

Modeling Time Series

The case of the Greek new-car sales sector.



Maria K. Voulgaraki

A thesis submitted for the degree of
Philosophiae Doctor (Ph.D.) in Economic Science.

Department of Economics,
University of Crete, UOC

September 2020
Rethymnon, Crete, Greece

Supervisor Committee members:

Supervisor: Assistant Professor *Georgios K. Tsiotas*, University of Crete,GR

Associate Professor *Dikaïos Tserkezos*, University of Crete,GR

Associate Professor *Margarita Genius*, University of Crete,GR

Professor *Raphael Markellos*, Norwich Business School, Univ.of East Anglia, UK

Doctoral Dissertation Oral Defense date: 20 November 2020

Doctoral Dissertation Defence Committee members:

Professor *Dimitris G. Kirikos*, Hellenic Mediterranean University, GR

Professor *Christos B. Floros*, Hellenic Mediterranean University, GR

Associate Professor *Ioannis A. Venetis*, University of Patras,GR

and all Supervisor Committee members

PhD Thesis written in L^AT_EX

© Copyright 2020 by Maria K. Voulgaraki

All rights reserved.

No part of this Thesis may be produced or transmitted in any form or by any means electronic or mechanical including photocopying recording or any information storage and retrieval system without the prior written permission of the publisher.

For permission contact Maria K.Voulgaraki

E-mail:maria.voulgaraki@uoc.gr

Dedication

This dissertation is lovingly dedicated to my family :

- my father, *Kallergos*, who believed in me and taught me to have faith in my inner-strength and work hard to make my dreams come true,
- my mother, *Despoina*, who brought me up to always think positive, never give up, dream big and always set new goals in life and
- my beloved husband *Yannis* and our precious twins *Stelios & Despoina* for being my source of super-power and my daily inspiration!

Their support, encouragement, and unconditional, generous love are proven to be my best assets in life.

Acknowledgments

This Ph.D. marks the successful completion of my doctoral studies in Economics at the Department of Economics at the University of Crete. I hereby would like to thank my advisor and co-advisors in this long but fulfilling research trip. I am most grateful to my supervisor Prof. Georgios Tsiotas (Assistant Professor in Quantitative Methods, Department of Economics, University of Crete, GR), for his ideas, continual collaboration, encouragement, and support. Without him, I would not have even started this research journey. His expertise in Quantitative Methods and his demand for excellence determine the quality of my research.

Special thanks to all the members of my committee members: Prof. Tserkezos Dikaios (Associate Professor of Econometrics, Department of Economics, University of Crete), Prof. Margarita Genius (Associate Professor of Econometrics, Department of Economics, University of Crete, GR), Prof. Raphael Markellos (Professor of Finance, Norwich Business School, University of East Anglia, UK) for their ideas, professional advice and encouragement.

Many thanks to Prof. Shabbar Jaffry (Professor of Economics and Finance, Portsmouth Business School, University of Portsmouth, UK) and Prof. Dimitris Kirikos (Professor of Economics, School of Management and Economics, Hellenic Mediterranean University, GR) my deepest gratitude to both of them, for believing in me and encouraging my early steps of this research journey.

I would like to thank Prof. Chrisoula Zacharopoulou (Professor of Economics, Depart-

ment of Economics, School of Economic Sciences, Aristotle University of Thessaloniki), who dug out my initial interest in econometrics as a bachelor student, giving me a different perspective of economic science and inspiring me through this long-life learning journey.

Many thanks to Prof. Rob J. Hyndman (Professor of Statistics at Monash University, Australia), that I was lucky to meet in one of his summer-schools in Switzerland, his expertise knowledge helped me solve various problems in my empirical research and his personality inspired me with his generous willingness to disseminate and share his expertise knowledge. Special thanks to the “A.G. Leventis Foundation” for sponsoring, twice, part of my travel expenses to London, for my participation in the Hellenic Observatory Ph.D. Symposium on Contemporary Greece and Cyprus, at London School of Economics and Political Science, European Institute in UK.

I want to express my gratitude to my family - my beloved husband Ioannis Tampakakis and our unique twins Stelios and Despoina, my parents Kallergos and Despoina Voulgaraki, my beloved sisters Vivian and Niki and my brothers in law Minas and Michalis Papadakis - for their support, encouragement, love, and help from day one until the last day of this process and over the years, that made success both possible and rewarding.

Many thanks to my colleagues and friends, who stood by me during the elaboration of this research, and to all the administration staff in the Department of Economics, University of Crete, especially to Mrs Ioanna Yiotopoulou and Mrs Eleni Kalogeraki for their constant support in all administrative matters.

Last but not least, I would like to warmly thank all personnel and institutional structures in Greece, that facilitate and support policies and initiatives for working mothers and females PhD students in general, to fulfill their research goals in academia, achieving gender equality and empowering all women and girls in academia and life in general.¹

¹ “My favourite popular quote: Those who do not dream of flying, will never grow any wings!”

Abstract

The *purpose* of this thesis is to study the time series modeling and make a comparative empirical analysis of the Greek new car-market, report and discuss the research results, and investigate whether time series methods can successfully be applied in marketing data and give reliable forecasts. Although the practice of forecasting marketing data, like sales, is a widely researched area, it hasn't been extensively applied for the Greek new-car sales sector. Therefore, this research is an attempt to report and discuss the new car sales forecasting practices concerning Greek companies in a turbulent economic time interval and a strictly supervised economic environment.

The *design* of this empirical research application reports on sales forecast practices, using monthly data of the Greek new-car sales sector, available from the Association of motor vehicles importers representatives (AMVIR) of Greece. The monthly new car registration number is assumed, to be equal to the monthly new car sales level in the Greek market. The *methodology* of this applied study uses an in-sample and an out-of-sample time series modeling and forecasting, in a variety of new-car sales firms in the Greek retail market, testing various data sets in a time frame of the last two decades (1998 till 2016).

Despite the difficult economic conditions in the Greek market, the researcher develops a comparative study, by modelling various time series and analysing them. Simple time series models are used, like the Mean or Average, the Naïve and the Seasonal Naïve models, but also some more sophisticated ones, like for example the Linear Models with Seasonal Dummies (LMSD), the Exponential Smoothing state space model (ETS), the Seasonal

Autoregressive Integrated Moving Average (SARIMA), and the family of Seasonal Autoregressive Integrated Moving Average-General Autoregressive Conditional Heteroscedastic Model (SARIMA- GARCH).

Furthermore, in addition to the original form of the variables and their log transformation, the use of the family of Box-Cox (1964) data transformation is applied in this research, where the transformation parameter λ is determined by the method of Guerrero (1993). All these alternative approaches improve the quality and performance of the data, which ultimately proves that transforming variables provides a powerful tool for developing models. Lastly, the approach of combined forecast, as a forecasting tool, of various time series forecasting models is reported and discussed. This technique uses information from various individual forecasting methods, assuming either equal weights or optimally weighted forecasts, depending on the limitations, set by each type of combined forecasts studied.

The empirical *findings* of this study gave evidence of an improvement in forecasting and confidence intervals, using the appropriate data transformation process. In addition, there is the prevailing conclusion that it is extremely difficult to find a single model, that can capture and forecast new car sales for all companies and at all times. Each car firm should be treated separately. However, research suggests that data transformation, the use time series models, and the combination of their predictions in the right way, prove to be beneficial in improving the accuracy of predictions for all cases of this research in the Greek car market.

The conclusions of this thesis are very interesting cause the research takes place over a difficult time for the Greek economy. During the research period the country signed three (3) Memorandums of Understanding, on specific economic policy conditions, which lead to an economic supervision of the country, by a decision group referred as “TROIKA”, which was formed by the International Monetary Fund, the European Central Bank and the European Commission. Greek government was forced to implement austerity measures to its citizens, as a consequence of the supervision. Greek consumers, due to economic uncertainty, liquidity shortage, and bank crises, prolonged buying durable goods, like cars,

which resulted in a sharp fall of new car sales in the Greek market.

Finally, the *originality* of this research and the *added value in economic science* is that it is the first time, to my knowledge, that new car sales sector of the Greek market is treated as time series, with an extensive application of time series modeling, for the first two decades of the 21st century. Additionally, this thesis research contributes in depicting the best way for predicting sales, and helps improve the quality and accuracy of sales forecasts, hedge against risk, and reduce failure, for the Greek new-car decisions makers. It also emphasizes the problems, that evolved from the diminishing new car sales, due to the economic crisis in the Greek new car market.

Keywords: Time Series, Forecasting, Greek market, new car retail sector, Naïve Model, Seasonal Naïve Model, Linear model with seasonal dummies, Exponential Smoothing state space model (ETS), Seasonal Autoregressive Integrated Moving Average (SARIMA), Seasonal Autoregressive Integrated Moving Average-General Autoregressive Conditional Heteroscedastic Model (SARIMA- GARCH), Box-Cox transformations, Guerrero method, forecasts combinations.

Contents

1	Introduction	1
1.1	Thesis Purpose, aims and objectives	3
1.2	Background and Motivation	4
1.3	Data Source and Used Software	6
1.4	Thesis Overview	7
1.5	List of publications & research presentations.	9
1.6	Research Contribution	10
2	Time Series Empirical Analysis.	13
2.1	Introduction	13
2.2	Overview of the Greek new car market.	15
2.3	Empirical Data Analysis.	21
2.3.1	Descriptive Statistics.	22
2.3.2	Summary statistics of Data	25
2.3.3	Graphical Presentation.	28
2.4	Related Work - Literature Review	38
2.5	Theoretical Framework of Time Series Models	44
2.5.1	Simple Time Series Forecasting Methods	44
2.5.2	Linear Models with Seasonal Dummies (LMSD)	46
2.5.3	Holt–Winters methods modeling	49

2.5.4	Exponential Smoothing State Space (ETS) Models	51
2.5.5	Autoregressive Integrated Moving Average (ARIMA)	54
2.5.6	Seasonal ARIMA Model (SARIMA)	58
2.5.7	Generalized Autoregressive Conditional Heteroscedastic	60
2.6	Forecasting Methodology.	64
2.6.1	Qualitative vs Quantitative Forecasting Methods.	65
2.7	Discussion	67
3	Time Series with-in-sample Modeling.	69
3.1	Introduction	69
3.2	Simple Time Series Models	71
3.3	Holt-Winters & Exponential Smoothing State Space Models	71
3.4	Linear Modelling with Seasonal Dummies Models	72
3.5	(S)ARIMA Modelling using Box-Jenkins Methodology.	79
3.5.1	STEP 1: Identification.	81
3.5.2	STEP 2: Estimation	93
3.5.3	STEP 3: Diagnostics checking.	94
3.6	SARIMA - GARCH Modelling	112
3.6.1	SARIMA GARCH Model Building	113
3.6.2	Step 1. Identification	114
3.6.3	Step 2. Estimation	120
3.6.4	Step 3. Diagnostic Checking.	122
3.7	Discussion	123
4	Time Series out-of-sample Forecasting.	125
4.1	Introduction.	125
4.2	Forecast Accuracy Measures.	128
4.3	Empirical out-of-sample Forecasting.	131
4.3.1	Graphical Presentation.	133
4.3.2	Accuracy Measures Evaluation.	134

4.4	Discussion	144
5	SARIMA-GARCH Forecasting Models.	147
5.1	Introduction.	147
5.1.1	Autocorrelation function of SARIMA model residuals.	150
5.1.2	Testing for ARCH/GARCH Effects	154
5.2	Fitting GARCH (1,1) Model	156
5.2.1	Model Selection and Analysis	156
5.2.2	Alternative Conditional Distributions	157
5.3	Diagnostic checking of SARIMA-GARCH Models	159
5.4	SARIMA-GARCH Prediction Intervals	164
5.5	Discussion	167
6	Data Transformations and Forecasting.	169
6.1	Introduction.	169
6.2	The Box-Cox Transformation Review.	170
6.3	Methodology of Box-Cox Transformations.	175
6.4	Empirical Results of Data Transformation.	178
6.4.1	ETS models and Box and Cox transformation (in-sample).	179
6.4.2	ETS models and Box and Cox transformation (out of-sample).	182
6.4.3	Time Series forecasting & transformation (out-of-sample).	191
6.5	Confidence Interval.	196
6.6	Discussion.	199
7	Forecast Combinations.	215
7.1	Introduction.	215
7.2	Methodology	219
7.2.1	Combining forecasts without training	220
7.2.2	Combining forecasts with training	220
7.2.3	Data generating procedure	223

7.3	Empirical Results	223
7.3.1	Graphical presentation of combination forecasting	225
7.3.2	Accuracy measures of combination forecasting.	229
7.4	Discussion	230
A	Appendix: New Car Sales Graphs	257
B	Appendix: Autocorrelations in Data	265
C	List of Abbreviations	275

List of Figures

2.1	Total New Car Sales in Greece 1998-2019 [Bar Chart].	16
2.2	Household Car Ownership in Greece [2011 Census].	19
2.3	Market Share for Greek new-cars [Pie Chart].	20
2.4	Toyota new-car sales. Panel (a)Line Plot, (b)Box Plot, (c) ACF, (d)PACF	35
2.5	Opel new-car sales. Panel(a)Line Plot, (b)Box Plot (c) ACF, (d)PACF . .	36
2.6	Fiat new-car sales. Panel(a)Line Plot, (b)Box Plot,(c)ACF, (d)PACF . . .	37
2.7	Box and Jenkins Methodology.	56
2.8	Forecasting Methods [Tree Graph].	65
2.9	Time Series Forecasting Methods [Tree Graph].	67
3.1	ACF and PACF for OPEL Sales (Log-values)	82
3.2	OPEL-SARIMA residual (a)Histogram,(b)Q-Q Plot	100
3.3	Opel-SARIMA residual(a) Standardized (b)ACF (c)Ljung-Box	100
3.4	Opel SARIMA model Squared and Standardized Residual Plots	101
3.5	ACF of SARIMA resid.and $resid.^2$ -Opel,Toyota,VW,Hyundai,Peugeot . . .	105
3.6	ACF of SARIMA resid.and $resid.^2$ -Ford,Nissan,Skoda,Fiat,Citroen	106
3.7	McLeod-Li Tests for SARIMA residuals.	116
4.2	Forecasting for Opel. Best Model:Naïve. Set A	136
4.3	Forecasting for Opel. Best Model:S.Naïve & ETS. Set B	137
4.4	Forecasting for Opel. Best Model:ETS. Set C	138

4.5	Forecasting for Opel. Best Model:ETS. Set D	139
4.1	SARIMA Model residual analysis. OPEL,TOYOTA,FIAT.Set A	140
5.1	SARIMA residuals Plot (Set A)	151
5.2	ACF of SARIMA residuals & squared residuals (Set A)	153
5.3	OPEL SARIMA-GARCH(1,1)STD stand.residuals & deviation-Set A . . .	161
5.4	OPEL SARIMA-GARCH(1,1)STD residuals ACF plots-A	163
5.5	Confidence intervals 95% SARIMA-GARCH model (Opel-Set A)	165
5.6	Confidence intervals SARIMA-GARCH model(TOYOTA-Set A)	166
6.1	Forecasting Opel new-car sales using ETS and data transformation	188
6.2	Forecasting Toyota new-car sales using ETS and Data transformation. . . .	189
6.3	Forecasting Fiat new-car sales using ETS and data transformation.	190
6.4	Opel new car sales forecasts.Set D-Actual	193
6.5	Opel new car sales forecasts. Set D,Box-Cox/Guerrero	195
6.6	ETS forecasts & 95%CI for Opel (Set D-Original)	197
6.7	Seasonal Naïve forecasts & 95%CI for Opel (Set D-Original)	197
6.8	LMSD forecasts & 95%CI (Opel-Set D-BoxCox/Guerrero).	199
6.9	ETS Model forecast & 95%CI for Opel (Set D-BoxCox/Guerrero)	199
6.10	Forecasting new-car sales and time series model selection.	201
6.11	Forecasting new-car sales and data transformation.	201
7.1	Simple Average Combination forecasts (Opel-Set D,BC/Guerrero)	225
7.2	Simple Average Combination forecasts (Toyota-Set D,BC/Guerrero)	226
7.3	Simple Average Combination forecasts (Fiat-Set D, BC/Guerrero)	226
7.4	Selection of the best Combined forecasts (Opel-Set D,BC/Guerrero)	227
7.5	Selection of the best Combined forecasts (TOYOTA-Set D, BC/Guerrero .	228
7.6	Selection of the best Combined forecasts (FIAT -Set D, BC/Guerrero) . . .	228
7.7	Forecast Combination for Opel, Toyota and Fiat new-car sales.	233

A.1	HYUNDAI new-car sales (a)Line Plot, (b)Box Plot (c) ACF, (d)PACF (original values)	257
A.2	FORD new-car sales (a)Line Plot, (b)Box Plot (c) ACF, (d)PACF (original values)	258
A.3	NISSAN new-car sales (a)Line Plot, (b)Box Plot (c) ACF, (d)PACF (orig- inal values)	259
A.4	CITROEN new-car sales (a)Line Plot, (b)Box Plot (c) ACF, (d)PACF (original values)	260
A.5	PEUGEOT new-car sales (a)Line Plot, (b)Box Plot (c) ACF, (d)PACF (original values)	261
A.6	VOLKSWAGEN new-car sales (a)Line Plot, (b)Box Plot (c) ACF, (d)PACF (original values)	262
A.7	SKODA new-car sales (a)Line Plot, (b)Box Plot (c) ACF, (d)PACF (orig- inal values)	263
B.1	ACF and PACF of the TOYOTA (logs)	266
B.2	ACF and PACF of the VOLKSWAGEN (logs)	267
B.3	ACF and PACF of the HYUNDAI (logs)	268
B.4	ACF and PACF of the PEUGEOT (logs)	269
B.5	ACF and PACF of the FORD (logs)	270
B.6	ACF and PACF of the NISSAN (logs)	271
B.7	ACF and PACF of the SKODA (logs)	272
B.8	ACF and PACF of the FIAT (logs)	273
B.9	ACF and PACF of the CITROEN (logs)	274

List of Tables

2.1	Descriptive Statistics for original and log values of monthly new-car sales of 10 firms.	26
2.2	Ljung-Box Test for car sales series (original data with $lag = 12$)	33
2.3	Seasonal Dummies Variables	48
2.4	The 30 possible different combinations for ETS models.	53
3.5	Indicator Variables of Linear Model with Seasonal Dummies	72
3.1	Simple models RMSE & MAPE (log data with-in-sample).	74
3.2	Diagnostic Tests for Seasonal Naïve model	75
3.3	Holt-Winters & ETS models RMSE & MAPE (log data with-in-sample). .	76
3.4	Diagnostic Tests for ETS model	77
3.6	Linear Model with Seasonal Dummies RMSE & MAPE (log values).	77
3.7	Diagnostic Tests for LMSD models	78
3.8	ADF & KPSS Stationarity Tests (log with-in-sample).	79
3.9	OCS D Seasonality Unit Root Test. Log data, in-sample.	80
3.10	RMSE and MAPE for different SARIMA models in log new-car sales series.	85
3.11	AIC information criterion for different SARIMA specifications (d=1,D=0).	91
3.13	AIC information criterion for different SARIMA specifications(d=1,D=1). .	92
3.15	SARIMA models for new-car sales "in sample" estimation.	95
3.16	ARCH effect test for SARIMA residuals.	107

3.17	ARCH effect test for squared residuals of selected SARIMA models	108
3.18	Shapiro Wilk & Jarque Bera test for selected SARIMA residuals.	111
3.19	Engle Test for absence of ARCH effect in squared residuals of SARIMA models.	119
3.20	AIC information criterion for SARIMA-GARCH models.	120
3.21	Residuals Diagnostic Tests for SARIMA-GARCH model	122
3.22	Time Series MAPE% metrics for with-in-sample modeling (log values) . . .	124
4.1	Empirical Analysis Training & Test Sets.	126
4.2	Forecast Accuracy Measures	129
4.3	SARIMA models for Forecasting.	132
4.4	Forecast Performance Evaluation for OPEL (Data:log values).	141
4.5	Forecast Performance Evaluation for TOYOTA (Data:log values).	142
4.6	Forecast Performance Evaluation for FIAT (Data:log values).	143
4.7	Summary of Best out-of-sample Forecasting Model(log values)	145
5.1	Testing SARIMA residuals for Heteroscedasticity.	155
5.2	Comparing alternative conditional distribution of variance equation.	158
5.3	Box-Ljung Q-Test of standard. residuals SARIMA-GARCH(1,1)-STD Set A.	160
5.4	Volatility Forecast for monthly new car sales (Opel).	166
5.5	Forecasting Results (RMSE).	168
6.1	Information Criteria for OPEL with Box-Cox & ETS models (in-sample) .	181
6.2	Information Criteria for TOYOTA with Box-Cox & ETS models (in-sample)	182
6.3	Information Criteria for FIAT with Box-Cox & ETS models (in-sample) . .	183
6.4	Forecasting Performance of OPEL using Box-Cox & ETS models (out-of- sample)	184
6.5	Forecasting Performance of TOYOTA with Box-Cox & ETS models (out-of- sample)	186
6.6	Forecasting Performance of FIAT with Box-Cox & ETS models (out-of-sample)	187
6.7	Forecasting Performance Comparison, Opel, Set A.	202
6.8	Forecasting Performance Comparison, Opel, Set B.	203

6.9	Forecasting Performance Comparison, Opel Set C.	204
6.10	Forecasting Performance Comparison, Opel, Set D.	205
6.11	Forecasting Performance Comparison, Toyota, Set A.	206
6.12	Forecasting Performance Comparison, Toyota, Set B.	207
6.13	Forecasting Performance Comparison, Toyota, Set C.	208
6.14	Forecasting Performance Comparison, Toyota, Set D.	209
6.15	Forecasting Performance Comparison, Fiat, Set A.	210
6.16	Forecasting Performance Comparison, Fiat, Set B.	211
6.17	Forecasting Performance Comparison, Fiat, Set C.	212
6.18	Forecasting Performance Comparison, Fiat, Set D.	213
7.4	Summary of Best Forecasting Models(*min value)	232
7.1	Combination Forecast Performance for OPEL.	234
7.2	Combination Forecast Performance for TOYOTA.	235
7.3	Combination Forecast Performance for FIAT.	236

Introduction

Time series analysis is an important tool in modeling and forecasting economic variables. Different time series analysis techniques are applied in this thesis in an attempt to measure, capture, and forecast marketing data of new car sales in the Greek market. Our research starts with analysing various time series variables of the Greek new car market and arguing on the elements that seems to influence the level of sales in that market, during a turbulent economic period for the Greek economy (1998-2016). The study focus initially in applying various simple models, like the average -or mean- model, the naïve, and seasonal naïve model, and continues with the Box-Jenkins modeling methods used in Autoregressive Integrated Moving Average (ARIMA) Models, and the Seasonal ARIMA models, the Linear Model with Seasonal Dummies (LMSD), the state space models based on exponential smoothing models (ETS), and the family of Generalized Autoregressive Conditional Heteroscedastic (GARCH) models.

The empirical evidence of this research shows that simple time series models, like the Seasonal Naïve models can be adequate for a short-term forecasting, while more complicated ones, like the state space models with exponential smoothing (ETS), the Linear Models with Seasonal Dummies (LMSD) and the seasonal autoregressive Integrated Moving Average (SARIMA) Models are better for long term predictions.

Additionally, we study the accuracy in time series modeling, when data used are transformed. So this research presents results when modeling the variables in original values,

alternatively in log values, and finally in transformed values, using the family of Box-Cox process. Furthermore, we emphasize the importance of combining forecasts of the various individual time series models. Thus, instead of trusting one single model for forecasting, we decrease risk by using all the available information, and study the technique of combining forecast from various time series models. This is done either assuming equal weights in each forecast or weight them optimally according to the weighting scheme of the combination method. The combined forecast outperform all single time series models in this research.

This thesis research covers a difficult time interval for the Greek economy at the beginning of the 21st century, from the year 1998 till the year 2016. Hence, this study has a parallel interest in the economic situation in the Greek market, and try to show indirectly, how the economic crisis, the austerity measures, that came into force after the implementation of the financial agreement of the Greek government with European Commission (EC), European Central Bank (ECB) and International Monetary Fund (IMF), changed the economic activity in the automobile retail sector in Greece, reflecting the prolong economic difficulties of the Greek market.

In more detail, the present chapter gives a brief introduction to the subject, followed by the research purpose, aims, and objectives, while it set the questions that puzzled the researcher during the period of study. The purpose of this study and the research objectives are clearly stated and will be fully developed throughout this thesis, background and motivation are explained, while research questions will be answered scientifically and efficiently based in applied econometric theory and practice used in Time Series data. The Data Source and the used software are given in detail, and the thesis overview helps in understanding the structure of this thesis. The research contribution, and the final conclusions of this thesis can give useful information, that can help the decision makers to establish strategies for proper inventory planning, human resource, and logistics applications in sector of retail new car representatives in Greece. Additionally, it gives a major contribution to the science, due to the shortage of research in this sector of the Greek market and the valuable application of a comparative study using time series theory and methods.

1.1 Thesis Purpose, aims and objectives

This thesis research work is in the field of applied econometrics of economic science. It focus in time series analysis, which is one of the major, necessary, and essential field of economics, and business administration. Forecasting time series contain a necessary source of information, which supports each business decision making. Strategic planning, and future actions in business world, are based in economic variables forecasting, that are produced from the econometric analysis of a set of variables, observed over time. People, working as chief executives officers or in other committee decision making groups, are asked to take important decisions, that determine the future of their business. For example, decisions considering the schedule of production line, human or natural resources plan, logistics plan, advertisements expenditures plan, and business strategy in general are based mainly in the forecasting of future product demand, which determines the products' future sales level.

Consequently, the forecasting of a product demand, sales level or services, consist one of the most important functions for the well being of business and organization operations, and considering that, time series analysis have become part of the mainstream statistics. However, relatively little work has been done for non-linear models or marketing data, like new car demand. In addition to that, marketing data like sales have always been the focus of considerable attention in the last decade, and it seems likely that this is only the beginning in terms of the exploration in new areas of statistical modeling. Furthermore, research in the Greek market of new car sales seems to be very basic without a lot of scientific work done for developing more the dynamics of these marketing sales data.

New car sales levels in the Greek market can easily be treated as Time Series data. The use of Time Series data for business analysis is, however, not new. What is new, in economic science, is the ability to collect and analyze high volume of data in sequence, at extremely high velocity to get the clearest picture to predict and forecast market changes, buyers' behavior, resource consumption, environmental conditions and much more.

1.2 Background and Motivation

The motivation for this research work was, firstly to explain, analyse and model new car sales levels in the Greek market based on time series theoretical background. Secondly to investigate the impact of Greek economic crisis and austerity measures on new car demand and how that had influenced the car representatives all over Greece. Thirdly to evaluate the accuracy of various time series forecasting models to predict car demand.

Furthermore, this thesis analyses the time series data for more than 15 years and select the “best” models and forecasting techniques to predict car demand before and after the inclusion of Greek economy recession. The important aspect of this research is that it uses quantitative methodology to estimate what were the implications in the Greek car market, during a period that the whole economy of the country had shrunk by a quarter and unemployment had stood at more than 25 percent.

Greece is currently under a prolong period of economic crisis. After the implementation of austerity measures Greek citizens postpone or delay the purchase of durable products, like cars and other goods. The great recession in the Greek market started in early 2008. It became even worse after the announcement of the financial agreement of the Greek government with the International Monetary Fund in May 2010 and the austerity measures. Greece negotiated and agreed by signing a memorandum with the International Monetary Fund and the European Central Bank in the early months of 2010 for austerity packages and reforms that last until now and for many years to come. The country’s three main creditors - the European Commission (EC), the European Central Bank (ECB) and the International Monetary Fund (IMF) dictate the terms of the bailout and monitor the application of austerity measures and the reform of public administration, while the country already had received two bailout programs and is expected to start bargaining for the third during this year (2015). Greece is struggling to stay in the eurozone as an equal member.

The economic and social effects of austerity measures were dramatic. The overall increase in the share of population living at “risk of poverty or social exclusion” was not

significant during the first 2 years of the crisis 2009-2010¹ but for 2011 the estimated figure rose sharply above 33% and above 35% by 2013. According to an IMF official, austerity measures have helped Greece bring down its primary deficit in 2011 but as a side-effect they contributed to a worsening of the Greek recession, which began in October 2008 and only became worse in 2010 and 2011. The reality is that by 2012, wages have been cut to the level of late 1990s while purchasing power equals that of 1986. That had an immediate reflection in car sales levels with total new car sales recording historically the lowest level of sales for the last decade.

Our research data in this thesis -monthly sales of new cars in the Greek automobile market- reflect the decline of economic activity covering a wide time period, before and during the economic crisis in Greece. The new car sales data are analyzed and modelled as time series data while conclusions will help us understand, on the one hand, consumer behavior in period of economic crisis, and on the other hand, how car representatives operating in Greece reacted to survive and which of them manage to stay in the market and at what cost.

The main **research questions** that have triggered the applied econometric thesis from the very beginning are:

What is the best econometric model that fit more appropriate to the new passenger car registration levels in the Greek market?

Can we measure the volatility of sales at the Motor Vehicle Passenger Car Segments in Greece?

Is it possible to make safe forecasts intervals for the future sales levels of different Motor vehicle Importers Representatives operating in the Greek market?

What are the impacts in new car sales levels after the implementation of austerity measures in Greece for both consumers and sales representatives?

¹The figure was measured to 27.6% in 2009 and 27.7% in 2010 (and only slightly worse than the EU27-average at 23.4)

Furthermore the **research objectives** of this thesis work was to learn and interpret the business time series data of a specific segment in the Greek market, study simple and more complicated models and methods for analysis of Greek business time series data, understand the proper use and limits of econometric methods in business data, analyze business data from the Greek auto-mobile market scientifically and efficiently, explain the volatility of the new car sales and revile a model that can efficiently measure and predict the new car sales fluctuations in the Greek market, empirically use of econometric models in forecasting future new car sales levels in the market, develop an intervention analysis of car sales data for the evaluation of the impact of the Greek economic crisis on the new car sales levels and understand the impact of austerity measures in Greece in the new car market sector and how they affect the consumer behavior and the sales representatives' decision making.

1.3 Data Source and Used Software

The car registration data are available from the Association of Motor Vehicle Importers Representatives (AMVIR) web site named <http://www.seaa.gr> where SEAA is the Greek abbreviation of the association which is translated as “Syndesmos Eisagogeon Antiprosopon Autokiniton” in the Greek language. The marketing data available at that page in the Statistics refer to the Public cars’(PC) and taxi cars’ registration by month. We take as a fact that the car registration number equals the new car sales number for this research and that is how we will refer to these data from now on. The first year of statistical records of car sales available in the year 1998.

This econometric analysis is completely done by programming based open-source software named as **R** started with R-studio Desktop version 0.98.1102 ended with version 1.2.5001 from “The R Foundation for Statistical Computing”. *R* is my favorite all-around statistics software, free open-source version of *S*, a program developed in the 1970s and 1980s at Bell Laboratories.

R is excellent for graphics, classical statistical modeling, and various non-parametric

methods. The additional library package used for our analysis are *forecast*, *tseries*, *fpp2*, *rugarch*, *fGarch*, *Metrics*, *FinTS*, *stats* and so on, that fit specific classes of models in R.

This thesis report is produced in the **L^AT_EX** document preparation system which is a high-quality typesetting system for the production of technical and scientific documentation from MiKTeX version 2, a trademark of the American Mathematical Society using the latest *TeXStudio* Desktop version at first, and later the *Overleaf*, which is a collaborative cloud-based LaTeX editor used for writing scientific documents that can easily be shared.

1.4 Thesis Overview

The present chapter gives a brief introduction to the subject, followed by the purpose and motivation, research question, and objectives of this study. The data source of the empirical study and the used software are given in detail. The chapter ends with the thesis overview which gives the structure of this research study and the final conclusions.

The *second* chapter refers to the empirical analysis of the time series data set. Firstly an overview of the Greek new car market is presented so that the reader can have a general idea of the market segment before going into further details needed for the empirical study to follow. Visualizing the car sales structure with the use of different graphical presentation helps the time series analysis. Statistical inferences of the data, autocorrelation properties and stationarity tests will complete the data empirical analysis. Given the wide time range covered by the data set and the separated research work that needs to be done, for each one of the 10 different automobiles retail firms operating in Greece, one can understand the complexity of this project.

There is a reference to the related work that has already been done on this subject and a literature review of applied time series analysis to various data.

The Time Series Models are given in a separate section of this first chapter covering the theoretical background of all Time Series models that will be used in the empirical analysis and gives the methodology for model selection. We consider different time series models, simple like the average, naïve and seasonal naïve models and more complicated like Sea-

sonal Autoregressive Integrative Moving Average (SARIMA), the General Autoregressive Conditional Heteroscedastic (GARCH) and the Exponential State Space Smoothing (ETS) models in this research study.

The *third* chapter analyses the *in sample* empirical modeling of the time series data. Time-series methods and models are applied to the set of data and are compared and contrasted for selecting the one that fits best to the data set of each firm. Models like the average, naïve and seasonal naïve, Autoregressive Integrative Moving Average (ARIMA) and Seasonal ARIMA (SARIMA), the family of the General Autoregressive Conditional Heteroscedastic (GARCH) models and a combination of SARIMA - GARCH model and finally the Exponential State Space Smoothing (ETS) models are estimated, evaluated, compared to find which one capture best the monthly movement of each firms' new car sales levels.

The *fourth* chapter considers the forecasting performance of the applied time series models used to fit the data in an *out of sample* estimation process. The time series models used to fit the data are now estimated, evaluated, and compared to find which one can predict best the monthly movement of each firms' new car sales levels in the Greek market. Different forecasting measures are used to evaluate the model with the best forecasting performance. Different time intervals are also used as an intervention analysis, to fit and forecast new car sales levels in an attempt to evaluate how the economic crisis had influenced the new car retail segment in Greece.

The *fifth* chapter of this thesis present a volatility forecasting comparative study within the autoregressive conditional heteroskedasticity class of models. The focus is in identifying if a SARIMA-GARCH model can successful predict the volatility of car sales level during a period of economic crisis in the Greek market. The study proceeds with GARCH model selection and analysis and end up with the diagnostic checking of the model, but before fitting the GARCH model the research try to diagnose if the SARIMA squared residuals are serial correlated and test for ARCH effects. Finally the SARIMA-GARCH (1,1) with student-t distribution was preferred only for Opel and Toyota that show serial correlated SARIMA residuals in a specific data set. SARIMA-GARCH(1,1) model can be an effective

way to improve forecasting accuracy against SARIMA model and to reduce the prediction intervals of the forecast.

The *sixth* chapter of this thesis contains several mathematical data transformations from the family of Box-Cox transformation but also time series modeling using the original data with no transformation at all. Additionally, the 95% confidence Interval in the forecasting is studied in order to find the best data transformation and the best forecasting model for our marketing data.

The *seventh* chapter of this thesis combines the forecasts in a way to improve the forecast accuracy and hedge against risk. We go on from the selection of one single model to the combination of several time series models ending up with a more accurate and less risky model that uses all the available information. We study the simple average and the ordinary least squared combination of the forecasts using equal weights and unequal ones accordingly. The empirical evidence suggests that the Simple average method with equal weights outperforms all the individual time series methods giving the best results.

1.5 List of publications & research presentations.

I am grateful and lucky that I was given the opportunity to present my thesis research with my participation in various conferences and Symposiums and get valuable feedback for the continuation of my studies. Special thanks to the “A.G. Leventis Foundation” for sponsoring, twice, part of my travel expenses to London, for my participation in the Hellenic Observatory Ph.D. Symposium on Contemporary Greece and Cyprus, at London School of Economics and Political Science, European Institute in UK.

The list of publications resulting from my participation in various conferences during this thesis study in chronological order is provided below:

- Voulgaraki K. Maria (2019), “Box Cox transformation in Forecasting sales. Evidence of the Greek market.” in *Proceedings* of the 39th International Symposium on Forecasting (ISF 2019), organized by the International Institute of Forecasters, in Thessaloniki, Greece.

- Voulgaraki K. Maria (2017), “Forecasting sales using switching regime models in the Greek market.” in *Proceedings* of the 8th Biennial Hellenic Observatory Ph.D. Symposium on Contemporary Greece and Cyprus, London School of Economics and Political Science, European Institute, The Hellenic Observatory, in London, UK.
- Voulgaraki K. Maria (2016), “Modeling and Forecasting sales in the Greek market. The case of the Greek new car sales sector.”, in *Proceedings* of the IMAEF 2016, Ioannina Meeting on Applied Economics and Finance, in Corfu, Greece.
- Voulgaraki K. Maria (2014), “Forecasting Sales of durable products in the Greek market. Empirical evidence from the new car retail sector”, *Proceedings* of the Annual Postgraduate Seminar Day (13-04-2014), Department of Economics, University of Crete, in Rethimnon, Greece.
- Voulgaraki K. Maria (2013), “Forecasting sales and intervention analysis of durable products in the Greek market. Empirical evidence from the new car retail sector”, in *Proceedings* of the 6th Biennial Hellenic Observatory Ph.D. Symposium on Contemporary Greece and Cyprus, London School of Economics and Political Science, European Institute, The Hellenic Observatory, in London, UK.
- Voulgaraki K. Maria (2013), “Exponential Smoothing Time Series Forecasts” , *Proceedings* of the Annual Postgraduate Seminar Day (09-01-2013), Department of Economics, University of Crete, in Rethimnon, Greece.

1.6 Research Contribution

This thesis argues that marketing series like car sales in the Greek market can successfully be measured and forecasted using scientific time series methodology and procedures. This study uniquely faces the problem of measuring the performance of car sales in the Greek market in a turbulent period of economy for Greece. Austerity measures due to strict implementation of the various memorandum agreements that affected the Greek economy

and more specific the levels of car sales, where very well reported and measured by time series theory and the use of autoregressive models.

The demonstration of more than seven time series models and the empirical implementation of forecast combination give a thorough presentation of econometric models that can be used in marketing data that can capture their performance and forecast future values. The exhaustive test of data transformation using Box-Cox methodology against the original values gives a better understanding of why transforming the data can always help us in forecasting. The research offers steps towards a better understanding of the variability and complexity of the Greek market and that the more information one can use the better it is. Combine information given by various forecasting models gives better results than using separate forecasting models results.

Time Series Empirical Analysis.

2.1 Introduction

The collection of observations of well-defined items, obtained through repeated measurements over time, like day, month, week, are called “Time Series Data”. In this thesis’s empirical research, the time series data are the *new passengers’ car registration numbers* for the top motor vehicle importers –representatives of the automobile retail sector in Greece. Car registration numbers are well defined and consistently measured at equally spaced intervals, collected regularly, on a monthly base. The researcher has, equivalently, assume that the monthly new passengers’ car registration data series is the monthly *new car sales* levels of each one of the selected motor vehicle importers representatives, operating in the Greek car market. These observations are treated as time series data, and the time series analysis, that will follow, can be considered having characteristics like an analysis done on regular marketing data. Furthermore, the study of the new car sales data encompasses methods for analyzing data, to extract meaningful statistics and other characteristics. It also focuses on a comparative analysis, that points out similarities and differences in the researched time-series data.

Our data analysis begins with an overview of the Greek new Car Market during the last two decades, alone with some crucial economic and political events during that period. Then the various time-series data of this thesis empirical study are presented in more detail.

In the data description section, summary statistics for our data, are presented, like mean, kurtosis, variance, skewness (see Table 2.1 page 26), describing ten (10) different data sets and giving useful statistical inferences for each one of the different time-series samples. The summary of statistics gives all statistical properties of the series along with two tests for normality and randomness in our sample data. The graphical presentation of the data, for this empirical research analysis, is presented by a line and a Box plot of the raw data alone with the autocorrelation function (ACF) and the partial autocorrelation function (PACF) plots (see Figure 2.4, Figure 2.5 and Figure 2.6).

Seasonality, which is the regularly repeated pattern of highs and lows related to calendar time, such as seasons, quarters, months, and so on, is observed in Box plots of our data along with the months of the highest and lowest sale level in a yearly base. Concluding remarks are given in the analysis of the data characteristics of each sample set and their graphical presentation features.

In short the plan of this chapter is the following. In Section 2, a general reference in the overall Greek new car market is discussed, alone with some political and economic events. In Section 3, by the empirical data analysis of the variables are presented which include the data description, summary statistics and graphical presentation of the series. In Section 4 the literature review give more details in the related work that has been done in the field of the Automobile sector and the applied methodology for univariate time series modeling in marketing data (like sales). In Section 5 the research methodology and the univariate time series forecasting methods and models that will be applied are briefly explained with specific focus in their theoretical background and their methodology. In Section 6 there is a short presentation of the forecasting methods with a discussion of the qualitative vs quantitative forecasting methods and techniques. Finally in the last section, there is a short discussion that summarizes the findings of this chapter.

2.2 Overview of the Greek new car market.

The automotive industry around the world is facing several challenges right now, such as economical crisis and tough competition. At the same time demand for alternative technologies is increasing and the legislation on emissions for light-duty vehicles is tightening. The European car market has been a prime causality of the continents' economic crisis as hard-pressed consumers differ purchases and several leading European makers have been forced into radical restructurings. Greece has no heavy industry so there are not any automotive manufactures in the country. All vehicles are imported from different countries and various car representatives are promoting, selling, and delivering various automotive vehicles to motor trade customers. However, Greek economy's new reality, during the last ten years, shows that the Greek society is undergoing a severe pressure due to austerity measures applied in an attempt to control and stabilize the country's enormous debt. The global economic recession has put nations under situations that were difficult to handle or control.

Greece, as a member of the European Union, at the year 2010 asked for the contribution of EU (European Union) the European Central Bank (ECB) and the International Monetary Fund (IMF), and due to that, the Greek government had implemented a sequence of unprecedented austerity measures in an attempt to control the country's enormous debt. The consequences of economic depression, in business and in all aspects of everyday lives, caused uncontrollable destabilization in Greek society. New measures, voted by the Greek government, and new laws were in force. Measures that have not been implemented before in any other EU country and therefore the economic and social consequences had not been effectively calculated or anticipated. The exhaustive austerity measures combined with political instability, change the economic and social reality in Greece. The country was experiencing the biggest economic crisis in its modern history. People had to confront a new way of living based on uncertainty, insecurity, disappointment, lack of trust, unemployment, and disorientation. Greek citizens were called upon to reform the society, the economic activities, their habits, and be reformed by this new reality.

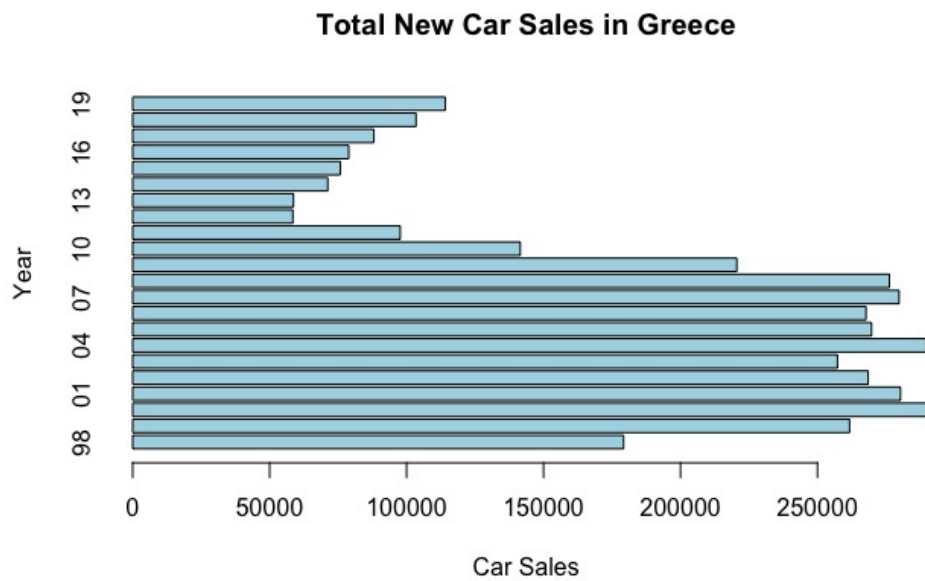


Figure 2.1: Total New Car Sales in Greece 1998-2019 [Bar Chart].

The economic situation in Greece had effected all aspects of economic life. Figure 2.1 (page 16)¹ visually gives the flow of the yearly total new car sales in the Greek market for the last 20 years (from 1998 till 2019). If we divide the period in half we come up with two different realities for the Greek car market, before and after 2009. In the period of the first 10 years (i.e. from 1998 till 2008) the total yearly car sales average is above 250.000 cars per year, while in the second period (i.e. from 2009 till 2019) the total yearly car sales average hardly reach 100.000 cars per year, which is less than half of the first period. The figure shows clearly the dramatic decrease in sales after the year 2010 and how the new reality influenced the sales levels of new cars. In periods of economic recession, like the one Greece was into, the first thing that the consumers do is to prolong or postpone buying durable consumer goods, like cars, furniture, and so on. The new car sales sector continues to experience difficult economic operating conditions. New vehicle registrations provide a measure of confidence and stability in the Greek economy and automobile retail sector from businesses and consumers alike. Low levels of sales revile how deep recession and instability is in the Greek Market even nowadays and how businesses (car representatives)

¹Data from the Association of Motor Vehicle Importers Representatives (AMVIR), <http://www.seaa.gr>

and customers (car buyers) react to this.

In the modern history of Greece there are some special events or elements that had influenced the economic situation of the country and therefore the total new car sales in Greece from 1998 till 2019. To understand and rationalize the flow of the new car sales it is useful to keep in mind what had happened during the period of the last 20 years in Greece. A short presentation follows:

- ★ During **1999-2008** Greek citizens could borrow money with low interest rates.
- ★ In the year **2000**, parliamentary elections were held in Greece in April 2000. The ruling Panhellenic Socialist Movement (PASOK) of Prime Minister Costas Simitis was re-elected and that gave stability to the economy (historical high level of car sales -peak 1).
- ★ In the year **2001**, the drachma, Greece's national currency, was replaced by Euro.
- ★ In the year **2004**, the Olympic Games were held in Greece, and parliamentary elections in March 2004, where New Democracy Party with Prime Minister Kostas Karamanlis won the elections (high level of car sales -peak 2).
- ★ In the year **2007**, parliamentary elections were held in Greece in September 2007 and the New Democracy Party of Kostas Karamanlis was re-elected (high level of car sales -peak 3).
- ★ In the year **2007-2008**, the worldwide financial crisis, also known as the global financial crisis (GFC), started with the depreciation in the subprime mortgage market in the United States, and it developed into an international banking crisis and a severe worldwide economic crisis (car sales starts to fall since car-loans were severely restricted).
- ★ In the year **2010**, parliamentary elections in Greece in November 2010 and the Panhellenic Socialist Movement with George Papandreou as Prime Minister was elected. Greece was in danger of bankruptcy. The first (1st) Memorandum of Understanding

with the “TROIKA” of foreign creditors — the European Commission, the European Central Bank, and the International Monetary Fund — was signed over the terms of a bailout agreement (car sales fall dramatically).

- ★ In the year **2012**, in March, the second (2nd) bailout package or the second memorandum was agreed. Parliamentary elections in Greece on June 2012 and Greek voters gave a narrow victory to Antonis Samaras, the leader of the New Democracy party (car sales reached a historical bottom record).
- ★ In the year **2015**, two snap parliamentary elections in Greece in January and September. Greece’s exhausted voters seem to have lost patience with the traditional parties of power and gave the victory for Alexis Tsipras’ Coalition of the Radical Left (SYRIZA) Party that promised to tear up bailout agreements that had created a “humanitarian crisis”. On 5 June 2015 the “Greek bailout referendum” was held where citizens had to decide whether Greece should accept the bailout conditions in the country’s government-debt crisis proposed. As a result, the bailout conditions were rejected by a majority of over 61% who voted “NO”. Despite that in July the third (3rd) bailout agreement/memorandum, was signed (car sales stay low but with an upward trend).
- ★ In the year **2019**, parliamentary elections in Greece in May where the liberal-conservative political party New Democracy and its new leader Kyriakos Mitsotakis won power over Alexis Tsipras’ Coalition of the Radical Left party (SYRIZA). This was largely been seen as a battle for the middle-class which was severely impacted by austerity measures following the country’s near-bankruptcy and assistance from international creditors (car sales develop a stable upward trend).

During the period of interest 1998-2019 it seems that there were three (3) peaks in 2000, 2004 and 2007 and then a down turning point in 2012 where car sales reached the lowest sales level of this period. In more detail, since the worldwide financial crisis in 2007-2008, the new passenger vehicle sales in Greece had a downward trend and reach the bottom line

in 2012. After that, the market seemed to regain some hope and a small growth of only 0.4% had occurred in 2013 car sales which were of course far from the historical peak of sale that happened in 2000.

The Greek car market was fueled by *low interest rates* courtesy of the euro, from 1999 to 2008. The lower the interest rate, the more willing people are to borrow money to make big purchases, such as a new car. During these years the Greek new car market exceeded 250.000 new car sales per year. That is a big number of small countries like Greece. Since the country has a population of around 10 million people the level of new car sales meant that Greeks bought one new car every 40 citizens each year. It is too hard to keep that pace or even increase it. The Greek statistic organization², reported at the Greek census of 2011 that :

30,4% of the total registered household had no car,

45,5% had one car

20,3% had two cars, and

3,8% had more than two cars.

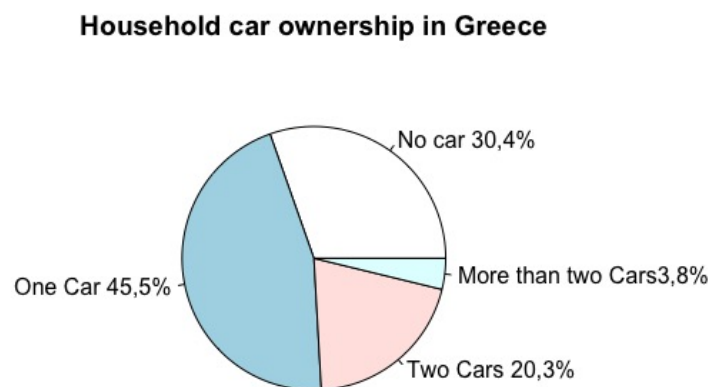


Figure 2.2: Household Car Ownership in Greece [2011 Census].

²ELSTAT

Furthermore, at the end of 2008, after the begging of the economic crisis, when the cheap credit ran out, the Greek car market crashed fast and deep. Car representatives' managers had to make budget cuts, sales personnel redundancies, or close down branches all around the country. The collapse of the Greek automobile market reflected the dire state of the Greek economy as a whole.

The new-car market pie was divided by various brand leaders in the Greek market. Toyota was at the top sales position as the bestselling car brand in Greece for many years and a leader ahead of other German brands like Volkswagen and Opel. On the other hand, the Japanese car-makers like Toyota, Mazda, Suzuki and Mitsubishi, all managed to grow in European countries, helped by a soft yen, new smaller diesel engines in their line-up and new models tailored to European tastes.

The market share of each firm's new car sales level in the Greek market change over the years, however, the differences are usually small. Graphically the sales of each firm, as a percentage of the total number of new car sales, are illustrated in Figure 2.3 as a pie chart, for the year 2011.

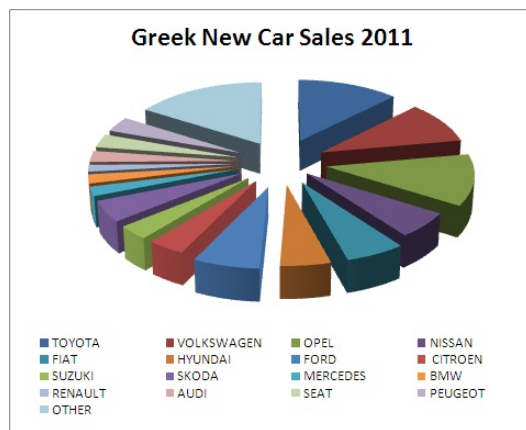


Figure 2.3: Market Share for Greek new-cars [Pie Chart].

In general, more than 35 different firms were operating in the Greek market and the market share was divided into many small players. Only 2 to 3 firms have each a market share around 10%, 6 to 8 firms have each around 5%, 7 to 8 firms have each less than 3%,

5 to 6 firms have each less than 0,5% and the rest have less than 0.1% market share during these years of research. Furthermore, sales activity can be divided into three different groups. The first group represent about 40% of total sales i.e. 1/3 of the total share and is dominated by the leading car representatives consisting of not more than 3 firms that have each around 10% of the total sales in the market (usually the players are Toyota, Opel, and Volkswagen). The second group manages to have only half the share of the first group i.e. around 20% of the total sales and consists of 6 to 8 firms with a share of less than 5%. The last group covers the share of 40% of the markets and includes many firms with small market share (3% to 1%). Throughout the years this taxonomy is not changing much.

2.3 Empirical Data Analysis.

This empirical analysis time series data includes a group of 10 different sample data sets. The sample data sets are comprised of monthly new car registration in Greece from 10 different firms. The time starts from January 1998 till December 2016, which gives a total of 228 observations for each one of the ten (10) data samples. These data sets are new car registrations from the largest Motor Vehicle importers-representatives operating in the Greek automobile retail market. The data are obtained from the Greek Motor Vehicle Importers - Representatives (AMVIR)³ statistical database. Data are the passenger car registration numbers for the last 19 years as officially recorded from the Greek authorities. More than 35 different passenger car representatives were operating in Greece but for this research, only a group of the top 10 firms operating in the Greek market was chosen. The criterion was the sales level of each firm operating in the Greek market in terms of cars registered ending up to a blend of leading companies operating in the Greek market.

The sample of the 10 leading companies of this retail market is a sample of more than 70 % of the total retail sector of the new car market. The pie chart representation of new car market sales in Greece for the year 2011, shows that the research sample covers more than 3/4 of the total market activity.

³<https://www.seaa.gr/el/statistics/registrations>

2.3.1 Descriptive Statistics.

In this section, a brief description of the 10 data sets is given, which represents a sample of the total new car sales in Greece. Descriptive statistics that summarize the data sets are broken down into measures of central tendency and measures of variability (spread) of each data set. The descriptive (summary) statistics for the monthly new-car sales series are given in original (raw) data and the logarithmic transformation of the original data. The mean, the median the maximum and minimum value of the series, the standard deviation, the skewness, and the kurtosis are given.

Under the time series analysis, the researcher is interested in the distributional properties of each series, which can be defined by the moments of each random variable. Generally, there are four central moments, which are important when examining a time series behavior:

1. The **first** moment is called the *mean* or *expectation* of the series and it measures the central location of the distribution μ or $\bar{x}$.
2. The **second** central moment measures the variability of the series and is called the *variance* of the series denoted as σ^2 . The positive square root of the variance is the *standard deviation* of the series.
3. The **third** central moment measures the symmetry of the series concerning the mean and is called *skewness*(S).
4. The **fourth** central moment measures the tail behavior of the series and is called *kurtosis*(K).

The first two moments of a random variable uniquely determine a normal distribution (i.e. $\bar{x} = 0$, $\sigma^2 = 1$) while for other distribution higher-order moments should be additionally considered as well. The third and fourth central moments of x are often used to summarize the extent of asymmetry and tail thickness. In a normal distribution, skewness equals zero ($S=0$), which indicates that the values are relatively evenly distributed on both sides of the mean, and kurtosis equals three ($K=3$), which is presented graphically as a bell-shaped diagram.

Furthermore a distribution with *negative skewness* ($S < 0$) indicates that the tail on the left side of the probability density function is longer than the right side and the bulk of the values (possibly including the median) lie to the right of the mean. On the opposite, a distribution with *positive skewness* ($S > 0$) indicates that the tail on the right side is longer than the left side and the bulk of the values lie to the left of the mean.

A distribution with *kurtosis more than three* ($K > 3$) is called *leptokurtic* and is said to have heavy tails, implying that the distribution puts more mass on the tails of its support than a normal distribution does. We tend to see kurtosis more than three in series where their distribution contains more extreme values. On the other hand, a distribution with *kurtosis less than three* ($K < 3$) is called *platykurtic* and is said to have short tails, which gives a uniform distribution over a finite interval.

Additionally, this study is interested to test the normality of the series, which means to see whether the sample series are normally distributed. This is important for the errors of a modelling process. The two popular formal statistical tests for normality that are referred in this research are: the *Shapiro and Wilk's test* (SW) and the *Jarque-Bera* (JB) tests.

The **Shapiro and Wilk's test** is a well known goodness of fit test for the normal distribution. It tests the null hypothesis that the sample $x_1, x_2, \dots, x_n$ come from a normal distributed population. It was published in 1965 by Samuel Shapiro and Martin Wilk.

The test statistic is

$$W = \frac{(\sum_{i=1}^n \alpha_i x_{(i)})^2}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

where $x_{(i)}$ (with parenthesis enclosing the subscript index i) is the i^{th} order statistic i.e. the i th smallest number in the sample; $\bar{x} = (x_1 + x_2 + \dots + x_n)/n$ is the sample mean, the constants α_i are given by $(\alpha_1, \alpha_2, \dots, \alpha_n) = \frac{m^T V^{-1}}{(m^T V^{-1} V^{-1} m)^{1/2}}$, where $m = (m_1, \dots, m_n)^T$ and $m_1, \dots, m_n$ are the expected values of the order statistics of independent and identically distributed random variables sampled from the standard normal distribution at time T , and V is the covariance matrix of those order statistics. The user may reject the null hypothesis if the value of the statistic W is too small.

The test of normality can additionally be interpreted via a Q-Q plot which can tell a lot

about the distribution non-normality and may point to solutions. On the other hand, the Shapiro-Wilk test is sample size-dependent and it does not indicate the degree of deviation from normality directly. Therefore in a small sample, someone almost always concludes normality and in a large enough sample, even a tiny deviation from normality will be significant.

The **Jarque and Bera test** (JB) is a goodness-of-fit test of the sample skewness($\hat{S}$) and sample kurtosis ($\hat{K}$) matching a normality distribution. It is based on the result that a normally distributed random variable has skewness equal to zero and kurtosis equal to three. The null hypothesis is a joint hypothesis of the skewness being zero (0) and the kurtosis equals three (3) (or the excess kurtosis being zero). For the normal distribution of x_t series, we get $S=0$, $K=3$, and $JB=0$, and any deviation from this increase the JB test.

The test is defined as:

$$JB = \frac{n}{6}(\hat{S}^2 + \frac{1}{4}(\hat{K} - 3)^2) \quad (2.1)$$

where n is the number of observations, $\hat{S}$ is the sample skewness, $\hat{K}$ is the sample kurtosis.

JB test is asymptotically distributed as a chi-squared random variable with 2 degrees of freedom, to test for the normality of x_t .

$H_0 : S = 0$ and $K=3$ against the alternative

$H_\epsilon : S \neq 0$ and $K \neq 3$ with error α .

Reject H_0 if $JB < x_{2,\alpha}^2$, in other words, reject the null hypothesis of normality if the p-value of the JB statistic is less than the significance level(α) [Tsay, 2005].

The main disadvantage of the JB test is the over-rejection of normality in the presence of serially correlated observations [Thomakos and Wang, 2003]. However, Thadewald and Buning [2007] provided simulated evidence that the JB test, generally works best in comparison to several alternative normality tests but for some cases, other tests should be preferred.

2.3.2 Summary statistics of Data

The summary statistics explain each one of the 10 data sets of the retail firms of the Greek market covering the period from January 1998 to December 2016 and are illustrated in Table 2.1 (page 26)⁴. In general, the first two moments (the mean and the variance) is said to determine the normal distribution. The last two moments (kurtosis and skewness) are used to summarize the extent of the asymmetry and tail thickness of the series. For a time series variable with normal distribution, the mean equals the median. In more detail, the median is the numerical value separating the higher half from the lower half of the observations. It can be found easy by simply arranging all the observations from the lowest to the highest value and picking the middle one. In our study case, we have an even number of observations so we take the two middle values, which corresponds to interpreting the median as the fully trimmed mid-range. More specifically in this empirical research sample series, the mean and the median appear to be close enough to each other but still do have a small difference which gives a hind for normality but also a small deviation from it as well (Table 2.1, 2nd & 3rd columns) which makes it difficult to decide whether we have a normal distribution or not.

Max represent the maximum value and *Min* the minimum value of the monthly new-car sales level (Table 2.1, 4th & 5th columns). By subtracting the minimum from the maximum value gives a new number, which is the monthly *range* of each firms extreme new car sales numbers (Table 2.1 6th column). The bigger that number, the wider the range of the extreme monthly sales of each firm. Notice that the higher range of the 2 extreme monthly car sales levels is in Toyota (3.640 cars) followed by Hyundai (3.059 cars) and Fiat (3.043 cars) while the Skoda (1.457 cars) and Peugeot (2.256 cars) have the lower range (in original data). This indicates how sharp the fall in sales was in each one of the firms and how wide was the sample of each firm.

Skewness(S), in Table 2.1 (8th) column, is positive and ground zero for the original data, but not equal to zero, as in normal distribution. It is giving a positive result at a range

⁴Note: SD=Standard Deviation, S=Skewness, K=Kurtosis, JB=Jarque Bera test, SW=Shapiro Wilk's test, p-values in brackets.

Statistics	Mean	Median	Max	Min	Range	SD	S	K	JB(p-value)	SW(p-value)
Monthly Original New-Car Sales										
Opel	1.357	1.293	3.154	263	2.891	738	0,39	-0,80	3,73(0,15)	0,98(0,04)
Toyota	1.577	1.612	3.938	234	3.704	839	0,34	-0,71	0,54(0,76)	0,98(0,19)
Volkswagen	1.473	1.497	2.634	219	2.415	492	0,06	-0,21	0,32(0,85)	0,99(0,31)
Hyundai	1.617	1.655	3.292	233	3.059	731	0,07	-0,75	3,70(0,15)	0,98(0,03)
Peugeot	1.062	1.089	2.369	113	2.256	506	0,15	-0,59	2,81(0,24)	0,98(0,02)
Ford	1.143	1.111	2.448	136	2.312	536	0,29	-0,84	6,78(0,03)	0,97(0,00)
Nissan	998	961	2.630	85	2.545	378	0,54	1,49	24,71(0,00)	0,97(0,00)
Fiat	996	880	3.260	94	3.260	670	0,86	0,19	17,84(0,00)	0,95(2,391e-05)
Skoda	610	631	1.524	67	1.457	302	0,27	-0,23	2,29(0,31)	0,98(0,03)
Citroen	1136	1101	2.729	88	2.641	577	0,28	-0,44	3,30(0,19)	0,97(0,00)
Monthly Log New-Car Sales										
Opel	7,34	7,42	8,06	5,95	2,11	0,42	-0,63	0,06	11,04(0,003)	0,96(0,00)
Toyota	7,48	7,55	8,28	5,67	2,61	0,44	-1,07	1,35	45,42(1,367e-10)	0,93(3,452e-07)
Volkswagen	7,23	7,31	7,88	5,39	2,49	0,4	-1,38	3,51	142,07(< 2,2e-16)	0,91(1,592e-8)
Hyundai	7,26	7,41	8,10	5,45	2,65	0,57	-0,95	0,27	26,10 (2,14e-06)	0,92(8,867e-08)
Peugeot	6,82	6,99	7,77	4,73	3,04	0,61	-1,01	0,51	30,76 (2,089e-07)	0,91(3,843e-08)
Ford	6,91	7,01	7,80	4,91	2,89	0,55	-0,74	0,27	16,08 (0,0003)	0,95(2,976e-05)
Nissan	6,82	6,87	7,87	4,44	3,43	0,46	-1,50	4,46	205,11 (< 2,2e-16)	0,90(4,913e-09)
Fiat	7,01	7,04	8,09	5,38	2,71	0,52	-0,37	-0,17	4,04 (0,132)	0,98(0,16)
Skoda	6,25	6,45	7,33	4,20	3,13	0,65	-1,11	0,74	38,71 (3,927e-09)	0,90(7,258e-09)
Citroen	6,86	7,00	7,91	4,48	3,43	0,66	-0,99	0,46	29,16 (4,649e-07)	0,91(3,95e-08)

Table 2.1: Descriptive Statistics for original and log values of monthly new-car sales of 10 firms.

between 0,06 to 0,79 in all sample series. This is indicating that we have an elongated tail at the right which means more data in the right tail than would be expected in a normal distribution. So the original series is skewed to the right, which implies that the left tail of the distribution is fatter than the right tail. On the other hand, the skewness becomes negative if we take the log values of each series indicating the opposite conclusion for the log values distribution. In the log values graph, there is an elongated tail at the left, which means more data in the left tail of the distribution, and that makes the right tail fatter than the left. So if we take the log transformation of the series, the log series is skewed to the left, while the original series is skewed to the right.

Kurtosis(K), in Table 2.1 (9th column), in original new-car sales series is much less than three (3) in all the sample series, indicating that the series have short tails and can be called platykurtic. (Notice: Results show kurtosis and not the excess kurtosis i.e. $kurtosis - 3$). Transforming the data, by taking their log values, increases the kurtosis but still keeps it around zero quite lower than 3. There are only two cases the opposite result (i.e. $(K > 3)$) which indicate distributions with heavy tails: Volkswagen ($K = 3,51$) and Nissan ($K = 4,46$).

Furthermore, the values of the Jarque-Bera test for normality of the original new-car sales values, rejects the null hypothesis that series are normality distributed only for the sales series of Ford, Nissan and Fiat, because their J-B statistics are high numbers away from zero (Table 2.1, 10th column) and their J-B statistic p-values are less than the 5% significance level. In all the rest series (i.e. Opel, Toyota, Hyundai, Peugeot, Volkswagen, Citroen, and Skoda) we can assume that they are normally distributed. However, in the log values of the same series, the results are different. All new-car sales series reject that the series are normality distributed due to a very high value for each one of the J-B statistics while the corresponding p-values are much less than the 5% level of confidence.

The Shapiro Wilk's test p-values (Table 2.1, 11th column) are less than $\alpha = 0,05$ (for a 95% confidence level) in case of Opel, Hyundai, Peugeot, Ford, Nissan, Fiat, Skoda and Citroen for original new-car sales series, which means that data seems to deviate from normal distribution i.e. reject Normality. On the other hand, the Shapiro Wilk's test for

the log values of the series indicates that all the series are rejecting normality. So for the original values of the series and the log-values of the series the conclusions is the same using either test.

2.3.3 Graphical Presentation.

Raw data are graphically presented in four different panels which present the time-series line plots, the Box plots, the time series Autocorrelation Function (ACF) plots and the time series Partial Autocorrelation Function (PACF) plots for three different Greek vehicles representatives for the sake of a concise presentation (Figure 2.5 on page 36, Figure 2.4 on page 35, Figure 2.6 on page 37).

The statistical graphs can summarize our data in a more readable and understandable way, for economists and non-economists. As many people say “a picture is worth a thousand words”. Therefore we consider the graphical presentation valuable for our analysis since it can indicate some very important features, useful in our analysis. Numerous graphical techniques can be used, but we will focus on just a few time series plots, to show how the variables change over time and what are the data characteristics. Firstly, we illustrate the line plots of the new car sales level that show the movements and fluctuations of sales over time and then we present the Box Plots of monthly sales of each firm in raw data. Secondly, we present the plots of Autocorrelation Function (ACF) and Partial Autocorrelation Function (PACF) of each sample data variable alone with the Ljung-Box autocorrelation test and finally the study gives the conclusions of the graphical presentation.

The line plot of the series illustrates the turbulent economic environment and some signs of seasonality. That gives us a hind that the variance of data may fluctuate and not be stable during the period that we examine. However, it is hard to identify a clear upward or downward trend throughout the years. Trend, in general, is the movements of sales in a prolonged period when they are rising or falling faster than their historical average. There is no consistent trend (upward or downward) over the entire time for none of the sample time series. The series appears to slowly wander up and down. However, it is identified a short-term trend if we split time into smaller segments since there is no evidence of a

clear long term direction (trend) of the series in the whole period. For example, generally speaking for all car representatives we have a slightly upward trend for the years 1998 till 2002 and a downward movement for the period 2008-2010, with a dramatic dropping trend till the year 2012 due to the economic crisis that affects all fields of economic life. After the year 2012 where sales reached a historically low level of sales, there was a slow upward trend and the market found a new equilibrium at a level about half of what it was during the years 1998-2008.

Box plots are a convenient way of graphically depicting groups of numerical data through their quartiles. They are useful for identifying outliers (which are observations far away from other data) and for comparing distributions. The Box plots display differences between different variables without making any assumptions about their statistical distribution. The different spacing in the parts of the box helps us to visually identify the degree of dispersion in other words the spread and skewness in the data and also identify outliers.

Additionally, marketing data like sales always exhibit calendar effects. The Box Plots graphical representation of the series gives clear evidence of seasonality. These plots support the seasonality of the data, reveals the underlined seasonal pattern, and indicate the customers' attitude in a yearly effect. In Box Plot figures observations for each month are collected together overall years in research and are presented in one time plot, as a separate time series for each of the firms. In other words, the January sales level for one car representative is collected for all years (1998-2016) and the horizontal line in each month is the average sales level for the sample period. Thus, the underlying pattern is seen and it enables us to visualize the seasonality over time. Furthermore, it is easy to identify the months with high or low sales levels throughout the year. According to these plots, the month with the highest sales level is *January* for all car representatives and July is the second month with the highest sales on a yearly base. Generally speaking for all firms autumn and winter months have low new-car sales levels. December, September, and August are months with quite low sales levels for all firms.

December, the last month of the year, has the lowest sales level. One reason for this

is that customers try to avoid the taxation claim⁵ placed by the Government, annual to all vehicles. It is due for each car yearly, on a calendar base. This tax bill depends on the engine power of the car and must be paid to the local Independent Authority for the Public Revenue office from the owner of the car. It is valid from January till December of each year and it is due to be legally be permitted to drive the car in Greece for the current calendar year. If one purchases a new car in December this tax bill should be paid for the year ending in 31st of December, while in January the same amount should be paid again for the coming year. Therefore a lot of potential new car buyers prefer to delay their purchase and make it in January instead of December. Customers try to avoid the yearly taxation amount that car owners have to pay to the Greek authorities by postponing their buy until January, so instead of December they buy in January and pay the yearly taxation only once, for the new coming year. This produces low sales in Autumn especially in December and the high sales level in January. The second month with low sales level is *September* where the school starts and people are focusing more on their family and their children and the families' educational expenditures rather than in purchasing durable products, like cars. *August* is the third month with a low sales level which is also the last month of summer. This comes not surprisingly mainly due to two reasons: Firstly the climatological circumstances. August, as the last month of summer, is the hottest month in Greece, and the period where consumers are on holiday. Secondly, it is the month where car manufactures are lowering their production or even closing their factories for holiday reasons. Car representatives are closing or work with a limited workforce since the demand is not high and the staff needs to take some days off so we notice a high level of holiday absence.

The most common feature exhibited by actual time series is the fact that the observations are not independent and this is a key point when research uses the previous history of the series to make predictions in particular conditional distributions on the past. One important tool for assessing the degree of dependence in observed data is the sample autocorrelation function (sample ACF) of the data.

⁵called Teli Kykloforias, in Greek

The autocovariance is the covariance⁶ of the variable with itself which actually means that it is the variance of the variable against a time-shifted version of itself. For a time series, the covariance and the autocorrelation functions are needed to model the dependence over an infinite number of random variables.

The sample autocorrelation function (ACF) for new-car sales time series, gives correlations between the x_t series and its lagged values⁷ In time series analysis a lag is defined as an event occurring at time $t+k$ where $k > 0$ and it is said to lag behind an event occurring at time t , the extent of the lag being k . In 1970, Box and Jenkins wrote, "...to obtain a useful estimate of the autocorrelation function, we would need at least 50 observations and the estimated autocorrelations would be calculated for $k = 0, 1, \dots, k$ lags, where k was not larger than $n/4$ ", where $n/4$ is the sum of observations divided by four. The ACF⁸ can be used to identify the possible structure of any time series data, because it gives the correlations between x_t and x_{t-1}, x_t and x_{t-2} and so on. In other words, it illustrates the similarity between observations as a function of the time separation between them. The only disadvantage is that there is often not one single clear-cut interpretation in the sample autocorrelation function.

In empirical research, the autocorrelation graph of the residuals is often used. For example, to detect seasonality, the researcher plot the autocorrelation function (ACF) by calculating and graphing the residuals (observed values minus mean for each data point). The graph of the residuals against a specified time interval is called a lagged autocorrelation function or a correlogram figure.

The null hypothesis for the ACF is that the time-series observations are not correlated to one another, which means that any pattern in the data is from random shocks only. When we plot the sample ACF of a model's residuals, the ideal is not to find any significant

⁶Sample *Autocovariance* Function for x_t time series with $\bar{x} = \frac{1}{n} \sum_{t=1}^n x_t$ and $Var(x_t) < \infty$ is: $\hat{\gamma}(h) = \frac{1}{n} \sum_{t=1}^{n-|h|} (x_{t+|h|} - \bar{x})(x_t - \bar{x})$, $-n < h < n$ where $\gamma_x(h) = Cov(x_{t+h}, x_t) = E[(x_{t+h} - \bar{x}_{t+h})(x_t - \bar{x}_t)] = E[x_{t+h}x_t] - \bar{x}_{t+h}\bar{x}_t$

⁷for lags of 1, 2, 3, ... k the lagged values are $x_{t-1}, x_{t-2}, x_{t-3}, \dots x_{t-k}$.

⁸Sample *autocorrelation* function (ACF) is: $\hat{\rho}(h) = \frac{\hat{\gamma}(h)}{\hat{\gamma}(0)}$, $-n < h < n$

correlations for any lag. As a rule of thumb for an ACF, if there are plotted residuals that are greater than 2 standard errors away from the zero mean, that indicates statistically significant autocorrelation. A given autocorrelation coefficient is classed as significant, if it is outside a $\pm 1.96 * 1/(n)^{1/2}$ band where n is the number of observations.

The partial autocorrelation function (PACF) is used to detect trends and seasonality. In general, PACF is the amount of correlation between a variable and its lag that is not explained by correlations at all lower-order lags. The partial autocorrelations at all lags can be computed by fitting a succession of autoregressive models with increasing numbers of lags. In particular, the partial autocorrelation at lag k is equal to the estimated coefficient in an autoregressive model (AR) with k terms noted as $AR(k)$. The $AR(k)$ is a multiple regression model in which x is regressed on $LAG(x,1)$, $LAG(x,2)$ and so on, up to $LAG(x,k)$. Thus, by mere inspection of the PACF we can determine how many autoregressive terms we need to use for the explanation of the autocorrelation pattern in our time series. If the partial autocorrelation is significant at lag k and not significant at any higher-order lags, for example, if the PACF cuts off at lag k , then we should try fitting an autoregressive model of order k for stationary series.

In ACF Plots there are 3 significant autocorrelation coefficients commonly observed in all series at lag 12, 24, and lag 36, which lay more than 2 standard errors (which is the approximate 95% confidence limits) from the zero mean. We interpret this as a twelve-month seasonal pattern that cycles yearly. So graphical presentation of ACF reveals the existence of seasonality in our sample data. The characteristics of the ACF and PACF of the 10 series of our research sample tend to show a strong peak at $k=12,24,36$ in the autocorrelation function. The fluctuations, around the vertical axis of zero means, indicate that there is a cycle of about 12 months in the new-car sales series indicating strong seasonality. There are of course many other smaller or higher peaks appearing at different lags in the autocorrelation function and the partial autocorrelation function. This may suggest the order of either a seasonal moving average (MA) or a seasonal autoregressive model (AR) or perhaps both, but we need to make sure first that our series is stationary before making any suggestions for their coefficient order. The “suspension bridge” pattern

in the ACF is typical of a series that is both non-stationary and strongly seasonal. In figures 2.4 on page 35 we have the ACF and the PACF of the sample time series plots of the Toyota raw series and Opel in Figure 2.5 page 36), which are obtained by plotting the “residuals” of an $ARIMA(0,0,0) \times (0,0,0)$ model with a constant: the autocorrelation plots indicate that the series may not be stationary, because the correlations die down very slowly in most of the cases.

Ljung–Box Autocorrelation Test

x_i	Q_{LB}	p-value
Opel	172,98*	$< 2, 2e - 16$
Toyota	96,99*	$< 2, 2e - 16$
Volkswagen	164,31*	$< 2, 2e - 16$
Hyundai	368,76*	$< 2, 2e - 16$
Peugeot	635,34*	$< 2, 2e - 16$
Ford	609,92*	$< 2, 2e - 16$
Nissan	148,47*	$< 2, 2e - 16$
Fiat	476,12*	$< 2, 2e - 16$
Skoda	484,12*	$< 2, 2e - 16$
Citroen	474,48*	$< 2, 2e - 16$

Table 2.2: Ljung-Box Test for car sales series (original data with $lag = 12$)

Note:*Original Series autocorrelated

Autocorrelation plots are one common method for testing the randomness of the series but the researcher additionally applies the Ljung–Box statistic test in the data. Ljung and Box’s test for autocorrelations of the series is based on the autocorrelation plot, but instead of testing randomness at each distinct lag, it tests the overall randomness based on several lags [Bowerman et al., 2005]. The tested hypothesis is H_o : all correlation coefficients up to lag 1 are 0 (i.e. data are random) and the alternative H_e : not all lags up to 1 are 0 (i.e.

data are not random)

The statistic for this test is $Q_{LB} = n(n+2) \sum_{\kappa=1}^m (n-\kappa)^{-1} r_{\kappa}^2$, where n is the sample size, r_{κ} is the autocorrelation at lag κ . H_o of randomness is rejected if $Q_{LB} > x_{\alpha^2, h}^2$, where x_{α}^2 denotes the $100(1 - \alpha)^{th}$ percentile of chi-squared distribution with m degrees of freedom and h is the number of lags being tested.

Ljung-Box Test in Table 2.2 on page 33 illustrates the results of our Ljung Box test of the car series raw data of each car representative. It rejects the null hypothesis of *no autocorrelation* at the 1 % level for all numbers of lags considered in each firm (lags=12). The p-values are so small ($P < 2, 2e - 16$), this means ($P < 2, 2 * 10 - 16$) i.e. equals 0,000000000000000022 which is effectively close to zero (actually numerically undistinguishable from 0) and therefore we reject the null hypothesis. Hence the results from the Ljung and Box's autocorrelation test give evidence that *original data series are not random* which means that they are *autocorrelated*.

Autocorrelation refers to the correlation of the time series with its past and future values and there is significant evidence that there is autocorrelation between the observation of each series. It is common to always use statistical tests to ensure that the graphical implied conclusions are correct. This test is not telling us whether the series is stationary or not and it is not adding much information, except that it is just confirming that there is some correlation in the data as illustrated in the ACF and PACF plots. However, this test will be much more useful later on, to test whether residuals of the estimated and selected models are correlated or not.

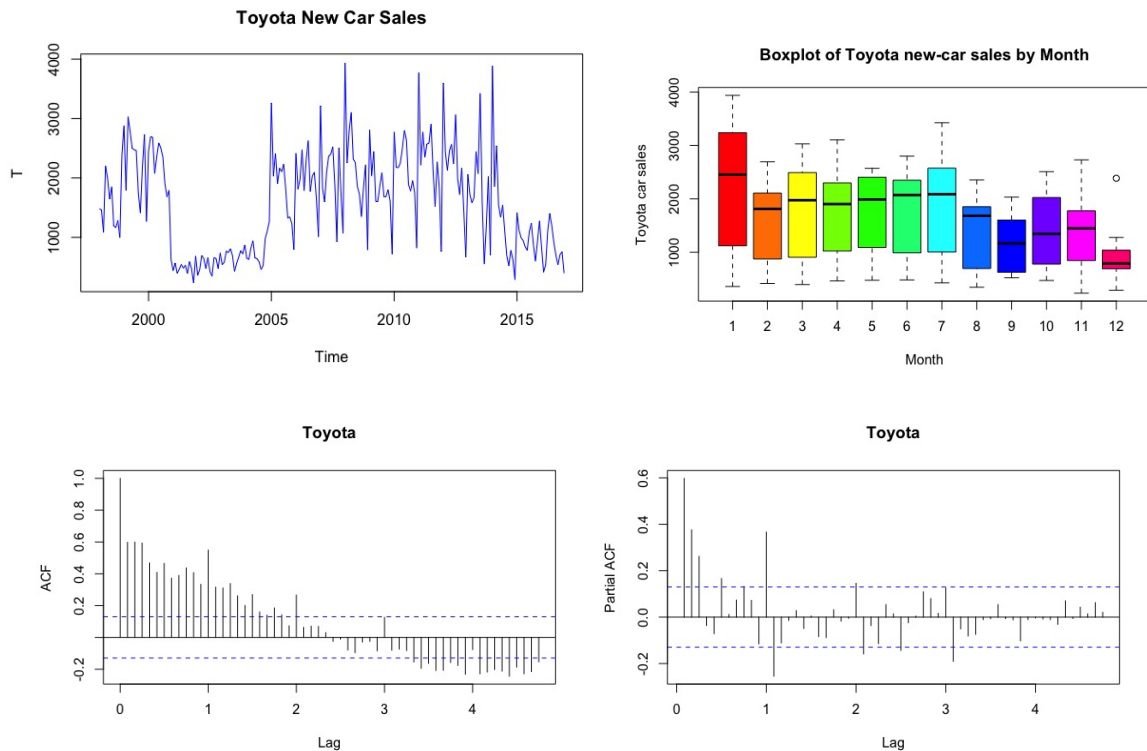


Figure 2.4: **Toyota** new-car sales. Panel (a) Line Plot, (b) Box Plot, (c) ACF, (d) PACF

Toyota line plot shows a rapid increase in their new car sales level for 2 years (1998-2000) and then the firm seems to maintain a stable high market share in sales with very small upward and downward movements throughout the years up until 2010. No extreme values are observed throughout the years. However, after 2010 the scenery change. A rapid downfall of the Toyota market share in new car sales seems to force the company to a much lower level of sales. *Seasonality* plot indicates the yearly seasonal movement of the sales with the highest sales occurring in January. The *random* plot indicate the change in the sales level around 1999 with the sharp increase and around 2009 with the dramatic fall (Figure 2.4 page 35).

Volkswagen car sales, had an extended period of rapid growth in new car sales from 1998 till 2002. After 2002 the firm generally maintained its market share at the same level for the next seven to eight years with some very small fluctuations. However, after 2020 there were high fluctuations which were caused by a dramatic decrease in sales that reached

very low levels.

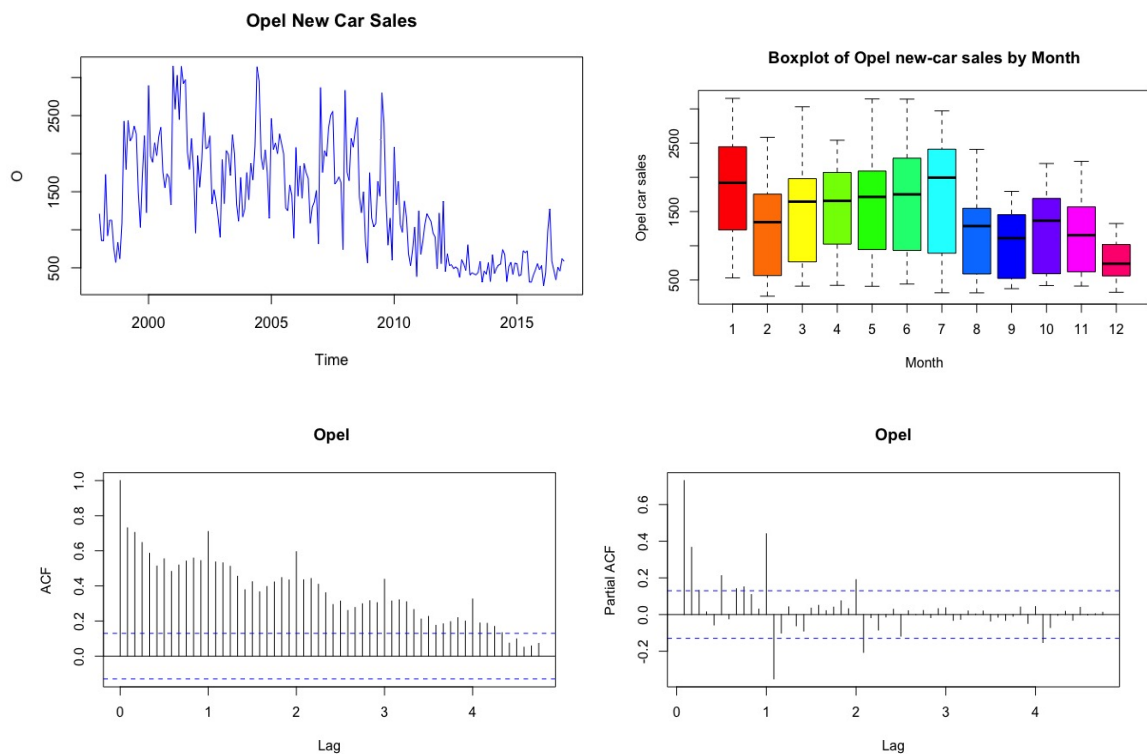


Figure 2.5: **Opel** new-car sales. Panel(a)Line Plot, (b)Box Plot (c) ACF, (d)PACF

Opel new car sales line plot during the period 1998 till 2016 indicates four picks each one occurred in January of 2001, 2005, 2007, and 2010. Sales are increasing rapidly from 1998 until 2001 but then started falling until 2004. The company managed to regain its market share in 2005 but sales decrease again. In 2007 the company managed to regain sales level at a much lower level but at a stable pace and kept it for the following two years. Unfortunately after 2009, the company started a downfall movement on its sales in the Greek market that became dramatic during 2011. The high seasonality of the series is obvious in the Box plot with the same pattern repeated on a yearly base. The Box plot gives evidence of big fluctuations in the years 1999, 2000, 2002, 2004, and 2010 but in the meanwhile fluctuations were more soft and stable (Figure 2.5 page 36).

Skoda seems to have an extended period with rapid growth in the new car sales market from 1998 till 2001. Additionally, the firm maintained its market share and surprisingly

increase it over the years until the beginning of 2010. During 2010 and 2011 sales decrease only slightly. This firm seems to be one of the favorite firms in periods of economic crisis among consumers. That is mainly due to the low purchasing values of Skoda cars and the perceived good value for money.

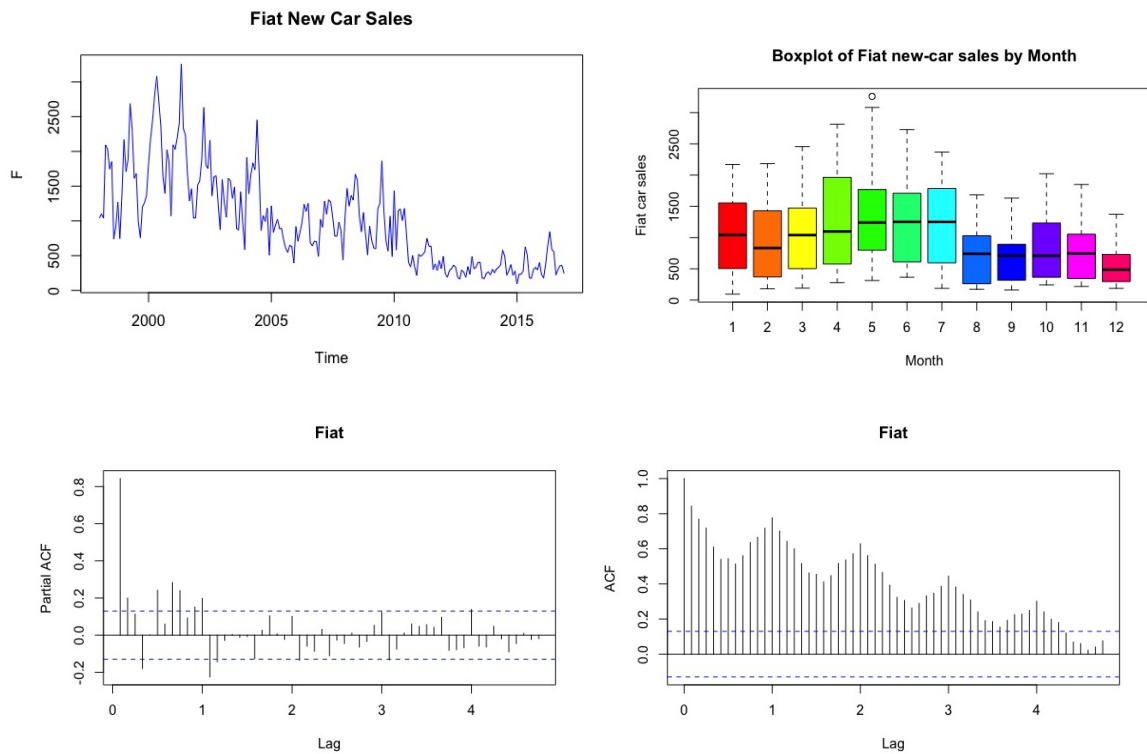


Figure 2.6: **Fiat** new-car sales. Panel(a)Line Plot, (b)Box Plot,(c)ACF, (d)PACF

Fiat had a good market share in the Greek market with a slightly increasing trend of its car sales level until 2002. Unfortunately, Fiat did not maintain its market share, which gradually decreased until 2005 and stabilized at a new lower level. After 2010 Fiat car sales decreased even more with less than 500 car sales per month (Figure 2.6 on page 37).

Hyundai started with quite a high sales level in 1998 and had a stable increase in sales until 2000 where sales had their first sales pick. The firm maintains a stable sales level with very small upward and downward movements and reaches the second sales pick in 2005 with about the same level of sales. After 2005 a downfall of sales starts giving a signal of temporal stability in 2010 but then starts to drop again. The highest sales were in January,

April, and June each year while the yearly pattern is stable except the beginning and end of the period where an increase followed by a decreasing trend is obvious respectively.

Peugeot sales had increased its new car sales level from 1998 until 2002 where the firm reached its highest sales level. After 2002 it started to gradually decrease until January 2010 where the fall in sales was sharp and drastic reaching very low levels. There is however an obvious repeated pattern on a yearly base but the fluctuations do not seem to be very wide. On the other hand, the randomness of series that is not explained by trend or seasonality seems to have quite big movements during the period and indicated the unpredicted and unexpected different sales levels in this firm.

Ford sales level had an extended period with gradual growth for six years (1998 till 2003). Ford seems to maintain its high level of sales for the next eight years (2003 till 2010) and finally started decreasing at an increasing speed. Seasonality is obvious in data with a repeated pattern throughout the years and randomness is only wide at the beginning of the observed period and stable for the rest of the period.

Nissan car sales reached their highest level in 2000 and ever since have a stable new car sales level for the next ten (10) years with very small fluctuations. After the year 2010 sales decrease rapidly. There were high seasonality and very small fluctuations concerning sales levels over each month.

Citroen starts with high levels of sales and keep them stable till 2008. It seems that the firm enjoyed a stable market share for 10 years but then sales started decreasing rapidly. In less than a year Citroen had lost more than half of its customers with a decreasing trend. In the second half of 2009, the firm seems to regain a small percentage of its share but then a more severe downfall movement started at the end of 2010 and made sales level drop to a minimum until the end of 2011 and then stabilize in a lower level.

2.4 Related Work - Literature Review

In recent years there has been a great deal of discussion on applications of various time series forecasting models and their performance in forecasting business activities. Several

time series models, original like the naïve models or the more complicated once like the family of Generalized Autoregressive Conditional Heteroskedasticity (GARCH) models and even further advanced model like the Exponential Smoothing (ETS) state-space models, have been applied to explain forecasting performance of preferred variables. However, there are paucities of studies with specific reference to the automobile retail segment in Greece. Few noteworthy studies relating to automobile retail segments and other businesses may be highlighted here. To my knowledge, there are no previous studies on the new car sales market in Greece that use different time series econometric models in estimating the volatility of sales levels.

Automobiles are highly differentiated durable products with variable lifetimes and the literature review of research in automobiles covers a wide range of modeling. Berkovec [1985] presents a short-run general equilibrium model for the automobile market combining a discrete model of consumer automobile demand with original models of new automobile production and used vehicle scrappage. Econometric estimates of the scrappage and demand functions are then used to create a simulation model of the automobile market, which is used to provide forecasts of automobile sales stocks and scrappage for the 1978-1990 period.

The Spanish Automobile Industry is investigated in Garcia-Ferrer et al. [1997] as a univariate forecasting comparison case study. This paper investigated the forecasting ability of two unobserved component models with fixed and time-varying parameters and compared it with the standard ARIMA, Box Jenkins, univariate approach. The data used are the monthly time series of automobile sales in Spain and the accuracy of the different methods was assessed by comparing several measures of forecasting performance based on the out-of-sample prediction for various horizons and different assumptions on the models' parameters. The unobserved components model seems to provide greater flexibility for adaptive applications while it is little to choose between the methods in forecasting performance terms. In Hülsmann et al. [2012] car sales forecast models are presented for the automobile markets of Germany and US –America based on time series analysis and Data Mining techniques. Lin Yuchen [2015] completed her thesis (with Honors, advised by Prof.

Ed.Rothman) at the University of Michigan, Department of Statistics, with research on Auto Car Sales Prediction as a statistical study using Functional Data Analysis and Time Series. She fitted an ordinary least square (OLS) Model and a Time Series model to the data since the error term was not independent and concluded that the unemployment rate and stock price are not the sensitive variables in the model. On the other hand interest rate, crude oil price, and consumer price index (CPI) for all items play a meaningful role in predicting auto car sales and therefore suggest including these variables in a foresting model.

In Barriera et al. [2013] presented a study describing the use of Internet search information to achieve an improved nowcasting ability with original autoregressive models using data from Portugal, Spain, France, and Italy for the unemployment rate and car sales. Especially for car sales, they have found that in some cases the volume of search queries helps to explain the variance of car sales data.

Nanaki [2018] measured the impact of economic crisis to the Greek Vehicle Market and her research findings show a positive relationship of net disposable income to car – sales and a negative relationship to unemployment rate, inflation rate and fuel price. She concluded that the implementation of austerity measures led car sales to a significant lower level and tax legislation of the Greek government affected car sales. On the one hand, she argues that the favorable loan terms during the period of 2003-2008 increased citizens' interest in buying new cars and on the other hand, the social and economic changes led to a 40% reduction during the period 2008-2014.

The automobile sector is among the most popular domains for application of choice modeling in general and in the design literature specifically in Train [2009]. The majority of researchers apply discrete choice methods to predict consumer-choice as a function of product attribute and price Wassernaar et al. [2005]. Applications of choice models within design implicitly rely on accurate choice predictions. In Lave and Train [1979] they proposed a disaggregate model of vehicle class purchase choice base on consumer characteristics and additional vehicle characteristics, like fuel economy, weight, size seat number, and horsepower. Logit models along with variants including nested and mixed logit mod-

els represent the most popular modeling approach by far. In Boyd and Mellman [1980] they proposed a random coefficient logit model, while Whitefoot and Skerlos [2012] investigated the effect of fuel economy standards on vehicle size and employ a logit model with coefficients.

The application of Time Series methods and techniques for business analysis is not new. Time Series models are widely used to better understand business data or to predict future points. One of the most important objectives in the analysis of a time series is to forecast its future values. In 1982 the Journal of Forecasting published the results of a forecasting competition organized by Makridakis et al. [1982]. That study used the 1001 time series to compare 21 extrapolative methods and their ex-ante forecasting errors using a variety of accuracy measures for different types of data and varying forecast horizons. Uni-variate time series analysis is used in Rothman [1998], Johnes [1999], Proietti [2001], Gli-Alana [2001] for forecasting the *employment rate* in US and UK.

Burman and Shumway [2006] considers the problem of prediction for stationary and non-stationary univariate time series using modification suggested by the usual exponential weighted moving average method. Empirical evidence shows that the method is competitive with autoregressive integrated moving average (ARIMA) models which usually lead to better forecasts.

Adding seasonality to business or science observations economic literature gave rise to the application of the Seasonal Autoregressive Integrated Moving Average (SARIMA) model which has been used in a wide variety of data for forecasting with successful applications in different fields of physical and social sciences.

In economic science, Wagner [2010] used SARIMA models in forecasting *daily demand in cash supply chains*. Antoniadis et al. [2006] compared the resulting predictions of two real-life data sets of wavelet -Kernel approach with those obtained by a smoothing spine method, a classical SARIMA model, and a Holt-Winters (HW) forecasting procedure. They concluded that although the SARIMA method performed better than method HW both were inferior to the predictions that were made by the other functional base methods. Furthermore Andrikopoulos and Markellos [2015] develop a model of dynamic interactions between price

variations in leasing and selling markets for automobiles monthly using US data from 2002 to 2011 showing that variations in selling market prices lead rapidly dissipating changes of leasing market prices in the opposite direction.

In tourism research, Kim [1999] and Lim and McAleer [1999, 2001] used SARIMA model to obtain short-term forecasts such as monthly or quarterly *outbound tourism* forecasts; Gonzalez and Moral [1995] applied SARIMA models in forecasting the *international demand* in Spain, Song and Wong [2003] found the model to be the best forecasting model for both China *foreign visitor arrivals and total visitors arrival* while Chang and Liao [2010] used SARIMA model for predicting monthly *outbound tourism departures* from Taiwan.

In environmental research, Li et al. [2003] applied SARIMA models for forecasting *soil dryness index* in southwest of Western Australia, while Mishra and Desai [2005], Modarres [2007] and Abebe and Foerch [2008] used the model to for *hydrological drought* forecasting. Momani [2009] used it for *rainfall prediction* in Jordan; Ibrahim et al. [2009] for *air pollutants* prediction in several area of Malaysia; Durdu [2010] used it for forecast *boron concentrations* in a river in Western Turkey and Yusof and Kane [2012] used SARIMA models and Exponential Smoothing (ETS) state-space models to predict the monthly rainfall.

In medical issues, Hu et al. [2004] applied the time series models for prediction of *Ross River virus* disease in Brisbane; Briet et al. [2008] for short term *malaria* prediction in Sri Lanka. In energy issues, Contreras et al. [2003] applied the SARIMA model for prediction of next-day electricity prices in Spain and California which they found reliable and with high accuracy; Ediger et al. [2006] for forecasting the production of fossil fuel in Turkey and Ediger and Akar [2007] for prediction of primary energy demand by fuel in Turkey.

Hall and Q.Yao [2003] shows that ARCH and GARCH models have proven valuable in modeling processes where a relatively large degree of fluctuations is present, just like in our data.

Taylor [2011] in his empirical case study paper evaluate a recently proposed seasonal exponential smoothing method previously considered only for forecasting daily supermarket sales to monthly sales data from a publishing company and show evidence that the method outperformed the other methods considered. Chatfield et al. [2001] reviews and compare

a variety of potential models for Exponential smoothing forecasting methods that allow for changing variance and is very interesting for our forecasting application because it allows for risk and uncertainty which is needed. It shows that the Exponential Smoothing procedures as an attractive proposition due to the availability of the software and the ease of interpretation of the forecasting results.

The dynamic linear models or State-Space models are very general models that seem to subsume a whole class of special cases of models and was introduced by Kalman [1960] and Kalman and Busy [1961] primarily for aerospace-related research. The model was applied in modelling economic data (Harrison and Stevens [1976], Harvey and Pierse [1984], Harvey and Todd [1983], Shumway and Stoffer [1982], Kitagawa and Gersch [1984]). Durbin and Koopman [2001] applied time series analysis in state-space models and Shumway and Stoffer [2006] gives extensive examples in applied time series based in the dynamic linear models.

Exponential Smoothing (ES) was originally developed in the late 1950s ES models rise from state-space models with only a single source of error and are called *innovation state-space models*. Gardner [1985] and Ord et al. [1997] developed more than modeling framework while Chatfield and Yar [1991], Ord et al. [1997] and Chatfield et al. [2001], established prediction intervals for exponential smoothing methods. Stellwagen and Goodrich [1999] used first the ES in automatic forecasting and Hyndman et al. [2002b] developed it further in a more general class of methods with a uniform approach to the calculation of prediction intervals, maximum likelihood estimation and the exact calculation of model selection criteria such as Akaike 's Information Criterion.

There is a variety of applications of time series data in forecasting in the literature. However, the number of efforts undertaken in the new car sales field of research is small, not to mention the complete lack of research in the field of Greek new car market, especially for time series models and methods applications that are covered in this thesis.

2.5 Theoretical Framework of Time Series Models

This section presents the way we will apply our research as a *Time series analysis*. Time series methodology is the process where the data series are analyzed to understand the underlying structure and function that produce the observations. Understanding the mechanisms of a time series allows a mathematical model to be developed that explains the data in such a way that prediction, monitoring, or control can occur. In this chapter, we are going to explain some of the major time series methods and in the following section, we empirically test our data using these methods. We explain some popular Time series models starting with the simple ones, like the Average, the Simple Naïve, and the Seasonal Naïve models, which are easily calculated, and we gradually explain more complicated ones. Research evidence shows, however, those very simple models are surprisingly effective in fitting time series or forecasting them [Hyndman and Athanasopoulos, 2013]. The Average, simple Naïve and seasonal Naïve time series models are used in this study to help us compare them with more advanced time series models and see if more complicated models can outperform these simple form models for our researched data.

2.5.1 Simple Time Series Forecasting Methods

Forecasting time series encompasses a variety of techniques, which rely primarily on the statistical properties of the data, either in isolated single time series or in groups of series. Time series models use only information on the variable to be forecast

$$x_{t+1} = f(x_t, x_{t-1}, x_{t-2}, x_{t-3}, \dots, \varepsilon)$$

where t is time, x is our variable under evaluation and ε is the error term. This equation means that the forecast value of x is a function of previous values and an error. These methods do not exploit the understanding of the time series behavior as an economic value. Therefore the objective of the forecasting process is *not* to build models, which are a good representation of the economy with all its complex interconnections, but rather to build simple models which capture the time-series behavior of the data and may be used to provide an adequate basis for forecasting (Hall Simon and Schuster, 1994).

Simple Time Series forecasting methods are proven to be quite effective sometimes in practice [Hyndman and Athanasopoulos, 2013]. Denoting n as the number of the observations, x_t is the observed time series, h is the forecasting horizon and ε_t is the white noise errors ($\varepsilon_t \sim N(0, 1)$) we have the following very simple forecasting methods.

Average or Mean Model

The Average models use a method that takes the mean of the historical data of a variable for a given time. In other words, the next value of a variable will be the mean of the last data at a given period, which is easily calculated. In the Average or Mean method all future forecasts are equal to a simple average of the observed data and therefore all observed data are of equal importance given equal weight when generating forecasts. The mathematical model can be written as:

$$\hat{x}_{n+1} = \frac{1}{n} \sum_{t=1}^n (x_t)$$

These simple forecasting methods may be suitable for forecasting, in some cases, but maybe insufficient in other cases.

Simple Naïve Model

The Simple Naïve models use a method that is even more simple as it needs no calculations at all. In this method, the next step forecasts equal to the last observed value so the next variable value will be the value of the last observation. Therefore only the most-current observation is of great importance while all previous observations provide no information for the future. This is so simple and easy that it is hard to believe that this method works remarkably well for many economics and financial time series as indicated by Hyndman and Athanasopoulos [2013]. It gives good predictions more specific in the short run because it can measure the behavior of a market better if it can be assumed to be efficient. Mathematically the model can be denoted as :

$$\hat{x}_{n+1} = x_n$$

where n is the number of the observed values. If the data follow a random walk process ($x_t = x_{t-1} + \varepsilon_t$) then this is the optimal method of forecasting. This method works remarkably well for many economic and financial time series, like stock price and stock index. It gives good predictions more specific in the short run because it can measure the behavior of a market better if it can be assumed to be efficient.

Seasonal Naïve Model

The seasonal Naïve models use a simple method that is suitable for seasonal data and therefore estimation becomes more complicated. For example, in case we have monthly data the forecast for a future month is the most recently observed value for the same month of the year i.e. the next March value will be equal to the last March value and so on. Adding the seasonal component to the simple Naïve method and we have forecast that equal to the last value from the same period. These simple forecasting methods may be suitable for forecasting, in some cases, but may also be insufficient for some other cases. So forecast x for time $n+h$ is denoted as:

$$\hat{x}_{n+h|n} = \hat{x}_{n+h-km}$$

where m is the seasonal period and $k = [(h-1)/m] + 1$. More commonly this means that when having monthly data, the forecast for all future, for example, April values is equal to the last observed April value.

2.5.2 Linear Models with Seasonal Dummies (LMSD)

Seasonal Dummies predictor is used as a special feature that adds to the model a seasonal indicator (or dummy variable) to serve as regressors for seasonal effects. A dummy variable is also known as an "indicator variable" and is a categorical variable that takes only two values (e.g. 0 or 1). Such a variable might arise, for example, when forecasting monthly car sales and you want to take account of the previous sales levels of the same month. So for forecasting January sales level the predictor takes value 1 when referring to month January and 0 otherwise. Since we have more than two categories, the variable

can be coded using several dummy variables but always one fewer than the total number of categories. In our case, we fit an ARIMA model that consists of an intercept and 11 seasonal dummy variables which is effectively a mean model with a separate mean for each month. Since our empirical data in this thesis are monthly the seasonal Dummies option added 11 seasonal dummy variables. These include a dummy regressor variable that is 1 for January and 0 for the other months, a regressor that is 1 only for February, and so forth through November. Because the model includes an intercept no dummy variable is added for December. The December effect is measured by the intercept while the effect of other seasons is measured by the difference between the intercept and the estimated regression coefficient for the season's dummy variable. The first seasonal dummy-parameter will always refer to the first period in the seasonal cycle (January for our monthly data) and since an intercept is present in the model there will be no seasonal dummy parameter for the last period in the seasonal cycle (December for our monthly data). In our case, we are forecasting monthly new car sales and we want to account for the month of the year as a predictor. In Table 2.3 we can see the dummy variable that can be created.

Table 2.3: Seasonal Dummies Variables

Month	D_{1t}	D_{2t}	D_{3t}	D_{4t}	D_{5t}	D_{6t}	D_{7t}	D_{8t}	D_{9t}	D_{10t}	D_{11t}	D_{12t}
January	1	0	0	0	0	0	0	0	0	0	0	0
February	0	1	0	0	0	0	0	0	0	0	0	0
March	0	0	1	0	0	0	0	0	0	0	0	0
April	0	0	0	1	0	0	0	0	0	0	0	0
May	0	0	0	0	1	0	0	0	0	0	0	0
June	0	0	0	0	0	1	0	0	0	0	0	0
July	0	0	0	0	0	0	1	0	0	0	0	0
August	0	0	0	0	0	0	0	1	0	0	0	0
September	0	0	0	0	0	0	0	0	1	0	0	0
October	0	0	0	0	0	0	0	0	0	1	0	0
November	0	0	0	0	0	0	0	0	0	0	1	0
December	0	0	0	0	0	0	0	0	0	0	0	0
January	1	0	0	0	0	0	0	0	0	0	0	0
.	.	.	.	.	.	.	.	.	.	.	.	.
.	.	.	.	.	.	.	.	.	.	.	.	.
.	.	.	.	.	.	.	.	.	.	.	.	.

Note: D_{1t} is the seasonal dummy for the 1_{st} month of the year i.e. January and $D_{1t} = 1$ if it is January, otherwise $D_{1t} = 0$ and so on.

The linear Model with Seasonal Dummies has a deterministic seasonality denoted S_t , that can be written as a function of seasonal dummy variables. The seasonal frequency for our monthly data is $s = 12$ and we denote as $D_{1t}, D_{2t}, D_{3t}, \dots, D_{12t}$ the seasonal dummies where:

- $D_{1t} = 1$ if s is the first period i.e. January, otherwise $D_{1t} = 0$

- $D_{2t} = 1$ if s is the second period i.e. February, otherwise $D_{2t} = 0$ and so on.

The Deterministic seasonality S_t is a linear function of the dummy variables and is denoted as

$$S_t = \sum_{i=1}^s \gamma_i D_{it} \quad (2.2)$$

$$S_t = \begin{cases} \gamma_1, & \text{if } t = \text{January} \\ \gamma_2, & \text{if } t = \text{February} \\ \dots, & \dots \\ \gamma_{12}, & \text{if } t = \text{December} \end{cases} \quad (2.3)$$

The estimation with least squares regression is

$$\hat{x}_{t+h} = \sum_{i=1}^s \gamma_i D_{it} + e_t \quad (2.4)$$

$$\hat{x}_{t+h} = \alpha + \sum_{i=1}^{s-1} \beta_i D_{it} + e_t \quad (2.5)$$

Since we regress x on an intercept and the seasonal dummies, we omit one dummy (one season i.e. December) because if we regress both the intercept plus all seasonal dummies there would be collinear and redundant.

Interpreting the coefficients α and β of the model

$$S_t = \alpha + \sum_{i=1}^{s-1} \beta_i D_{it} \quad (2.6)$$

the intercept $\alpha = \gamma_s$ is the seasonality in the omitted season while the coefficients $\beta = \gamma_i - \gamma_s$ are the difference in the seasonal component from the s^{th} period.

2.5.3 Holt–Winters methods modeling

The basic Exponentially Weighted Moving Average (EWMA) model was firstly introduced by Holt [1957a] forecasting trends in production, inventories and labor force and was later improved by Winters [1960], adding seasonality. If we define f_t to be the forecast of

x_t using only past information, then the Holt procedure uses the following formulate to forecast x_{t+1} :

$$f_{t+1} = m_t + g_t$$

where g is the expected rate of increase of the series and m is our best estimate of the underlying value of the series. We can then develop a recursion to produce a set of estimates for g and m through time

$$m_{t+1} = \lambda_0 x_{t+1} + (1 - \lambda_0)(m_t + g_t)$$

$$g_{t+1} = \lambda_1(m_{t+1} - m_t) + (1 - \lambda_1)g_t$$

Then we can perform the recursion conditional on prior values of the two smoothing parameters.

Holt's linear method [Holt, 1957b] is an extension of simple exponential forecasting that allows a locally linear trend to be extrapolated. Forecasts are given by $\hat{x}_{t+h|t} = \ell_t + hb_t$, where

$$\ell_t = \alpha x_t + (1 - \alpha)(\ell_{t-1} + b_{t-1})$$

$$b_t = \beta(\ell_t - \ell_{t-1}) + (1 - \beta)b_{t-1}$$

and the two parameters α and β must lie in $[0,1]$. Here ℓ_t denotes the level of the series and b_t the slope of the trend at time t .

A popular method for seasonal data is the Holt-Winters method for seasonal data which is introduced in Holt [1957b], and extends the Holt's method to include seasonal terms.

Then $\hat{x}_{t+h|t} = \ell_t + hb_t + s_{t-m} + h_m^+$, where

$$\ell_t = \alpha(x_t - s_{t-m}) + (1 - \alpha)(\ell_{t-1} + b_{t-1}),$$

$$b_t = \beta(\ell_t - \ell_{t-1}) + (1 - \beta)b_{t-1},$$

$$s_t = \gamma(x_t - \ell_t) + (1 - \gamma)s_{t-m},$$

$$h_m^+ = [(h - 1) \bmod m] + 1$$

and the three parameters α , β and γ all lie in $[0,1]$.

There is also a multiplicative version of the Holt-Winters method, and damped trend versions of both Holt's linear method and the Holt-Winters method [Makridakis et al., 1998]. None of these methods are explicitly based on underlying time series models, and as a result, the estimation of parameters and the computation of prediction intervals is often not done. However, all the above methods have recently been shown to be optimal for some state-space models [Hyndman and Khandakar, 2008], and maximum likelihood estimation of parameters, statistical model selection and computation of prediction intervals is now becoming more widespread.

2.5.4 Exponential Smoothing State Space (ETS) Models

One of the fundamental concepts of system theory is the state of the system which is meant to be the summary of the past behavior of the system, in our case, the past observations. According to Kitagawa and Gersch [1996] the state taken together with the future system inputs determines all future states and system outputs while the current state and the current input values determine the current outputs.

Exponential Smoothing (ES) modeling was developed in the 1950s [Brown, 1959, Holt, 1957b, Winters, 1960] and has been widely used ever since. It is a technique that can be applied to time series data either to produce smoothed data for presentation or to make forecasts. It is a broadly sensible approach to forecasting and not the result of a particular economic or statistical view about the way data was generated. In time-series data, which are a sequence of observations, exponential smoothing assigns exponentially decreasing weights over time. These models initiated research that gave rise to some of the most successful forecasting methods. Nowadays exponential smoothing has been revolutionized with the introduction of a complete modeling framework incorporating innovations state-space models, likelihood calculation, prediction intervals, and procedures for model selection Hyndman and Khandakar [2008].

In more detail forecasts produced using exponential smoothing methods are weighted averages of past observations, with the weights decaying exponentially as the observed values get older. Thus, the more recent the observations, the higher the associated weight.

Consequently, only the most recent data point and most recent forecast need to be stored.

According to Hyndman et al. [2002b] this method provides forecast accuracy comparable to the best forecasting methods and it appears to be particularly good for short forecast horizons with seasonal data. Some of its benefits are:

- the easy of calculations of the likelihood, the AIC, and other model selection criteria,
- the opportunity that it gives in computing prediction intervals for each method and
- the random simulation from the underlying state-space model.

The type of ES model has two equations the **Observed equation** and the **State space equation** as follows:

Observed equation $y_t : w'x_{t-1} + \epsilon_t$ where $\epsilon_t \sim N(0, \sigma_\epsilon^2)$

State equation $x_t : Fx_{t-1} + g\epsilon_{t-1}$

where

x_t is the state vector (unobserved)

y_t is the observed time series

ϵ_t is the white noise series

Exponential Smoothing State Space Models can improve predictability. The ETS refers to error (E), trend (T), and seasonal (S) components. The error (E) component is either additive (A) or multiplicative (M). The trend (T) and seasonal (S) component may be A, M, or inexistent (N). Trend (T) component may be dampened additively (A_d) or multiplicatively (M_d). That makes a total of thirty (30) possible ETS combinations within the forecasting framework comprising linear and non-linear ones [Hyndman and Khandakar, 2008].

Originally exponential Smoothing methods were classified by Pegels [1969] taxonomy. This was later extended by Gardner [1985], modified by Hyndman et al. [2002a] and extended again by Taylor [2003], giving a total of fifteen methods seen in Table 2.4 simply by considering variations in the combination of trend and seasonal components. Each method

is labeled by a pair of letters defining the type of Trend and Seasonal components. These 15 ETS models with multiplicative error structure (heteroskedastic) were considered for time series analysis [Medina et al., 2008]. Evidence shows that the models yield more realistic 95% prediction intervals values while the reduction on the number of ETS methods evaluated has the advantage that it diminishes the extensive computational time.

Table 2.4: The 30 possible different combinations for ETS models.

Trend component	Seasonal Component		
	N(none)	A(Additive)	M(multiplicative)
N(none)	N,N	N,A	N,M
A(Additive)	A,N	A,A	A,M
A_d (Additive damped)	A_d ,N	A_d , A	A_d , M
M(multiplicative)	M,N	M,A	M,M
M_d (Multiplicative damped)	M_d , N	M_d , A	M_d ,M

Source: "Forecasting Principle and Practice" by Hyndman and Athanasopoulos [2013]

In more detail forecasts produced using exponential smoothing methods are weighted averages of past observations, with the weights decaying exponentially as the observed values get older. Thus, the more recent the observations the higher the associated weight. Consequently, only the most recent data point and most recent forecast need to be stored.

There are many similarities between Exponential Smoothing (ES) and Moving Average (MA) models. Both models assume stationary (not trending) time series, both have roughly the same distribution of forecast error but they differ in that exponential smoothing takes into account all past data, whereas moving average only takes into account k past data points. Technically speaking, they also differ in that moving average requires that all past data points be kept, whereas exponential smoothing only needs the most recent forecast value to be kept. In the near past, this was an attractive feature of exponential smoothing method, when computer storage was expensive and it has proved remarkably robust or even optimal to a wide range of time series processes [Chatfield et al., 2001].

2.5.5 Autoregressive Integrated Moving Average (ARIMA)

Autoregressive (AR) models were first introduced by Yule [1927] and were consequently supplemented by Slutsky [1937] with the illustration of Moving Average (MA) schemes. The combination of both AR and MA schemes was made one year later by Word [1938] who showed that Autoregressive Moving Average (ARMA) process can be used to model all stationary time series as long as the appropriate order of p (the number of autoregressive terms) and q (the number of moving average terms) was appropriately specified. In other words, this means that any series x_t can be modeled as a combination of past x_t values and/or past errors ε_t :

$$x_t = \phi_1 x_{t-1} + \phi_2 x_{t-2} + \cdots + \phi_p x_{t-p} + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \theta_2 \varepsilon_{t-2} - \cdots - \theta_q \varepsilon_{t-q} \quad (2.7)$$

According to Makridakis and Hibon [2000] modeling real life time series requires four steps:

1. The original series x_t must be transformed to become stationary around its mean and its variance.
2. The appropriate order of p and q must be specified.
3. The value of the parameters $\phi_1, \phi_2, \dots, \phi_p$ and/or $\theta_1, \theta_2, \dots, \theta_q$ which represent the autoregressive and moving average components, respectively, must be estimated using some non-linear optimization procedure that minimizes the sum of the squared errors or some other appropriate loss function.
4. Practical ways of modeling seasonal series must be envisioned and the appropriate order of such a model should be specified.

The ARIMA model is considered suitable when one wants to forecast continuous data and exogenous variables Wen-hsien [2002]. However, this theoretical approach did not become possible to model real-life series at that time due to the complicated calculations involved. In the mid 1960s, the theory became popular and economical mainly due to the technological revolution of computers. All the difficult required calculations for optimization procedures of equation 2.7 were done by the computers and that made the theory

easily applicable and quite popular. Additionally, a methodology introduced by Box and Jenkins [1976] (original edition 1970) popularized the use of ARMA models through the following:

1. Provide guidelines for making the series stationary in both its means and variance,
2. Suggested the use of autocorrelations and partial autocorrelations coefficients for determining appropriate values of p and q (and their seasonal equivalent P and Q when the series exhibited seasonality)
3. Provide a set of computer programs to help users identify appropriate values for p and q as well as P and Q and estimate the parameters involved.
4. After the estimation of the parameters a diagnostic test was proposed to determine whether or not the residuals ε_t were white noise (i.e. mean equal to zero and variance equal one), in which case the order of the model was considered final. If the residuals were not white noise another model was entertained in step 2 and steps 3 and 4 were repeated. If the diagnostic check showed random residuals then the model developed was used for forecasting or control purposes assuming of course remain the same during the forecasting, or control, phase.

Strictly speaking, there is no such thing as “the best” model to fit or forecasting model. Thus, the most important problem to be solved when modeling is that of trying to match the appropriate model to the pattern of the available time-series data. Box and Jenkins [1976], therefore, proposed three practical stages for finding a good model, namely identification, estimation, and diagnostic checking. The approach proposed by Box and Jenkins came to be known as the Box-Jenkins methodology to ARIMA models where the letter I between AR and MA stood for the word Integrated. ARIMA models and the Box-Jenkins methodology became highly popular in the 1970s among academics in particular when it was shown through empirical studies that they could outperform the large and complex econometric models, popular at that time in a variety of situations Armstrong [1978].

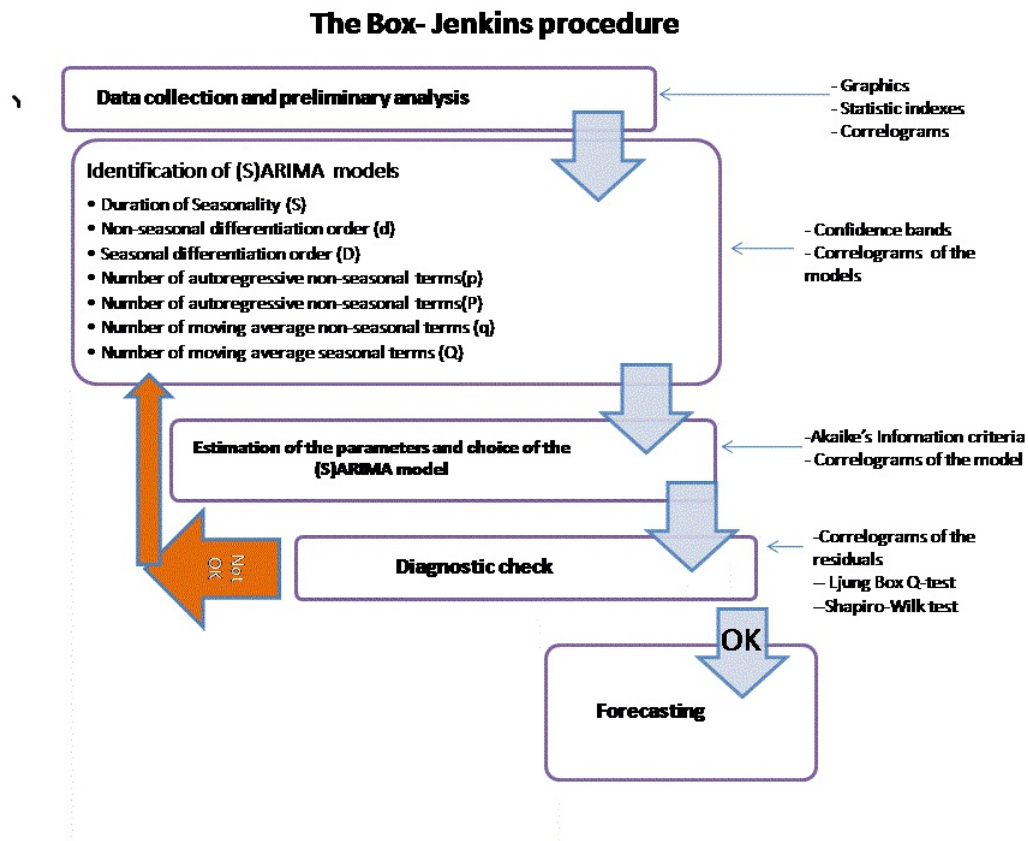


Figure 2.7: Box and Jenkins Methodology.

An ARIMA(p,d,q) process is given by:

$$\underbrace{(1 - \phi_1 B - \dots - \phi_p B^p)}_{AR(p)} \underbrace{(1 - B)^d}_{d-differences} x_t = c + \underbrace{(1 + \theta_1 B + \dots + \theta_q B^q)}_{MA(q)} \varepsilon_t \quad (2.8)$$

or in a more compact presentation as:

$$\phi(B)(1 - B)^d x_t = c + \theta(B)\varepsilon_t \quad (2.9)$$

where x_t represent the time series, ε_t is assumed to be uncorrelated error term with white noise process i.e. zero mean and variance σ^2 , B is the back-shift operator, c is the constant term and d takes values 1 or 0 depending on whether there is a need for first difference or not to ensure stationarity.

The polynomial of the Autoregressive (AR) order p is :

$$\phi(B) = (1 - \phi_1 B + \phi_2 B - \dots - \phi_p B^p)$$

and the polynomial of the moving average (MA) of order q is :

$$\theta(B) = (1 + \theta_1 B + \theta_2 B + \cdots + \theta_q B^q)$$

Furthermore to ensure causality and invertibility it is assumed that $\phi(B)$ and $\theta(B)$ have no roots for $|B| < 1$ [Brockwell and Davis, 1991].

The methodology for time series analysis proposed in the famous book by Box and Jenkins [1970] aims to find the most appropriate ARIMA (p,d,q) model that fits the data series and to find the model that is best for forecasting. In more details the Box- Jenkins methodology uses a six-stage scheme which are:

- A priori identification of the differentiation order d (or choice of another transformation)
- A priori identification of the orders p and q
- Estimation of the parameters $(\alpha_1, \dots, \alpha_p, \beta_1, \dots, \beta_q$ and $\sigma^2 = \text{Var } \varepsilon_t)$
- Validation
- Choice of a model
- Prediction

For the first step we examine the graph of the series for signs of non-stationarity, but there are also a variety of official unit root tests developed in the last 30 years, like the Dickey-Fuller test (DF) or the Augmented Dickey-Fuller test (ADF) Dickey and Fuller [1981], the Philips Perron Test (PP) Philips and Perron [1988], the Kwiatkowski, Phillips, Schmidt and Shin test (KPSS) Kwiatkowski et al. [1992] e.c.t. that can be applied to examine the stationarity of the series. If the data series exhibit apparent deviations from stationarity, we will take the first difference of the series (so $d=1$), and this may be an indication that the underlying process is heteroscedastic.

However the order selection for the ARIMA models is usually considered subjective and difficult to apply. Thus, researchers tried to automate the ARIMA modeling for a

stationary series during the last 30 years. An early review of the automatic ARIMA implementation can be viewed in Ord and Lowe [1996]. Hannan and Rissanen [1982] proposed a method where innovations (ε_t) could be obtained by fitting a long autoregressive model to the data and then the likelihood of potential models could be computed via a series of standard regressions. This method was later extended to include multiplicative seasonal ARIMA models identification Tsay [2005]. Many researchers implement different algorithms using different software (Gómez and Maravall [1998], Liu [1989], Mèlard and Pasteels [2000], Goodrich [2000], Makridakis and Hibon [2000], Reilly [2000]). Hyndman and Khandakar [2008] in the “forecast” package of R uses an algorithm for step-wise forecasting with ARIMA models and that is partially used in this thesis research along with a random selection method.

2.5.6 Seasonal ARIMA Model (SARIMA)

For a given time series (x_t) with seasonal period s , we can combine the seasonal and non –seasonal operators into a multiplicative seasonal autoregressive moving average model denoted as:

$$\underbrace{SARIMA}_{\uparrow \text{(Non-seasonal part of the model)}}(\underbrace{p, d, q}_{\uparrow \text{(Seasonal part of the model)}})x(\underbrace{P, D, Q}_s)_s$$

The SARIMA model equation is given by :

$$\underbrace{(1 - \phi_1 B - \dots - \phi_p B^p)}_{AR(p)} \underbrace{(1 - \Phi_1 B - \dots - \Phi_P B^s)}_{SAR(P)} \underbrace{(1 - B)}_d \underbrace{(1 - B^s)}_D x_t = \underbrace{(1 + \theta_1 B + \dots + \theta_q B^q)}_{MA(q)} \underbrace{(1 + \Theta_1 B + \dots + \Theta_Q B^Q)}_{SMA(Q)} \varepsilon_t \quad (2.10)$$

where t is the point in time which is $t=1,2,3,\dots,n$,

c is the constant term,

B is the lag operator,

ε_t is the error term at time t ,

AR(p) is the autoregressive model of order p ,

SAR(P) is the seasonal autoregressive model of order P ,

MA(q) is the moving average model of order q ,

SMA(Q) is the seasonal moving average model,

d is the non-seasonal differences,

D is the seasonal differences,

$\phi_p(B)$ and $\theta_q(B)$ of order p and q represent the ordinary autoregressive and moving average components,

$\Phi_P(B^s)$ and $\Theta_Q(B^s)$ of order P and Q represent the seasonal autoregressive and moving average components and all the polynomials should have roots outside the unit circle ensuring an invertible representation for the differenced data.

The SARIMA model can be written in a more compact way like:

$$\phi_p(B)\Phi_P(B^s)\nabla^d\nabla_s^D x_t = c + \theta_q(B)\Theta_Q(B^s)\varepsilon_t \quad (2.11)$$

where $t=1,2,3,\dots,n$, B is the lag operator, $\phi_p(B)$ and $\theta_q(B)$ of order p and q represent the ordinary autoregressive and moving average components, $\Phi_P(B^s)$ and $\Theta_Q(B^s)$ of order P and Q represent the seasonal autoregressive and moving average components and all the polynomials should have roots outside the unit circle ensuring an invertible representation for the differenced data. Furthermore

$$\nabla^d = (1 - B)^d$$

and

$$\nabla_s^D = (1 - B^s)^D$$

are the ordinary and seasonal difference components respectively where d is the consecutive differencing and D the seasonal differencing. To ensure causality and invertability, we

assume that $\phi_p(B)$ and $\theta_q(B)$ have no roots for $|B| < 1$ Brockwell and Davis [1991]. The model may contain a constant term c . The ordinary autoregressive (AR) characteristic polynomial $\phi_p(B)\Phi_P(B)$ and moving average (MA) characteristic polynomial $\theta_q(B)\Theta_Q(B)$ components are represented by:

$$\begin{aligned}\phi_p(B) &= (1 - \phi_1 B - \phi_2 B^2 - \dots - \phi_p B^p) \\ \Phi_P(B^s) &= (1 - \Phi_1 B^s - \Phi_2 B^{2s} - \dots - \Phi_P B^{Ps}) \\ \theta_q(B) &= (1 + \theta_1 B + \theta_2 B^2 + \dots + \theta_q B^q) \\ \Theta_Q(B^s) &= (1 + \Theta_1 B^s + \Theta_2 B^{2s} + \dots + \Theta_Q B^{Qs})\end{aligned}\tag{2.12}$$

Note that we have a special ARIMA model with AR order $p + Ps$ and MA order $q + Qs$ and integration of order d for the series and D for the seasonal component. The coefficients are being determined by only $p + P + q + Q$ coefficients and are not completely general. If $s = 12$ then $p + P + q + Q$ will be considerably smaller than $p + Ps + q + Qs$ and will allow a much more parsimonious model Cryer and Chan [2008].

2.5.7 Generalized Autoregressive Conditional Heteroscedastic

The Autoregressive conditionally heteroscedastic (ARCH) models were introduced by Engle [1982]. Bollerslev [1986] extended these models to the Generalized ARCH (GARCH) model as an alternative to the usual time series process. Empirical evidence for some kinds of data shows that the disturbance variances in time series models were less stable than usually assumed (Engle [1982], Engle [1983], and Cragg [1982]). Volatility is not constant over time, and sometimes errors appeared to occur in clusters, suggesting a form of heteroscedasticity in which the variance of the forecast error depends on the size of the previous disturbance. This means that large movements in the series tend to be followed by further large movements. Thus the economy has cycles with volatility and low volatility periods, just like in the financial return series. The basic concept in these models is the conditional variance, that is the variance conditional on the past. In the classical GARCH models, we notice that the conditional variance is expressed as a linear function of the squared past values of the series which enables the capture of the main characteristic

features of time series.

GARCH models modification

The classical GARCH modeling has an important drawback by its construction. The fact that the conditional variance depends only on the modulus of the past variables, is the major reason that negative and positive innovations have the same effect on the current volatility. On the other hand, empirical research, in most of the time-series data, shows a negative correlation between the squared current innovation and the past innovations, which gives a signal that the conditional distribution is asymmetric. If the conditional distribution were symmetric in the past variables, such a correlation would be equal to zero. In a conditional asymmetric distribution, the volatility increase due to a data level decrease is generally stronger than that resulting from a price increase of the same magnitude and this is also referred to as a leverage effect. In order to allow this asymmetry property and the leverage effect to be incorporated we consider two groups of GARCH models modifications:

1. GARCH models with alternative conditional error distributions (Gaussian, t-Student and generalized).
2. GARCH models with asymmetric conditional volatility (EGARCH, GJR, APARCH).

Alternative Conditional Distributions

Time series χ_t observations have a distribution that one often assumes to be normal (Gaussian) but in reality they usually tend to be leptokurtic (fat tailed). In this thesis the fat tailed Student-t distribution (STD) and the generalized error distributions (GED) are considered. Further information about the different types of distributions are given in this section.

Normal Distribution. Usually the default choice for the distribution (D) of the innovations z_t of a GARCH process is the Standardized Normal Probability Function:

$$f^*(y) = \frac{1}{\sqrt{2\pi}} e^{-\frac{y^2}{2}} \quad (2.13)$$

The probability function or density is named standardized, marked by a star $\star$, because $f^\star(z)$ has zero mean and unit variance. This can easily be verified computing the moments:

$$\mu_r^\star = \int_{-\infty}^{\infty} y^r f^\star(y) dy \quad (2.14)$$

Note, that $\mu_0^\star \equiv 1$ and $\sigma_1^\star \equiv 1$ are the normalization conditions, that the first moment $\mu_1^\star$ defines the mean $\mu \equiv 0$, and the second moment $\mu_2^\star$ the variance $\sigma^2 \equiv 1$.

An arbitrary Normal distribution located around a mean value μ and scaled by the standard deviation σ can be obtained by introducing a location and a scale parameter through the transformation

$$f(y)dy \rightarrow \frac{1}{\sigma} f^\star\left(\frac{y - \mu}{\sigma}\right) dy = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(y - \mu)^2}{2\sigma^2}} dy \quad (2.15)$$

The central moments μ_r of $f(y)$ can simply be expressed in terms of the moments $\mu_r^\star$ of the standardized distribution $f^\star(y)$. Odd central moments are zero and those of even order can be computed from

$$\mu_{2r} = \int_{-\infty}^{\infty} (y - \mu)^{2r} f(y) dy = \sigma^{2r} \mu_{2r}^\star = \sigma^{2r} \frac{2r}{\sqrt{\pi}} \Gamma\left(r + \frac{1}{2}\right) \quad (2.16)$$

yielding $\mu_2 = 0$, and $\mu_4 = 3$. The degree of asymmetry γ_1 of a probability function, named skewness, and the degree of peakedness γ_2 , named excess kurtosis, can be measured by normalized forms of the third and fourth central moments

$$\gamma_1 = \frac{\mu_3}{\mu_2^{3/2}} = 0, \gamma_2 = \frac{\mu_4}{\mu_2^2} - 3 = 0 \quad (2.17)$$

On the other hand, if we like to model an asymmetric and/or leptokurtic shape of the innovations we have to draw or to model y_t from a standardized probability function which depends on additional shape parameters which modify the skewness and kurtosis. However, it is important that the probability has still zero mean and unit variance. Otherwise, it would be impossible, or at least difficult, to separate the fluctuations in the mean and variance from the fluctuations in the shape of the density. In a first step we consider still symmetric probability functions but with an additional shape parameter which models the kurtosis. As examples we consider the generalized error distribution and the Student-t distribution with unit variance, both relevant in modelling GARCH processes.

Student-t distribution. Bollerslev [1988], Hsieh [1989], Baillie and Bollerslev [1989], Bollerslev et al. [1992], Palm [1996], Pagan [1996] and Palm and Vlaar [1997], among others showed that the Student-t distribution better captures the observed kurtosis in empirical log-return time series. The density $f^*(y|\nu)$ of the Standardized Student-t Distribution can be expressed as:

$$\begin{aligned} f^*(y|\nu) &= \frac{\Gamma(\frac{\nu+1}{2})}{\sqrt{\pi(\nu-2)\Gamma(\frac{\nu}{2})}} \frac{1}{(1 + \frac{y^2}{\nu-2})^{\frac{\nu+1}{2}}} \\ &= \frac{1}{\sqrt{\nu-2}B(\frac{1}{2}, \frac{\nu}{2})} \frac{1}{(1 + \frac{y^2}{\nu-2})^{\frac{\nu+1}{2}}} \end{aligned} \quad (2.18)$$

where $\nu > 2$ is the shape parameter and $B(\alpha, b) = \Gamma(\alpha)\Gamma(b)/\Gamma(\alpha + b)$ the Beta function. Note, when setting $\mu = 0$ and $\sigma^2 = \nu/(\nu - 2)$ formula (2.18) results in the usual one-parameter expression for the Student-t distribution as implemented in the S function `dt`.

Again, arbitrary location and scale parameters μ and σ can be introduced via the transformation $y \rightarrow \frac{y-\mu}{\sigma}$. Odd central moments of the standardized Student-t distribution are zero and those of even order can be computed from:

$$\mu_{2r} = \sigma^{2r} \mu_{2r}^* = \sigma^{2r} (\nu - 2)^{\frac{r}{2}} \frac{B(\frac{r+1}{2}, \frac{\nu-r}{2})}{B(\frac{1}{2}, \frac{\nu}{2})} \quad (2.19)$$

Skewness γ_1 and kurtosis γ_2 are given by

$$\gamma_1 = \frac{\mu_3}{\mu_2^{3/2}} = 0, \gamma_2 = \frac{\mu_4}{\mu_2^2} - 3 = \frac{6}{\nu - 4} \quad (2.20)$$

This result was derived using the recursion relation $\mu_{2r} = \mu_{2r} - 2\frac{2r-1}{\nu-2r}$.

Generalized Error Distributions. D.B.Nelson [1991] suggested to consider the family of Generalized Error Distributions, GED, already used by Box and Tiao [1973], and Harvey [1981]. $f^*(y|\nu)$ can be expressed as:

$$\begin{aligned} f^*(y|\nu) &= \frac{\nu}{\lambda_\nu 2^{1+1/\nu} \Gamma(1/\nu)} e^{-\frac{1}{2}|\frac{y}{\lambda}|^\nu}, \\ \lambda_\nu &= \left(\frac{2^{(-2/\nu)} \Gamma(\frac{1}{\nu})}{\Gamma(\frac{3}{\nu})} \right)^{1/2} \end{aligned} \quad (2.21)$$

with $0 < \nu < \infty$. Note, that the density is standardized and thus has zero mean and unit variance. Arbitrary location and scale parameters μ and σ can be introduced via the

transformation $y \rightarrow \frac{y-\mu}{\sigma}$. Since the density is symmetric, odd central moments of the GED are zero and those of even order can be computed from

$$\mu_{2r} = \sigma^{2r} \mu_{2r}^* = \sigma^{2r} \frac{(2^{1/\nu} \lambda_\nu)^{2r}}{\Gamma(\frac{1}{\nu})} \Gamma(\frac{2r+1}{\nu}) \quad (2.22)$$

Skewness γ_1 and kurtosis γ_2 are given by

$$\gamma_1 = \frac{\mu_3}{\mu_2^{3/2}} = 0, \gamma_2 = \frac{\mu_4}{\mu_2^2} - 3 = \frac{\Gamma(\frac{1}{\nu})\Gamma(\frac{5}{\nu})}{\Gamma(\frac{3}{\nu})} - 3 \quad (2.23)$$

For $\nu = 1$ the GED reduces to the Laplace distribution, for $\nu = 2$ to the Normal distribution, and for $\nu \rightarrow \infty$ to the uniform distribution as a special case. The Laplace distribution takes the form $f(y) = e^{-\sqrt{2}|y|/\sqrt{2}}$, and the uniform distribution has range $\pm 2\sqrt{3}$. The Laplace distribution is a GED with shape parameter $\nu = 1$.

2.6 Forecasting Methodology.

All forecasting methods can be divided into two broad categories:

- Qualitative and
- Quantitative

which are divided further into more categories as illustrated below in the following figure.

Qualitative forecasting models are useful in developing forecasts with a limited scope and these models are usually highly reliant on expert opinions and are most beneficial in the short term. Some examples of qualitative forecasting models include market research, polls, and surveys that apply the Delphi method or the Scenario writing or the subjective approach.

Quantitative methods of forecasting exclude expert opinions and utilize statistical data based on quantitative information and available historical time series data. Therefore quantitative forecasting models include time series methods and other econometric modeling methods.

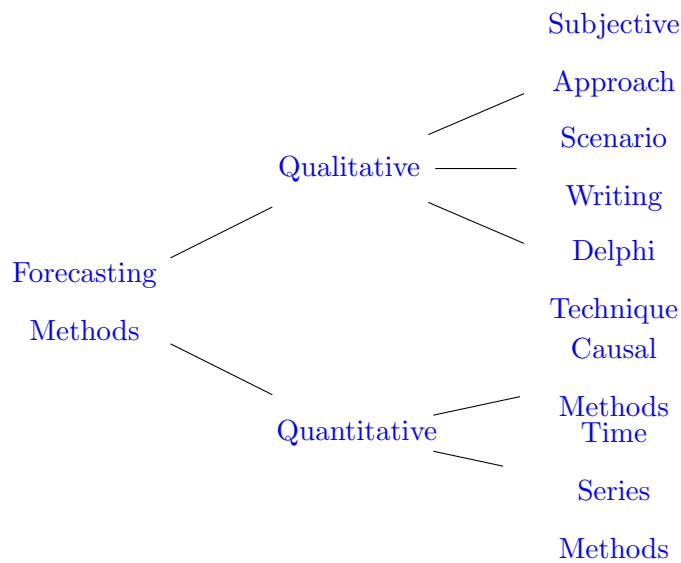


Figure 2.8: Forecasting Methods [Tree Graph].

2.6.1 Qualitative vs Quantitative Forecasting Methods.

Qualitative forecasting techniques generally employ the judgment of experts in the appropriate field to generate forecasts. A key advantage of these procedures is that they can be applied in situations where historical data are simply not available. Moreover, even when historical data are available, significant changes in environmental conditions affecting the relevant time series may make the use of past data irrelevant and questionable in forecasting future values of the time series. There are three important qualitative forecasting methods are the Delphi technique, scenario writing, and the subject approach.

In the Delphi technique, an attempt is made to develop forecasts through "group consensus". Usually, a panel of experts is asked to respond to a series of questionnaires. The experts, physically separated from and unknown to each other, are asked to respond to an initial questionnaire (a set of questions). Then, a second questionnaire is prepared to incorporate the information and opinions of the whole group. Each expert is asked to reconsider and to revise his or her initial response to the questions. This process is continued until some degree of consensus among experts is reached. It should be noted that the objective of the Delphi technique is not to produce a single answer at the end. Instead, it attempts to

produce a relatively narrow spread of opinions the range in which opinions of the majority of experts lie.

Under the Scenario Writing approach, the forecaster starts with different sets of assumptions. For each set of assumptions, a likely scenario of the business outcome is charted out. Thus, the forecaster would be able to generate many different future scenarios, corresponding to the various sets of assumptions. The decision-maker is presented with different scenarios and has to decide which scenario is most likely to prevail.

The Subjective Approach allows individuals to participate in the forecasting decision to arrive at a forecast based on their subjective feelings and ideas. This approach is based on the premise that a human mind can arrive at a decision based on factors that are often very difficult to quantify. "Brainstorming sessions" are frequently used as a way to develop new ideas or to solve complex problems. In loosely organized sessions, participants feel free from peer pressure and, more importantly, can express their views and ideas without fear of criticism.

Quantitative forecasting methods are used, when historical data on variables of interest are available. These methods are based on an analysis of historical data concerning the time series of the specific variable of interest and possibly other related time series. Many forecasting techniques use past or historical data in the form of time series. There are two major categories of quantitative forecasting methods.

- Time series methods, are based on past data that are being forecast.
- Other Econometric modeling methods, which may use historical data, like the causal methods, which examine the cause and effect relationships of the variable with other relevant variables.

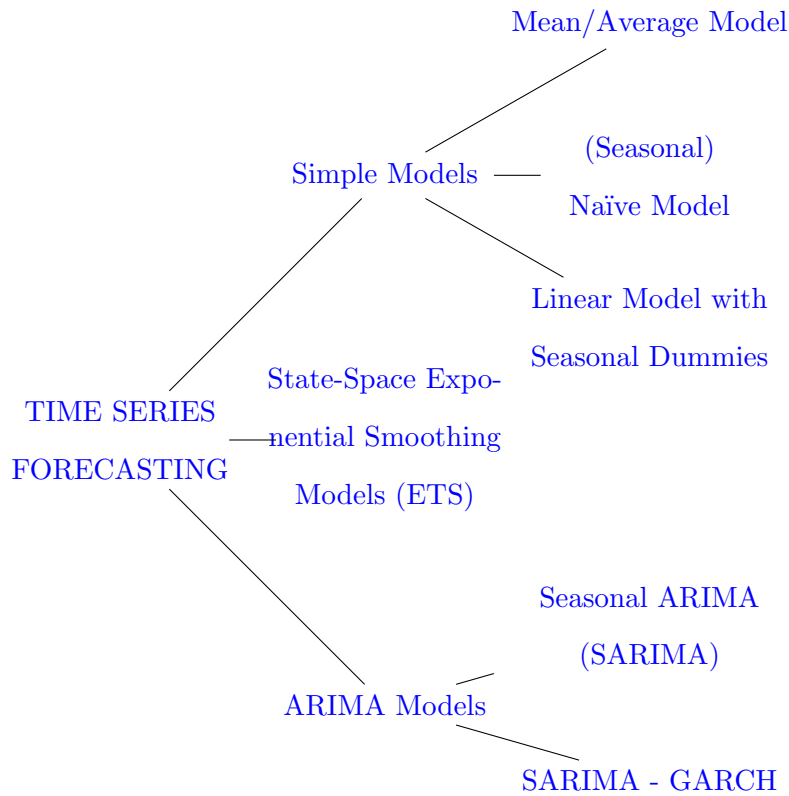


Figure 2.9: Time Series Forecasting Methods [Tree Graph].

There are various time series forecasting methods, some of them are simple and some of them are more complicated, but they all have the same goal, accurate predictions. Time-series methods are going to be examined in-depth and are going to be used in this thesis applied data research.

2.7 Discussion

During the last two decades new car market in Greece has been undergoing some severe changes due to the difficult economic condition in the country. New car sales have shrunk and never managed to overcome the economic crises especially after the worldwide financial crisis 2007-2008.

In the Greek car market there are more than 35 different firms operating in a national level but only 2-3 of them have a share around 10% of the total market. We studied ten (10) of the top firms of the market. The results in Table 2.1 (page 26) shows that when using

the original values of the series the mean of the series it is not equal to the median so that indicate that we do not have a normal distribution in none of the ten different time series. Additionally, the skewness and the kurtosis of the series differ from the ones normally found in normally distributed series (i.e. $S=0$, $K=3$). Results show that the original data are all right skewed (positive) platykurtic ($K<3$) while the log data are left skewed platykurtic with only two exceptions (VolksWagen and Nissan). However, when we consider the log values of the series then the differences between the mean and the median become smaller so there is a kind of normality in the data but we still have a small deviation. Furthermore, the Jarque -Bera test for normality, rejects normality for all series in log values (except the Fiat case) while in the original values only Toyota and VolksWagen seem to be normally distributed.

After presenting the related work done (literature review) there is a short theoretical framework of the various time series models that are going to be empirically implemented in the next chapter along with the forecast methodology that we are going to follow. This research continues with only a small group of the firms in their log values for more accurate, and quick data processing, and for the facilitation and easiness of results presentation, in an in-sample, and out-of sample, model estimation and forecasting.

Time Series with-in-sample Modeling.

3.1 Introduction

In this thesis chapter, we initiate the empirical analysis using a with-in-sample time series modeling. The monthly new car sales of the top firms operating, as a retailer of new car representative, in the Greek market are treated as time series, and the empirical study starts with simple econometric models, like the Average (or Mean), the Naïve and the Seasonal Naïve and then continues with some more advanced models like the Exponential Smoothing state space (ETS) models and the Linear models with seasonal dummies (LMSD). Additionally, the research focus in relating the present value of series, to past values and past prediction errors so it uses time series models called Autoregressive Integrated Moving Average (ARIMA) models, and since our data have seasonality the models are called Seasonal Autoregressive Integrated Moving Average (SARIMA) models. They are constructed using the Box and Jenkins methodology for time series data. This empirical research is completed with the estimation of the hybrid model of the linear seasonal autoregressive moving average (SARIMA) and the non-linear generalized autoregressive conditional heteroscedasticity (GARCH) in an in-sample modeling of the series. The goodness of the fit of each model and its performance is measured using the root mean square error (RMSE) and the mean absolute square error (MASE), while the analysis is carried out using the R software.

The idea of using various model types is due to the fact that different types of econometric models are designed to capture different characteristics that are commonly associated with time series, for example varying variance, fat tails, volatility clustering, leverage effects, and so on. However there is a problem of how to choose the correct model that best fit the series. This problem is solved by the calculation of variance at the series of each model. The variance is the Minimum Squared Errors (MSE) that results from fitting the various models. The Root of the Minimum Squared Errors (RMSE) is actually the standard deviation of the series and it is given as :

$$RMSE = \sqrt{\frac{1}{n} \sum_{t=1}^n (x_t - \hat{x}_t)^2} \quad (3.1)$$

where n is the number of observations x_t is the values of the series at time t and $\hat{x}_t$ is the estimated values of the series from the chosen model at time t .

Additionally another measure is the Minimum Absolute Percentage Errors (MAPE) which is calculated as:

$$MAPE = \frac{1}{n} \sum_{t=1}^n \frac{|x_t - \hat{x}_t|}{x_t} \times 100 \quad (3.2)$$

According to Wang and Lim [2005], RMSE is the square root of the average of all squared errors and it ignores any over or under estimation, but it does not allow comparison across different time series and different time intervals. On the other hand MAPE does allow comparison across different time series and different time intervals and is particularly useful when the units of measurements of x_t are relatively large. In the results that follow both metrics are calculated for a variety of models using the log values of all series. Finally, all models will be compared using the MAPE accuracy metrics that allows for comparison across different time series in order to discuss the finding of this with-in sample empirical study.

3.2 Simple Time Series Models

This with-in-sample empirical analysis starts with the most simple Time series models like the Mean or Average model, the Naïve and Seasonal Naïve model. The RMSE and the MAPE is calculated in Table 3.1 on page 74 and shows that the Seasonal Naïve model fit best to our data among the simple time series models of all firms since it gives the minimum value of RMSE and MAPE performance measures with no exception.

However the diagnostic checking of the Seasonal Naïve model (see Table 3.2, page 75) which is employed by examining if the fitted model specification is adequate do not give good results. The outcome of Box-Pierce and L-Jung Box test shows that the residual series are correlated so we need to develop a better model for analysis of new car sales series.

3.3 Holt-Winters & Exponential Smoothing State Space Models

The with-in-sample empirical analysis continues with the Holt-Winters time series models and the Exponential Smoothing State Space (ETS) models. The research considers three types of Holt-Winters which is the simple exponential model: a)with Level b)with Level and Trend (L, T) and c)with Level Trend and Seasonality (L, T, S). Additionally, it calculates the best ETS model from various possible combinations. In this research, the ETS(A, N, A) is proven to be the best which means the exponential smoothing with additive errors and additive seasonality but with no trend. The RMSE and MAPE performance measures were calculated in Table 3.3 on page 76. Results show that the Exponential Smoothing State-Space model with no Trend but with an additive seasonal component and error fits best among these models¹ to the time series data.

Furthermore the diagnostic tests for the ETS models (see Table 3.4 page 77) shows that the residuals series after fitting the ETS models are uncorrelated and that means that ETS

¹Note L: Level, T: Trend, S: Season

models give a good fit to our data and be consider to handle the new car series effectively.

3.4 Linear Modelling with Seasonal Dummies Models

In modeling monthly demand for new car and the researcher wants to account for the month of the year, as a predictor. Therefore the following monthly dummy variables can be created with frequency $s = 12$ and denoted as seasonal dummies or indicators as $D_{1t}, D_{2t}, D_{3t}, \dots, D_{12t}$ according to Hyndman and Athanasopoulos [2013].

Table 3.5: Indicator Variables of Linear Model with Seasonal Dummies

Month	
January D_{1t}	$x_{1t} = \alpha + \beta_1 + e_t$
February D_{2t}	$x_{2t} = \alpha + \beta_2 + e_t$
March D_{3t}	$x_{3t} = \alpha + \beta_3 + e_t$
April D_{4t}	$x_{4t} = \alpha + \beta_4 + e_t$
May D_{5t}	$x_{5t} = \alpha + \beta_5 + e_t$
June D_{6t}	$x_{6t} = \alpha + \beta_6 + e_t$
July D_{7t}	$x_{7t} = \alpha + \beta_7 + e_t$
August D_{8t}	$x_{8t} = \alpha + \beta_8 + e_t$
September D_{9t}	$x_{9t} = \alpha + \beta_9 + e_t$
October D_{10t}	$x_{10t} = \alpha + \beta_{10} + e_t$
November D_{11t}	$x_{11t} = \alpha + \beta_{11} + e_t$
December D_{12t}	$x_{12t} = \alpha + e_t$

In the Table 3.5 on page 72 the dummy or indicator variables were defined and there were only eleven (11) dummy variables² that were needed to code twelve (12) categories

²Note: D_{1t} is the seasonal dummy for the 1_{st} month of the year i.e. January and x_{1t} are the new car

i.e. twelve (12) months of each year. That was because the 12th category, which was month December in our study, was specified when the dummy variables were all set to zero (0).

The Linear Model with seasonal Monthly data is given by³:

$$x_t = \alpha + \beta_1 D_{1t} + \beta_2 D_{2t} + \beta_3 D_{3t} + \beta_4 D_{4t} + \beta_5 D_{5t} + \beta_6 D_{6t} + \beta_7 D_{7t} + \beta_8 D_{8t} + \beta_9 D_{9t} + \beta_{10} D_{10t} + \beta_{11} D_{11t} + e_t \quad (3.3)$$

According to Hyndman and Athanasopoulos [2013] many beginners try to add a twelve dummy variable for the twelve category and that mistake is known as the "dummy variable trap" because it will cause the regression to fail because of the big amount of parameters to estimate. Therefore the general rule is to use one fewer dummy variables than categories. So for our yearly data we use eleven dummy variables. Furthermore Hyndman and Athanasopoulos say that the interpretation of each of the coefficients associated with the dummy variables is the measure of the effect of that category relative to the omitted category. In the above example, the coefficient associated with January will measure the effect of January compared to February on the forecast variable and so on.

The results from the diagnostic tests of Box-Pierce and Box Ljung for the Linear models with Seasonal Dummies - LMSD (see Table 3.7 page 78) shows that the residuals series of this model are correlated so we need to develop a better model to fit new car sales series.

sales for January at year t and so on.

³Note: D_{12t} is not included

Table 3.1: Simple models RMSE & MAPE (log data with-in-sample).

	Mean	Naïve	Seasonal Naïve
RMSE			
Opel	0,43	0,41	0,36
Toyota	0,44	0,49	0,37
VW	0,41	0,40	0,37
Hyundai	0,61	0,44	0,39
Peugeot	0,65	0,43	0,40
Ford	0,56	0,39	0,40
Fiat	0,45	0,42	0,38
Nissan	0,60	0,45	0,43
Skoda	0,58	0,41	0,42
Citroen	0,55	0,39	0,40
MAPE%			
Opel	4,76	4,30	3,75
Toyota	4,61	4,68	3,63
VW	4,46	4,08	3,60
Hyundai	6,78	4,68	4,16
Peugeot	6,95	4,37	4,29
Ford	7,71	4,66	4,59
Fiat	5,65	4,67	4,15
Nissan	7,93	4,78	4,64
Skoda	7,03	4,55	4,50
Citroen	6,98	4,40	4,45

Table 3.2: Diagnostic Tests for **Seasonal Naïve** model

	Box-Pierce	p-value	Box-Ljung	p-value
Opel	67,76	(2,2e-16)	68,76	(2,2e-16)
Toyota	101,19	(2,2e-16)	102,6	(2,2e-16)
VW	39,00	(4,23e-10)	39,75	(2,87e-10)
Hyundai	45,37	(1,63e-11)	46,25	(1,04e-11)
Peugeot	62,66	(2,44e-15)	63,88	(1,33e-15)
Ford	44,12	(3,08e-11)	44,97	(1,99e-11)
Nissan	43,24	(4,83e-11)	44,08	(3,14e-11)
Fiat	57,92	(2,62e-14)	58,80	(1,743e-16)
Skoda	83,31	(2,2e-16)	84,92	(2,2e-16)
Citroen	40,63	(1,83e-10)	41,42	(1,22e-10)

Table 3.3: Holt-Winters & ETS models RMSE & MAPE (log data with-in-sample).

	Level	L,T	L,T,S	ETS
RMSE				
Opel	0,35	4,25	5,24	0,21
Toyota	0,39	4,33	5,33	0,25
VW	0,33	4,19	5,18	0,24
Hyundai	0,38	4,21	5,15	0,24
Peugeot	0,38	3,96	4,86	0,25
Ford	0,33	4,00	4,96	0,26
Fiat	0,36	4,18	5,87	0,23
Nissan	0,40	4,10	5,10	0,27
Skoda	0,39	3,80	4,70	0,25
Citroen	0,40	4,01	4,95	0,26
MAPE%				
Opel	3,77	35,96	52,04	2,12
Toyota	4,02	33,36	52,03	2,40
VW	3,39	35,85	52,16	2,41
Hyundai	4,37	36,40	52,14	2,52
Peugeot	6,95	34,37	54,29	2,86
Ford	3,83	35,92	52,11	2,76
Fiat	3,25	34,57	50,09	2,35
Nissan	7,25	38,10	55,03	2,95
Skoda	6,98	37,09	54,05	2,90
Citroen	7,28	38,05	53,10	2,95

Table 3.4: Diagnostic Tests for **ETS model**

ETS	Box-Pierce	p-value)	Box-Ljung	(p-value)
Opel (A,Ad,A)	2,36	(0,12)	2,39	(0,12)
Toyota(A,A,A)	0,015	(0,89)	0,01	(0,89)
VolksWagen (A,Ad,A)	0,10	(0,74)	0,10	(0,74)
Hyundai (A,N,A)	0,08	(0,76)	0,08	(0,76)
Peugeot (A,N,A)	0,09	(0,92)	0,09	(0,92)
Ford (A,N,A)	0,29	(0,58)	0,30	(0,58)
Nissan (A,N,A)	0,54	(0,45)	0,55	(0,45)
Fiat (A,A,A)	5,23	(0,99)	5,30	(0,99)
Skoda (A,N,A)	0,02	(0,88)	0,02	(0,88)
Citroen (A,N,A)	1,82	(0,17)	1,85	(0,17)

Table 3.6: Linear Model with Seasonal Dummies RMSE & MAPE (log values).

	RMSE	MAPE%
Opel	0,58	7,45
Toyota	0,59	7,16
VW	0,68	8,11
Hyundai	0,71	6,28
Peugeot	0,65	7,56
Ford	0,67	8,40
Fiat	0,71	9,41
Nissan	0,70	8,50
Skoda	0,69	8,20
Citroen	0,66	7,60

Table 3.7: Diagnostic Tests for **LMSD** models

ETS	Box-Pierce	p-value)	Box-Ljung	(p-value)
Opel	178,34	(2,2e-16)	2,39	(2,2e-16)
Toyota	140,56	(2,2e-16)	142,42	(2,2e-16)
Volkswagen	70,33	(2,2e-16)	71,59	(2,2e-16)
Hyundai	118,45	(2,2e-16)	120,57	(2,2e-16)
Peugeot	124,21	(2,2e-16)	126,44	(2,2e-16)
Ford	112,68	(2,2e-16)	114,70	(2,2e-16)
Nissan	83,68	(2,2e-16)	85,19	(2,2e-16)
Fiat	185,65	(2,2e-16)	188,10	(2,2e-16)
Skoda	139,44	(2,2e-16)	141,95	(2,2e-16)
Citroen	125,52	(2,2e-16)	127,78	(2,2e-16)

3.5 (S)ARIMA Modelling using Box-Jenkins Methodology.

The stationarity of the series should be checked first before evaluating any (S)ARIMA models. Therefore the Augmented Dickey Fuller (ADF) test and the KPSS test for stationarity of the series was evaluated (Table 3.8 on page 79). The KPSS test often select fewer differences than the ADF test. A KPSS test has the null hypothesis of stationarity whereas ADF test assume that the data have $I(1)$ non –stationarity. Consequently the KPSS test will only select one or more differences if there is enough evidence to overturn the stationarity assumption while the other tests will select at least one difference unless there is enough evidence to overturn the non-stationarity assumption. According to Hyndman and Athanasopoulos [2013] KPSS test led to models with better forecast.

Table 3.8: ADF & KPSS Stationarity Tests (log with-in-sample).

	Stationarity Tests			
	ADF	Num.Diff.	KPSS	Num.diff.
Opel	-3,98(0,01)	0	1,18(0,01)	1
Toyota	-4,00(0,01)	0	1,22(0,01)	1
VW	-2,97(0,17)	1	0,71(0,01)	1
Hyundai	-2,62(0,31)	1	2,84(0,01)	1
Peugeot	-2,75(0,26)	1	2,92(0,01)	1
Ford	-1,41(0,82)	1	1,41(0,01)	1
Fiat	-2,50(0,57)	1	1,32(0,01)	1
Nissan	-1,58(0,70)	1	1,35(0,01)	1
Skoda	-2,10(0,60)	1	1,25(0,01)	1
Citroen	-2,60(0,75)	1	2,70(0,01)	1

The seasonality Unit root test of OCSD is tested to our data but considering the ADF of the residuals of the model we will make the final decision for taking or not a difference

due to seasonality.

Table 3.9: OCSD Seasonality Unit Root Test. Log data, in-sample.

	OCSD Tests
	Seasonal Num.Diff.
Opel	0
Toyota	1
VW	1
Hyundai	0
Peugeot	0
Ford	0
Fiat	0
Nissan	1
Skoda	0
Citroen	1

The “with-in-sample” estimation continues with a methodology proposed in the famous book by Box and Jenkins [1976], that aims to find the most appropriate ARIMA (p,d,q), model. The Box and Jenkins [1976] methodology for the analysis and forecasting of time series is also extended to

$\text{SARIMA}(p, d, q)(P, D, Q)_s$

models as well, and generally it is widely regarded to be the most efficient technique for fitting a model and use it for forecasting. In practice, it is used extensively, especially for univariate time series. The Box-Jenkins technique requires, for each series, three steps:

Step 1. Identification.

Analyze the series and recognize a possible appropriate model.

Step 2. Estimation.

Model estimation and evaluation of the model’s coefficients.

Step 3. Diagnostic checking.

Test if the model is the best-estimated model.

These steps are going to be applied to find which among the ARIMA or seasonal ARIMA (SARIMA) models fits best our sample data. After the model identification process, the in-sample estimations for each-one of the 10 firms will result in the models' parameters. In-sample estimation means that we are considering all available observations in the data sample leaving nothing out of consideration. Finally, we proceed to the diagnostic checking of the models' residuals to identify if the models are adequate for fitting out the 10 different sample series.

3.5.1 STEP 1: Identification.

Identification is the process of recognizing a possibly appropriate model for time series after analyzing them Box and Jenkins [1976] provided both a theoretical framework and practical rules for determining appropriate values for their autoregressive (p) and moving average terms (q) as well as their seasonal counterparts P and Q. The only difficulty is that often more than one model could be entertained, requiring the researcher to choose one of them without any knowledge of the implications of that choice on the model's fitting or forecasting accuracy.

The order of the ARIMA model is found by examining the autocorrelations (ACF) and partial autocorrelations (PACF) functions of the stationary series. Results from the data plots of the ACF and PACF (see Figures in Appendix A) of the original series and additionally the unit root tests give evidence that the series is non-stationarity. Thus, we take the log values of the data to minimize the variations and the first difference of the series (i.e. $\Delta(x_t) = x_t - x_{t-1} = (1 - B)x_t$) to have stationary series and lower means and standard deviations.

We are graphically presenting the autocorrelation function (ACF) and partial autocorrelation function (PACF) correlograms of the log data series in all 10 different sample series for three different cases:

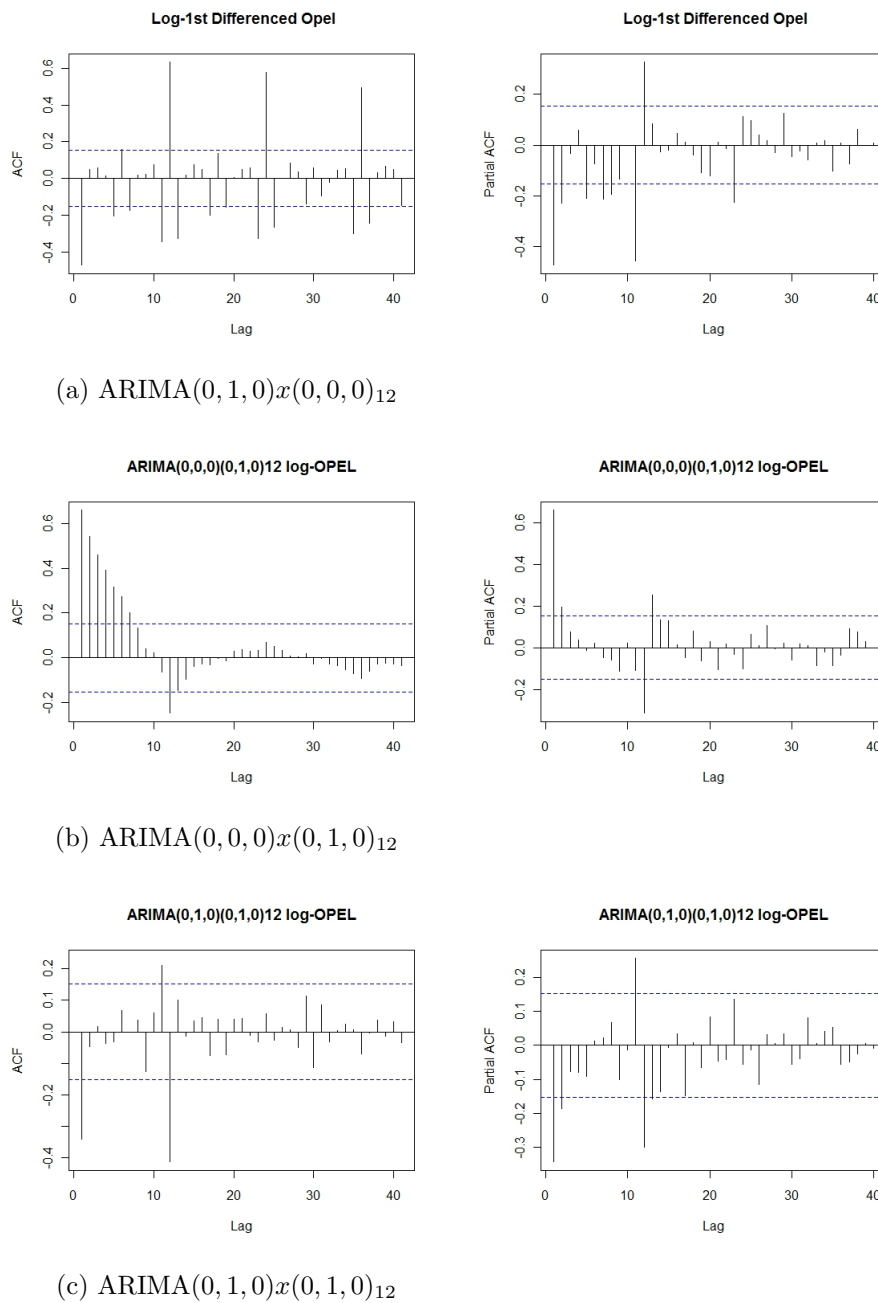


Figure 3.1: ACF and PACF for OPEL Sales (Log-values)

***Panel (a)** ARIMA(0, 1, 0) x (0, 0, 0)₁₂ models with constant, i.e. first difference data (d=1) with non-seasonal difference (D=0)

***Panel (b)** ARIMA(0, 0, 0) x (0, 1, 0)₁₂ models with constant, i.e. simple raw data (d=0) with non-seasonal difference (D=0) and

***Panel (c)** ARIMA(0, 1, 0) x (0, 1, 0)₁₂ models with constant, i.e. first difference data (d=1) with seasonal difference (D=1)

The ACF and PACF of the 3 different cases are illustrated for OPEL in Figure 3.1 in page 82 and in Appendix A for all 10 new car firms. Generally experienced researchers notice that:

- If the ACF plot cut off (or fast decline) the time series appears stationary.
- If the ACF plot shows a very slow linear decay pattern the time series appears non-stationary.
- If the ACF/PACF plot presents repeated waves the series indicate seasonality.
- If the ACF presents a pattern of changes of sign from one observation to the next the series shows signs of over differences.
- If the ACF plot exhibit a definite tendency to return to it mean the time series appears stationary.

The ACF and PACF plots of the non-seasonal difference are showed in Panel (a) for all series suggests that since all the ACF cut off (or fast decline) or tend to return to it mean the differenced series are stationary. Additionally, the ACF and PACF plots of the seasonal difference are shown in Panel (b) for all series and the repeated waves which are more obvious in the ACF of Peugeot, Nissan, Fiat, and Hyundai indicate the presence of seasonality in the new-car sales series. The log-differenced series (the residuals of a random-walk-with-growth model) look (more-or-less) stationary, but there is still very strong autocorrelation at the seasonal period (lag 12,24,36). The seasonal pattern is strong

and stable. Knowing that if the series has a strong and consistent seasonal pattern, we should use the order of seasonal differencing in the model we proceed in doing so. We must note that researchers suggest never to use more than one order of seasonal differencing or more than 2 orders of total differencing (seasonal plus non-seasonal).

The seasonal and non-seasonal differences of the series are shown in Panel (c) for all series and the results give signs of mild over-differencing. The positive spikes in the ACF and PACF have become negative. In more details the four types of the $SARIMA(p, d, q)x(P, D, Q)_s$ models are:

- * **Type A** $\rightarrow SARIMA(0, 0, 0)x(0, 0, 0)_{12}$ with constant.
- * **Type B** $\rightarrow SARIMA(0, 1, 0)x(0, 0, 0)_{12}$ with constant
- * **Type C** $\rightarrow SARIMA(0, 0, 0)x(0, 1, 0)_{12}$ with constant
- * **Type D** $\rightarrow SARIMA(0, 1, 0)x(0, 1, 0)_{12}$ with constant

In other words we have four different models.

- In model type A we have the original data in log values.
- In model type B we have the log data in first difference.
- In model type C we have the log data in first seasonal differences and
- In model type D we have the log-data in first differences and first seasonal differences.

Data order of differencing selection criteria for modelling. The problem of choosing the correct order of differencing is solved by the calculation of the variance of the series at each level of differencing. The variance is just the minimum squared errors (MSE) that results from fitting the various difference only ARIMA models. The root of the minimum squared errors (RMSE) is actually the standard deviation of the series at each level of differencing.

Table 3.10: RMSE and MAPE for different SARIMA models in log new-car sales series.

	RMSE			
TYPE	A	B	C	D
SARIMA	$(0, 0, 0)x(0, 0, 0)_{12}$	$(0, 1, 0)x(0, 0, 0)_{12}$	$(0, 0, 0)x(0, 1, 0)_{12}$	$(0, 1, 0)x(0, 1, 0)_{12}$
Opel	0,4131	0,4067	0,3672	0,2950
Toyota	0,4366	0,4988	0,3621	0,3550
VW	0,3999	0,4022	0,3673	0,3440
Hyundai	0,5664	0,4434	0,3958	0,3622
Peugeot	0,6098	0,4338	0,4072	0,3405
Ford	0,5502	0,3971	0,4006	0,3805
Nissan	0,4559	0,4267	0,4023	0,3817
Fiat	0,5169	0,3822	0,3848	0,3318
Skoda	6,5223	0,4329	0,4634	0,3256
Citroen	6,5589	0,5559	0,4016	0,3780
	MAPE%			
Opel	3,3700	4,1643	3,7301	2,8970
Toyota	4,5928	4,6716	3,5155	3,4765
VW	4,2442	3,9956	3,4953	3,4232
Hyundai	6,5209	4,6415	4,1317	3,6541
Peugeot	7,4702	4,5758	4,6189	3,7568
Ford	6,7221	4,4245	4,2560	4,0985
Nissan	4,9813	4,7380	4,3350	4,1643
Fiat	6,0464	4,2483	4,2582	3,6764
Skoda	8,6561	5,2911	1,0149	4,0343
Citroen	7,9280	5,6541	4,5652	4,0012

The results are given in Table 3.10 at page 85, show that the lowest RMSE is obtained by model D which uses one difference of each type, i.e. first difference for the non-seasonal

and first difference for the seasonal counterparts. This, together with the appearance of the relevant plots of ACF and PACF in Figure 3.1 page 82 strongly suggests that we should use both a seasonal and a non-seasonal difference in our data. The analysis of the series using unit root tests (Table 3.8 provide evidence for the need of the first difference in our data. On the other hand the official seasonal unit root test results (Table 3.9 based on the OCSB test, yield to first seasonal difference only for four (4) of the ten(10) new car sales series i.e. Toyota, Volkswagen, Nissan, Citroen new-car sales series.

It is worth noticing that, except for the gratuitous constant term, model D is the *Seasonal Random Trend* (SRT) model, whereas model B is just the *Seasonal Random Walk* (SRW) model. Comparing these two models, the SRT model appears to fit better than the SRW model. Model C (SARIMA(0, 0, 0) x (0, 1, 0)₁₂) is not going to be considered seriously because it has the original series without any differences (d=0) which is not correct since we have proven that we should take the first difference (d=1) of the series, based on KPSS unit root test.

Furthermore if we notice the ACF of model C (Panel b in Figure 3.1 are slowly decreasing for almost all series, which is an indication that the mean of the series is not stationary. On the contrary, in model (a), SARIMA(0, 1, 0) x (0, 0, 0)₁₂) the ACF functions (Panel b in Figures 3.1 appears approximately stationary, since the autocorrelation functions exhibit a definite tendency to return to the series mean and exhibit no long-term trend.

The ACF for model D (Panel c in Figures 3.1) shows signs of over-differencing especially for the time series of new car sales of Peugeot, Hyundai, Fiat, and Skoda. We notice a pattern of changes of sign from one observation to the next (positive-negative -positive-negative and so on) at the ACF plot. Additionally ACF in model D has a negative spike at lag 1 that is close to 0.5 in magnitude.

Conclusion Remarks. In the analysis that follows, we will try to improve two models:

1. Model B: SARIMA(0, 1, 0) x (0, 0, 0)₁₂for Opel, Hyundai, Peugeot, Ford, Fiat, Skoda new-car sales series and,
2. Model D:SARIMA(0, 1, 0) x (0, 1, 0)₁₂ for Toyota, Volkswagen, Nissan, Citroen new-car

sales series,

through the addition of seasonal and non-seasonal ARIMA terms, comparing which model fit better to each time series in an "in sample" estimation.

For each of the series we will try to investigate which among a variety of Seasonal Autoregressive Moving Average (SARIMA) models best fit our data. In SARIMA(p, d, q) $x(P, D, Q)_s$ model general equation (2.11) the error term ε_t is initially assumed to be a Gaussian white noise process, with zero mean and a constant variance σ^2 . We allow for a non-seasonal first differences ($d=1$) in all 10 different time series, according to the results from Table 3.8 (page 79), and a seasonal first difference ($D=1$) only in the case of Toyota, Volkswagen, Nissan and Citroen due to the results from Table 3.9 on page 80.

SARIMA Model Selection Criteria. Since we are dealing with monthly data we set $s=12$ and allow the autoregressive (p) and moving average (q) orders to take values from 0 to 5, and the equivalent seasonal P and Q to take values from 0 to 2. If the values of p, q, P , and Q are allowed to range more widely, the number of possible models increases rapidly.

Since it is sometimes not possible to identify the parameters p, d, q and P, D, Q using visualization tools, such as ACF and PACF plots, researchers use different selection criteria, like the Akaike Information Criteria (AIC) and/or the Schwartz Bayesian Criteria (SBC or BIC). Using these techniques, the problem of over-fitting in the modeling time series, which leads to less precise estimators and bad forecasts, is controlled:

- Using BIC we select the (S)ARIMA model with the lowest value of the BIC
- Using AIC we select the (S)ARIMA model with the lowest value of the AIC.

In this empirical study, the AIC criterion is used and our scope is to identify which model better fits our data and proceed to the model selection that generates the best "in sample" estimates. The number of parameters that yield the minimum specifies the best model for each one of the car representatives. The AIC information criterion which

helps in selecting the orders of p , q , P and Q for every series in different SARIMA models specification is calculated as:

$$AIC = -2\log(L) + 2k \quad (3.4)$$

where L is the maximized likelihood of model and k is the number of estimated parameters (including the variance).

Calculating the AIC values for each model with the same data set will give us the “best” model which is the one with the minimum AIC value. The value of the AIC depends on the data series x_t , which leads to model selection uncertainty. Table 3.11 on page 91 gives the best model for Opel, Hyundai, Peugeot, Ford, Fiat, and Skoda series while Table 3.13 on page 92 gives the best model for Toyota, Volkswagen, Nissan, and Citroen. Briefly, the best SARIMA models for our series are:

Opel Model:	SARIMA(0, 1, 1)(1, 0, 1) ₁₂
Toyota Model:	SARIMA (2, 1, 3)(1, 1, 1) ₁₂
Volkswagen Model:	SARIMA (0, 1, 1)(2, 1, 0) ₁₂
Hyundai Model:	SARIMA (1, 1, 1)(2, 0, 1) ₁₂
Peugeot Model:	SARIMA (0, 1, 1)(1, 0, 1) ₁₂
Ford Model:	SARIMA (0, 1, 1)(1, 0, 1) ₁₂
Nissan Model:	SARIMA (0, 1, 4)(0, 1, 1) ₁₂
Fiat Model:	SARIMA (2, 1, 1)(1, 0, 1) ₁₂
Skoda Model:	SARIMA (0, 1, 1)(2, 0, 1) ₁₂
Citroen Model:	SARIMA (0, 1, 2)(0, 1, 1) ₁₂

We notice that each new vehicle representative has a different SARIMA model specification. That is due to the different operational and marketing strategy of each car firm. Each car retailer in the market may maintain, increase or decrease their market share in the Greek new car market. However, there are some car firms that share the same type of model specification for their sales level. Opel, Peugeot and Ford new-car sales in Greece are best modelled by the SARIMA (0, 1, 1)(1, 0, 1)₁₂. That means that these three firms share the same general formula for estimating there new-car sales levels, produced from the equation (2.11).

On the other hand we have different coefficients estimations for each one of the car representatives, due to the data variety, which brings about 10 different models as illustrated in Table 3.15. Grouping the SARIMA models formulation for the 10 different new car representatives, leads to only 8 different types of models formulations (since Opel Peugeot and Ford share the same type of model) which are illustrated below:

- Model 1

Opel & Peugeot & Ford new-car sales $\implies SARIMA(0, 1, 1)(1, 0, 1)_{12}$:

$$(1 - B)(1 - \Phi_1 B^{12})x_t = (1 + \theta_1 B)(1 + \Theta_1 B^{12})\varepsilon_t \quad (3.5)$$

- Model 2

Toyota new-car sales $\implies SARIMA(2, 1, 3)(1, 1, 1)_{12}$:

$$(1 - B)(1 - B^{12})x_t = (1 + \theta_1 B + \theta_2 B^2 + \theta_3 B^3)(1 + \Theta_1 B^{12})\varepsilon_t \quad (3.6)$$

- Model 3

Volkswagen new-car sales $\implies SARIMA(0, 1, 1)(2, 1, 0)_{12}$:

$$(1 - B)(1 - B^{12})(1 - \Phi_1 B^{12} - \Phi_2 B^{24})x_t = (1 + \theta_1 B)\varepsilon_t \quad (3.7)$$

- Model 4

Hyundai new-car sales $\implies (1, 1, 1)(2, 0, 1)_{12}$:

$$(1 - \phi_1 B)(1 - B)(1 - \Phi_1 B^{12} - \Phi_2 B^{24})x_t = (1 + \theta_1 B)(1 + \Theta_1 B^{12})\varepsilon_t \quad (3.8)$$

- Model 5

Nissan new-car sales $\implies SARIMA(0, 1, 4)(0, 1, 1)_{12}$:

$$(1 - B)(1 - B^{12})x_t = (1 + \theta_1 B + \theta_2 B^2 + \theta_3 B^3 + \theta_4 B^4)(1 + \Theta_1 B^{12})\varepsilon_t \quad (3.9)$$

- Model 6

Citroen new-car sales $\implies SARIMA(0, 1, 2)(0, 1, 1)_{12}$:

$$(1 - B)(1 - B^{12})x_t = (1 + \theta_1 B + \theta_2 B^2)(1 + \Theta_1 B^{12})\varepsilon_t \quad (3.10)$$

- Model 7

Skoda new-car sales $\implies SARIMA(0, 1, 1)(2, 0, 0)_{12}$:

$$(1 - B)(1 - \Phi_1 B^{12} - \Phi_2 B^{24})x_t = (1 + \theta_1 B)\varepsilon_t \quad (3.11)$$

- Model 8

Fiat new-car sales $\implies SARIMA(2, 1, 1)(1, 0, 1)_{12}$:

$$(1 - \phi_1 B - \phi_2 B^2)(1 - B)(1 - \Phi_1 B^{12})x_t = (1 + \theta_1 B)(1 + \Theta_1 B^{12})\varepsilon_t \quad (3.12)$$

where B is the lag operator and ε_t is the error term at time t (i.e. the residuals of each model). After completing the estimation process, we will rewrite the above models replacing the unknown parameters with the ones estimated using the Maximum Likelihood estimation process. So ϕ_p will be replaced with the estimated parameters for the $AR(p)$, θ_q with the estimated parameters of $MA(q)$, Φ_P with the estimated parameters of $SAR(P)$, and Θ_Q with the estimated parameters of $SMA(Q)$.

Table 3.11: AIC information criterion for different SARIMA specifications ($d=1, D=0$).

SARIMA	Opel	Hyundai	Peugeot	Ford	Fiat	Skoda
$(0, 1, 1)(0, 0, 1)_{12}$	80,76	117,90	110,92	95,28	88,54	83,96
$(0, 1, 2)(0, 0, 1)_{12}$	82,76	119,77	112,44	96,64	89,99	85,96
$(2, 1, 0)(0, 0, 1)_{12}$	83,02	121,27	110,71	99,02	89,01	88,84
$(3, 1, 1)(1, 0, 0)_{12}$	45,03	78,83	74,56	84,08	63,85	66,74
$(1, 1, 1)(1, 0, 1)_{12}$	4,89	43,30	52,15	-	-	34,23
$(1, 1, 2)(1, 0, 1)_{12}$	2,75	45,90	53,77	69,37	22,12	36,20
$(3, 1, 2)(1, 0, 1)_{12}$	8,21	46,92	50,91	71,71	21,74	34,20
$(2, 1, 2)(1, 0, 1)_{12}$	6,72	46,90	55,86	-	-	37,34
$(1, 1, 1)(1, 0, 2)_{12}$	6,40	40,92	54,03	69,03	21,58	35,71
$(0, 1, 1)(2, 0, 1)_{12}$	5,24	40,46	52,01	68,26	19,82	34,23
$(2, 1, 1)(2, 0, 0)_{12}$	15,89	45,43	54,22	74,54	38,05	59,48
$(1, 1, 2)(2, 0, 2)_{12}$	6,09	43,30	-	72,4	23,77	36,64
$(0, 1, 2)(2, 0, 2)_{12}$	8,06	43,04	-	70,78	21,93	34,47
$(2, 1, 3)(1, 0, 2)_{12}$	9,99	45,34	53,96	-	23,88	39,99
$(0, 1, 1)(2, 0, 0)_{12}$	14,68	44,89	55,02	72,38	37,78	56,43
$(2, 1, 3)(1, 0, 1)_{12}$	8,68	46,92	51,85	72,62	22,45	38,21
$(0, 1, 1)(1, 0, 1)_{12}$	3,60	45,14	50,16	66,87	18,42	32,52
$(0, 1, 2)(1, 0, 2)_{12}$	6,66	41,57	54,03	69,03	21,59	35,69
$(0, 1, 1)(1, 0, 2)_{12}$	5,28	40,84	52,04	68,17	19,61	34,06
$(1, 1, 1)(2, 0, 1)_{12}$	7,38	40,39	-	-	22,99	40,12
$(1, 1, 1)(2, 0, 0)_{12}$	15,51	44,33	56,93	72,54	39,39	57,91

Table 3.13: AIC information criterion for different SARIMA specifications($d=1, D=1$).

SARIMA	Toyota	Volkswagen	Nissan	Citroen
$(0, 1, 1)(0, 1, 1)_{12}$	61,33	34,93	79,24	66,09
$(0, 1, 4)(0, 1, 1)_{12}$	60,57	39,84	72,23	66,67
$(2, 1, 0)(0, 1, 1)_{12}$	60,74	41,11	82,62	76,27
$(3, 1, 1)(1, 1, 0)_{12}$	76,86	38,34	89,38	85,85
$(1, 1, 1)(1, 1, 1)_{12}$	63,30	29,36	75,24	65,18
$(1, 1, 2)(1, 1, 1)_{12}$	60,99	30,76	75,92	66,68
$(3, 1, 2)(1, 1, 1)_{12}$	63,17	33,93	77,51	70,73
$(2, 1, 2)(1, 1, 1)_{12}$	61,81	32,75	77,89	68,73
$(1, 1, 1)(1, 1, 2)_{12}$	65,16	31,16	76,58	65,45
$(0, 1, 1)(2, 1, 1)_{12}$	64,67	29,53	79,84	69,51
$(2, 1, 1)(2, 1, 0)_{12}$	70,14	31,02	74,44	-
$(1, 1, 2)(2, 1, 2)_{12}$	64,96	34,22	77,14	70,25
$(0, 1, 2)(0, 1, 1)_{12}$	61,39	36,67	76,81	62,72
$(2, 1, 3)(1, 1, 2)_{12}$	59,39	30,41	80,65	-
$(0, 1, 1)(2, 1, 0)_{12}$	71,70	27,77	78,42	80,84
$(2, 1, 3)(1, 1, 1)_{12}$	57,66	29,26	79,21	70,68
$(0, 1, 3)(0, 1, 1)_{12}$	59,96	38,28	78,28	64,69
$(0, 1, 2)(1, 1, 2)_{12}$	64,09	31,10	78,57	65,03
$(0, 1, 1)(1, 1, 2)_{12}$	64,67	29,72	81,17	68,06
$(1, 1, 1)(2, 1, 1)_{12}$	65,29	29,72	75,30	66,74
$(1, 1, 1)(2, 1, 0)_{12}$	72,36	30,09	73,60	77,03

3.5.2 STEP 2: Estimation

In the estimation process we follow the Box-Jenkins methodology. This methodology supports the principle of parsimony, meaning that a simpler (having fewer parameters) model should be selected in case more than one model are possible. The final models are selected via the AIC information criterion and computed using maximum likelihood estimation. The results are the following for each one of the new-car firm:

- OPEL (S)ARIMA $(0, 1, 1)(1, 0, 1)_{12}$ Model

$$(1 - B)(1 - 0.97B^{12})x_t = (1 - 0.54B)(1 - 0.71B^{12})\varepsilon_t$$
- TOYOTA (S)ARIMA $(2, 1, 3)(1, 1, 1)_{12}$ Model

$$(1 - 1.64B - 0.86B^2)(1 - B)(1 - B^{12})(1 - 0.12B^{12})x_t = (1 + 1.13B - 0.23B^2 - 0.67B^3)(1 - 0.75B^{12})\varepsilon_t$$
- VOLKSWAGEN (S)ARIMA $(0, 1, 1)(2, 1, 0)_{12}$ Model

$$(1 - B)(1 - B^{12})(1 + 0.84B^{12} + 0.31B^{24})x_t = (1 - 0.62B)\varepsilon_t$$
- HYUNDAI (S)ARIMA $(1, 1, 1)(2, 0, 1)_{12}$ Model

$$(1 - 0.21B)(1 - B)(1 - 0.64B^{12} - 0.29B^{24})x_t = (1 - 0.74B)(1 - 0.47B^{12})\varepsilon_t$$
- PEUGEOT (S)ARIMA $(0, 1, 1)(1, 0, 1)_{12}$ Model

$$(1 - B)(1 - 0.95B^{12})x_t = (1 - 0.55B)(1 - 0.63B^{12})\varepsilon_t$$
- FORD (S)ARIMA $(0, 1, 1)(1, 0, 1)_{12}$ Model

$$(1 - B)(1 - 0.92B^{12})x_t = (1 - 0.66B)(1 - 0.68B^{12})\varepsilon_t$$
- NISSAN (S)ARIMA $(0, 1, 4)(0, 1, 1)_{12}$ Model

$$(1 - B)(1 - B^{12})x_t = (1 - 0.49B - 0.09B^2 - 0.007B^3 - 0.22B^4)(1 - 0.63B^{12})\varepsilon_t$$
- FIAT (S)ARIMA $(2, 1, 1)(1, 0, 1)_{12}$ Model

$$(1 - 0.47B - 0.34B^2)(1 - B)(1 - 0.98B^{12})x_t = (1 - 0.98B)(1 - 0.77B^{12})\varepsilon_t$$
- SKODA (S)ARIMA $(0, 1, 1)(2, 0, 0)_{12}$ Model

$$(1 - B)(1 - 0.42B^{12} - 0.27B^{24})x_t = (1 - 0.50B)\varepsilon_t$$

- CITROEN (S)ARIMA $(0, 1, 2)(0, 1, 1)_{12}$ Model

$$(1 - B)(1 - B^{12})x_t = (1 - 0,54B - 0,19B^2)(1 - 0,93B^{12})\varepsilon_t$$

The results from the estimation process are given in more details in Table 3.15, page 95.

After the estimation of the parameters the SARIMA models general formulation is given in Equations (3.5) to (3.12) can now be transformed into a more accurate presentation as the one already developed above for each one of the new-car firm.

3.5.3 STEP 3: Diagnostics checking.

The diagnostic check is a procedure that is used to check the residuals of a model. After the model specification, the diagnostic checking is employed by examining the residuals from the fitted model to see if the model specification is adequate. The residuals should fulfill the model's assumptions.

For SARIMA models residuals should be random, independent, and normally distributed. Suppose these assumptions were not fulfilled then the researcher choose another model for the series. In terms of the Box-Jenkins methodology, any model which results in random residuals, is an appropriate one.

Once an appropriate model had been entertained and its parameters estimated, the Box-Jenkins methodology required examining the residuals of the actual values minus those estimated through the model. If such residuals are random, it is assumed that the model is appropriate. If not, another model is entertained, its parameters estimated, and its residuals checked for randomness.

In practically all instances, a model should be found to result in random residuals since it is a standard statistical procedure not to use models whose residuals are not random. The implementation of the diagnostic checking can be done both by graphical and statistical testing.

Graphical Testing of SARIMA models Residuals. The Graphical testing includes a variety of plots like :

SARIMA	AR1	AR2	MA1	MA2	MA3	MA4	SAR1	SAR2	SMA1	σ^2	Loglik
MODELS	ϕ_1	ϕ_2	θ_1	θ_2	θ_3	θ_4	Φ_1	Φ_2	Θ_1		
Opel	-	-	-0,543	-	-	-	0,976	-	-0,717	0,054	-2,47
(0, 1, 1)(1, 0, 1) ₁₂	-	-	(0,07)	-	-	-	(0,01)	-	(0,08)		
Toyota	-1,645	-0,861	1,138	-0,233	-0,678	-	-0,122	-	-0,758	0,068	-20,83
(2, 1, 3)(1, 1, 1) ₁₂	(0,048)	(0,052)	(0,075)	(0,129)	(0,073)	(0,073)	(0,126)		(0,114)		
VW	-	-	-0,621	-	-	-	-0,847	-0,311	-	0,062	-9,89
(0, 1, 1)(2, 1, 0) ₁₂			(0,070)				(0,093)	(0,099)			
Hyundai	0,212	-	-0,745	-	-	-	0,644	0,296	-0,474	0,062	-13,3
(1, 1, 1)(2, 0, 1) ₁₂	(0,132)		(0,745)				(0,136)	(0,118)	(0,140)		
Peugeot	-	-	-0,555	-	-	-	0,956	-	-0,631	0,068	-20,29
(0, 1, 1)(1, 0, 1) ₁₂			(0,072)				(0,024)		(0,083)		
Ford	-	-	-0,660	-	-	-	0,929	-	-0,681	0,079	-29,30
(0, 1, 1)(1, 0, 1) ₁₂			(0,069)				(0,044)		(0,103)		
Nissan	-	-	-0,495	-0,092	-0,007	-0,223	-	-	-0,631	0,082	-30,11
(0, 1, 4)(0, 1, 1) ₁₂			(0,078)	(0,094)	(0,080)	(0,079)			(0,072)		
Fiat	0,473	0,341	-0,986	-	-	-	0,983	-	-0,778	0,055	-4,23
(2, 1, 1)(1, 0, 1) ₁₂	(0,078)	(0,078)	(0,016)				(0,013)		(0,080)		
Skoda	-	-	-0,506	-	-	-	0,429	0,272	-	0,074	-23,84
(0, 1, 1)(2, 0, 0) ₁₂			(0,069)				(0,078)	(0,085)			
Citroen	-	-	-0,547	-0,190	-	-	-	-	-0,934	0,071	-27,36
(0, 1, 2)(0, 1, 1) ₁₂			(0,080)	(0,081)					(0,210)		

Table 3.15: SARIMA models for new-car sales "in sample" estimation.

- Histogram of Residuals
- Q-Q plot of Residuals
- Standardize Residuals Plot
- ACF plot of Residuals
- p-value plot of Ljung-Box test
- Squared Standardized Residuals Plot
- ACF and PACF plots of squared Residuals

Visually a histogram of the residuals, a normal probability plot or a normal quantile-quantile plot (Q-Q plot) of the squared residuals can help in identifying departures from normality and investigates marginal normality [Johnson and Wichern, 1992].

Histograms. The overall pattern of the *histogram of the residuals* should be similar to the bell-shaped pattern observed when plotting a histogram of normally distributed data. Departures from these shapes usually mean that the residuals contain a structure that is not accounted for in the model.

Identifying that structure and altering the model may lead to a better model. The histogram of the residuals of the SARIMA models fit each firms' new car sales data indicate that the residuals are close to normality as they are bell-shaped except for a few extreme values in the tails. The histograms of the SARIMA model residuals are graphically presented in the first figure at the top left angle of Fig.3.2 (page 100) for Opel.

Q-Q plot of Residuals. The normal *quantile-quantile plot (Q-Q plot)* of the residuals is shown in the top right angle of the Figures indicated above for all 10 different series and since the variance of the residuals σ_t^2 , is unobservable, the squared residuals serve as a proxy.

The Q-Q plot compares the ordered values of each variable with the quantiles of a specific theoretical distribution (for example the normal distribution). If the two distributions

match the points on the plot, it will form a linear pattern passing through the origin with a unit slope. The Q-Q plots are used to see how well a theoretical distribution models the empirical data.

All car firms appear to have a small variation in their new car sales level with a wider range of outliers in the upper and lower extreme values. In other words, we observe some variations in the sales levels for the months that we have the higher and the lower sales levels for each firm. However, the linearity of the points of the normal Q-Q plots suggests that the data are very close to normal distribution but also give some evidence of heteroskedasticity.

Standardize Residuals Plot. Firstly we observe the time plot of the innovations (residuals) $\varepsilon_t = x_t - \hat{x}_{t-1}$ or the time plot of the standardized innovations. The *standardized* residual (ε_t^*) is the residual divided by its standard deviation:

$$\varepsilon_t^* = \frac{x_t - \hat{x}_{t-1}}{\hat{\sigma}}$$

where $\hat{x}_{t-1}$ is the “one-step-ahead” prediction of $\hat{x}_{t-2}$, based on the fitted SARIMA models and $\hat{\sigma}$ is the estimated squared root error variance (or the standard deviation) from the selected SARIMA model.

In this way we are scaling the residuals by simply standardizing them, meaning scale the residuals by dividing with the constant standard error. If a model fits well, the standardized residuals should behave as an identical independently distributed (iid) sequence, with mean zero ($\mu = 0$) and variance equal to one ($\sigma^2 = 1$). The time plot of the series is inspected for any obvious departures from this assumption. There is evidence that the SARIMA models do not fit very well; since there are large and small deviations from the mean in different periods.

We observe that the standardized residuals plot analysis shows signs of heteroskedasticity with a mean around zero. The time plots of the standardized residuals are illustrated in the first out of the three plots in the second half of the Figure 3.2 (page 100) for Opel.

ACF of Residuals. Since the third stage, in building SARIMA models, consists of validating the model through an examination of the one-step prediction residuals ε_t , a

basic visual diagnostic technique should also be to examine the autocorrelation function (ACF) of the residuals.

Checking the correlation structure of the residuals, we plot the sample autocorrelations of the residuals illustrated in the second out of the three plots in the second half of the Figure 3.3 (page 100) for Opel and so on for the other firms.

Generally the presence of large autocorrelations indicates that the models may be inadequate. It is well known that for random and independent series of length n , the lag k autocorrelation coefficient is normally distributed with a mean of zero and variance of $\frac{1}{n}$ and the 95% confidence limits are given by $\pm \frac{1.96}{\sqrt{n}}$. For a white noise sequence, the sample autocorrelations are approximately independently and normally distributed with zero means and variances $\frac{1}{n}$.

However, residuals from a model fit, may not have the exact properties of a white noise sequence and the variance of $\hat{r}_k(\varepsilon)$ can be much less than $\frac{1}{n}$ [Box and Pierce, 1970, McLeod and Hipel, 1978].

The detection of the obvious departures from the independence assumption can be visually inspected. The ACF plots of the SARIMA squared residuals in all ACF figures show that there is no significant autocorrelation left in the residuals from the SARIMA type models of the monthly new vehicle sales. In other words, the autocorrelation function give the impression that residuals are purely random.

According to Shumway and Stoffer [1982] evidence that the residuals are uncorrelated is enough only if we know that time series is Gaussian i.e normally distributed. Because there are examples where residuals are uncorrelated but time series are non-Gaussian, like in the case of the GARCH family models.

Ljung-Box p-value plot. The levels of p-values of the Ljung-Box p-statistic in the third out of the three plots in the second half of the figure 3.2 (page 100) for Opel and so on for the other firms give evidence of :

1. High p-values which indicate no autocorrelation in the residuals (in case of Opel, Toyota, Hyundai, Ford, and Fiat) and

2. Low-level p-values or p-values with fluctuations as the level of lags increases, which indicates that there is some autocorrelation left in the residuals (cases of Volkswagen, Peugeot, Skoda, Nissan, and Citroen).

In general the p-value's exceedance of 0,05 indicate the acceptance of the null hypothesis of the model adequacy at significance level 0,05 [Wang et al., 2005]. Results show that the SARIMA models may not be adequate for all the sample series.

Squared Standardized Residuals Plots and their ACF and PACF Plots. There is a tendency in the data that large (small) absolute values of residual process are followed by other large (small) values of unpredictable sign, which is a common behaviour of GARCH processes. It was Granger and Andersen [1978] that have found that some of the series modelled by Box and Jenkins [1976] exhibit autocorrelated squared residuals even though the residuals themselves do not seem to be correlated over time and therefore suggested that the ACF of the squared time series could be useful in identifying non-linear time series. Bollerslev [1986] stated that the ACF and PACF of squared process are useful in identifying and checking GARCH behaviour. In Figure 3.4(page101) and in Figure 3.3(page 100) we have the ACF and the PACF of the squared standardized residual series from the selected SARIMA models of monthly new car sales in Greece.

It is shown that although the residuals are almost uncorrelated as shown in the plot of the simple residuals Figure 3.2(page 100), the squared residuals series illustrated in Figure 3.4 are correlated illustrating high volatility especially during the last 3 years. Additionally the ACF and PACF plot of the squared standardized residual series exhibit seasonality effect, which has been taken care by the seasonal ARIMA modelling and structure that is not modelled by the selected SARIMA models. As a consequence, we can assume that given the structure of the ACF and PACF there is an indication that the variance of residuals series may exhibit an ARCH effect.

The Statistical testing of SARIMA models residuals includes a variety of tests like the followings:

- Box - Pierce Test

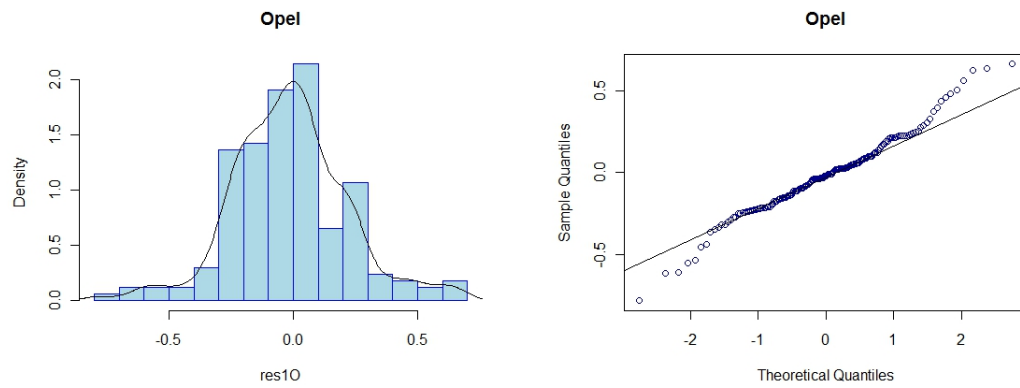


Figure 3.2: OPEL-SARIMA residual (a)Histogram,(b)Q-Q Plot

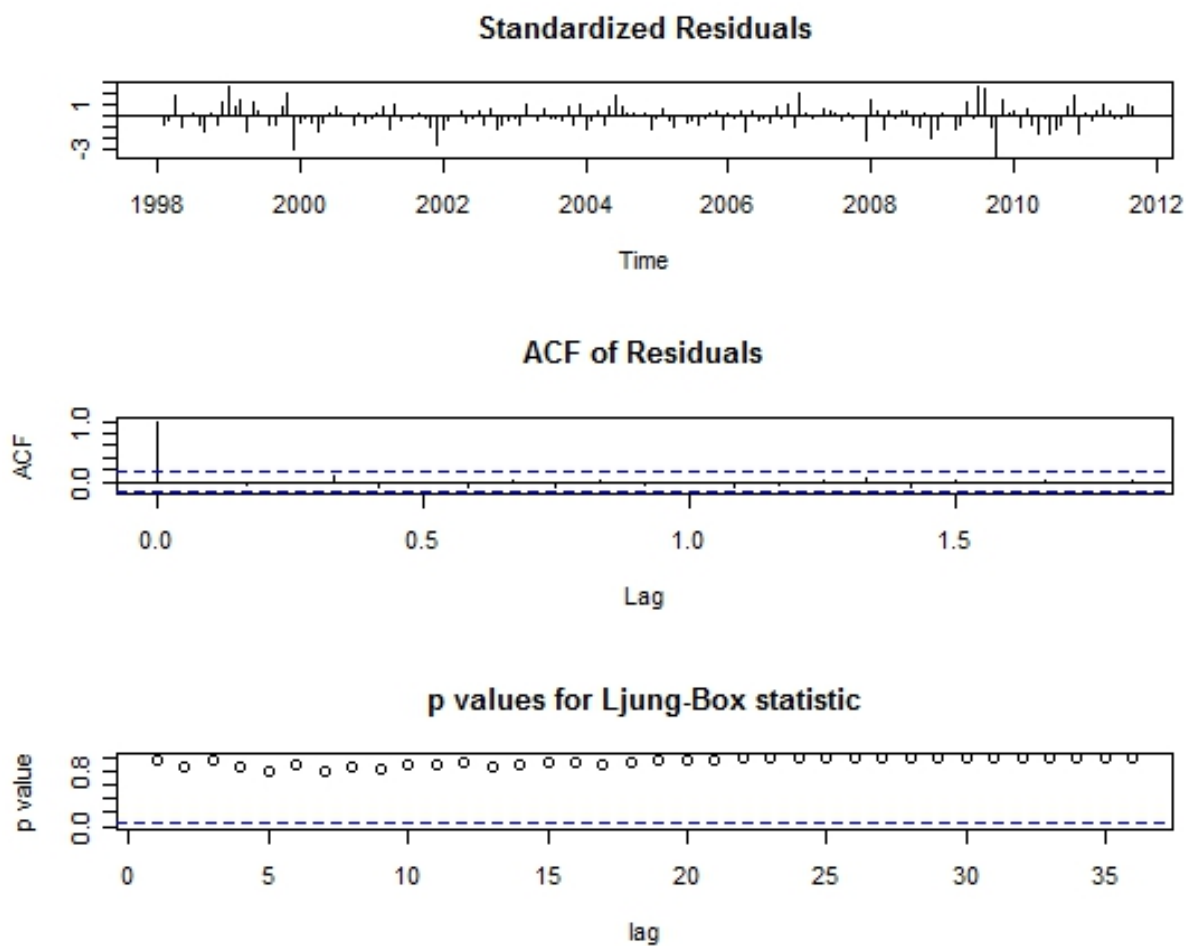


Figure 3.3: Opel-SARIMA residual(a) Standardized (b)ACF (c)Ljung-Box

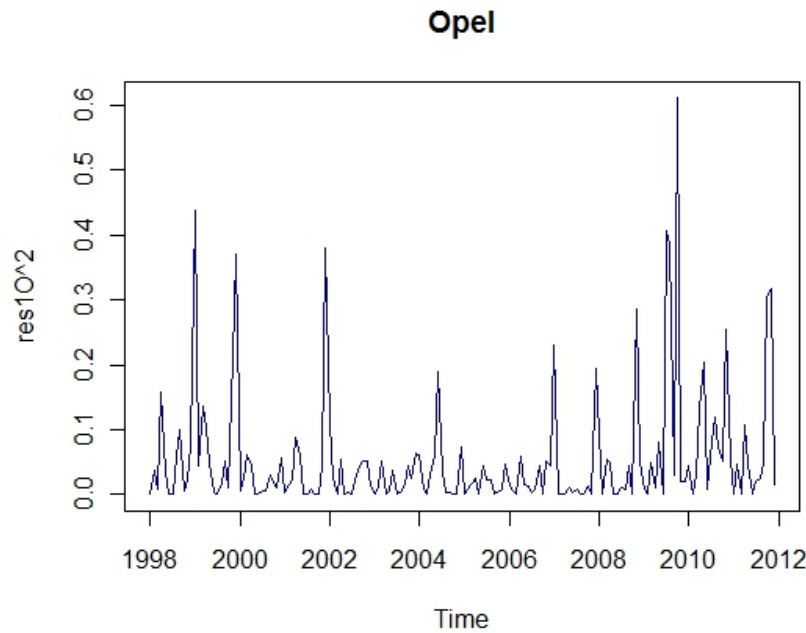


Figure 3.4: Opel SARIMA model Squared and Standardized Residual Plots

- Ljung - Box Test
- Shapiro - Wilk Test
- Jarque - Bera Test

The Box and Pierce [1970] and Ljung and Box [1978] Q -statistics is often adopted to test the adequacy of the model. After the graphical examination of the residuals, we continue and instead of testing the significance of any individual autocorrelation coefficient via the ACF, we test the joint hypothesis that all the autocorrelation coefficients (ρ_k) up to a certain lag are simultaneously equal to zero. Box-Pierce and the Ljung-Box test statistic for examining the null hypothesis of independence in a given time series is based on the autocorrelation plot. However, instead of testing the randomness at each distinct lag, it tests the overall randomness based on several lags. For this reason, they are sometimes known as a “portmanteau” test.

The Box - Pierce Test (Q_{BP} -statistic) developed by Box and Pierce [1970] is defined

as:

$$Q_{BP} = n \sum_{k=1}^m \hat{\rho}_k^2 \quad (3.13)$$

where n =sample size and m =the number of autocorrelations included in the sample and $\hat{\rho}_k^2$ is the squared sample autocorrelation of residual series ε_t at lag k . The Box and Pierce [1970] Q_{BP} -statistics test has the null hypothesis that there is no autocorrelation up to order 6, 12, and 24 computed on $\hat{\varepsilon}_t^2 \hat{\sigma}_t^{-2}$. The hypothesis of no autocorrelation up to lag m is rejected if the Q statistic is greater than the appropriate percentile of the χ^2 -square distribution with $m - p - q$ degrees of freedom (where $m > p + q$, p the autocorrelation order and q the moving average order of the SARIMA model).

The Ljung and Box [1978] Test (Q_{LB} -statistics) is a variant of the Box and Pierce [1970] Test (Q_{BP} statistic) and is defined as:

$$Q_{LB} = n(n+2) \sum_{k=1}^m \left(\frac{\hat{\rho}_k^2}{n-k} \right) \quad (3.14)$$

where n =sample size and m =the number of autocorrelations included in the sample and $\hat{\rho}_k^2$ is the squared sample autocorrelation of residual series ε_t at lag k .

Under the null hypothesis of model adequacy, the Q_{LB} test statistic is asymptotically $\chi^2(m - p - q)$ distributed. The null hypothesis at level α is rejected if the value of Q_{LB} exceeds the $(1 - \alpha)$ quantile of the $\chi^2(m - p - q)$ distribution. In other words, if the computed Q_{BP} exceeds the critical Q value from the chi-square distribution at the chosen level of significance (α), one can reject the null hypothesis that all the (true) ρ_k are zero, meaning that at least some of them must be non-zero. In this test, a small p -value is an evidence that there is dependence. So we want to see large p -values. However, a large p -value is not evidence of independence, merely a lack of evidence of dependence.

In a large sample, both Q_{BP} and Q_{LB} follow the chi-square distribution with m degrees of freedom. In small samples, Q_{LB} has been found to have more powerful, in a statistical sense, 'small sample properties' than the Q_{BP} statistic [Gujarati, 2003].

BP and LB Test of SARIMA Residuals. The Q -statistics tests and their p -values for the non-autocorrelation of residuals of the selected SARIMA models in modeling the

new car sales for Q(6), Q(12) and Q(24) statistic are given in Table 3.16 on page 107. The sum of the Q(6), Q(12) and Q(24)squared autocorrelations of the selected SARIMA residuals as shown by the Box-Pierce Q_{BP} and Ljung Box Q_{LB} statistics is not statistically significant for a level $\alpha = 0,10$ since the Q statistics calculated are less than the critical values of the χ^2 distribution 10,64 for Q(6), 18,54 for Q(12) and 33,19 for Q(24).

However, if we increase the error term to $\alpha = 0,5$ then the critical values decrease and become 5,34 for Q(6) 11,34 for Q(12) and 23,33 for Q(24). In this case, both Box-Pierce Q_{BP} and Ljung Box Q_{LB} statistics reject the null hypothesis that all autocorrelations are zero for the cases of Volkswagen, Peugeot, Skoda while the null hypothesis is also rejected for Nissan and Citroen when the Q statistic considers 24 lags.

The levels of p-values are also illustrated graphically in the Ljung–Box p-values plots in Figure, in the third plot (c) on the row. High p-values indicate no autocorrelation in the residuals as shown in the case of Opel, Toyota, Hyundai, Ford, and Fiat. On the other hand, low-level p-values or p-values with fluctuations as the level of lags increases indicate that there is some autocorrelation left in the residuals like in the case of Volkswagen, Peugeot, Skoda, Nissan, and Citroen.

In general the p-value' exceedance of 0,05 indicates the acceptance of the null hypothesis of the model adequacy at significance level 0,05 [Wang et al., 2005] Conclusively, we can say that there is evidence that the selected SARIMA model is adequate for at least half of the 10 different new car representatives in the sample while the other half firms' SARIMA models seem to have autocorrelation in the residuals of the selected SARIMA models. The SARIMA models may not be adequate for all the sample series.

BP and LB Test of squared SARIMA Residuals. However the application of the Box Pierce and Ljung Box test to the *squared* standardized residuals of the SARIMA models illustrated in Table 3.17 page 108 show evidence of heteroskedasticity in all residual series. The Q-statistics have higher values and the associated p-values are less than the results given in the case of standardized residuals. They are now significant up to lag 24 at the 5% level in all squared SARIMA residuals indicating that there is heteroskedasticity in the

SARIMA squared residuals. That is suggesting that there can be some improvement on the current models through volatility modeling.

We also notice that the firms that are suffering more from the serial correlation are the Volkswagen and the Nissan, since their selected SARIMA model's squared residual give a quite high value in the ARCH effect test statistics results in Table 3.17. Hence it is necessary for this in sample testing to develop a better model for analysis of the car sales series which is the SARIMA - GARCH model that can handle heteroscedasticity in the series.

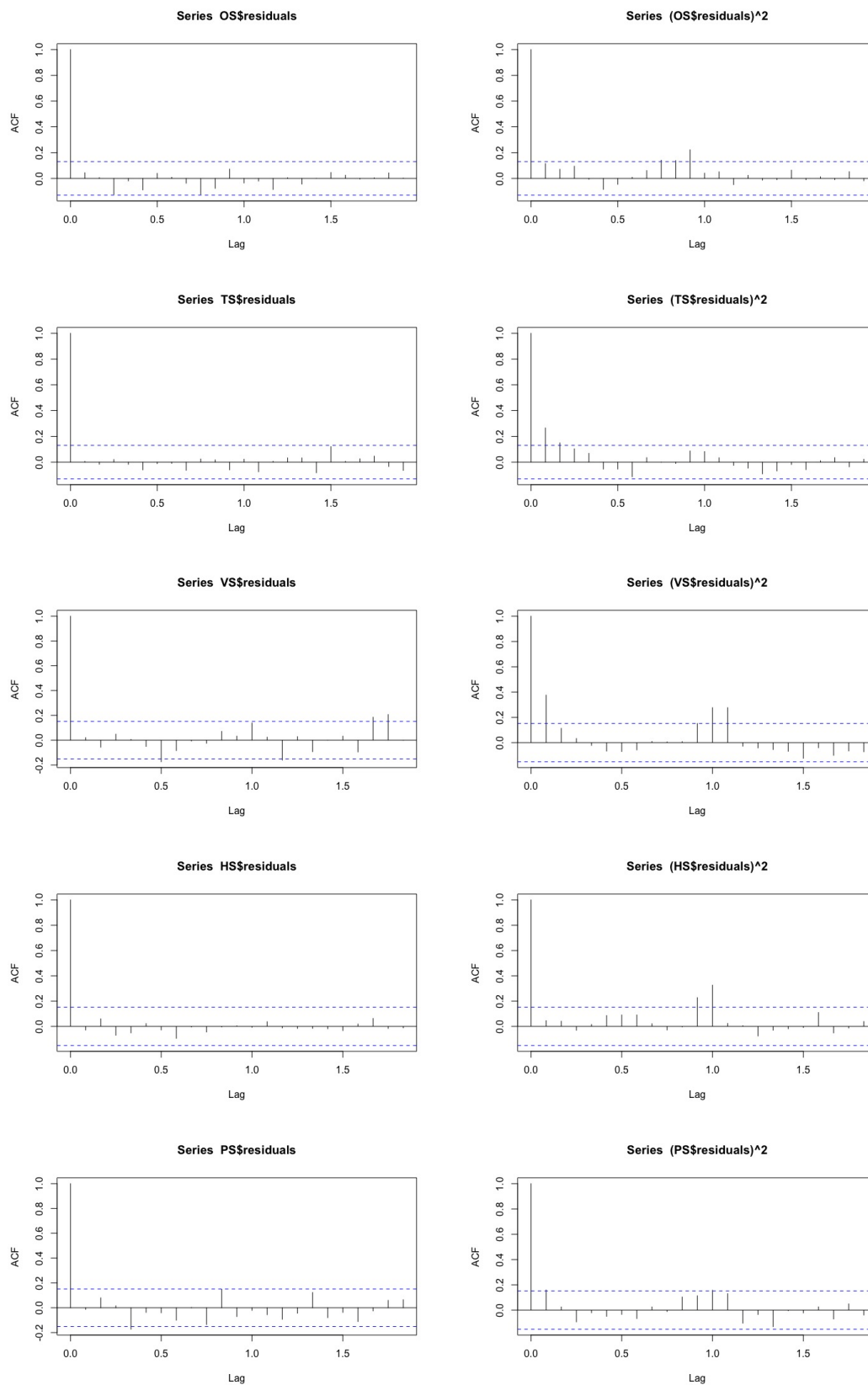


Figure 3.5: ACF of SARIMA resid.and $resid.^2$ -Opel,Toyota,VW,Hyundai,Peugeot

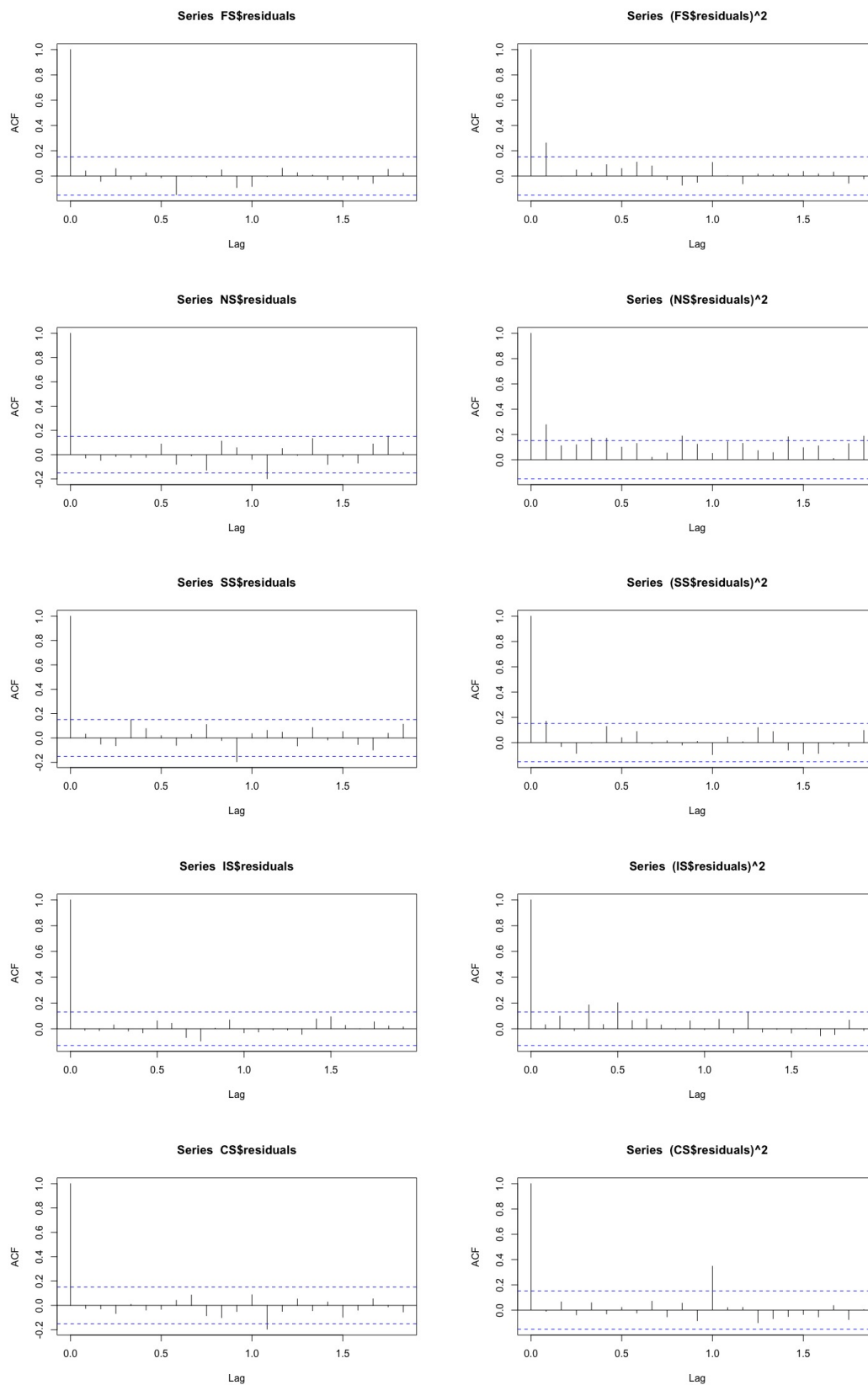


Figure 3.6: ACF of SARIMA resid.and $resid.^2$ -Ford,Nissan,Skoda,Fiat,Citroen

SARIMA residuals		Box–Pierce	Ljung–Box
Opel	6 lags	2,78 (0,83)	2,89(0,82)
	12 lags	6,72 (0,87)	7,13(0,84)
	24 lags	11,90 (0,98)	13,06(0,96)
Toyota	6 lags	1,48 (0,96)	1,54(0,95)
	12 lags	7,46 (0,82)	7,96(0,78)
	24 lags	19,51 (0,72)	21,55(0,60)
Volkswagen	6 lags	6,54(0,36)	6,83(0,33)
	12 lags	12,13(0,43)	12,87(0,37)
	24 lags	33,57 (0,09)	37,32(0,04)
Hyundai	6 lags	2,23(0,89)	2,30(0,88)
	12 lags	4,10 (0,98)	4,28(0,97)
	24 lags	6,98 (0,99)	7,60(0,99)
Peugeot	6 lags	6,68 (0,35)	6,92(0,32)
	12 lags	16,08 (0,18)	16,98(0,15)
	24 lags	26,66 (0,32)	28,95(0,22)
Ford	6 lags	1,39(0,96)	1,43(0,96)
	12 lags	8,00(0,78)	8,49(0,74)
	24 lags	11,10(0,98)	12,04(0,97)
Nissan	6 lags	2,02((0,91)	2,10(0,90)
	12 lags	8,76 (0,72)	9,32(0,67)
	24 lags	26,24(0,34)	28,95(0,22)
Skoda	6 lags	5,13(0,52)	5,32(0,50)
	12 lags	19,10(0,08)	20,37(0,60)
	24 lags	33,95(0,08)	37,52(0,03)
Fiat	6 lags	1,46(0,96)	1,50(0,95)
	12 lags	5,19 (0,95)	5,53(0,93)
	24 lags	10,45(0,99)	11,58(0,98)
Citroen	6 lags	1,44(0,96)	1,49(0,96)
	12 lags	7,59(0,81)	8,10(0,77)
	24 lags	23,78(0,47)	26,46(0,33)

Table 3.16: ARCH effect test for SARIMA residuals.

SARIMA residuals		Box–Pierce	Ljung–Box
Opel	6 lags	9,68 (0,13)	9,92(0,12)
	12 lags	24,57 (0,01)	26,00(0,01)
	24 lags	32,24 (0,12)	34,94(0,06)
Toyota	6 lags	13,74 (0,03)	14,21(0,02)
	12 lags	17,99 (0,11)	18,77(0,09)
	24 lags	22,72 (0,53)	24,24(0,44)
Volkswagen	6 lags	27,75(0,00)	28,31(8,178e-05)
	12 lags	45,04(1,013e-05)	47,10(4,464e-06)
	24 lags	67,66 (4,911e-06)	72,41(9,378e-07)
Hyundai	6 lags	3,45(0,75)	3,59(0,73)
	12 lags	31,58 (0,00)	34,13(0,00)
	24 lags	34,64 (0,04)	39,70(0,02)
Peugeot	6 lags	6,49 (0,37)	6,64(0,35)
	12 lags	15,49 (0,21)	16,38(0,17)
	24 lags	30,78 (0,16)	33,76(0,08)
Ford	6 lags	13,89(0,03)	14,20(0,02)
	12 lags	20,40(0,05)	21,18(0,04)
	24 lags	22,22(0,56)	23,23(0,50)
Nissan	6 lags	28,78((6,694e-05)	29,60(4,671e-05)
	12 lags	40,98(4,934e-05)	42,69(2,541e-05)
	24 lags	67,40(5,361e-06)	72,67(8,535e-07)
Skoda	6 lags	6,827(0,33)	6,98(0,32)
	12 lags	9,35(0,67)	9,70(0,64)
	24 lags	30,75(0,16)	34,07(0,08)
Fiat	6 lags	12,74(0,04)	13,21(0,03)
	12 lags	18,11 (0,11)	18,98(0,08)
	24 lags	25,93(0,35)	27,86(0,26)
Citroen	6 lags	1,77(0,93)	1,83(0,93)
	12 lags	25,09(0,01)	27,19(0,00)
	24 lags	35,63(0,05)	39,37(0,02)

Shapiro Wilk Test. Running the *Shapiro–Wilk test* (SW) [Shapiro and Wilk, 1965, Royston, 1982] we examine whether the population being sampled has a specified distribution and more specifically in our case we examine if the sample (residuals of each firm’s SARIMA model) came from a normally distributed population. The Shapiro–Wilk test statistic which can indicate if the residuals are normally distributed is denoted as:

$$SW = \frac{(\sum_{i=1}^n \alpha_i x_{(i)})^2}{\sum_{i=1}^n (x_i - \bar{x})^2} \quad (3.15)$$

where

- x_i is the sample series,
- $x_{(i)}$ (with parentheses enclosing the subscript index i) is the i^{th} smallest number in the sample,
- $\bar{x} = \frac{(x_1 + x_2 + \cdots + x_n)}{n}$ is the series mean,
- the constants α_i are given by $\alpha_1, \dots, \alpha_n = \frac{m^T V^{-1}}{(m^T V^{-1} V^{-1} m)^2}$, where $m = (m_1, \dots, m_n)^T$ and $m_1, \dots, m_n$ are the expected values of the order statistics of independent and identically – distributed random variables sampled from the standard normal distribution, and
- V is the covariance matrix of those order statistics.

SW test Decision Rule: If the p-value of the SW test is greater than the chosen alpha level then we do not reject the null hypothesis which means that we conclude that the data are from a normally distributed population. If the p-value of the SW test is less than the chosen alpha level, then one does reject the null hypothesis that the data came from a normally distributed population.

The results from the Shapiro–Wilk test applied in the selected SARIMA residuals shown in Table 3.18 (page 111) yields that for an alpha level of 0.10 (i.e for a 90% confidence interval) we cannot reject the hypothesis that the data are from a normally distributed population i.e. the residuals are normal for Hyundai, Nissan, Fiat and Skoda while it

indicates that the residuals are not normally distributed for Opel, Toyota, Volkswagen, Peugeot, Ford and Citroen. Thus, the (S)ARIMA models for Opel, Toyota, Volkswagen, Peugeot, Ford and Citroen appear to fit well, but for Hyundai, Nissan, Fiat and Skoda a distribution with heavier tails than the normal distribution should be employed.

Jarque - Bera Test. There is also no evidence of no normality using the *Jarque-Bera (JB) test* [Jarque and Bera, 1987] for Hyundai, Fiat, Skoda and Nissan. The JB test has the null hypothesis that the SARIMA models residuals are drawn from a normally distributed population and hence the test can be regarded as a test of goodness of fit. The results from this test illustrated in Table 3.18 (page 111) confirms that the SARIMA residuals have skewness and kurtosis matching normal distribution only for Hyundai, Nissan, Skoda, and Fiat. However the majority of the car representatives have JB statistic p-values sufficiently low which is happening because the value of the statistic is very different from zero (0) and so we reject the hypothesis that the residuals are normally distributed for Opel, Toyota, Volkswagen, Peugeot, Ford, and Citroen. Of course, we keep in mind that the JB test is a large sample test and our sample may not be necessarily large.

Table 3.18: Shapiro Wilk & Jarque Bera test for selected SARIMA residuals.

	Shapiro–Wilk	p–value	Jarque Bera	p–value
Opel	0,97	0,01	8,52	0,01
Toyota	0,96	0,00	34,76	2,827e-08
VW	0,97	0,00	15,00	0,00
Hyundai	0,99	0,96	0,24	0,88
Peugeot	0,97	0,00	9,71	0,00
Ford	0,98	0,02	9,48	0,00
Nissan	0,99	0,294	1,81	0,40
Fiat	0,98	0,19	0,31	0,85
Skoda	0,99	0,76	0,94	0,62
Citroen	0,96	0,00	28,41	6,766e-07

3.6 SARIMA - GARCH Modelling

Since GARCH family models have been introduced to the world, people often use them to analyze volatility and they usually fit well. Nowadays, with the development of economics all around the world and the economic recession of Europe and Greece more and more people are considering less investing their money in durable products like cars. This situation gives rise to volatility in car sales over time. For a definition, volatility is a measure for variation of the price of a financial instrument over time, according to Lin C. 1996. Therefore the future new car sales levels uncertainty could be presented as volatility. Modeling and forecasting volatility need new models other than the traditional SARIMA models. The well Known SARIMA models have their limitations in the application since they always ignore the heteroskedasticity of the monthly new car sales data.

The Seasonal Autoregressive Integrated Moving Average with Generalized Autoregressive Conditional Heteroscedastic Errors (SARIMA-GARCH) models are called volatility models. These models were developed after relaxing the hypothesis of constant variance in the series - which was needed for SARIMA model estimation. In this way, researchers created other types of models like a SARIMA with generalized autoregressive conditional heteroscedastic errors (SARIMA-GARCH) models.

The SARIMA models are good for modeling homoscedastic time series, meaning time series with constant variances. However most of the financial and sales time series are leptokurtic with fat tails and if we wish to relax the hypothesis of homoscedasticity and assume that the variance is not constant, we model the series using heteroscedastic models such as the Auto-Regressive Conditional Heteroscedastic (ARCH) models [Engle, 1982] or the Generalized ARCH (GARCH) Models [Bollerslev, 1986] and their derivatives which are the IGARCH, GARCH-M, GARCH-t, asymmetric GARCH models and so on (for more details Brooks [2008], Enders [1995], Mills [1999], Tsay [2010])

In this chapter, we are going to introduce the SARIMA GARCH models and then we are going to follow the Box Jenkins methodology to model our series. In the first step, the Identification process will be done by checking if there is an ARCH effect in the SARIMA

residuals. Step 2 is going to be the estimation of a variety of volatility models and Step 3 will end the process with a diagnostic checking.

3.6.1 SARIMA GARCH Model Building

The SARIMA-GARCH model may be interpreted as a combination of a SARIMA model which is used to model mean behavior and an ARCH model which is used to model the ARCH effect in the residuals series from the SARIMA model. The number of GARCH models is immense, but the most influential models were the first. The standard ARCH model was introduced by Engle [1982] and the GARCH model by Bollerslev [1986]. There exist a collection of review articles by Bollerslev et al. [1992], Higgins and Bera [1993], Bollerslev et al. [1994], Engle [2001], Engle and Patton [2001], and Ling and McAleer [2002] that give a good overview of the scope of the ARCH and GARCH models. In applied research, these models are especially useful when the goal of the study is the analysis and the forecast-volatility. These volatility models are able to forecast volatility and incorporate in a model some facts like persistence, mean reversion, asymmetry, and the possibility of exogenous or predetermined variables influencing volatility.

The SARIMA-GARCH model building procedure proceeds in the following steps:

- Logarithmize the original new-car sales series for the 10 different car representatives.
- Fit a SARIMA model to the logarithmized new-car sales series
- Calculate seasonal standard deviations of the residuals obtained from SARIMA model and standardize the residuals.
- Fit a GARCH model to the standardized residual series.

The SARIMA model has the general form as specified in equations 2.11, page 59 and the GARCH (p,q) model has the form [Bollerslev, 1986] :

$$\begin{aligned} \varepsilon_t | \psi_{t-1} &\sim N(0, h_t) \\ h_t &= \omega + \sum_{i=1}^q \alpha_i \varepsilon_{t-i}^2 + \sum_{i=1}^p \beta_i h_{t-i} \end{aligned} \tag{3.16}$$

where ε_t denotes a real-valued discrete-time stochastic process, ψ_t the available information set ω is the constant term, while α_i and β_i are the coefficients with the restrictions $p \geq 0, q > 0, \omega > 0, \alpha_i \geq 0$, and $\beta_i \geq 0$.

When $p=0$ the GARCH(p,q) model reduces its form to the ARCH(q) model. Under the GARCH (p,q) model the conditional variance(h_t) of ε_t , depends on the squared residuals in the previous q time steps and the conditional variance in the previous p time steps. Since GARCH models can be treated as SARIMA models for squared residuals the order of GARCH can be determined with the method for selecting the order of SARIMA models and traditional model selection criteria such as Akaike information criterion (AIC) or Bayesian information criterion (BIC) for the selected models [Wang et al., 2005].

3.6.2 Step 1. Identification

Before analyzing and building SARIMA-(G)ARCH models, which are a combination of SARIMA modeling for the mean equation and GARCH modeling for the volatility equation, the researcher check if there is an ARCH effect in the residuals of the selected SARIMA models.

Test of Arch effect. There are some formal methods to test for the ARCH effect of a process such as the McLeod-Li test [McLeod and Li, 1983], the Engle's Lagrange Multiplier test [Engle, 1982], the BDS test [A. et al., 1996], Tsay [1986] e.t.c.

All these tests share the principle that once any linear structure is removed from the data, any remaining structure should be due to a non-linear data generating mechanism. The linear structure is removed from the data through the SARIMA models and the residuals of the preferred SARIMA models, which are by construction serially uncorrelated, are then tested for non-linear independence using each of the procedures in turn. All the procedures embody the null hypothesis that the series under consideration is an independently and identically distributed (IID) process. McLeod-Li and Engle's Lagrange Multiplier test are used in this study to check the existence of an ARCH effect in the new car registration series. Details of the tests are as follows.

McLeod-Li test for the ARCH effect A formal test for ARCH effect based on the Ljung-Box test was proposed by McLeod and Li [1983]. The test looks at the autocorrelation function of the squares of the residuals and tests whether the first L autocorrelations for the squared residuals are collectively small in magnitude. Similar to the Ljung-Box test equation

$$Q_{ML} = n(n+2) \sum_{k=1}^L \frac{\hat{\rho}_k^2(\varepsilon^2)}{n-k} \quad (3.17)$$

where n is the sample size, L =the number of autocorrelations included in the statistic and

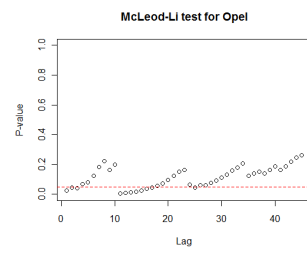
$$\hat{\rho}_k^2 = \frac{\sum_{t=k+1}^n \varepsilon_t^2 \varepsilon_{t-k}^2}{\sum_{t=1}^n \varepsilon_t^2}$$

$\hat{\rho}_k^2$ is the squared sample autocorrelation of squared residuals series at lag k , obtained from fitting the selected SARIMA model to the data. If the series ε_t is independently and identically distributed (IID) then the asymptotic distribution of Q_{ML} is $\chi^2(L)$ distributed with L degrees of freedom. Thus, under the null hypothesis $-H_0$:No ARCH effect in the data- the test statistic is asymptotically $\chi^2(L)$ distributed. Figure 3.7, in page 116 shows the results of the McLeod-Li test for monthly new car sales SARIMA residuals series. It illustrates that the null hypothesis of no ARCH effect is clearly rejected only for the case of Volkswagen, Hyundai, and Nissan SARIMA residual series.

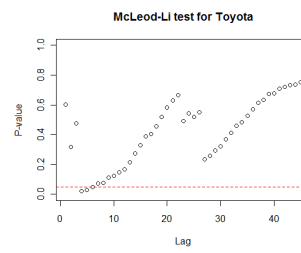
Engle's Lagrange Multiplier test for the ARCH effect Given the fact that the ARCH model has the form of an autoregressive model, Engle [1982] proposed the Lagrange Multiplier (LM) test, in order to test for the existence of ARCH behavior based on the regression. The Lagrange multiplier Test (or LM Test) for the ARCH effect is applied to examine if the residuals are higher-order series correlated. ARCH LM tests whether coefficients in the following regression are zero.

$$\varepsilon_t^2 = \omega + \alpha_1 \varepsilon_{t-1}^2 + \cdots + \alpha_q \varepsilon_{t-q}^2$$

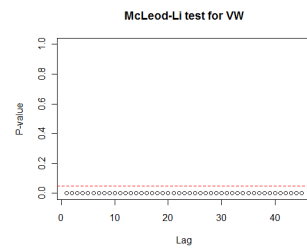
The ε_t are the residuals from the selected SARIMA models, which we want to test for ARCH effects. The null hypothesis, which states that there is no ARCH effects in the



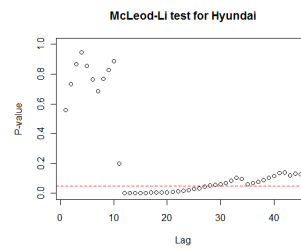
(a) Opel



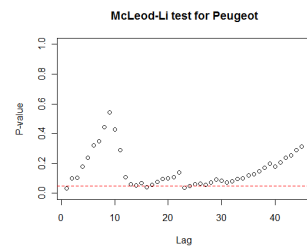
(b) Toyota



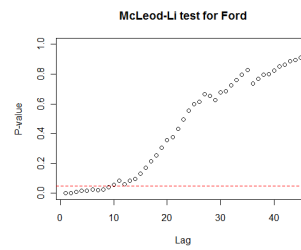
(c) VW



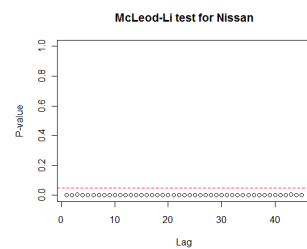
(d) Hyundai



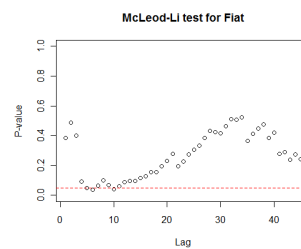
(e) Peugeot



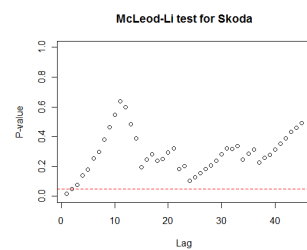
(f) Ford



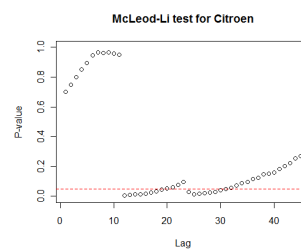
(g) Nissan



(h) Fiat



(i) Skoda



(j) Citroen

residuals is :

$$H_0 : \alpha_1 = \alpha_2 = \dots = \alpha_q = 0$$

If the hypothesis is accepted then we can say that series has no ARCH effects because the volatility becomes equal to $\varepsilon_t^2 = \omega$ meaning that we have constant volatility or in other words homoscedasticity. If the null hypothesis is rejected then one or more coefficients are non zero and we say that there is an ARCH effect in the SARIMA residuals, heteroskedasticity. In more detail the testing procedure for ARCH effect has the following steps:

1. Run the selected SARIMA model and obtain the residual ($\hat{\varepsilon}_t$) and the squared residuals ($\hat{\varepsilon}_t^2$).
2. Regress the squared residuals $\hat{\varepsilon}_t^2$ on past values $\hat{\varepsilon}_{t-1}^2, \hat{\varepsilon}_{t-2}^2, \dots, \hat{\varepsilon}_{t-q}^2$ where

$$\hat{\varepsilon}_t^2 = c + c_1\hat{\varepsilon}_{t-1}^2 + c_2\hat{\varepsilon}_{t-2}^2 + \dots + c_q\hat{\varepsilon}_{t-q}^2 + \text{error}.$$
3. Calculate the test statistic given by TR^2 , where R is the sample multiple correlation coefficient computed from the regression of ε_t^2 on a constant and $\varepsilon_{t-1}^2, \dots, \varepsilon_{t-q}^2$, while T is the sample size. The test statistic $LM_{test} = TR^2 \sim x_q^2$ is asymptotically distributed as chi-square distribution with q degrees of freedom.
4. Decision Rule: Reject the null of no ARCH effects if $LM_{test} > x_q^2$ critical values.

This test, as Bollerslev [1986] suggested, should have power against the GARCH model alternatives.

Using the squared residuals of the identified SARIMA model estimation result, we test for the ARCH effect and use 1-2-3-4-5-6-7-8-9-10 lag for this test. Engle's LM test was carried out and the null hypothesis that there are no (G)ARCH effects in the residuals is rejected at the 10% level. The results as presented in details in Table 3.19, page 119 show that the Engle LM test statistics for the absence of ARCH effect in the squared residuals of the selected SARIMA models, is strongly rejected for the case of Opel, Peugeot, Ford, and Fiat at a 10% level of significance at lag one. So we should consider the use of heteroscedastic models,

especially for these cases. Additionally, the test shows weaker evidence of ARCH effect for Hyundai and Nissan while on the other hand the rest of the firms (Toyota, Volkswagen, Skoda, and Citroen) show no ARCH effect in the SARIMA residuals. This suggests that there can be some improvement in the current model though volatility modeling especially for the case of Opel, Peugeot, Ford, and Fiat. Therefore conditional volatility models will be built to investigate the ARCH effects in new car sales series. We do not test directly for GARCH effects because if ARCH effects exist, GARCH models can be considered.

On the whole, shreds of evidence from both the McLeod-Li test and Engle's LM test about the existence of conditional heteroskedasticity in the residual series, from SARIMA fitted to the logarithmized monthly new vehicles' registration series in Greece suggest that conditional volatility models should be built. The McLeod-Li and the Engle test are both tests for GARCH effects. Although the picture is not clear in all new car firms representatives in the sample, one could argue that GARCH effects are present in the majority of the new car sales residuals series.

The results of the Engle LM test show an ARCH effect in the SARIMA residuals of Opel, Peugeot, Ford, Fiat and the results from the McLeod-Li test show evidence of an ARCH effect in Volkswagen Nissan and Hyundai. These evidences are reinforced by the plots of the squared standardized residuals Figure 3.4 and Figure 3.3, in page 101 and page 100 respectively which shows that the volatility of the SARIMA models is not constant.

This empirical study fits the GARCH (p, q) models to the SARIMA residual series for $p = 0, 1, 2$ and $q = 1, 2, 3$ with a mean equation the chosen SARIMA model for each car representative and compare the models AIC values. In table 3.20 on page 120 the Akaike information criterion (AIC) is illustrated for nine (9) different SARIMA-GARCH models. In general, the SARIMA-GARCH (1,1) model is the most popular model for our sample series since half of our sample series seems to have the lowest AIC value at this model.

The other half of the series seems to fit better with SARIMA-ARCH(1) or SARIMA-ARCH(2) models More specific Opel and Skoda seems to fit better using a SARIMA-ARCH(2) model while Peugeot Ford and Citroen prefer the SARIMA-ARCH(1) model while the all the rest car firms (Toyota, Volkswagen, Hyundai, Nissan, Fiat) better fit their

Table 3.19: Engle Test for absence of ARCH effect in squared residuals of SARIMA models.

		1	2	3	4	5	6	7	8	9	10
O	LM_m	3.00*	3.99	4.15	5.78	7.56	7.56	8.44	10.72	12.33	16.83*
	p-value	(0.08)	(0.13)	(0.24)	(0.21)	(0.18)	(0.27)	(0.29)	(0.21)	(0.19)	(0.07)
T	LM_m	0.31	1.20	2.23	3.64	3.97	5.69	6.02	7.66	8.21	8.85
	p-value	(0.57)	(0.54)	(0.52)	(0.45)	(0.55)	(0.45)	(0.53)	(0.46)	(0.51)	(0.54)
VW	LM_m	0.01	2.49	3.49	3.44	4.52	5.84	9.21	12.21	12.81	12.51
	p-value	(0.90)	(0.28)	(0.32)	(0.48)	(0.47)	(0.44)	(0.23)	(0.14)	(0.17)	(0.25)
H	LM_m	2.44	2.59	2.78	3.01	6.69	10.01	10.29	10.37	10.78	10.97
	p-value	(0.11)	(0.27)	(0.42)	(0.55)	(0.24)	(0.12)	(0.17)	(0.23)	(0.29)	(0.35)
P	LM_m	3.65*	6.42*	7.87*	8.32*	9.99*	10.00	11.26	13.81*	14.56	16.58*
	p-value	(0.05)	(0.04)	(0.04)	(0.08)	(0.07)	(0.12)	(0.12)	(0.08)	(0.10)	(0.08)
Fo	LM_m	2.89*	2.97	2.97	3.46	4.07	7.09	9.11	9.03	8.99	9.30
	p-value	(0.08)	(0.22)	(0.39)	(0.48)	(0.53)	(0.31)	(0.24)	(0.33)	(0.43)	(0.50)
N	LM_m	1.09	1.35	1.22	3.61	3.60	3.83	4.70	5.99	6.18	6.74
	p-value	(0.29)	(0.50)	(0.74)	(0.46)	(0.60)	(0.69)	(0.69)	(0.64)	(0.72)	(0.74)
Fi	LM_m	2.90*	3.04	2.87	3.23	5.30	8.66	9.03	10.28	11.65	11.76
	p-value	(0.08)	(0.21)	(0.41)	(0.51)	(0.38)	(0.19)	(0.25)	(0.24)	(0.23)	(0.30)
S	LM_m	0.24	0.60	0.80	1.78	5.00	5.43	8.40	9.03	9.83	10.05
	p-value	(0.61)	(0.73)	(0.84)	(0.77)	(0.41)	(0.48)	(0.29)	(0.33)	(0.36)	(0.43)
C	LM_m	0.20	0.21	1.92	3.39	3.30	3.49	3.61	4.67	4.59	4.75
	p-value	(0.65)	(0.89)	(0.58)	(0.49)	(0.65)	(0.74)	(0.82)	(0.79)	(0.86)	(0.90)
cv	$x^2_{m,\alpha=0.05}$	3.84	5.99	7.81	9.48	11.07	12.59	14.06	15.50	16.91	18.30
cv	$x^2_{m,\alpha=0.10}$	2.70	4.60	6.25	7.77	9.23	10.64	12.01	13.36	14.68	15.98

Note: Decision Rule $LM > x^2_m$ Reject H_0 :No ARCH effect, (*)sign for ARCH effect.

new car sales series using a SARIMA-GARCH(1,1) model.

Table 3.20: AIC information criterion for SARIMA-GARCH models.

	(0,1)	(0,2)	(0,3)	(1,1)	(1,2)	(1,3)	(2,1)	(2,2)	(2,3)
Opel	-15,81	-17,53	-13,28	-16,93	-15,02	-10,26	-13,65	-12,88	-8,62
Toyota	24,21	17,22	24,47	15,87	19,51	24,90	19,60	18,65	26,05
VW	-15,63	-14,73	-11,91	-16,31	-12,93	-8,64	-11,88	-9,87	-4,88
Hyundai	13,76	15,04	17,66	13,59	17,10	19,64	17,84	19,07	21,64
Peugeot	25,85	28,88	29,10	27,86	31,05	31,27	31,36	33,25	33,61
Ford	48,12	51,06	51,96	50,19	53,16	54,05	51,37	54,93	55,67
Nissan	33,98	33,11	34,11	29,01	32,78	35,96	31,85	34,17	37,62
Fiat	-5,76	-3,65	-1,49	-7,59	-5,16	-2,32	-4,60	-3,55	-0,46
Skoda	40,24	39,89	43,47	41,95	42,27	45,50	42,25	44,07	47,59
Citroen	25,14	26,18	29,27	27,14	28,28	31,62	30,00	30,10	33,61

3.6.3 Step 2. Estimation

The unknown parameters $\alpha_i (i = 1, 2, \dots, q)$ and $\beta_j (j = 1, 2, \dots, p)$ of the Equation 3.16 can be estimated using the (conditional) maximum likelihood estimation (MLE). Estimates of the conditional standard deviation $\sqrt{h_t}$ are obtained as a side product standard deviation. The results from the estimation process are the following:

- OPEL Model: $SARIMA(0, 1, 1)(1, 0, 1)_{12} - ARCH(2)$

$$(1 - B)(1 - 0.97B^{12})x_t = (1 - 0.54B)(1 - 0.71B^{12})\varepsilon_t$$

$$h_t = 0,03 + 0,33\varepsilon_{t-1}^2 + 0,17\varepsilon_{t-2}^2$$

- TOYOTA Model: $SARIMA(2, 1, 3)(1, 1, 1)_{12} - GARCH(1, 1)$

$$(1 - 1,64B - 0,86B^2)(1 - B)(1 - B^{12})(1 - 0,12B^{12})x_t = (1 + 1,13B - 0,23B^2 - 0,67B^3)(1 - 0,75B^{12})\varepsilon_t$$

$$h_t = 0,01 + 0,34\varepsilon_{t-1}^2 + 0,54h_{t-1}$$

- VOLKSWAGEN Model: $SARIMA(0, 1, 1)(2, 1, 0)_{12} - GARCH(1, 1)$
 $(1 - B)(1 - B^{12})(1 + 0,84B^{12} + 0,31B^{24})x_t = (1 - 0,62B)\varepsilon_t$
 $h_t = 0,02 + 0,39\varepsilon_{t-1}^2 + 0,29h_{t-1}$
- HYUNDAI Model: $SARIMA(1, 1, 1)(2, 0, 1)_{12} - GARCH(1, 1)$
 $(1 - 0,21B)(1 - B)(1 - 0,64B^{12} - 0,29B^{24})x_t = (1 - 0,74B)(1 - 0,47B^{12})\varepsilon_t$
 $h_t = 0,03\varepsilon_{t-1}^2 + 0,90h_{t-1}$
- PEUGEOT Model: $SARIMA(0, 1, 1)(1, 0, 1)_{12} - ARCH(1)$
 $(1 - B)(1 - 0,95B^{12})x_t = (1 - 0,55B)(1 - 0,63B^{12})\varepsilon_t$
 $h_t = 0,05 + 0,18\varepsilon_{t-1}^2$
- FORD Model: $SARIMA(0, 1, 1)(1, 0, 1)_{12} - ARCH(1)$
 $(1 - B)(1 - 0,92B^{12})x_t = (1 - 0,66B)(1 - 0,68B^{12})\varepsilon_t$
 $h_t = 0,06 + 0,14\varepsilon_{t-1}^2$
- NISSAN Model: $SARIMA(0, 1, 4)(0, 1, 1)_{12} - GARCH(1, 1)$
 $(1 - B)(1 - B^{12})x_t = (1 - 0,49B - 0,09B^2 - 0,007B^3 - 0,22B^4)(1 - 0,63B^{12})\varepsilon_t$
 $h_t = 0,01 + 0,30\varepsilon_{t-1}^2 + 0,53h_{t-1}$
- FIAT Model: $SARIMA(2, 1, 1)(1, 0, 1)_{12} - GARCH(1, 1)$
 $(1 - 0,47B - 0,34B^2)(1 - B)(1 - 0,98B^{12})x_t = (1 - 0,98B)(1 - 0,77B^{12})\varepsilon_t$
 $h_t = 0,08\varepsilon_{t-1}^2 + 0,78h_{t-1}$
- SKODA Model: $SARIMA(0, 1, 1)(2, 0, 0)_{12} - ARCH(2)$
 $(1 - B)(1 - 0,42B^{12} - 0,27B^{24})x_t = (1 - 0,50B)\varepsilon_t$
 $h_t = 0,05 + 0,23\varepsilon_{t-1}^2 + 0,05\varepsilon_{t-2}^2$
- CITROEN Model: $SARIMA(0, 1, 2)(0, 1, 1)_{12} - ARCH(1)$
 $(1 - B)(1 - B^{12})x_t = (1 - 0,54B - 0,19B^2)(1 - 0,93B^{12})\varepsilon_t$
 $h_t = 0,06 + 0,0049\varepsilon_{t-1}^2$

3.6.4 Step 3. Diagnostic Checking.

In order to check the goodness of the fit of the SARIMA -GARCH models we apply the Jarque-Bera (JB) and the Ljung Box test illustrated in Table 3.21, page 122. In Ljung Box test we have high p-values in all firms and that gives evidence of lack of dependence in the residuals of the SARIMA-GARCH models. The goodness of fit Jarque-Bera (JB) test confirms that the SARIMA-GARCH residuals have skewness and kurtosis matching normal distribution only for the cases, where the JB test p-value is sufficiently high, like for example Hyundai, Fiat, Skoda, Nissan, Opel. We reject the hypothesis that the SARIMA-GARCH residuals are normally distributed for the Toyota, Volkswagen, Peugeot, Ford and Citroen.

Table 3.21: Residuals Diagnostic Tests for SARIMA-GARCH model

firm	Jarque Bera	p-value	Box-Ljung	p-value
Opel	2,54	0,27	0,02	0,88
Toyota	6,03	0,04	9e-04	0,97
VW	5,34	0,06	0,001	0,96
Hyundai	0,14	0,93	0,007	0,93
Peugeot	12,89	0,00	0,04	0,84
Ford	6,94	0,03	0,28	0,59
Nissan	2,82	0,24	0,36	0,54
Fiat	0,28	0,86	0,36	0,54
Skoda	1,32	0,51	6e-04	0,98
Citroen	27,44	1,1e-06	0,021	0,88

3.7 Discussion

To sum up, in this chapter we present the empirical evidence of a with-in-sample empirical analysis of the top ten vehicle representatives in the Greek market (see Table 3.22 page 124). The research shows that after fitting more than six time series models to all the data set of each series, the model that fit best is the ETS model. That evidence is the same in all time series, since the model has the smallest value of MAPE performance measure in all series. The reason for that is the fact that, since the economic environment is pretty unstable during these last two decades, and ETS model give high importance to the last observations into the model, these both are crucial and give as a result better forecasts for our study.

However, other models like the SARIMA, and Seasonal Naïve models, that come as a second best choice, also give very good results in comparison to other time series models. The goodness of the fit to the data for these models makes sense, since the researched data do have seasonality effects that are very well explained with the SARIMA and Seasonal Naïve models that do take that fact into highly consideration.

Furthermore, as the accuracy measure, MAPE, can also allow comparison across different time series, we can conclude that if we examine the fit of ETS model across all time series, we get the best with-in-sample fit for the model in the case of Opel, Toyota, Fiat and then VolksWagen, Hyundai, Ford and so on. For the SARIMA model we have similar results, the best with-in-sample fit appears to the VolksWagen, Hyundai, Citroen, Opel, Peugeot, Nissan, Skoda, Ford, Fiat and Toyota. We notice that the firms that fit better to the models are the ones that keep the biggest share in the market place during the years. Therefore evidence show us that the firms that had a biggest share (almost 10% of the total sales) in the Greek new car market (like Opel, Toyota, VolksWagen) keep a more stable level of sales and is easier to better fit a time series model to those data than to firms with smaller market share and a less solid position in the new car market. This study concluded that ETS and SARIMA models are more appropriate models in fitting new car sales series.

This with-in-sample empirical research also examined the residuals of the SARIMA

Table 3.22: Time Series MAPE% metrics for with-in-sample modeling (log values)

	Mean	Naïve	S.Naïve	ETS	LMSD	SARIMA*
Opel	4,76	4,30	3,75	2,12	7,45	2,84
Toyota	4,61	4,68	3,65	2,40	7,16	3,27
VW	4,46	4,08	3,60	2,41	8,11	2,44
Hyundai	6,78	4,68	4,16	2,52	6,28	2,76
Peugeot	6,95	4,37	4,29	2,86	7,56	2,92
Ford	7,75	4,66	4,59	2,76	8,40	3,18
Fiat	5,65	4,67	4,15	2,35	9,41	3,22
Nissan	7,93	4,78	4,65	2,95	8,50	3,13
Skoda	7,03	4,55	4,50	2,90	8,20	3,15
Citroen	6,98	4,40	4,45	2,95	7,60	2,77

*SARIMA selected see Table 3.15 (page 95)

model estimation. There was evidence of heteroskedasticity in the SARIMA squared residuals which suggested that the research can improve the SARIMA model through volatility modeling and so we estimated the SARIMA-GARCH models for all series. While the Engle LM Test showed an ARCH effect in only a small group of firms (Opel, Peugeot, Ford, Fiat) and Box-Pierce and Ljung-Box statistics show that squared residuals of SARIMA models of Volkswagen and Nissan are mostly suffering from serial correlation, this study continues the in-sample-estimation research with different types of SARIMA-GARCH models. We estimated various specification of these models and conclude that, in general, the SARIMA-GARCH(1,1) model is the most popular one, for in-sample-estimations since it has the lowest AIC value at these models. We can assume that the volatility of the series is better explained using these models, but we can not be sure that these models specification can improve the forecasting ability. For the estimation of the forecasting accuracy of the models we continue with an out of sample research in the next chapters.

Time Series out-of-sample Forecasting.

4.1 Introduction.

Forecasting is an attempt to predict the future, as accurate as possible, given all the available information, including historical data and knowledge of past, present, or future events that may be influencing the forecast valuables. It is a common statistical task in economics that provides useful information for the decision-makers of every field in the public or private sector. In issues like production scheduling, personnel management, strategic planning, the use of prediction methods is of great importance. In practice, however, the business or public unit forecasting is often done poorly using simple forecasting methods that are often not based on statistical modeling. In this research the use of statistical forecasting is developed, using some of the most commonly used time series methods of forecasting.

To test which is the appropriate model for forecasting, the researcher can use a within-sample forecast or an out-of-sample forecast. *In-sample* forecasts are those generated for the same set of data that was used to estimate the model's parameters. *Out-of-sample* forecasts are those generated by not using all of the observations in estimating the model parameters but rather hold some back and then use the hold-out-observation sample to compare how close the forecasts values are relative to their actual values. In this study the researcher is using an out-of sample forecast methodology.

The time interval 1998-2016, in which our time series expands, is divided into four (4) smaller different data sets for this research. The various data sets of our empirical study are defined following the 8:2 ratio between the training and test set. In more details the total data set is divided into the “*Training Set*” and the “*Test Set*”. The *Training Set* is used to estimate the model parameters and covers about 80% of the whole sample, while the *Test Set* is used to measure the forecasting accuracy and covers the remaining 20% of the whole data sample. In this empirical *out-of-sample* forecasts analysis the researcher created A,B,C, and D sets which are divided as follows¹:

Table 4.1: Empirical Analysis Training & Test Sets.

SET	Training Set	80%	Test set	20%
	Estimation Period		Forecasting Period	
A	Jan.1998-Dec.2012	180 obsv. (15y)	Jan.2013-Dec.2016	48 obsv.(4y)
B	Jan.2006-Dec.2013	96 obsv. (8y)	Jan.2014-Dec.2015	24 obsv. (2y)
C	Jan.2006-Dec.2009	48 obsv. (4y)	Jan.2010-Dec.2010	12 obsv. (1y)
D	Jan.2002-Dec.2009	96 obsv. (8y)	Jan.2010-Dec.2011	24 obsv. (2y)

The first set (A) of this study covers 19 years. It starts with a Training Set of January 1998 till December 2012 that covers 15 years with 180 monthly data observations and a test set that covers the next four (4) years from January 2013 till December 2016 with 48 observations.

The second set (B) of this study covers 10 years. The Training set starts in January 2006 till December 2013 and covers 8 years with 96 monthly observations while the test set has 24 observations and covers the following two (2) years from January 2014 till December 2015.

The third set (C) of this study starts in January 2006 till December 2009 and covers four (4) years with 48 monthly data that specify each time the model and test its performance

¹Table Note: obsv.= number of observations in each set and y=number of years.

in a set of 12 observations that covers the next year from January till December 2010.

The fourth set (D) of this study starts in January 2002 till December 2009 and covers eight (8) years with 96 monthly data that specify each time the model and test its performance in a set of 24 observations that covers the following two (2) years from January 2010 till December 2011.

Furthermore we consider that the great recession in the Greek market started in early 2008 along with the global financial crisis (GFC) of 2007-2008 and became even worse after the announcement of the financial agreement of the Greek government with the International Monetary Fund (IMF) in the May 2010. Sales reached their historical bottom in the year 2012 and until the end of 2016 the Greek market was under a lot of pressure and a prolong period of economic crisis. The implementation of austerity measures forced Greek citizens to postpone or delay the purchase of cars, banks gave almost no loans for such a purchase and therefore there was a dramatic change in the sales levels over the last years.

In this chapter, after this small introduction, the researcher is going to explain some of the major forecasting accuracy measures used in this empirical study. Furthermore we are going to test these measures empirically in our time series data using an out-of-sample estimation. More specifically, after a definition and a brief reference to the literature for forecasting accuracy measures used, the research continues with an empirical implementation of them in various time series models.

The empirical results of the out-of-sample forecast are given for different time series models and with various forecasting accuracy measures. The calculation process is the following: estimation of each model parameters using the log values of the training set, generate point forecast for the forecast period in each set, and estimate the forecasting measures using the back-transform forecast values of the model with the original values of the test set which are held out for comparison reasons in each set.

Additionally, monthly sales point forecasts are graphically presented for the various time series forecasting methods, in comparison with the actual sales levels with a 95 % confidence interval for the variance of each model, for Opel's new-car sales in Greece. That is an attempt to visually show the results of our empirical study. Lastly we discuss the

finding of this empirical study using different forecasting time series methods and models and enrich the findings with economic analysis of the specific retail sector of the Greek market.

4.2 Forecast Accuracy Measures.

There are different measures of forecast accuracy. These measures are used because the forecasting models may be bias and there may be a need to find the absolute size of the forecast errors which is defined as: *Error = Actual – Forecast* i.e. $\epsilon = x_t - f_t$

In the forecasting process it is important the forecast values of a time series generated from the training set model, not to be too far from the actual future outcomes as indicated by the test set, which is left for comparison reasons, and that is important for the accuracy of the forecasting model. So, the difference between the forecast value and the associated actual value of the variable, at time t , is the forecast error. However, if we simply take the average of forecast errors over time, this will eventually not capture the true magnitude of forecast errors, since large positive errors may simply cancel out large negative errors, giving a misleading impression about the accuracy of forecasts generated. As a result, researchers commonly use accuracy measures for the evaluation of the forecasting performance of models.

Analogous to the estimation of time series models, where the parameters are chosen such that the residuals variance are minimized, forecasting is considered desirable to choose the model which minimizes a chosen error function. The measures however can also be used to compare alternative forecasting models and also to identify the forecasting models that may need adjustments. There is a variety of the forecasting accuracy measures. The Measures²used in this study are illustrated below (see Table 4.2 on page 129):

²Table Notes:n =number of time periods, x_t =actual value in time t, f_t =forecast value in time period t.

Table 4.2: Forecast Accuracy Measures

Definition	Equation of Error Magnitude
Mean Squared Error	$MSE = \frac{1}{n} \sum_{t=1}^n (x_t - f_t)^2$
Mean Absolute Error	$MAE = \frac{\sum_{i=1}^n x_t - f_t }{n}$
Root Mean Squared Error	$RMSE = \sqrt{n^{-1} \sum_{t=1}^n (x_t - f_t)^2}$
Mean Absolute Percentage Error	$MAPE = 100n^{-1} \sum_{t=1}^n x_t - f_t / x_t $

The **mean squared error (MSE)**, also called mean squared deviation (MSD), is a measure that gives the average of the squares errors, which is the average squared difference between the estimated values and the actual value. It is always non-negative, and values closer to zero are better. This is an easily computable quantity for a particular sample and hence is sample-dependent. It is denoted as:

$$MSE = \frac{1}{n} \sum_{t=1}^n (x_t - f_t)^2$$

One of the simplest measures of forecast accuracy is the **Mean Absolute Error (MAE)** also called Mean Absolute Deviation (MAD) indicates the absolute size of the errors and is denoted as:

$$MAE \text{ or } MAD = \frac{\sum_{i=1}^n |x_t - f_t|}{n}$$

The **Root Mean Squared Error (RMSE)** is a frequently used measure for evaluating a models performance for fitting the data or forecasting them. It is calculated as:

$$RMSE = \sqrt{n^{-1} \sum_{t=1}^n (x_t - f_t)^2}$$

It calculates the differences between values predicted by a model and the values observed. It is usually best to report the root mean squared error (RMSE) rather than mean squared error (MSE), because the RMSE is measured in the same units as the data, rather than in squared units, and is representative of the size of a typical error.

The root of the minimum squared errors (RMSE) is the standard deviation of the series at each level of differences. According to Wang and Lim [2005], RMSE is the square root of the average of all squared errors and it ignores any over- or under- estimation, but it does not allow comparison across different time series and different time intervals. Furthermore RMSE gives greater weight to large errors than to smalls because the errors are squared before summed.

A typical measure for forecasting accuracy is the **Mean Absolute Percentage Error (MAPE)** which is :

$$MAPE = 100n^{-1} \sum_{t=1}^n |x_t - f_t|/|x_t|$$

The MAPE evaluation criterion is a great tool for model evaluation and forecasting accuracy because it is not scale-dependent, while MAE and RMSE are both scale dependent, and what is more important is that it allows comparison across different time series and different time intervals.

However MAPE has some limitations. For example if the data contain zero values the MAPE can be infinite or if the data contain very small numbers MAPE can be huge. The car series data of this research do not have zero or very small values so those limitations are of minor importance. Additionally MAPE of monthly car sales level makes some kind of sense, when expressed in generic percentage terms, even to someone who has no idea what constitutes a big or small error.

Another limitation of MAPE measurement is that it puts a heavier penalty on negative than on positive errors i.e. when actual values are less than forecast values ($x_t < f_t$). It was Armstrong [1985] that first argue that MAPE has a bias favoring estimates that are below the actual values and later Armstrong and Collopy [1992] argued that MAPE “puts a heavier penalty on forecasts that exceed the actual than those that are less the actual”. For this asymmetry of MAPE Makridakis [1993] argue that “equal errors above the actual value results in a greater absolute percentage error than those below the actual value”. To avoid asymmetry on MAPE Armstrong [1985].proposed the “adjusted MAPE” while

Makridakis [1993] proposed the symmetric MAPE (sMAPE) which are :

$$MAPE_{adjusted} = 100\text{mean}(2|x_t - f_t|/(x_t + f_t))$$

$$MAPE_{symmetric} = 100\text{mean}(2|x_t - f_t|/|x_t + f_t|)$$

In the process of in sample model estimation and the model forecast evaluation two different criteria are used: the Root Mean Squared Error (RMSE) and the Mean Absolute Percentage Error (MAPE). Both RMSE and MAPE were calculated for evaluating the in sample performance of the estimated models, but also for one step out of sample forecast errors of the models. It is more common to use the RMSE and MAPE accuracy measures not only in selecting the best forecasting model but also in comparing the fit of several models. For more information about the different error functions given in the literature see Poon and Granger [2005].

Three forecasting Measures are used in this chapter; the root mean square error (RMSE) and the minimum absolute percentage error (MAPE) and the Mean Absolute Error (MAE) which usually give similar outcomes. However due to the turbulence time of sales levels they sometimes might come up with conflicting or different results.

4.3 Empirical out-of-sample Forecasting.

The empirical Study of this chapter focused in three (3) different car representatives: Opel, Toyota and Fiat. It test the data forecasting performance in four (4) different period using six (6) different time series models, from the simple ones like the Mean/Average, the Naïve, the Seasonal Naïve model to more complicated ones like the Exponential Smoothing State Space Model (ETS) and the Seasonal Autoregressive Integrative Moving Average Model (SARIMA). The researcher keep as a rule that 80% of the data are for the training set and 20% for the test set in each time series.

For preparing the data it is sufficient to take the log values of all the original observations and then calculate the appropriate model based on the observations of the training set. The researcher use the model to estimate point forecast for the next values (equal to range of

test set observations) and then accuracy measures (RMSE, MAPE, and MAE) are estimated to calculate the accuracy of the forecast of each models. The results gave different model specification for each one of the sets given in Table 4.1 and for each one of the different firms. The use of the appropriate software made it feasible and sufficient to calculate the various models.

Furthermore preparing the appropriate SARIMA models the researcher had to follow the steps given by Box-Jenkins methodology: identification, estimation, diagnostic checking. The focus was to find an appropriate SARIMA models based on the ACF and PACF³. Spikes on PACF give the AR number, spikes on ACF gives the MA number while spikes of PACF and ACF around lags 12,24 show evidence of seasonality and lead to AR_s and MA_s . The model selection⁴ for the best SARIMA model for Opel, Toyota and Fiat with the smallest AIC and BIC in all 4 different data sets as stated in Table 4.3.:

Table 4.3: SARIMA models for Forecasting.

SET	Opel	Toyota	Fiat
A	$(1, 0, 1)(2, 0, 0)_{12}$	$(4, 0, 0)(1, 0, 0)_{12}$	$(1, 0, 1)(2, 0, 0)_{12}$
B	$(1, 0, 1)(2, 0, 0)_{12}$	$(1, 0, 0)(1, 0, 0)_{12}$	$(1, 0, 1)(2, 0, 0)_{12}$
C	$(1, 0, 1)(1, 0, 0)_{12}$	$(0, 0, 0)(1, 0, 0)_{12}$	$(1, 0, 0)(1, 0, 0)_{12}$
D	$(2, 0, 0)(2, 0, 0)_{12}$	$(1, 0, 1)(1, 0, 0)_{12}$	$(1, 0, 1)(2, 0, 0)_{12}$

The estimation of the parameters of SARIMA models are calculated and diagnostic checking of the best models, directly related to whether residuals analysis is preformed well, is done for the fitted models of Opel, Toyota and Fiat and plots of standardized residuals, sample ACF of residuals and normal Q-Q plots are estimated (see Figure 4.1 on

³Identification of an AR model is often best done with PACF and identification of a MA model is often best calculated with ACF.

⁴Model selection=strategy used to select the appropriate model from a variety of possible models. Selection criteria: Akaike Information Criteria (AIC) and Bayesian information criterion (BIC) or Schwarz criterion (SBC). The model with the lowest AIC/BIC is preferred based on the likelihood function.

page 4.1).

All firms ACF and PACF of the SARIMA residuals shows all spikes within the significance limits, the QQ plots of standard residuals, which explore the distributional stage, suggest that the distribution may have a tail thicker than that of a normal distribution and may be somewhat skewed to the right and show some non-normality on the tails of the distribution and a nearly straight line suggesting that the residuals follow approximately normal distribution and that the models seems to fit quite well.

4.3.1 Graphical Presentation.

Since graphics always help to picture best the empirical results of a table full of numbers or evidences the researcher illustrate the forecasting models for all four (4) different cases in our empirical result for Opel retail representative in Greece. The first Set A is illustrated in Figure 4.2 on page 136, Set B is illustrated in Figure 4.3 on page 137, Set C is illustrated in Figure 4.4 on page 138 while Set D is illustrated in Figure 4.5 on page 139. Each figure has eight (8) different graphs. The first two graphs show the original and the log data for each period of the data set. The vertical line in the first two graphs, separates the training set from the test set for the easy of understanding. Then the forecast values of the six (6) different forecasting methods used in this thesis (i.e. the Mean, Naïve, seasonal Naïve, LMSD, ETS and SARIMA forecasting models) alone with the 95% CI and the real data for that period are presented. It is noticeable that actual data are not always into the 95% CI of the forecast estimation. When the actual values are close to the forecast values and in the 95%CI with a squeezed range between the upper and the lower level then that usually is the best forecast model.

The study has very bad results for the Mean model as a forecast technique since it simply takes the mean of the historical data and set that as the next forecast value. The Naïve model, on the other hand, takes the last observation and set it as the next forecast value, while the Seasonal Naïve model takes seasonality into account and takes the last observation of the same month of last year and set it as the forecast value. They are both used with some success especially in the long run, and when the economic period

is turbulent and unstable. ETS models are also preferred, especially from the Fiat data, and takes into account the error, the trend and the seasonality while weight more the last observations given to the model. Linear models with seasonal dummies (LMSD) do not give very good results. However as data show seasonality SARIMA models are sometimes hard to win, like in the case of Toyota at data set A. However there is no single model that can fit all cases and all data so the researcher has to evaluate each case study separately.

This research present graphically in the following pages the original data time series and their log transformation values, the estimated forecast values using different time series models and different data sets and their 95% confidence interval. We graphically present all data after back transforming the results to the original units. So in new car sales example that we study, if we are 95% confident that $\alpha < \log \mu < \beta$ then we are also 95% confident that $e^\alpha < \mu < e^\beta$ so we back transform the forecast mean values of each of the six (6) time series models in all four (4) different data sets and the results are presented in Figures.

4.3.2 Accuracy Measures Evaluation.

In this study we approach the forecasting by finding the “best possible” models using models estimations over the mean, Naïve, seasonal Naïve, ETS, LMSD and SARIMA function outputs in R, and select the best estimated models based on the lowest RMSE or MAPE. There is no such a thing as a good RMSE or MAE, because they are scale-dependent i.e. dependent on the dependent variable and therefore one can not claim a number, as a good RMSE or a good MAE. On the other hand, MAPE is a scale-free measures of fit, but we cannot claim a threshold of being good. However, the smaller the root mean square error (RMSE) or the mean absolute percentage error (MAPE) or the mean absolute error (MAE), the better, although small differences between forecasting measures may not be relevant or even significant.

For each candidate model we test the out-of-sample RMSE, MAPE and MAE of each set and each firm, and pick the one with the best out-of-sample metric. The resulting model is the best, in the sense that it gives us a good in-sample fit, associated with low error measures and white noise residuals and avoids over-fitting by giving us the best out-

of-sample forecast accuracy.

According to the results presented in Table 4.4 on page 141 for Opel, Table 4.5 on page 142 for Toyota, and Table 4.6 on page 143 for Fiat, the forecasting accuracy measures RMSE, MAPE and MAE are calculated for six time series models⁵ in all the data sets.

In Set A, the best forecasting model for Opel seems to be the Naïve model, for Toyota the SARIMA model, while for Fiat the ETS model, selected according to the forecast accuracy measures.

In Set B and C, the best forecasting model for Opel, Toyota and Fiat seems to be the ETS Model. This model seems to be very good when the economic environment face instability especially in the a short-run forecasting.

Finally in Set D, the best forecasting model for Opel seems to be the ETS model, for Toyota the Seasonal Naïve model and the Naïve or SARIMA model for FIAT.

⁵Forecasting Models: SARIMA=Seasonal Autoregressive Integrated Mean Average Model,
LMSD=Linear Model with Seasonal Dummies,
ETS=Exponential State space smoothing Model,
S.Naïve=Seasonal Naïve Model.

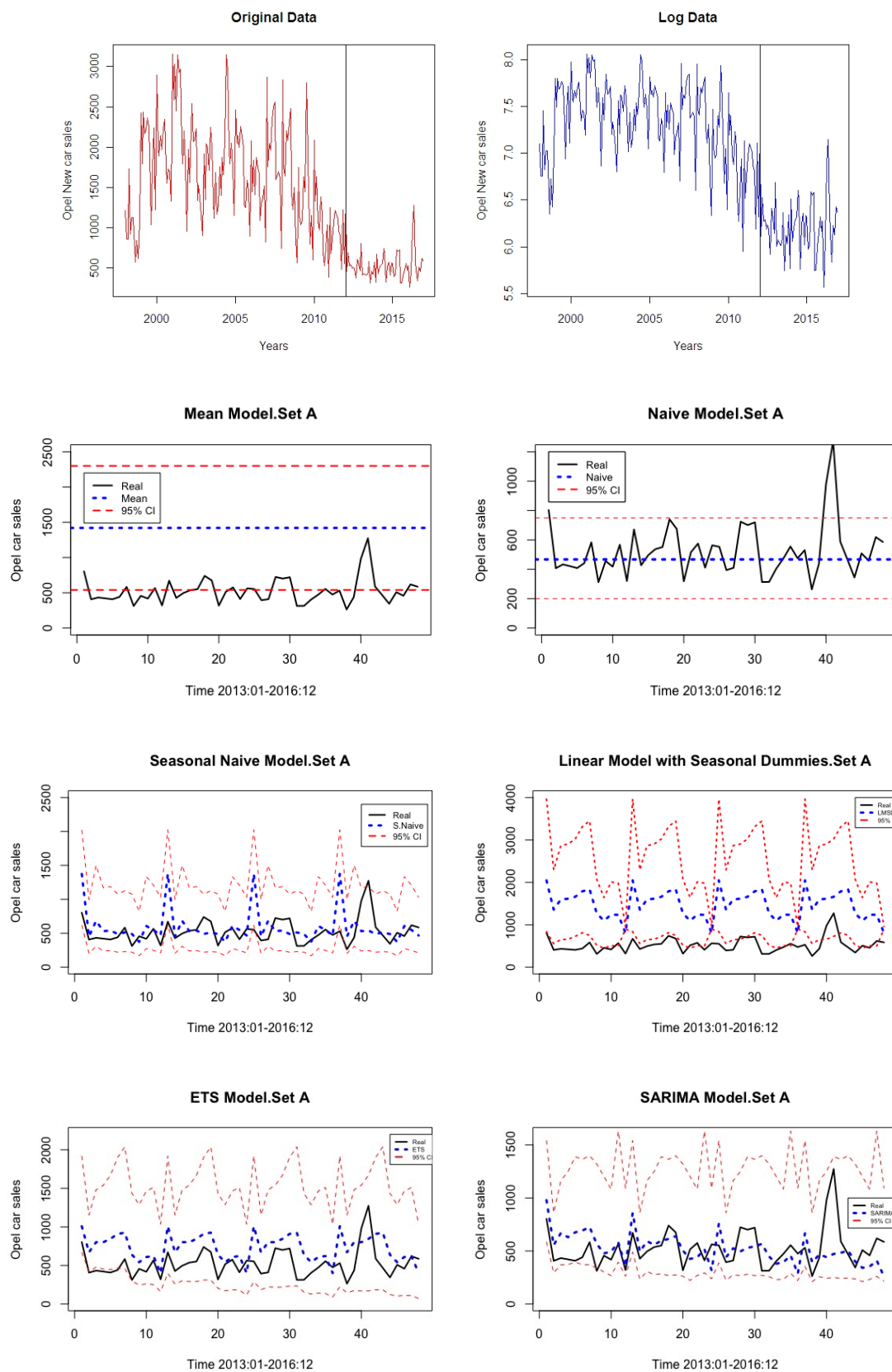


Figure 4.2: Forecasting for Opel. Best Model:Naïve. Set A

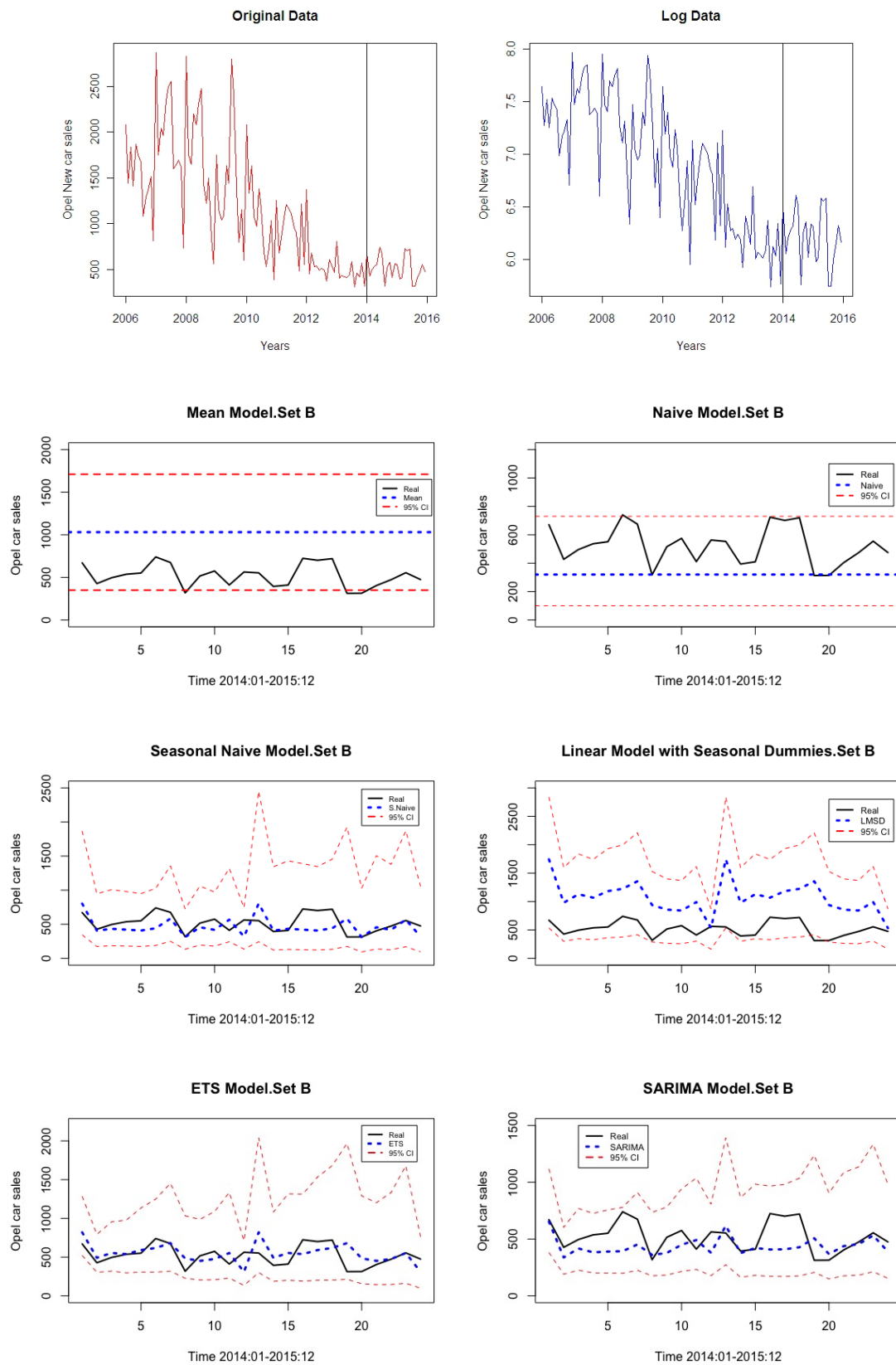


Figure 4.3: Forecasting for Opel. Best Model:S.Naïve & ETS. Set B

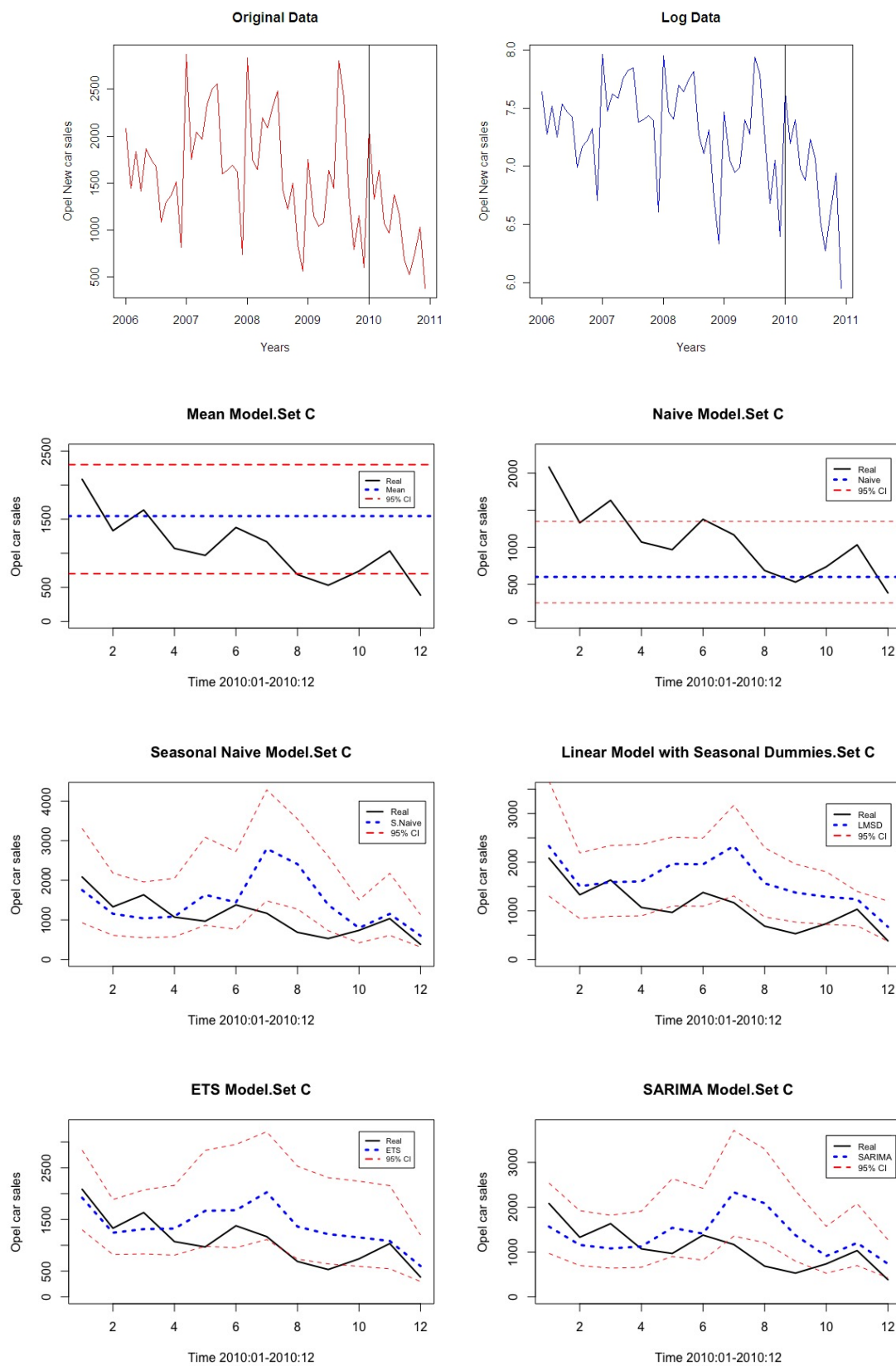


Figure 4.4: Forecasting for Opel. Best Model:ETS. Set C

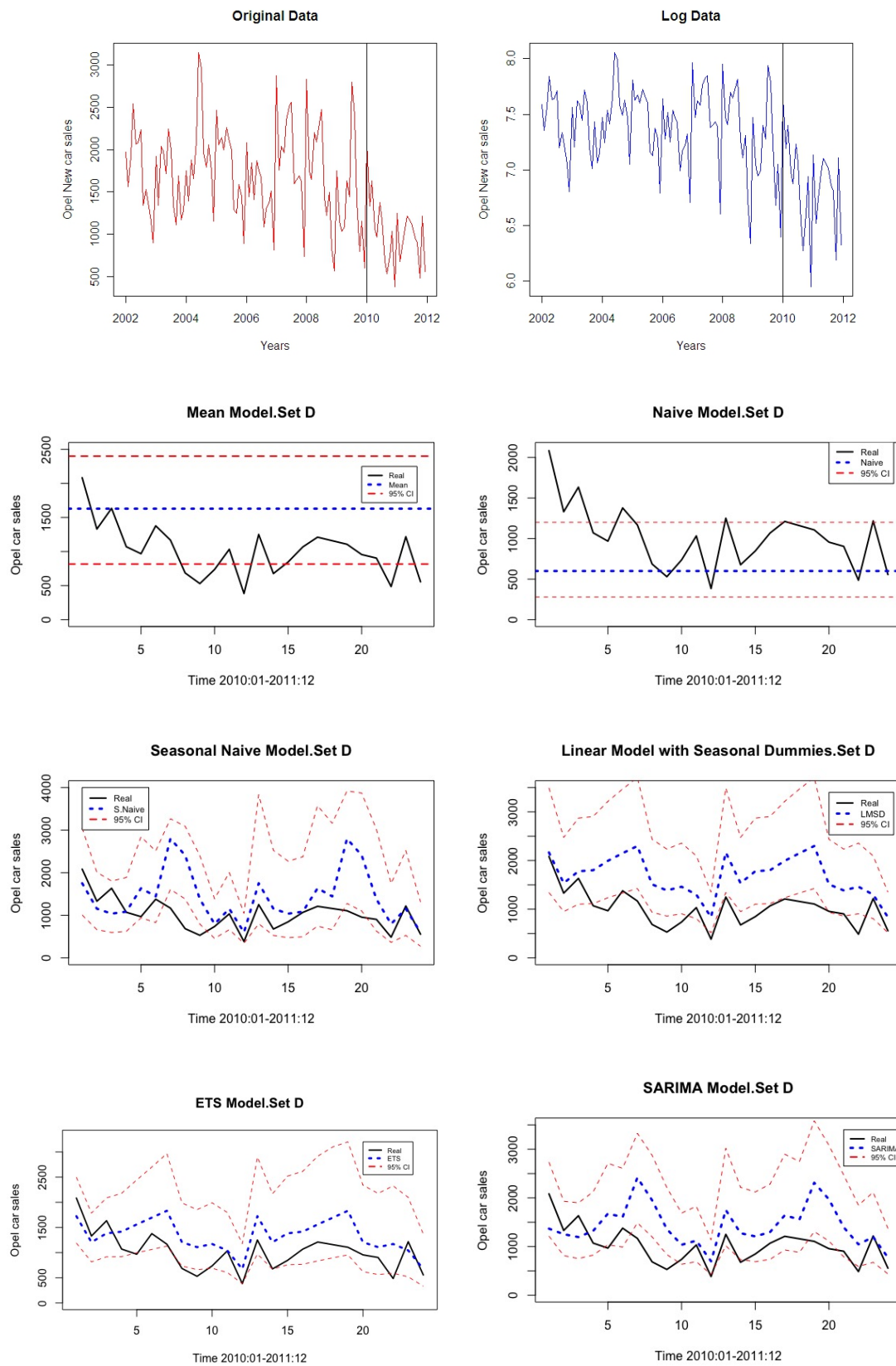


Figure 4.5: Forecasting for Opel. Best Model:ETS. Set D

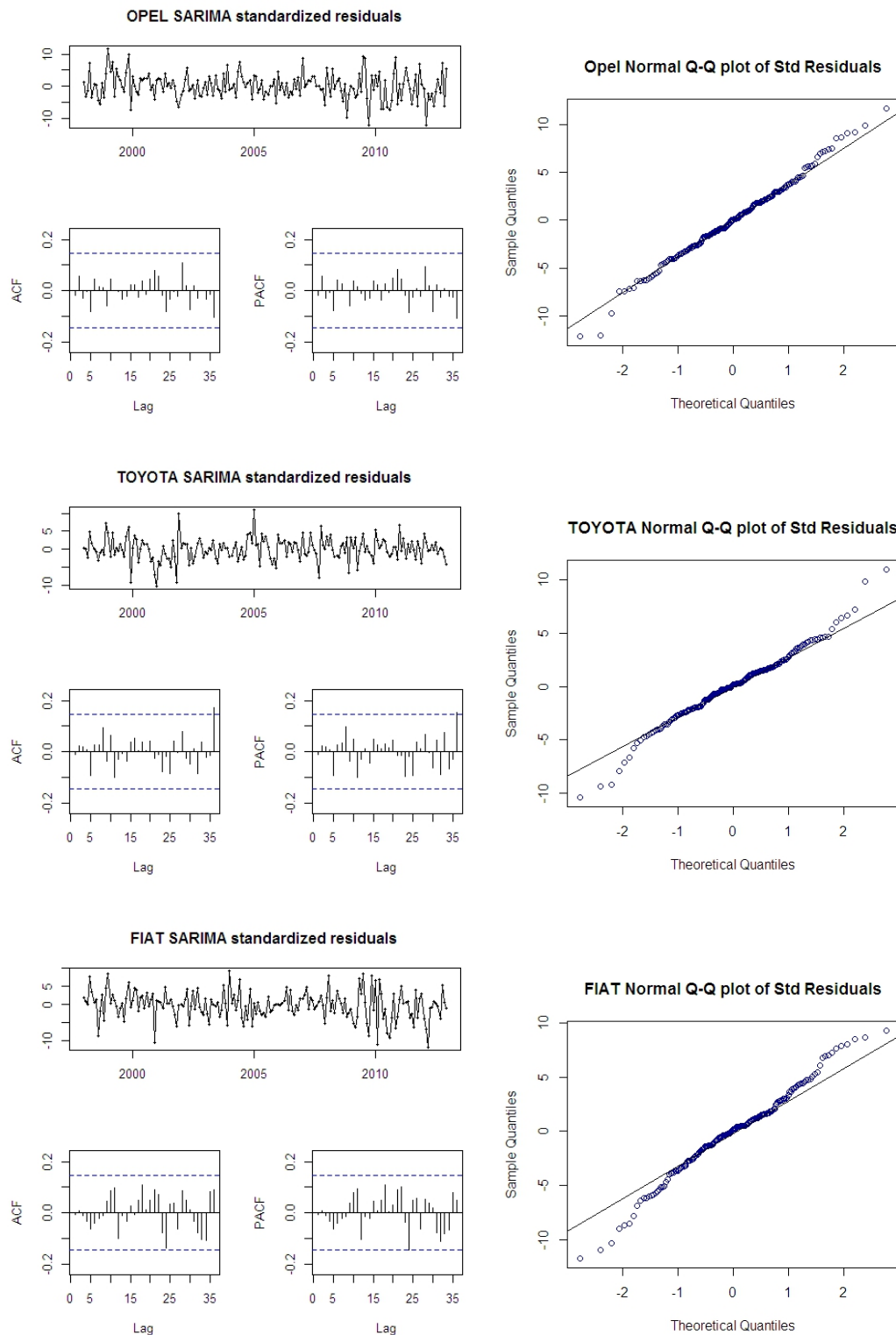


Figure 4.1: SARIMA Model residual analysis. OPEL,TOYOTA,FIAT.Set A

Table 4.4: Forecast Performance Evaluation for OPEL (Data:log values).

-OPEL-	RMSE	MAPE	MAE
Data Set A			
Mean	914	63	896
Naïve	188	27	127
S.Naïve	282	29	193
LMSD	993	62	938
ETS	314	34	273
SARIMA	225	40	156
Data Set B			
Mean	526	49	509
Naïve	240	63	202
S.Naïve	169	29	133
LMSD	613	48	551
ETS	147	22	117
SARIMA	234	64	196
Data Set C			
Mean	652	36	566
Naïve	668	88	531
S.Naïve	782	31	537
LMSD	644	34	543
ETS	471	28	393
SARIMA	655	33	501
Data Set D			
Mean	718	39	649
Naïve	563	75	455
S.Naïve	782	30	515
LMSD	745	39	666
ETS	436	29	392
SARIMA	617	32	501

Table 4.5: Forecast Performance Evaluation for TOYOTA (Data:log values).

-TOYOTA-	RMSE	MAPE	MAE
Data Set A			
Mean	755	42	599
Naïve	927	94	631
S.Naïve	1225	48	1064
LMSD	688	39	561
ETS	784	40	667
SARIMA	567	32	396
Data Set B			
Mean	783	24	481
Naïve	843	57	402
S.Naïve	771	21	422
LMSD	642	22	419
ETS	533	20	335
SARIMA	731	23	420
Data Set C			
Mean	569	23	447
Naïve	1502	195	1407
S.Naïve	464	20	347
LMSD	642	11	220
ETS	285	11	229
SARIMA	450	18	336
Data Set D			
Mean	1151	82	1038
Naïve	1627	63	1498
S.Naïve	533	23	426
LMSD	1067	73	962
ETS	948	60	854
SARIMA	951	56	818

Table 4.6: Forecast Performance Evaluation for FIAT (Data:log values).

-FIAT-	RMSE	MAPE	MAE
Data Set A			
Mean	664	65	648
Naïve	209	82	155
S.Naïve	164	42	116
LMSD	125	32	93
ETS	123	30	97
SARIMA	230	124	176
Data Set B			
Mean	307	47	283
Naïve	157	51	114
S.Naïve	134	34	104
LMSD	348	47	326
ETS	122	28	95
SARIMA	155	56	122
Data Set C			
Mean	430	41	383
Naïve	482	78	382
S.Naïve	524	48	433
LMSD	387	38	344
ETS	382	38	354
SARIMA	443	46	391
Data Set D			
Mean	534	46	484
Naïve	355	51	250
S.Naïve	536	45	437
LMSD	518	45	344
ETS	441	42	407
SARIMA	401	41	363

4.4 Discussion

In this chapter the research was in the out-of-sample forecast performance of six (6) different time series models: Mean or Average, Naïve, Seasonal Naïve, State space Exponential Smoothing (ETS), Linear Model with Seasonal Dummies (LMSD) and the Seasonal Autoregressive Moving Average (SARIMA) models.

The data time was from the beginning of 1998 till the end of 2016 and it was divided in four (4) different data sets, based on the rule of 80:20 for training and testing sets. The aim was to examine how good those time series models are in predicting new car sales levels under difficult economic situations like the one the Greek economy was facing in the last two decades.

Empirical evidence for three different firms Opel, Toyota and Fiat were studied. Results show that the basic Mean/Average model and the simple Naïve model were very poor in predicting sales levels. In this empirical application each firm corresponded differently and that is normal since each firm reports different sales levels and fluctuations thought out the years.

Additionally, the forecasting accuracy measures did not always result at the same model which is expected since the root mean squared error (RMSE) measures the average magnitude of the error (i.e. is the square root of the average of squared difference between prediction and actual observation), while the mean absolute percentage error (MAPE) measures the average magnitude of the errors in the set of predictions as a percentage (%). The main difference between those two measures is that RMSE gives a relatively high weight to large errors. This means that RMSE might be more useful when large errors are particularly undesirable because it penalizes large errors more.

Table 4.7: Summary of Best out-of-sample Forecasting Model(log values)

Set{Estimate}: Forecast	<u>RMSE</u>			<u>MAPE</u>		
	Opel	Toyota	Fiat	Opel	Toyota	Fiat
A {1998-2012}:2013-16	Naïve	SARIMA	ETS	Naïve	SARIMA	ETS
B {2006-2013}:2014-15	ETS	ETS	ETS	ETS	ETS	ETS
C {2006-2009}:2010-10	ETS	ETS	ETS	ETS	ETS	ETS
D {2002-2009}:2010-11	ETS	S.Naïve	Naïve	ETS	S.Naïve	SARIMA

In Table 4.7 on page 145 the summary of the empirical research evidences in choosing the best model in an out of sample forecasting is presented, according to the forecasting accuracy measures Root Mean Square Error (RMSE) and Mean Absolute Percentage Error (MAPE). Results from the Mean Absolute Error (MAE) metric are not presented because they are similar to the results given by the MAPE metric.

Empirical results for Opel give the choice of Naïve (in both RMSE and MAPE metrics) for data set A, which has 15 years of monthly observations in the training set and gives a long term forecast for four (4) years. However, in data set B, C and D, Opel sales are better predicted by ETS models, mainly because the training set covers a period of turbulent movements and deep recession for car sales due to the economic crisis in the market that was obvious after 2010. Furthermore, these models are better in short-run forecasting.

The study for Toyota new car sales give a preference in SARIMA model for the long-term forecast in data set A and the short-term forecast at the beginning of the deep recession period and prefer the ETS models for the data set B and C, during the deep recession period. Seasonality determine new car sales data ad therefore the S. Naïve model is selected as the best forecasting model for the data set D.

For Fiat new car sales the long-term and short-term forecast is estimated more accurate by the ETS model for data Set A,B and C. However, in data set D we get better results when using the Naïve or the SARIMA model according to RMSE and MAPE accordingly.

In general, there is no single model that is the best forecasting model for all the cases studied in this empirical research. Each firm has its own sales level, marketing program, goodwill and customers perception in the market place, so each one responds slightly different, in sales level, during a period of economic crisis. The training and testing set are seem to be crucial when forecasting, therefore the researcher, should make carefully the choice of the training and test sets.

Furthermore accuracy measures do not always agree and empirical results do not give always the same “best” model. Usually the RMSE measure give the same results as the MAPE but there are cases, few exceptions, that they do not agree and give different results. For example, for Fiat in Set D we can choose Naïve or SARIMA model according to RMSE or MAPE accordingly.

In general the results found in the calculations of the accuracy metrics agree and can be visually observed in the graphical presentation of the series forecast values in comparison with the actual values in each data set.

As far as confidence intervals (CI) are under consideration, it seems that if we have a short-term forecast period (like one year as in Set C) then the forecast values fail in predicting the real values of the variable as the forecasts go beyond the forecast 95% confidence interval area. Having two years ahead forecast seems better as forecast values are within the confidence interval boundaries and that seems to be safer for forecasting.

Additionally the narrower the CI is the better, especially when actual values and forecast values are within the 95% CI. In these cases forecast values are close to the actual ones and additionally the error may be possible but is limited.

SARIMA-GARCH Forecasting Models.

5.1 Introduction.

In this chapter we present a volatility forecasting comparative study within the autoregressive conditional heteroskedasticity class of models. Our goal is to identify if a SARIMA-GARCH model can successfully predict the volatility of car sales level during a period of economic crisis in the Greek market.

Forecast volatility is important for three main purposes : risk management, asset allocation, and for taking bets on future volatility of inventory. In 1982, Robert Engle developed the autoregressive conditional heteroskedasticity (ARCH) models, to model the time-varying fluctuations, often observed in economical time series data. For this contribution, he won the 2003 Nobel Prize in Economics. The ARCH models assume, the variance of the current error term or innovation to be a function of the actual sizes of the previous time periods' error terms, and often the variance is related to the squares of the previous innovations.

There is a vast literature on volatility and Poon and Granger [2003, 2005] provide an extensive survey of the literature's main findings, while T.G.Andersen et al. [2006] provide a comprehensive theoretical overview on forecasting. Brownlees et al. [2011] present a volatility forecasting comparative study within the autoregressive conditional heteroskedasticity (ARCH) class of models and find that volatility, during the 2008 crisis, was well approxi-

mated by predictions made with these models in the short-run.

Having explored the general theory of Seasonal Autoregressive Integrative Moving Average Models (SARIMA) and Generalized Autoregressive Conditional Heteroskedasticity (GARCH) models in the preceding chapters, this study introduce the univariate SARIMA-GARCH models, in an attempt to examine their forecast ability. The family of GARCH models are useful because one can predict better the volatility of the variables with those models. However they are found to be perfect for predicting a few periods ahead but not very good for long terms predictions. These models are helpful in expanding the phenomenon of volatility clustering. This phenomenon has periods of relative calm and periods of high volatility which is common in market data, like sales.

In our study, car sales volatility clearly moves around though time and its seasonality depends on the particular market, where trading happens. So the GARCH model view is that volatility spikes upwards, and then decays until there is another spike.

An autoregressive approach helps to build more accurate and reliable volatility models. According to Tsay [2005] market volatility is known to cluster, which means that highly volatile periods tends to persist for sometime before the market returns to a more stable environment. The ARCH models structure refers to an autoregressive model since ϵ_t clearly depends on previous ϵ_{t-i} and is conditionally heteroscedastic. The GARCH family of models is widely used in practice for prediction of financial market volatility and returns.

Briefly the steps for building a SARIMA–GARCH model are introduced in this research:

- Step One: Specify a mean equation by testing for serial dependence in the data and building a SARIMA model for the sales series to remove any linear dependence and seasonality.
- Step Two: Use the residuals of the mean equation to test for ARCH effects (Ljung-Box Test for residuals or Lagrange Multiplier Test).
- Step Three: Specify a volatility model if ARCH effects are statistically significant and perform a joint estimation of the mean and volatility equations.
- Step Four: Finally check the fitted model carefully and refine it if necessary.

This study proceeds with GARCH model selection and analysis and end up with the diagnostic checking of the model, but before fitting the GARCH model the research try to diagnose a 2^{nd} order dependency using ACF of squared residuals and test for ARCH effects.

Briefly, the plan of this chapter is the following: after the short introduction we first plot the residuals and then examine the autocorrelation function of the SARIMA residuals and SARIMA squared residuals, to see if they give some evidence of autocorrelation. Then we apply tests on these SARIMA residuals a) for autocorrelation the Ljung-Box and b) for an ARCH effects the Lagrange Multiplier test. They show evidence of Autocorrelation only in Set A for Opel and Toyota and evidence of an Arch effect only for Opel in data set D.

In section two, we fit the SARIMA-GARCH(1,1) model only in Opel and Toyota at Data Set A where there was evidence of autocorrelation in the SARIMA residuals. We specify the mean equation with the pre-selected SARIMA models in our research in the former chapter (see chapter 4) and continue with the volatility equation, using the simple GARCH(1,1) model. This study also examine, three alternative distributions on the SARIMA-GARCH (1,1) models to find the most suitable one. The empirical evidence show that GARCH with student-t distribution is the best case against the normal and the generalized error distribution.

Furthermore in section three we make the diagnostic checking of the SARIMA-GARCH(1,1) models residuals and research findings confirms that there is no correlation left in the residuals and that the fit of the models is good. In section four, we focus on SARIMA-GARCH prediction intervals and show how they are reduced, especially in the short-run, and provide better forecasts. Lastly this chapter ends up with a discussion for the findings of this empirical research on SARIMA-GARCH models in the Greek new-car market. According to RMSE forecasting results the SARIMA-GARCH(1,1) models can produce better forecasts against the SARIMA, when SARIMA residuals give evidence of serial correlation.

5.1.1 Autocorrelation function of SARIMA model residuals.

The first step to identify the volatility clustering is to plot the SARIMA residual series and then test for autocorrelation and ARCH effect. In volatility clustering, we expect high fluctuations to be followed by high fluctuation, and calm movements to be followed by calm movements, and that is what can be seen in Figure 5.1 on page 151. However the flow over the years is turbulent and do not exhibit persistence in the high or low levels of the series which is typical when series have an ARCH effect.

The identification for an ARCH effect can also be explored via the autocorrelation (ACF) plot of the SARIMA residuals and squared residuals. According to the literature, it is remarkable how squared residuals can magnify the effect in the series that exhibit autocorrelation. If SARIMA squared residuals are serial correlated, then ARCH effects are present. In general, if a time series exhibit conditional heteroskedasticity - or autocorrelation in squared series - it is said to have autoregressive conditional heteroskedastic (ARCH) effects.

On evaluating autocorrelations of residuals and squared residuals of the fitted SARIMA models, for each data set of Opel, Toyota and Fiat new cars' sales in the Greek market, we found that SARIMA residuals give no evidence of autocorrelation. However the *squared* SARIMA residuals give some evidence of serial correlation. That lead us to the conclusion that conditional heteroskedastic behaviour *might* be present since we have slightly high level of ACF at some specific lags of the squared SARIMA residuals, as can be seen in Figure 5.2(on page 153) for example :

Opel: lag 11 for squared residuals of the fitted SARIMA(1, 0, 1)(2, 0, 0)₁₂ model

TOYOTA: lag 11 for squared residuals of the fitted SARIMA(4, 0, 0)(1, 0, 0)₁₂ model

FIAT: lag 5 for squared residuals of the fitted SARIMA(1, 0, 1)(2, 0, 0)₁₂ model

Generally, in research if no autocorrelation in the squared standardize residuals can be found, then volatility is properly explained by the SARIMA model. In our case study, there are some weak evidence that an ARCH/GARCH model might be needed because the

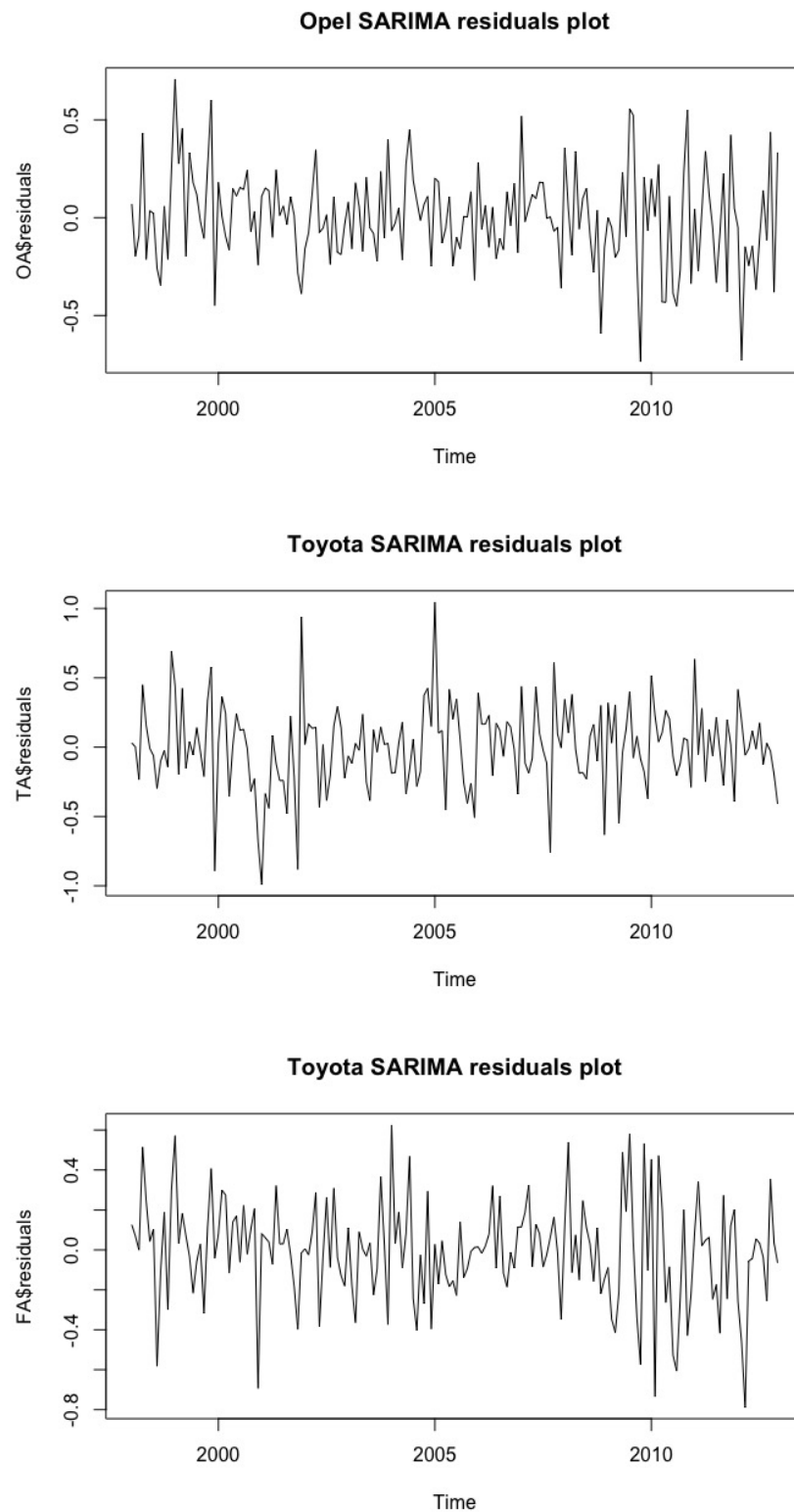


Figure 5.1: SARIMA residuals Plot (Set A)

squared standardized residuals of the SARIMA models have autocorrelation i.e. evidence of volatility clustering and a need for GARCH model specification.

However, the case of FIAT new car sales the ACF of the squared residuals looks more like a realisation of a discrete white noise process, indicating that the serial correlation present in the squared residuals have all been explained with the appropriate mixture of SARIMA(1, 0, 1)(2, 0, 0)₁₂ model.

So there is no strong evidence of a GARCH effect in our data series according to the ACF of the squared residuals, however we choose to proceed further the research into the SARIMA-GARCH models, in order to investigate whether they add usefully information. and determine later what value they add in this study.

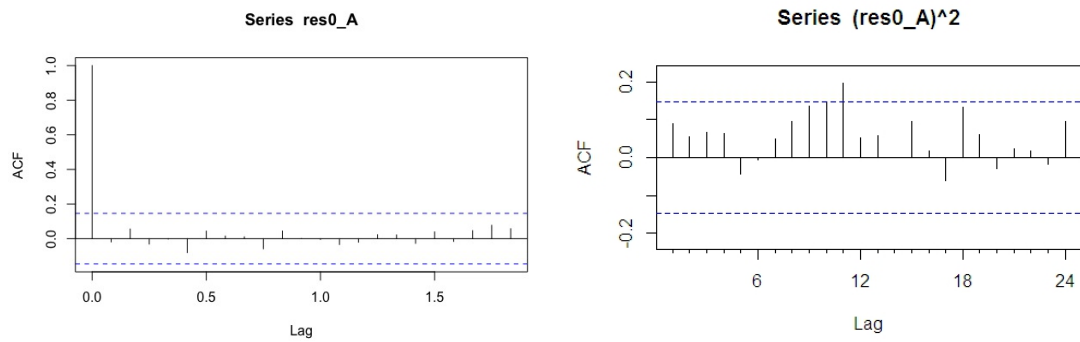
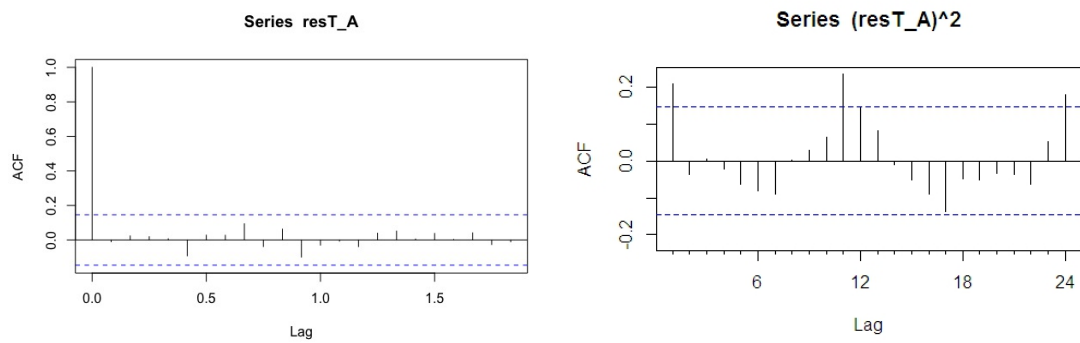
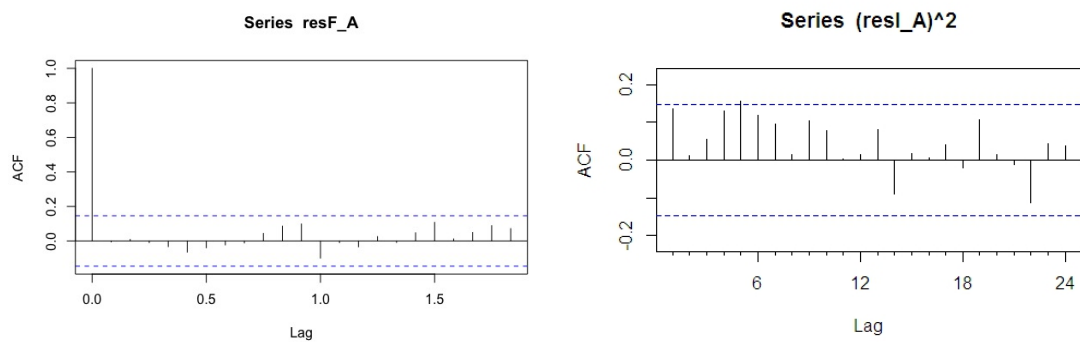
OPEL ACF of SARIMA residuals & squared residuals**TOYOTA ACF of SARIMA residuals & squared residuals****FIAT ACF of SARIMA residuals & squared residuals**

Figure 5.2: ACF of SARIMA residuals & squared residuals (Set A)

5.1.2 Testing for ARCH/GARCH Effects

There are two test used for the error variance in order to test the evidence of an ARCH effect, before fitting a GARCH model. If the error variance is not constant the data are said to be heteroscedastic and hence have an ARCH effect.

The two test for ARCH effect are :

- Ljung-Box (LB) test statistic (LB) which is based on *squared* residuals and is used to test for independence of the series.

In this test the null hypothesis is *Ho: There is no autocorrelation*

Therefore if the test give $p - value < 0,05$ then we reject H_0 , hence there is autocorrelation i.e. volatility clustering and ARCH effect.

- Lagrange Multiplier (LM) test for autoregressive conditional heteroscedasticity (ARCH) LM test statistic of Engle [1982], as described by Tsay [2005] is widely used as a specification test in univariate time series models. It is a test of no conditional heteroskedasticity against an ARCH model.

In this test the null hypothesis is *Ho: There is no ARCH effect*

Consequently if $p - value < 0,05$ then we reject H_0 , hence there is an ARCH effect.

Results of the Ljung Box and Lagrange Multiplier Test are recorded in Table 5.1 (page 155). The Ljung Box tests, on squared residuals of the selected SARIMA forecasting models, are computed and applied in our series for twelve (12) lags - since our data are monthly observations exhibiting seasonality. According to Ljung-Box Test in SARIMA squared residuals are uncorrelated, except Opel and Toyota SARIMA squared residuals that are suffering from serial correlation, in data Set A. So all residuals are serially uncorrelated except those two cases. Hence it is necessary to develop a better time series model for analysis of new car sales and a GARCH model is proposed to handle heteroskedasticity in both series. On the other hand the ARCH Lagrange Multiplier Test on levels of residuals of the selected SARIMA forecasting models, are computed and show no ARCH effect in the residuals except in the case of OPEL in data set D.

Table 5.1: Testing SARIMA residuals for Heteroscedasticity.

Ljung-Box and Lagrange Multiplier Tests (lag order=12)				
OPEL				
Data Set	Ljung-Box Test	<i>p – values</i>	LM Test	<i>p – values</i>
A	21,71	(0,04)*	15,72	(0,20)
B	6,65	(0,87)	5,36	(0,94)
C	9,94	(0,62)	19,84	(0,07)
D	16,00	(0,19)	38,34	(0,00)**
TOYOTA				
A	27,47	(0,00)*	20,16	(0,06)
B	2,40	(0,99)	3,90	(0,98)
C	8,78	(0,72)	9,05	(0,69)
D	3,42	(0,99)	2,91	(0,99)
FIAT				
A	19,07	(0,08)	13,74	(0,31)
B	13,36	(0,34)	12,23	(0,42)
C	7,49	(0,82)	7,34	(0,83)
D	10,88	(0,53)	8,51	(0,74)

Note: * Autocorrelation **Arch Effect

LB Test Ho:No Autocorrelation. If p-value< 0,05 Reject Ho

LM Test Ho:No ARCH Effect. If p-value< 0,05 Reject Ho

Surprisingly, the two test -Ljung Box and Lagrange Multiplier Test - do not seem to support each other results. In Set A, Opel and Toyota LB tests show serial correlation in SARIMA squared residuals and it was expected that LM tests would support those evidence giving an ARCH effect in SARIMA residuals, but LM tests outcome do not confirm those results. Additionally, in Set D, Opel seems to have an ARCH effect (according to LM test) but squared residuals are serial uncorrelated (according to LB test). So we conclude that this give evidence of a very weak ARCH effect in the series.

In this study it is clear that there is no constant variance, therefore the research continues with the estimation of a GARCH(p,q) model, where p stand for number of period of past squared returns and q stands for the previous variance into consideration. It is however not acceptable to apply ARCH model of high order, like for example of order 11 or more as suggested in the ACF plot of Squared residuals for OPEL and TOYOTA selected SARIMA models (see Figure 5.2 page 153). The extraordinarily large number of parameters and an order of this degree would generally not be considered a good fit since it would be inconvenient and difficult to work with such a high order process. The order can be restricted and used to estimate the coefficients of a desired process of lower order, so the parsimonious GARCH(1,1) model is applied that generally gives a good fit.

5.2 Fitting GARCH (1,1) Model

5.2.1 Model Selection and Analysis

Our strategy in choosing the appropriate GARCH model from competing models is based on the Akaike information criterion (AIC) and Bayesian information criterion (BIC) while the idea is to have a parsimonious model that captures as much variation in the data as possible. Therefore a quick comparison of the residuals from R software to fit GARCH models, using the squared residuals of SARIMA models of the series, are used to determine the best fitting model. The suggested models with their respective fit statistics are given in Table 5.2 (page 158) only for the data set A and for three firms: Opel and Toyota that

showed evidence of autocorrelations in the SARIMA squared residuals.

The study specifies the mean equation with the pre-selected SARIMA models (see Table 4.3 on page 132) and continue with the specification of a volatility model since there is autocorrelation and therefore signs of an ARCH effects. A joint estimation is performed of the mean and the volatility equation. Usually, the simple GARCH model captures most of the variability in most series, therefore small lags for p and q are common in applications. A typical GARCH (1,1) model is adequate for modeling volatility even over long sample periods with normal conditional distribution, however alternative conditional distributions are also explored (See Table 5.2).

5.2.2 Alternative Conditional Distributions

In this comparative study we have included SARIMA - GARCH(1,1) model estimation for Opel and Toyota using alternative conditional distributions¹ like the normal one, the generalized error distribution and the fat tailed student-t distribution for data set A and the pre-selected SARIMA models for each firm and check which conditional distribution is more appropriate for modelling time varying variances of our data according to their AIC and BIC performance metrics.

Results in Table 5.2 illustrate the AIC and BIC of the various models. We know that small AIC and BIC values are better than large AIC and BIC values i.e. makes the model unfavourable. However since all values are negative, the best model, is the one with the highest absolute value. As a conclusion, we notice that the selection criteria AIC and BIC for Opel and Toyota in data set A, have their lowest values at the SARIMA-GARCH(1,1) model with a fat tailed Student-t distribution (STD). Hence, these time series observations have a distribution that one often assumes to be normal (Gaussian) but in reality they usually tend to be **leptokurtic** (i.e. fat tailed).

Parameters estimation of SARIMA(1, 0, 1)(2, 0, 0)₁₂-GARCH(1,1) models with Student-

¹Alternative conditional distributions: Normal, Generalized Error (GED) and Student-t (STD) distribution.

Table 5.2: Comparing alternative conditional distribution of variance equation.

Set A	Volatility Forecasting Models	Cond.		
		Distrib.	AIC	BIC
OPEL	$SARIMA(1, 0, 1)(2, 0, 0)_{12} - GARCH(1, 1)$	Normal	-1.9067	-1.8358
	$SARIMA(1, 0, 1)(2, 0, 0)_{12} - GARCH(1, 1)$	GED	-2.5662	-2.4775
	$SARIMA(1, 0, 1)(2, 0, 0)_{12} - GARCH(1, 1)$	STD	-2.7940	-2.7053
TOYOTA	$SARIMA(4, 0, 0)(1, 0, 0)_{12} - GARCH(1, 1)$	Normal	-0.8661	-0.7952
	$SARIMA(4, 0, 0)(1, 0, 0)_{12} - GARCH(1, 1)$	GED	-1.6776	-1.5890
	$SARIMA(4, 0, 0)(1, 0, 0)_{12} - GARCH(1, 1)$	STD	-1.9746	-1.8859

t distribution derive using the method of maximum likelihood and estimate² the coefficients of the SARIMA-GARCH models. The coefficients for OPEL in set A are given, as an example of model analysis, below for the mean and volatility equations:

ar1	ma1	sar1	sar2	intercept
0.9160	-0.4442	0.3349	0.4637	6.9055
s.e.	0.0375	0.0847	0.0721	0.0762
	0.4352			

mu	omega	alpha1	beta1	shape
0.02137236	0.00038873	0.11904089	0.93258634	2.09758114

Notice that in the coefficients of Opel the sum of $\alpha_1 + \beta_1 = 0.11904089 + 0.93258634 = 1.0515$ is close to unity, which is a standard GARCH (1,1) model restriction.

The full SARIMA(1, 0, 1)(2, 0, 0)₁₂-GARCH(1,1) model for OPEL with its estimated coefficients is represented as:

$$(1 - 0.9160B)(1 - 0.3349B^{12} - 0.4637B^{24})x_t = 6.9055 + (1 - 0.4442B)\epsilon_t \quad (5.1)$$

²estimated with R package “fGarch” and “garchFit” function

$$\sigma_t^2 = 0.02137236 + 0.11904089\epsilon_{t-1}^2 + 0.93258634\sigma_{t-1}^2 \quad (5.2)$$

Notice that ϵ_{t-1}^2 is the ARCH term i.e. the news about volatility from the previous period measured, as the lag of the squared residuals from the mean equation. The σ_{t-1}^2 is the GARCH term i.e. it is the last period's forecast variance.

The value of $\alpha_1 + \beta_1$ is close to unity and this agrees that the volatility shocks are totally continual. The coefficients of the squared residuals are positive and statistically significant showing that GARCH effect is evidence for our data.

5.3 Diagnostic checking of SARIMA-GARCH Models

Diagnostic checking of Opel's SARIMA(1, 0, 1)(2, 0, 0)₁₂-GARCH(1,1) model with t-Student distribution is crucial before accepting it as a fitted model and interpret its findings. It is essential to check, if the model is correctly specified, whether the model assumptions are supported by the data, because if any key assumptions seem to be violated, then we should specify a new model, that will be fitted and checked again, until the model is found adequate to fit the data.

A useful tool for checking the model specification, is the standardized residuals where the model checking is done through analyzing the residuals from the fitted model. In time series model analysis, the selection of the best model to fit the data, is directly related to whether residuals analysis is performed well. One of the GARCH model assumptions is that, for a good model the residuals must follow a white noise process.

For SARIMA(1, 0, 1)(2, 0, 0)₁₂-GARCH(1,1) model with Student-t distribution fitted for OPEL new car sales we get the following results:

Jarque-Bera test statistic=1329.181 (too large) with p-value =0 and

Shapiro Wilk test statistic=0.6313004 with p-value =0.

Both tests are significant and both reject the null hypothesis at 5% level, that means the distribution is normal. If the volatility clustering is properly explained by the model, then there will be no autocorrelation in the squared standardized residuals.

It is common in research to do a Box-Ljung test to test for autocorrelation. Table 5.3

on page 160 gives the output for Box-Ljung on a fit assuming a normal distribution on returns for Opel new car sales. It shows the p-values of residuals of Box-Ljung Q statistics of the SARIMA(1, 0, 1)(2, 0, 0)₁₂-GARCH(1,1) model with Student-t distributional fitted for OPEL and Toyota, all are well above 0.05, indicating “non-significance”. This is a desirable research result, as it shows no significant serial correlation in the residuals. So we assume that the residuals are uncorrelated.

Table 5.3: Box-Ljung Q-Test of standard. residuals SARIMA-GARCH(1,1)-STD Set A.

OPEL		
SARIMA(1, 0, 1)(2, 0, 0) ₁₂ -GARCH(1,1)		
Lag	Q-Statistic	p-value
10	5.0832	0.8855
15	12.6624	0.6283
20	16.080	0.71161
TOYOTA		
SARIMA(4, 0, 0)(1, 0, 0) ₁₂ -GARCH(1,1)		
Lag	Q-Statistic	p-value
10	8.2313	0.6062
15	18.6938	0.2279
20	25.3744	0.1874

Additionally, the standardized residuals of the SARIMA(1, 0, 1)(2, 0, 0)₁₂-GARCH(1,1) model fitted for OPEL (i.e. the residuals divided by their conditional standard deviation) are plotted in Figure 5.3 (page 161) for a fat tailed student-t distribution, and appear to be random.

The plot in Figure 5.3 looks like a white noise expect for the change in spread (variation) of the series of standardized residuals. Such heteroskedasticity would most likely not be evident in a truly random data set. Notice the relationship between the standard residuals

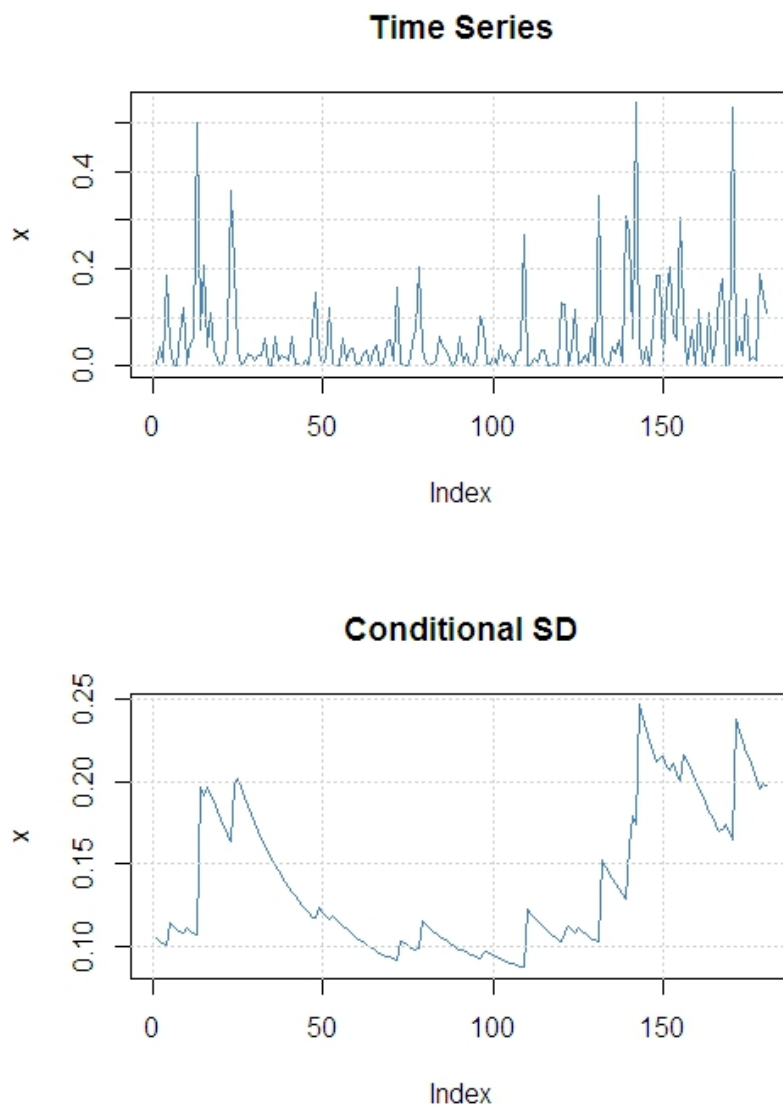


Figure 5.3: OPEL SARIMA-GARCH(1,1)STD stand.residuals & deviation-Set A

derived from the fitted data and the corresponding conditional standard deviations of Opel new car sale plots. The Figures show that there are some extreme values especially in the year 2011 and 2012. The residuals in this Figure display the observed volatility of modeling sales with a univariate SARIMA(1, 0, 1)(2, 0, 0)₁₂-GARCH(1,1) model.

The researcher observed some aspects of volatility clustering phenomenon, which covers some periods of relative calm and some periods of high volatility, that is a universal attribute of market data and becomes even more intense in periods of financial crisis.

Figure 5.3 shows an Opel SARIMA (1, 0, 1)(2, 0, 0)₁₂ - GARCH(1,1) model with Student-t distribution model of volatility as it moves around through time, and this indicated that there is an ARCH effect which means in other words some stationary parts and some more changeable parts. That is the volatility clustering, the phenomenon of there being periods of relative calm and periods of high volatility, which is a common attribute of market data, like car sales.

The GARCH model view is that volatility spikes upwards and then decays away until there is another spike. The peaks of the model residuals coincide with the peaks of the standard deviation shown in the graphs in Figure 5.3. Moreover the standard deviation of the SARIMA (1, 0, 1)(2, 0, 0)₁₂- GARCH(1,1) Opel model process shows that there is a high volatility in the beginning of year 2000 and at the year 2010 and 2012.

Moreover if the model is successful at modeling the serial correlation structure in the conditional mean and conditional variance then there should be no autocorrelation left in the standard residuals and squared standard residuals.

The ACF plot of squared standard residuals Figure 5.4 on page 163, show no correlation left and some peaks of squared standardized Residuals are reduced. Therefore we proceed to use the model to forecast future values of the Opel new car sales. The second plot in Figure 5.4 give the ACF of standardized residuals and of squared standardized residual of the SARIMA (1, 0, 1)(2, 0, 0)₁₂ -GARCH(1,1) Opel model process. More specific the graph of the ACF of standardized residuals (figure show there are peaks i.e. autocorrelations but within ACF boundary and die out slowly and this indicates that there is correlation between the magnitude of change in the residuals. Meaning that there is serial dependence

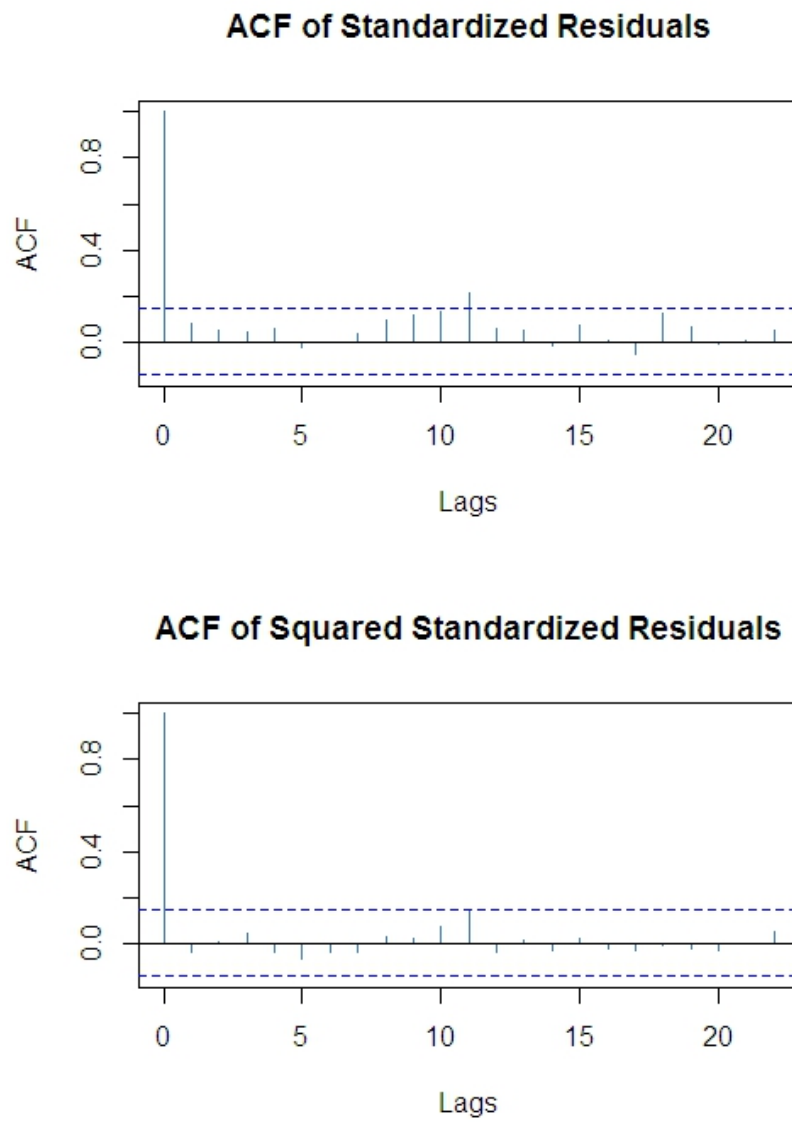


Figure 5.4: OPEL SARIMA-GARCH(1,1)STD residuals ACF plots-A

in the variance of the data that is the variance of residuals is conditional on its past history and may change over time.

Generally, SARIMA is a model for the realization of a stochastic process imposing a specific structure of the conditional mean of the process. GARCH is a model for the realization of a stochastic process imposing a specific structure of the conditional variance of the process. In this research we use GARCH models to capture volatility dynamics of SARIMA type models. Fitting SARIMA in mean equation of GARCH model helps correct the problem of serial correlation and seasonality in the residuals. Once the absence of serial correlation and seasonality is confirmed by adding required SARIMA terms, the conditional volatility can be modeled using GARCH. The residuals of GARCH model confirm the absence of serial correlation and Arch effect.

5.4 SARIMA-GARCH Prediction Intervals

Volatility forecasting assessments are commonly structured to hold the test asset and estimation strategy fixed, focusing on model choice. Our pragmatic approach in this research, consider the SARIMA-GARCH model family model for the new car sales and try to predict volatility and the confidence intervals (CI) of the forecast values within 1.96 standard deviations of the mean. This is done for a range of SARIMA models which differ in accordance with the different car firm and their new car sales level.

Volatility forecasting focus almost exclusively in how much the confidence intervals are reduced and provide a better forecast value and how volatility levels can escalate dramatically especially in crisis periods. This thesis research draws attention to the dramatically improvement of the prediction and the decrease of confidence intervals by the implementation of SARIMA-GARCH models in forecasting. This is perhaps the most challenging application of volatility forecasting, however, it is used for developing a volatility inventory strategy. Decision makers often develop their own forecast of sales volatility, and based on this forecast they compare their estimate sales level with the market demand of new car models. The simplest approach to estimating volatility is to use historical standard

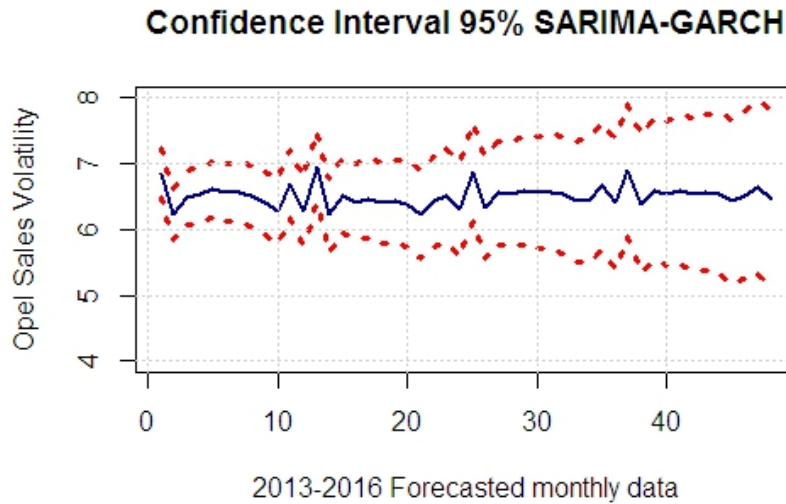


Figure 5.5: Confidence intervals 95% SARIMA-GARCH model (Opel-Set A)

deviation, but there is some empirical evidence, that this can also be improved.

The researcher calculated the SARIMA models, and form the mean equation with the critical values that delineated the region of rejection, continue with the GARCH(1,1) estimation with t-student conditional distribution that predict volatility. For a two-tailed test the distance to these critical values is also called the margin of error and the region between critical values is called the confidence interval. Such a confidence interval is commonly formed when we want to estimate a population parameter, like sales level, where the interval estimate contains a range of reasonable or tenable values both the population mean and population standard deviation.

Since 95.0% of a normally distributed population is within 1.96 (95% is within about 2) standard deviations of the mean, we can often calculate an interval around the statistic of interest which for 95% of all possible samples would contain the population parameter of interest. Our 95% confidence intervals are then formed with $z=\pm 1.96$

A prediction interval gives an interval, within which we expect our predicted value to lie, with a specified probability. In the SARIMA-GARCH models, for example, if we assume that the forecast errors are normally distributed, a 95% prediction interval for h -step forecast is: $\hat{y}_{n+1} + Z_{1-\alpha/2} * \hat{\sigma}_{n+1}$ and for the **lower** prediction interval: $\hat{y}_{n+1} - Z_{1-\alpha/2} * \hat{\sigma}_{n+1}$

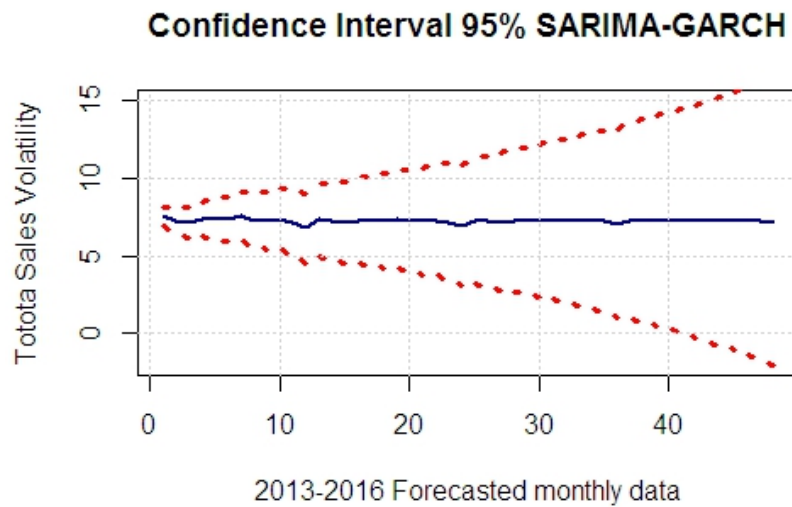


Figure 5.6: Confidence intervals SARIMA-GARCH model(TOYOTA-Set A)

where $\hat{y}_{n+1}$ is an estimate of the mean from the SARIMA model and $\hat{\sigma}_{n+1}$ is an estimate of the standard deviation of the $n+1$ step forecast distribution.

In this study, we calculate 95% intervals and that means that the value of $Z_{1-\alpha/2}$ will equal 1,96. The important aspects that prediction intervals cover is that they better express the uncertainty in the forecast and they describe in a better way, how accurate the forecast are. We must add prediction intervals in our SARIMA-GARCH study as we did in the rest of the six models we examined, because it helps to make clear how much uncertainty is associated with each forecast.

Table 5.4: Volatility Forecast for monthly new car sales (Opel).

Horizon	1	2	3	4	5	6	7	8
Mean	6.858	6.223	6.485	6.527	6.607	6.569	6.577	6.506
Volatility	0.194	0.200	0.206	0.212	0.218	0.225	0.231	0.238

According to the results in Table 5.4 its is obvious that we have very low volatility in one step ahead horizon forecasts and at the 95%CI it gets wider as time horizon increases.

More specific in SARIMA-GARCH models, we take the mean equation from the SARIMA model and the variance equation from the GARCH model, it is important to estimate the prediction intervals of the volatility, in order to show that we end up with a closer estimation for the variance fluctuations using GARCH models and a more accurate confidence interval for it.

The plots of 95% prediction intervals illustrate in red lines the upper and the lower interval for the data set A for Opel and Toyota SARIMA-GARCH models (see figures 5.5 and 5.6) as an example in this section. It is obvious that the SARIMA GARCH models can be more accurate in the short run than in the long run.

The SARIMA-GARCH models encounter the same weaknesses as the SARIMA-ARCH models. Firstly, they respond equally to positive and negative shocks. Secondly, the tail behavior of GARCH models remain too short even with standardized Student-t innovations as recent empirical studies of high-frequency financial time series indicate.

5.5 Discussion

This chapter is dedicated in addressing the problem of forecasting and volatility forecasting in new car sales with the use of SARIMA-GARCH models alone with several other methods that are heavily used in practice and tested their accuracy using real data from the Greek market sector (i.e. new car sales of Opel and Toyota from the year 1998 till 2016). The family of SARIMA-GARCH methods has its advantages and disadvantages which are described in this study. Some models are simple easy to implement, yet yield good results. Other methods are more difficult to implement but do not yield always in good results. In short, there is no single preference approach.

The research found that there was not a very strong ARCH effect on the SARIMA residuals autocorrelation function (ACF), yet we decided to continue our study since the Ljung-Box Test gave evidence of autocorrelation in the SARIMA residuals for Opel and Toyota in Data Set A. The SARIMA-GARCH (1,1) with student-t distribution was preferred, according to the AIC and BIC information criterion among the normal and the

generalized alternative distribution. However, the results given in testing the data set A were encouraging. Based on Table 5.5 the forecasts produced by SARIMA-GARCH(1,1) model are better, since the root mean squared error (RMSE) metrics are lower than those produced by SARIMA. So it can be concluded that in the case of monthly car sales of Opel and Toyota, the SARIMA-GARCH(1,1) model can be an effective way to improve forecasting accuracy.

Table 5.5: Forecasting Results (RMSE).

R M S E	SARIMA	SARIMA-GARCH(1,1)
OPEL (1, 0, 1)(2, 0, 0) ₁₂	225	110
TOYOTA (4, 0, 0)(1, 0, 0) ₁₂	565	230

It is also incredible, how SARIMA-GARCH models help so much to reduce the prediction intervals of the forecast, and hence can provide a better forecast. On the other hand, according to the prediction intervals in Figure 5.5 and 5.6, it seems that SARIMA-GARCH models should better be used for a quick approximation of the volatility forecast in the short run and not in the long run, since they can produce more accurate forecasts in the short run. Additionally this is a valuable research experience of implementing the family of the SARIMA-GARCH models to the new car sales in the Greek market and add value to science and to the research of the vehicle market sales in Greece.

Data Transformations and Forecasting.

6.1 Introduction.

Historical data can often be adjusted or transformed from their initial values, according to Box and Cox [1964], to lead forecasting research into a simpler task. Research, in everyday reality, shows that almost all analyses benefit from improved normality of the variables, particularly in cases where substantial non-normality is present. Until this point, we have selected a traditional transformation - log values of the original data - which is frequently used in research, for improving normality of our data and producing relationships with more homogeneous residuals, in other words stable variance (Nelson and Granger [1979]).

However, using a broader class of power transformation, introduced by Box and Cox, will help us easily find the optimal normalizing transformation for our variables. This represents a family of power transformations that incorporates and extends the traditional options. Therefore, we continue our research using the general Box-Cox transformation, which represents a potential best practice, because normalizing data and equalizing variance is desired. Additionally this thesis examine the case of using the original data, which means no transformation at all, and compare the “in-sample” estimation fitting results and forecasting “out-of-sample” ability of various time series models.

Our main research focus, in this chapter, is on how data transformation effect time

series modelling and forecasting process. As empirical results show, transformation often considered to stabilize the variance of a series but can also be used to make highly skewed distributions less skewed or reduce the influence of outliers.

The plan of this chapter is the following. In Section 2, after an short introduction, a detailed literature review of Box and Cox data transformation is discussed, followed, in Section 3, by a section of its methodology process and its back transformation methodology. In Section 4, the empirical results of data transformation are presented initially for the Exponential smoothing (ETS) model, in an “in-sample” and an “out-of-sample” estimation. The research goes deeper and examine on the one hand, the case of none transformation at all, which means the use of original values, and on the other hand the research continues using several mathematical transformations of the family of Box-Cox transformation for various values of λ . Additionally the research evolves in more time series models and various data transformation. In Section 5 the confidence interval of the best models are presented and conclusions are drawn at Section 6 as a sum up of this chapter.

6.2 The Box-Cox Transformation Review.

Tukey [1957] is usually credited with presenting the initial idea that transformations can be thought of, as class or family of similar mathematical functions. However the theoretical foundations was set by Box and Cox [1964], who introduced transformation in time series research, as a way to allow for non-linearity between the original and transformed values, and to ensure that the disturbance can be well approximated by normal errors, using the properties of constant variance and uncorrelated errors.

It was Draper and Cox [1969], who first observed the high degrees of inconsistency in the estimator of λ , coming from the skewed observed values of the transformed process. Later, Box and Jenkins [1970] suggested that, using the Box-Cox transformation to validate not only the constant variance assumption, but all the underlying assumptions of an Auto - Regressive Integrated Moving Average (ARIMA) model, can be done by estimating the transformation index (λ) together with the model parameters. From the theoretical

standpoint, Granger and Newbold [1976] provided a general analytical approach, based on the Hermite polynomials series expansion, to forecasting transformed series.

A more extensive investigation was carried out by Nelson and Granger [1979], who considered a data set consisting of twenty-one time series. After fitting a linear ARIMA model to the power transformed series, and using 20 observations for post-sample evaluation, they concluded that the Box-Cox transformation does not lead to an improvement of the forecasting performance. Another important conclusion, supported also by simulation evidence, is that the Naïve forecasts, which are obtained by simply reversing the power transformation, perform better than the optimal forecasts based on the conditional expectation. The explanation is that the conditional expectation underlying the optimal forecast assumes that the transformed series is normally distributed, but this assumption may not be realistic. In contrast to Nelson and Granger's results, Hopwood et al. [1981] research, found for a range of quarterly earnings per share series that the Box-Cox transformation can improve forecast efficiency.

The family of Box-Cox transformations are a key element in regression analysis and a useful tool in research [Atkinson, 1985] including both logarithms and power transformations, and depend on the parameter λ that can take both negative and positive values. Analytic expressions for the minimum mean squared error predictors were provided by Pankratz and Dudley [1987] for specific values of the Box-Cox power transformation parameter.

Simple model structure, normal errors and constant error variance of a distribution are often improved using Box-Cox data transformation, for example, for forecasting volatility according to Higgins and Bera [1992], but opponents of this application, like Sakia [1992] and other researchers, have noted that it does not always manage these challenging goals. Chen and Lee [1997] proposed a Bayesian method to choose the value of λ for a given model structure but opponents present the problem that the model form may depend on the transformation selected.

However, Box and Cox observed that these are all special cases of power transformation and proposed a more flexible method of transformation for researchers to optimise align-

ment with assumptions. An extensive literature review of the Box Cox transformation can be found in Tsiotas [2001] where it is argued that the Box Cox transformation using its properties of constant variance and uncorrelated error increments can give rise to a normal analysis.

Gourieroux and Jasiak [2002] have shown that the autocorrelations (hence the ARIMA model structure) change as function of the nonlinear transformation index could be inappropriate in some cases. However there is a solid theoretical foundation in Exponential smoothing state space models in Hyndman et al. [2002b] and in Hyndman and Khandakar [2008] and with the use of the programming language R, the attempt to improve forecast accuracy was made both worthwhile and easy. More recently, Pascual et al. [2005] have proposed a bootstrap procedure for constructing prediction intervals for a series when an ARIMA model is fitted to its power transformation, while Fernandes and Grammig [2006] applied it in financial time series analysis, like price duration.

According to Shumway and Stoffer [2006] and Hyndman and Athanasopoulos [2013] logarithms are useful because they are interpretable, so changes in a log value are relative (or percentage) changes on the original scale. For example, if we use the log base of 10, an increase of 1 on the log scale of sales corresponds to a multiplication of 10 on the original scale. Additionally, the log transformations constrain the forecasts to stay positive on the original scale.

According to Osborne [2008] few researchers appear to use Box and Cox transformations or report data cleaning of any kind. Proietti and Riani [2009] argue that transformations of seasonal time series are both feasible and relevant, since they can be easily computed and often result in relevant different estimates from those obtained when logarithms or original data are used. According to their research, this is usually the case, when we consider sales, tourism or industrial production, where seasonality is the most prominent source of variation of the data.

Osborne [2010] argues, that given the potential benefits of utilizing transformation (like meeting the assumptions of analyses, improving generalized ability of the results, improving effect sizes) the drawbacks do not seem compelling, in the age of modern computing.

Osborne paper presents an overview of traditional normalizing transformations and how Box-Cox incorporates, extends and improves on traditional approaches to normalizing data, presenting Box-Cox as a potential best practice technique. According to his research Box-Cox not only does easily normalize skewed data¹ but normalizing data also can have a dramatic impact on effect sizes in analyses. Generally, for right-skewed positive skew—data (i.e. tail is on the right) a common transformations may include square root, cube root, and log. On the other hand for left-skewed negative skew—data (i.e. tail is on the left) a common transformations may include square, cube root, and logarithmic.

Additionally according to Goodwin [2010] exponential smoothing state space (ETS) is still today one of the most practically relevant forecasting methods, available even after more than fifty (50) years of widespread use. These models have the ability to adapt to many situations and they are simple in use and transparent in their results. This is the reason this research starts with the use of the data transformations in ETS models and then expands in other time series models.

According to Proietti and Luetkepohl [2011] transformations aim at improving the statistical analysis of time series, by finding a suitable scale for which a model belonging to a simple and well known class has the best performance. Their research focused in assessing whether transforming a variable lead to an improvement in forecasting accuracy. Their empirical evidence shows (in a ratio 1:5) that Box-Cox transformation produces forecasts significantly better than the non transformed data at a one-step-ahead horizon and in most cases the logarithmic transformation is the relevant one. However evidence show that as the forecast horizon increases the evidence in favour of a transformation becomes less strong.

In related work, Lütkepohl and Xu [2011] have investigated whether the logarithmic transformation (as a special case of a power transformation) leads to improved forecasting accuracy over the non-transformed series; the target variables are annual inflation rates computed from seasonally un-adjusted price series. The overall conclusion is that forecasts based on the original variables are characterized by a lower mean square forecast error.

¹If skewness value lies above +1 or below -1, data is highly skewed. If it lies between +0.5 to -0.5, it is moderately skewed. If the value is 0, then the data is symmetric

Furthermore Goncalves and N.Meddahi [2011] used Box-Cox transformation for forecasting volatility with good results .

On the other hand, based on data on a range of monthly stock price indices as well as quarterly consumption series Lütkepohl and Xu [2012] conclude that using logarithmic data transformation can be quite beneficial for forecasting, but can also be damaging for the forecast precision if a stable variance is not achieved. Their paper also points out that there does not appear to be a reliable criterion, for deciding between logs and levels to maximize forecast accuracy.

Bergmeir et al. [2016] have presented a novel method of bagging for exponential smoothing methods using Box Cox transformation, STL decomposition and the moving block bootstrap which is highly recommended to be used for monthly data in practice.

Our main research purpose using Box Cox data transformations is to simplify the patterns in our historical data, by removing variation or by making the pattern more consistent across the whole data set. We know that some values in business are normally distributed, but other variables are log-normally distributed, which means there may be many low values with fewer high values, and even fewer, very high values. There is evidence in the literature that simpler patterns usually lead to more accurate forecasts.

This chapter contributes to the debate in two ways: first, we propose to use the Box Cox transformation in estimating Time series Models of sales data of a turbulent economic period in the Greek market. Our procedure has the advantage that it be compared with the same models that will use no data transformation at all. Our second contribution is to assess the empirical relevance of the choice of the transformation parameter by performing a test whether using Guerrero's method for choosing the best λ is better than using the log transformation or the original seasonal monthly time series. Furthermore, in the previous studies only much more limited data sets were used for Greece and by considering the retail sector of new car sales we hope to get a better overall picture of the sector, and the situation, and may explain some of our previous results. The challenge is to identify if a power transformation may help and which model can give the better forecasts.

In the early chapters of this thesis, we stated our research by using the log transfor-

mation of the original data, which is a relatively strong transformation process. Since our data show monthly variation, as new car sales increase and decrease monthly, we conclude that a data transformation would be useful. Therefore, we decided to use the logarithmic transformation² which was very easy in application and gave empirically very good results. Researchers agree that there is nothing illicit in transforming variables, however we must be careful about how the results from analyses with transformed variables will be reported and analyzed.

6.3 Methodology of Box-Cox Transformations.

The well known Box and Cox [1964] transformation for simultaneously correcting : normality, linearity and homoscedasticity are presented below.

The Box-Cox transformed values are defined as follows:

$$y_t^{(\lambda)} = \begin{cases} \log(x_t) & \text{if } \lambda = 0 \\ \frac{(x_t^\lambda - 1)}{\lambda} & \text{if } \lambda \neq 0. \end{cases} \quad (6.1)$$

What makes a good value of λ is the one that *minimizes* the size of seasonal variation across the whole data set, as that makes the forecasting model simpler.

After the estimation of the one step ahead forecast for the values, the predictions of y_t^λ are considered on its original scale of measurement levels.

The *back-transformed* mean for Box-Cox transformations, is obtained as a naïve procedure simply by reversing Box-Cox transformation [Nelson and Granger, 1979] as follows:

$$\tilde{y}_t^{(\lambda)} = \begin{cases} \exp(y_t^{(\lambda)}) & \text{if } \lambda = 0 \\ (\lambda * y_t^{(\lambda)} + 1)^{1/\lambda} & \text{if } \lambda \neq 0. \end{cases} \quad (6.2)$$

There are methods that can help any researcher to choose the correct λ needed for research. In this study the “forecast” package in **R** is used. This method produces a point estimate of the index λ by minimizes a coefficient of variation for sub-series of the variable.

²We simply denoted the original observations as $x_1, \dots, x_t$ and the transformed observations as $y_1, \dots, y_t$ in a mathematical logarithmic transformation as $y_t = \log(x_t)$.

As already mentioned the Box and Cox family of transformations incorporates many traditionally transformations, like for example:

- $\lambda = 2.00$: square transformation, i.e. x^2
- $\lambda = 1.00$:no transformation needed, produces results identical to original data
- $\lambda = 0.50$: square root transformation,i.e. $\sqrt{x}$
- $\lambda = 0.33$: cube root transformation, i.e. $\sqrt[3]{x}$
- $\lambda = 0.25$: fourth root transformation, i.e. $\sqrt[4]{x}$
- $\lambda = 0.00$: natural log transformation, i.e. $\log(x)$
- $\lambda = -0.50$: reciprocal square root transformation, i.e. $1/\sqrt{x}$
- $\lambda = -1.00$: reciprocal (inverse) transformation, i.e. $1/x$
- $\lambda = -2.00$: reciprocal (inverse) square transformation,i.e. $1/x^2$

There are two methods for calculating λ , the “Guerrero” and “loglik” method³. In the Guerrero [1993] paper the researcher have developed an automatic technique to determine Box and Cox transformation parameter λ and most importantly this procedure is *model independent*, according to Hyndman and Athanasopoulos [2013]. This method is done by minimizing the coefficient of variation of time series. On the other hand in the “loglik” method, the value of λ is chosen to maximize the profile log-likelihood of a linear model fitted to data. According to many researchers, the “Guerrero” method gives good values of λ as compared to the “loglik” method, which helps in better forecasting results. Consequently, this study is going to implement the “Guerero” method in its empirical testing.

According to Guerrero and Parera [2004] underlying that method is the theoretical result that states that the choice of the transformation index is done in such a way that :

$$\frac{[var(x_t)]^{1/2}}{[E(x_t)]^{1-\lambda}} = c$$

³reference is the BoxCox.lambda function in the forecast package

holds valid for all t and some constant $c > 0$ (where x_t is the original series). To use these results it is necessary to estimate both the mean and the variance involved. In this thesis applied research we deal with time series, which means that there is only one observation at each time t , therefore $var(x_t)$ cannot be applied directly. To operationalize the result, we work with the observations grouped into $H \geq 2$ sub-series. This enables the calculations of pairs of samples means and standard deviations, for example, $(\bar{x}_h, S_h)$ for $h=1, \dots, H$ and then search for the λ value that produces:

$$\frac{S_h}{\bar{x}_h^{1-\lambda}} = c \quad (6.3)$$

for $h=1, \dots, H$ for some constant $c > 0$. The elements in this equation are given by $\bar{x}_h = \frac{\sum_{r=1}^R x_{h,r}}{R}$ and $S_h^2 = \frac{\sum_{r=1}^R (x_{h,r} - \bar{x}_h)^2}{R-1}$, where $x_{h,r}$ denotes the r th observation of the subseries h . The subseries $x_{h,1}, \dots, x_{h,r}, \dots, x_{h,R}$, for $h=1, \dots, H$, are formed by grouping R consecutive observations of the original series $x_t : t = 1, \dots, N$, trying to keep homogeneity between the subseries. For this to happen they must be equal-sized. Therefore some number (n) of observations, with $0 \leq n < R$, will have to be left out of the length of the seasonality, if such an effect is present in the series.

The proposed methods stemmed from two empirical interpretations of the equation 6.3. The first one led to minimizing the coefficient of variation of $\frac{S_h}{\bar{x}_h^{1-\lambda}}$ as a function of λ . This method is not linked to a formal statistical model and therefore no assumptions need to be validated to be applied correctly in practice. The second empirical interpretation led to a method based on a simple linear regression in logarithms. The assumption of zero error autocorrelation that underlies this method needs careful attention as it is seldom valid when working with time series. Thus the main method, because of its robustness against violation of assumptions, is the one that minimizes relative variation.

Our empirical research, in predicting new car sales levels by using the Box-Cox data transformation process, goes through the following steps for each one of the data set, firm and model:

Step 1: Transform original data using Box and Cox and different values of λ .

Step 2: Fit Time Series Models (in-sample) and forecasting (out-of-sample).

Step 3: Back-transform the outcome of the forecast values in original values.

Step 4: Estimate forecasting accuracy measures (like RMSE, MAPE, MAE) for all time series models using the back-transformed values.

6.4 Empirical Results of Data Transformation.

In the applied research process of this thesis, the aim is to examine whether applying Box and Cox data transformation gives a better fit and a better forecast for our time series models.

Firstly, the researcher examines the fit of the exponential smoothing state space (ETS) model and for that purpose three (3) groups of exponential smoothing state space (ETS) models are specified, for each one of our four (4) data sets (A, B, C, and D) and for three (3) different firms (Opel, Toyota, and Fiat) in an in-sample estimation process, ending up into thirty-six (36) different models in total which are going to be compared based on each model information criteria results. The three groups are specified as :

- The first model group generates ETS models using the original values with no transformation at all i.e. $\lambda = 1$
- The second model group estimates ETS models using the log values of the data i.e. $\lambda = 0$
- The third model group calculates ETS models using Box-Cox data transformation where λ values will be estimated using Guerrero's method (1993).

For the purpose of this study, various information criteria are calculated, in order to compare the fit of the models like the Akaike (AIC) and Schwarz Bayesian (BIC) and the corrected Akaike (AICc). Thus the ETS model that minimizes those information criteria is chosen to have the best fit for the data of the whole data sample.

In addition, this research study continues with an out-of-sample estimation, by examining the forecast accuracy of the ETS models when using Box and Cox data transformation, original and log values. So each ETS model forecasting results will be compared with the

actual values in an out-of-sample estimation process, which means that for each data set 80% of the data is selected as a training set to fit the best ETS model and the forecasting results are compared with the remaining 20% actual data of the test set. Furthermore, for comparison reasons, in the case data were transformed before fitting the model, then one more step is needed: the back-transformation of the estimated forecast values of the ETS models. After that process, the forecast values can be compared with the actual ones. Accuracy measures, like root mean square error (RMSE), mean absolute percentage error (MAPE), mean absolute error (MAE) are calculated for comparison reasons. For each data period, the model that minimizes the accuracy measures give the best forecast values for new car sales levels and it is interesting to see if the three groups of transformation come up with the same result in this study.

6.4.1 ETS models and Box and Cox transformation (in-sample).

Since exponential smoothing state space (ETS) models are proven to be one of the best models in forecasting this thesis various variables, this research continues with the in-sample estimation of ETS models for 4 different data sets, using Box and Cox transformation, for three (3) different firms of new car sales i.e. Opel, Toyota, and Fiat. For comparison reasons the Akaike⁴(AIC) , the Schwarz Bayesian⁵(BIC) and the corrected Akaike⁶ (AICc) information criteria are estimated. After comparing the estimated models the best one is selected, according to Hyndman et al. [2002a], as the one that minimizes the information criteria.

The ETS methods used in this research are algorithms that return point forecasts. It is also possible to use the “Innovation state-space models” that generate the same point forecast but additionally can generate forecast intervals. That is a stochastic (or random) data generating process that can generate an entire forecast distribution and can allow for “proper” model selection according to Hyndman et al. [2002b]. Each model has an

⁴ $AIC = -2\log(L) + 2k$

⁵ $BIC = AIC + k(\log(T) - 2)$

⁶ $AICc = AIC + \frac{2(k+1)(k+2)}{T-k}$

observation equation and a transition equation, one for each state (level, trend, seasonal) i.e. state-space models. The taxonomy of exponential smoothing method according to Hyndman et al. [2005] is used, which is based on characterizing each model against three dimensions:

E underlined error model: additive (A), multiplicative (M)

T type of trend : none (N), additive (A), damped (A_d)

S type of seasonal: none (N), additive (A), multiplicative (M)

Hence the model ETS refers to error (E), trend (T), and seasonality (S), and these three components can either be additive (A), multiplicative (M), inexistent/none (N) or dampened additively (A_d). Therefore when specifying an ETS model, for example, ETS(A, N, A) then the first letter denotes the error type, the second letter denotes the trend type and the third denotes the season type.

In Table 6.1 (page 181) it shows that Opel is in favor of the ETS models with additive errors and additive seasonality while there is no evidence of a clear trend. So the ETS (A, N, A) is dominant in the in-sample estimation of the ETS model in data set B, C, D while only in the data set A the models have a dampened additive trend. Furthermore, the model that takes the log values seems to give a better fit of the ETS model for data set A, C, D while the Box-Cox transformation with the Guerrero method gives the best results in the case of Data set B.

In Table 6.2 (page 182) there is evidence that Toyota prefers the ETS(A,N,A) models with additive errors, additive seasonality and no trend in the in-sample estimation. The only differentiation is in Data Set A where trend sometimes is additive or damped additive. From the empirical evidence, the log transformation of the data in data set B gives the best results while for the data set A, C, and D the Box-Cox transformation with the Guerrero method is preferred. In Set C the Information criteria are all negative which can happen.

In Table 6.3 (page 183) there is evidence that Fiat is in favor of the ETS(A, N, A) models i.e. additive errors, additive seasonality and no trend with only one exception in

Table 6.1: Information Criteria for OPEL with Box-Cox & ETS models (in-sample)

OPEL	ETS	AIC	AICc	BIC
Data Set A:1998-2016				
$\lambda = 1$ Actual	(A, A_d ,A)	3917	3920	3079
$\lambda = 0$ Logs	(A, A_d ,A)	646	649	707
$\lambda = 0, 18$ Box Cox	(A, A_d ,A)	1228	1231	1290
Data Set B:2006-2015				
$\lambda = 1$ Actual	(A,N,A)	1986	1990	2028
$\lambda = 0$ Logs	(A,N,A)	282	287	324
$\lambda = -0, 10$ Box Cox	(A,N,A)	116	121	158
Data Set C:2006-2010				
$\lambda = 1$ Actual	(A,N,A)	968	978	999
$\lambda = 0$ Logs	(A,N,A)	87	98	118
$\lambda = 0, 59$ Box Cox	(A,N,A)	610	621	641
Data Set D:2002-2011				
$\lambda = 1$ Actual	(A,N,A)	1985	1990	2027
$\lambda = 0$ Logs	(A,N,A)	231	236	273
$\lambda = 0, 42$ Box Cox	(A,N,A)	963	968	1005

the log transformation of Data Set A where additive trend is noticed. It is also important to stress that the log transformation is preferred in all data sets in the case of Fiat in sample estimations with the Box-Cox, Guerrero method as the second-best choice for our models.

In general, the empirical results as illustrated in Table 6.1, Table 6.2 and Table 6.3, in page 181, 182 and 183 accordingly, shows that the estimated ETS model specification with the log values of the series or the Box-Cox transformed data using the Guerrero method are the two best competing methods for the three different car representatives in the four different data sets in an in-sample estimation. On the other hand, the estimated ETS model specification using the original data i.e $\lambda = 1$, which means with no transformation at all, was always the last case scenario. That comes as a result of our empirical study since these are the cases where the estimated values of the information criteria are minimized. Additionally, the value of the best λ according to Box-Cox and Guerrero [1993] method

Table 6.2: Information Criteria for TOYOTA with Box-Cox & ETS models (in-sample)

TOYOTA	ETS Model	AIC	AICc	BIC
Data Set A:1998-2016				
$\lambda = 1$ Actual	(A,Ad,A)	4089	4091	4140
$\lambda = 0$ Logs	(A,A,A)	776	778	834
$\lambda = 0, 18$ Box Cox	(A,N,A)	486	489	538
Data Set B:2006-2015				
$\lambda = 1$ Actual	(A,N,A)	2053	2057	2094
$\lambda = 0$ Logs	(A,N,A)	265	270	307
$\lambda = -0, 10$ Box Cox	(A,N,A)	2291	2296	2333
Data Set C:2006-2010				
$\lambda = 1$ Actual	(A,N,A)	972	983	1003
$\lambda = 0$ Logs	(A,N,A)	57	68	88
$\lambda = 0, 59$ Box Cox	(A,N,A)	-829	-818	-798
Data Set D:2002-2011				
$\lambda = 1$ Actual	(A,N,A)	2053	2058	2095
$\lambda = 0$ Logs	(A,N,A)	287	292	329
$\lambda = 0, 42$ Box Cox	(A,N,A)	30	35	72

for the data transformation differs according to the size of the selected sample in research. Lastly, the specification of the ETS model changes as data sets alter and also differs if λ changes (i.e. data are transformed or not).

The empirical research, in this part of our study, concludes that in general the ETS models using log values transformation and the Box-Cox transformation with Guerreros's method in calculating the λ seems both to give better in-sample fit to our data than in the case of using the original data with no transformation at all. Hence transforming the time series data used in research always give better empirical results in model specification.

6.4.2 ETS models and Box and Cox transformation (out of-sample).

This research is interesting not only for the model that best fits the data in research but also for the model that gives the best forecasts. Therefore the study proceeds next to

Table 6.3: Information Criteria for FIAT with Box-Cox & ETS models (in-sample)

FIAT	ETS Model	AIC	AICc	BIC
Data Set A:1998-2016				
$\lambda = 1$ Actual	(A,N,A)	3822	3825	3874
$\lambda = 0$ Logs	(A,A,A)	672	675	730
$\lambda = 0, 14$ Box Cox	(A,N,A)	1093	1095	1144
Data Set B:2006-2015				
$\lambda = 1$ Actual	(A,N,A)	1864	1868	1905
$\lambda = 0$ Logs	(A,N,A)	309	314	351
$\lambda = -0, 02$ Box Cox	(A,N,A)	346	351	388
Data Set C:2006-2010				
$\lambda = 1$ Actual	(A,N,A)	918	929	949
$\lambda = 0$ Logs	(A,N,A)	96	107	128
$\lambda = 1, 38$ Box Cox	(A,N,A)	1243	1254	1254
Data Set D:2002-2011				
$\lambda = 1$ Actual	(A,N,A)	1918	1923	1960
$\lambda = 0$ Logs	(A,N,A)	232	237	274
$\lambda = 0, 26$ Box Cox	(A,N,A)	671	675	712

see if ETS models and data transformations can lead to an improvement of the forecasting performance.

Firstly, the ETS models are estimated for the actual values, the log values, and the transformed values using Box-Cox with the Guerrero method for all data sets A, B, C, D, and for Opel, Toyota and Fiat in an out-of-sample estimation. Each model is estimated using 80% of the data in each set and keeping the remaining 20% of the observations as the test set, which are needed for measuring the forecasting accuracy of each model. More specific the ETS models using 80% of the Opel observations of each data set are used (transformed or not) to give point forecasts for “h” next periods, then these estimations, if transformed data were used, are back-transformed into original values otherwise they are used as estimated and they are compared with the actual values using various forecasting performance measures.

The forecasting accuracy measures used for this purpose are the Root Mean Square

Error (RMSE), the Mean Absolute Error (MAE), and the Mean Absolute Percentage Error (MAPE). The same process is used in the case of the original values with no transformation at all but in that case, there is no need for the step of the back transformation of the forecast values before calculating the forecast accuracy measures.

Table 6.4: Forecasting Performance of OPEL using Box-Cox & ETS models (out-of-sample)

OPEL	ETS Models	RMSE	MAE	MAPE
A^{train}:1998-2012, A^{test}:2013-2016				
$\lambda = 1$ Actual	(A, A_d ,A)	614	533	113
$\lambda = 0$ Logs	(A, A_d ,A)	314	273	34
$\lambda = 0, 36$ Box Cox	(A,N,A)	521	466	102
B^{train}:2006-2013, B^{test}:2014-2015				
$\lambda = 1$ Actual	(A,N,A)	157	125	27
$\lambda = 0$ Logs	(A,N,A)	147	117	22
$\lambda = 0, 31$ Box Cox	(A,N,A)	299	244	55
C^{train}:2006-2009, C^{test}:2010-2010				
$\lambda = 1$ Actual	(A,N,A)	467	392	47
$\lambda = 0$ Logs	(A,N,A)	471	393	28
$\lambda = 0, 77$ Box Cox	(A,N,A)	558	468	56
D^{train}:2002-2009, D^{test}:2010-2011				
$\lambda = 1$ Actual	(A,N,A)	546	472	55
$\lambda = 0$ Logs	(A,N,A)	436	392	29
$\lambda = 0, 86$ Box Cox	(A,N,A)	539	471	55

In Table 6.4 (page 184) empirical results show the estimated ETS models for Opel with the original values the log values and the Box-Cox transformed data using the Guerrero method for the three different car representatives in the four different data sets in an in-sample estimation. It is worth noticing that the best value of λ , estimated with Guerrero's method [Guerrero, 1993] in an out-of-sample estimation is different from what was estimated in the in-sample estimation for each one of the data set (see for comparison Table 6.1, page 184). This is expected due to the different size of samples, even if the Guerrero method is model-independent. For example in data set A in Table 6.1 the monthly

observations from 1998 till 2016 (i.e. 228 observations) are considered to fit the model in an in-sample estimation while in Table 6.4 the same data set is studied but we use fewer observations to fit the model in an out-of-sample estimation considering observations from 1998 till 2012 (i.e. 180 observations). The last 48 observations are left in the test set, which is used to estimate prediction error for the ETS model selected.

Additionally it is interesting that the Exponential Smoothing state-space models specifications do not change in both tables 6.1 and 6.4. The model that this algorithm is chosen is the same ETS for each data set for each firm in the in-sample and out of sample estimation. ETS(A, N, A) is preferred in all cases except the case of actual and log values where ETS (A, A_d , A) is chosen as the best one. The same specifications are valid also for tables 6.2 and 6.5 for Toyota and tables 6.3 and 6.6 for Fiat for their in-sample and out-of-sample estimation. The accuracy measures of Root Mean Square Error (RMSE), Mean Absolute Error (MAE), and Mean Absolute Percentage Error (MAPE) are calculated for comparison reasons of the out-of sample forecasting.

Our empirical research, in this part of our study, conclude that for Opel, Toyota and Fiat, the ETS models using log values or Box-Cox transformed values with Guerreros's method in calculating the best λ value, seems to give better out-of-sample forecast accuracy measures in forecasting against the case of using the original data, with no transformation at all. However, both cases are all beaten from the log transformation case study, which seems to give the best forecasts for all data sets according to all the accuracy measures.

Table 6.5: Forecasting Performance of TOYOTA with Box-Cox & ETS models (out-of-sample)

TOYOTA	ETS Model	RMSE	MAE	MAPE
A^{train} :1998-2012, A^{test} :2013-2016				
$\lambda = 1$ Actual	(A, A_d ,A)	742	628	76
$\lambda = 0$ Logs	(A,N,A)	785	667	40
$\lambda = -0,14$ Box Cox	(A,N,A)	1638	1336	163
B^{train} :2006-2013, B^{test} :2014-2015				
$\lambda = 1$ Actual	(A,N,A)	933	856	95
$\lambda = 0$ Logs	(A,N,A)	533	335	20
$\lambda = 0,51$ Box Cox	(A,N,A)	840	803	87
C^{train} :2006-2009, C^{test} :2010-2010				
$\lambda = 1$ Actual	(A,N,A)	403	308	13
$\lambda = 0$ Logs	(A,N,A)	285	229	11
$\lambda = -0,99$ Box Cox	(A,N,N)	283	234	0,10
D^{train} :2002-2009, D^{test} :2010-2011				
$\lambda = 1$ Actual	(A,N,A)	865	793	35
$\lambda = 0$ Logs	(A,N,A)	948	854	60
$\lambda = -0,18$ Box Cox	(A,N,A)	648	534	24

In the case of Toyota the results are not so general for all data sets (see Table 6.5, page 186). The log transformation is preferred for all data sets according to MAPE metrics, while Box-Cox transformation is preferred in data set C and D according to RMSE metrics, leaving data set A with the best forecasting results when the original value are used, according to MAE and RMSE metrics. The ETS(A, N, A) model is in favor in most of the cases and data sets with two exceptions: the first one is in Data Set A when original values are used where ETS(A, A_d , A) is chosen, and in data set C in the case where the Box-Cox transformation is used where ETS(A, N, N) is chosen. However all model specifications results come in accordance with the in-sample estimation of Toyota (see Table 6.2 page 182).

Table 6.6: Forecasting Performance of FIAT with Box-Cox & ETS models (out-of-sample)

FIAT	ETS Model	RMSE	MAE	MAPE
A^{train} :1998-2012, A^{test} :2013-2016				
$\lambda = 1$ Actual	(A,N,A)	128	97	31
$\lambda = 0$ Logs	(A,N,A)	125	96	30
$\lambda = 0, 10$ Box Cox	(A,N,A)	123	92	37
B^{train} :2006-2013, B^{test} :2014-2015				
$\lambda = 1$ Actual	(A,N,A)	121	96	42
$\lambda = 0$ Logs	(A,N,A)	122	94	28
$\lambda = -0, 08$ Box Cox	(A,N,A)	135	100	48
C^{train} :2006-2009, C^{test} :2010-2010				
$\lambda = 1$ Actual	(A,N,A)	391	348	68
$\lambda = 0$ Logs	(A,N,A)	382	354	38
$\lambda = 0, 63$ Box Cox	(A,N,A)	413	376	77
D^{train} :2002-2009, D^{test} :2010-2011				
$\lambda = 1$ Actual	(A,N,A)	477	435	95
$\lambda = 0$ Logs	(A,N,A)	441	407	42
$\lambda = 0, 07$ Box Cox	(A,N,A)	509	466	101

In the case of Fiat the results are interesting (see Table 6.6, page 187). All data sets and cases in the out of sample estimation are in favor of the ETS(A,N,A) just like in the equivalent cases of the in-sample estimation (see Table 6.3). Data A and B give better forecasts when data are transformed with Box-Cox and Guerrero method according to MAE metric, while data C and D give better forecasting results when log data are used, according to MAPE and RMSE metrics. In general, it seems that more than half of the empirical results in the out of sample estimations give preference to the ETS model with transformed data, with logs or with Box and Cox transformations with the Guerrero method.

Graphical Presentation of ETS models.

Additionally, the forecasting performance of the ETS models can also be presented graphically. The ETS model's forecast values of Opel, Toyota, and Fiat new car sales levels, for data set D, using Box-Cox transformation, log values, and original data, are

visually presented in comparison with the actual values of new car sales levels in a forecast time horizon of two years (i.e. $h=24$ months).

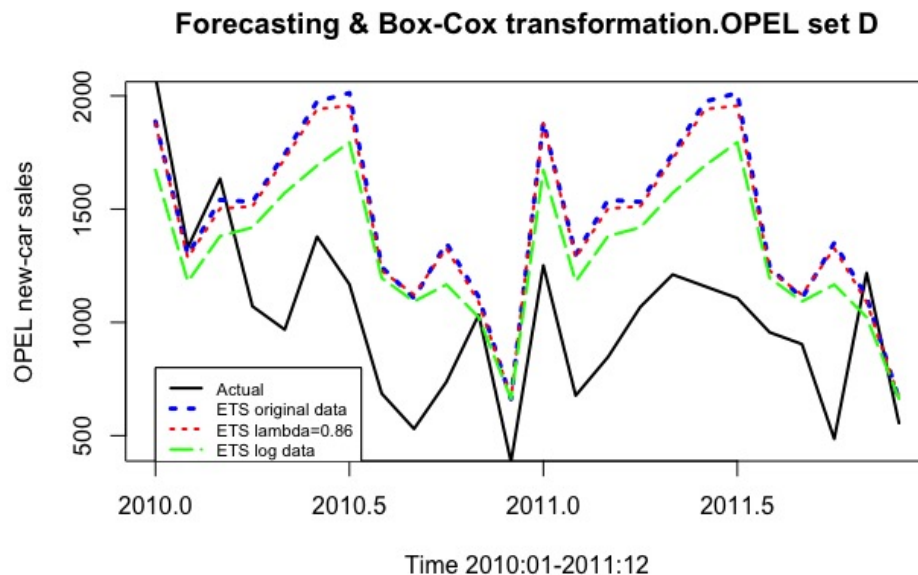


Figure 6.1: Forecasting Opel new-car sales using ETS and data transformation

In Figure 6.1 (page 188) the ETS (A, N, A) models for Opel, in data set D are presented for the original data the log data and the Box-Cox transformed data. Evidence shows that ETS with Box-Cox transformation moves almost the same as the ETS model that uses the actual values and only the ETS that uses the log values appears to slightly give forecasts closer to the line of the true sales values.

So the graphical presentation confirms the results from Table 6.4 (page 184) in data set D which gives ETS with log values as the best model in forecasting performance despite the very small difference between the models. It is also noticeable that all the models overestimate the actual sales level of the series and seem to be very close to the actual sales in the short-run i.e. the 1st quarter of the year.

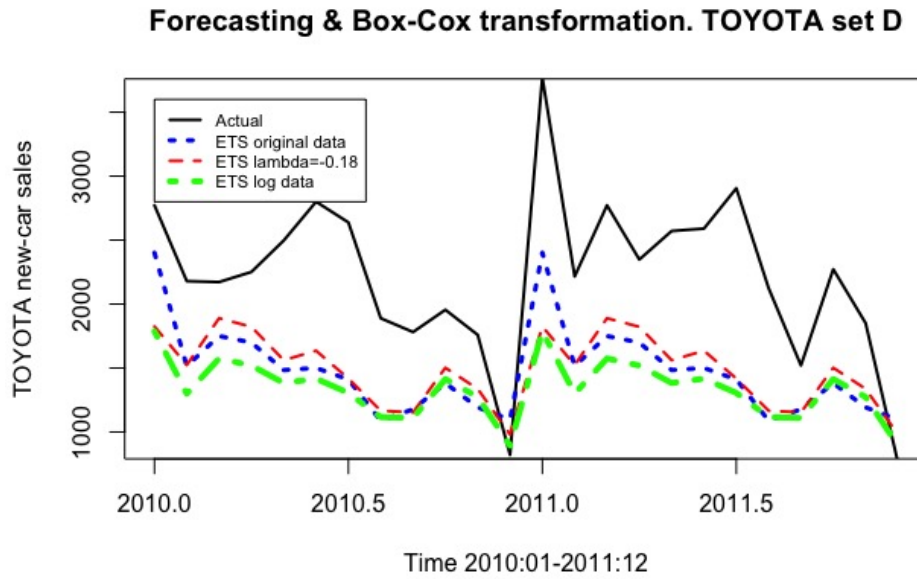


Figure 6.2: Forecasting Toyota new-car sales using ETS and Data transformation.

In Figure 6.2 (page 189) the ETS (A, N, A) models for Toyota in data set D are presented for the original data the log data and the Box-Cox transformed data. In this case, the ETS model with Box-Cox transformed data with $\lambda = -0,18$ gives the best forecasting results, then the ETS estimated based on the actual values and last the ETS with the log values. The graphical presentation confirms the results of forecast performance given by the Table 6.5 (page 186) for data set D and Toyota.

However all the ETS model has the same model type (A, N, A) and the differences are small but they do exist and that is visually observed in the graphical presentation of the forecasting performance of the models. Furthermore, all ETS models under-estimate the actual level of the sales series and never really capture the sharp changes in the sales level.

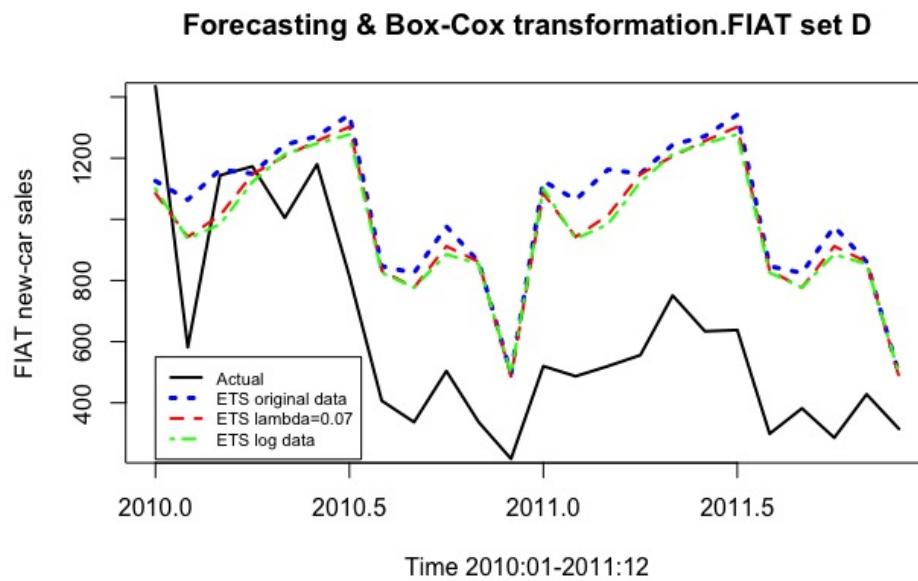


Figure 6.3: Forecasting Fiat new-car sales using ETS and data transformation.

In Figure 6.3 (page 190) the ETS (A,N,A) models for Fiat in data set D are presented for original data, log data and data transformed with Box-Cox. In this data set, according to the forecast performance in Table 6.6 (page 187), the ETS estimated with the actual values is the best one then the ETS with the Box-Cox transformation is the second best one and last the ETS with the log values. However, differences in the graphical presentation are small and not very clear. There are only some differentiations when sharp changes happen in the series. Additionally, all models overestimate the actual level of the sale series and as time horizon becomes bigger the values range between the actual and the forecast values become bigger as well.

As a consequence, it is generally observed that:

- ETS models can fit very well with the variations of the actual sales level.
- ETS models can give very good forecasts when data are transformed either with log values or using the Box-Cox transformation with λ chosen by Guerrero method.
- Each data set and time series react differently so research much be done to find the best transformation of the series

- When sharp changes are recorded in real data due to economic reasons, ETS models are very good in following the changes, hence forecast become very efficient due to the ability of the ETS models to follow the data variations and variability.

These results are well in line with ETS model characteristics as this model can capture very well trend and seasonality of our seasonal data. In general ETS models usually give better forecasting results when transformed data are used in model estimation either using Box-Cox or log values. So in most of the cases, it is better to fit an ETS model after transforming data using Box and Cox, and Guerrero method but the researcher should also study log values since they might give very good results. It is hard to decide which method is the best one since different time series give different results that do not agree with each other.

6.4.3 Time Series forecasting & transformation (out-of-sample).

Our research continues with an out-of-sample estimation using the Opel, Toyota and Fiat, new car sales levels for all data sets, in original values ($\lambda = 1$) and when using two types of Box-Cox transformation, the log values ($\lambda = 0$) and the case where λ is chosen based in the Guerrero's method ($\lambda = 0,86$). So the research estimates three different cases for each of our data and for various types of time series models [Mean, Naïve, seasonal Naïve, Linear model with seasonal Dummies, Exponential smoothing state space models-ETS, and SARIMA]. The models were calculated, the forecast values were estimated and then back-transformed in original values.

Furthermore the accuracy measures were evaluated, like the root mean squared error (RMSE), the mean absolute percentage error (MAPE) and the mean absolute error (MAE) for comparison reasons. The model type that minimizes the accuracy measures, gives the best forecast values for new car sales levels of the firm.

We notice that the accuracy measures in original values and log values for Naïve and Seasonal Naïve models give the same results (see Tables 6.7-Table 6.18, pages 209-213). That makes sense, since the two models basically and give the same forecast values, if we

take the original or the log values of the series. The fact that, we back transform the log forecast values of Naïve and seasonal Naïve models in original values, before calculating the accuracy measures, results in the calculation of the same values in all accuracy metrics for original and log data. However this is not valid for other models. For example the Mean/Average model, which is a basic model as well, but gives different forecast mean values for original and log series and so on.

The empirical research for Opel in Table 6.10 (page 205) shows that no model can capture the turbulence movement of the original series. In data set D and Opel new car sales, if the original data are used to estimate the time series models then the ETS - Exponential state-space model seems to forecast better and the comes the Seasonal Naïve model as a second choice. The turbulence in the scale of sales is so vast that the Seasonal Naïve seems to be a good forecasting model choice in Opel case using the original data with no transformation at all.

However when data are transformed either by taking their log values or using the Box-Cox and Guerrero method the results stay the same with ETS model as a first choice and the Seasonal Naïve model as a second best model in forecasting ability. On the other hand, if we had to choose which one of the cases: original, log or Box-Cox transformed data is the best, the log transformed data is giving the best results and the minimum values in all metrics of forecast performance.

Graphical Presentation of new car sales forecasts.

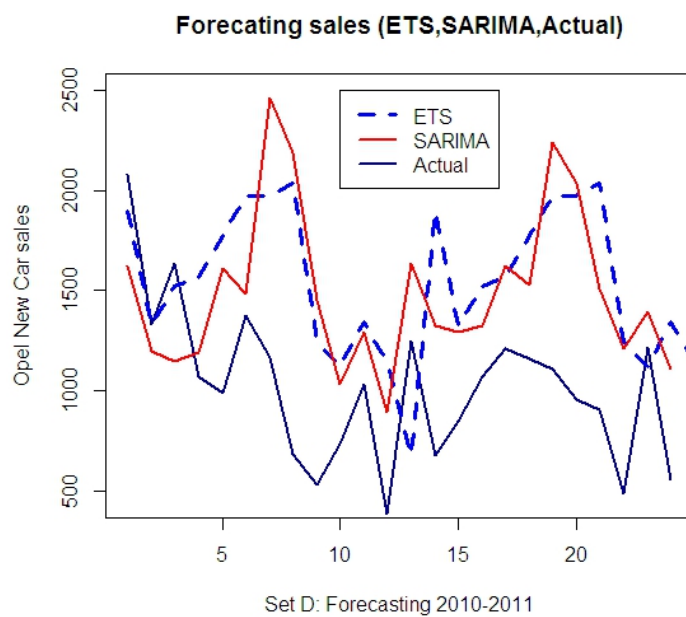
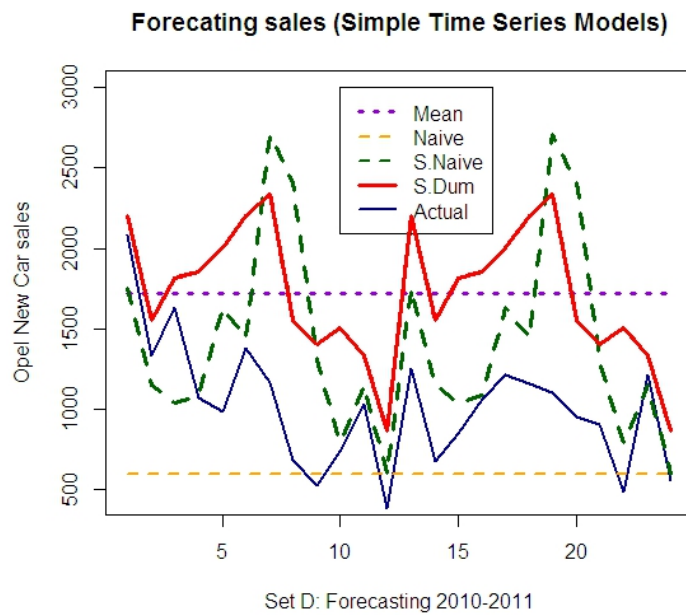


Figure 6.4: Opel new car sales forecasts.Set D-Actual

Furthermore, the results of the empirical research that is presented algebraically in Table 6.10 can also be presented graphically. In Figure 6.4 (page 193) each time series models' forecast values are illustrated, for $h=24$ months, using original data in the estimating

process for Data Set D, and then in Figure 6.5 (page 195) the forecast levels of Opel new car sales are all illustrated when Opel new car sales data are transformed with Box-Cox and Guerrero method in data set D. Both figures are divided into two different graphs that represent two groups of forecast values using various time series models and each time the forecast values line is compared with the actual line of the sales level.

More specific in the first graph of Figure 6.4 a group of simple Time Series Models is presented i.e. the Mean, the Naïve, the Seasonal Naïve, the Linear Model with Seasonal Dummies LMSD forecast values, along with the original values of Opel new car sales. It is easily noticed that in the short-run ($h < 6$ months) the linear model with seasonal dummies and the seasonal Naïve model gives very good forecasts for sales levels but in the long run they both fail.

Furthermore the second graph in Figure 6.4, present graphically the forecast values of ETS and SARIMA time series models for data set D using original values for a time horizon of $h=24$. Optically it is noticed that in the short-run (at least 3 months ahead) the ETS model gives very good forecasts for sales levels while in the long run, it overestimates the level of sales. On the other hand, the SARIMA model starts with underestimating the level of sales in the short-run (4 months ahead), and then it continues with overestimating them. Another interesting point between ETS and SARIMA model in this study is that in our research there are periods where the ETS model corresponds more quickly in sales levels fluctuations and consequently gives better forecasts.

In Figure 6.5 (page 195) the forecast levels of Opel new car sales are presented when the Box-Cox data transformation and the Guerrero method to choose the best λ , are used in data set D. According to the results given by Table 6.10 (page 205) the best time series model for forecasting, when data are transformed using Box-Cox and Guerrero's method, is the Linear Model with Seasonal Dummies (LMSD) while the Exponential smoothing state-space model (ETS) comes as a second-best choice. Both models are presented in the first part of Figure 6.5 and it seems that they are quite accurate in the short-run but in the long run they overestimate the level of sales but manage somehow to follow the fluctuations during the forecast period. On the other hand, the second half of the figure

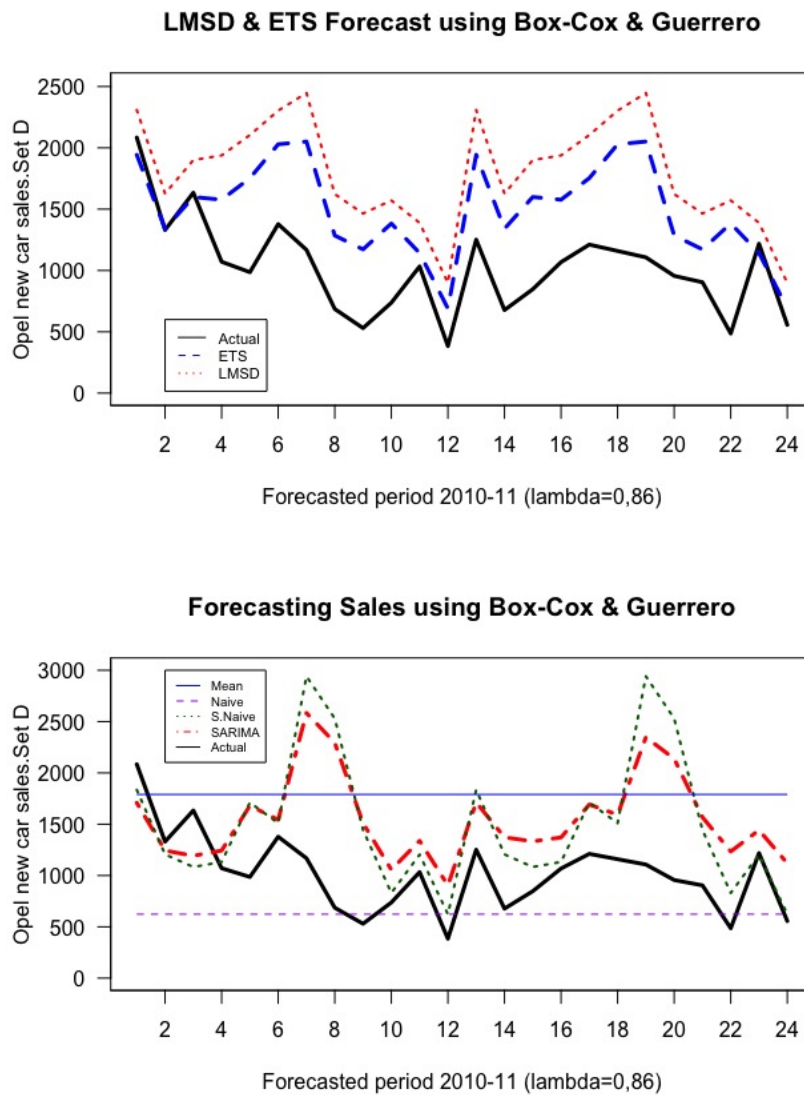


Figure 6.5: Opel new car sales forecasts. Set D,Box-Cox/Guerrero

shows the forecast values of sales when using the mean and naive method Seasonal naive, and SARIMA models with transformed data. The simple mean and naive models are both out of consideration since they give flat results that do not come close to reality. However seasonal naive and SARIMA models seem to follow the movement of the actual sales level with the SARIMA models to perform better forecast than the Seasonal Naive models. Due to the seasonality of our data, it is very important for the model used to take seasonality as a major characteristic of the series. However, both SARIMA and Seasonal Naive models seem to overestimate the level of sales and exaggerate in a slight increase in sales during

specific periods.

6.5 Confidence Interval.

A confidence interval is how much uncertainty there is with any particular forecasting method. Confidence intervals (CI) are often used in percentages (for example, a 95% or 90% confidence level) because they include an error margin. This means that if the study could be repeated over and over again, 95 or 90 percent of the time the outcome results will match the results of the entire population. Actually, the CI tells us how confident one is that the results from our empirical study reflect what one would expect to find if it was possible to test the entire market i.e. the population.

Confidence intervals are the results one gets that are intrinsically connected to confidence levels. In this thesis research work, the interest is to see whether the forecasts from the best time series models are within the 95% confidence interval and which model performs best and within these limits. Therefore a graphical presentation for the original data, using the two best time series models to forecast values for $h=24$ next months, in a 95% confidence interval is presented.

The forecast values of Opel using the original data for data set D and the model, that gives the best forecasts according to this research, which is the Exponential smoothing state space (ETS) model is presented in Figure 6.6 (page 197). This ETS model has the best forecasting accuracy, for $h=24$, and is presented along with the actual sales values at a 95% confidence interval. The researcher noticed that, when the original values in estimating an ETS model are used, the forecast values are always within the 95% confidence interval (CI). Additionally, they follow the flow of the variations of the actual value during the time and at the same pace, for the short-run ($h \leq 3$). However, they keep a slightly higher level of forecasts in the long-run. In this graph, it is remarkable how the actual values stay within the confidence interval of the forecast values that the ETS model produces for all the forecast period ($h=24$).

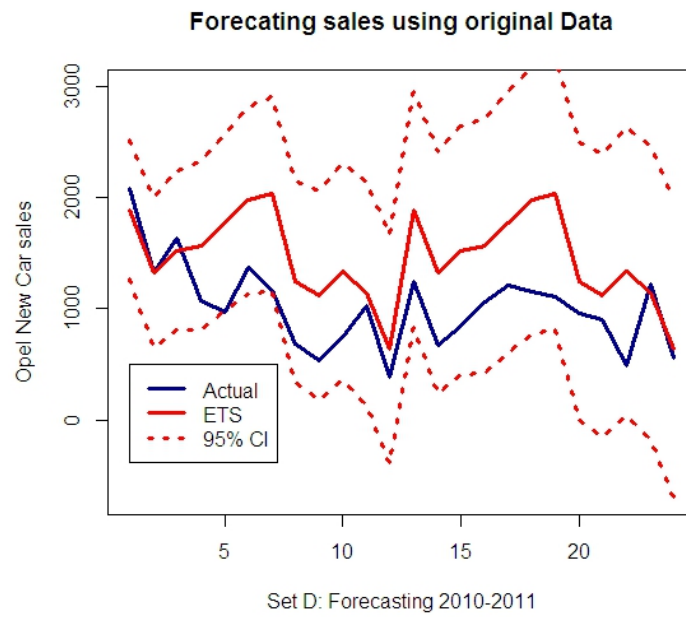


Figure 6.6: ETS forecasts & 95%CI for Opel (Set D-Original)

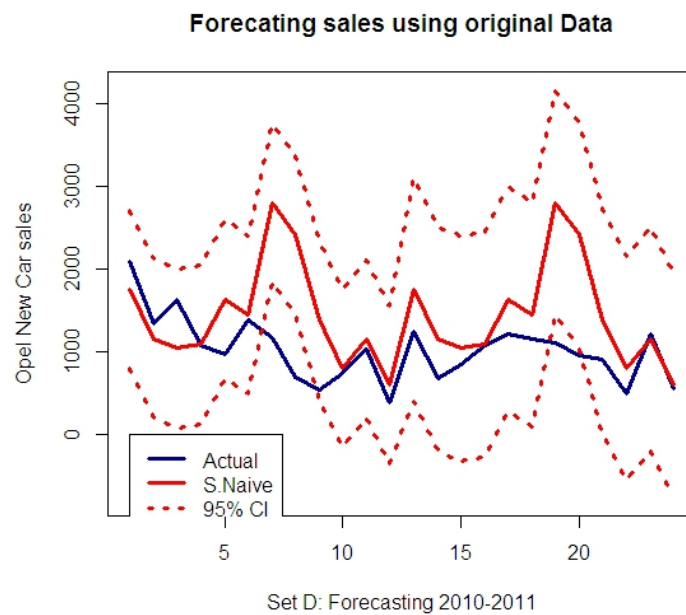


Figure 6.7: Seasonal Naïve forecasts & 95%CI for Opel (Set D-Original)

Furthermore, the Seasonal Naïve model for Opel in data set D using the original values, which is the second-best forecast model is presented in Figure 6.7 (page 197) for $h=24$ with

a 95% confidence interval, along with the actual sales values. The researcher can notice that when using the Seasonal Naïve model, the forecast values underestimate the actual values in the short-run ($h \leq 4$) while in the long run, the forecast values overestimate the actual values and in many cases, there are big variations from the actual data. It is most noticeable that the actual values do not always stay within the 95% confidence interval (CI) and while the actual values decrease and drop below an estimated 95% CI the forecast values on the contrary have a sharp increase (like a shock).

Additionally in Figure 6.8 (page 199) the forecast values of the Linear Model with Seasonal Dummies is illustrated, which is the best model when using Box-Cox transformation and Guerrero's method for choosing the best λ value ($\lambda = 0,86$) and also its 95% CI and the actual values for comparison reasons. There is a very good short-run forecast ($h \leq 4$ months) but in the long run, the model seems to overestimate the actual sales levels. It is worth mentioning that the actual sales are not always within the 95% CI of this model especially when changes in the sales level are sharp.

Furthermore, the Exponential state-space smoothing model (ETS), is presented in Figure 6.9 (page 199), is the best forecasting model, when using Box-Cox transformation and Guerrero's method for choosing the best λ value. It appears that the model is very good in the short-run, but overestimates long-term forecasts. The forecast values and the 95% CI do not always capture the actual sales level. Hence the graph gives evidence of how turbulent the level of the new car sales is during that period, affected by the unstable economic Greek market. In general, the out of sample estimation and forecasting led to sensible results and it was relatively easy to apply the time series models.

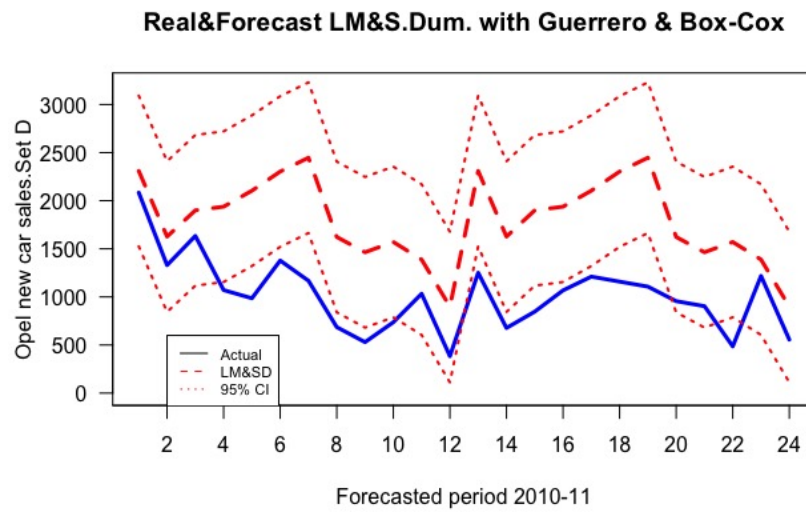


Figure 6.8: LMSD forecasts & 95%CI (Opel-Set D-BoxCox/Guerrero).

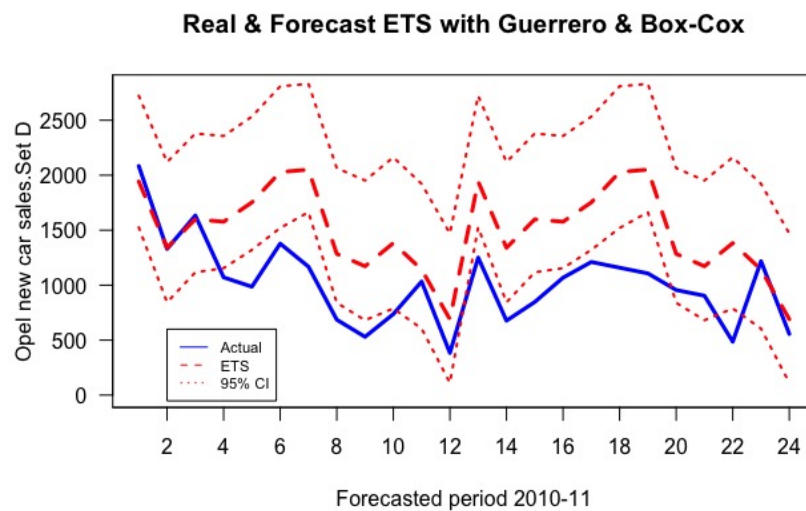


Figure 6.9: ETS Model forecast & 95%CI for Opel (Set D-BoxCox/Guerrero)

6.6 Discussion.

This thesis empirical research, concludes that in general the Exponential smoothing state space (ETS) models using data transformations give quite good (in-sample) fit to the data, and are more accurate in forecasting (out-of-sample) than the ETS models using

new car sales data (see Figure 6.11, page 6.10). The ETS models perform very good in forecasting this thesis's time series new car sales levels and become even better when the data are transformed in log values or use Box-Cox and Guerreros's method in calculating the best λ value. Therefore the mathematical transformations in data can be trusted and used because they give good results.

Furthermore in the short-run (≤ 6 months) ETS models forecast values can capture the sales levels movement but in the long-run, they have a tendency to overestimate sales level. However, the researcher can be 95% confident that it will give a good forecast of the new car sales as forecast values and real values are all in the range of the 95% confident interval zone. In some cases, the 95% confidence interval (CI) can not capture the real movement of the series especially when changes in the sales level are sharp and there is evidence of a quite turbulent economic environment at that period of time.

In examining the three cases of data (actual, log, BC-Guerrero transformed) in the four different data sets and the three different firms we can conclude that the log transformation of the time series is giving the majority of empirical results with the minimum forecasting accuracy metrics, and should be preferred (see Figure 6.10, page 201). There are of course some exceptions that lead more towards the Box-Cox transformation with Guerrero case study. So examining the variable with transformations was worthwhile in this chapter, cause we can indicate that logarithm and Box-cox transformation are useful, easy to implement and give interpretable results that can benefit the decisions makers in a firm.

Finally, it is hard to be sure that one single type of time series model or one specific type of data transformation is the best one for all-time series in general and can capture the movement of the series or make the best forecasts. Therefore the researcher must carefully consider the nature of the series and the empirical results obtained here must be interpreted with attention. This may lead to further research in the area of time series models and data transformation and more research in the retail sector of the Greek market place in general.

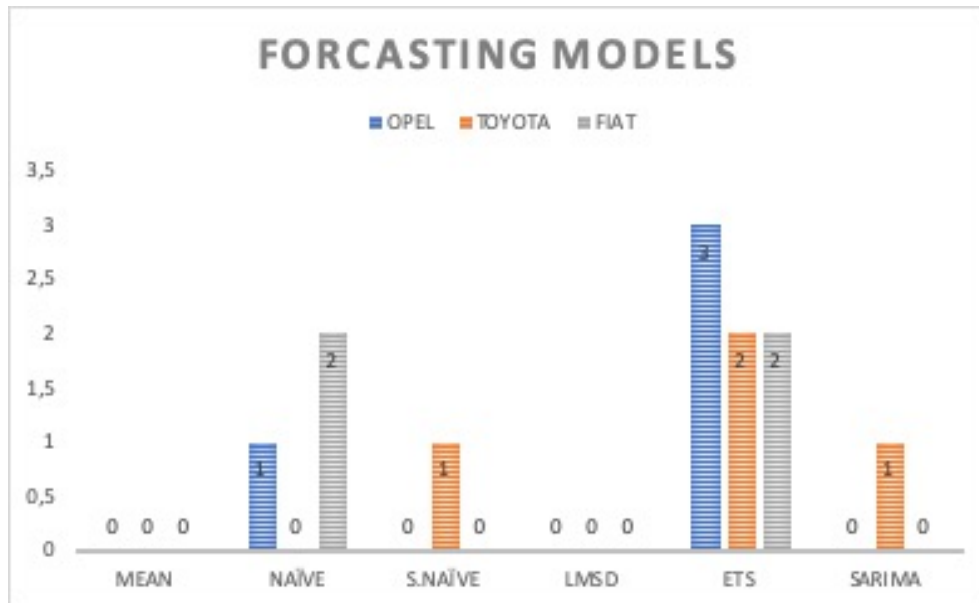


Figure 6.10: Forecasting new-car sales and time series model selection.

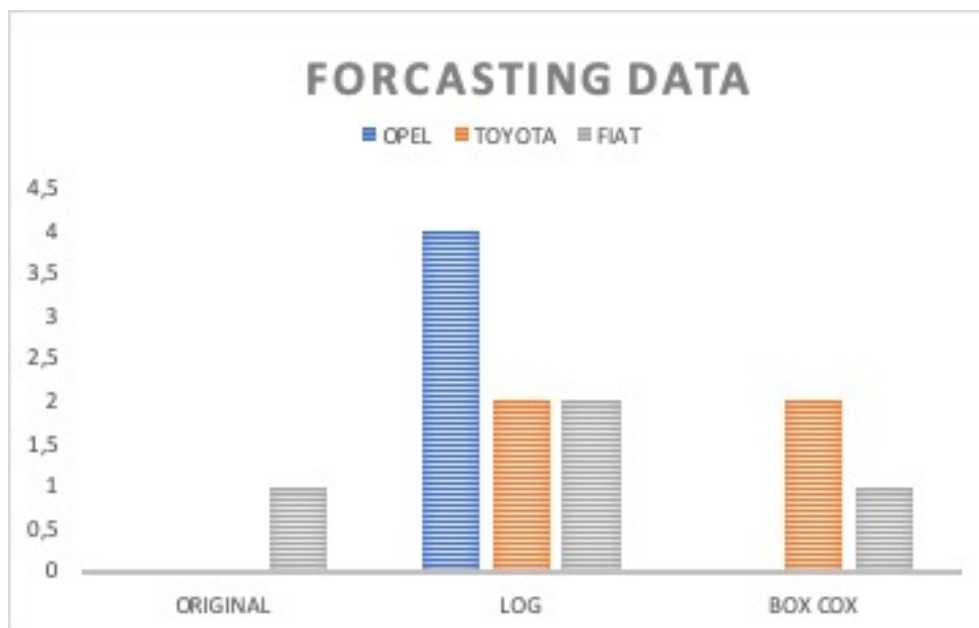


Figure 6.11: Forecasting new-car sales and data transformation.

Table 6.7: Forecasting Performance Comparison, Opel, Set A.

OPEL - SET A (h=48)			
CASE 1: ORIGINAL data	$\lambda = 1$		
Forecasting Model	RMSE	MAPE	MAE
Mean	1070	231	1055
Naïve	188	27	127
Seasonal Naïve	282	29	193
Linear Model with Seasonal Dummies	1307	281	1256
Exponential Smoothing state space	614	113	533
SARIMA(2, 1, 1)(2, 0, 0) ₁₂	234	29	171
CASE 2: LOG Values	$\lambda = 0$		
Forecasting Model	RMSE	MAPE	MAE
Mean	914	63	896
Naïve	188	27	127
Seasonal Naïve	282	29	193
Linear Model with Seasonal Dummies	993	62	938
Exponential Smoothing state space	314	34	273
SARIMA(1, 0, 1)(2, 0, 0) ₁₂	225	40	156
CASE 3: BOX-COX /Guerrero	$\lambda = 0.36$		
Forecasting Model	RMSE	MAPE	MAE
Mean	1074	232	1059
Naïve	1497	278	1304
Seasonal Naïve	347	52	249
Linear Model with Seasonal Dummies	446	91	427
Exponential Smoothing state space	521	102	466
SARIMA(2, 1, 1)(2, 0, 0) ₁₂	234	29	171

Table 6.8: Forecasting Performance Comparison, Opel, Set B.

OPEL - SET B (h=24)			
CASE 1: ORIGINAL data	$\lambda = 1$		
Forecasting Model	RMSE	MAPE	MAE
Mean	705	149	693
Naïve	240	63	202
Seasonal Naïve	169	29	133
Linear Model with Seasonal Dummies	309	52	239
Exponential Smoothing state space	157	27	125
SARIMA(1, 1, 1)(2, 0, 0) ₁₂	304	49	266
CASE 2: LOG Values	$\lambda = 0$		
Forecasting Model	RMSE	MAPE	MAE
Mean	526	49	509
Naïve	240	63	202
Seasonal Naïve	169	29	133
Linear Model with Seasonal Dummies	613	48	551
Exponential Smoothing state space	147	22	117
SARIMA(1, 0, 1)(2, 0, 0) ₁₂	234	64	196
CASE 3: BOX-COX /Guerrero	$\lambda = 0.31$		
Forecasting Model	RMSE	MAPE	MAE
Mean	712	150	700
Naïve	423	77	342
Seasonal Naïve	174	27	139
Linear Model with Seasonal Dummies	178	31	138
Exponential Smoothing state space	299	55	244
SARIMA(1, 1, 1)(2, 0, 0) ₁₂	304	49	266

Table 6.9: Forecasting Performance Comparison, Opel Set C.

OPEL - SET C (h=12)			
CASE 1: ORIGINAL data	$\lambda = 1$		
Forecasting Model	RMSE	MAPE	MAE
Mean	736	92	645
Naïve	668	88	531
Seasonal Naïve	782	31	537
Linear Model with Seasonal Dummies	723	77	637
Exponential Smoothing state space	467	47	392
SARIMA(1, 0, 0)(1, 0, 0) ₁₂	695	69	549
CASE 2: LOG Values	$\lambda = 0$		
Forecasting Model	RMSE	MAPE	MAE
Mean	652	36	566
Naïve	668	88	531
Seasonal Naïve	782	31	537
Linear Model with Seasonal Dummies	644	34	543
Exponential Smoothing state space	471	28	393
SARIMA(1, 0, 1)(1, 0, 0) ₁₂	655	33	501
CASE 3: BOX-COX /Guerrero	$\lambda = 0.77$		
Forecasting Model	RMSE	MAPE	MAE
Mean	737	92	647
Naïve	644	60	528
Seasonal Naïve	790	65	548
Linear Model with Seasonal Dummies	1166	98	1071
Exponential Smoothing state space	594	58	493
SARIMA(1, 0, 0)(1, 0, 0) ₁₂	695	69	549

Table 6.10: Forecasting Performance Comparison, Opel, Set D.

OPEL - SET D (h=24)			
CASE 1: ORIGINAL data	$\lambda = 1$		
Forecasting Model	RMSE	MAPE	MAE
Mean	797	98	732
Naïve	563	75	455
Seasonal Naïve	782	30	515
Linear Model with Seasonal Dummies	1079	99	1011
Exponential Smoothing state space	546	55	472
SARIMA(2, 0, 0)(1, 0, 0) ₁₂	674	69	563
CASE 2: LOG Values	$\lambda = 0$		
Forecasting Model	RMSE	MAPE	MAE
Mean	718	39	649
Naïve	563	75	455
Seasonal Naïve	782	30	515
Linear Model with Seasonal Dummies	745	39	666
Exponential Smoothing state space	436	29	392
SARIMA(2, 0, 0)(2, 0, 0) ₁₂	617	32	501
CASE 3: BOX-COX /Guerrero	$\lambda = 0.86$		
Forecasting Model	RMSE	MAPE	MAE
Mean	797	98	732
Naïve	518	45	392
Seasonal Naïve	755	60	528
Linear Model with Seasonal Dummies	581	57	495
Exponential Smoothing state space	593	55	471
SARIMA(2, 0, 1)(1, 0, 0) ₁₂	674	69	563

Table 6.11: Forecasting Performance Comparison, Toyota, Set A.

TOYOTA - SET A (h=48)			
CASE 1: ORIGINAL data	$\lambda = 1$		
Forecasting Model	RMSE	MAPE	MAE
Mean	852	93	724
Naïve	952	42	631
Seasonal Naïve	1225	123	1064
Linear Model with Seasonal Dummies	955	98	844
Exponential Smoothing state space	741	76	628
SARIMA(2, 1, 0)(2, 0, 0) ₁₂	582	55	489
CASE 2: LOG Values	$\lambda = 0$		
Forecasting Model	RMSE	MAPE	MAE
Mean	755	45	599
Naïve	927	94	631
Seasonal Naïve	1225	48	1064
Linear Model with Seasonal Dummies	688	39	561
Exponential Smoothing state space	784	40	667
SARIMA(4, 0, 0)(1, 0, 0)₁₂	565	32	396
CASE 3: BOX-COX /Guerrero	$\lambda = -0.14$		
Forecasting Model	RMSE	MAPE	MAE
Mean	873	96	749
Naïve	1677	184	1478
Seasonal Naïve	2719	262	2310
Linear Model with Seasonal Dummies	749	48	559
Exponential Smoothing state space	1638	1630	1336
SARIMA(2, 1, 0)(2, 0, 0) ₁₂	582	55	489

Table 6.12: Forecasting Performance Comparison, Toyota, Set B.

TOYOTA - SET B (h=24)			
CASE 1: ORIGINAL data	$\lambda = 1$		
Forecasting Model	RMSE	MAPE	MAE
Mean	1121	130	1045
Naïve	885	40	566
Seasonal Naïve	1073	87	868
Linear Model with Seasonal Dummies	926	90	842
Exponential Smoothing state space	933	95	859
SARIMA	1015	99	909
CASE 2: LOG Values	$\lambda = 0$		
Forecasting Model	RMSE	MAPE	MAE
Mean	783	24	481
Naïve	843	57	402
Seasonal Naïve	771	21	422
Linear Model with Seasonal Dummies	642	22	419
Exponential Smoothing state space	533	20	335
SARIMA(1, 0, 0)(1, 0, 0) ₁₂	731	23	420
CASE 3: BOX-COX /Guerrero	$\lambda = 0.51$		
Forecasting Model	RMSE	MAPE	MAE
Mean	1123	1311	1047
Naïve	1630	165	1440
Seasonal Naïve	1103	91	899
Linear Model with Seasonal Dummies	1692	99	1386
Exponential Smoothing state space	872	87	803
SARIMA(0, 0, 3)(2, 0, 0) ₁₂	1015	99	909

Table 6.13: Forecasting Performance Comparison, Toyota, Set C.

TOYOTA - SET C (h=12)			
CASE 1: ORIGINAL data	$\lambda = 1$		
Forecasting Model	RMSE	MAPE	MAE
Mean	539	25	423
Naïve	1502	62	1407
Seasonal Naïve	464	15	347
Linear Model with Seasonal Dummies	378	13	299
Exponential Smoothing state space	403	13	308
SARIMA(0, 0, 0)(1, 0, 0) ₁₂	435	16	327
CASE 2: LOG Values	$\lambda = 0$		
Forecasting Model	RMSE	MAPE	MAE
Mean	569	23	447
Naïve	1502	195	1407
Seasonal Naïve	464	20	347
Linear Model with Seasonal Dummies	642	11	220
Exponential Smoothing state space	285	11	229
SARIMA(0, 0, 0)(1, 0, 0) ₁₂	450	18	336
CASE 3: BOX-COX /Guerrero	$\lambda = -0.99$		
Forecasting Model	RMSE	MAPE	MAE
Mean	527	25	405
Naïve	1326	55	1227
Seasonal Naïve	453	14	336
Linear Model with Seasonal Dummies	292	10	229
Exponential Smoothing state space	283	10	234
SARIMA(0, 0, 0)(1, 0, 0) ₁₂	435	16	327

Table 6.14: Forecasting Performance Comparison, Toyota, Set D.

TOYOTA - SET D (h=24)			
CASE 1: ORIGINAL data	$\lambda = 1$		
Forecasting Model	RMSE	MAPE	MAE
Mean	964	38	842
Naïve	1627	63	1498
Seasonal Naïve	533	23	426
Linear Model with Seasonal Dummies	549	19	469
Exponential Smoothing state space	865	35	793
SARIMA	862	30	735
CASE 2: LOG Values	$\lambda = 0$		
Forecasting Model	RMSE	MAPE	MAE
Mean	1151	82	1038
Naïve	1627	63	1498
Seasonal Naïve	533	23	426
Linear Model with Seasonal Dummies	1067	73	962
Exponential Smoothing state space	948	60	854
SARIMA(1, 0, 1)(1, 0, 0) ₁₂	951	56	818
CASE 3: BOX-COX /Guerrero	$\lambda = -0.18$		
Forecasting Model	RMSE	MAPE	MAE
Mean	953	37	831
Naïve	1131	47	986
Seasonal Naïve	371	13	295
Linear Model with Seasonal Dummies	543	20	463
Exponential Smoothing state space	648	24	534
SARIMA(0, 1, 1)(2, 0, 0) ₁₂	862	30	735

Table 6.15: Forecasting Performance Comparison, Fiat, Set A.

FIAT - SET A (h=48)			
CASE 1: ORIGINAL data	$\lambda = 1$		
Forecasting Model	RMSE	MAPE	MAE
Mean	848	314	835
Naïve	209	39	155
Seasonal Naïve	164	37	116
Linear Model with Seasonal Dummies	224	71	185
Exponential Smoothing state space	128	31	97
SARIMA(1, 1, 2)(2, 0, 0) ₁₂	138	43	110
CASE 2: LOG Values	$\lambda = 0$		
Forecasting Model	RMSE	MAPE	MAE
Mean	664	247	648
Naïve	209	39	155
Seasonal Naïve	164	37	116
Linear Model with Seasonal Dummies	125	32	93
Exponential Smoothing state space	123	30	97
SARIMA(1, 0, 1)(2, 0, 0) ₁₂	230	46	176
CASE 3: BOX-COX /Guerrero	$\lambda = 0.10$		
Forecasting Model	RMSE	MAPE	MAE
Mean	858	317	846
Naïve	360	116	295
Seasonal Naïve	156	40	103
Linear Model with Seasonal Dummies	752	253	722
Exponential Smoothing state space	123	37	92
SARIMA(1, 1, 2)(2, 0, 0) ₁₂	138	43	110

Table 6.16: Forecasting Performance Comparison, Fiat, Set B.

FIAT - SET B (h=24)			
CASE 1: ORIGINAL data	$\lambda = 1$		
Forecasting Model	RMSE	MAPE	MAE
Mean	419	166	400
Naïve	157	33	114
Seasonal Naïve	134	40	104
Linear Model with Seasonal Dummies	743	27	658
Exponential Smoothing state space	121	42	96
SARIMA(0, 1, 3)(2, 0, 0) ₁₂	153	36	119
CASE 2: LOG Values	$\lambda = 0$		
Forecasting Model	RMSE	MAPE	MAE
Mean	307	47	283
Naïve	157	51	114
Seasonal Naïve	134	34	104
Linear Model with Seasonal Dummies	342	47	326
Exponential Smoothing state space	122	28	121
SARIMA(1, 0, 1)(2, 0, 0) ₁₂	155	56	122
CASE 3: BOX-COX /Guerrero	$\lambda = -0.80$		
Forecasting Model	RMSE	MAPE	MAE
Mean	425	168	406
Naïve	215	77	179
Seasonal Naïve	148	46	105
Linear Model with Seasonal Dummies	124	43	92
Exponential Smoothing state space	134	48	100
SARIMA(0, 1, 3)(2, 0, 0) ₁₂	153	36	119

Table 6.17: Forecasting Performance Comparison, Fiat, Set C.

FIAT - SET C (h=12)			
CASE 1: ORIGINAL data	$\lambda = 1$		
Forecasting Model	RMSE	MAPE	MAE
Mean	454	95	392
Naïve	482	48	382
Seasonal Naïve	524	79	433
Linear Model with Seasonal Dummies	404	73	367
Exponential Smoothing state space	391	68	348
SARIMA(0, 0, 0)(1, 0, 0) ₁₂	454	84	405
CASE 2: LOG Values	$\lambda = 0$		
Forecasting Model	RMSE	MAPE	MAE
Mean	430	41	383
Naïve	482	78	382
Seasonal Naïve	524	48	433
Linear Model with Seasonal Dummies	387	38	359
Exponential Smoothing state space	382	38	360
SARIMA(1, 0, 0)(1, 0, 0) ₁₂	443	46	391
CASE 3: BOX-COX /Guerrero	$\lambda = 0.63$		
Forecasting Model	RMSE	MAPE	MAE
Mean	455	95	392
Naïve	481	74	425
Seasonal Naïve	529	81	442
Linear Model with Seasonal Dummies	411	75	375
Exponential Smoothing state space	413	77	376
SARIMA(1, 0, 0)(1, 0, 0) ₁₂	854	158	819

Table 6.18: Forecasting Performance Comparison, Fiat, Set D.

FIAT - SET D (h=24)			
CASE 1: ORIGINAL data	$\lambda = 1$		
Forecasting Model	RMSE	MAPE	MAE
Mean	605	135	542
Naïve	355	38	250
Seasonal Naïve	536	87	437
Linear Model with Seasonal Dummies	482	94	443
Exponential Smoothing state space	477	95	435
SARIMA(1, 0, 0)(1, 0, 0) ₁₂	556	117	516
CASE 2: LOG Values	$\lambda = 0$		
Forecasting Model	RMSE	MAPE	MAE
Mean	534	46	484
Naïve	355	51	250
Linear Model with Seasonal Dummies	518	45	344
Exponential Smoothing state space	441	42	407
SARIMA(1, 0, 0)(1, 0, 0) ₁₂	401	41	363
CASE 3: BOX-COX /Guerrero	$\lambda = 0.07$		
Forecasting Model	RMSE	MAPE	MAE
Mean	607	135	544
Naïve	612	125	544
Seasonal Naïve	615	104	514
Linear Model with Seasonal Dummies	525	104	483
Exponential Smoothing state space	509	101	466
SARIMA(0, 1, 1)(2, 0, 0) ₁₂	415	82	380

Forecast Combinations.

7.1 Introduction.

The theoretical foundation of forecast combination started five decades ago, initiated by the seminar papers of Crane and Crotty [1967] and Bates and Granger [1969]. Since at least 1969, when Bates and Granger wrote their famous paper on “The Combination of Forecasts” it has been well-known, that combining forecasts often leads to better forecast accuracy. Therefore, an easy way to improve forecast accuracy is to use a combination of several methods on the same time series, like for example to average the resulting forecasts or to use weights for the forecasts and so on.

In response to the criticisms of the idea of combining, Newbold and Granger [1974] agreed: “...that combination is not a valid proposition if one of the individual forecasts does not differ significantly from the optimum”. However, combining forecasts from very similar models is also important. Until today there has been considerable research on using weighted averages or some other more complicated combination approach. An extensive review of the literature, techniques, and applications of forecast combinations can be found in Clemen [1989] where he wrote: “The results have been virtually unanimous: combining multiple forecasts leads to increased forecast accuracy. In many cases, one can make dramatic performance improvements by simply averaging the forecasts”.

Instability of a model selection has been recognized in statistics and related literature,

as Breiman [1996] argue. Therefore, when multiple models are considered for estimation and forecasting, the term “model uncertainty” has been used by several authors to capture the difficulty in identifying the correct model, according to Chatfield [1996].

Combining forecasts has been studied for the past three decades and various methods have been proposed. The focus has been on the case where the forecasts to be combined are distinct in nature (i.e. based on very different methods). For example, Clement and Hendry [1998] stated that “When forecasts are all based on econometric models, each of which has access to the same information set, then combining the resulting forecasts will rarely be a good idea. It is better to sort out the individual models—to derive a preferred model that contains the useful features of the original models”.

Generally, the ‘true’ model may or may not be in the candidate research list and even if the true model happens to be included, the task of finding the true model can be very different from that of finding the best model for the purpose of prediction. Hoeting et al. [1999] argues that finding the ‘best’ model may be defined sometimes in terms of an appropriate loss function (e.g. square error loss in prediction).

Additionally, it is also accepted that different forecasting models deliver different results at different time periods. There is strong empirical support that the performance of different models, change over time, according to Elliott and Timmermann [2005] and many others. Following the advice in Hansen [2005] we abandon the conceptual error of assuming one true single data generating process, so we are free to include information from different models. According to Hansen [2005] “Models should be viewed as approximations and econometric theory should take this seriously”. Thus, selecting a single forecasting model as the “best one” bears the risk of ending up with a model, which is only accurate when evaluated using some validation sample, yet might prove unreliable, when applied to new data.

In general, the combination reduces the information in a vector of forecasts to a single summary measure using a set of combination weights. The optimal combination chooses weights that minimize the expected loss of the combined forecast. The technique gives larger weights to more accurate forecasts and small estimation errors. In a word with

no model misspecification, infinite data samples (i.e. no estimation error and complete access to the information sets underlying the individual forecasts) there is no need for forecast combination. Techniques and applications of forecast combinations can be found in Timmermann [2006].

There has been in the past and until today, many research papers discussing new combinations techniques and stimulate further research, like for example Hansen [2007, 2008] and Hansenn and Racine [2012]. Furthermore, there is a research of forecast combinations not only in the "first moments" but also for higher moments as well, like for volatility forecasting in Christiansen et al. [2012].

In many cases, according to Hyndman and Athanasopoulos [2013] research evidence, one can make dramatic performance improvements by simply averaging the forecasts using a simple average, and furthermore this method has been proven hard to beat. Combining has great potential to reduce the variability that arises in the forced action of selecting a single model. The simple combining methods in the literature attempt to improve the individual forecasts, while the more advanced target is always on the performance of the best candidate model according to Elliott et al. [2013].

Generally, it is accepted that a combination of forecasts from different models is an appealing strategy to hedge against forecast risk. According to Graefe et al. [2014] these are often cases when combined forecasts are more accurate than even their best component. Additionally, Opschoor et al. [2014] research of forecast combinations is for Value-at-Risk forecasting, while Morana [2015] introduces a new "Frequentist" model averaging estimation procedure by minimizing the squared Euclidean distance between actual and predicted value vectors (MSE metric) that yields more accurate and more efficient estimation. Cheng and Yang [2015] argues that combining forecasts can also minimize the occurrence of forecast outliers and proposed a synthetic loss function to achieve both the usual accuracy and outlier-protection simultaneously.

Kourentzes et al. [2019] argues that selecting a reasonable pool of forecast is fundamental in the modeling process and considers forecast selection and combination as two extreme pools of forecasts thus propose a model to construct forecast pools so as to improve

performance and reduce computational effort.

In macroeconomics and finance, there are many applications using combining different forecast methods in order to hedge against risk. For example, Avramov [2002], Ravazzolo et al. [2007] and Rapach et al. [2010] predict stock returns, Stock and Watson [2004] use forecast combination for output forecasting, Wright [2008] research the exchange rate forecasts, Kapetanios et al. [2008] and Wright [2009] focus in inflation forecasting research, Andrawis et al. [2011] consider forecasting unbound tourism figures, Magnus and Wang [2014] explore growth determinants, Nowotarski et al. [2014] and Raviv et al. [2015] research electricity price forecasting and Weiss [2017] study health demand forecasting.

In this chapter, the researcher will provide comprehensive implementation of common ways in which forecasts can be combined. Various estimation methods are going to be explained for creating a combined forecast and implemented to various data sets in order to rationalize and visualize the combination results.

The plan of this chapter is the following: we start with the introduction of the forecast combination theory and an extensive literature review. Section two makes a reference to the methodology used and divide the combination forecast techniques into two categories: combination forecast with or without a training set. This training set is needed for the weights estimation of the individual forecast. So the Simple Average Combination technique, that works without a training data set is introduced in 4 different combinations which are easy to implement and hard to beat, due to their excellent results.

On the other hand more complicated techniques, are explained and applied that need a training set for their calculation, like Bates & Granger (1969) and Newbold & Granger (1974). Additionally we explain the data generating process of this empirical research. In section three we illustrate the empirical results of the research both numerically and graphically for the six (6) different combination forecasts. Lastly we discuss the conclusion of the forecast combination empirical research.

7.2 Methodology

There are various frequently used schemes for forecast combinations since the seminal paper by Bates and Granger [1969]. Research by Batchelor and Dua [1995] showed that combining was more effective when data and methods differed substantially. In this method, the averaging is done by using a rule that can be replicated, i.e. take the simple average of the forecasts. Opponents of this method believe more in traditional statistical procedures and that there is one right way to forecast and develop a comprehensive model that can incorporate all relevant information and be more effective than others.

According to Armstrong [2001] combining forecasts sometimes referred to as composite forecasts, refers to the averaging of independent forecasts. These forecasts can be based on different data or different methods or both. Some researchers even suggest to combine “combined forecasts” like the so-called hierarchical forecast combinations of Andrawis et al. [2011]. The question is which is the best way to combine different forecasts. This unfortunately has no theoretical underpinning and it mainly depends on the data in research according to Weiss et al. [2018].

However, in practice, combining forecast is an approach with very good results. Based on the assumption that each model has something to contribute and to improve forecast accuracy it usually wins the single best method or frame of the research. The combining method is even more relevant when there is uncertainty about the method or situation and when it is important to avoid large errors in the future, like for example in inventory management or sales forecasts.

Due to the lack of theoretical foundation in this area and the empirical lack of evidence for one single way, which dominates the way forecasts should be combined, various forecast combinations are presented and applied in this research.

Forecast can be combined in a very simple way and these simple methods there is no need to exactly estimate the weight each forecast should be given in the overall combination. Location measures of the cross-sectional distribution of the individual forecasts are used, as the average, the median, or the mean. However those location measures are all the same

if the cross-sectional distribution is symmetric, but if the distribution is asymmetric then one can choose one of these simple measures.

7.2.1 Combining forecasts without training

In this technique measures of central tendency are used, like mean (or median), which is the most straight forward way of combining forecasts from multiple forecasting methods. This simple approach requires no additional input besides just the forecasts being combined which is an appealing practical advantage. On the other hand, combining forecasts with the help of training data is a more cumbersome and computational expensive approach compared to the one without training data to produce a more accurate combined forecast.

- Simple Average Combination

The most natural approach to combine forecasts is using the mean of all those forecasts. Over the years this innocent approach has been proven as an excellent benchmark despite or perhaps because of its simplicity Genre et al. [2013]. This approach uses the average of all the forecasts to combine them giving equal weights in each component forecast. The combined forecast is given by :

$$f^{combined} = \frac{1}{P} \sum_{i=1}^p f_i \quad (7.1)$$

where $f^{combined}$ is the combined forecast, f_i is the forecast obtained using model i , where $i \in (1, \dots, P)$ and P are forecast at each point in time.

7.2.2 Combining forecasts with training

This technique uses information about how the individual forecasting methods performed and can be utilized and not ignored, to find a more optimal weighting scheme. So instead of just assuming equal weights, these weights can be obtained by using the historical forecasts made for the q periods before t , i.e. $i=t-1, t-2, \dots, t-q$, where q is a positive integer less than t . This is done because the data set could contain valuable information about how the individual forecasts at t , $f_{t(j)}$ should be weighted optimally.

There is a two steps procedure in this method:

- Step 1. Actual values and individual forecasts are used, from periods $t-q$ to $t-1$, to fit an optimal model and weights.
- Step 2. The fitted model is utilized to construct the predicted combined forecast f_t by using the individual forecasts as input, at period t .

- Bates/Granger(1969) Combination

A weighting scheme based on the individual performance of each of the forecast method can be computed using the mean squared distance between actual value (y_i) and forecast value ($f_{i(j)}$ where $i = t-1, t-2, \dots, t-q$) and is called the Variance based weighting method. The mean squared error (MSE) for a particular forecasting method is computed as,

$$MSE_j = \frac{1}{q} \sum_{i=t-q}^{t-1} (x_i - f_{i(j)})^2 \quad (7.2)$$

The forecast weights (w_j) are then obtained as :

$$w_j = \frac{\frac{1}{MSE_j}}{\sum_{j=1}^k \frac{1}{MSE_j}} \quad (7.3)$$

where the reciprocal of MSE_j can be viewed as an accuracy measurement meaning that higher accuracy generates a higher relative weight and a lower and less and $\sum_{j=1}^k w_j = 1$. The combined forecast at period t is therefore computed as:

$$f_t^{combined} = w_1 f_{t(1)} + w_2 f_{t(2)} + \dots + w_k f_{t(k)} \quad (7.4)$$

This method is one of the methods mentioned in the seminar paper by Bates and Granger [1969]. Their approach builds on portfolio diversification theory and uses the diagonal elements of the estimated mean squared prediction error matrix in order to compute combination weights.

- Newbold/Granger (1974) Combination

The methodology of Newbold and Granger [1974] extracts the combination weights from the estimated mean squared prediction error matrix. Suppose x_t is the variable of interest, there are k not perfectly collinear predictors,

$$f_t = (f_{1t}, \dots, f_{t(k)})' \quad (7.5)$$

Σ is the (positive definite) mean squared prediction error matrix of f_t and e is an $k \times 1$ vector of $(1, \dots, 1)'$. Building on the Bates and Granger [1969] early research this method is a constrained minimization of the mean squared prediction error using the normalization condition $e'w = 1$. This yields the following combination weights:

$$w = \frac{\Sigma^{-1}e}{e' \Sigma^{-1}e} \quad (7.6)$$

The combined forecast is then obtained by:

$$f_t^{combined} = (f_t)'w \quad (7.7)$$

This method according to Timmermann [2006] ignores correlations across forecast errors, just like the Bates and Granger [1969] method but on the other hand is more robust to outliers, since total rankings are not likely to change dramatically by the presence of extreme forecasts. While the method dates back to Newbold and Granger (1974), the variant of the method used here does not impose the prior restriction that Σ is diagonal. This approach is used by Hsiao and Wan [2014] as a generalization of the original method and may be viewed as a way to make the forecast more robust against misspecification biases and measurement errors in the data set.

There are of course the regression-based and the eigenvector-based combination methods which are a natural extension to the previous approaches viewed through the lens of regression or based on the idea of minimizing the mean squared prediction error subject to a normalization condition accordingly and are not going to be further discussed for this study.

7.2.3 Data generating procedure

According to Weiss et al. [2018] the research function required as inputs a vector of the actual data and a matrix of the set of component forecasts to be combined. Observation values refer to sales so they are all non-negative. New-car sales refer to Opel, Toyota and Fiat operating in the Greek market, and values in their original form, in log values and in Box-Cox with Guerrero transformed values are tested.

For each series that is included in the evaluation, the time horizon is divided into the training and test set with an 80:20 proportion. Forecasts from the various univariate forecasting methods are generated using the historical values of the training set and these forecasts are going to be later used as input to combined forecasts. Each test set data has a different forecast horizon in this study i.e. in A data set there are 48 steps ahead ($h=48$), in B data set there are 24 steps ahead ($h=24$), in C data set there are 12 steps ahead ($h=12$), in D data set there are 24 steps ahead ($h=24$), forecasts horizon.

In the combining forecast methods where test sets are needed, the test set outputs are divided into half (50%) and the first half is used to compute the combined forecast models and estimate the weights for each individual forecast, and the other half is used to evaluate the forecasting results from the combined forecasting models.

7.3 Empirical Results

A set of components forecast is already obtained in our previous research using various statistical techniques on the new car sales data from the Greek market, and now we seek to improve accuracy by combining those component forecasts into one. The univariate forecast models used until this point are the Mean, the Naïve, the Seasonal Naïve, the Linear model with seasonal dummies (LMSD), the Exponential smoothing state space (ETS) models, the Seasonal ARIMA (SARIMA) and the SARIMA-GARCH models with and without data transformation.

The research continues by choosing the three (3) best performed individual models. In this research data will have three different forms: original, logs, and transformed using

Guerrero Method for Box-Cox transformation, as it has been proven to be very effective. The research will examine if the combined models can produce better forecasts values from the individual forecasting models, for three different firms (Opel, Toyota and Fiat) in four different time intervals. The component models to be combined in our empirical study are the following :

- Seasonal Naïve - SN
(produced using the *snaive* function in forecast package)
- Linear Model with seasonal Dummies - LMSD
(produced using the *tslm* function in forecast package)
- Exponential smoothing Space state models - ETS
(produced using the *ets* function in forecast package)

To implement the models we divide the monthly observation of the data sets A, B, C, and D in a Training set and a Test set, as specified in Table 4.1 page 126. Then estimate each time series model using the training set observations and produce point forecast for the next forecast horizon which are compared with the test set actual values for the estimation of forecasting accuracy measures. Furthermore, to illustrate the combination methodology we apply the combination techniques of the simple average, the Bates and Granger [1969] methodology, and the Newbold and Granger [1974] combined forecast model.

The empirical research starts with a variety of simple average combined models, which combine the forecast values of the various time series forecasting models in equal weights. For the three (3) individual forecasting models the researcher created four (4) different combination models: three (3) combinations of two (2) models and one (1) combination of all three models. In more detail the following models that combined forecast error as a weighted average of the individual forecast errors are created:

$$- Combo_1 \Rightarrow (SN + ETS) / 2$$

$$- Combo_2 \Rightarrow (ETS + LMSD) / 2$$

- $Combo_3 \Rightarrow (SN + LMSD)/2$
- $Combo_4 \Rightarrow (SN + LMSD + ETS) /3$

Furthermore the research also focuses on the two more combinations that extract combination weights from the estimated mean squared prediction error matrix :

- $Combo_5 \Rightarrow$ Bates and Granger [1969] forecast combination or Variance based weighting method and
- $Combo_6 \Rightarrow$ Newbold and Granger [1974] forecast combination or Variance based constrained weighting method

7.3.1 Graphical presentation of combination forecasting

We present a visual plot of the forecast values from the combination forecast models alone with the actual for Opel, Toyota and Fiat values for data set D when the series are transformed with Box-Cox and Guerrero method.

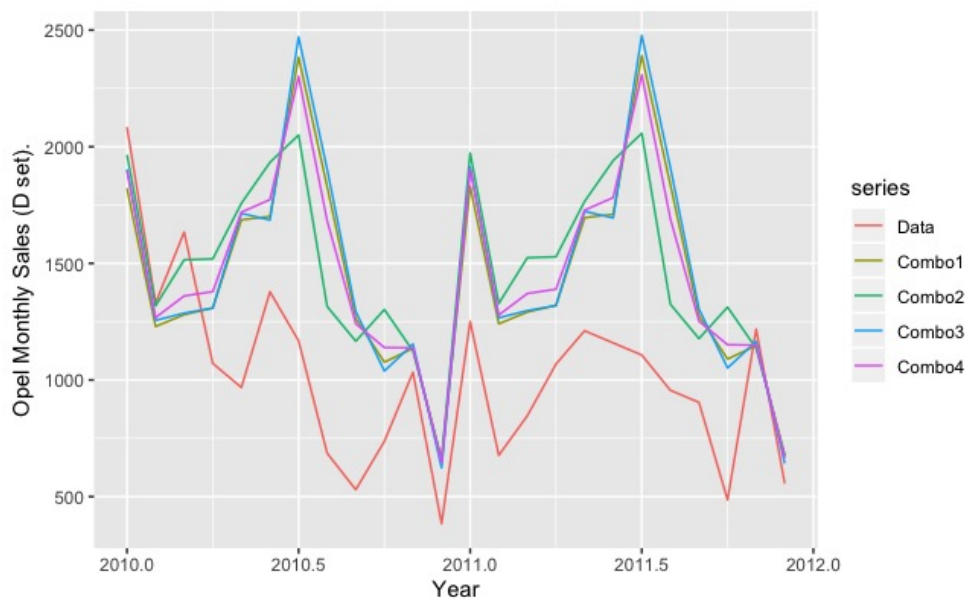


Figure 7.1: Simple Average Combination forecasts (Opel-Set D,BC/Guerrero)

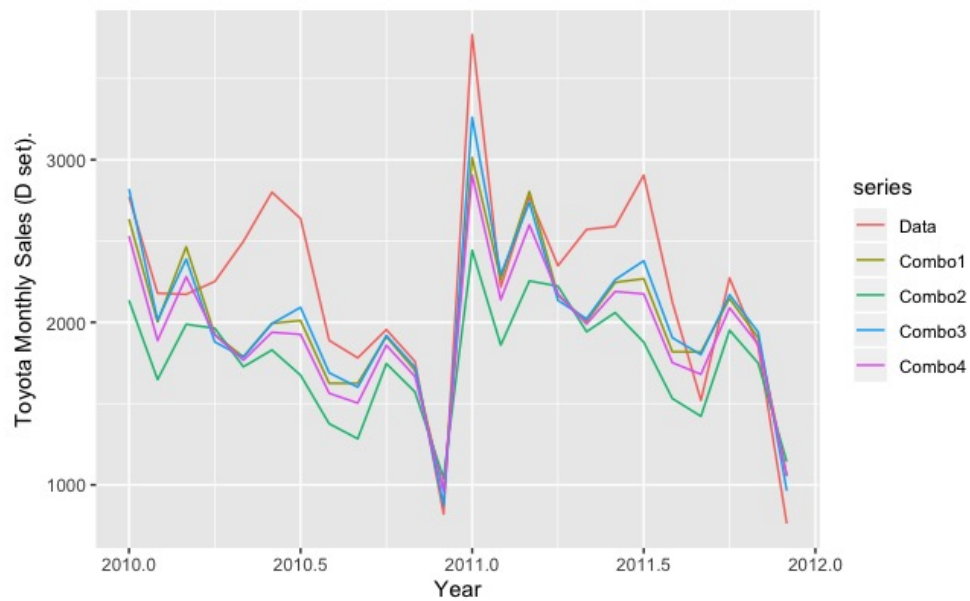


Figure 7.2: Simple Average Combination forecasts (Toyota-Set D,BC/Guerrero)

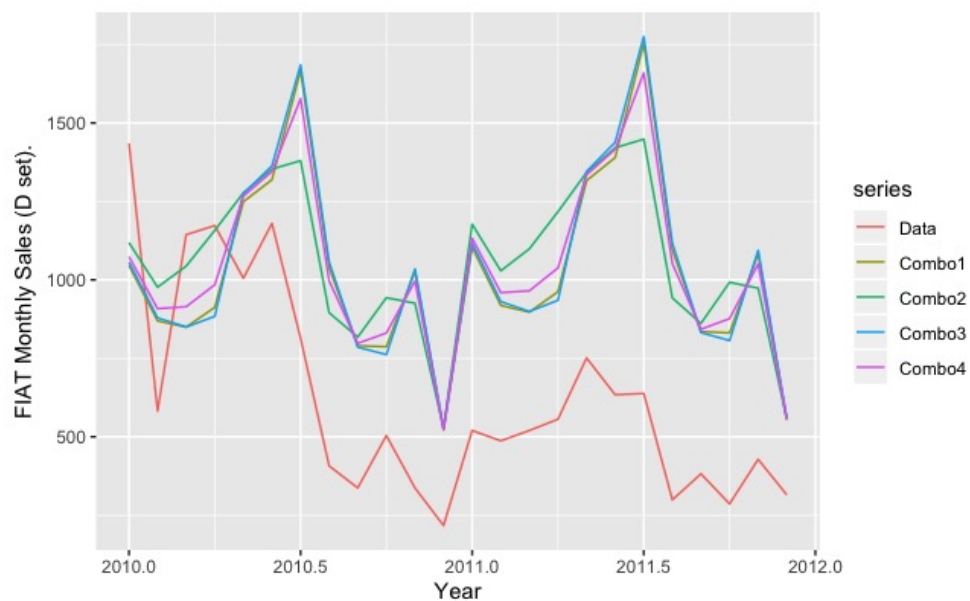


Figure 7.3: Simple Average Combination forecasts (Fiat-Set D, BC/Guerrero)

The simple average combined forecast methods in the four (4) different combinations are presented graphically for Opel, Toyota and Fiat for the D data set in Figure 7.1 (page

225), Figure 7.2 (page 226), and Figure 7.3 (page 226) respectively for the forecast horizon of 2 years. The actual data are presented in the red line and each Combination has a different color. It can be visually noticed that in the case of Opel and Fiat *Combo*₂ (i.e. the combination num.2 -green line), which is $(ETS + LMSD)/2$ simple average combined forecast model, seems to perform better than the other simple average models. However, in Toyota case study, *Combo*₃ (i.e combination num.3 -blue line), which is the $(Seasonal\ Naive + LMSD)/2$ simple average combined forecast models seem to perform best for Toyota cars.

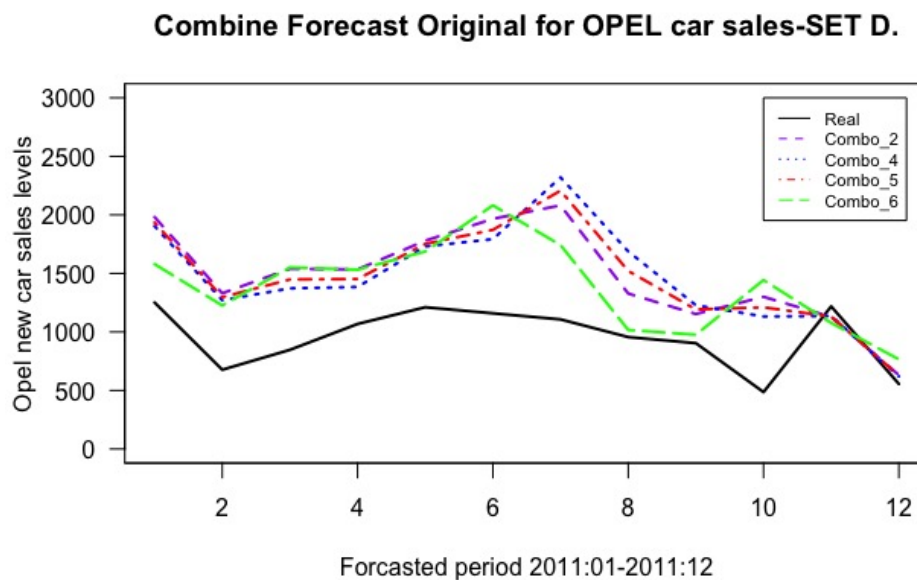


Figure 7.4: Selection of the best Combined forecasts (Opel-Set D,BC/Guerrero)

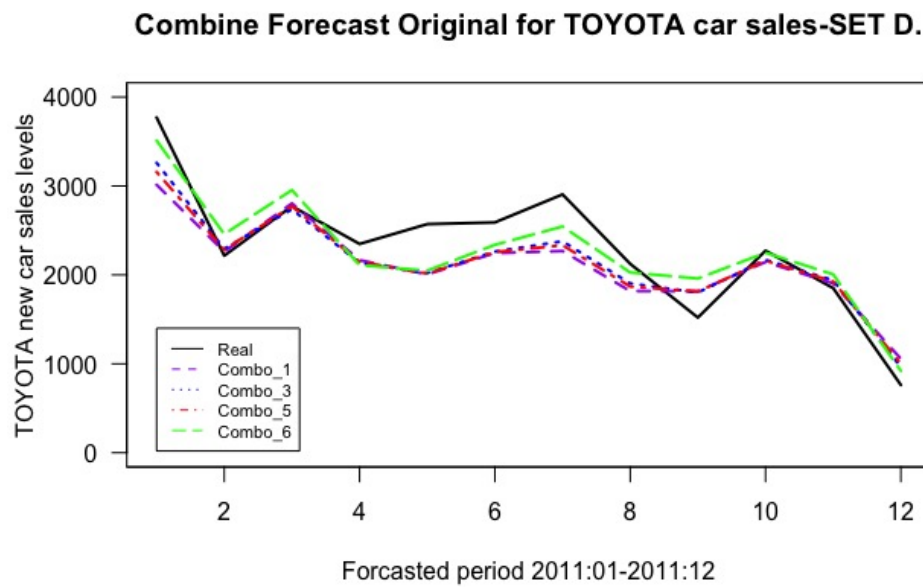


Figure 7.5: Selection of the best Combined forecasts (TOYOTA-Set D, BC/Guerrero)

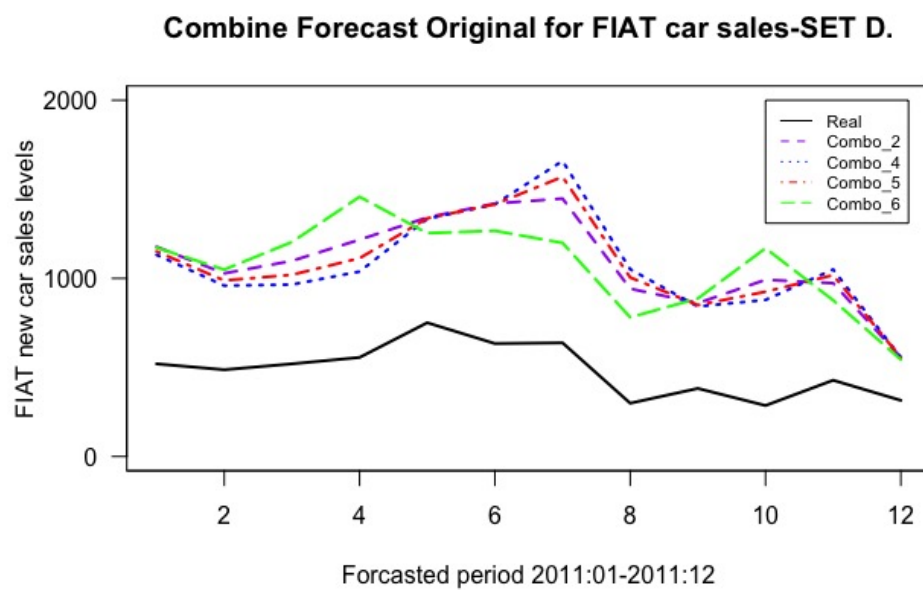


Figure 7.6: Selection of the best Combined forecasts (FIAT -Set D, BC/Guerrero)

The four (4) best combined forecasting models of the Opel, Toyota and Fiat for the Data set D when series are transformed, according to Box-Cox with Guerrero, are illustrated in

Figure 7.4 (page227), Figure 7.5 (page228) and Figure 7.6 (page 228). The actual values of the series and the forecast values of the various combined forecasting models in the test set are presented in a one-year forecast horizon. We notice that the downfall of the sales is hard to predict so the various models usually overestimate the sales levels. However, the model that gives the best forecast, visually is the line that comes closer to the line of the actual values line (black line) during that period. In other words it is *Combo*₆ for Opel and Toyota and *Combo*₂ for Fiat.

7.3.2 Accuracy measures of combination forecasting.

The forecasting performance measures of the various combination of different time series models are illustrated for the Opel, Toyota and Fiat for *Combo*₁ till *Combo*₆ in Table 7.1 till Table 7.3 (page 234 - 236). When forecasts are from the same series and the same data set, it is reasonable to use the root mean squared error (RMSE), the mean absolute error (MAE) and the Mean absolute percentage error (MAPE) as metrics in order to compare the different combination forecast models. However if we need to compare different time series and different time intervals it is better to use the MAPE metric.

For forecasting purposes, the model with the minimum accuracy measures gives the better forecasts. Similar empirical studies, often find that simple weighted forecast combinations perform very well, compared with more sophisticated combination schemes that rely on estimated combination weights. This rise a question, as Smith and Wallis [2009] wonder: “Why is it that, in comparisons of combinations of point forecast based on mean-squared forecast errors ... a simple average with equal weights, often outperforms more complicated weighting schemes.” According to Timmermann [2006] errors introduced by estimation of the combination weights could overwhelm any gains from setting the weights to their optimal values overusing equal weights. Furthermore, evidence shows that estimation error might be large and additionally gains from setting the combination weights to their optimal values might be small relative to using equal weights.

For the case study of Opel (see Table 7.1 p.234) it is clear that *Combo*₂ (ETS+LMSD)/2 is in favour for the original and all series transformation in set C and also for the log values

in set D. $Combo_3$ (SN+LMSD)/2 is in favour for log and Box-Cox values in set B and Box-Cox values in set A, while Bates and Granger model is in favour for the original values in set A,B and D and the log values in set A. Newbold Granger variance-based constrained weighted forecast combination method is in favour only for the Box-Cox set D values.

For the case study of Toyota (see Table 7.2, p. 235) it is clear that, $Combo_2$ (ETS + LMSD) /2 is in favour for, the original and all series transformation, in set A and B. Newbold Granger method, is in favour for all cases, in set C and D, except the Box-Cox transformation in set C, that prefers the Bates and Granger approach.

For the case study of Fiat (see Table 7.3, p.236) it is clear that $Combo_2$ (ETS + LMSD) /2 is in favour for, the original and all series transformation, in set C, and the case of Box-Cox transformation of set A and D. $Combo_3$ (SN + LMSD) /2 is in favour for, the original and all series transformation, in set B, the original values in set A and D and the log values of A. Newbold Granger method is in favour only for the case of log values in set D.

7.4 Discussion

The results from this evaluation study clearly shows that it is possible to combine univariate forecasts to achieve better forecast accuracy compared to just selecting the best individual forecast model. This evidence goes hand in hand with previous investigations Clemen [1989]. However, the choice of method is important since some of the combining methods perform much worse than even the worse univariate forecasting methods.

Various time series models are considered in this thesis for fitting time-series data. The task of choosing the most appropriate one for forecasting is proved to be very difficult. In this last chapter, the use of a combining method is proposed to convexly combine the candidate models, instead of selecting one of them. The idea is that when there is much uncertainty in finding the best model, as is the case study of the Greek new car market empirical application, combining the models may reduce the instability of the forecast and therefore improve prediction accuracy.

According to our research there is substantial evidence that forecast combination is beneficial in terms of reducing the forecast errors as well as toning down modeling uncertainty as the researcher is not forced to choose a single model. Furthermore, it is a good strategy to hedge against model risk. Combined forecasting data results and actual data examples indicate the potential advantage of this method over model selection for this case study.

Combinations of forecasts are motivated by misspecified forecasting models due to diversification across forecasts and uncertain economic conditions. According to Elliott and Timmermann [2016] simple, robust estimation schemes tend to work well in small samples where the estimation of errors is done in combination weight. There is evidence that even if they do not always deliver the most precise forecasts, forecast combinations, particularly equal-weighted ones, generally do not deliver poor performance and so from a "risk" perspective, they represent a relatively safe choice. Empirically, this thesis research proved that simple combination forecasts work well for sales of new cars in the Greek market.

The results of the combined forecasting methods in our research data example can be observed in Figure 7.7 (page 233) and Table 7.4 (page 232) can be summarized as follows:

- A Combination of time series models seems to produce forecasts closer to the actual values of the series and therefore give a better forecast for our variables.
- Simple average combination with equal weights is usually the best model for forecasting purposes. Easy to calculate hard to beat!
- Log transformation in Toyota and Fiat give better combination forecasts (lower MAPE metric) while original values for Opel data.
- It is hard to specify one sole model for all cases, each case and each time frame has to be examined separately and with caution.

This study suggests that combination forecasts are almost certain to outperform the individual forecasts and avoid the risk of complete forecast failure. Therefore, in circumstances where forecasting models are available and the researcher has to generate forecasts

but is uncertain as to which model is likely to generate the best forecasts, combining the forecasts from various alternative models would be the best and safest way forward.

To sum up, time series methods, especially combined forecasting is proven to be successfully applied in new car sales i.e. marketing data in Greece and give reliable forecasts.

Table 7.4: Summary of Best Forecasting Models(*min value)

M A P E	Original	Logs	Box-Cox
OPEL			
A {1998-2012}:2013-16	Naïve/ <i>Combo</i> ₅	Naïve*	SARIMA
B {2006-2013}:2014-15	<i>Combo</i> ₅ *	ETS	S. Naïve
C {2006-2009}:2010-10	S.Naïve	ETS*	<i>Combo</i> ₂
D {2002-2009}:2010-11	S.Naïve	ETS*	Naïve
TOYOTA			
A {1998-2012}:2013-16	Naïve/ <i>Combo</i> ₅	SARIMA*	LMSD
B {2006-2013}:2014-15	Naïve	ETS*	ETS/ <i>Combo</i> ₂
C {2006-2009}:2010-10	<i>Combo</i> ₂	<i>Combo</i> ₃	<i>Combo</i> ₅ *
D {2002-2009}:2010-11	<i>Combo</i> ₆	<i>Combo</i> ₆ *	<i>Combo</i> ₆
FIAT			
A {1998-2012}:2013-16	ETS	ETS*	ETS
B {2006-2013}:2014-15	LMSD*	ETS	SARIMA
C {2006-2009}:2010-10	Naïve	ETS*	Naïve
D {2002-2009}:2010-11	Naïve*	SARIMA	SARIMA

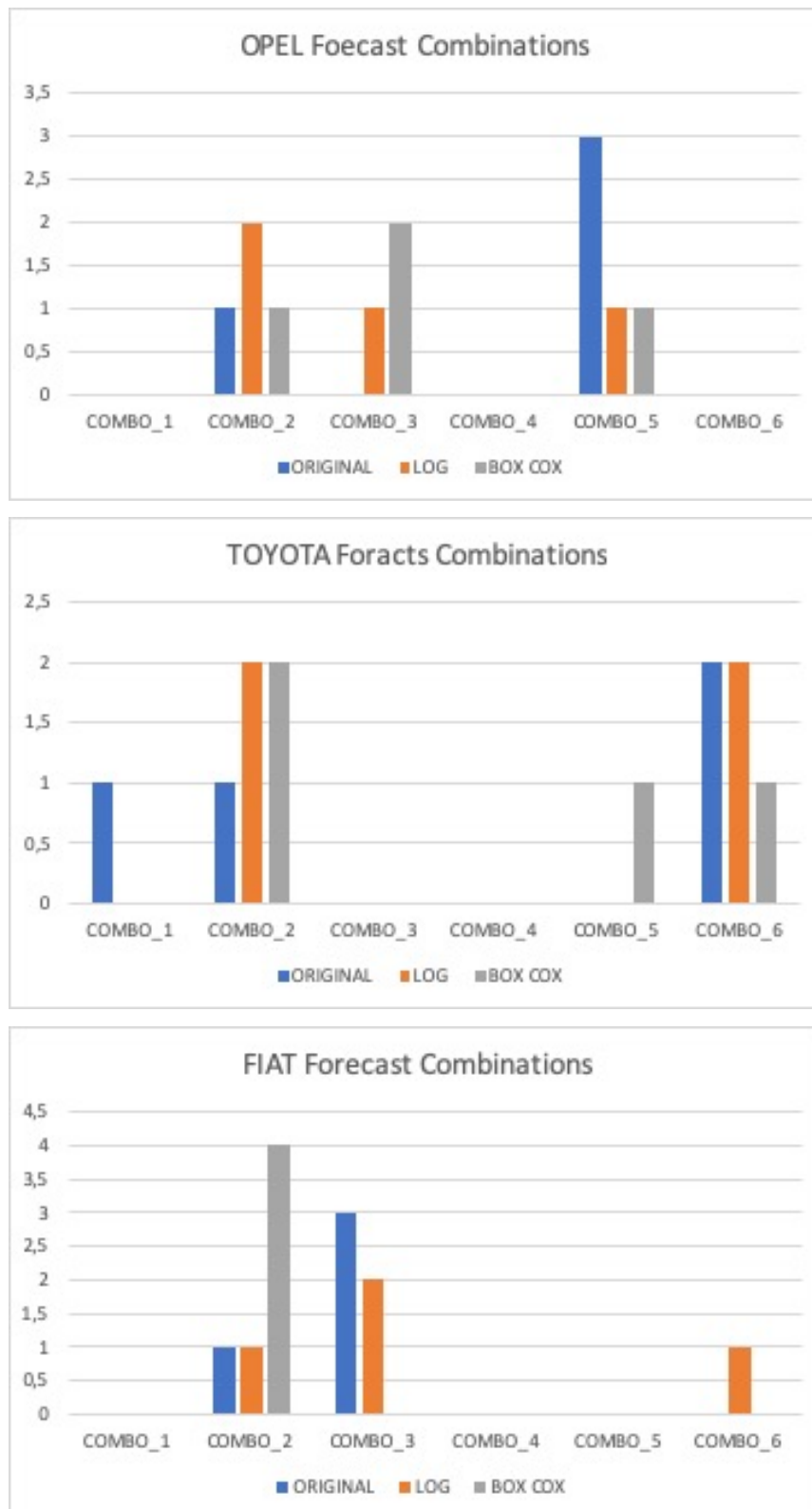


Figure 7.7: Forecast Combination for Opel, Toyota and Fiat new-car sales.

Table 7.1: Combination Forecast Performance for OPeL.

OPeL Forecasting Model	Original Values $\lambda = 1$			Log Values $\lambda = 0$			Box-Cox/Guerrero $\lambda = 0.36$		
Data Set A (h=48)	RMSE	MAE	MAPE	RMSE	MAE	MAPE	RMSE	MAE	MAPE
<i>Combo₁(SN + ETS)</i>	394.22	325.41	69.22	341.91	271.87	59.41	406.01	344.76	75.59
<i>Combo₂(ETS + LMSD)</i>	567.72	494.58	105.98	381.83	336.61	74.77	497.62	443.91	97.37
<i>Combo₃(SN + LMSD)</i>	357.83	289.55	61.62	354.18	283.40	61.81	388.38	326.18	71.36
<i>Combo₄(SN + ETS + LMSD)</i>	433.53	368.32	78.65	353.64	237.29	65.33	426.11	370.15	81.29
<i>Combo₅(Bates/Granger)</i>	213.70	141.72	27.90	242.20	197.83	41.62	450.30	376.92	84.49
<i>Combo₆(Newbold/Granger)</i>	333.28	229.51	41.02	360.67	273.09	59.52	399.77	296.72	61.52
Data Set B (h=24)	$\lambda = 1$			$\lambda = 0$			$\lambda = 0.03$		
<i>Combo₁(SN + ETS)</i>	242.24	181.95	40.31	192.09	147.94	33.11	212.92	162.55	36.96
<i>Combo₂(ETS + LMSD)</i>	364.44	293.34	63.63	206.23	159.52	36.98	263.11	207.92	47.29
<i>Combo₃(SN + LMSD)</i>	201.95	154.33	33.16	168.60	133.23	28.43	187.10	143.94	31.93
<i>Combo₄(SN + ETS + LMSD)</i>	263.64	199.87	44.21	185.61	143.46	32.26	217.91	166.12	37.97
<i>Combo₅(Bates/Granger)</i>	129.43	96.48	17.71	218.34	177.00	41.75	226.45	182.14	43.31
<i>Combo₆(Newbold/Granger)</i>	192.03	161.75	34.10	187.21	152.93	34.85	194.95	160.18	36.43
Data Set C (h=12)	$\lambda = 1$			$\lambda = 0$			$\lambda = 0.77$		
<i>Combo₁(SN + ETS)</i>	637.37	490.37	58.94	670.61	513.79	62.51	647.71	508.52	61.04
<i>Combo₂(ETS + LMSD)</i>	578.24	474.24	56.55	544.39	443.69	54.80	573.90	473.05	57.02
<i>Combo₃(SN + LMSD)</i>	662.45	486.03	59.25	677.71	504.92	62.00	665.85	493.91	60.75
<i>Combo₄(SN + ETS + LMSD)</i>	617.35	481.28	57.59	625.32	486.36	59.69	621.69	485.57	59.23
<i>Combo₅(Bates/Granger)</i>	802.52	654.90	90.15	801.43	669.63	93.68	809.32	670.17	93.11
<i>Combo₆(Newbold/Granger)</i>	1508.03	1140.48	152.36	1091.72	815.49	110.69	1195.74	866.87	117.98
Data Set D (h=24)	$\lambda = 1$			$\lambda = 0$			$\lambda = 0.86$		
<i>Combo₁(SN + ETS)</i>	604.40	495.33	56.93	626.52	513.16	59.28	605.58	500.18	57.95
<i>Combo₂(ETS + LMSD)</i>	567.27	485.02	56.66	497.05	444.19	52.49	558.77	482.88	56.88
<i>Combo₃(SN + LMSD)</i>	632.36	505.60	57.92	636.15	512.19	59.09	634.74	511.18	58.93
<i>Combo₄(SN + ETS + LMSD)</i>	590.69	495.30	57.17	577.38	489.84	56.95	589.48	498.08	57.92
<i>Combo₅(Bates/Granger)</i>	481.66	421.68	48.91	555.22	500.73	57.53	596.57	533.42	60.63
<i>Combo₆(Newbold/Granger)</i>	547.25	460.36	55.94	564.24	481.93	57.32	535.39	454.94	56.23

Table 7.2: Combination Forecast Performance for TOYOTA.

TOYOTA Forecasting Model	Original Values $\lambda = 1$			Log Values $\lambda = 0$			Box-Cox/Guerrero $\lambda = -0.14$		
	RMSE	MAE	MAPE	RMSE	MAE	MAPE	RMSE	MAE	MAPE
Data Set A (h=48)									
<i>Combo</i> ₁ (<i>SN</i> + <i>ETS</i>)	991.84	882.65	104.50	1812.86	1529.17	182.86	2158.27	1801.55	212.24
<i>Combo</i> ₂ (<i>ETS</i> + <i>LMSD</i>)	877.32	777.36	95.63	1370.78	1163.75	141.90	1529.64	1274.37	155.50
<i>Combo</i> ₃ (<i>SN</i> + <i>LMSD</i>)	1082.21	951.40	110.78	1736.94	1474.55	176.84	2047.87	1725.00	204.00
<i>Combo</i> ₄ (<i>SN</i> + <i>ETS</i> + <i>LMSD</i>)	978.68	870.01	102.62	1634.50	1382.20	166.98	1904.62	1588.76	190.19
<i>Combo</i> ₅ (<i>Bates/Granger</i>)	126.99	1051.09	134.78	1907.88	1782.91	226.57	2154.00	1998.78	252.42
<i>Combo</i> ₆ (<i>Newbold/Granger</i>)	984.65	900.04	116.37	1550.12	1399.60	182.91	1705.25	1516.62	197.09
Data Set B (h=24)	$\lambda = 1$			$\lambda = 0$			$\lambda = 0.51$		
<i>Combo</i> ₁ (<i>SN</i> + <i>ETS</i>)	959.17	835.86	87.31	962.86	834.73	87.33	961.63	840.44	88.21
<i>Combo</i> ₂ (<i>ETS</i> + <i>LMSD</i>)	915.87	835.98	90.19	825.70	755.94	81.37	878.35	805.73	87.11
<i>Combo</i> ₃ (<i>SN</i> + <i>LMSD</i>)	973.17	840.54	87.48	977.55	842.44	87.87	973.04	842.53	88.21
<i>Combo</i> ₄ (<i>SN</i> + <i>ETS</i> + <i>LMSD</i>)	942.21	837.25	88.31	915.90	810.52	85.47	931.55	829.57	87.85
<i>Combo</i> ₅ (<i>Bates/Granger</i>)	1003.51	925.80	98.18	917.32	844.75	90.06	972.23	897.34	95.35
<i>Combo</i> ₆ (<i>Newbold/Granger</i>)	1176.74	1083.97	113.73	974.38	824.81	86.10	1091.33	972.89	102.23
Data Set C (h=12)	$\lambda = 1$			$\lambda = 0$			$\lambda = -0.99$		
<i>Combo</i> ₁ (<i>SN</i> + <i>ETS</i>)	354.54	264.23	11.08	359.35	264.05	11.07	339.53	256.56	10.58
<i>Combo</i> ₂ (<i>ETS</i> + <i>LMSD</i>)	327.52	238.66	10.08	304.84	233.38	10.32	285.18	232.22	10.45
<i>Combo</i> ₃ (<i>SN</i> + <i>LMSD</i>)	404.33	315.35	13.97	336.17	245.73	10.08	346.38	262.56	10.89
<i>Combo</i> ₄ (<i>SN</i> + <i>ETS</i> + <i>LMSD</i>)	356.60	270.64	11.61	327.54	240.70	10.09	315.41	243.32	10.24
<i>Combo</i> ₅ (<i>Bates/Granger</i>)	302.30	262.69	10.30	323.13	295.54	11.69	190.45	164.08	8.81
<i>Combo</i> ₆ (<i>Newbold/Granger</i>)	203.96	190.35	12.30	255.22	218.84	12.42	268.14	231.26	13.40
Data Set D (h=24)	$\lambda = 1$			$\lambda = 0$			$\lambda = -0.18$		
<i>Combo</i> ₁ (<i>SN</i> + <i>ETS</i>)	574.04	479.08	20.14	483.96	393.35	17.33	389.91	307.53	14.04
<i>Combo</i> ₂ (<i>ETS</i> + <i>LMSD</i>)	597.89	524.54	22.97	647.79	561.84	24.75	592.53	499.04	22.22
<i>Combo</i> ₃ (<i>SN</i> + <i>LMSD</i>)	527.58	429.09	17.51	400.68	313.72	13.53	349.73	271.08	12.04
<i>Combo</i> ₄ (<i>SN</i> + <i>ETS</i> + <i>LMSD</i>)	563.14	474.74	20.02	501.71	415.34	18.14	427.93	342.84	15.29
<i>Combo</i> ₅ (<i>Bates/Granger</i>)	634.68	560.64	22.60	317.41	250.24	11.19	340.70	276.40	12.82
<i>Combo</i> ₆ (<i>Newbold/Granger</i>)	406.43	367.37	17.17	274.82	231.02	9.43	279.03	244.07	11.72

Table 7.3: Combination Forecast Performance for FIAT.

FIAT Forecasting Model	Original Values $\lambda = 1$			Log Values $\lambda = 0$			Box-Cox/Guerrero $\lambda = 0.10$		
Data Set A (h=48)	RMSE	MAE	MAPE	RMSE	MAE	MAPE	RMSE	MAE	MAPE
<i>Combo₁</i> (<i>SN</i> + <i>ETS</i>)	193.11	150.13	60.02	134.93	95.02	38.98	132.61	94.49	38.16
<i>Combo₂</i> (<i>ETS</i> + <i>LMSD</i>)	290.32	239.75	89.26	124.91	93.01	38.66	125.57	94.45	39.11
<i>Combo₃</i> (<i>SN</i> + <i>LMSD</i>)	151.61	116.17	46.07	135.19	93.55	37.36	136.91	95.99	39.37
<i>Combo₄</i> (<i>SN</i> + <i>ETS</i> + <i>LMSD</i>)	201.72	161.26	63.69	129.75	92.35	37.90	129.90	93.53	38.47
<i>Combo₅</i> (<i>Bates</i> /Granger)	180.77	136.22	54.34	162.90	116.91	52.99	163.15	117.23	53.55
<i>Combo₆</i> (<i>Newbold</i> /Granger)	187.52	138.39	53.01	157.89	113.47	47.38	163.14	117.26	53.57
Data Set B (h=24)	$\lambda = 1$			$\lambda = 0$			$\lambda = -0.08$		
<i>Combo₁</i> (<i>SN</i> + <i>ETS</i>)	159.88	126.73	56.70	142.55	106.25	50.08	136.02	100.75	46.80
<i>Combo₂</i> (<i>ETS</i> + <i>LMSD</i>)	208.22	168.93	74.14	138.89	104.71	50.43	129.29	96.41	46.04
<i>Combo₃</i> (<i>SN</i> + <i>LMSD</i>)	141.83	109.08	48.73	133.62	98.69	45.54	132.68	97.85	44.88
<i>Combo₄</i> (<i>SN</i> + <i>ETS</i> + <i>LMSD</i>)	166.28	133.30	59.58	136.80	102.42	48.47	131.45	98.10	45.84
<i>Combo₅</i> (<i>Bates</i> /Granger)	182.41	143.23	73.36	178.29	140.64	75.96	171.45	134.20	72.10
<i>Combo₆</i> (<i>Newbold</i> /Granger)	168.13	130.41	65.53	160.27	123.66	66.11	162.60	125.75	68.02
Data Set C (h=12)	$\lambda = 1$			$\lambda = 0$			$\lambda = 0.63$		
<i>Combo₁</i> (<i>SN</i> + <i>ETS</i>)	456.95	409.46	79.42	483.43	432.16	83.18	456.95	409.46	79.42
<i>Combo₂</i> (<i>ETS</i> + <i>LMSD</i>)	408.80	371.85	75.42	432.88	395.28	79.35	412.20	375.82	76.47
<i>Combo₃</i> (<i>SN</i> + <i>LMSD</i>)	453.86	405.22	77.82	475.94	425.48	81.67	457.35	409.19	78.80
<i>Combo₄</i> (<i>SN</i> + <i>ETS</i> + <i>LMSD</i>)	436.18	395.51	77.55	460.73	417.64	81.40	438.52	398.16	78.25
<i>Combo₅</i> (<i>Bates</i> /Granger)	529.07	504.03	125.58	563.65	532.92	132.15	533.12	508.75	127.13
<i>Combo₆</i> (<i>Newbold</i> /Granger)	699.26	680.46	181.16	435.75	428.41	109.66	656.62	673.08	159.35
Data Set D (h=24)	$\lambda = 1$			$\lambda = 0$			$\lambda = 0.07$		
<i>Combo₁</i> (<i>SN</i> + <i>ETS</i>)	479.17	437.09	90.75	592.25	534.50	112.86	541.81	489.89	103.23
<i>Combo₂</i> (<i>ETS</i> + <i>LMSD</i>)	480.10	442.10	94.40	570.01	520.10	112.43	517.03	474.44	102.81
<i>Combo₃</i> (<i>SN</i> + <i>LMSD</i>)	487.36	440.33	91.45	558.29	503.38	105.19	552.19	499.04	104.47
<i>Combo₄</i> (<i>SN</i> + <i>ETS</i> + <i>LMSD</i>)	476.13	439.15	92.14	568.13	518.58	110.10	532.20	487.79	103.50
<i>Combo₅</i> (<i>Bates</i> /Granger)	548.81	526.57	114.60	677.10	653.20	142.43	617.25	594.86	130.47
<i>Combo₆</i> (<i>Newbold</i> /Granger)	570.96	546.24	118.11	257.30	214.65	47.26	613.67	587.55	130.86

Bibliography

- Brock W. A., W.D. Dechert, J. A. Scheinkman, and B. LeBaron. A test for independence based on the correlation dimension. *Econometric Reviews*, 15:197–235, 1996.
- A. Abebe and G. Foerch. Stochastic simulation of the severity of hydrological drought. *Water Environ. Journal*, 22:2–10, 2008.
- R.R. Andrawis, A.F. Atiya, and H. El-Shishiny. Combination of long term and short term forecasts, with application to tourism demand forecasting. *International Journal of Forecasting*, 27(3):870–886, 2011.
- A. Andrikopoulos and R.N. Markellos. Dynamic interaction between markets for leasing and selling automobiles. *Journal of Banking and Finance*, 50:260–270, 2015.
- A. Antoniadis, E. Paparoditis, and T. Sapatinas. A functional wavelet-kernel approach for time series prediction. *Journal of the Royal Statistical Society. Series B (Statistical Methodology)*, 68(5):837–857, 2006.
- J.S. Armstrong. Forecasting with econometric methods folklore versus fact with discussion. *Journal of Business*, 51:549–600, 1978.
- J.S. Armstrong. *Long-Range Forecasting: Form Crystal Ball to Computer*. Wiley, New York, 2nd ed. edition, 1985.
- J.S. Armstrong. *Combining Forecasts. Principles of Forecasting: A Handbook for Researchers*

- and Practitioners*. J. Scott Armstrong (ed.): Norwell, ma: kluwer academic publishers edition, 2001.
- J.S. Armstrong and F. Collopy. Error measures for generalizing about forecasting methods: empirical comparison. *International Journal of Forecasting*, 08(1):69–80, 1992.
- A.C. Atkinson. *Plots, Transformations and Regression: An Introduction to Graphical Methods of Diagnostic Regression Analysis*. New York: Oxford University Press., 1985.
- D. Avramov. Stock returns predictability and model uncertainty. *Journal of Financial Economics*, 64(3):423–458, 2002.
- R.T. Baillie and T. Bollerslev. Common Stochastic Trends in a System of Exchange Rates. *Journal of Finance, American Finance Association*, 44(1):167–81, 1989.
- N. Barriera, P. Godinho, and P. Melo. Nowcasting unemployment rate and new car sales in south-western europe with google trends. *Netnomics*, 14:129–165, 2013.
- D. Batchelor and P. Dua. Forecaster diversity and the benefits of combining forecasts. *Management Science*, 45:68–75, 1995.
- J.M. Bates and C.W.J. Granger. The combination of forecasts. *Journal of Operational Research Society*, 20(4):451–468, 1969.
- C. Bergmeir, R.J. Hyndman, and J.M. Benítez. Bagging exponential smoothing methods using stl decomposition and box-cox transformation. *International Journal of Forecasting*, 32(2):303–312, 2016.
- J. Berkovec. New car sales and used car stocks: A model for the automobile market. *The RAND journal of Economics*, 2:195–214, 1985.
- T. Bollerslev. Generalised autoregressive conditional heteroskedasticity. *Journal of Econometrics*, 31:307–327, 1986.
- T. Bollerslev. On the Correlation Sstructure for the Generalized Autoregressive Conditional Heteroskedastic Process. *Journal of Time Series Analysis*, 9:121–131, 1988.

- T. Bollerslev, R.Y. Chou, and F.K. Kroner. Arch modeling in finance: a review of the theory and empirical evidence. *Journal of Econometrics*, 52(5):5–59, 1992.
- T. Bollerslev, R.F. Engle, and D.B. Nelson. *ARCH Model Handbook of econometrics, volume IV, chapter 49, edited by Engle and McFadden*. Elsevier Science, 1994.
- B.L. Bowerman, R.T. O’Connell, and A.B. Koehler. *Forecasting, Time Series and Regression: an applied approach*. Thomson Brooks/Cole, Belmont CA, 4 edition, 2005.
- G. E. P. Box and G. Jenkins. *Time series analysis: forecasting and control*. Holden-Day, San Francisco, 2nd edition, 1976.
- G.E.P. Box and D.R. Cox. An analysis of transformations. *Journal of the Royal Statistical Society, Series B*, 26:211–252, 1964.
- G.E.P. Box and G. Jenkins. *Time series analysis: forecasting and control*. Holden-Day, San Francisco, 1970.
- G.E.P. Box and D.A. Pierce. Distribution of residuals autocorrelations in autoregressive interated moving average time series models. *Journal of American Statistical Association*, 65:1509–1526, 1970.
- G.E.P. Box and G.C. Tiao. *Balayesian Inference in Statistical Analysis*. Addison-Wesley, Publishing Reading, M.A., 1973.
- H.J. Boyd and E.R. Mellman. The effect of fuel economy standards on the us automotive market : A hedonic demand analysis. *Transportation Research Part A: General*, 14: 367–378, 1980.
- L. Breiman. Bagging predictors. *Machine Learning*, 24:123–140, 1996.
- O.J.T. Briet, P. Vounatsou, D.M. Gunawardena, G.N.L. Galappaththy, and P.H. Amerasinghe. Models for short term malaria prediction in sri lanka. *Malaria Journal*, 7:76–76, 2008.

- P.J. Brockwell and R.A. Davis. *Time series: theory and methods*. Springer-Verlag, New York, 2nd edition, 1991.
- C. Brooks. *Introductory Econometrics for Finance*. Cambridge University Press, 2008.
- R.G. Brown. *Statistical Forecasting for inventory control*. New York, Mc Graw-Hill, 1959.
- C. Brownlees, R. Engle, and B. Kelly. A practical guide to volatility forecasting through calm and storm. *The Journal of Risk*, 14(2):3–22, 2011.
- P. Burman and R.H. Shumway. Generalized exponential predictors for time series forecasting. *Journal of the American Statistical Association*, 101(476):1598–1606, 2006.
- Y. Chang and M. Liao. A seasonal tourism demand modelling ARIMA Model of Tourism Forecasting: The case of Taiwan. *Asia Pacific Journal of Tourism*, 15(2):215–221, 2010.
- C. Chatfield. Model uncertainty and forecast accuracy. *Journal of Forecasting*, 15:495–508, 1996.
- C. Chatfield and M. Yar. Prediction intervals for multiplicative Holt-Winters. *International Journal of Forecasting*, 7:31–37, 1991.
- C. Chatfield, A.B. Koehler, J.K. Ord, and R.D. Snyder. A new look at models for exponential smoothing. *Journal of the Royal Statistical Society. Series D (The Statistician)*, 50(2):147–159, 2001.
- C.W.S. Chen and J.C. Lee. On selecting a power transformation in time series analysis. *Journal of Forecasting*, 16:343–354, 1997.
- G. Cheng and Y. Yang. Forecast combination with outlier protection. *International Journal of Forecasting*, 31(2):223–237, 2015.
- C. Christiansen, M. Schmeling, and A. Schrimpf. A comprehensive look at financial volatility prediction by economic variables. *Journal of Applied Econometrics*, 27(6):956–977, 2012.

- R. Clemen. Combining forecasts a review and annotated bibliography with discussion. *International Journal of Forecasting*, 5:559–608, 1989.
- M. P. Clement and D. F. Hendry. *Forecasting economic times series*. Cambridge University Press, 1998.
- J. Contreras, R. Espínola, F.J. Nogales, and A. J. Conejo. Arima models to predict next-day electricity prices. *IEEE transactions on power systems*, 18:1014–1020, 2003.
- J. Cragg. Estimation and Testing in Testing in Time Series Regression Models with Heteroskedastic Disturbances. *Journal of Econometrics*, 20:135–157, 1982.
- D.B. Crane and J.R. Crotty. A two-stage forecasting model: Exponential smoothing and multiple regression. *Management Science*, 13(8):501–507, 1967.
- J.D. Cryer and K.S. Chan. *Time series Analysis with Applications in R. Series:Springer Texts in Statistics*. Springer, New York, 2nd edition, 2008.
- D.B.Nelson. Conditional heteroscedasticity in asset returns: A new approach. *Econometrica*, 59:347–370, 1991.
- A.D. Dickey and A.W. Fuller. Distribution of the estimators for autoregressive time series with a unit root. *Econometrica*, 49:1057–72, 1981.
- N.R. Draper and D.R. Cox. On distributions and their transformation to normality. *Journal of the Royal Statistical Society, Series B(Methodological)*, 31(3):472–476, 1969.
- J. Durbin and J. Koopman. *Time Series Analysis by State Space Methods*. Oxford University Press, Oxford, 2001.
- O.F. Durdu. Stochastic approaches for time series forecasting of boron: A case study of western turkey. *Environ. Monitor. Assessment*, 169:687–701, 2010.
- V.S. Ediger and S. Akar. ARIMA forecasting of primary energy demand by fuel in turkey. *Energy Policy*, 35:1701–1020, 2007.

- V.S. Ediger, S. Akar, and B. Ugurlu. Forecasting production of fossil fuel sources in turkey using a comparative regression and ARIMA model. *Energy Policy*, 34(18):3836–3846, 2006.
- G. Elliott and A. Timmermann. Optimal forecast combinations under regime switching. *International Economic Review*, 46(4):1081–1102, 2005.
- G. Elliott and A. Timmermann. *Economic Forecasting*. Princeton University Press, 2016.
- G. Elliott, A. Gargano, and A. Timmermann. Complete subset regression. *Journal of Econometrics*, 2013.
- W. Enders. *Applied Econometric Time Series*. John Wiley & Sons Inc. , New York, 3 edition, 1995.
- F.R. Engle. Autoregressive conditional heteroskedasticity with estimates of variance of united kindom inflation. *Econometrica*, 50:987–1007, 1982.
- F.R. Engle. Estimates of the variance of u.s. inflation based on the arch model. *Journal of Money, Credit and Banking*, 15:286–301, 1983.
- F.R. Engle and A.J. Patton. What good is a volatility model? *Quantitative Finance*, 1: 237–245, 2001.
- R.F. Engle. Garch 101:an introduction to the use of arch/garch models in applied econometrics. Working paper FIN-01-030, NYU, 2001.
- M. Fernandes and J. Grammig. A family of autoregressive conditional duration models. *Journal of Econometrics*, 130:1–23, 2006.
- A. Garcia-Ferrer, J. Del Hoyo, and A.S. Martin-Arroyo. Univariate forecasting comparisons: The case of the spanish automobile industry. *Journal of Forecasting*, 16:1–17, 1997.
- E.S. Gardner. Exponential smoothing: the state of the art. *Journal of Forecasting*, 4:1–28, 1985.

- V. Genre, G. Kenny, A. Meyler, and A. Timmermann. Combining expert forecasts: Can anything beat the simple average? *International Journal of Forecasting*, 29(1):108–121, 2013.
- L.A. Gli-Alana. A fractionally integrated exponential model for uk unemployment. *Journal of Forecasting*, 20:329–340, 2001.
- V. Gómez and A. Maravall. Programs tramo and seats, instructions for the users. Technical report, Ministerio de Economía y Hacienda, Dirección General de Análisis y Programación Presupuestaria, 1998. Working paper 97001.
- S. Goncalves and N. Meddahi. Box-cox transforms for realized volatility. *Journal of Econometrics*, 160(1):129–144, 2011.
- P. Gonzalez and P. Moral. An analysis of the international tourism demand in Spain. *International Journal of Forecasting*, 11:233–251, 1995.
- R.L. Goodrich. The forecast pro methodology. *International Journal of Forecasting*, 16(4):533–535, 2000.
- P. Goodwin. The holt-winters approach to exponential smoothing: 50 years old and going strong foresight. *The International Journal of Applied forecasting*, 19:30–33, 2010.
- Ch. Gouriéroux and J. Jasiak. Nonlinear autocorrelograms: an application to inter-trade durations. *Journal of Time Series Analysis*, 23:127–154, 2002.
- A. Graefe, J.S. Armstrong, R.J. Jones, and A.G. Cuzan. Combining forecasts: an application to elections. *International Journal of Forecasting*, 30(1):43–55, 2014.
- C.W.J. Granger and A.P. Andersen. *An introduction to bilinear time series models*. Vandenhoeck and Ruprecht, Gottingen, 1978.
- C.W.J. Granger and P. Newbold. Forecasting transformed time series. *Journal of the Royal Statistical Society*, B(38):189–203, 1976.

- V.M. Guerrero. Time-series analysis supported by power transformations. *Journal of Forecasting*, 12:37–48, 1993.
- V.M. Guerrero and R. Parera. Variance stabilizing power transformation for tim series. *Journal of Modern Applied Statistical Methods*, 3(2):–, 2004.
- D.N. Gujarati. *Basic Econometrics*. McGraw Hill, 2003.
- P. Hall and Q.Yao. Inference in ARCH and GARCH Models with Heavy-Tailed Errors. *Econometrica*, 71(1):285–317, 2003.
- E.J. Hannan and J. Rissanen. Recursive estimation of mixed autoregressive- moving average order. *Biometrika*, 69:81–94, 1982.
- B.E. Hansen. Challenges for econometric model selection. *Econometric Theory*, 21(1): 60–68, 2005.
- B.E. Hansen. Least-squares model averaging. *Econometrica*, pages 1175–1187, 2007.
- B.E. Hansen. Least-squares forecast averaging. *Journal of Econometrics*, 146(2):342–350, 2008.
- B.E. Hansenn and J.R. Racine. Jackknife model averaging. *Journal of Econometrics*, 167 (1):38–46, 2012.
- P.J. Harrison and C.F. Stevens. Bayesian forecasting (with discussion). *Journal of the Royal Statistical Society Series, B*(38):205–247, 1976.
- A.C. Harvey. *The Econometric Analysis of Time Series*. Philip Allan, Deddington, Paperback, 1981.
- A.C. Harvey and R.G. Pierse. Estimating missing observations in economic time series. *Journal of the American Statistical Association*, 79(385), 1984.
- A.C. Harvey and P.H.J. Todd. Forecasting economic time series with structural and box-jenkins models (with discussion). *Journal of Business and Economic Statistics*, 1:299–315, 1983.

- M.I. Higgins and A. K. Bera. A class of nonlinear arch models. *International Economic Review*, 33:137–158, 1992.
- M.L. Higgins and A. K. Bera. A survey of arch models: Properties estimation and testing. *Journal of Economic Surveys*, 7:305–366, 1993.
- J. A. Hoeting, D. Madigan, A.E. Raftery, and C. T. Volinsky. Bayesian model averaging: A tutorial (with discussion). *Statistical Science*, 14:382–417, 1999.
- C. C. Holt. *Planning production, inventories and work force*. Englewood Cliffs, New Jersey: Prentice Hall, 1957a.
- C.E. Holt. Forecasting trends and seasonals by exponentially weighted averages. O.N.R. Memorandum 52/1957, Carnegie Institute of Technology, 1957b.
- W.S. Hopwood, J.C. McKeown, and P. Newbold. Power transformation in time series models of quarterly earnings per share. *The Accounting Review*, 56(4):927–933, 1981.
- C. Hsiao and S.K. Wan. Is there an optimal forecast combination? *Journal of Econometrics*, 178(2):294–309, 2014.
- D.A. Hsieh. Testing for non-linear dependence in daily foreign exchange rates. *Journal of Business*, 62(3):339–368, 1989.
- W.B. Hu, N. Nicholls, M. Lindsay, P. Dale, A.J. McMichael, and J.S. Mackenzie. Development of a predictive model for ross river virus disease in brisbane, australia. *Am. J. Tropical Med. Hygiene*, 71:129–137, 2004.
- M. Hülsmann, D. Borscheid, C.M. Friedrich, and D. Reith. General Sales forecast models for Automobile Markets and their Analysis. *Transactions on Machine Learning and Data Mining*, 5(2):65–86, 2012.
- R.J. Hyndman and G. Athanasopoulos. *Forecasting Principles and Practice*. John Wiley & Sons, New York, 2013.

- R.J. Hyndman and Y. Khandakar. Automatic time series forecasting: the forecast package for r. *Journal of Statistical Software*, 27(3), 2008.
- R.J. Hyndman, A.B. Koehler, R.D. Snyder, and S. Grose. A state space framework for automatic forecasting using exponential smoothing methods. *International Journal of Forecasting*, 18(3):439–454, 2002a.
- R.J. Hyndman, A.B. Koehler, R.D. Snyder, and S. Grose. A state space framework for automatic forecasting using exponential smoothing methods. *International Journal of Forecasting*, 18:439–454, 2002b.
- R.J. Hyndman, A.B. Koehler, J.K. Ord, and R.D. Snyder. Prediction intervals for exponential smoothing state space models. *Journal of Forecasting*, 24:17–37, 2005.
- M.Z. Ibrahim, R. Zailan, M. Ismail, and M.S. Lola. Forecasting and time series analysis of air pollutants in several area of malaysia. *Am. J. Environ. Sci.*, 5:625–632, 2009.
- C. Jarque and A. Bera. A test for normality of observations and regression residuals. *International Statistical Review*, 55:163–172, 1987.
- G. Johnes. Forecasting unemployment. *Applied Financial Economics*, 6:605–607, 1999.
- R.A. Johnson and D.W. Wichern. *Applied Multivariate Statistical Analysis*. Prentice –Hall 1992, 1992.
- R.E. Kalman. A new approach to linear filtering and prediction problems. *Journal of Basic Engineering*, 82(1):35–45, 1960.
- R.E. Kalman and R.S. Busy. New results in linear filtering and prediction theory. *Journal of Basic Engineering*, 83(1):95–108, 1961.
- G. Kapetanios, V. Labhard, and S. Price. Forecasting using bayesian and information-theoretic model averaging: an application to uk inflation. *Journal of Business and Economic Statistics*, 26(1):33–41, 2008.

- J.H. Kim. Forecasting Monthly Tourist Department from Australia. *Tourism Economics*, 5(3):277–291, 1999.
- G. Kitagawa and W. Gersch. A smoothness prior—state space modelling of time series with trend and seasonality. *Journal of the American Statistical Association*, 79:378–89, 1984.
- G. Kitagawa and W. Gersch. *Smoothness Priors Analysis of Time Series (Lecture Notes on Statistics)*. Springer-Verlag, New York, 1996.
- N. Kourentzes, D. Barrow, and F. Petropoulos. Another look at forecast selection and combination:evidence from forecast pooling. *International Journal of Production Economics*, 209:226–235, 2019.
- D. Kwiatkowski, P.C.B. Philips, P. Schmidt, and Y. Shin. Testing the null hypothesis of stationarity against the alternative of a unit root. *Journal of Econometrics*, 54:159–178, 1992.
- C. Lave and K. Train. A disaggregate model of auto –type choice. *Transportation Research Part A: General*, 13(1):1–9, 1979.
- Y. Li, E.P. Campbel, D. Haswell, R.J. Sneeuwjagt, and W.N. Venables. Statistical Forecasting of soil dryness index in the Southwest of Western Australia. *Forest Ecological Management*, 183:147–157, 2003.
- C. Lim and M. McAleer. A Seasonal Analysis of Malaysian Tourist Arrivals to Australia. *Mathematics and Computers in Simulation*, 48:573–583, 1999.
- C. Lim and M. McAleer. Monthly Seasonal Variations Asian Tourism to Australia. *Annals of Tourism Research*, 28(1):68–82, 2001.
- S. Ling and M. McAleer. Necessary and sufficient momment conditions for the garch(r,s) and asymmetric power garch(r,s) models. *Econometric Theory*, 18:722–729, 2002.

- L.M. Liu. Identification of Seasonal Arima Models Using a Filtering Method. *Communications in Statistics Theory and Methods*, 18:2279–2288, 1989.
- M.G. Ljung and E.P.G. Box. On a measure of lack of fit in time series models. *Biometrika*, 66:66–72, 1978.
- H. Lütkepohl and F. Xu. Forecasting annual inflation with seasonal monthly data: using levels versus logs of the underlying price index. *Journal of time Series Econometrics*, 3(1):–, 2011.
- H. Lütkepohl and F. Xu. The role of the log transformation in forecasting economic variables. Working paper, CESifo (Center for Economic Studies and Ifo Institute), 2012.
- J.R. Magnus and W. Wang. Concept-based bayesian model averaging and growth empirics. *Oxford Bulletin of Economics and Statistics*, 76(6):874–897, 2014.
- M. Makridakis and M. Hibon. The M3-competition: results, conclusions and implications. *International Journal of Forecasting*, 16:451–476, 2000.
- S. Makridakis. Accuracy measures: theoretical and practical concerns. *International Journal of Forecasting*, 9(4):527–529, 1993.
- S. Makridakis, A. Anderson, R. Carbone, R. Fildes, M. Hibon, R. Lewandowski and J. Newton, E. Parzen, and R. Winkler. The accuracy of extrapolation (time series) methods: results of a forecasting competition. *Journal of Forecasting*, 1:111–153, 1982.
- S. Makridakis, S. C. Wheelwright, and R. J. Hyndman. *Forecasting: methods and applications*. John Wiley & Sons, New York, 3rd edition, 1998.
- A.I. McLeod and K.W. Hipel. Diagnostic checking ARMA time series models using squared residual autocorrelations. *AGU Journal Water Resources Research*, 14(5):969–975, 1978.
- A.I. McLeod and W.K. Li. Diagnostic checking ARMA time series models using squared residual autocorrelations. *Journal of time Series Analysis*, 4:269–273, 1983.

- D.C. Medina, S.E. Findley, and D. Seydou. State–space forecasting of schistosoma haematobium time-series in niono. *Mali. www.plosntds.org*, 2(8):–, 2008.
- G. Mèlard and J.M. Pasteels. Automatic ARIMA Modeling Including intervention, using Time Series expert software. *International Journal of Forecasting*, 16:497–508, 2000.
- T. C. Mills. *The Econometric Modelling of Financial Time Series*. Cambridge University Press, 2 edition, 1999.
- A.K. Mishra and V.R. Desai. Drought forecasting using stochastic models. *Stochastic Environ. Res. Risk Assessment*, 19:326–339, 2005.
- R. Modarres. Streamflow drought time series forecasting. *Stochastic Environ. Res. Risk Assessment*, 21:223–233, 2007.
- P.E.N.M. Momani. Time series analysis model for rainfall data in jordan: Case study for using time series analysis. *Am. J. Environ. Sci.*, 5:599–604, 2009.
- C. Morana. Model averaging by stacking. Working paper, University of Milan, Bicocca, Department of Economics, Management and Statistics, 2015.
- E.A. Nanaki. Measuring the impact of economic crisis to the greek vehicle market. *Sustainability*, 10(2):510, 2018.
- L.H. Nelson and C.W.J. Granger. Experience with using the box-cox tranformation when forecasting economic time series. *Journal of Econometrics Pacific*, 10:57–69, 1979.
- P. Newbold and C.W.J. Granger. Experience with forecasting univariate time series and the combination of forecasts (with discussion). *Journal of the Royal Statistical Society (A)*, 137:131–165, 1974.
- J. Nowotarski, E. Raviv, S. Trück, and R. Weron. An empirical comparison of alternative schemes for combining electricity spot price forecasts. *Energy Economics*, 46:395–412, 2014.

- A. Opschoor, D.J. Van Dijk, and M. Van der Wel. Improving density forecast and value-at-risk estimates by combining densities. Technical report, Tinbergen Institution, 2014. Discussion Paper 14-090/III.
- J.K. Ord, A.B. Koehler, and R.D. Snyder. Estimation and prediction for a class of dynamic nonlinear statistical models. *Journal of American Statistical Association*, 92:1621–1629, 1997.
- K. Ord and S. Lowe. Automatic forecasting. *The American Statistician*, 50(1):88–94, 1996.
- J.W. Osborne. Sweating the small stuff in educational psychology: how effect size and power reporting failed to change from 1969 to 1999, and what that means for the future of changing practices. *Educational Psychology*, 28(2):1–10, 2008.
- J.W. Osborne. Improving your data transformations: Applying the box-cox transformation. practical assessment, research and evaluation. Technical Report 12, North Carolina State University, 2010.
- A. Pagan. The econometrics of financial markets. *Journal of Empirical Finance*, 3:15–102, 1996.
- F.C. Palm. *GARCH models of volatility*. Elsevier, North-Holland, in: maddala, g.s., rao, c.r. (eds.), handbook of statistics, vol. 14, pp. 209– 240. chap. 4 edition, 1996.
- F.C. Palm and P.J.G. Vlaar. Simple diagnostics procedures for modelling financial time series. *Allgemeines Statistisches Archiv*, 81:85–101, 1997.
- A. Pankratz and U. Dudley. Forecast of power-transformed series. *Journal of Forecasting*, 6(4):239–248, 1987.
- L. Pascual, J. Romo, and E. Ruiz. Bootstrap prediction intervals for power-transformed time series. *International Journal of Forecasting*, 21(2):219–235, 2005.
- C.C. Pegels. Exponential smoothing: some new variations. *Management Science*, 12:311–315, 1969.

- P.C.B. Philips and P. Perron. Testing for a unit root in time series regression. *Biometrika*, 75(2):335–346, 1988.
- S. Poon and C.W.J. Granger. Forecasting volatility in financial markets: a review. *Journal of Economic Literature*, 51(2):478–539, 2003.
- S. Poon and C.W.J. Granger. Practical issues in forecasting volatility. *Financial Analysts Journal*, 61(1):45–56, 2005.
- T. Proietti. Forecasting the us unemployment rate. Working paper, Dipartimento di Scienze Statistiche, Universita di Udine , Udine Italy, 2001.
- T. Proietti and H. Luetkepohl. Does the box-cox transformation help in forecasting macroeconomic time series? Technical report, Munich Personal RePEc Archive, 2011.
- T. Proietti and M. Riani. Transformations and seasonal adjustment. *Journal of Time Series Analysis*, 30(1), 2009.
- D.E. Rapach, J.K. Strauss, and G. Zhou. Out-of-sample equity premium prediction: combination forecasts and links to the real economy. *Review of Financial Studies*, 23(2):821–862, 2010.
- F. Ravazzolo, M. Verbeek, and H.K. VanDijk. Predictive gains from forecasting combinations using time varying model weights. Report 2007-26, Econometric Institute, 2007.
- E. Raviv, K.E. Bouwman, and D. VanDijk. Forecasting day-ahead electricity prices: utilizing hourly prices. *Energy Economics*, 50:227–239, 2015.
- D. Reilly. The autobox system. *International Journal of Forecasting*, 16(4):531–533, 2000.
- P. Rothman. Forecasting asymmetric unemployment rates. *Review of Economics and Statistics*, 80:104–168, 1998.
- P. Royston. An extension of shapiro and wilk w test for normality to large samples. *Applied Statistics*, 31:115–124, 1982.

- R.M. Sakia. The box-cox tranformation technique: A review. *The statistician*, 41:169–178, 1992.
- S.S. Shapiro and M.B. Wilk. An analysis of variance test for normality -complete samples. *Biometrika*, 52(34):591–611, 1965.
- R.H. Shumway and D.S. Stoffer. An approach to time series smoothing and forecasting using the em algorithm. *Journal Time Series Analysis*, 3:253–264, 1982.
- R.H. Shumway and D.S. Stoffer. *Time Series Analysis and its Applications with R Examples*. Springer Verlag, 2nd edition, 2006.
- E. Slutsky. The summation of random causes as the sources of cyclic processes. *Econometrica*, 5:105–46, 1937.
- J. Smith and K.F. Wallis. A simple explanation of the forecast combination puzzle. *Oxford Bulletin of Economic and Statistics*, 71(3):495–508, 2009.
- H. Song and K.K.F. Wong. Tourism demand modelling : a time varying parameter approach. *Journal of Travel Research*, 42:57–64, 2003.
- E.A. Stellwagen and R.L. Goodrich. Forecast pro estimation and prediction for a class of dynamic 4.0 manual. Technical report, Business Forecast Systems: Belmont, Massachusetts., 1999. Manual.
- J.H. Stock and M.W. Watson. Combination forecasting of output growth in seven-country data set. *Journal of Forecasting*, 23(6):405–430, 2004.
- J.W. Taylor. Exponential smoothing with a damped multiplicative trend. *International journal of Forecasting*, 19:715–725, 2003.
- J.W. Taylor. Multi-item sales forecasting with total and split exponential smoothing. *The Journal of the Operational Research Society*, 66(3):555—563, 2011.

- T.G.Andersen, T. Bollerslev, P.F.Christoffersen, and F.X. Diebold. *Volatility and correlation forecasting. In Handbook of Economic Forecasting.* Elliott, G., Granger, C. W. J., and Timmermann, A., north-holland, amsterdam edition, 2006.
- T. Thadewald and H. Buning. Jarque-bera test and its competitors for testing normality-A power comparison. *Journal of Applied Statistics*, 34(1):87–105, 2007.
- D.D. Thomakos and T. Wang. Realized volatility in the futures markets. *Journal of Empirical Finance*, 10:321–353, 2003.
- A. Timmermann. *Forecast combinations.* Handbook of economic forecasting, 2006.
- E. K. Train. *Discrete Choice Models With Simulation.* Cambridge University, New York, 2009.
- R.S. Tsay. Nonlinearity tests of Time Series. *Biometrika*, 73:461–466, 1986.
- R.S. Tsay. *Analysis of Financial Time Series.* John Wiley and Sons Inc. Publications, New Jersey, 2nd edition, 2005.
- R.S. Tsay. *Analysis of Financial Time Series.* John Wiley and Sons Inc. Publications, New Jersey, 3rd edition, 2010.
- G. Tsiotas. *Nonlinearities in stochastic Volatility Models (PhD thesis).* University of Essex, 2001.
- J.W. Tukey. The comparative mean after a fitted power transformations. *Annals of Mathematical Statistics*, 28:602–632, 1957.
- M. Wagner. Forecasting daily demand in cash supply chains. *American Journal of Economics and Business Administration*, 2:377–383, 2010.
- W. Wang, Van Gelder P.H.A.J.M., Vrijling J.K., and Ma J. Testing and modelling autoregressive conditional heteroskedasticity of steamflow processes. *Nonlinear Processes in Geophysics*, 12:55–66, 2005.

- Y. Wang and C. Lim. Using time series models to forecast tourist flows. *Proceedings of the 2005 International Conference on Simulation and Modeling*, www.mssanz.org.au, 2005.
- H.J. Wassenaar, W. Chen, J. Cheng, and A. Sudjianto. Enhancing discrete choice demand modeling for decision based design. *ASME Journal of Mechanical Design*, 127(4):214–523, 2005.
- C.E. Weiss. Hierarchical time series and forecast combination in healthcare demand forecasting. Working paper, Judge Business School, University of Cambridge, 2017.
- C.E. Weiss, E. Raviv, and G. Roetzer. Forecast combinations in R using the forecastcomb package. *The R Journal*, 10(2):262–281, 2018.
- C. Wen-hsien. *Econometrics and Time Series Analysis*. Best-Wise Inc, 2002.
- S.K. Whitefoot and J.S. Skerlos. Design incentives to increase vehicle size created from the u. s. footprint –based fuel economy standards. *Energy Policy*, 41(1):402–411, 2012.
- P. R. Winters. Forecasting sales by exponentially weighted moving averages. *Management Science*, 6:324–342, 1960.
- H. Word. *A study in the Analysis of Stationary Time Series (PhD thesis)*. Uppsala:Almqvist and Wiksell, 1938.
- J.H. Wright. Bayesian model averaging and exchange rates forecasts. *Journal of Econometrics*, 146(2):329–341, 2008.
- J.H. Wright. Forecasting US inflation by bayesian model averaging. *Journal of Forecasting*, 28(2):131–144, 2009.
- L. Yuchen. Auto car sales prediction. a statistical study using functional data analysis and time series. PhD, 2015.
- G.U. Yule. On a method of investigating the periodicities of disturbed series, with special reference to wolfer’s sunspot numbers. *Philosophical Transactions of the Royal Society*, A(226):226–298, 1927.

F. Yusof and I.L. Kane. Modeling monthly rainfall time series using ETS state space and SARIMA models. *International Journal of Current Research*, 4(09):195–200, 2012.

Appendix A

Appendix: New Car Sales Graphs

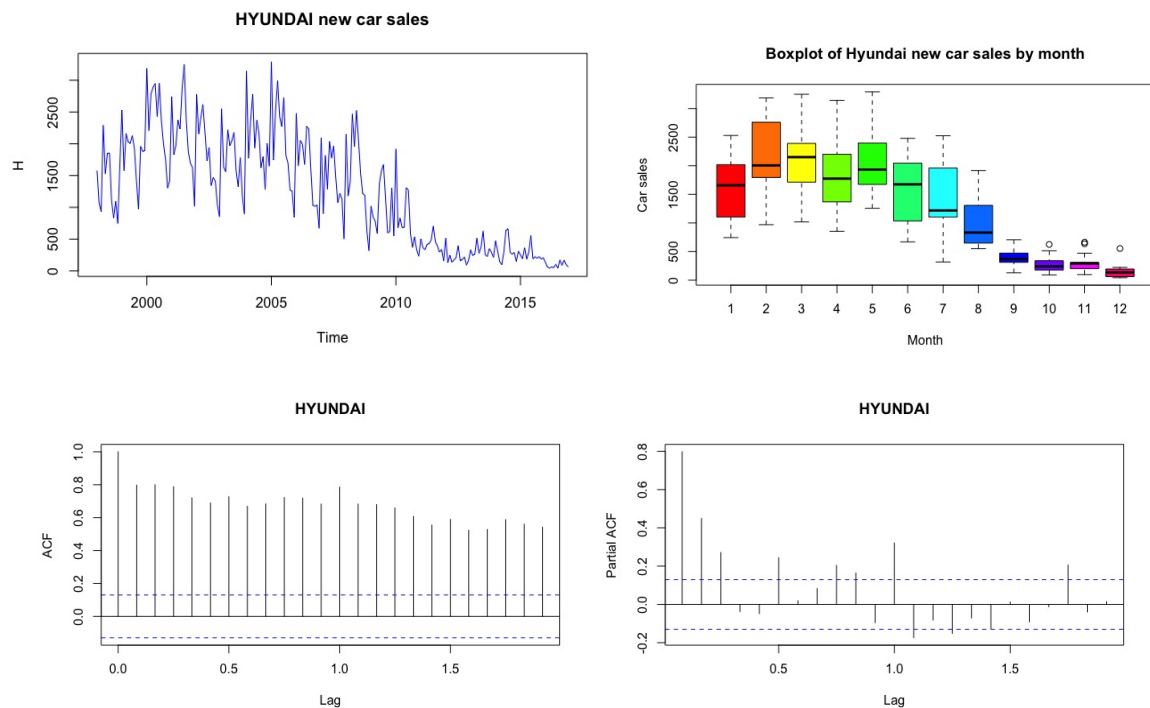


Figure A.1: **HYUNDAI** new-car sales (a)Line Plot, (b)Box Plot (c) ACF, (d)PACF (original values)

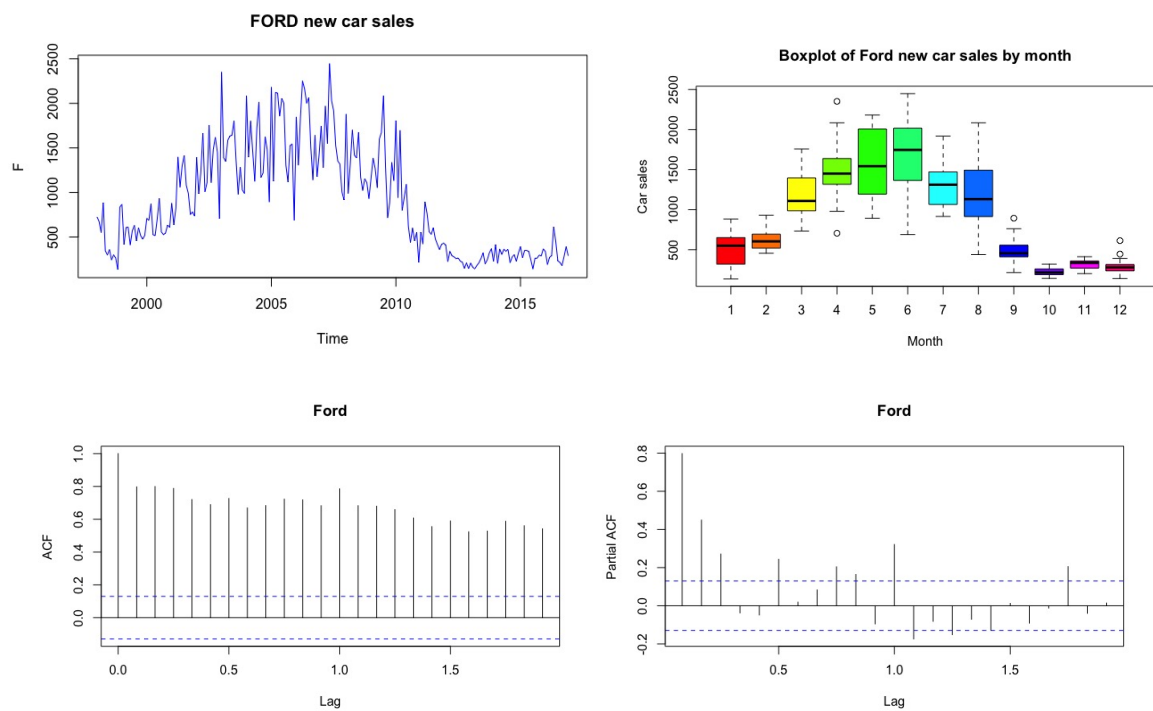


Figure A.2: **FORD** new-car sales (a)Line Plot, (b)Box Plot (c) ACF, (d)PACF (original values)

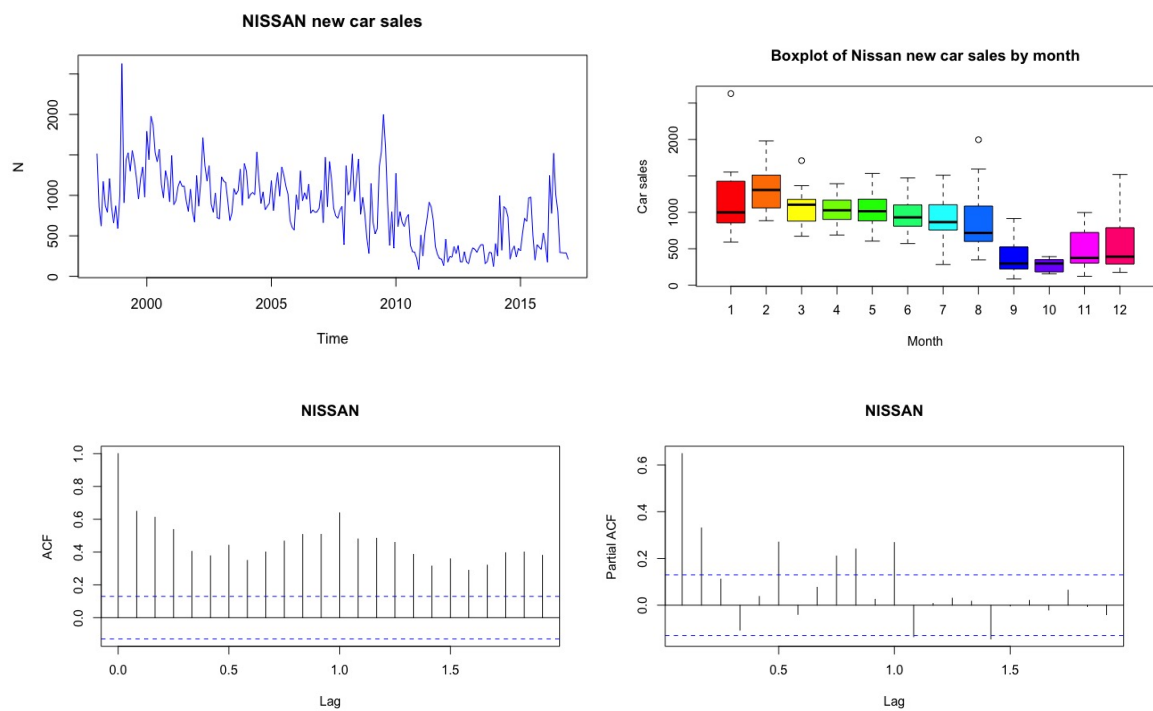


Figure A.3: **NISSAN** new-car sales (a)Line Plot, (b)Box Plot (c) ACF, (d)PACF (original values)

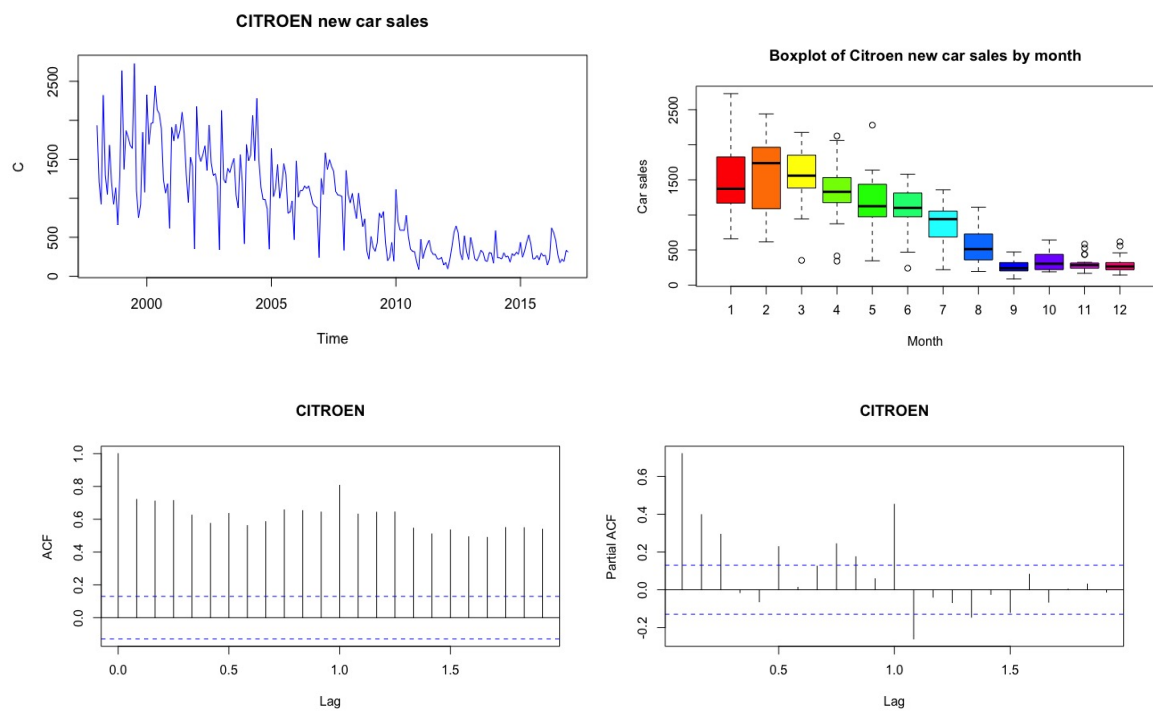


Figure A.4: **CITROEN** new-car sales (a)Line Plot, (b)Box Plot (c) ACF, (d)PACF (original values)

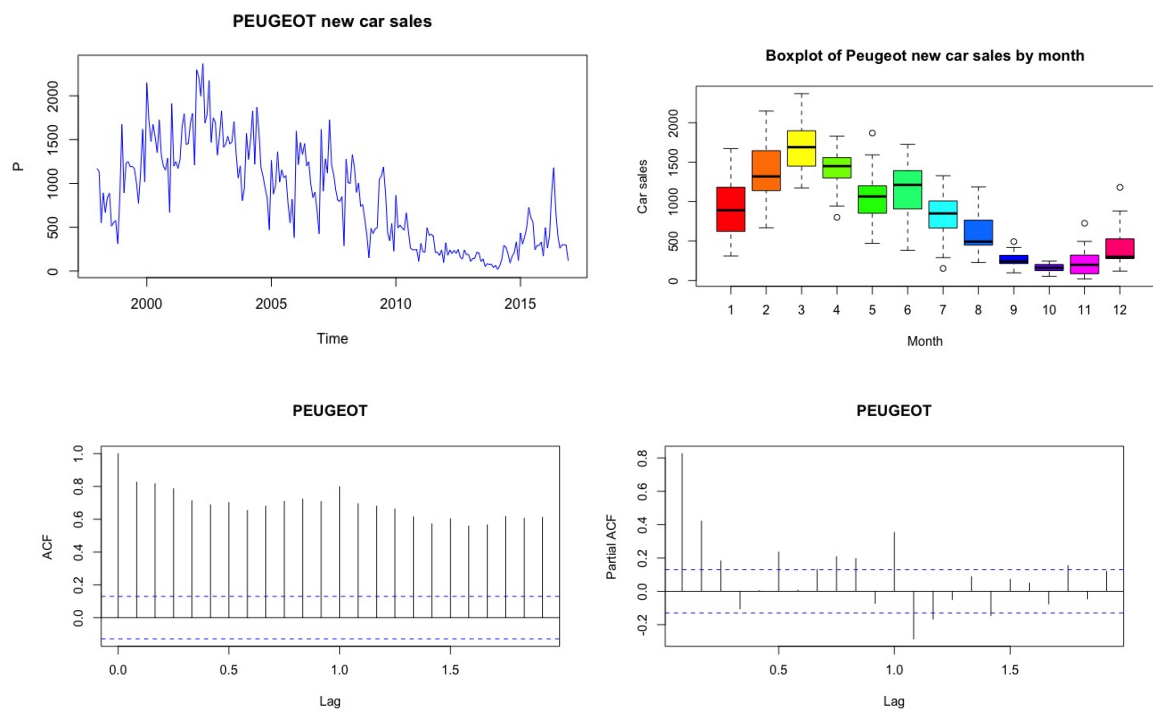


Figure A.5: **PEUGEOT** new-car sales (a)Line Plot, (b)Box Plot (c) ACF, (d)PACF (original values)

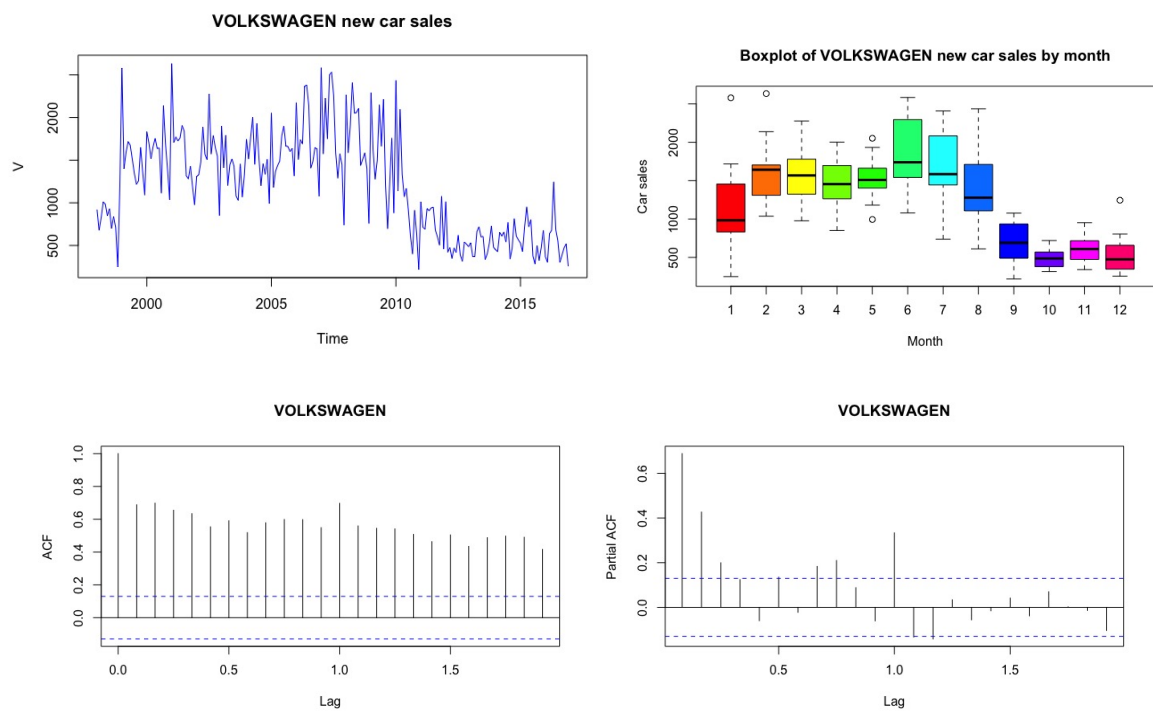


Figure A.6: **VOLKSWAGEN** new-car sales (a)Line Plot, (b)Box Plot (c) ACF, (d)PACF (original values)

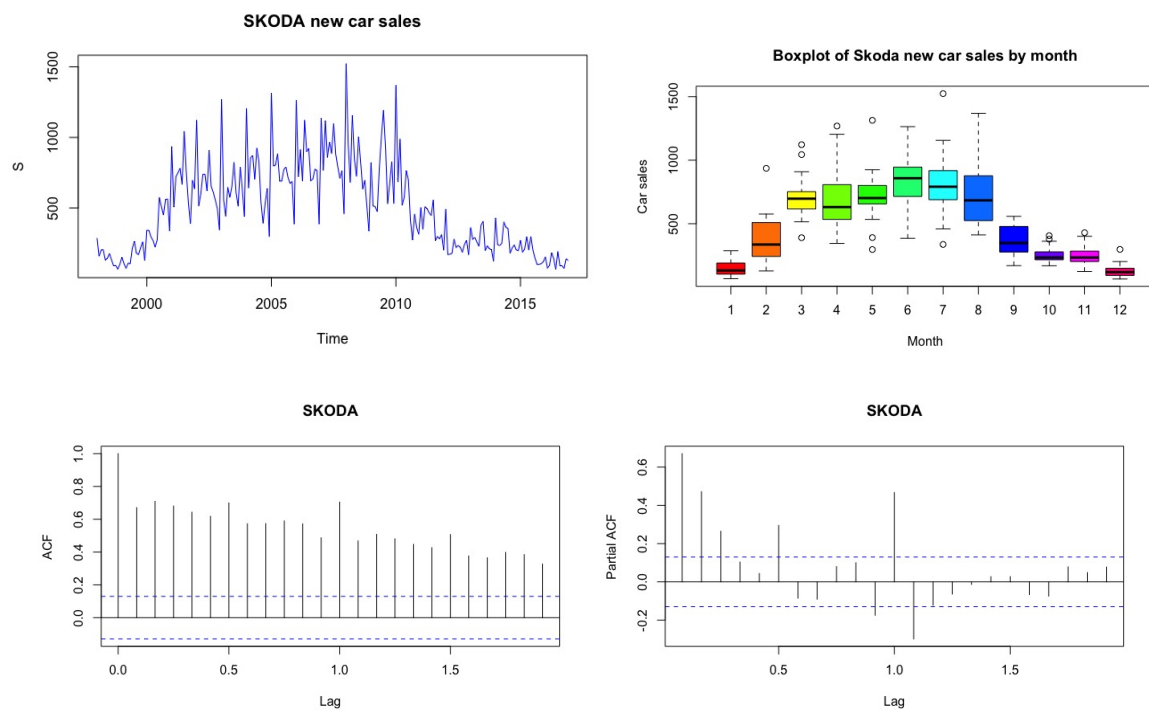


Figure A.7: **SKODA** new-car sales (a)Line Plot, (b)Box Plot (c) ACF, (d)PACF (original values)

Appendix **B**

Appendix: Autocorrelations in Data

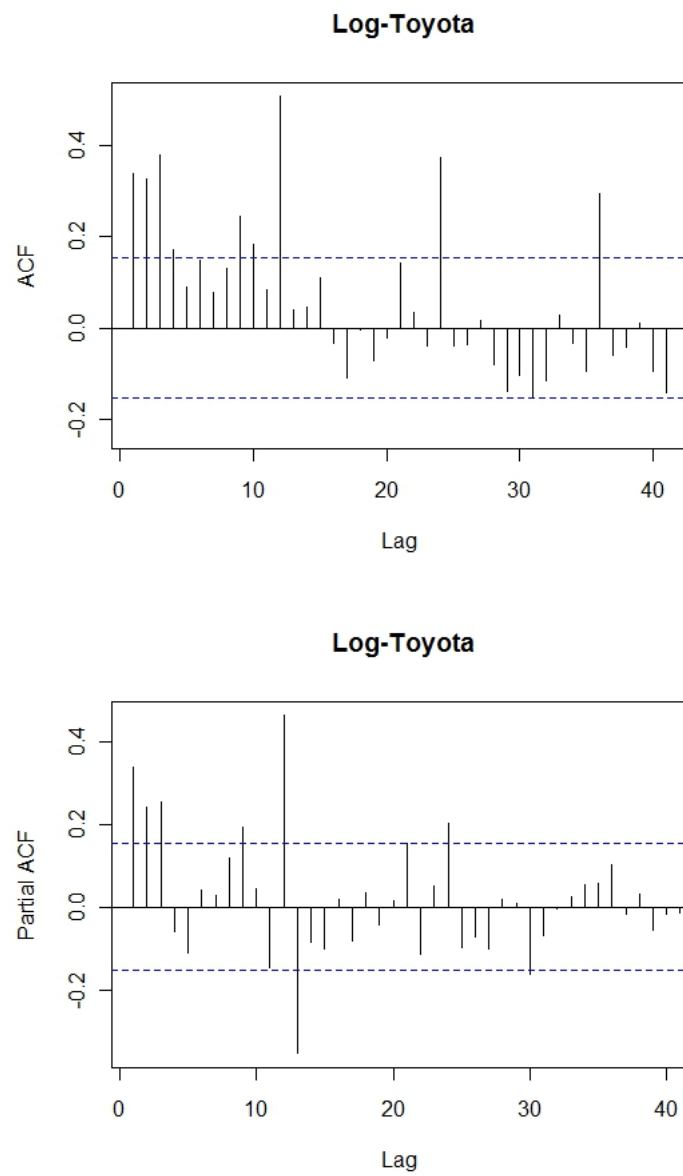


Figure B.1: ACF and PACF of the TOYOTA (logs)

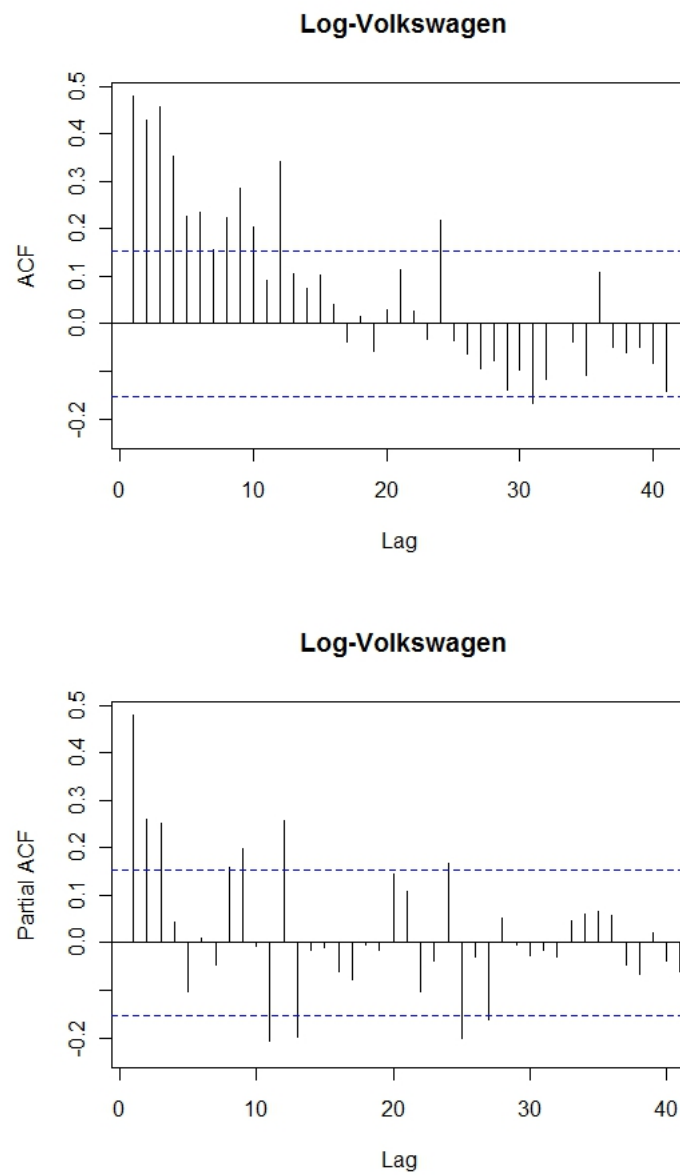


Figure B.2: ACF and PACF of the VOLKSWAGEN (logs)

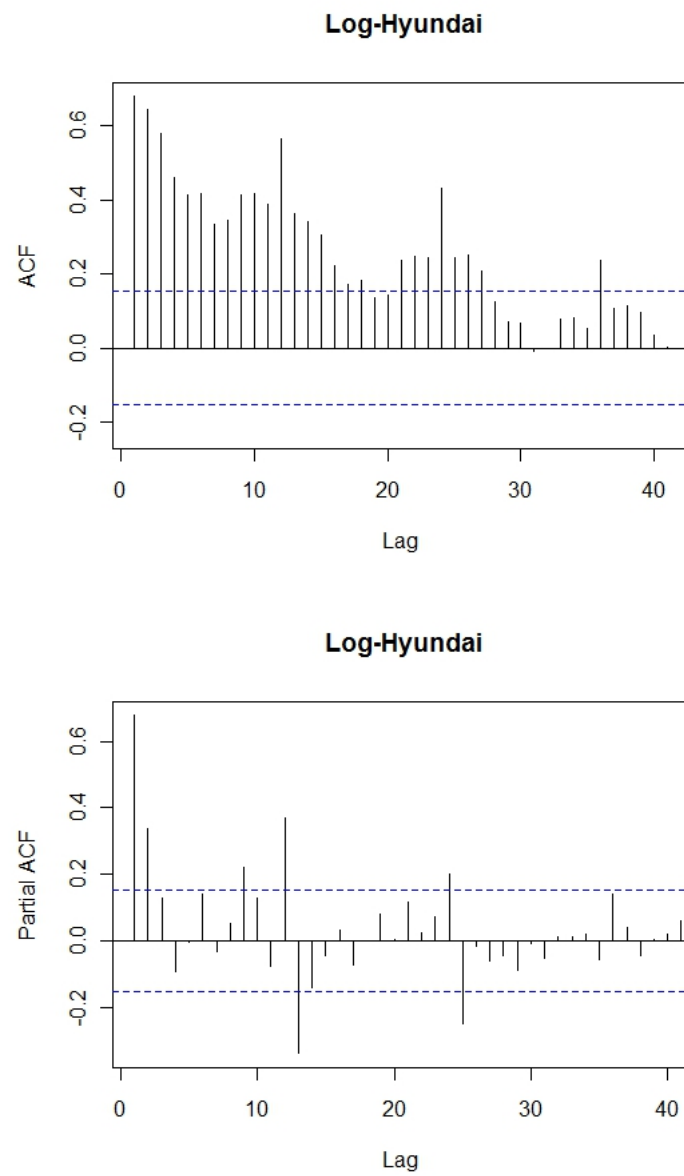


Figure B.3: ACF and PACF of the HYUNDAI (logs)

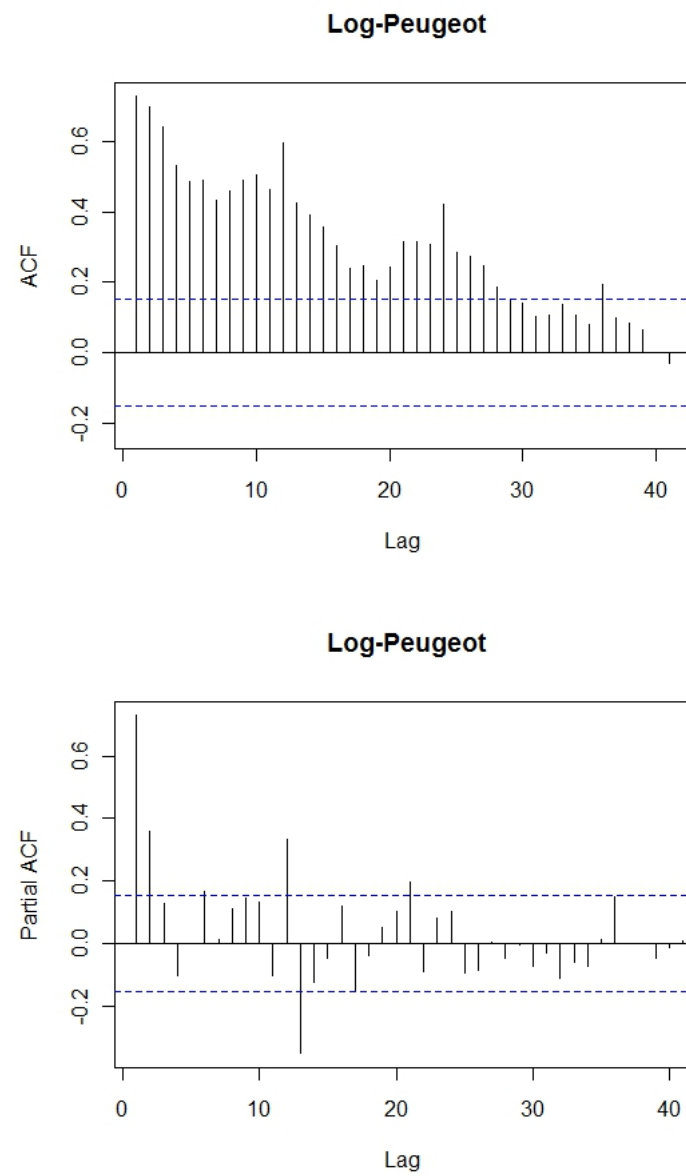


Figure B.4: ACF and PACF of the PEUGEOT (logs)

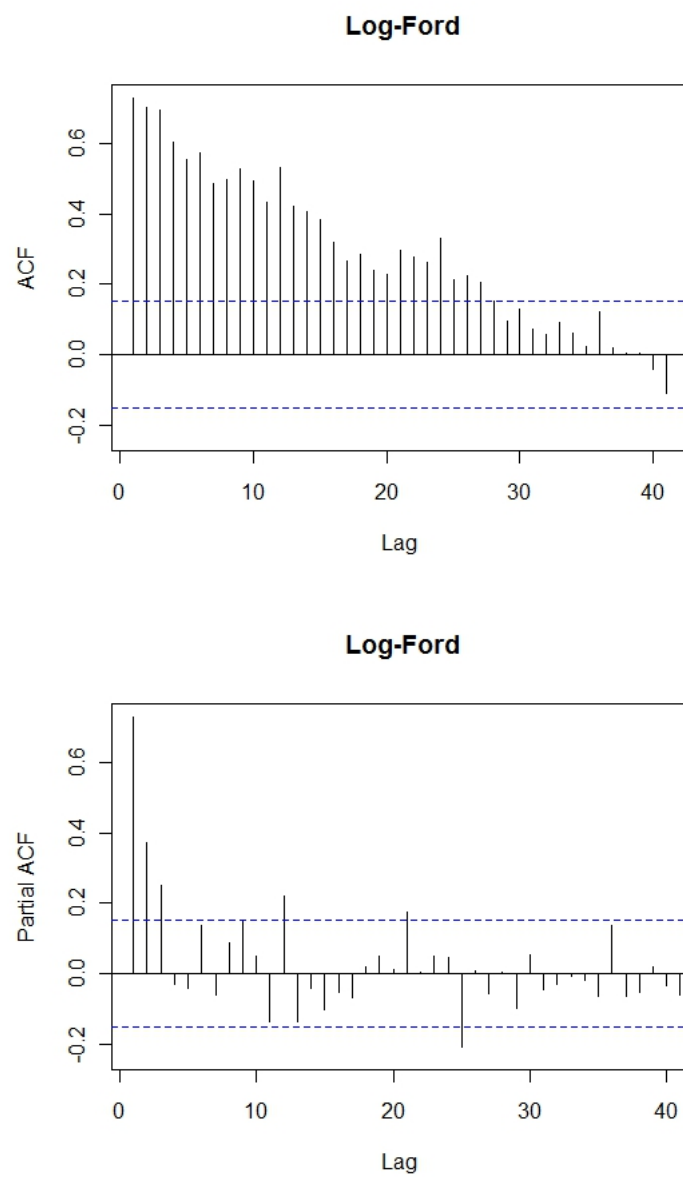


Figure B.5: ACF and PACF of the FORD (logs)

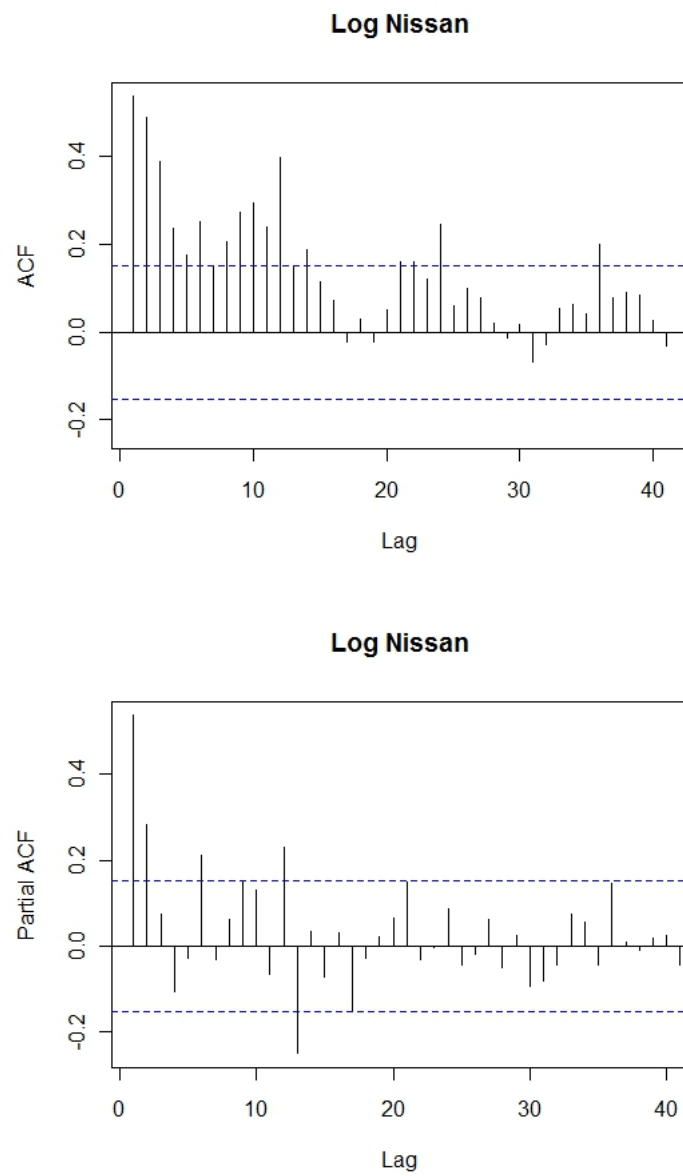


Figure B.6: ACF and PACF of the NISSAN (logs)

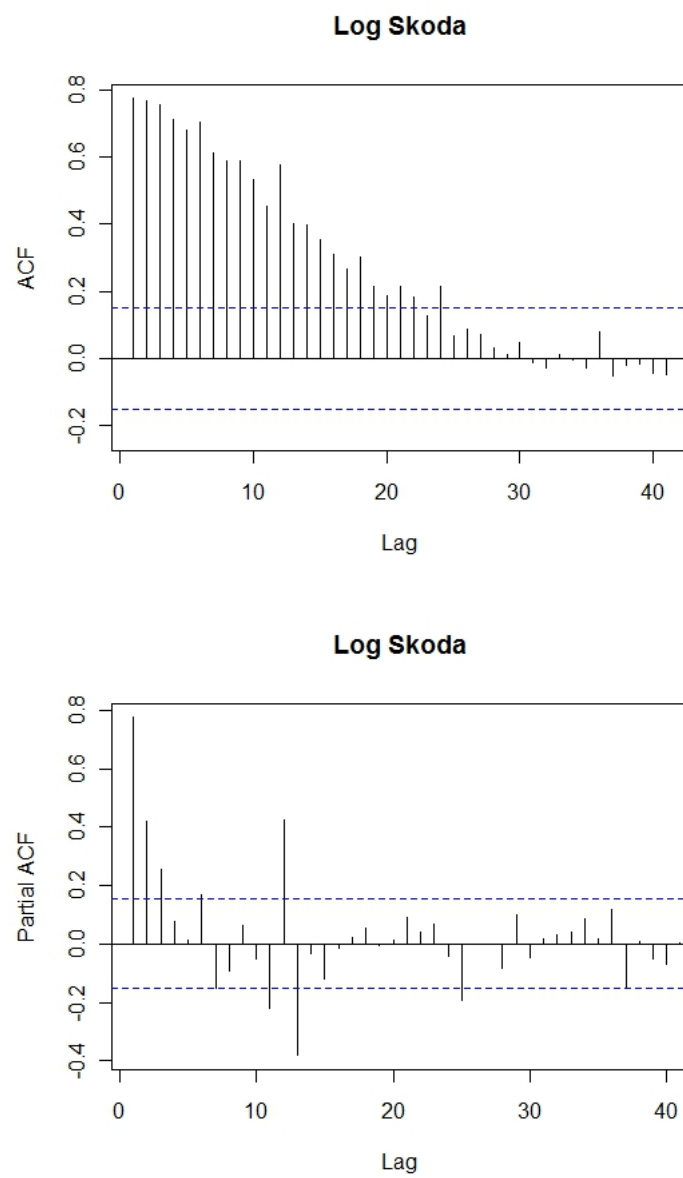


Figure B.7: ACF and PACF of the SKODA (logs)

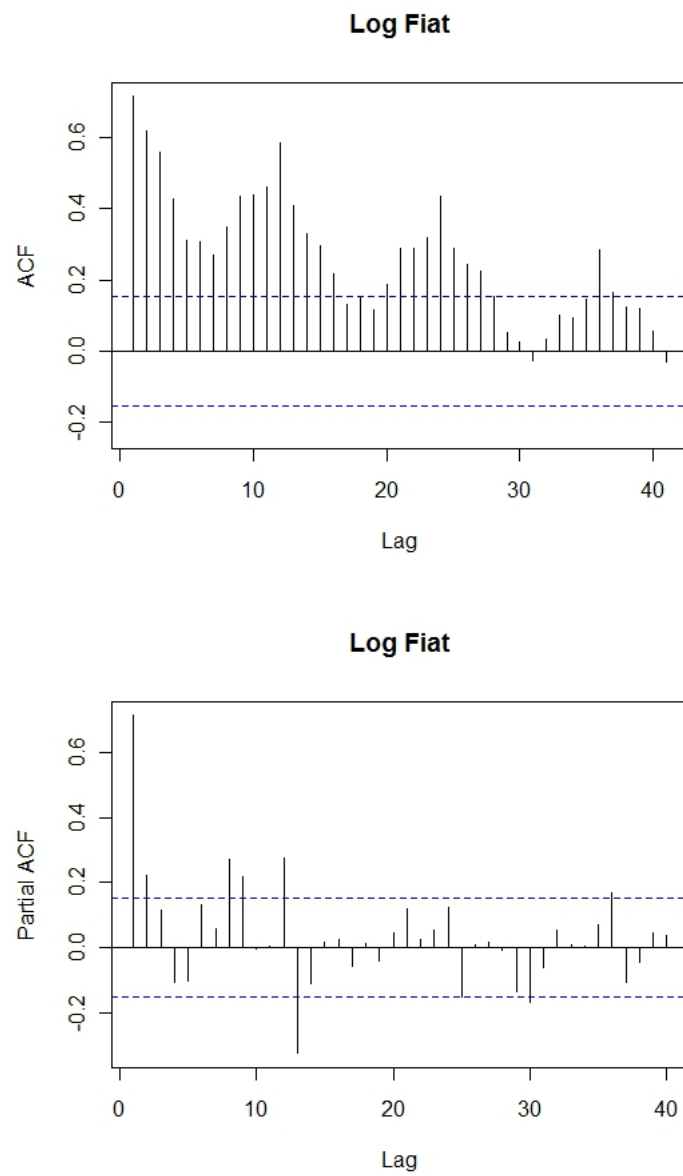


Figure B.8: ACF and PACF of the FIAT (logs)

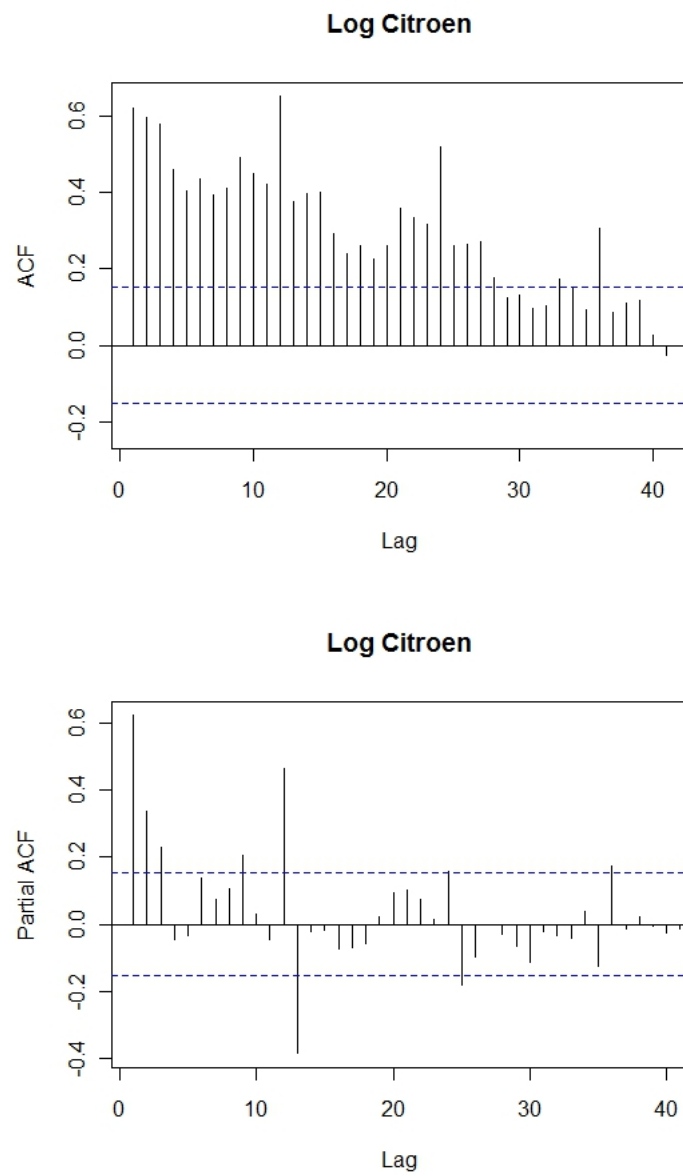


Figure B.9: ACF and PACF of the CITROEN (logs)

List of Abbreviations

- International Monetary Fund (IMF)
- European Central Bank (ECB)
- European Commission (EC)
- Association of Motor Vehicle Importers Representatives (AMVIR)
- Autocorrelation Function (ACF)
- Partial Autocorrelation Function (PACF)
- Autoregressive Intergrated Moving Average Models (ARIMA)
- Autoregressive (AR) models
- Moving Average (MA)
- Autoregressive Moving Average (ARMA)
- p (the number of autoregressive terms)
- q (the number of moving average terms)
- seasonal ARIMA (or SARIMA) models
- Exponential Smoothing State space models (ETS)

Υποδειγματοποίηση
Χρονολογικών Σειρών
Εμπειρική διερεύνηση της αγοράς
αυτοκινήτων στην Ελλάδα.

Μαρία Καλλέργου Βουλγαράκη

Σχολή Κοινωνικών Επιστημών
Τμήμα Οικονομικών Επιστημών

Πανεπιστήμιο Κρήτης

Εκτενής Περίληψη στα ελληνικά, διατριβής που υποβλήθηκε στα αγγλικά, για την απονομή
Διδακτορικού Διπλώματος στην Οικονομική Επιστήμη, στον τομέα Εφαρμοσμένης
Οικονομετρίας.

Σεπτέμβριος 2020

Ρέθυμνο - Κρήτης

Αυτή η διατριβή είναι αφιερωμένη με αγάπη και απέραντη ευγνωμοσύνη στην οικογένειά μου και ιδιαίτερα :

- * τον πατέρα μου, **Καλλέργο**, που με δίδαξε να δουλεύω άοκνα για να πραγματοποιώ τα όνειρά μου και πιστεύει σε μένα,
- * στη μητέρα μου, **Δέσποινα**, που με μεγάλωσε να σκέφτομαι θετικά, να μην τα παρατάω ποτέ, να ονειρεύομαι και να θέτω νέους στόχους στη ζωή,
- * στον αγαπημένο μου σύζυγο, **Γιάννη** και τα πολυαγαπημένα μας παιδιά **Στέλιο & Δέσποινα**, που είναι η καθημερινή πηγή άντλησης της δύναμης και της έμπνευσης μου.

Η υποστήριξη, η ενθάρρυνση και η άνευ όρων γενναιόδωρη αγάπη τους, αποδεικνύονται τα καλύτερα εφόδια στη ζωή μου.

Συνοπτική Περίληψη

(Abstract in Greek)

Ο σκοπός της παρούσας διδακτορική διατριβή είναι η πραγματοποίηση μιας ποιοτικής εμπειρικής έρευνας στην ανάλυση και υποδειγματοποίηση χρονολογικών σειρών στην ελληνική αγορά. Η ποσοτική εφαρμοσμένη έρευνα που αναπτύσσεται, αναζητά λύσεις και εξετάζει μεθόδους χρονολογικών σειρών, που μπορούν να εφαρμοστούν με επιτυχία σε δεδομένα μάρκετινγκ και να δώσουν αξιόπιστες προβλέψεις. Παρόλο που η πρακτική της πρόβλεψης δεδομένων μάρκετινγκ, όπως οι πωλήσεις, είναι μια ευρέως μελετημένη περιοχή, δεν έχει εφαρμοστεί εκτενώς στον ελληνικό τομέα πωλήσεων και συγκεκριμένα στον τομέα πωλήσεων αυτοκινήτων. Επομένως, αυτή η διατριβή αποτελεί μια πρωτότυπη προσπάθεια αναφοράς και συζήτησης διαφόρων νέων πρακτικών πρόβλεψης πωλήσεων, και αφορά πωλήσεις αυτοκινήτων στην ελληνική αγορά, ενώ εξετάζει μια ταραχώδη οικονομική χρονική περίοδο για την Ελλάδα, όπου η οικονομική δραστηριότητα ήταν υπό επιτήρηση σε ένα αυστηρά εποπτευόμενο περιβάλλον.

Ο σχεδιασμός αυτής της εμπειρικής, ερευνητικής, συγκριτικής μελέτης αφορά σε πρακτικές πρόβλεψης πωλήσεων, με τη χρήση μηνιαίων στοιχείων του ελληνικού τομέα πωλήσεων αυτοκινήτων, που διατίθενται από το Σύνδεσμο Εισαγωγέων Αντιπροσώπων Αυτοκινήτων (ΣΕΑΑ) της Ελλάδας. Ο μηνιαίος αριθμός ταξινόμησης νέων αυτοκινήτων υποθέτουμε ότι ισούται με το μηνιαίο επίπεδο πωλήσεων νέων αυτοκινήτων στην ελληνική αγορά. Συνεπώς με τη χρήση μονομετάβλητων υποδειγμάτων χρονολογικών σειρών γραμμικών και μη γραμμικών, με σταθερή ή κυμαινόμενη διακύμανση γίνεται μια προσπάθεια να μετρηθούν, να αξιολογηθούν και να προβλεφθούν τα επίπεδα πωλήσεων αυτοκινήτων στην Ελλάδα. Αναπτύσσεται για τα δεδομένα μια συγκριτική ανάλυση η οποία επισημαίνει ομοιότητες και διαφορές.

Η μεθοδολογία εύρεσης του καλύτερου μοντέλου χρησιμοποιεί μεθόδους υποδειματοποίησης και πρόβλεψη χρονολογικών σειρών εντός και εκτός δείγματος, οι οποίες εφαρμόζονται εμπειρικά σε μια ποικιλία εταιρειών πωλήσεων νέων αυτοκινήτων στην ελληνική λιανική αγορά, και σε διάφορα υποσύνολα δειγμάτων, κατά το χρονικό διάστημα των δύο τελευταίων δεκαετιών (1998 έως 2016). Γίνεται εμπειρική εφαρμογή σε μια ποικιλία υποδειγμάτων χρονολογικών σειρών ενώ εξετάζονται και τα πλεονεκτήματα του μετασχηματισμού των τιμών των μεταβλητών και της εφαρμογής συνδυασμού των προβλέψεων διαφορετικών υποδειγμάτων.

Στην εμπειρική έρευνα χρησιμοποιούνται απλά μοντέλα χρονολογικών σειρών, όπως το μέσο [Mean/Average], το Αφελές [Naïve] και το Εποχιακό Αφελές [Seasonal Naïve] αλλά και πιο εξελιγμένα, όπως τα Γραμμικά μοντέλα με εποχιακές ψευδομεταβλητές [Linear with Seasonal Dummies (LMSD)], τα μοντέλα Εκθετικής Εξομάλυνσης χώρου - κατάστασης, [Space Exponential Smoothing state space (ETS)] το εποχιακό αυτοπαλίνδρομο μοντέλο κινητού μέσου [SARIMA], και το εποχιακό αυτοπαλίνδρομο κινητού μέσου υπό συνθήκη ετεροσκεδαστικό SARIMA-GARCH καθώς και διάφοροι συνδυασμοί των προβλέψεων τους. Επιπλέον, διερευνάται η χρήση μετασχηματισμών δεδομένων τύπου Box-Cox, σε μια προσπάθεια βελτίωσης της ποιότητας και της απόδοσης των δεδομένων, καθώς και της προβλεπτικής ικανότητας των υποδειγμάτων, όπου αποδεικνύεται ότι παρέχει ένα ισχυρό εργαλείο για την ανάπτυξη τους. Τέλος, αναφέρεται και μελετάται η αξιοπιστία των προβλέψεων των υποδειγμάτων ενώ εξετάζεται και η προσέγγιση του συνδυασμού προβλέψεων από διάφορα μοντέλα χρονολογικών σειρών.

Τα εμπειρικά ευρήματα αυτής της μελέτης έδωσαν ενδείξεις βελτίωσης των προβλέψεων και των διαστημάτων εμπιστοσύνης, με τη χρήση της κατάλληλης διαδικασίας μετασχηματισμού των δεδομένων. Επιπλέον, υπάρχει ένα κυρίαρχο συμπέρασμα ότι δεν είναι δυνατόν να υπάρχει ενιαίο μοντέλο, για όλες τις εταιρείες, που να μπορεί πραγματικά να συλλάβει και να προβλέψει τις νέες πωλήσεις αυτοκινήτων στην ελληνική αγορά. Κάθε εταιρεία πρέπει να αντιμετωπίζεται ξεχωριστά σε συνάρτηση με το χρονικό διάστημα που εξετάζεται κάθε φορά. Παρόλα αυτά, ο μετασχηματισμός δεδομένων και οι μέθοδοι συνδυασμένης πρόβλεψης αποδεικνύονται επωφελείς για τη βελτίωση της ακρίβειας των προβλέψεων για όλες τις περιπτώσεις αυτής της έρευνας.

Το ξεχωριστό ενδιαφέρον αυτής της έρευνας είναι ότι πραγματοποιείται σε μια δύσκολη χρονική περίοδο για την ελληνική οικονομία. Κατά τη διάρκεια της ερευνητικής περιόδου, από το έτος 1998 έως τα τέλη του 2016 η Ελλάδα υπέγραψε τρία (3) Μνημόνια Συνεργασίας, με συγκεκριμένους όρους οικονομικής πολιτικής, που οδήγησαν στην οικονομική εποπτεία της χώρας, από την ομάδα αποφάσεων, που αναφέρεται ως ΤΡΟΙΚΑ και αποτελείται από το Διεθνές Νομισματικό Ταμείο (ΔΝΤ), την Ευρωπαϊκή Κεντρική Τράπεζα (ΕΚΤ) και την Ευρωπαϊκή Επιτροπή (ΕΕ). Η ελληνική κυβέρνηση αναγκάστηκε να εφαρμόσει μέτρα λιτότητας, που έπληξαν τους πολίτες και την οικονομική ζωή της χώρας, ως συνέπεια της εποπτείας αυτής. Οι Έλληνες καταναλωτές, λόγω της οικονομικής αβεβαιότητας, της έλλειψης ρευστότητας και της τραπεζικής κρίσης, παρατείνανε την αγορά σε διαρκή καταναλωτικά αγαθά (όπως αυτοκίνητα, έπιπλα κ.τ.λ.) με αποτέλεσμα την απότομη πτώση των νέων πωλήσεων αυτοκινήτων στην ελληνική αγορά.

Η πρωτοτυπία αυτής της διατριβής είναι ότι, για πρώτη φορά οι πωλήσεις αυτοκινήτων στον τομέα της ελληνικής αγοράς καινούργιων αυτοκινήτων αντιμετωπίζονται ως χρονολογικές σειρές με δεδομένα πωλήσεων των δυο τελευταίων δεκαετιών για την ελληνική αγορά. Επιπροσθέτως, η μελέτη αυτή αναδεικνύει τα προβλήματα, που δημιουργήθηκαν λόγω της οικονομικής κρίσης στον κλάδο πωλήσεων καινούργιων αυτοκινήτων την Ελλάδα, ενώ ανοίγει νέους ερευνητικούς ορίζοντες σε εφαρμογή οικονομετρικών υποδειγμάτων σε δεδομένα αγοράς στον ελλαδικό χώρο. Εν κατακλείδι, η *συμβολή στην επιστήμη* αυτής της διατριβής είναι στην καλύτερη κατανόηση της αγοράς αυτοκινήτων στην Ελλάδα και την μελέτη του καλύτερου τρόπου πρόβλεψης του επιπέδου πωλήσεων και στη βελτίωση της ποιότητας και της ακρίβειας του, ώστε να επιτευχθεί η αντιστάθμιση πιθανού κινδύνου και η μείωση της αποτυχία προβλέψεων, για τις ελληνικές εταιρείες νέων αυτοκινήτων. Οι μετασχηματισμοί των δεδομένων βελτιώνουν τα αποτελέσματα της έρευνας ενώ δεν υπάρχει ένα κοινό μοντέλο ως το καλύτερο για όλες τις χρονολογικές σειρές και για όλα τα χρονικά διαστήματα.¹

¹ Λέξεις Κλειδιά: Χρονολογικές Σειρές, Προβλέψεις, Μετασχηματισμός Box-Cox, Συνδυασμός Προβλέψεων, Ελληνική Αγορά, Πωλήσεις Αυτοκινήτων, Αφελές, Εποχιακό Αφελές, Γραμμικό με Εποχιακές Ψευδομεταβλητές, Εκθετικής Εξομάλυνσης Χώρου-Χρόνου, Εποχιακό Αυτοπαλίνδρομο Κινητού Μέσου υπό συνθήκη ετεροσκεδαστικό.

Περιεχόμενα

1	Εισαγωγικό Κεφάλαιο	3
1.1	Εισαγωγή	3
1.2	Λίστα δημοσιεύσεων & ερευνητικών παρουσιάσεων.	5
1.3	Συνεισφορά διατριβής στην επιστήμη.	7
2	Ανάλυση Χρονολογικών Σειρών	9
2.1	Εισαγωγή	9
2.2	Συμπεράσματα	11
3	Υποδειματοποίηση εντός-δείγματος	13
3.1	Εισαγωγή	13
3.2	Συμπεράσματα	15
4	Υποδειματοποίηση και Προβλέψεις εκτός-δείγματος	17
4.1	Εισαγωγή	17
4.2	Συμπεράσματα	20
5	Προβλέψεις με SARIMA-GARCH	25
5.1	Εισαγωγή	25
5.2	Συμπεράσματα	27
6	Μετασχηματισμοί δεδομένων και Προβλέψεις.	29
6.1	Εισαγωγή	29
6.2	Συμπεράσματα	30

7 Συνδιασμός Προβλέψεων.	33
7.1 Εισαγωγή	33
7.2 Συμπεράσματα	35

Εκτενής Περίληψη

(Extensive Summary in Greek)

Διευκρινιστικές παρατηρήσεις της περίληψης στα ελληνικά.

Η συνοπτική και εκτενής περίληψη που ακολουθεί είναι μια ελεύθερη μετάφραση από το αγγλικό κείμενο στην ελληνική γλώσσα σε μία προσπάθεια να αποδοθεί όσο το δυνατόν καλύτερα το περιεχόμενο της διδακτορικής διατριβής στην ελληνική γλώσσα. Η διατριβή τιτλοφορείται ως: “Υποδειματοποίηση Χρονολογικών Σειρών. Εμπειρική διερεύνηση της αγοράς αυτοκινήτων στην Ελλάδα (Modeling Time Series. The case of the Greek new-car sales sector)”. Επιπλέον στην περίληψη στα ελληνικά που ακολουθεί, δεν γίνονται βιβλιογραφικές αναφορές ούτε μαθηματικές αναλύσεις των υποδειγμάτων, που χρησιμοποιούνται, όμως για όλα αυτά ο αναγνώστης μπορεί να ανατρέξει στο αγγλικό κείμενο και στη λίστα βιβλιογραφικών αναφορών της διδακτορικής διατριβής. Επίσης οι αναφορές σε πίνακες δεδομένων ή σχεδιαγραμμάτων γίνεται μόνο όσο αναφορά στον αριθμό τους, ενώ δεν αναπαράγονται στο ελληνικό κείμενο εφόσον υπάρχουν σε εκτενή αναφορά στο αγγλικό περιεχόμενο και μπορούν εύκολα να μελετηθούν από εκεί για όποιον επιθυμεί περαιτέρω ενασχόληση ή εμβάθυνση.

Καλή Ανάγνωση.

Κεφάλαιο 1

Εισαγωγικό Κεφάλαιο

1.1 Εισαγωγή

Η ανάλυση χρονολογικών σειρών είναι ένα σημαντικό εργαλείο στη υποδειγματοποίηση και πρόβλεψη οικονομικών μεταβλητών. Διαφορετικές τεχνικές ανάλυσης χρονολογικών σειρών εφαρμόζονται σε αυτή τη διατριβή σε μια προσπάθεια αποδοτικότερης μέτρησης και αποτελεσματικότερης πρόβλεψης δεδομένων μάρκετινγκ για νέες πωλήσεις αυτοκινήτων στην ελληνική αγορά. Η έρευνά μας ξεκινά με την ανάλυση διαφόρων μεταβλητών χρονολογικών σειρών της ελληνικής αγοράς αυτοκινήτων και μελετούνται τα στοιχεία που φαίνεται να επηρεάζουν το επίπεδο των πωλήσεων αυτοκινήτων, κατά τη διάρκεια μιας ταραχώδους οικονομικής περιόδου για την ελληνική οικονομία (1998 – 2016). Το επίκεντρο της μελέτης μας είναι αρχικά η εφαρμογή διαφόρων απλών οικονομετρικών υποδειγμάτων, όπως το μέσο, το αφελές και το εποχιακά αφελές υπόδειγμα και η έρευνα συνεχίζεται με τις μεθόδους υποδειγματοποίησης κατά Box-Jenkins που χρησιμοποιούνται στα αυτοπαλίνδρομα μοντέλα κινητού μέσου (ARIMA) και τα εποχιακά αυτοπαλίνδρομα κινητού μέσου SARIMA, το γραμμικό υπόδειγμα με εποχικές ψευδομεταβλητές (LMSD), τα υποδείγματα εκθετικής εξομάλυνσης χώρου χρόνου (ETS) και την οικογένεια γενικευμένων υποδειγμάτων αυτοπαλίνδρομου κινητού μέσου υπό συνθήκη ετεροσκεδαστικού (GARCH).

Τα εμπειρικά στοιχεία αυτής της έρευνας δείχνουν ότι τα απλά υποδείγματα χρονολογικών σειρών, όπως τα εποχιακά αφελές υποδείγματα μπορεί να είναι επαρκή για βραχυπρόθεσμη

πρόβλεψη, ενώ πιο περίπλοκα υποδείγματα, όπως τα μοντέλα εκθετικής εξομάλυνσης χώρου χρόνου (ETS), τα γραμμικά μοντέλα με εποχικές ψευδομεταβλητές (LMSD) και τα εποχιακά αυτοπαλίνδρομα μοντέλα κινητού μέσου (SARIMA) είναι καλύτερα για μακροπρόθεσμες προβλέψεις.

Επιπλέον, μελετάμε την ακρίβεια στην υποδειματοποίησης χρονολογικών σειρών, όταν πραγματοποιούμε μετατροπή στα αρχικά δεδομένα. Έτσι, αυτή η έρευνα παρουσιάζει αποτελέσματα για τρεις διαφορετικές περιπτώσεις κατά τη υποδειματοποίησης των μεταβλητών α) σε αρχικές τιμές, β) σε λογαριθμημένες τιμές και γ) τιμές με τη χρήση μετασχηματισμού Box-Cox. Επιπροσθέτως, τονίζεται η σημασία του συνδυασμού διαφορετικών υποδειγμάτων χρονολογικών σειρών για την πραγματοποίηση προβλέψεων. Έτσι, αντί να βασιζόμαστε σε ένα μόνο υπόδειγμα για τις προβλέψεις, χρησιμοποιούμε ένα συνδυασμό υποδειγμάτων και με τον τρόπο αυτό μειώνουμε τον κίνδυνο χρησιμοποιώντας όλες τις διαθέσιμες πληροφορίες. Αυτό γίνεται είτε υποθέτοντας μια ενιαία στάθμιση σε κάθε πρόβλεψη είτε με τη χρήση συντελεστών στάθμισης ανάλογα με την ακρίβεια της πρόβλεψης του κάθε υποδείγματος. Η τεχνική της συνδυασμένης πρόβλεψης ξεπερνά τις προβλέψεις κάθε ενός από τα μεμονωμένα υποδείγματα σε αυτήν την εμπειρική μελέτη.

Αυτή η διατριβή καλύπτει με την έρευνα της χρονικά ένα δύσκολο διάστημα για την ελληνική οικονομία στις αρχές του 21ου αιώνα, από το 1998 έως το έτος 2016. Ως εκ τούτου, αυτή η μελέτη ενδιαφέρεται παράλληλα για την οικονομική κατάσταση στην ελληνική αγορά, και προσπαθεί να αναδείξει έμμεσα, ότι η οικονομική κρίση και τα μέτρα λιτότητας, που τέθηκαν σε ισχύ μετά την εφαρμογή των Μνημονίων Συνεργασίας της ελληνικής κυβέρνησης με την Ευρωπαϊκή Επιτροπή (ΕΚ), την Ευρωπαϊκή Κεντρική Τράπεζα (ΕΚΤ) και το Διεθνές Νομισματικό Ταμείο (ΔΝΤ)), άλλαξαν την οικονομική δραστηριότητα στον τομέα λιανικής πώλησης αυτοκινήτων στην Ελλάδα, αντικατοπτρίζοντας τις παρατεταμένες οικονομικές δυσκολίες της ελληνικής αγοράς.

Αναλυτικότερα, το πρώτο κεφάλαιο της διατριβής δίνει μια σύντομη εισαγωγή στο θέμα, ακολουθούμενο από τους σκοπούς της έρευνας, ενώ θέτει τα ερευνητικά ερωτήματα που προβληματίσαν την ερευνήτρια κατά την περίοδο της μελέτης. Ο σκοπός αυτής της μελέτης και οι ερευνητικοί στόχοι δηλώνονται με σαφήνεια και θα αναπτυχθούν πλήρως στα επόμενα κε-

φάλαια αυτής της διατριβής, εξηγούνται τα κίνητρα της έρευνας, ενώ οι ερευνητικές ερωτήσεις θα απαντηθούν επιστημονικά και αποτελεσματικά βάσει της εφαρμοσμένης οικονομετρικής θεωρίας και της πρακτικής που χρησιμοποιείται σε δεδομένα χρονολογικών σειρών. Αναφέρουμε την πηγή των δεδομένων, που είναι ο δικτυακός τόπος του Συνδέσμου Εισαγωγέων αντιπροσώπων αυτοκινήτων και δικύκλων (<https://seaa.gr>), το λογισμικό R, που χρησιμοποιήθηκε για την επεξεργασία των δεδομένων, γραφικά, στατιστική ανάλυση υποδειγματοποίηση και τεστ, και το πρόγραμμα L^AT_EX που χρησιμοποιήθηκε για την επιστημονική σύνταξη της παρούσας διατριβής.

Η επισκόπηση της διατριβής βοηθά στην κατανόηση της δομής της. Η ερευνητική συμβολή και τα τελικά συμπεράσματα αυτής της διατριβής μπορούν να δώσουν χρήσιμες πληροφορίες, οι οποίες μπορούν να βοηθήσουν τους υπεύθυνους λήψης αποφάσεων να καταρτίσουν στρατηγικές για την ορθή διαχείριση των αποθεμάτων, για εφαρμογές ανθρώπινου δυναμικού και για την διαχείριση της εφοδιαστικής αλυσίδας στον τομέα του λιανικού εμπορίου νέων αυτοκινήτων στην Ελλάδα. Επιπλέον, συμβάλλει σημαντικά στην οικονομική επιστήμη, λόγω της εμφανής έλλειψης εμπειριστατωμένης έρευνας σε αυτόν τον τομέα της ελληνικής αγοράς και της εφαρμογής συγκριτικής μελέτης χρησιμοποιώντας το θεωρητικό υπόβαθρο και τις μεθόδους χρονολογικών σειρών.

1.2 Λίστα δημοσιεύσεων & ερευνητικών παρουσιάσεων.

Είμαι ευγνώμων και αισθάνομαι ιδιαίτερα τυχερή, που μου δόθηκε η ευκαιρία να παρουσιάσω μέρος της έρευνας μου κατά τη διάρκεια υλοποίησης της, σε συνέδρια, συμπόσια και συναντήσεις και να λάβω πολύτιμα και εποικοδομητικά σχόλια για τη συνέχιση της έρευνάς μου. Ιδιαίτερες ευχαριστίες στο Ίδρυμα Α. Γ. Λεβέντη (A.G Leventis Foundation) για την ευγενική χορηγία του και την κάλυψη μέρους των εξόδων μου, για τη συμμετοχή μου σε δύο Συμπόσια για διδακτορικούς φοιτητές Ελλάδας και Κύπρου, που διοργάνωσε το Ελληνικό Παρατηρητήριο του Ευρωπαϊκού Ινστιτούτου του Πανεπιστημίου London School of Economics and Political Science- LSE , στο Λονδίνο, της Αγγλίας.

Ο κατάλογος των παρουσιάσεων ή δημοσιεύσεων, που προέκυψαν κατά τη διάρκεια αυτής της διδακτορικής διατριβής με χρονολογική σειρά είναι οι εξής:

- Voulgaraki K. Maria (2019), Box Cox transformation in Forecasting sales. Evidence of the Greek market. Παρουσίαση στα πλαίσια του Συνεδρίου 39th International Symposium on Forecasting (ISF 2019), που διοργανώθηκε από τη διεθνή επιστημονική οργάνωση International Institute of Forecasters-IIF, στη Θεσσαλονίκη, στην Ελλάδα.
- Voulgaraki K. Maria (2017), Forecasting sales using switching regime models in the Greek market. Παρουσίαση στα πλαίσια του Συνεδρίου 8th Biennial Hellenic Observatory Ph.D. Symposium on Contemporary Greece and Cyprus, στο Πανεπιστήμιο London School of Economics and Political Science, European Institute, The Hellenic Observatory, στο Λονδίνο, της Αγγλίας.
- Voulgaraki K. Maria (2016), Modeling and Forecasting sales in the Greek market. The case of the Greek new car sales sector. Παρουσίαση στα πλαίσια του Συνεδρίου IMAEF 2016-Ioannina Meeting on Applied Economics and Finance, στην Κέρκυρα, στην Ελλάδα.
- Voulgaraki K. Maria (2014), Forecasting Sales of durable products in the Greek market. Empirical evidence from the new car retail sector. Παρουσίαση στα πλαίσια Ετήσιου Μεταπτυχιακού Σεμιναρίου (13-04-2014), στο Τμήμα Οικονομικών Επιστημών του Πανεπιστημίου Κρήτης, στο Ρέθυμνο, Κρήτης στην Ελλάδα.
- Voulgaraki K. Maria (2013), Forecasting sales and intervention analysis of durable products in the Greek market. Empirical evidence from the new car retail sector. Παρουσίαση στα πλαίσια του Συνεδρίου 6th Biennial Hellenic Observatory Ph.D. Symposium on Contemporary Greece and Cyprus, στο Πανεπιστήμιο London School of Economics and Political Science, European Institute, The Hellenic Observatory, στο Λονδίνο, της Αγγλίας.
- Voulgaraki K. Maria (2013), Exponential Smoothing Time Series Forecasts. Παρουσίαση στα πλαίσια Ετήσιου Μεταπτυχιακού Σεμιναρίου (09-01-2013), στο Τμήμα Οικο-

νομικών Επιστημών του Πανεπιστημίου Κρήτης, στο Ρέθυμνο, Κρήτης στην Ελλάδα.

1.3 Συνεισφορά διατριβής στην επιστήμη.

Με την παρούσα διατριβή υποστηρίζουμε ότι δεδομένα οικονομικής δραστηριότητας, όπως οι πωλήσεις αυτοκινήτων στην ελληνική αγορά, μπορούν εύκολα να μετρηθούν και να προβλεφθούν επιτυχώς χρησιμοποιώντας μεθοδολογία και διαδικασίες που χρησιμοποιούνται για την αξιολόγηση χρονολογικών σειρών στην οικονομική επιστήμη. Αυτή η μελέτη αντιμετωπίζει με μοναδικό και πρωτότυπο τρόπο το πρόβλημα της παρακολούθησης και πρόβλεψης των πωλήσεων καινούργιων αυτοκινήτων στην ελληνική αγορά, σε μια πολύ ταραχώδη οικονομική περίοδο για την Ελλάδα. Τα μέτρα λιτότητας που εφαρμόστηκαν, λόγω της αυστηρής εφαρμογής των μνημονικών συμφωνιών, επηρέασαν την ελληνική οικονομία και δραστηριότητα, συνεπώς και τα επίπεδα πωλήσεων αυτοκινήτων. Με τη χρήση μονομεταβλητών υποδειγμάτων χρονολογικών σειρών γραμμικών και μη γραμμικών, με σταθερή ή κυμαινόμενη διακύμανση έγινε μια προσπάθεια να μετρηθούν, να αξιολογηθούν και να προβλεφθούν τα επίπεδα πωλήσεων αυτοκινήτων στην Ελλάδα.

Η εμπειρική εφαρμογή περισσότερων από επτά υποδειγμάτων χρονολογικών σειρών και επιπροσθέτως η εφαρμογή συνδυασμού προβλέψεων διαφορετικών υποδειγμάτων ολοκληρώνουν μια εκτενή και λεπτομερή μελέτη των οικονομετρικών υποδειγμάτων χρονολογικών σειρών, που μπορούν να εφαρμοστούν σε δεδομένα αγοράς και που μπορούν να αποτυπώσουν με αξιοπιστία την πορεία των τιμών των μεταβλητών αλλά και να την προβλέψουν.

Η διεξοδική εμπειρική εφαρμογή του μετασχηματισμού δεδομένων χρησιμοποιώντας τη μεθοδολογία Box-Cox έναντι των αρχικών τιμών δίνει μια καλύτερη κατανόηση του λόγου που ερμηνεύει γιατί ο μετασχηματισμός των δεδομένων μπορεί πάντα να βοηθήσει στην πρόβλεψη μίας μεταβλητής. Η έρευνα προσφέρει βήματα προς μια καλύτερη κατανόηση της μεταβλητότητας και της πολυπλοκότητας της ελληνικής αγοράς και συμπεραίνει ότι όσο περισσότερες πληροφορίες μπορεί κανείς να χρησιμοποιήσει τόσο καλύτερη θα είναι η πρόβλεψη της τιμής. Καταλήγουμε στο συμπέρασμα ότι ο συνδυασμός πληροφοριών που παρέχονται από διαφορετικά μοντέλα πρόβλεψης δίνει καλύτερα αποτελέσματα από τη χρήση αποτελεσμάτων μεμονω-

μένων υποδειγμάτων. Παρόλα αυτά είναι δύσκολο να προσδιοριστεί ένα μόνο υπόδειγμα για όλες τις χρονολογικές σειρές και όλα τα χρονικά διαστήματα. Για το λόγο αυτό απαιτείται πάντα προσεκτική επιλογή του υποδείγματος που είναι κατάλληλο για τα συγκεκριμένα κάθε φορά στοιχεία και το δεδομένο χρονικό διάστημα που εξετάζεται.

Κεφάλαιο 2

Ανάλυση Χρονολογικών Σειρών

2.1 Εισαγωγή

Η συλλογή παρατηρήσεων σε προκαθορισμένες μεταβλητές κατά τη διάρκεια του χρόνου σε προκαθορισμένες επαναλαμβανόμενες στιγμές, όπως κάθε ημέρα, εβδομάδα, μήνα κτλ. ονομάζεται χρονολογική σειρά. Σε αυτή τη διατριβή η εμπειρική μελέτη καλύπτει την ανάλυση χρονολογικών σειρών της ταξινόμησης νέων επιβατικών αυτοκινήτων στην Ελλάδα για ορισμένες από τις μεγαλύτερες αντιπροσωπίες αυτοκινήτων που εισάγουν επιβατικά αυτοκίνητα που προορίζονται για το λιανικό εμπόριο στην ελληνική αγορά. Η ταξινόμηση των νέων επιβατικών αυτοκινήτων (μέσα από το σύστημα της Ανεξάρτητης Αρχής Δημοσίων Εσόδων -ΑΑΔΕ) είναι σαφέστατα προκαθορισμένη, υποχρεωτική για όλα τα νέα αυτοκίνητα, που κυκλοφορούν στο ελλαδικό χώρο και συστηματικά καταγεγραμμένη, σε ίσα χρονικά διαστήματα ενώ δίνεται για περαιτέρω επεξεργασία σε δεδομένα μηνιαίας βάσης. Συνεπώς μπορούμε να υποθέσουμε με ασφάλεια ότι η μηνιαία ταξινόμηση των νέων αυτοκινήτων ισοδυναμεί με τις μηνιαίες πωλήσεις αυτοκινήτων για κάθε μια επιλεγμένη αντιπροσωπία που λειτουργεί στην ελληνική αγορά. Αυτές οι παρατηρήσεις αντιμετωπίζονται ως δεδομένα χρονολογικών σειρών, και η ανάλυση τους, έχει χαρακτηριστικά όπως μια ανάλυση που γίνεται σε κανονικά δεδομένα μάρκετινγκ. Επιπλέον, η μελέτη των νέων δεδομένων πωλήσεων αυτοκινήτων περιλαμβάνει μεθόδους ανάλυσης δεδομένων, εξαγωγής σημαντικών στατιστικών και άλλων χαρακτηριστικών. Αναπτύσσεται σε αυτό το κεφάλαιο μια συγκριτική ανάλυση, η οποία επισημαίνει ομοιότητες

και διαφορές στα δεδομένα των χρονολογικών σειρών που ερευνούνται.

Η ανάλυση των δεδομένων ξεκινάει από την παρουσίαση της ελληνικής αγοράς αυτοκινητών και συνεχίζει με μια πιο εμπεριστατωμένη περιγραφή των επιλεγμένων δειγμάτων. Στην εμπειρική έρευνα, η περιγραφική στατιστική, περιλαμβάνει τον υπολογισμό του μέσου, της διακύμανση, της κύρτωσης, της λοξότητας/ασυμμετρίας για κάθε μία από τις δέκα (10) μεγαλύτερες ελληνικές αντιπροσωπίες αυτοκινήτων. Όλα τα στατιστικά χαρακτηριστικά των σειρών δίνονται μαζί με δύο (2) τεστ ελέγχου της κανονικότητας και της τυχειότητας /σποραδικότητας των δειγμάτων μας. Η γραφική απεικόνιση των δεδομένων αυτής της εμπειρικής ανάλυσης παρουσιάζεται σε γραμμογραφήματα τιμών και γραφήματα κουτιών Box- Plots σε μηνιαία βάση των αρχικών δεδομένων μαζί με τα κορελογράμματα των συναρτήσεων αυτοσυσχέτισης και μερικής αυτοσυσχέτισης (δείτε αναλυτικά στο Παράρτημα Α).

Η εποχικότητα, η οποία είναι η συχνά επαναλαμβανόμενη αύξηση ή μείωση των μεταβλητών σε συγκεκριμένα χρονικά διαστήματα κατά την διάρκεια του έτους, μπορεί να παρατηρηθεί στα σχεδιαγράμματα και σε συνάρτηση με τους μήνες ανά έτος, εύκολα έχουμε και οπτικά συμπεράσματα για τους μήνες με τις υψηλότερες και χαμηλότερες πωλήσεις σε ετήσια βάση ανά εταιρεία.

Εν συντομία, το κεφαλαίο αυτό αναπτύσσεται ως εξής : Αρχικά έχουμε μία μικρή εισαγωγή όπου συζητείται μια γενική αναφορά στη συνολική ελληνική αγορά αυτοκινήτων, με μια σύντομη αναφορά σε πολιτικά και οικονομικά γεγονότα. Στην συνέχεια, παρουσιάζονται τα συνοπτικά στατιστικά από την εμπειρική ανάλυση των μεταβλητών, στοιχεία περιγραφικής στατιστικής ανάλυσης δεδομένων, και γραφική παρουσίαση των σειρών. Ακολουθεί η βιβλιογραφική ανασκόπηση που παρέχει περισσότερες λεπτομέρειες σχετικά με τη σχετική εργασία που έχει γίνει στον τομέα του αυτοκινήτου και την εφαρμοσμένη μεθοδολογία για υποδειγματοποίηση χρονολογικών σειρών σε δεδομένα μάρκετινγκ (όπως πωλήσεις). Παρουσιάζεται, εν συντομία, η ερευνητική μεθοδολογία και οι μέθοδοι πρόβλεψης χρονολογικών σειρών και υποδειγμάτων, με ιδιαίτερη έμφαση στο θεωρητικό τους υπόβαθρο και τη μεθοδολογία τους. Στη συνέχεια γίνεται παρουσίαση των μεθόδων πρόβλεψης και αναλύεται έννοια της ποιοτικής έναντι ποσοτικής πρόβλεψης και στην τελευταία ενότητα, έχουμε μια σύνοψη των πορισμάτων αυτού του κεφαλαίου.

2.2 Συμπεράσματα

Κατά τη διάρκεια των δύο τελευταίων δεκαετιών, η νέα αγορά αυτοκινήτων στην Ελλάδα υπέστη ορισμένες σοβαρές αλλαγές λόγω της δύσκολης οικονομικής κατάστασης στη χώρα. Οι πωλήσεις καινούργιων αυτοκινήτων συρρικνώθηκαν και δεν κατάφεραν ποτέ να ξεπεράσουν την οικονομική κρίση, ιδίως μετά την παγκόσμια οικονομική κρίση του 2007-2008.

Στην ελληνική αγορά αυτοκινήτων λειτουργούν περισσότερες από 35 διαφορετικές εταιρείες σε εθνικό επίπεδο, αλλά μόνο 2-3 από αυτές έχουν μερίδιο αγοράς σε ποσοστό περίπου 10% του συνόλου της αγοράς. Μελετήσαμε δέκα (10) από τις κορυφαίες εταιρείες σε πωλήσεις στην ελληνική αγορά. Τα αποτελέσματα περιγραφικής στατιστικής ανάλυσης σε αρχικές και λογαριθμημένες τιμές (Πίνακα 2.1 διατριβής) δείχνουν ότι όταν χρησιμοποιούνται οι αρχικές τιμές της σειράς, ο μέσος όρος της σειράς δεν είναι ίσος με τον διάμεσο, οπότε υποδηλώνει ότι δεν έχουν κανονική κατανομή, σε καμία από τις δέκα διαφορετικές χρονολογικές σειρές. Επιπλέον, η ασυμμετρία (S) και η κυρτότητα (K) της σειράς διαφέρουν από αυτά που συνήθως απαντώνται σε κανονικά κατανομημένες σειρές (δηλαδή $S=0$, $K=3$). Τα αποτελέσματα δείχνουν ότι τα αρχικά δεδομένα έχουν θετική ασυμμετρία (δεξιά) και είναι πλατύκυρτα ($K < 3$), ενώ οι λογαριθμημένες τιμές τους έχουν αρνητική ασυμμετρία (αριστερά) και είναι πλατύκυρτα (με δύο εξαιρέσεις την Volkswagen και την Nissan).

Ωστόσο, όταν εξετάζουμε τις λογαριθμημένες τιμές της σειράς, τότε οι διαφορές μεταξύ του μέσου και του μέσου όρου γίνονται μικρότερες, οπότε υπάρχει μια κανονικότητα στα δεδομένα, αλλά διατηρούν ακόμα μια μικρή απόκλιση. Επιπλέον, το τεστ των Jarque-Bera για τον έλεγχο της κανονικότητας, απορρίπτει την κανονικότητα για όλες τις λογαριθμημένες τιμές των σειρών (εκτός από την περίπτωση Fiat), ενώ στις αρχικές τιμές μόνο η Toyota και η Volkswagen φαίνεται να κατανέμονται κανονικά.

Μετά την παρουσίαση της σχετικής βιβλιογραφίας παρατίθεται ένα σύντομο θεωρητικό πλαίσιο των διαφόρων μοντέλων χρονολογικών σειρών, που πρόκειται να εφαρμοστούν εμπειρικά στα επόμενα κεφάλαια με τη μεθοδολογία πρόβλεψης, που θα ακολουθήσουμε. Η εμπειρική έρευνα συνεχίζεται με την χρήση των λογαριθμημένων τιμών και αργότερα και άλλων μετασχηματισμών των τιμών τους, για μια πιο ακριβή και γρήγορη επεξεργασία των δεδομένων και

για τη διευκόλυνση και σε βάθος παρουσίασης των αποτελεσμάτων σε μια υποδειγματοποίηση εντός- και εκτός- δείγματος και την προσπάθεια εκτίμησης του βέλτιστου μοντέλου και την πραγματοποίηση πιο αποτελεσματικών προβλέψεων.

Κεφάλαιο 3

Υποδειγματοποίηση εντός-δείγματος

3.1 Εισαγωγή

Σε αυτό το κεφάλαιο της διατριβής, ξεκινάμε την εμπειρική ανάλυση με υποδειγματοποίηση εντός δείγματος των χρονολογικών σειρών που εξετάζουμε. Οι μηνιαίες πωλήσεις αυτοκινήτων από τις κορυφαίες εταιρείες, που λειτουργούν στην ελληνική αγορά αντιμετωπίζονται ως χρονολογικές σειρές, και η μελέτη ξεκινά με απλά οικονομετρικά υποδείγματα, όπως το μοντέλο Μέσου(Mean), το Αφελές (Naïve) και το εποχικά αφελές (Seasonal Naïve) και στη συνέχεια προχωράει σε μερικά πιο προηγμένα μοντέλα όπως τα υποδείγματα Εκθετικής Εξομάλυνσης χώρου και χρόνου (Exponential Smoothing state space -ETS) και τα γραμμικά υποδείγματα με εποχιακές ψευδομεταβλητές ((Linear Model with Seasonal Dummies-LMSD)).

Επιπλέον, η έρευνα επικεντρώνεται στη συσχέτιση της παρούσας τιμής της σειράς, σε προηγούμενες τιμές και σε σφάλματα προηγούμενης πρόβλεψης, οπότε χρησιμοποιεί μοντέλα χρονολογικών σειρών όπως τα αυτοπαλλίνδρομα κινητού μέσου (AutoRegressive Integrated Moving Average - ARIMA) και δεδομένου ότι τα δεδομένα μας έχουν εποχικότητα, τα υποδείγματα αυτά λαμβάνουν υπόψη και την εποχικότητα (Seasonal AutoRegressive Integrated Moving Average -SARIMA). Για την δημιουργία των υποδειγμάτων χρησιμοποιείται η μεθοδολογία Box and Jenkins. Αυτή η εμπειρική έρευνα ολοκληρώνεται με την εκτίμηση του υβριδικού υποδείγματος που έχει το γραμμικό εποχιακό αυτοπαλλίνδρομο υπόδειγμα κινητού μέσου (SARIMA) για την εκτίμηση του μέσου και το μη γραμμικό υπόδειγμα της γενικευμένης

αυτοπαλλίνδρομής ετεροσκεδαστικής υπό συνθήκης εκτίμησης (GARCH) για την υποδειγματοποίηση της διακύμανσης της σειράς. Ο έλεγχος της αποτελεσματικότητας και της καλής εφαρμογής των υποδειγμάτων ποσοτικοποιείται με τη χρήση του ριζικού μέσου τετραγωνικού σφάλματος (Root Mean Square Error - RMSE), το μέσο απόλυτο τετράγωνο σφάλμα (Mean Absolute Square Error - MASE) και το κριτήριο αξιολόγησης Akaike (AIC) σε ορισμένες περιπτώσεις. Η ανάλυση πραγματοποιείται χρησιμοποιώντας το λογισμικό R.

Το σκεπτικό της χρήσης διαφορετικών τύπων υποδειγμάτων έγκειται στο γεγονός ότι διαφορετικοί τύποι οικονομετρικών μοντέλων έχουν μελετηθεί και σχεδιαστεί για να αποτυπώνουν διαφορετικά χαρακτηριστικά των δεδομένων που συνήθως συσχετίζονται με χρονολογικές σειρές. Για παράδειγμα συχνά οι χρονολογικές σειρές εμφανίζουν ποικίλες διακυμάνσεις, υψηλές συχνότητες σε ακραίες τιμές (fat tails), διακύμανση με ομαδοποιημένη συμπεριφορά (volatility clustering) και ούτω καθεξής κατά τη διάρκεια του χρονικού διαστήματος που εξετάζεται. Ωστόσο, υπάρχει πάντα το πρόβλημα του τρόπου επιλογής του σωστού υποδείγματος, που θα ταιριάζει καλύτερα στη κάθε μεταβλητή. Αυτό το πρόβλημα επιλύεται με τον υπολογισμό της διακύμανσης στη σειρά κάθε υποδείγματος. Υπολογίζουμε λοιπόν τα ελάχιστα τετραγωνικά σφάλματα (Mean Square Error -MSE) που προκύπτουν από την εφαρμογή των διαφόρων υποδειγμάτων αλλά και τις ρίζες τους (Root Mean Square Error - RMSE) που είναι στην πραγματικότητα οι τυπικές αποκλίσεις των σειρών. Επιπλέον, ένα άλλο μέτρο σύγκρισης στην εντός δείγματος ανάλυση είναι τα ελάχιστα σφάλματα σε απόλυτο ποσοστό (MAPE).

Όσο αφορά τα μέτρα σύγκρισης για την καλύτερη απόδοση και ακρίβεια των υποδειγμάτων, το RMSE, που είναι η τετραγωνική ρίζα του μέσου όρου όλων των τετραγώνων σφαλμάτων, αγνοεί τυχόν υπερεκτιμήσεις ή υποεκτίμησης της σειράς, ενώ δεν επιτρέπει σύγκριση μεταξύ διαφορετικών χρονικών σειρών και διαφορετικών χρονικών διαστημάτων. Από την άλλη πλευρά, το MAPE, επιτρέπει τη σύγκριση μεταξύ διαφορετικών χρονολογικών σειρών και διαφορετικών χρονικών διαστημάτων και είναι ιδιαίτερα χρήσιμο όταν οι μονάδες μέτρησης της μεταβλητής είναι σχετικά μεγάλες. Στα εμπειρικά αποτελέσματα που παρουσιάζονται σε αυτό το κεφάλαιο και τα δύο μέτρα σύγκρισης υπολογίζονται για μια ποικιλία υποδειγμάτων χρησιμοποιώντας τις τιμές των χρονολογικών σειρών για το σύνολο του δείγματος. Τέλος, όλα τα υποδείγματα συγκρίνονται χρησιμοποιώντας τις μετρήσεις του μέτρου ακρίβειας MAPE,

που επιτρέπει σύμφωνα με τη βιβλιογραφία τη σύγκριση μεταξύ διαφορετικών χρονολογικών σειρών, προκειμένου να αξιολογηθεί ή εντός δείγματος καλύτερη απόδοση των διαφορετικών υποδειγμάτων.

3.2 Συμπεράσματα

Συνοψίζοντας, σε αυτό το κεφάλαιο παρουσιάζουμε τα αποτελέσματα της εμπειρικής ανάλυσης εντός δείγματος των δέκα κορυφαίων εταιρειών πωλήσεων αυτοκινήτων στην ελληνική αγορά. Η έρευνα δείχνει ότι μετά την αξιολόγηση περισσότερων από έξι μοντέλων χρονολογικών σειρών σε κάθε μία από τις σειρές δεδομένων που εξετάζουμε, το υπόδειγμα που ταιριάζει καλύτερα και έχει μεγαλύτερη ακρίβεια είναι το υπόδειγμα εκθετικής εξομάλυνσης ETS. Αυτό το αποτέλεσμα είναι κοινό σε όλες τις χρονολογικές σειρές υπό εξέταση, καθώς το μοντέλο έχει τη μικρότερη τιμή μέτρου ακρίβειας MAPE σε όλες τις σειρές. Ο λόγος είναι το γεγονός ότι, το οικονομικό περιβάλλον είναι αρκετά ασταθές κατά τη διάρκεια αυτών των δύο τελευταίων δεκαετιών, και το μοντέλο ETS έχει το πλεονέκτημα να δίνει μεγάλη σημασία στις τελευταίες παρατηρήσεις της σειράς, επομένως δίνει ως αποτέλεσμα καλύτερες προβλέψεις για τη χρονολογική σειρά.

Ωστόσο άλλα υποδείγματα όπως το SARIMA και τα εποχιακά μοντέλα Naïve, έρχονται ως δεύτερη καλύτερη επιλογή, και δίνουν επίσης πολύ καλά αποτελέσματα σε σύγκριση με άλλα μοντέλα χρονολογικών σειρών. Το συμπέρασμα αυτό έχει νόημα δεδομένου ότι τα υπό εξέταση δεδομένα έχουν εποχικότητα που εξηγείται πολύ καλά με τα υποδείγματα SARIMA και Seasonal Naïve που λαμβάνουν υπόψη το ιδιαίτερο αυτό χαρακτηριστικό.

Επιπλέον, καθώς το μέτρο ακρίβειας, το MAPE, μπορεί να επιτρέψει τη σύγκριση μεταξύ διαφορετικών χρονολογικών σειρών, μπορούμε να συμπεράνουμε ότι εάν εξετάσουμε την εφαρμογή του μοντέλου ETS σε όλες τις χρονολογικές σειρές, έχουμε την καλύτερη εφαρμογή του στην περίπτωση των Opel, Toyota, Fiat και στη συνέχεια Volkswagen, Hyundai, Ford και ούτω καθεξής. Για το μοντέλο SARIMA έχουμε παρόμοια αποτελέσματα, η καλύτερη εφαρμογή όμως είναι στο δείγμα των Volkswagen, Hyundai, Citroen, Opel, Peugeot, Nissan, Skoda, Ford, Fiat και Toyota.

Παρατηρούμε ότι οι εταιρείες που δίνουν καλύτερα αποτελέσματα με τις πιο μικρές τιμές στα μέτρα ακρίβειας είναι αυτές που διατηρούν το μεγαλύτερο μερίδιο (ποσοστό πωλήσεων) στην αγορά, κατά τη διάρκεια των ετών που εξετάζουμε. Επομένως, τα στοιχεία μας δείχνουν ότι οι εταιρείες που είχαν το μεγαλύτερο μερίδιο (σχεδόν το 10% των συνολικών πωλήσεων) στην ελληνική αγορά αυτοκινήτων (όπως Opel, Toyota, VolksWagen) διατηρούν ένα πιο σταθερό επίπεδο πωλήσεων και είναι ευκολότερο να αναδειχθεί από την έρευνα το καλύτερα μοντέλο χρονολογικών σειρών σε αυτά τα δεδομένα παρά σε εταιρείες με μικρότερο μερίδιο αγοράς και λιγότερο σταθερή θέση στη ελληνική αγορά αυτοκινήτων. Αυτή η μελέτη κατέληξε στο συμπέρασμα ότι τα υποδείγματα ETS και SARIMA είναι καταλληλότερα για την εντός δείγματος μελέτη των πωλήσεων αυτοκινήτων.

Αυτή η εμπειρική εντός δείγματος έρευνα εξέτασε επίσης τα κατάλοιπα της εκτίμησης του μοντέλου SARIMA. Υπήρχαν ενδείξεις ετεροσκεδιστικότητας στα τετράγωνα των καταλοίπων των υποδειγμάτων SARIMA που υποδηλώνουν ότι η έρευνα μπορεί να βελτιώσει το μοντέλο SARIMA μέσω υποδειγματοποίησης της μεταβλητότητας και έτσι εκτιμήσαμε τα μοντέλα SARIMA-GARCH για όλες τις σειρές. Ενώ ο έλεγχος Engle LM Test έδειξε ένα φαινόμενο ARCH μόνο σε μια μικρή ομάδα εταιρειών (Opel, Peugeot, Ford, Fiat) και τα στατιστικά στοιχεία Box-Pierce και Ljung-Box δείχνουν ότι τα τετράγωνα των καταλοίπων των υποδειγμάτων SARIMA των εταιρειών Volkswagen και Nissan υποφέρουν κυρίως από σειριακή συσχέτιση, αυτή η μελέτη συνεχίζει την έρευνα της με την εκτίμηση διαφορετικών τύπων υποδειγμάτων SARIMA-GARCH. Εκτιμήσαμε διάφορες προδιαγραφές αυτών των μοντέλων και καταλήγουμε στο συμπέρασμα ότι, γενικά το μοντέλο SARIMA-GARCH (1,1) είναι το πιο δημοφιλές, για εκτιμήσεις σε δείγματα, καθώς έχει τη χαμηλότερη τιμή AIC από άλλα υποδείγματα. Μπορούμε να υποθέσουμε ότι η μεταβλητότητα της σειράς εξηγείται καλύτερα χρησιμοποιώντας αυτήν την ομάδα υποδειγμάτων, αλλά δεν μπορούμε να είμαστε σίγουροι ότι αυτές οι προδιαγραφές υποδειγμάτων μπορούν να βελτιώσουν την ικανότητα πρόβλεψης. Για την εκτίμηση της ακρίβειας της πρόβλεψης των μοντέλων συνεχίζουμε με μια εκτός δείγματος έρευνα στα επόμενα κεφάλαια.

Κεφάλαιο 4

Υποδειγματοποίηση και Προβλέψεις εκτός-δείγματος

4.1 Εισαγωγή

Η πρόβλεψη είναι μια προσπάθεια πρόγνωσης της μελλοντικής τιμής μιας μεταβλητής, με όσο το δυνατόν μεγαλύτερη ακρίβεια, λαμβανομένων όλων των διαθέσιμων πληροφοριών, συμπεριλαμβανομένων των ιστορικών δεδομένων και των γνώσεων για παρελθόντα, παρόντα ή μελλοντικά γεγονότα, που δύναται να επηρεάσουν τις τιμές των προβλέψεων. Αυτό είναι ένας σύνηθες στατιστικό έργο στην οικονομική επιστήμη, που παρέχει χρήσιμες πληροφορίες για τους υπεύθυνους λήψης αποφάσεων, κάθε τομέα του δημόσιου ή του ιδιωτικού τομέα. Είναι δε μείζονος σημασίας σε ζητήματα, όπως ο προγραμματισμός παραγωγής, η διαχείριση προσωπικού, ο στρατηγικός σχεδιασμός, η χρήση μεθόδων πρόβλεψης. Στην πράξη, ωστόσο, η πρόβλεψη πωλήσεων σε συγκεκριμένες εταιρείες, γίνεται συχνά με τη χρήση απλών μεθόδων πρόβλεψης, που συχνά δεν βασίζονται σε στατιστικά υποδείγματα. Σε αυτήν την έρευνα αναπτύσσεται η χρήση στατιστικών προβλέψεων στην αγορά αυτοκινήτων στην Ελλάδα, με τη χρήση μεθόδων πρόβλεψης χρονολογικών σειρών.

Για να αξιολογήσουμε ποιο είναι το κατάλληλο υπόδειγμα πρόβλεψης, μπορούμε να χρησιμοποιήσουμε μια εντός-δείγμα ή μια εκτός-δείγματος διαδικασία εκτίμησης των υποδειγμάτων. Στην εντός-δείγματος τεχνική λαμβάνουμε υπόψη όλα τα δεδομένα για τον υπολογισμό του

18ΚΕΦΑΛΑΙΟ 4. ΥΠΟΔΕΙΓΜΑΤΟΠΟΙΗΣΗ ΚΑΙ ΠΡΟΒΛΕΨΕΙΣ ΕΚΤΟΣ-ΔΕΙΓΜΑΤΟΣ

υποδείγματος. Στην εκτός-δείγματος τεχνική εκτιμάτε το υπόδειγμα με τη χρήση μέρους των παρατηρήσεων, ενώ οι υπόλοιπες παρατηρήσεις που μένουν εκτός δείγματος χρησιμοποιούνται για την αξιολόγηση των τιμών των προβλέψεων σε σύγκριση με τις πραγματικές τους τιμές. Σε αυτή τη μελέτη η ερευνήτρια χρησιμοποιεί μια μεθοδολογία πρόβλεψης εκτός- δείγματος.

Το χρονικό διάστημα 1998-2016, στο οποίο επεκτείνονται τα δεδομένα των χρονολογικών σειρών που εξετάζουμε, χωρίζεται σε τέσσερα (4) μικρότερα διαφορετικά υποσύνολα για τις ανάγκες της παρούσας έρευνας. Τα διάφορα υποσύνολα δεδομένων της εμπειρικής μελέτης μας καθορίζονται σύμφωνα με την αναλογία 8:2 μεταξύ του υποσυνόλου εκτίμησης του υποδείγματος και του υποσυνόλου ελέγχου. Αναλυτικότερα, το υποσύνολο εκτίμησης των δεδομένων χρησιμοποιείται για την εκτίμηση των παραμέτρων του μοντέλου και καλύπτει περίπου το 80% του συνόλου του δείγματος, ενώ το υποσύνολο ελέγχου χρησιμοποιείται για τη μέτρηση της ακρίβειας της πρόβλεψης και καλύπτει το υπόλοιπο 20% του συνολικού δείγματος των δεδομένων. Σε αυτήν την εμπειρική μελέτη εκτός-δείγματος δημιουργούμε συνολικά τέσσερα υποσύνολα A, B, C και D που χωρίζονται σε χρονικά διαστήματα (αναλυτικά στον Πίνακα 4.1 της διατριβής).

Το πρώτο σετ (A) αυτής της μελέτης καλύπτει 19 χρόνια. Ξεκινά με ένα υποσύνολο εκτίμησης υποδείγματος από τον Ιανουάριο 1998 έως τον Δεκέμβριο του 2012 που καλύπτει 15 χρόνια με 180 μηνιαίες παρατηρήσεις, και ένα σύνολο ελέγχου που καλύπτει τα επόμενα τέσσερα (4) έτη από τον Ιανουάριο του 2013 έως τον Δεκέμβριο του 2016 με 48 παρατηρήσεις.

Το δεύτερο υποσύνολο (B) αυτής της έρευνας καλύπτει 10 χρόνια. Το υποσύνολο εκτίμησης ξεκινά από τον Ιανουάριο του 2006 έως τον Δεκέμβριο του 2013 και καλύπτει 8 χρόνια με 96 μηνιαίες παρατηρήσεις, ενώ το υποσύνολο ελέγχου έχει 24 παρατηρήσεις και καλύπτει τα επόμενα δύο (2) έτη από τον Ιανουάριο του 2014 έως τον Δεκέμβριο του 2015.

Το τρίτο υποσύνολο (C) αυτής της εμπειρικής μελέτης ξεκινά τον Ιανουάριο του 2006 έως τον Δεκέμβριο του 2009 και καλύπτει τέσσερα (4) έτη με 48 μηνιαία δεδομένα, που καθορίζουν κάθε φορά το μοντέλο και ελέγχουν την απόδοσή του σε ένα υποσύνολο 12 παρατηρήσεων που καλύπτουν τον επόμενο χρόνο από τον Ιανουάριο έως τον Δεκέμβριο του 2010.

Το τέταρτο υποσύνολο (D) αυτής της μελέτης ξεκινά τον Ιανουάριο του 2002 έως τον Δεκέμβριο του 2009 και καλύπτει οκτώ (8) έτη με 96 μηνιαία δεδομένα που καθορίζουν κάθε

φορά το υπόδειγμα και ελέγχουν την απόδοσή του με ένα υποσύνολο 24 παρατηρήσεων που καλύπτει τα επόμενα δύο (2) έτη από τον Ιανουάριο του 2010 έως τον Δεκέμβριο του 2011.

Επιπλέον θεωρούμε ότι η μεγάλη ύφεση στην ελληνική αγορά ξεκίνησε στις αρχές του 2008 μαζί με την παγκόσμια χρηματοπιστωτική κρίση (GFC) 2007-2008 και έγινε ακόμη χειρότερη μετά την ανακοίνωση της οικονομικής συμφωνίας της ελληνικής κυβέρνησης με το Διεθνές Νομισματικό Ταμείο (ΔΝΤ/(IMF) τον Μάιο του 2010. Οι πωλήσεις έφτασαν στο ιστορικό τους κατώτατο σημείο το 2012 και μέχρι το τέλος του 2016 η ελληνική αγορά δέχθηκε μεγάλη πίεση και μια παρατεταμένη περίοδο οικονομικής κρίσης. Η εφαρμογή μέτρων λιτότητας ανάγκασε τους Έλληνες πολίτες να αναβάλουν ή να καθυστερήσουν την αγορά αυτοκινήτων, οι τράπεζες δεν έδιναν σχεδόν κανένα δάνειο για αγορά καταναλωτικών αγαθών και αυτοκινήτων και ως εκ τούτου υπήρξε μια δραματική αλλαγή στα επίπεδα πωλήσεων τα τελευταία χρόνια.

Σε αυτό το κεφάλαιο, μετά από αυτήν τη μικρή εισαγωγή, αναλύονται μερικά από τα βασικότερα μέτρα αξιολόγησης της ακρίβειας των προβλέψεων που χρησιμοποιούνται στην παρούσα εμπειρική μελέτη (Πίνακας 4.2). Επιπλέον, εκτιμούνται και αξιολογούνται αυτά τα μέτρα ακρίβειας των προβλέψεων εμπειρικά στα δεδομένα χρονολογικών σειρών μας χρησιμοποιώντας μια εκτίμηση εκτός-δείγματος. Πιο συγκεκριμένα, μετά από έναν ορισμό και μια σύντομη αναφορά στη βιβλιογραφία για τα μέτρα ακρίβειας πρόβλεψης που χρησιμοποιήθηκαν, η έρευνα συνεχίζεται με μια εμπειρική εφαρμογή τους σε διάφορα μοντέλα χρονολογικών σειρών.

Τα εμπειρικά αποτελέσματα της πρόβλεψης πωλήσεων εκτός-δείγματος δίδονται για διαφορετικά μοντέλα χρονολογικών σειρών με διάφορα μέτρα ακρίβειας πρόβλεψης. Η διαδικασία υπολογισμού είναι η ακόλουθη: εκτίμηση κάθε παραμέτρου του υποδείγματος χρησιμοποιώντας τις τιμές παρατηρήσεων του υποσυνόλου για την εκτίμηση των υποδειγμάτων, δημιουργία πρόβλεψης σημείου για την περίοδο πρόβλεψης σε κάθε σύνολο και εκτίμηση των μέτρων ακρίβειας πρόβλεψης χρησιμοποιώντας τις τιμές εκτιμώμενες τιμές πρόβλεψης του μοντέλου και τις πραγματικές τιμές για κάθε υποσύνολο που έχουν κρατηθεί εκτός δείγματος.

Επιπλέον, οι μηνιαίες προβλέψεις πωλήσεων παρουσιάζονται και γραφικά για τις διάφορες μεθόδους πρόβλεψης χρονολογικών σειρών, σε σύγκριση με τα πραγματικά επίπεδα πωλήσεων με διάστημα εμπιστοσύνης 95% για τη διακύμανση κάθε μοντέλου, για τις πωλήσεις νέων

αυτοκινήτων της Opel στην Ελλάδα (σχεδιάγραμμα 4.2-4.5).

Αυτή είναι μια προσπάθεια να δείξουμε οπτικά τα αποτελέσματα της εμπειρικής μας μελέτης. Τέλος, συζητάμε το εύρημα αυτής της εμπειρικής μελέτης χρησιμοποιώντας διαφορετικές μεθόδους και μοντέλα χρονολογικών σειρών πρόβλεψης και εμπλουτίζουμε τα ευρήματα με οικονομική ανάλυση του συγκεκριμένου τομέα λιανικής της ελληνικής αγοράς.

4.2 Συμπεράσματα

Σε αυτό το κεφάλαιο, η έρευνά εστίασε στην απόδοση των προβλέψεων εκτός-δείγματος έξι (6) διαφορετικών μοντέλων χρονολογικών σειρών: του Μέσου όρου, του Αφελές, του εποχιακά Αφελές μοντέλου, του υποδείγματος εκθετικής εξομάλυνσης χώρου χρόνου, του γραμμικού υποδείγματος με εποχιακές ψευδομεταβλητές, τα εποχιακά αυτοπαλλίνδρομα υποδείγματα κινητού μέσου και τα γενικευμένα υπο συνθήκη ετεροσκεδαστικά υποδείγματα. Ο χρόνος συλλογής των μηνιαίων δεδομένων ήταν από τις αρχές του 1998 έως το τέλος του 2016 και χωρίστηκε σε τέσσερα (4) διαφορετικά υποσύνολα δεδομένων, με βάση τον κανόνα του 80:20 για τα διευκολύνει τα υποσύνολα ελέγχου και δοκιμών. Ο στόχος ήταν να εξεταστεί πόσο καλά είναι αυτά τα υποδείγματα χρονολογικών σειρών στην πρόβλεψη των επιπέδων μηνιαίων πωλήσεων αυτοκινήτων σε δύσκολες οικονομικές καταστάσεις, όπως αυτή που αντιμετώπιζε η ελληνική οικονομία τις τελευταίες δύο δεκαετίες.

Μελετήθηκαν εμπειρικά στοιχεία για τρεις διαφορετικές εταιρείες Opel, Toyota και Fiat. Τα αποτελέσματα δείχνουν ότι το βασικά υποδείγματα Μέσου όρου και το Αφελές υπόδειγμα ήταν πολύ λίγα στην πρόβλεψη των επιπέδων πωλήσεων. Σε αυτήν την εμπειρική εφαρμογή κάθε εταιρεία αντιδρούσε διαφορετικά και αυτό είναι φυσιολογικό, καθώς κάθε εταιρεία παρουσιάζει διαφορετικά επίπεδα πωλήσεων και διακυμάνσεις. Επιπλέον, τα μέτρα ακρίβειας των προβλέψεων δεν έδιναν πάντα το ίδιο αποτέλεσμα. Αυτό ήταν αναμενόμενο, καθώς η ρίζα του μέσου τετραγωνικού σφάλματος (RMSE) μετρά το μέσο μέγεθος του σφάλματος (δηλαδή είναι η τετραγωνική ρίζα του μέσου όρου της τετραγωνικής διαφοράς μεταξύ πρόβλεψη και πραγματική παρατήρηση), ενώ το μέσο απόλυτο ποσοστό σφάλματος (MAPE) μετρά το μέσο μέγεθος των σφαλμάτων στο σύνολο των προβλέψεων ως ποσοστό. Η κύρια διαφορά μεταξύ

αυτών των δύο μέτρων είναι ότι το RMSE δίνει ένα σχετικά υψηλό βάρος σε μεγάλα σφάλματα. Αυτό σημαίνει ότι το RMSE μπορεί να είναι πιο χρήσιμο όταν τα μεγάλα λάθη είναι ιδιαίτερα ανεπιθύμητα επειδή τιμωρεί περισσότερο τα μεγάλα λάθη.

Στον Πίνακα 4.7 παρουσιάζεται η σύνοψη των εμπειρικών ερευνητικών αποτελεσμάτων επιλογής του καλύτερου υποδείγματος, σε μελέτη της προβλεπτικής ικανότητας με τη μεθολογία εκτός-δείγματος ελέγχου, σύμφωνα με τα μέτρα ακρίβειας πρόβλεψης της ρίζας του μέσου τετραγωνικού σφάλματος (RMSE) και μέσου απόλυτου ποσοστού σφάλματος (MAPE).

Τα εμπειρικά αποτελέσματα για την εταιρεία Opel δίνουν ως πρώτη επιλογή το Αφελές μοντέλο στις μετρήσεις RMSE και MAPE για το υποσύνολο δεδομένων A, το οποίο έχει 15 χρόνια μηνιαίες παρατηρήσεις στο υποσύνολο εκτίμησης του μοντέλου και δίνει μακροπρόθεσμη πρόβλεψη για τα επόμενα τέσσερα (4) χρόνια. Ωστόσο, στο υποσύνολο δεδομένων B, C και D, οι πωλήσεις της εταιρείας Opel προβλέπονται καλύτερα από τα μοντέλα εκθετικής εξομάλυνσης ETS, κυρίως επειδή το υποσύνολο εκτίμησης καλύπτει μια περίοδο αρκετά ταραχώδους κίνησης και βαθιάς ύφεσης για τις πωλήσεις αυτοκινήτων, λόγω της οικονομικής κρίσης στην αγορά, που ήταν προφανής κυρίως μετά το 2010. Επιπλέον, αυτό το μοντέλο εκθετικής εξομάλυνσης δίνει πολύ καλύτερα αποτελέσματα όταν η πρόβλεψη είναι βραχυπρόθεσμη.

Η μελέτη για τις πωλήσεις νέων αυτοκινήτων της εταιρείας Toyota δίνει προτίμηση στο υπόδειγμα SARIMA για τη μακροπρόθεσμη πρόβλεψη στο υποσύνολο δεδομένων A και τη βραχυπρόθεσμη πρόβλεψη στην αρχή της βαθιάς περιόδου ύφεσης και προτιμά το υπόδειγμα εκθετικής εξομάλυνσης ETS για το υποσύνολο δεδομένων B και C, κατά τη διάρκεια της βαθιάς περιόδου ύφεσης. Η εποχικότητα είναι πολύ καθοριστική στις νέες πωλήσεις αυτοκινήτων, και το εποχιακό Αφελές υπόδειγμα βλέπουμε να επιλέγεται ως το καλύτερο υπόδειγμα πρόβλεψης για το σύνολο δεδομένων D.

Για τις πωλήσεις νέων αυτοκινήτων της εταιρείας Fiat, τα πράγματα είναι πιο ξεκάθαρα, η μακροπρόθεσμη και βραχυπρόθεσμη πρόβλεψη εκτιμάται ως πιο ακριβής από το μοντέλο εκθετικής εξομάλυνσης (ETS) για τα δεδομένα των υποσυνόλων A, B και C. Ωστόσο, για το υποσύνολο δεδομένων D έχουμε καλύτερα αποτελέσματα όταν χρησιμοποιούμε το Αφελές ή το εποχιακό αυτοπαλίνδρομο υπόδειγμα κινητού μέσου (SARIMA), σύμφωνα με τα μέτρα ακρίβειας προβλέψεων RMSE και το MAPE αντίστοιχα.

22ΚΕΦΑΛΑΙΟ 4. ΥΠΟΔΕΙΓΜΑΤΟΠΟΙΗΣΗ ΚΑΙ ΠΡΟΒΛΕΨΕΙΣ ΕΚΤΟΣ-ΔΕΙΓΜΑΤΟΣ

Σε γενικές γραμμές, δεν υπάρχει ένα μόνο υπόδειγμα, που να είναι το καλύτερο υπόδειγμα πρόβλεψης για όλες τις περιπτώσεις, που μελετήθηκαν σε αυτήν την εμπειρική έρευνα. Κάθε εταιρεία έχει το δικό της επίπεδο πωλήσεων, το δικό της πρόγραμμα μάρκετινγκ, την δική της ξεχωριστή καλή φήμη και πελατεία η οποία καθορίζει την αντίληψη και τη συμπεριφορά των πελατών στην αγορά. Επομένως, η κάθε εταιρεία ανταποκρίνεται διαφορετικά ακόμα και σε μία δύσκολη περίοδο οικονομικής κρίσης. Ο διαχωρισμός των υποσυνόλων εκτίμησης και δοκιμής της μελέτης, φαίνεται να είναι καθαριστικά κρίσιμος για την επιτυχία της πρόβλεψης και τον καθορισμό του καλύτερου υποδείγματος. Επομένως ο ερευνητής, πρέπει να κάνει προσεκτικά την επιλογή των υποσυνόλων αυτών, διότι καθορίζουν το τελικό αποτέλεσμα.

Επιπλέον, τα διαφορετικά μέτρα ακρίβειας δεν συμφωνούν πάντα μεταξύ τους ως προς τα εμπειρικά αποτελέσματα επομένως δεν δίνουν πάντα το ίδιο υπόδειγμα ως το καλύτερο. Συνήθως το μέτρο RMSE δίνει τα ίδια αποτελέσματα με το MAPE, αλλά υπάρχουν και εξαιρέσεις, δηλαδή περιπτώσεις που δεν συμφωνούν και δίνουν διαφορετικά αποτελέσματα. Όπως για παράδειγμα, στην εταιρεία Fiat στο υποσύνολο D δίνεται η δυνατότητα επιλογής του Αφελές υποδείγματος ή του Εποχιακού αυτοπαλίνδρομου υποδείγματος κινητού μέσου σύμφωνα με το RMSE ή το MAPE ανάλογα.

Γενικά, τα αριθμητικά αποτελέσματα που υπολογίστηκαν στις μετρήσεις για την ακρίβεια των υποδειγμάτων (Πίνακες 4.4 - 4.6), συμφωνούν και με την οπτική παρουσίαση που δίνεται από τα γραφήματα (Σχεδιαγράμματα 4.2 - 4.5) των προβλεπόμενων τιμών των διαφόρων υποδειγμάτων σε σύγκριση με τις πραγματικές τιμές κάθε εταιρείας.

Όσον αφορά τα διαστήματα εμπιστοσύνης (CI), φαίνεται ότι εάν έχουμε μια βραχυπρόθεσμη περίοδο πρόβλεψης (π.χ. ένα έτος, όπως στο υποσύνολο C), οι τιμές πρόβλεψης αποτυγχάνουν στην προσέγγιση των πραγματικών τιμών της μεταβλητής, καθώς οι προβλέψεις πηγαίνουν πέρα από την προβλεπόμενη περιοχή διαστήματος εμπιστοσύνης. Έχοντας δύο χρόνια ως περίοδο πρόβλεψης, η πρόβλεψη φαίνεται καλύτερη, καθώς οι τιμές πρόβλεψης βρίσκονται εντός των ορίων του διαστήματος εμπιστοσύνης και αυτό φαίνεται να είναι ασφαλέστερο για προβλέψεις.

Επιπλέον, όσο πιο στενό είναι το εύρος του διαστήματος εμπιστοσύνης, τόσο το καλύτερο, ειδικά όταν οι πραγματικές τιμές και οι τιμές των προβλέψεων είναι εντός του διαστήματος

εμπιστοσύνης. Σε αυτές τις περιπτώσεις, οι τιμές πρόβλεψης είναι αρκετά κοντά στις πραγματικές τιμές και επιπλέον το σφάλμα μπορεί να είναι πιθανό αλλά περιορισμένο.

Κεφάλαιο 5

Προβλέψεις με SARIMA-GARCH

5.1 Εισαγωγή

Σε αυτό το κεφάλαιο παρουσιάζουμε μια συγκριτική μελέτη πρόβλεψης των μεταβλητών πωλήσεων νέων αυτοκινήτων με τη χρήση των εποχιακών αυτοπαλίνδρομων υποδειγμάτων κινητού μέσου (SARIMA) στην συνάρτηση του μέσου και τα γενικευμένα αυτοπαλίνδρομα υπο συνθήκη ετεροσκεδαστικά υποδείγματα (GARCH) για την συνάρτηση της διακύμανσης. Στόχος μας είναι να προσδιορίσουμε εάν ένα μοντέλο SARIMA-GARCH μπορεί να προβλέψει επιτυχώς την αστάθεια του επιπέδου πωλήσεων αυτοκινήτων σε μια περίοδο οικονομικής κρίσης στην ελληνική αγορά.

Η μεταβλητότητα της πρόβλεψης είναι σημαντική για τρεις βασικούς σκοπούς: διαχείριση κινδύνων, κατανομή περιουσιακών στοιχείων και για τη λήψη αποφάσεων για το μελλοντικό επίπεδο του αποθέματος. Ο Robert Engle το 1982 ανέπτυξε τα αυτοπαλίνδρομα ετεροσκεδαστικά μοντέλα υπό όρους (ARCH), για να υποδειγματοποιήσει τις διακυμάνσεις που ποικίλλουν χρονικά, που παρατηρούνται συχνά σε οικονομικά δεδομένα χρονολογικών σειρών. Για αυτή τη συνεισφορά του, κέρδισε το βραβείο Νόμπελ στα Οικονομικά του 2003. Τα μοντέλα ARCH υποθέτουν ότι η διακύμανση του τρέχοντος όρου σφάλματος ή της καινοτομίας είναι συνάρτηση των πραγματικών μεγεθών των όρων σφάλματος των προηγούμενων χρονικών περιόδων

και συχνά η διακύμανση σχετίζεται με τα τετράγωνα των προηγούμενων καταλοίπων.

Έχοντας διερευνήσει τη γενική θεωρία των εποχιακών αυτοπαλίνδρομων υποδειγμάτων κινητού μέσου (SARIMA) και τα γενικευμένα αυτοπαλίνδρομα υπο συνθήκη ετεροσκεδαστικά υποδείγματα (GARCH) στα προηγούμενα κεφάλαια, αυτή η μελέτη εισάγει τα μονο-μεταβλητά υποδείγματα SARIMA-GARCH, σε μια προσπάθεια να εξετάσει την προβλεπτική τους ικανότητά τους. Η οικογένεια των μοντέλων GARCH είναι χρήσιμη επειδή μπορεί ο ερευνητής να προβλέψει καλύτερα την αστάθεια των μεταβλητών με αυτά τα υποδείγματα. Ωστόσο, είναι ιδανικό για την βραχυπρόθεσμη πρόβλεψη, δηλαδή μερικών μόνο χρονικών περιόδων μπροστά, αλλά όχι τόσο για τη μακροπρόθεσμη πρόβλεψη. Αυτά τα μοντέλα βοηθούν στην επέκταση του φαινομένου της ομαδοποίησης της μεταβλητότητας της διακύμανση της τιμής ιδιαίτερα εάν αυτή εμφανίζει σε περιόδους σχετικής ηρεμίας και περιόδους υψηλής μεταβλητότητας που είναι συχνές στα δεδομένα της αγοράς, όπως οι πωλήσεις.

Τα μοντέλα ARCH υποθέτουν ότι η διακύμανση του τρέχοντος όρου σφάλματος είναι συνάρτηση των πραγματικών μεγεθών των όρων σφάλματος των προηγούμενων χρονικών περιόδων και συχνά η διακύμανση σχετίζεται με τα τετράγωνα των προηγούμενων σφαλμάτων. Στη μελέτη μας, η διακύμανση των τιμών των πωλήσεων αυτοκινήτων κινείται με αστάθεια, αν και ο χρόνος και η εποχικότητά του εξαρτάται από τη συγκεκριμένη αγορά, όπου πραγματοποιούνται οι συναλλαγές. Έτσι, η σύνθεση ενός υποδείγματος GARCH θα βοηθήσει στην καταγραφή της διακύμανσης της τιμής που αυξάνεται προς τα πάνω και στη συνέχεια μειώνεται έως ότου υπάρξει άλλη μία επόμενη ανοδική πορεία.

Μια αυτοπαλίνδρομη προσέγγιση βοηθά στη δημιουργία αξιόπιστων υποδειγμάτων με μεγάλη ακρίβεια. Σύμφωνα με το Tsay (2005) η μεταβλητότητα της αγοράς είναι γνωστό ότι συσσωρεύεται (cluster), πράγμα που σημαίνει ότι, οι εξαιρετικά ευμετάβλητες περίοδοι τείνουν να παραμένουν για κάποιο χρονικό διάστημα προτού η αγορά επιστρέψει σε ένα πιο σταθερό περιβάλλον. Η οικογένεια μοντέλων GARCH χρησιμοποιείται ευρέως στην πράξη για την πρόβλεψη της μεταβλητότητας και των αποδόσεων της χρηματοπιστωτικής αγοράς.

Εν συντομία, το σχέδιο αυτού του κεφαλαίου είναι το ακόλουθο: μετά τη σύντομη εισαγωγή εξετάζουμε πρώτα εάν υπάρχει το φαινόμενο της αυτοσυσχέτισης των καταλοίπων στα SARIMA υποδείγματα, τα οποία δίνουν κάποια ένδειξη αυτοσυσχέτιση. Στη συνέχεια,

εξετάζουμε εάν τα κατάλοιπα των SARIMA υποδειγμάτων εμφανίζουν φαινόμενο ARCH, χρησιμοποιώντας το τεστ Ljung-Box και το Lagrange Multiplier. Τα δύο τεστ δείχνουν κάποια στοιχεία επίδρασης Arch στα κατάλοιπα. Συνεχίζουμε με τον προσδιορισμό υποδειγμάτων SARIMA-GARCH (1,1) για κάθε εταιρεία. Προσδιορίζουμε την εξίσωση του μέσου από τα προεπιλεγμένα μοντέλα SARIMA και συνεχίζουμε με την εξίσωση διακύμανσης, χρησιμοποιώντας το απλό μοντέλο GARCH (1,1). Αυτή η μελέτη εξετάζει επίσης, τρεις εναλλακτικές κατανομές στα μοντέλα SARIMA-GARCH (1,1) ώστε να βρεθεί η πιο κατάλληλη. Τα εμπειρικά στοιχεία δείχνουν ότι η κατανομή t-student είναι η καλύτερη περίπτωση έναντι της κανονικής και της γενικευμένης κατανομής σφαλμάτων.

Επιπλέον κάνουμε τον διαγνωστικό έλεγχο των υπολειμμάτων των υποδειγμάτων SARIMA-GARCH και τα ευρήματα της έρευνας επιβεβαιώνουν ότι δεν υπάρχει συσχέτιση στα υπολείμματα και ότι η εφαρμογή των υποδειγμάτων είναι σχετικά καλή. Η μελέτη κάνει σύγκριση των υποδειγμάτων SARIMA-GARCH (1,1) με άλλα υποδείγματα χρονολογικών σειρών, που έχουν ερευνηθεί για την πρόβλεψη των πωλήσεων. Τέλος εστιάζουμε στα διαστήματα πρόβλεψης SARIMA-GARCH και εξετάζουμε εάν μειώνονται για να παρέχουν καλύτερες προβλέψεις. Το κεφάλαιο καταλήγει σε μια συζήτηση για τα ευρήματα αυτής της εμπειρικής έρευνας για τα μοντέλα SARIMA-GARCH στην ελληνική αγορά αυτοκινήτων.

5.2 Συμπεράσματα

Αυτό το κεφάλαιο είναι αφιερωμένο στην αντιμετώπιση του προβλήματος της πρόβλεψης τιμών και της πρόβλεψης μεταβλητότητας στις πωλήσεις νέων αυτοκινήτων με τη χρήση μοντέλων (SARIMA-GARCH σε σύγκριση με πολλές άλλες μεθόδους που χρησιμοποιούνται σε μεγάλο βαθμό στην πράξη και έχει δοκιμαστεί η ακρίβειά τους χρησιμοποιώντας πραγματικά δεδομένα από τον τομέα της ελληνικής αγοράς (δηλ. πωλήσεις νέων αυτοκινήτων (Opel και Toyota από το 1998 έως το 2016). Η οικογένεια των μεθόδων (SARIMA-GARCH έχει πλεονεκτήματα και μειονεκτήματά της. Ορισμένα υποδείγματα χρονολογικών σειρών είναι απλά εύκολα στην εφαρμογή, αλλά αποδίδουν καλά αποτελέσματα. Άλλες μέθοδοι είναι πιο δύσκολο να εφαρμοστούν αλλά δεν αποδίδουν πάντα καλά αποτελέσματα. Εν ολίγοις, δεν υπάρχει καμία ξεκάθαρη

προσέγγιση προτιμήσεων.

Η έρευνα διαπίστωσε ότι δεν υπήρχε πολύ ισχυρή επίδραση ARCH στον έλεγχο αυτοσυσχέτισης καταλοίπων SARIMA μελετώντας το κορελόγραμμα της συνάρτηση αυτοσυσχέτισης ACF , ωστόσο αποφασίσαμε να συνεχίσουμε τη μελέτη μας, καθώς το τεστ Ljung-Box έδωσε στοιχεία αυτοσυσχέτισης στα κατάλοιπα του υποδείγματος SARIMA για την Opel και την Toyota στο σύνολο δεδομένων A. Προτιμήθηκε το SARIMA-GARCH (1,1) με κατανομή μαθητή-τ, σύμφωνα με το κριτήριο πληροφοριών AIC και BIC μεταξύ της κανονικής και της γενικευμένης εναλλακτικής διανομής. Ωστόσο, τα αποτελέσματα που δόθηκαν στη δοκιμή του συνόλου δεδομένων A ήταν ενθαρρυντικά. Με βάση τον Πίνακα 5.5 οι προβλέψεις που παράγονται από το μοντέλο SARIMA-GARCH (1,1) είναι καλύτερες, καθώς οι μετρήσεις του RMSE είναι χαμηλότερες από αυτές που παράγονται από το SARIMA. Έτσι μπορούμε να συμπεράνουμε ότι στην περίπτωση των μηνιαίων πωλήσεων αυτοκινήτων Opel και Toyota, το μοντέλο (SARIMA-GARCH(1,1) μπορεί να είναι ένας αποτελεσματικός τρόπος βελτίωσης της ακρίβειας των προβλέψεων.

Είναι επίσης απίστευτο, πως τα μοντέλα γενική θεωρία των εποχιακών αυτοπαλίνδρομων υποδειγμάτων κινητού μέσου SARIMA-GARCH βοηθούν τόσο πολύ στη μείωση των διαστημάτων πρόβλεψης της πρόβλεψης, και ως εκ τούτου μπορεί να προσφέρει μια καλύτερη πρόβλεψη. Από την άλλη πλευρά, σύμφωνα με τα διαστήματα πρόβλεψης στο σχήμα 5.5 και 5.6 , φαίνεται ότι τα μοντέλα SARIMA-GARCH θα πρέπει καλύτερα να χρησιμοποιηθούν για μια βραχυπρόθεσμη προσέγγιση της πρόβλεψης μεταβλητότητας και όχι μακροπρόθεσμα, καθώς μπορούν να παράγουν πιο ακριβείς προβλέψεις βραχυπρόθεσμα. Επιπλέον, αυτή είναι μια πολύτιμη ερευνητική εμπειρία για την εφαρμογή της οικογένειας των μοντέλων SARIMA-GARCH στις νέες πωλήσεις αυτοκινήτων στην ελληνική αγορά και την προστιθέμενη αξία στην επιστήμη και στην έρευνα των πωλήσεων στην αγορά οχημάτων στην Ελλάδα.

Κεφάλαιο 6

Μετασχηματισμοί δεδομένων και Προβλέψεις.

6.1 Εισαγωγή

Τα ιστορικά δεδομένα μπορούν συχνά να προσαρμοστούν ή να μετατραπούν από τις αρχικές τους τιμές, για να οδηγήσουν την έρευνα πρόβλεψης σε μια πιο απλούστερη εργασία. Η έρευνα, στην καθημερινή πραγματικότητα, δείχνει ότι σχεδόν όλες οι αναλύσεις επωφελούνται από τη βελτιωμένη ομαλότητα των μεταβλητών, ιδιαίτερα σε περιπτώσεις όπου υπάρχει ουσιαστική μη κανονικότητα. Μέχρι αυτό το σημείο, έχουμε επιλέξει έναν παραδοσιακό μετασχηματισμό – τις λογαριθμημένες τιμές καταγραφής των αρχικών δεδομένων - η οποία χρησιμοποιείται συχνά στην έρευνα, για τη βελτίωση της ομαλότητας των δεδομένων μας και την παραγωγή σχέσεων με πιο ομοιογενή υπολείμματα, με άλλα λόγια σταθερή διακύμανση.

Ωστόσο, χρησιμοποιώντας μια ευρύτερη κατηγορία μετασχηματισμού ισχύος, που εισήχθη από την Box-Cox (1964), θα μας βοηθήσει να βρούμε εύκολα το βέλτιστο μετασχηματισμό που θα ομαλοποιεί τις μεταβλητές μας. Αυτό αντιπροσωπεύει μια οικογένεια μετασχηματισμών ισχύος που ενσωματώνει και επεκτείνει τις παραδοσιακές επιλογές. Συνεπώς, συνεχίζουμε την έρευνά μας χρησιμοποιώντας τον γενικό μετασχηματισμό Box-Cox, ο οποίος αντιπροσωπεύει μια πιθανή βέλτιστη πρακτική, επειδή είναι επιθυμητή η ομαλοποίηση των δεδομένων και η εξισορρόπηση της απόκλισης. Επιπλέον, αυτή η διατριβή εξετάζει την περίπτωση χρήσης των

αρχικών δεδομένων, πράγμα που σημαίνει καθόλου μετασχηματισμό και συγκρίνει την ικανότητα εκτιμήσεων εντός-δείγματος και εκτός – δείγματος στις προβλέψεις και την ικανότητα διαφόρων μοντέλων χρονολογικών σειρών.

Το κύριο ερευνητικό μας επίκεντρο, σε αυτό το κεφάλαιο, είναι ο τρόπος με τον οποίο η μετατροπή δεδομένων επηρεάζει τη διαδικασία υποδειγματοποίησης και πρόβλεψης χρονολογικών σειρών. Όπως δείχνουν τα εμπειρικά αποτελέσματα, ο μετασχηματισμός θεωρείται συχνά ότι σταθεροποιεί τη διακύμανση μιας σειράς, αλλά μπορεί επίσης να χρησιμοποιηθεί για να δώσει τις κλίση ή να μειώσει την επιρροή των ακραίων τιμών.

Το σχέδιο αυτού του κεφαλαίου είναι το ακόλουθο. Μετά από μια σύντομη εισαγωγή, συζητείται μια λεπτομερής βιβλιογραφική ανασκόπηση του μετασχηματισμού δεδομένων Box-Cox, ακολουθούμενη, από μια ενότητα της διαδικασίας μεθοδολογίας και της μεθοδολογίας του ανασχηματισμού των δεδομένων. Τα εμπειρικά αποτελέσματα του μετασχηματισμού δεδομένων παρουσιάζονται αρχικά για το μοντέλο εκθετικής εξομάλυνσης (ETS), σε εκτίμηση εντός και εκτός δείγματος. Η έρευνα προχωράει βαθύτερα και εξετάζει αφενός, την περίπτωση καθόλου μετασχηματισμού, που σημαίνει τη χρήση των αρχικών τιμών και από την άλλη πλευρά συνεχίζει να χρησιμοποιεί αρκετούς μαθηματικούς μετασχηματισμούς της οικογένειας μετασχηματισμών Box-Cox για διάφορες τιμές του λ . Επιπλέον, η έρευνα αναπτύσσεται σε περισσότερα μοντέλα χρονολογικών σειρών και σε διάφορους μετασχηματισμούς δεδομένων. Και τέλος τα διαστήματα εμπιστοσύνης των καλύτερων υποδειγμάτων παρουσιάζονται μαζί με τα συμπεράσματα της έρευνας.

6.2 Συμπεράσματα

Αυτή η εμπειρική έρευνα, καταλήγει στο συμπέρασμα ότι γενικά τα υποδείγματα εκθετικής εξομάλυνσης χώρου-χρόνου (ETS), που χρησιμοποιούν μετασχηματισμούς δεδομένων, δίνουν ενδείξεις μίας αρκετά καλής εφαρμογής στα δεδομένα, και είναι πιο ακριβή στην πρόβλεψη (εκτός-δείγματος) από τα ίδια υποδείγματα που όμως χρησιμοποιούν τα αρχικά δεδομένα, χωρίς καθόλου μετασχηματισμό. Τα υποδείγματα ETS αποδίδουν πολύ καλά, και προβλέπουν αρκετά σωστά τα χρονοδιαγράμματα αυτής της διατριβής με νέα επίπεδα πωλήσεων αυτοκι-

νήτων και γίνονται ακόμη καλύτερα όταν τα δεδομένα μεταμορφώνονται σε τιμές καταγραφής ή χρησιμοποιούν τη μέθοδο Box-Cox και Guerrero για τον υπολογισμό της καλύτερης αξίας λ. Επομένως, οι μαθηματικοί μετασχηματισμοί στα δεδομένα μπορούν να είναι αξιόπιστοι και να χρησιμοποιηθούν επειδή δίνουν καλά αποτελέσματα.σχέσεων με πιο ομοιογενή υπολείμματα, με άλλα λόγια σταθερή διακύμανση.

Επιπλέον, βραχυπρόθεσμα (για έως 6 μήνες) τα μοντέλα ETS προβλέπουν ότι οι τιμές προβλέπουν την κίνηση των επιπέδων πωλήσεων, αλλά μακροπρόθεσμα, έχουν την τάση να υπερεκτιμούν το επίπεδο πωλήσεων. Ωστόσο, μέσω της έρευνας μπορούμε να είμαστε 95 τις εκατό βέβαιοι ότι θα δώσει μια καλή πρόβλεψη για τις νέες πωλήσεις αυτοκινήτων, καθώς οι τιμές πρόβλεψης και οι πραγματικές τιμές είναι όλα στο εύρος της ζώνης του 95 της εκατό διαστήματος εμπιστοσύνης. Σε ορισμένες περιπτώσεις, το διάστημα εμπιστοσύνης (CI) δεν μπορεί να αποτυπώσει την πραγματική κίνηση της σειράς, ειδικά όταν οι αλλαγές στο επίπεδο των πωλήσεων είναι απότομες και υπάρχουν ενδείξεις για ένα αρκετά ταραχώδες οικονομικό περιβάλλον κατά τη συγκεκριμένη χρονική περίοδο.

Εξετάζοντας τις τρεις περιπτώσεις δεδομένων (πραγματικές τιμές, λογαριθμημένες τιμές και μετασχηματισμός τιμών με Box-Cox με τη μέθοδο Guerrero για την επιλογή του λ) στα τέσσερα διαφορετικά υποσύνολα και στις τρεις διαφορετικές εταιρείες, μπορούμε να συμπεράνουμε ότι ο λογαριθμικός μετασχηματισμός των χρονολογικών σειρών δίνει στην πλειοψηφία των εμπειρικών αποτελεσμάτων την ελάχιστη τιμή στα μέτρα αρίβειας των προβλέψεων και θα πρέπει να προτιμάτε. Ωστόσο υπάρχουν και ορισμένες εξαιρέσεις, που οδηγούν περισσότερο στην προτίμηση της επιλογής του μετασχηματισμού τω τιμών με τη μέθοδο του Box-Cox και την μέθοδο του Guerrero. Επομένως, η μελέτη των διαφορετικών μετασχηματισμών σε σύγκριση με την μελέτη των αρχικών τιμών ήταν χρήσιμη σε αυτό το κεφάλαιο, γιατί αποδεικνύεται τελικά ότι ο λογαριθμικός μετασχηματισμός και ο μετασχηματισμός Box-cox είναι χρήσιμοι, εύκολοι στην εφαρμογή και δίνουν ερμηνεύσιμα αποτελέσματα, που μπορούν να ωφελήσουν τους υπεύθυνους λήψης αποφάσεων σε μια εταιρεία.

Τέλος, είναι δύσκολο να είμαστε σίγουροι ότι ένας μοναδικός τύπος υποδείγματος χρονολογικών σειρών ή ένας συγκεκριμένος τύπος μετασχηματισμού δεδομένων είναι ο καλύτερος για όλες τις σειρές όλων των εποχών που να μπορεί γενικά να αποτυπώσει την κίνηση της

ροής των πωλήσεων ή να κάνει τις καλύτερες προβλέψεις. Επομένως, ο ερευνητής πρέπει να εξετάζει προσεκτικά τη φύση της σειράς κάθε φορά αλλά και τα εμπειρικά αποτελέσματα που λαμβάνονται από αυτή τη σειρά και θα πρέπει να ερμηνευθούν τα αποτελέσματα με προσοχή. Αυτό μπορεί να οδηγήσει σε περαιτέρω έρευνα στον τομέα των υποδειγμάτων χρονολογικών σειρών και μετασχηματισμού των δεδομένων και σε περισσότερη έρευνα στον τομέα της ελληνικής λιανικής αγοράς γενικότερα.

Κεφάλαιο 7

Συνδιασμός Προβλέψεων.

7.1 Εισαγωγή

Ο θεωρητικό θεμέλιος λίθος του συνδυασμού προβλέψεων ξεκίνησε πριν από πέντε δεκαετίες, όταν το έτος 1969, οι Bates-Granger έγραψαν το διάσημο άρθρο τους για το συνδυασμό των προβλέψεων. Ως γνωστό, ο συνδυασμός προβλέψεων από διάφορα μεμονωμένα υποδείγματα, συχνά οδηγεί σε καλύτερη ακρίβεια των προβλέψεων. Επομένως, ένας εύκολος τρόπος για να βελτιωθεί η ακρίβεια των προβλέψεων μας, είναι να χρησιμοποιηθεί ένας συνδυασμός προβλέψεων πολλών υποδειγμάτων στις ίδιες χρονικές σειρές, όπως για παράδειγμα να συνδυάσουμε με ίδια βαρύτητα τις προκύπτουσες προβλέψεις ή να χρησιμοποιηθεί διαφορετική βαρύτητα σε αυτές σταθμίζοντας με υψηλότερο συντελεστή τις καλύτερες προβλέψεις και ούτω καθεξής.

Οι επικριτές αυτής της μεθόδου υποστηρίζουν ότι ο συνδυασμός δεν είναι μια έγκυρη πρόταση εάν μία από οι μεμονωμένες προβλέψεις δεν διαφέρουν σημαντικά από το βέλτιστο αποτέλεσμα. Ωστόσο, ο συνδυασμός των προβλέψεων από πολύ παρόμοια υποδείγματα έχει αποδειχθεί σημαντικός. Μέχρι σήμερα έχει γίνει σημαντική έρευνα σχετικά με τη χρήση σταθμισμένων μέσων όρων ή κάποια άλλη πιο περίπλοκη προσέγγιση συνδυασμού. Μια εκτενής ανασκόπηση της βιβλιογραφίας, των τεχνικών και των εφαρμογών συνδυασμών προβλέψεων μπορεί να βρεθεί στο Clemen, 1989 όπου ισχυρίζεται ότι τα αποτελέσματα την έρευνας του σχεδόν ομόφωνα δείχνουν ότι ο συνδυασμός πολλαπλών προβλέψεων οδηγεί σε αυξημένη ακρίβεια των προβλέψεων. Σε πολλές περιπτώσεις, μπορεί κανείς να κάνει δραματικές βελτιώσεις

απόδοσης απλώς με τον μέσο όρο των προβλέψεων.

Υπάρχει ωστόσο μια δυσκολία στον προσδιορισμό του σωστού μοντέλου συνδυασμού προβλέψεων και αυτό δημιουργεί μια μικρή αβεβαιότητα. Κάποιοι υποστηρίζουν ότι οι προβλέψεις που βασίζονται σε οικονομετρικά μοντέλα, καθένα από τα οποία έχει πρόσβαση στο ίδιο σύνολο πληροφοριών, δίνουν όχι και τόσο καλούς συνδυασμούς προβλέψεων. Είναι ίσως καλύτερα να ταξινομούνται τα μεμονωμένα μοντέλα και να δημιουργείται ένα υπόδειγμα που να περιέχει τα χρήσιμα χαρακτηριστικά των αρχικών υποδειγμάτων .

Γενικά, το «καλύτερο» υπόδειγμα μπορεί να βρίσκεται στη λίστα υποψηφίων υποδειγμάτων της έρευνας ή όχι και ακόμη και αν το πραγματικό μοντέλο τυχαίνει να συμπεριληφθεί, το έργο της εύρεσης του πραγματικού μοντέλου μπορεί να είναι πολύ διαφορετικό από αυτό της εύρεσης του καλύτερου μοντέλου για τους σκοπούς ακριβών προβλέψεων.

Επιπλέον, είναι επίσης αποδεκτό ότι διαφορετικά μοντέλα πρόβλεψης παρέχουν διαφορετικά αποτελέσματα σε διαφορετικές χρονικές περιόδους. Έτσι, η επιλογή ενός μοντέλου πρόβλεψης ως «καλύτερου» φέρει τον κίνδυνο να καταλήξει σε ένα μοντέλο, το οποίο είναι ακριβές μόνο όταν αξιολογείται χρησιμοποιώντας κάποιο δείγμα επικύρωσης, αλλά μπορεί να αποδειχθεί αναξιόπιστο, όταν εφαρμόζεται σε νέα δεδομένα.

Γενικά, ο συνδυασμός προβλέψεων μειώνει τις πληροφορίες σε ένα φορέα προβλέψεων, σε ένα μόνο συνοπτικό μέτρο χρησιμοποιώντας ένα σύνολο σταθμισμένων βαρών συνδυασμού των προβλέψεων. Ο βέλτιστος συνδυασμός προβλέψεων πέρα από την ίση στάθμιση των προβλέψεων μπορεί να επιλέξει τη στάθμιση βαρών που ελαχιστοποιούν την αναμενόμενη απώλεια της συνδυασμένης πρόβλεψης. Η τεχνική αυτή δίνει μεγαλύτερα βάρη σε πιο ακριβείς προβλέψεις και μικρά λάθη εκτίμησης, και μικρότερα βάρη σε ανακριβείς προβλέψεις και μεγάλα λάθη εκτίμησης με άλλα λόγια σε κακή προδιαγραφή μοντέλου.

Σε πολλές περιπτώσεις, μπορεί κανείς να κάνει δραματικές βελτιώσεις απόδοσης απλά χρησιμοποιώντας έναν απλό μέσο όρο στις προβλέψεις, και επιπλέον αυτή η μέθοδος έχει αποδειχθεί δύσκολο να ξεπεραστεί. Ο συνδυασμός έχει μεγάλες δυνατότητες μείωσης της αβεβαιότητας που προκύπτει από την αναγκαστική επιλογή ενός μεμονωμένου υποδείγματος. Οι απλοί συνδυασμοί μεθόδων στη βιβλιογραφία προσπαθούν να βελτιώσουν τις μεμονωμένες προβλέψεις, ενώ ο πιο προχωρημένος στόχος είναι πάντα η εύρεση του υποδείγματος με την

καλύτερη προβλεπτική ικανότητα.

Σε αυτό το κεφάλαιο, ο ερευνητής θα παρέχει ολοκληρωμένη εφαρμογή κοινών τρόπων με τους οποίους μπορούν να συνδυαστούν οι προβλέψεις. Διάφορες μέθοδοι εκτίμησης θα εξηγηθούν για τη δημιουργία μιας συνδυασμένης πρόβλεψης και θα εφαρμοστούν σε διάφορα σύνολα δεδομένων προκειμένου να εξορθολογιστούν και να απεικονιστούν τα αποτελέσματα του συνδυασμού.

Το σχέδιο αυτού του κεφαλαίου είναι το εξής: ξεκινάμε με την εισαγωγή της θεωρίας συνδυασμού προβλέψεων και μια εκτεταμένη ανασκόπηση της βιβλιογραφίας. Συνεχίζουμε με μία αναφορά στη μεθοδολογία που χρησιμοποιείται και τη ομαδοποίηση των τεχνικών συνδυασμού προβλέψεων σε δύο κατηγορίες: πρόβλεψη συνδυασμού με ή χωρίς σύνολο ελέγχου. Αυτό το υποσύνολο ελέγχου απαιτείται για την εκτίμηση των σταθμισμένων βαρών των μεμονωμένων προβλέψεων. Έτσι, η τεχνική του απλού μέσου συνδυασμού προβλέψεων, που λειτουργεί χωρίς ένα υποσύνολο ελέγχου, εισάγεται σε τέσσερις (4) διαφορετικούς συνδυασμούς που είναι εύκολο να εφαρμοστούν και είναι δύσκολο να νικηθούν, λόγω των εξαιρετικών αποτελεσμάτων τους.

Από την άλλη πλευρά, εξηγούνται και εφαρμόζονται δύο (2) πιο περίπλοκες τεχνικές που χρειάζονται ένα υποσύνολο ελέγχου για τον υπολογισμό τους, όπως οι τεχνικές των Bates-Granger (1969) και των Newbold-Granger (1974). Επιπλέον εξηγούμε τη διαδικασία δημιουργίας δεδομένων αυτής της εμπειρικής έρευνας. Στην συνέχεια παρουσιάζονται τα εμπειρικά αποτελέσματα της έρευνας, τόσο αριθμητικά όσο και γραφικά, για τους έξι (6) διαφορετικούς συνδυασμούς προβλέψεων. Τέλος, δίνονται τα συμπεράσματα της εμπειρικής έρευνας του συνδυασμού προβλέψεων.

7.2 Συμπεράσματα

Τα αποτελέσματα αυτής της μελέτης αξιολόγησης δείχνουν ξεκάθαρα ότι είναι δυνατό να συνδυαστούν μοντέλα προβλέψεων μίας μεταβλητής, για να επιτευχθεί καλύτερη ακρίβεια πρόβλεψης, σε σύγκριση με την απλή επιλογή του καλύτερου μεμονωμένου μοντέλου πρόβλεψης. Αυτά τα στοιχεία συμβαδίζουν με προηγούμενες έρευνες (Clemen 1989). Ωστόσο, η επιλογή

μεθόδου συνδυασμού προβλέψεων είναι σημαντική δεδομένου ότι ορισμένες από τις μεθόδους αποδίδουν πολύ χειρότερα από τις χειρότερες μεθόδους μεμονωμένων προβλέψεων.

Στην παρούσα διατριβή εξετάζονται διάφορα υποδείγματα χρονολογικών σειρών ως προς την καλύτερη προσαρμογή τους στα δεδομένα αλλά και την προβλεπτική τους ικανότητα. Το έργο της επιλογής του καταλληλότερου υποδείγματος για την πρόβλεψη αποδεικνύεται πολύ δύσκολο. Σε αυτό το τελευταίο κεφάλαιο, προτείνεται η χρήση μιας συνδυαστικής μεθόδου για τον συνδυασμό διαφορετικών υποψήφιων μοντέλων, αντί της επιλογής ενός μεμονωμένου μοντέλου. Γνωρίζοντας ότι υπάρχει μεγάλη αβεβαιότητα στην εύρεση του καλύτερου μεμονωμένου υποδείγματος, στην περίπτωση της εμπειρικής εφαρμογής στην νέα ελληνική αγορά αυτοκινήτων, ο συνδυασμός των υποδειγμάτων μπορεί να μειώσει την αστάθεια της πρόβλεψης και επομένως να βελτιώσει την ακρίβεια των προβλέψεων μας.

Σύμφωνα με την έρευνά μας, υπάρχουν σημαντικές ενδείξεις ότι ο συνδυασμός προβλέψεων είναι επωφελής, όσον αφορά τη μείωση των σφαλμάτων πρόβλεψης καθώς και τη μείωση της αβεβαιότητας επιλογής του καταλληλότερου υποδείγματος, καθώς ο ερευνητής δεν υποχρεούται να επιλέξει ένα μόνο μοντέλο. Επιπλέον, είναι μια καλή στρατηγική για την αντιστάθμιση του κινδύνου. Τα συνδυασμένα υποδείγματα πρόβλεψης καθώς και τα εμπειρικά αποτελέσματα δείχνουν το πλεονέκτημα αυτής της μεθόδου σε σχέση με την επιλογή ενός μεμονωμένου υποδείγματος για την περίπτωση του λιανικού εμπορίου αυτοκινήτων που εξετάζουμε.

Οι συνδυασμοί των προβλέψεων, δημιουργούνται από διαφορετικά υποδείγματα προβλέψεων χρονολογικών σειρών, που είτε έχουν σταθμίσει ίσα τις προβλέψεις είτε έχουν δώσει διαφορετική βαρύτητα σε κάθε μια από αυτές. Υποκινήθηκαν κυρίως λόγω των διαφοροποιήσεων που προκύπτουν στις προβλέψεις κάθε μοντέλου και λόγω των αβέβαιων οικονομικών συνθηκών.

Ερευνητές υποστηρίζουν ότι απλά, ισχυρά σχήματα εκτίμησης τείνουν να λειτουργούν καλά σε μικρά δείγματα, όπου η εκτίμηση των σφαλμάτων γίνεται με συντελεστές βαρύτητας. Υπάρχουν ενδείξεις ότι ακόμη και αν δεν παρέχουν πάντα τις πιο ακριβείς προβλέψεις, οι συνδυασμοί προβλέψεων, ιδιαίτερα όσοι σταθμίζουν ίσα τα υποδείγματα, γενικά προσφέρουν μια καλή απόδοση και έτσι από άποψη κινδύνου, αντιπροσωπεύουν μια σχετικά ασφαλή επιλογή. Εμπειρικά, αυτή η ερευνητική εργασία απέδειξε ότι οι απλές προβλέψεις συνδυασμού υποδειγμάτων, λειτουργούν καλά για τις πωλήσεις νέων αυτοκινήτων στην ελληνική αγορά.

Τα αποτελέσματα της συνδυασμένης μεθόδου πρόβλεψης στο παράδειγμα των ερευνητικών μας δεδομένων συνοψίζονται ως εξής:

- Ένας συνδυασμός μοντέλων χρονολογικών σειρών φαίνεται να παράγει προβλέψεις πιο κοντά στις πραγματικές τιμές της σειράς και ως εκ τούτου δίνει μια καλύτερη πρόβλεψη για τις μεταβλητές μας από τις μεμονωμένες προβλέψεις.
- Ο απλός μέσος συνδυασμός με ίση στάθμιση των προβλέψεων διαφορετικών υποδειγμάτων είναι συνήθως το καλύτερο μοντέλο για σκοπούς πρόβλεψης. Εύκολο να υπολογιστεί δύσκολο να νικηθεί !

Αυτή η μελέτη υποδηλώνει ότι οι συνδυαστικές προβλέψεις είναι σχεδόν βέβαιες ότι θα ξεπεράσουν τις επιμέρους προβλέψεις και θα αποφύγουν τον κίνδυνο πλήρους αποτυχίας των προβλέψεων. Επομένως, σε περιπτώσεις όπου τα μοντέλα πρόβλεψης είναι διαθέσιμα και ο ερευνητής πρέπει να δημιουργήσει προβλέψεις, αλλά δεν είναι βέβαιο ως προς το ποιο μοντέλο είναι πιθανό να δημιουργήσει τις καλύτερες προβλέψεις, ο συνδυασμός των προβλέψεων από διάφορα εναλλακτικά μοντέλα θα ήταν ο καλύτερος και ασφαλέστερος τρόπος προόδου.

Συνοψίζοντας, οι μέθοδοι πρόβλεψης χρονολογικών σειρών , και ειδικά ο συνδυασμός των προβλεψεων των μεμονομένων υποδειγμάτων, αποδεικνύεται ότι εφαρμόζεται με επιτυχία στις νέες πωλήσεις αυτοκινήτων, δηλαδή σε δεδομένα μάρκετινγκ στην ελληνική αγορά και παρέχει αξιόπιστες προβλέψεις.