of interest reproduces the sample covariance matrix. These indices measure presumes that the best fitting model has a fit of zero. The measure of fit then determines how far the model is from perfect fit. The most commonly used absolute fit indices RMSEA, GFI, AGFI and SRMR are defined in this section.

3.1.1 Standardized Root Mean Square Residual (SRMR)

The SRMR is an absolute measure of fit and is defined as the standardized difference between the observed correlation and the predicted correlation. Unlike incremental fit indices, their calculation does not rely on comparison with a baseline model but is instead a measure of how well the model fits in comparison to no model at all (Jöreskog & Sörbom, 1993). It is a positively biased measure and that bias is greater for small n and for low degrees of freedom (df) studies. Because the SRMR is an absolute measure of fit, a value of zero indicates perfect fit. The SRMR has no

penalty for model complexity. A value less than 0.08 is generally considered a good fit (Hu & Bentler, 1999).

3.1.2 Goodness of Fit Index (GFI)

The goodness-of-fit statistic (GFI) was created by Jöreskog and Sorbom as an alternative to the Chi-Square test and calculates the proportion of variance that is accounted for by the estimated population covariance. The formula of GFI index is given in Equation 3.2.

$$GFI = 1 - \left(\frac{\chi_m^2}{\chi_i^2} \right) \quad (3.2)$$

where

χ_m^2 : χ^2 test statistic of the model,

χ_i^2 : χ^2 test statistic of the independence model.

GFI index is affected by sample size. It ranges from 0 to 1 with larger samples increasing its value. When there are a large number of degrees of freedom in comparison to sample size, the GFI has a downward bias. In addition, it has also been found that the GFI increases as the number of parameters increases and also has an upward bias with large samples. Generally in literature cut-off point of 0.90 has been recommended for the GFI. Given the sensitivity of this index, it has become less popular in recent years and it has even been recommended that this index should not be used.

3.1.3 Adjusted Goodness of Fit Index (AGFI)

Related to the GFI is the AGFI which adjusts the GFI based upon degrees of freedom, with more saturated models reducing fit. The formula of the AGFI is given in Equation 3.3.

$$AGFI = 1 - \frac{k(k+1)}{2df} (1 - GFI) \quad (3.3)$$

where k is the number of observed variables.

AGFI tends to increase with sample size. As with the GFI, values for the AGFI also range between 0 and 1 and it is generally accepted that values of 0.90 or greater indicate well fitting models. Given the often detrimental effect of sample size on GFI and AGFI indices they are not relied upon as a standalone index, however given their historical importance they are often reported in covariance structure analyses.

3.1.4 Root Mean Square Error of Approximation (RMSEA)

This absolute measure of fit is based on the non-centrality parameter. It tells us how well the model, with unknown but optimally chosen parameter estimates would fit the populations covariance matrix. In recent years it has become regarded as ‘one of the most informative fit indices’ because of its sensitivity to the number of estimated parameters in the model. Its computational formula is:

$$RMSEA = \sqrt{\frac{(\chi^2 / df - 1)}{n - 1}} \quad (3.4)$$

where n the sample size and df the degrees of freedom of the model. If χ^2 is less than df , then the RMSEA is set to zero. Recommendations for RMSEA cut-off points have been reduced considerably in the last fifteen years. More recently, a cut-off value close to .06 or a stringent upper limit of 0.07 seems to be the general consensus amongst authorities in this area. One of the greatest advantages of the RMSEA is its ability for a confidence interval to be calculated around its value. It is generally reported in conjunction with the RMSEA and in a well-fitting model the lower limit is close to 0 while the upper limit should be less than 0.08.

3.2 Incremental Fit Indices

Incremental fit indices, also known as comparative or relative fit indices, are a group of indices that do not use the chi-square in its raw form but compare the chi-square value to a baseline model. The common used incremental fit indices, Normed Fit Index, Non-Normed Fit Index and Comparative Fit Index are defined in this subsection.

3.2.1 Normed Fit Index (NFI)

NFI is proposed by Bentler and Bonett in 1980. This statistic assesses the model by comparing the χ^2 value of the model to the χ^2 of the null model. The best model is defined as model with a χ_m^2 of zero and the worst model by the χ_i^2 of the null model. It is calculated :

$$NFI = \frac{\chi_i^2 - \chi_m^2}{\chi_i^2} \quad (3.5)$$

Values for this statistic range between 0 and 1. Bentler and Bonnet (1980) recommends values greater than 0.90 indicating a good fit. More recent suggestions state that the cut-off criteria should be $NFI \geq .95$. A major disadvantage of this index is sensitive to sample size. It cannot be smaller if more parameters are added to the model. Its “penalty” for complexity is zero. Thus, the more parameters added to the model, the larger the index. Because of this reason it is not recommended to be solely relied on (Kline, 2005).

3.2.2 Non-Normed Fit Index (NNFI) or Tucker and Lewis Index (TLI)

Let χ^2 / df be the ratio of chi square to its degrees of freedom, and the TLI is computed as follows:

$$NNFI = \frac{\frac{\chi_i^2}{df_i} - \frac{\chi_m^2}{df_m}}{\frac{\chi_i^2}{df_i} - 1} \quad (3.6)$$

A major disadvantage of NFI measure is that it cannot be smaller if more parameters are added to the model. Its “penalty” for complexity is zero. Thus, the more parameters added to the model, the larger the index. The Tucker-Lewis index, another incremental fit index, does have such a penalty.

If the index is greater than one, it is set at one. It is interpreted as the NFI. Note that for a given model, a lower chi square to df ratio implies a better fitting model. Its penalty for complexity is χ^2/df . That is, if the chi square to df ratio does not change, the TLI does not change. TLI depends on the average size of the correlations in the data. If the average correlation between variables is not high, then the TLI will not be very high.

3.2.3 Comparative Fit Index (CFI)

CFI is directly based on the non-centrality measure. It is a revised form of the NFI which takes into account sample size that performs well even when sample size is small. The Comparative Fit Index (CFI) equals

$$CFI = 1 - \frac{\chi_m^2 - df_m}{\chi_i^2 - df_i} \quad (3.7)$$

If the index is greater than one, it is set at one and if less than zero, it is set to zero. As with the NFI, values for this statistic range between 0.0 and 1.0 with values closer to 1.0 indicating good fit. A cut-off criterion of CFI ≥ 0.90 is initially advanced however, recent studies have shown that a value greater than 0.90 is needed in order to ensure that misspecified models are not accepted. From this, a value of $CFI \geq 0.95$ is presently recognized as indicative of good fit.

3.3 Fit Indexes Based On Information Criterion

3.3.1 Akaike Information Criterion (AIC)

The AIC is a comparative measure of fit and so it is meaningful only when two different models are estimated. Small values indicate a better fit and so the model with the smallest AIC is the best fitting model. The formula of AIC is given in Equation 3.8.

$$AIC = \chi^2 + k(k-1) - 2df \quad (3.8)$$

where k is the number of variables in the model and df is the degrees of freedom of the model.

One advantage of the AIC, BIC, and CAIC measures is that they can be computed for models with zero degrees of freedom, i.e., just-identified models.

3.3.2 Bayes Information Criterion (BIC)

The other comparative fit indices is the BIC. Whereas the AIC has a penalty of 2 for every parameter estimated, the BIC increases the penalty as sample size increases. It is computed by following equation.

$$BIC = \chi^2 + [k(k+1)/2 - df][\ln(n)] \quad (3.9)$$

where $\ln(n)$ is the natural logarithm of the number of cases in the sample. The BIC places a high value on parsimony.

One advantage of the AIC and BIC measures is that they can be computed for models with zero degrees of freedom.

The summary of fit indices and their acceptable thresholds is given in Table 3.1.

Table 3.1 Fit indices and their acceptable thresholds

Goodness of Fit Criterion	Good Fit	Acceptable Fit	Interpretation
χ^2	$0 \leq \chi^2 \leq 2df$	$2df \leq \chi^2 \leq 3df$	Nonsignificant and small values shows good fit, significant and large values show poor fit.
χ^2 / df	$0 \leq \chi^2 / df \leq 2df$	$2 \leq \chi^2 / df \leq 5$	Value close to 2 reflects good fit.
GFI	$0.95 \leq GFI \leq 1$	$0.90 \leq GFI \leq 0.95$	GFI shows the amount of variances explained by the model.
AGFI	$0.90 \leq AGFI \leq 1$	$0.85 \leq AGFI \leq 0.90$	Values greater than 0.90 shows good fit.

Table 3.1 Fit indices and their acceptable thresholds (cont.)

SRMR	$0 \leq SRMR \leq 0.05$	$0.05 \leq SRMR \leq 0.1$	Values close to 0 shows good fit.
NNFI	$0.97 \leq NNFI \leq 1$	$0.95 \leq NNFI \leq 0.97$	Values below 0.90 need to respecify model.
NFI	$0.95 \leq NFI \leq 1$	$0.90 \leq NFI \leq 0.95$	Values below 0.90 need to respecify model.
RMSEA	$0 \leq RMSEA \leq 0.05$	$0.05 \leq RMSEA \leq 0.08$	RMSEA estimates how well the fitted model approximates the population covariance matrix per degree of freedom.
CFI	$0.97 \leq CFI \leq 1$	$0.95 \leq CFI \leq 0.97$	Values close to 1 shows good fit.
AIC CAIC BIC	-	-	Hypothesized model values should be smaller than independent and saturated model.

CHAPTER FOUR

APPLICATION

Banks find an opportunity to provide several services without time and place limitation by means of internet one of the most important technological developments. But some customers still do not use it, because of different reasons. This study adopts a model to investigate factors that determine an individual's intention to use online banking by bank customers among academicians in Dokuz Eylül University campuses. By using simple random sampling, two hundred fifty two academicians from Dokuz Eylül University will took part in the survey. Data was collected via survey which is developed by Jahangir and Parvez (2012). Scaling in this survey was examined under five subscales titles and the subscales were named "Importance of Internet Banking Needs", "Compatibility", "Convenience", "Communication" and "Customer Adaptation to Internet Banking". Survey consisted of 20 questions. In the first part, there was a section where demographical features were asked. This part was excluded from the 20 questions. Moreover, each question was evaluated with 1 to 5 scores in such a way that it would be one of the scales of "never, sometimes, neutral, frequently and always". The descriptive statistics for survey question score computed separately and totally are given in Table 4.1 and Table 4.2, respectively.

Table 4.1 Descriptive Statistics for each survey question score

Item Statistics			
Questions	Mean	Standard Deviation	N
Question 1	4.50	0.572	30
Question 2	4.60	0.814	30
Question 3	4.33	0.711	30
Question 4	4.33	0.758	30
Question 5	4.13	0.900	30
Question 6	4.10	0.759	30
Question 7	4.50	0.820	30
Question 8	4.37	0.615	30
Question 9	4.70	0.750	30
Question 10	4.47	0.776	30

Table 4.1 Descriptive Statistics for each survey question score (cont.)

Question 11	4.47	0.681	30
Question 12	4.47	1.074	30
Question 13	4.27	1.081	30
Question 14	3.90	0.885	30
Question 15	3.77	0.898	30
Question 16	3.87	0.860	30
Question 17	4.70	0.702	30
Question 18	4.67	0.758	30
Question 19	4.20	0.847	30
Question 20	2.47	1.074	30

Table 4.2 Descriptive Statistics for all survey question score

Summary Item Statistics							
	Mean	Minimum	Maximum	Range	Max/Min	Variance	N of items
Item means	4.240	2.467	4.700	2.233	1.905	0.248	20
Item variances	0.686	0.328	1.168	0.840	3.565	0.059	20
Inter-item correlations	0.352	-0.211	0.874	1.085	-4.146	0.035	20

4.1 Survey Reliability and Validity Analysis

After a survey is given, there are various statistical ways to assess the reliability of a survey. This occurs by assessing the reliability of the answers to all the questions collectively and the answers to each question on a survey. The most common method of assessing reliability is through what is called internal consistency. The approach analyzes the correlation or relationship of the answers to the questions on a survey with one another. The analysis gives a reliability score generally ranging from a negative number to 1.0, with a score closer to 1.0 suggesting that the survey does produce reliable answers. An acceptable survey reliability score is 0.70, and a

strong survey reliability score is 0.80 and above. The internal consistency analysis of our survey is given Table 4.3.

Table 4.3 The internal consistency analysis

Item-Total Statistics				
Questions	Scale Mean If Item Deleted	Scale Variance If Item Deleted	Corrected Item-Total Correlation	Cronbach's Alpha If Item Deleted
Question 1	80.30	93.183	0.596	0.904
Question 2	80.20	89.131	0.671	0.901
Question 3	80.47	89.844	0.723	0.900
Question 4	80.47	89.154	0.724	0.900
Question 5	80.67	87.678	0.689	0.900
Question 6	80.70	90.700	0.610	0.903
Question 7	80.30	93.872	0.349	0.909
Question 8	80.43	93.289	0.541	0.905
Question 9	80.10	91.266	0.577	0.903
Question 10	80.33	89.195	0.703	0.900
Question 11	80.33	90.575	0.698	0.901
Question 12	80.33	93.609	0.257	0.914
Question 13	80.53	84.740	0.713	0.899
Question 14	80.90	91.197	0.480	0.906
Question 15	81.03	90.654	0.505	0.905
Question 16	80.93	88.064	0.699	0.900
Question 17	80.10	92.852	0.499	0.905
Question 18	80.13	92.257	0.499	0.905
Question 19	80.60	92.041	0.452	0.906
Question 20	82.33	92.161	0.329	0.912

One of the most popular reliability statistics in use today is Cronbach's alpha. It determines the internal consistency or average correlation of items in a survey instrument to gauge its reliability. An acceptable Cronbach's alpha score is 0.70, The reliability statistics of our survey is given in Table 4.4.

Table 4.4 Reliability statistics of the survey

Reliability Statistics		
Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
0.908	0.916	20

Table 4.4 shows that the Cronbach's Alpha coefficient is 0.908. It is demonstrated that survey is highly reliable.

Table 4.5 ANOVA table and Tukey's test for nonadditivity

ANOVA with Tukey's Test for Nonadditivity							
			Sum of Squares	<i>df</i>	Mean Square	F	Sig
Between People			145.140	29	5.005		
Between Items			141.440	19	7.444	16.221	0.000
Nonadditivity			0.122	1	0.122	0.266	0.606
Within	Residual		252.738	550	0.460		
People	Balance		252.860	551	0.459		
	Total		394.300	570	0.692		
	Total		539.440	599	0.901		
Total							

Table 4.5 shows the ANOVA Table for Tukey's Nonadditivity Test. This test is the criterion to understand whether the scale is prepared in additive form or not. The hypothesis are given below (Özdamar, 2013).

H_0 : There is no non-additivity.

H_1 : There is non-additivity.

According to the Table 4.5, H_0 cannot be rejected because p-value is greater than alpha. So, we conclude that scale is additive.

Table 4.6 Hotelling's t-squared test

Hotelling's T-Squared Test				
Hotelling's T-Squared	F	df_1	df_2	Sig
292.302	5.835	19	11	0.002

Table 4.6 shows the Hotelling's T-squared Test results. This test examines whether the response means are equal to each other or not. The hypothesis are given below.

H_0 : Response means are equal to each other.

H_1 : Response means are not equal to each other.

According to the results, H_0 is rejected (p-value ≤ 0.05). We conclude that the response means are not equal to each other.

4.2 Factor Analysis

Firstly, the data were input into SPSS version 20 software program. The relationship of each questions to the underlying factor (sub-scale) is expressed by the so-called factor loading. Sometimes, the estimated loadings from a factor analysis model can give a large weight on several factors for some of the measured variables, making it difficult to interpret what those factors represent. The goal of factor rotation is to find a solution for which each variable has only a small number of large loadings, i.e., is affected by a small number of factors, preferably only one. Rotations can be orthogonal or oblique. Varimax is one common criterion for orthogonal rotation. Using Varimax rotation of these data Factor Analysis (FA) was conducted for obtain the sub-scales questions of survey. The output of a Factor Analysis looking at indicators of 20 questions and five resulting factors (sub-scales) are given in Table 4.7.

Table 4.7 Rotated component matrix for factor analysis

Rotated Component Matrix					
	Component				
Questions	1	2	3	4	5
Question 1	0.353	0.412	0.148	0.047	-0.132
Question 2	0.027	0.254	0.671	-0.139	0.110
Question 3	0.072	0.324	0.358	-0.091	0.663
Question 4	0.187	0.613	0.205	0.013	0.283
Question 5	0.098	0.700	0.052	0.135	0.322
Question 6	0.153	0.777	0.045	0.157	-0.193
Question 7	-0.208	0.232	0.676	0.324	-0.286
Question 8	0.015	0.373	0.456	-0.092	0.219
Question 9	0.149	0.047	0.193	0.194	0.071
Question 10	0.583	0.067	0.202	0.404	0.216
Question 11	0.393	-0.047	0.296	0.324	0.173
Question 12	0.141	0.122	-0.061	0.669	0.168
Question 13	0.560	0.066	0.116	0.454	-0.041
Question 14	0.738	0.186	0.201	-0.053	-0.044
Question 15	0.753	0.184	-0.085	0.093	0.021
Question 16	0.664	0.335	-0.083	0.171	0.077
Question 17	0.253	-0.072	0.728	0.025	0.138
Question 18	0.469	-0.139	0.216	0.470	0.161
Question 19	0.126	0.441	-0.049	0.646	-0.110
Question 20	0.002	-0.007	0.028	0.192	0.672

According to the result of the FA, final form of the survey is given in Table 4.8.

Table 4.8 Final survey form after FA

1. Communication						
10	My bank is efficient in providing IB service.	1	2	3	4	5
11	The IB service provided by my bank is convenient in terms of location.	1	2	3	4	5
13	My bank always provides prompt IB service.	1	2	3	4	5
14	The IB service of my bank is very open in case of communicating with its customer.	1	2	3	4	5
15	The IB service is very prompt in case of responding to its customer.	1	2	3	4	5

Table 4.8 Final survey form after FA (cont.)

16	The IB service of my bank always provides me quality information.	1	2	3	4	5
18	My bank provides all the documentary evidences for all the transactions performed online.	1	2	3	4	5
2. Importance of Internet Banking Needs						
1	IB is an easier way to solve my banking needs.	1	2	3	4	5
4	IB is a convenient way to manage my finance.	1	2	3	4	5
5	IB allows me to manage my finance efficiently.	1	2	3	4	5
6	IB is useful to manage my financial resources.	1	2	3	4	5
3. Benefits of Internet Banking						
2	IB gives me greater control over my finances	1	2	3	4	5
7	IB is compatible with my lifestyle.	1	2	3	4	5
8	IB fits well with the way I like to manage my finances.	1	2	3	4	5
17	I am able to access the internet at any time at work and at home.	1	2	3	4	5
4. Convenience						
3	IB allows me to manage my finance effectively.	1	2	3	4	5
20	My bank provides trial period of doing online banking transaction before registering for the IB service.	1	2	3	4	5
5. Customer Adaptation to Internet Banking						
9	Using IB fits into my working style.	1	2	3	4	5
12	My bank provides 24 hours IB service.	1	2	3	4	5
19	The terms and conditions of IB service set by my bank are acceptable.	1	2	3	4	5

The scale that is used for the survey can be considered in five sub-scales. These are; Communication, Importance of Internet Banking Needs, Benefits of Internet Banking, Convenience and Customer Adaptation to Internet Banking.

Table 4.9 Cronbach's Alpha for internal consistency checking

Factors	Questions	Scale Mean If Item Deleted	Scale Variance If Item Deleted	Corrected Item-Total Correlation	Cronbach's Alpha	Cronbach's Alpha If Item Deleted
Communication	Question 10	80.33	89.195	0.703	0.870	0.900
	Question 11	80.33	90.575	0.698		0.901
	Question 13	80.53	84.740	0.713		0.899
	Question 14	80.90	91.197	0.480		0.906
	Question 15	81.03	90.654	0.505		0.905
	Question 16	80.93	88.064	0.699		0.900
	Question 18	80.13	92.257	0.499		0.905
Importance of Internet Banking Needs	Question 1	80.30	93.183	0.596	0.819	0.904
	Question 4	80.47	89.154	0.724		0.900
	Question 5	80.67	87.678	0.689		0.900
	Question 6	80.70	90.700	0.610		0.903
Benefits of Internet Banking	Question 2	80.20	89.131	0.671	0.698	0.901
	Question 7	80.30	93.872	0.349		0.909
	Question 8	80.43	93.289	0.541		0.905
	Question 17	80.10	92.852	0.499		0.905
Convenience	Question 3	80.47	89.844	0.723	0.559	0.900
	Question 20	82.33	92.161	0.329		0.912
Customer Adaptation to Internet Banking	Question 9	80.10	91.266	0.577	0.563	0.903
	Question 12	80.33	93.609	0.257		0.914
	Question 19	80.60	92.041	0.452		0.906

The level of internal consistency of the questions in each of the five factors was verified. As a result, the five factors of questions are appropriate for further use in data processing.

4.3 Construct SEM Model Using AMOS

In this section, the shapes that represents observed and latent variables are drawn by using the menu bar. And, all the questions, which are also the observed variables of the application, are assigned to the rectangular/square boxes one by one. According to the Table 4.8, grouping process is performed. Because the factors' name are too long, short forms like a, b, c, d and e are given each of them, respectively. Also the error terms of the observed variables are defined and short names are given them, too. Figure 4.1 shows the view of AMOS program for the initial model of the analysis.

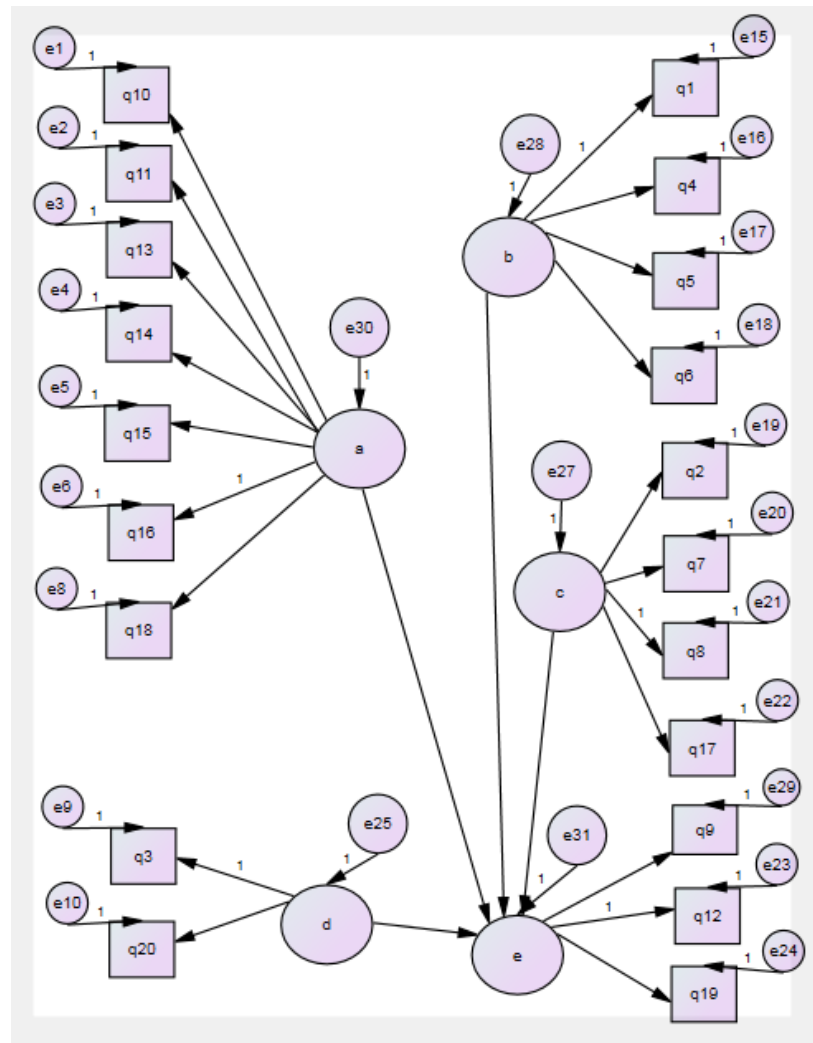


Figure 4.1 Initial model for the analysis

The relationships among latent variables are depicted in the overall measurement model shown in Figure 4.1. The five latent variables were measured by respective questions. These latent variables can correlate with each other. The selected fit indices of this model are given in Table 4.10. These indices indicate that this model has not achieved a good fit.

Table 4.10 Fit indices for the initial model

Fit Indices For the Initial Model		
Fit index	Structural Model	Recommended Value
χ^2 / df	3.751	< 3
GFI	0.806	> 0.8
AGFI	0.756	> 0.8
RMSEA	0.105	< 0.08
CFI	0.651	> 0.9

AIC, BIC and CAIC should be smaller than saturated and independent models. Table 4.11 shows their values.

Table 4.11 Information criterions for the initial model

Information Criterions			
Model	AIC	BIC	CAIC
Default model	712.464	864.229	907.227
Saturated model	420.000	1161.180	1371.180
Independence model	1544.645	1615.223	1635.233

As it can be seen in the Table 4.11, AIC doesn't satisfy the condition, whereas the other ones do. Therefore these results are not sufficient. Since the hypothesized model is rejected based on goodness-of-fit statistics, we must modified this model that fits the data. In modification step ,the causal paths can be evaluated in terms of statistical significance and strength using standardized path coefficient that range between -1 and +1. Based on α of 0.05, the test statistic generated from the AMOS output should be greater than ± 1.96 to indicate that the null hypothesis can be rejected. The rejection of the null hypothesis means that the structural coefficient is

not zero. The hypotheses to test the relationships among the factors for this study are given below.

I. H_0 : There is a positive relationship between benefits of internet banking and importance of internet banking needs.

H_1 : There is not a positive relationship between benefits of internet banking and importance of internet banking needs.

II. H_0 : There is a positive relationship between communication and importance of internet banking needs.

H_1 : There is not a positive relationship between communication and importance of internet banking needs.

III. H_0 : There is a positive relationship between customer adaptation to internet banking and communication.

H_1 : There is not a positive relationship between customer adaptation to internet banking and communication.

IV. H_0 : There is a positive relationship between customer adaptation to internet banking and convenience.

H_1 : There is not a positive relationship between customer adaptation to internet banking and convenience.

After reviewing the statistical significance of the standardized paths, the strength of relationships among the variables can be observed. According to results modified model is obtained. The modified model fit indices are shown Table 4.12 and Table 4.13.

Table 4.12 Fit indices for the last model

Fit Indices For the Last Model		
Fit index	Structural Model	Recommended Value
χ^2 / df	2.207	<3
GFI	0.886	>0.8
AGFI	0.851	>0.8
RMSEA	0.069	<0.08
CFI	0.852	>0.9

Table 4.13 Information criterions for the last model

Information Criterions For the Last Model			
Model	AIC	BIC	CAIC
Default model	453.326	626.268	675.268
Saturated model	420.000	1161.180	1371.180
Independence model	1544.645	1615.223	1635.233

Table 4.12 shows the goodness of fit of generated or modified structural model is better compared to the initial model Figure 4.2 shows the residual covariance matrix. This matrix is the difference between observed and estimated variance-covariance. If these values are small, model has a good fit. As it can be seen, all the values are about zero. Therefore, the observed and estimated values are close to each other.

Residual Covariances (Group number 1 - Default model)

	q18	q6	q17	q3	q9	q16	q15	q14	q13	q11	q10	q20	q19	q12	q5	q8	q7	q2	q4	q1
q18	,000																			
q6	-,066	,000																		
q17	,050	-,053	,000																	
q3	,029	-,037	-,003	,014																
q9	,000	,019	,029	,023	,000															
q16	,005	,024	-,006	,025	-,002	,000														
q15	-,013	-,014	,003	-,019	,033	,011	,000													
q14	-,030	,011	,066	,025	-,002	,026	,058	,000												
q13	,032	-,006	,053	-,037	-,030	-,012	-,008	-,014	,000											
q11	,045	-,033	,013	,002	,026	-,035	-,011	,005	,021	,011										
q10	,000	-,026	,047	,033	,024	-,013	-,030	-,009	,012	,070	,000									
q20	,044	-,005	,060	,041	,005	,055	,079	,037	,084	,129	,090	,000								
q19	,029	-,006	-,047	-,035	-,021	,013	-,011	-,037	,020	-,023	,007	,051	-,004							
q12	,000	-,040	,002	,020	,000	-,014	-,013	-,045	,031	,001	-,015	,036	,011	,000						
q5	-,037	,028	,006	,046	,049	,041	-,038	-,036	-,031	-,016	,001	,111	,015	,016	,000					
q8	,004	,006	,005	,013	,042	-,011	-,012	,051	-,044	,029	,022	,058	-,003	-,023	,065	,000				
q7	,002	,011	-,002	,004	,031	-,048	-,073	-,010	-,008	,018	-,009	,054	,029	,014	-,040	,000	-,006			
q2	-,004	-,031	,011	,019	,027	-,021	-,041	,020	-,002	,009	-,018	,035	-,025	-,039	-,018	-,029	-,003	,000		
q4	-,013	-,002	-,013	,016	,045	,001	-,007	-,008	,009	,004	,013	,032	-,027	-,008	-,003	,025	,016	,029	,003	
q1	,000	-,008	,011	-,012	,035	,050	,040	,015	,034	,006	,030	,033	,001	,000	-,012	,003	,006	-,004	-,006	,000

Figure 4.2 Residual covariance matrix

Since the initial model did not achieve model fit ($p < .000$), therefore, the explanation of hypotheses result is based on modified model which achieved model fit of ($p > .000$). The revised model produces regression standardized estimates direct effects readings as shown in Table 4.14.

Table 4.14 Standardized regression weights

Standardized Regression Weights			
	Estimate		Estimate
b ← c	0.542	q19 ← e	0.419
a ← b	0.592	q20 ← d	0.210
q17 ← c	0.515	q10 ← a	0.510
q6 ← b	0.680	q11 ← a	0.413
q18 ← a	0.500	q15 ← a	0.646
e ← a	0.736	q16 ← a	0.683
e ← d	0.052	q14 ← a	0.632
q1 ← b	0.464	q9 ← e	0.477
q4 ← b	0.630	q13 ← a	0.634
q2 ← c	0.636	q3 ← d	0.787
q7 ← c	0.570	q3 ← c	0.617
q8 ← c	0.492	q11 ← q17	0.329
q5 ← b	0.645	q19 ← q6	0.336
q12 ← e	0.590	q4 ← d	0.241
q7 ← d	0.355	q10 ← q18	0.294

As a consequence, the structural equations can be written as;

$$b = 0.542 \times c$$

$$a = 0.592 \times b$$

$$e = 0.736 \times a + 0.052 \times d$$

And the letters are replaced with the factor names. The last model is obtained as follows.

Importance of the internet banking needs = $0.542 \times$ Benefits of internet banking

Communication = $0.592 \times$ Importance of the internet banking needs

Customer adaptation to internet banking = $0.736 \times$ Communication + $0.052 \times$ Convenience

Writing these equations, the results of hypothesis testing is given in. Table 4.15.

Table 4.15 Results of the hypothesis testing

Results of the Hypothesis Testing		
Hypotheses	p-value	Result
I	***	Not Rejected
II	***	Not Rejected
III	***	Not Rejected
IV	***	Not Rejected

*** p<0.00

According to the results in Table 4.15, All hypotheses are supported i.e. all direct paths are significant and positive (C.R. values $>+/-1.96$; p-value < 0.05).

CHAPTER FIVE

CONCLUSION

This study contains the theory of SEM, computer programs for its analysis, model fit indices and the application. The aim of application is to identify the factors that determine academicians to use Internet banking services at Dokuz Eylül University. The analysis of the application focused on the reliability and validity of the data and measurement model, and the path coefficients and goodness of fit of the structural model. According to analysis result we obtain the model given in below.

Importance of the internet banking needs = $0.542 \times$ Benefits of internet banking

Communication = $0.592 \times$ Importance of the internet banking needs

Customer adaptation to internet banking = $0.736 \times$ Communication + $0.052 \times$ Convenience

This model shows, there are positive relationships among the sub-scales of the survey. The results of this study provide bank managers and policy makers, an insight into the most influential factors that determine Dokuz Eylül University academicians' intention to use Internet banking. For instance, communication was found to be one of the most important factor in determining customer adaptation to internet banking as it influences on importance of internet banking needs, benefits of internet banking and convenience. Banks should plan awareness campaigns especially for minimizing the risk perceptions of the customers and to increase their confidence on the system, by communicating its benefits and advantages over other traditional channels.

The applicability of the proposed model can be evaluated in other internet banking customer where similar conditions are prevailing.

REFERENCES

- Anderson, T.W., & Amemiya, Y. (1988). The asymptotic normal distribution of estimates in factor analysis under general conditions. *Annals of Statistics* 16, 759-771.
- Anderson, T.W., & Amemiya, Y. (1990). Asymptotic chi-square tests for a large class of factor analysis models. *Annals of Statistics* 18 (3), 1453-1463.
- Austin J.T., & Caldero'n, R.F. (1996). Theoretical and technical contributions to structural equation modeling: An updated annotated bibliography. *Structural Equation Modeling*, 3, 105–175.
- Austin J.T., & Wolfle, L.M. (1991). Annotated bibliography of structural equation modeling: technical work. *British Journal of Mathematical and Statistical Psychology*, 44, 93–152.
- Bayram, N. (2010). *Yapısal eşitlik modellemesine giriş AMOS uygulamaları* (1. Baskı), Bursa: Ezgi Kitabevi.
- Bentler, P. M. (1980). Multivariate analysis with latent variables: Causal modeling. *Annual Review of Psychology*, 31, 419-456. Reprinted in C. Fornell (Ed.), *A second generation of multivariate analysis*, 1. New York: Praeger (1982).
- Bentler, P. M. (1986). Structural modeling and Psychometrika: An historical perspective on growth and achievements. *Psychometrika*, 51, 35-51.
- Bentler, P. M. (1990). Comparative fit indexes in structural models. *Psychological Bulletin*, 107, 238–246.
- Bentler, P. M. (1995). *EQS Structural equations program manual*. Encino, CA: Multivariate Software, Inc.

- Bentler, P. M., & Bonett, D. G. (1980). Significance tests and goodness of fit in the analysis of covariance structures. *Psychological Bulletin*, 88, 588–606.
- Bentler, P. M., & Dudgeon, P. (1996). Covariance structure analysis: Statistical practice, theory, directions. *Annual Review of Psychology*, 47, 563-592.
- Bollen K.A., & Noble, M.D. (2011). Structural equation models and the quantification of behavior. *Proceeding of the National Academy of Sciences*, 108 (3), 15639-15646.
- Demir, M. (2011). İşgörenlerin çalışma yaşamı kalitesi algılamalarının işte kalma niyeti ve işe devamsızlık ile ilişkisi. *Ege Akademik Bakış*, 11 (3), 453-464.
- Hu, L., & Bentler, P.M. (1999). Cut off criteria for fit indexes in covariance structure analysis: Conventional criteria versus new alternatives. *Structural Equation Modeling: A Multidisciplinary Journal*, 6(1) , 1-55.
- Iriondo, J.M., Albert, M.J., & Escudero, A. (2003). Structural equation modeling: an alternative for assessing causal relationships in threatened plant populations. *Biological Conservation*, 113, 367-377.
- Jahangir, N., & Parvez, N. (2012). Factors determining customer adaptation to internet banking in the context of private commercial banks of Bangladesh. *Business Perspectives and Research*, 25-36.
- Jöreskog, K. G., & Sörbom, D. (1979). *Advances in factor analysis and structural equation models*. New York: University Press of America.
- Jöreskog, K., & Sörbom, D. (1993). LISREL 8: *Structural Equation Modeling with the SIMPLIS Command Language*. Chicago, IL: Scientific Software International Inc.
- Kline, R.B. (2011). *Principles and practice of Structural Equation Modeling* (3rd ed.). New York: The Guilford Press.

- McDonald, J.A., & Clelland, D.A. (1984). Textile workers and union sentiment. *Social Forces*, 63, 502-521.
- McDonald, R. P., & Marsh, H.W. (1990). Choosing a multivariate model: Noncentrality and goodness of fit. *Psychological Bulletin*, 107, 247–255.
- McQuitty, S. (2004). Statistical power and structural equation models in business research. *Journal of Business Research*, 57 (2), 175–83.
- Özdamar, K. (2013). *Paket programlar ile istatistiksel veri analizi* (9. Baskı). Ankara: Nisan Kitabevi.
- Reisinger Y., & Mavondo F. (2006). Structural equation modeling: critical issues and new developments. *Journal of Travel & Tourism Marketing*, 21 (4), 41-67.
- Schumacker R.E., & Lomax R.G. (2010). *A beginners guide to structural equation modeling*. New York: Routledge.
- Tucker, L. R., & Lewis, C. (1973). A reliability coefficient for maximum likelihood factor analysis. *Psychometrika*, 38, 1–10.
- Ustasüleyman, T. & Eyüboğlu, K. (2010). Bireylerin internet bankacılığını benimsemesini etkileyen faktörlerin yapısal eşitlik modeli ile belirlenmesi. *BDDK Bankacılık ve Finansal Piyasalar*, 4 (2), 11-31.
- Yuan, K.H., Bentler, P.M., & Zhang,W. (2005). The effect of skewness and kurtosis on mean and covariance structure analysis. *Sociological Methods & Research*, 34(2), 240-258.