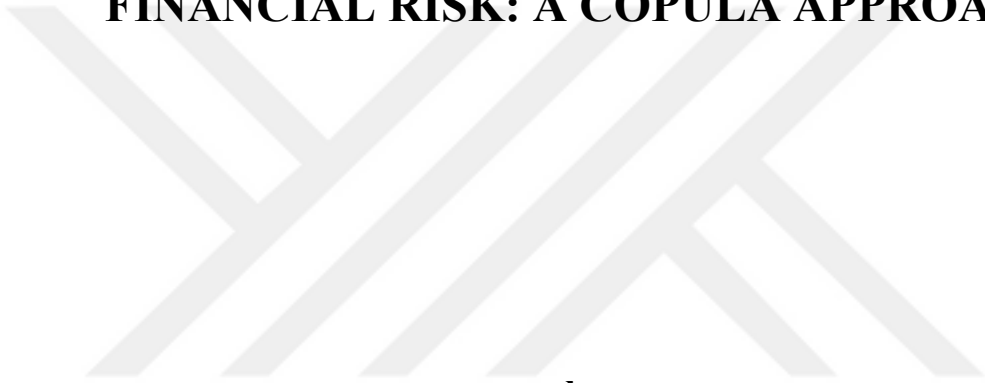


DOKUZ EYLÜL UNIVERSITY
GRADUATE SCHOOL OF NATURAL AND APPLIED SCIENCES

**MODELLING DEPENDENCE STRUCTURE FOR
FINANCIAL RISK: A COPULA APPROACH**



by
Tolga YAMUT

July, 2018
İZMİR

MODELLING DEPENDENCE STRUCTURE FOR FINANCIAL RISK: A COPULA APPROACH

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

**by
Tolga YAMUT**

**July, 2018
İZMİR**

M.Sc THESIS EXAMINATION RESULT FORM

We have read the thesis entitled “**MODELLING DEPENDENCE STRUCTURE FOR FINANCIAL RISK: A COPULA APPROACH**” completed by **TOLGA YAMUT** under supervision of **PROF. DR. BURCU ÜÇER** and we certify that in our opinion it is fully adequate, in scope and in quality, as a thesis for the degree of Master of Science.



Prof. Dr. Burcu ÜÇER

Supervisor



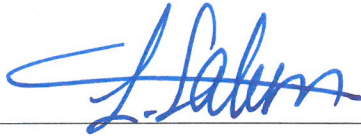
Assoc. Prof. Dr. Tüba YILDIZ

Jury Member



Assoc. Prof. Dr. Güler KENALBAY

Jury Member



Prof. Dr. Latif SALUM

Director

Graduate School of Natural and Applied Sciences

ACKNOWLEDGEMENTS

I would like to show my gratitude to my supervisor, Prof. Dr. Burcu ÜÇER for her vital assistance. Her heartwarming attitude and encouragement has escalated the finishing process of this work.

I wish to present my warm thanks to my dear friends; Ayça ÖLMEZ, Aylin GÖÇÖĞLU, Çağla ÇINAR, Hasan KIRAN, İrem ÖZALP, Sami AKDENİZ and Yusuf SEVİL for their morally and emotionally support.

Lastly, I would like to pay special thankfulness and appreciation to my Mother and Father for their endless support and motivation.

Tolga YAMUT

MODELLING DEPENDENCE STRUCTURE FOR FINANCIAL RISK: A COPULA APPROACH

ABSTRACT

Copula distributions are frequently being used in financial studies. The reason of this preference is that unlike correlation based models, copulas provide a comprehensive approach to interpret the dependence structures. Financial time series have heavy tailed characteristics and volatility. This situation requires a copula approach to analyze a market's dependence structure and forecast its risks while considering extremes. Also, joint behaviour of multiple markets hence multiple investment potentials can be examined without the concern of marginal distributions with the help of copulas. In order to clarify the risks of an investment, financial structure of the targeted market requires sensitive modeling that acknowledges the volatile nature. Value-at-Risk (VaR) models help the investors to forecast the consequences of the investment for a particular market. In this study, we explicitly analyze the relation between changes in oil prices and stock market returns for emerging markets. Our goal is to investigate how Brent oil and emerging markets are affected by each other and compute risk measures for these two assets. All datasets are tested via goodness of fit test to determine proper copula distributions. We propose to forecast the Value-at-Risk of the portfolios on the basis of bivariate copulas using nonparametric estimates of the coefficient of tail dependence which is estimated to fit the data according to its extreme observations. Changes in tail dependence by time are visualized by comparing parametric and non-parametric estimations through a simulation study. Then, VaR for the scenario of equally weighted assets are computed by simulating cumulative returns. Furthermore, validity of VaR models are tested by applying backtesting. Lastly, Conditional Value-at-Risk (CoVaR) estimation is carried out to observe market dynamics when Brent market is under stress.

Keywords: Copula, value at risk, conditional value at risk, tail dependence, emerging markets

FINANSAL RİSKLER İÇİN BAĞIMLILIK YAPISINI MODELLEME: BİR KOPULA YAKLAŞIMI

ÖZ

Kopula dağılımları sıklıkla finansal çalışmalarda kullanılmaktadır. Korelasyona dayalı modellerin aksine bağımlılık yapılarını yorumlamak üzere daha kapsamlı bir yaklaşım sağlaması bu tercihin temelini oluşturmaktadır. Finansal zaman serileri ağır kuyruklu ve oynak bir yapıya sahiptir. Bu durum, bir piyasanın bağımlılık yapısını analiz etmek ve uç gözlemleri hesaba katarak piyasa risklerini öngörmek için bir kopula yaklaşımı gerektirmektedir. Ayrıca birçok piyasanın dolayısıyla birçok yatırım potansiyelinin ortak davranışları kopula sayesinde marjinal dağılımlardan bağımsız olarak incelenebilmektedir. Riske maruz değer, belirli bir pazar için yatırımın sonuçlarını tahmin etmede yatırımcıya yardım eder. Bir yatırımın risklerini açıklığa kavuşturmak için, hedeflenen piyasanın finansal yapısı, değişken yapıyı tanıyan hassas modelleme gerektirir. Bu çalışmada, gelişen piyasalar için, petrol ve hisse senedi getirileri arasındaki değişimin detaylı analizi yapılmıştır. Çalışmanın amacı, Brent petrol fiyatı ile gelişmekte olan piyasaların birbirinden nasıl etkilendiğini araştırmak ve ikili varlıkların risk ölçümlerini hesaplamaktır. Tüm veri setleri için uygun kopula dağılımları uyum iyiliği testi ile belirlenmiştir. Uç-değer gözlemlerine göre model tahmini yapılan veri seti için kuyruk bağımlılığı yapısının parametrik olmayan kestirimi kullanılarak iki değişkenli kopulalar temelinde portföylerin riske maruz değerlerinin öngörümlemesi önerilmiştir. Benzetim çalışması ile parametrik ve parametrik olmayan tahminler karşılaştırılarak zamana bağlı kuyruk yapısındaki değişim gözlenmiştir. Daha sonra, eşit ağırlıklı varlıklara ait portföy, birikimli getirilerin benzetimi ile riske maruz değer hesaplanmış ve riske maruz değer için geriye dönüş testi ile geçerliliği sınanmıştır. Son olarak, Brent piyasası stres altında olduğunda market dinamiklerini gözlemlemek için Koşullu Riske Maruz Değer kestirimi uygulanmıştır.

Anahtar kelimeler: Kopula, riske maruz değer, koşullu riske maruz değer, kuyruk bağımlılığı, gelişmekte olan piyasalar

CONTENTS

	Page
M.Sc THESIS EXAMINATION RESULT FORM.....	ii
ACKNOWLEDGEMENTS	iii
ABSTRACT	iv
ÖZ.....	v
LIST OF FIGURES	ix
LIST OF TABLES.....	xi
 CHAPTER ONE – INTRODUCTION.....	 1
 CHAPTER TWO – GENERAL ASPECTS OF COPULA.....	 3
2.1 Preliminaries	3
2.2 Copula.....	4
2.2.1 Sklar’s Theorem.....	7
2.2.2 Copula Concept With Random Variables	9
2.2.2.1 Probability Density Function of a Copula	10
2.2.2.2 Useful Features of Copula	11
2.2.3 Survival Copula	13
 CHAPTER THREE – FAMILIES OF COPULA.....	 16
3.1 Archimedean Copula	16
3.1.1 Gumbel Copula.....	19
3.1.2 Frank Copula	20
3.1.3 Clayton Copula.....	20
3.2 Elliptical Copula	20
3.2.1 Normal Copula	21
3.2.2 Student’s t copula.....	22
3.3 Visualization Of Copula Types	22

CHAPTER FOUR – DEPENDENCE	26
4.1 Dependence Structures.....	26
4.2 Concordance Measures	27
4.2.1 Kendall’s tau.....	27
4.3 Properties of Dependence.....	33
4.3.1 Quadrant Dependence	33
4.3.2 Tail Monotonicity	34
4.3.3 Tail Dependence	35
4.3.3.1 Non-parametric Approach to Tail Dependence.....	37
CHAPTER FIVE – METHODS OF PARAMETER ESTIMATION AND GOODNESS OF FIT TESTING.....	39
5.1 Parameter Estimation.....	39
5.2 Parametric Ways.....	40
5.2.1 Exact Maximum Likelihood Estimation.....	40
5.2.2 IFM Method	41
5.3 Non-Parametric Ways	42
5.3.1 Maximum Pseudolikelihood Estimator	42
5.3.2 Estimation of Parameter by Kendall’s Tau	43
5.4 A Fundamental Problem of Modeling	43
5.4.1 Kolmogorov-Smirnov Test Statistic.....	45
5.4.2 Cramer von Mises Test Statistic	45
5.4.2.1 Bootstrap Procedure.....	46
CHAPTER SIX – MANAGING AND MEASURING RISKS	47
6.1 Overview	47
6.2 Financial Time Series	49
6.2.1 GARCH Models	52
6.2.1.1 GARCH Process	52
6.2.1.2 EGARCH Process.....	53

6.3 Foremost Goals of Risk Measurement.....	54
6.4 Ways of Measuring Risks.....	54
6.4.1 Notional Amount Approach.....	55
6.4.2 Factor-Sensitivity Measures.....	55
6.4.3 Risk Measures Based on Loss Distributions.....	55
6.4.4 Scenario-Based Risk Measures	56
6.5 Value-at-Risk	56
6.6 Backtesting	57
6.6.1 POF-Test.....	58
6.6.2 Mixed Kupiec Test.....	58
6.7 Conditional Value at Risk.....	59
CHAPTER SEVEN – APPLICATION	63
7.1 Data and Empirical Results	63
7.2 Tail Dependence Analysis	74
7.3 Value-at-Risk	74
7.4 Conditional Value-at-Risk	77
CHAPTER EIGHT – CONCLUSION	82
REFERENCES.....	84

LIST OF FIGURES

	Page
Figure 2.1 Probability within the rectangle: $P(u_1 \leq x_1 \leq u_2, v_1 \leq x_2 \leq v_2)$	6
Figure 3.1 Gumbel density via $\rho = 0.3$	23
Figure 3.2 Gumbel level curves of density via $\rho = 0.3$	23
Figure 3.3 Gumbel density via $\rho = 0.7$	23
Figure 3.4 Gumbel level curves of density via $\rho = 0.7$	23
Figure 3.5 Frank density via $\rho = 0.3$	23
Figure 3.6 Frank level curves of density via $\rho = 0.3$	23
Figure 3.7 Frank density via $\rho = 0.7$	23
Figure 3.8 Frank level curves of density via $\rho = 0.7$	23
Figure 3.9 Clayton density via $\rho = 0.3$	23
Figure 3.10 Clayton level curves of density via $\rho = 0.3$	23
Figure 3.11 Clayton density via $\rho = 0.7$	24
Figure 3.12 Clayton level curves of density via $\rho = 0.7$	24
Figure 3.13 Normal density via $\rho = 0.3$	24
Figure 3.14 Normal level curves of density via $\rho = 0.3$	24
Figure 3.15 Normal density via $\rho = 0.7$	24
Figure 3.16 Normal level curves of density via $\rho = 0.7$	24
Figure 3.17 Student's t density via $\rho = 0.3$ and d.o.f.=4	24
Figure 3.18 Student's t level curves of density via $\rho = 0.3$ and d.o.f.=4	24
Figure 3.19 Student's t density via $\rho = 0.7$ and d.o.f.=4	24
Figure 3.20 Student's t level curves of density via $\rho = 0.7$ and d.o.f.=4	24
Figure 7.1 Daily closings through 2010-2015	65
Figure 7.2 Autocorrelation graphics of returns and squared returns of datasets	66
Figure 7.3 Autocorrelation graphics after EGARCH(1,1) process	69
Figure 7.4 Probability plots of each datasets for Kolmogorov-Smirnov Test	71
Figure 7.5 Visuals of parametric and non parametric tail coefficients	75
Figure 7.6 Parametric and nonparametric estimations of VaR for 300 trading days	77

Figure 7.7 CoVaR estimations of through 300 trading days with $\alpha = 0.05$ and $\beta = 0.05$, y-axis being CoVaR values and risk contribution and x-axis being trading days 80



LIST OF TABLES

	Page
Table 7.1 Descriptive Statistics.....	65
Table 7.2 Estimation of EGARCH(1,1) parameters (standard errors).....	68
Table 7.3 Kendall's tau values of the financial dataset.....	71
Table 7.4 Parameter Estimations (Standard Errors) [Maximized log likelihood].....	72
Table 7.5 p-values (Test statistics) of Cramer-von Mises Test	73
Table 7.6 Simulated VaR estimates for fifty-fifty Brent and Index assets with 90% confidence and p-values for the VaR models determined by POF test	73
Table 7.7 Descriptive Statistics for VaR Estimations	78
Table 7.8 Backtesting results with 0.05 significance	78
Table 7.9 Descriptive Statistics for $\Delta CoVaR$ in Percentage.....	81
Table 7.10 Individual and joint violation numbers where joint violations investigated among the VaR violations of Brent which has number of 34 trading days	81

CHAPTER ONE

INTRODUCTION

The word copula has the meaning "link", "tie", "bond" in latin. These distributions help us describe a joint distribution by using its marginal distributions. They construct a bridge between joint distribution and its marginals. As another sense, one can say copula distributions are joint distributions whose margins are uniform and lay inside the interval $[0, 1]$. As Fisher (1997) stated there are two main goals in copula; firstly, they help us to investigate dependencies without using scale measures of dependence and secondly, by using them bivariate distribution families can be constructed.

To build up main points before the formal definition of copula first consider (X, Y) as a pair of random variables. The marginal distribution functions are $F(x) = P(X \leq x)$ and $G(y) = P(Y \leq y)$, respectively and a joint distribution $H(x, y) = P(X \leq x, Y \leq y)$. One can see for each observation, we have three numbers of $F(x)$, $G(x)$ and $H(x, y)$ which are inside $[0, 1]$. As another point of view, each pair of real numbers correspond to a point $(F(x), G(y))$ in the unit square $[0, 1] \times [0, 1]$ and this point carries us to a real value of $H(x, y)$ in $[0, 1]$. This observation helps to see the linking between a joint distribution and their marginal distribution which copula provides. This feature helps in the matters of joint analysis greatly since one has not to deal with marginal distributions one by one. Copula's modern ways of measuring dependencies and practical usage for analysing comovements made them a solid tool of analysis in financial affairs. They mostly used in financial analysis and risk assessments.

In order to make a reliable risk assessment for a particular portfolio, one needs to establish an accurate dependence structures for the prices. Especially when one plans to invest on multiple assets, the model of dependence structure helps to make connections between variables and interpret the risk situation. Oil and stock assets are very active market areas which require careful examination. Volatilities in the portfolio of oil markets and stock indices have a very dynamic nature. Capturing the patterns of their co-movements of both portfolio is a necessity to forecasting both of

their returns. Analyzing potential of an investment demands a rigorous statistical inference since being sure is very important when one deals with the very volatile market areas. According to the widely accepted idea in empirical literature, financial asset returns have a nonlinear and time varying dependency Longin & Solnik (2001), Ang & Chen (2002). Indeed flexibility of copula modeling works great instead of models based on correlation when one deals with nonlinear dependence structures Embrechts et al. (2002), Li (2000). In this study, we examine the risk environment between various emerging markets and Brent oil. In order to achieve this, after modeling volatility of financial data with EGARCH process, proper copula distributions are calibrated, then based on the simulation of returns VaR and CoVaR models are conducted. Also performances of methods are compared; parametric and nonparametric approach for VaR, different condition setups for CoVaR.

In the second chapter, we touch upon main concept and features of theory of copula. Then, we represent the well known copula types and their properties in the third chapter. Fourth one is devoted to important measure of associations and main dependence types with copula. The general methodologies of inference is revisited in chapter five. The key ideas of modelling return series and risk measures are given in chapter six. All of the details and results are mentioned in chapter seven.

CHAPTER TWO

GENERAL ASPECTS OF COPULA

2.1 Preliminaries

In this section we give the notion of the term ‘2-increasing’. Let $\mathbf{R}$ be ordinary real line $(-\infty, \infty)$ and $\overline{\mathbf{R}}$ denotes the extended real line $[-\infty, \infty]$. Hence we have the extended real plane $\overline{\mathbf{R}}^2$, a rectangle B in the extended real plane is $B = [x_1, x_2] \times [y_1, y_2]$, where vertices are (x_1, y_1) , (x_1, y_2) , (x_2, y_1) and (x_2, y_2) .

On the other hand, the unit square is denoted by $\overline{\mathbf{I}}^2$. It is a product of $I \times I$ where $I = [0, 1]$. Bivariate real function H has the domain, $\text{Dom}H$ subset of $\overline{\mathbf{R}}^2$ and has the range, $\text{Ran}H$ subset of $\overline{\mathbf{R}}$.

Definition 2.1.1. Let S_1 and S_2 be nonempty subsets inside $\overline{\mathbf{R}}^2$ and let H be a bivariate real function where $\text{Dom}H = S_1 \times S_2$. Additionally let $B = [x_1, x_2] \times [y_1, y_2]$ be a rectangle whose vertices lay in $\text{Dom}H$. The H-volume of B is given as

$$V_H(B) = H(x_2, y_2) - H(x_2, y_1) - H(x_1, y_2) + H(x_1, y_1) \quad (2.1)$$

Definition 2.1.2. A bivariate real function is said to be 2-increasing if $V_H(B) \geq 0$ for all rectangles whose vertices are inside $\text{Dom}H$. A 2-increasing H is referred as the H-measure of B. Instead of the term 2-increasing, quasi-monotone is also used by some authors. Note that for H , there is no if and only if condition between the statements 2-increasing and nondecreasing. The following lemmas are important and help us for the establishment of definitions of subcopulas and copulas.

Lemma 2.1.1. Let S_1 and S_2 be nonempty subsets inside $\overline{\mathbf{R}}$ and let H be a 2-increasing function whose $\text{dom}H = S_1 \times S_2$. Also, let x_1, x_2 be in S_1 via $x_1 \leq x_2$ and let y_1, y_2 be in S_2 via $y_1 \leq y_2$. Then the function $t \mapsto H(t, y_2) - H(t, y_1)$ is nondecreasing on S_1 and $t \mapsto H(x_2, t) - H(x_1, t)$ is nondecreasing on S_2 .

We can improve the scope of this lemma by introducing the notion of grounded

function. Assume S_1 has a least element l_1 and S_2 has a least element l_2 . We say that H is grounded if $H(l_1, y) = H(x, l_2) = 0$ for all real pairs of (x, y) in $S_1 \times S_2$.

Lemma 2.1.2. Let S_1 and S_2 be nonempty subsets in $\overline{\mathbf{R}}$, and let H be a grounded 2-increasing function which has the domain $S_1 \times S_2$. Then H is nondecreasing on each subsets.

Consider S_1 has a greatest element g_1 and S_2 has a greatest element g_2 then we can say that a bivariate function H has margins, and marginal functions F and G given as

$$\text{Dom}F = S_1 \quad \text{and} \quad F(x) = H(x, g_2) \quad \forall x \in S_1$$

$$\text{Dom}G = S_2 \quad \text{and} \quad G(y) = H(g_1, y) \quad \forall y \in S_2$$

Example 2.1.1. Let H has the domain $[-1, 1] \times [0, \infty]$ and it is given as Nelsen (2006)

$$H(x, y) = \frac{(x+1)(e^y - 1)}{x + 2e^y - 1}.$$

Observe that, -1 is the least and 1 is the greatest element in the first subset, 0 is the least and ∞ is the biggest element in the second subset. The function H is grounded since $H(-1, y) = 0$ and $H(x, 0) = 0$, furthermore the margins of H are obtained by

$$G(y) = H(1, y) = 1 - e^{-y} \quad \text{and} \quad F(x) = H(x, \infty) = \frac{x+1}{2}.$$

2.2 Copula

Definition 2.2.1. The function C' is a bivariate subcopula which has the properties:

1. $\text{Dom}C' = S_1 \times S_2$, where the individual sets are the subsets of I .
2. C' is grounded and 2-increasing function.
3. $\forall u \in S_1$ and $\forall v \in S_2$

$$C'(u, 1) = u \quad \text{and} \quad C'(1, v) = v$$

Remark 2.2.1. For each pair of (u, v) in $\text{Dom}C'$, $0 \leq C'(u, v) \leq 1$ such that $\text{Ran}C'$ is subset which belongs to I .

Definition 2.2.2. The function C is a bivariate copula whose domain is I^2 . Or in another words, copula is a function $C : I^2 \rightarrow I$ which has the properties:

$$1. \forall u, v \in I,$$

$$C(u, 0) = C(0, v) = 0 \tag{2.2}$$

and

$$C(u, 1) = u \quad \text{and} \quad C(1, v) = v \tag{2.3}$$

$$2. \forall u_1, u_2, v_1, v_2 \in I \text{ with the conditions } u_1 \leq u_2 \text{ and } v_1 \leq v_2,$$

$$C(u_2, v_2) - C(u_2, v_1) - C(u_1, v_2) + C(u_1, v_1) \geq 0 \tag{2.4}$$

Remark 2.2.2. Copula can be described as $C(u, v) = V_C([0, u] \times [0, v])$ so in this case we can see $C(u, v)$ as assigned number in I to the rectangle $[0, u] \times [0, v]$. The 2-increasing property provides that the number assigned by C to all rectangles $[u_1, u_2] \times [v_1, v_2]$ in I^2 must be nonnegative. The property (2.4) also can be rewritten as,

$$\begin{aligned} &P(0 \leq x_1 \leq u_2, 0 \leq x_2 \leq v_2) - P(0 \leq x_1 \leq u_2, 0 \leq x_2 \leq v_1) - \\ &- P(0 \leq x_1 \leq u_1, 0 \leq x_2 \leq v_2) + P(0 \leq x_1 \leq u_1, 0 \leq x_2 \leq v_1) \geq 0 \end{aligned}$$

by visualizing Figure (2.1) it can be seen that the probability that it is between the rectangle and indeed nonnegativeness is preserved.

Theorem 2.2.1. Fréchet-Hoeffding Bounds Let C be a copula then $\forall (u, v) \in \text{Dom}C$

$$\max(u + v - 1, 0) \leq C(u, v) \leq \min(u, v)$$

Proof. Let (u, v) be any point in $\text{Dom}C$. Since we have the nondecreasing feature,

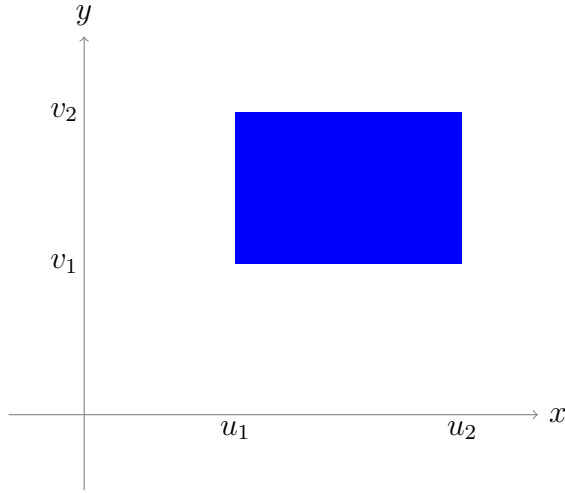


Figure 2.1 Probability within the rectangle: $P(u_1 \leq x_1 \leq u_2, v_1 \leq x_2 \leq v_2)$

$C(u, v) \leq C(u, 1) = u$ and $C(u, v) \leq C(1, v) = v$ holds. Thus clearly, we have upper bound, $C(u, v) \leq \min(u, v)$. Further, consider the volume $V_C([u, 1] \times [v, 1]) = C(1, 1) - C(1, v) - C(u, 1) + C(u, v) = 1 - v - u + C(u, v) \geq 0$ hence $C(u, v) \geq u + v - 1$. Surely $C(u, v) \geq 0$ so, we have the lower bound $C(u, v) \geq \max(u + v - 1, 0)$.

□

Remark 2.2.3. The bounds we examined above are also indeed copulas. They are often denoted by $M(u, v) = \min(u, v)$ and $W(u, v) = \max(u + v - 1, 0)$. As a result for every pair $(u, v) \in I^2$, any copula distribution C lays between

$$W(u, v) \leq C(u, v) \leq M(u, v) \quad (2.5)$$

The inequality (2.5) is known as *Fréchet-Hoeffding bounds*. M is referred as the Fréchet-Hoeffding upper bound and the W is referred as the Fréchet-Hoeffding lower bound. And another important copula is called by the name *product copula* that is $\Pi(u, v) = uv$. These three types of copulas are very generic and show up frequently in the subject of the copula.

2.2.1 Sklar's Theorem

The theorem holds a crucial role for the subject that it exists in the core of the theory of copula. Also, the theorem furnishes a starting point in many areas of Statistics. By the Sklar's Theorem, we can see the connection of copula with the multivariate distributions and their marginal distributions. In this section first, we recall the general concept of the distribution function.

Definition 2.2.3. A joint distribution function H has the domain $\overline{\mathbf{R}}^2$ satisfies the properties:

1. H is 2-increasing.
2. $H(x, -\infty) = H(-\infty, y) = 0$ and $H(\infty, \infty) = 1$.

Furthermore again we can get the distribution functions; $F(x) = H(x, \infty)$ and $G(y) = H(\infty, y)$.

Theorem 2.2.2. Sklar's Theorem

Let H be a joint distribution function whose marginal distributions are F and G . Then for the joint distribution function H , there exists a copula distribution $C \forall x, y \in \overline{\mathbf{R}}$ so that

$$H(x, y) = C(F(x), G(y)) \quad (2.6)$$

Additionally, C is a unique copula if both F and G are continuous. If the marginal are discontinuous, C depends on the $\text{Ran}F \times \text{Ran}G$. One can see that in (2.2.5) how copula distributions connect a joint distribution function to its marginals. With the other way around the function H is determined by the copula C and the marginals F and G in the equality (2.2.5).

Proof. The subcopula we defined earlier only a little different than the copula but in the proof of Sklar's Theorem it holds a rather important role (See Nelsen 2006). $\square$

Remark 2.2.4. In (2.6) the joint distribution H is determined in terms of C and margins F, G . The converse is also possible that by taking the inverses of each univariate marginals, a copula C can be expressed in terms of joint distribution and inverses of its marginals. This basic idea helps one to find a copula for a particular joint distribution.

Definition 2.2.4. Let F be a distribution function. A quasi-inverse of F whose domain is I , denoted as $F^{(-1)}$ and has the properties:

1. If $a \in \text{Ran}F$ then $F^{(-1)}(a) = x$ where $x \in \bar{\mathbf{R}}$ such that $F(x) = a$. Also $\forall a \in \text{Ran}F$ the identity relation holds; $F(F^{(-1)}(a)) = a$

2. If a is not the element of $\text{Ran}F$, the inverse is determined by

$$F^{(-1)}(a) = \inf(x \mid F(x) \geq a) = \sup(x \mid F(x) \leq a)$$

If F is strictly increasing, it has only one quasi-inverse which means we are just taking standard inverse operation, denoted as F^{-1} .

Corollary 2.2.1. Again let H be a joint distribution, with marginals F and G , and C be a copula where (2.6) holds. Further, let $F^{(-1)}$ and $G^{(-1)}$ quasi-inverses of margins. We have

$$C(u, v) = H(F^{(-1)}(u), G^{(-1)}(v)) \text{ where } u = F(x) \text{ } v = G(y) \quad (2.7)$$

This method is very useful to generate copula distributions from joint distribution functions.

Example 2.2.1. Let us investigate the copula of the joint distribution from Example 2.1.1: (See Nelsen 2006)

$$H(x, y) = \begin{cases} \frac{(x+1)(e^y-1)}{x+2e^y-1}, & (x, y) \in [-1, 1] \times [0, \infty], \\ 1 - e^{-y}, & (x, y) \in (1, \infty] \times [0, \infty], \\ 0, & \text{otherwise} \end{cases}$$

with the marginal distributions,

First take the inverses of marginal distributions that is $F^{-1}(u) = 2u - 1$ and $G^{-1}(v) = -\ln(1 - v)$ where all

$$F(x) = \begin{cases} 0, & x < -1, \\ \frac{x+1}{2}, & x \in [-1, 1], \\ 1, & x > 1 \end{cases} \quad G(y) = \begin{cases} 0, & y < 0, \\ 1 - e^{-y}, & y \geq 0, \end{cases}$$

$u, v \in [0, 1]$. Since $\text{Ran}F = \text{Ran}G = I$ plugging the inverses into the joint distribution as shown in (2.7) will give us the copula C :

$$C(u, v) = H(F^{(-1)}(u), G^{(-1)}(v)) = \frac{(2u - 1 + 1)(e^{-\ln(1-v)} - 1)}{2u - 1 + 2e^{-\ln(1-v)} - 1}$$

$$C(u, v) = \frac{uv}{u + v - uv}.$$

Example 2.2.2. As an another example, consider Gumbel's bivariate exponential distribution. We have the joint distribution H_θ such that

$$H_\theta = \begin{cases} 1 - e^{-x} - e^{-y} + e^{-(x+y+\theta xy)}, & x \leq 0, y \leq 0 \\ 0, & \text{otherwise} \end{cases}$$

The parameter θ lies in $[0, 1]$. First find the marginal distributions; $H_\theta(x, \infty) = F(x) = 1 - e^{-x}$ and $H_\theta(\infty, y) = G(y) = 1 - e^{-y}$. By taking the inverses of the marginals we have; $F^{-1}(u) = -\ln(1 - u)$ and $G^{-1}(v) = -\ln(1 - v)$ $\forall u, v \in I$. Again by plugging inverses, we find the copula for H_θ given by

$$C_\theta = u + v - 1 + (1 - u)(1 - v)e^{-\theta \ln(1-u)\ln(1-v)}.$$

2.2.2 Copula Concept With Random Variables

Now let us revisit the Sklar's Theorem with the concept of random variables and their distribution functions.

Theorem 2.2.3.

Let X and Y be the random variables and have the distribution functions F and G respectively, and have a joint distribution H . There exists a copula C which satisfies (2.6). Again C is unique if F and G are continuous. Otherwise, it is dependent on the $\text{Ran}F \times \text{Ran}G$. When the usage of the random variables is useful, a copula C of random variables X, Y can be denoted as C_{XY} .

2.2.2.1 Probability Density Function of a Copula

In this part, we deduce that the joint pdf of a copula is represented by the product of marginal pdfs and the density of copula.

Definition 2.2.5. Let F and G be continuous marginals of joint distribution H , bounded by the copula C , then the joint pdf of X and Y is described as

$$h(x, y) = f(x)g(y)c(F(x), G(y)) \quad \forall (x, y) \in \mathbf{R}^2 \quad (2.8)$$

where

$$f(x) = \frac{dF(x)}{dx}$$

$$g(y) = \frac{dG(y)}{dy}$$

And the density c is defined as

$$c(u, v) = \frac{\partial^2}{\partial u \partial v} C(u, v), \quad \forall (u, v) \in I^2 \quad (2.9)$$

Notice that the joint pdf involves a part which corresponds to the independence, i.e. $f(x)g(y)$. Then the joint pdf of H can also be interpreted as independent pdf $f(x)g(y)$ times $c(F(x), G(y))$. If $c = 1$ then joint pdf is described by the factors of marginals i.e. X and Y are independent.

Remark 2.2.5. The definition of density comes from the observation that $\frac{\partial^2}{\partial u \partial v} C(u, v)$ exists almost everywhere in I^2 . For details, see Michel (2005).

Definition 2.2.6. (Absolutely Continuous)

If $\forall u, v \in [0, 1]$ while $\frac{\partial C(u, v)}{\partial u \partial v}$ being the density of C

$$C(u, v) = \int_0^v \int_0^u \frac{\partial C(u, v)}{\partial u \partial v} du dv$$

then we say C is absolutely continuous. After this part, we can move on to the important results of the copula which makes it a useful tool in statistical inferences.

2.2.2.2 Useful Features of Copula

Theorem 2.2.4. Probability Integral Transform

Let $F_X(x)$ be a continuous distribution function and introduce another random variable as $Y = F_X(x)$, then Y is distributed uniformly on the interval $(0, 1)$.

Proof. Since $Y = F_X(x)$,

$$\begin{aligned} F_Y(y) &= P(Y \leq y) = P(F_X(x) \leq y) \\ &= P(X \leq F_X^{-1}(y)) \\ &= F_X(F_X^{-1}(y)) = y \quad 0 < y < 1 \end{aligned}$$

□

Remark 2.2.6. In the light of Sklar's theorem we investigated that by the Corollary 2.2.1, we can describe the copula in terms of inverse of marginals and their joint distribution. To do so, we assigned $U = F(x)$ and $V = G(y)$. As a result of probability integral transform U and V distribute uniformly on the interval $(0, 1)$. This aspect of copula provides that whatever distribution of the marginals are, one can

treat them as a uniform distribution. This outcome brings a computational ease at studies and way to deal with the marginal distributions without too much effort. For the following theorem, idea of independence for the random variables is given in terms of the product copula.

Theorem 2.2.5. *Let X and Y be continuous random variables. X and Y are independent if and only if $C_{XY} = \prod$*

Proof. Recall the product copula is; $\prod(u, v) = uv$. Trivially, this independence comes from the Theorem 2.2.3 and independence condition of random variables X and Y that is $H(x, y) = F(x)G(y)$. $\square$

In nonparametric processes of statistics, copula proves to be very helpful in the matter of *strictly monotone transformations* of random variables. Because under such transformations, copulas are either *invariant* or their change can be tracked easily. Remember the fact that if distribution function with random variable X is continuous and if α is a strictly monotone function whose domain has $\text{Ran}X$ inside, then distribution function of random variable $\alpha(X)$ is again continuous.

Theorem 2.2.6. *Let C_{XY} be the copula with continuous random variables X and Y . Further, if α and β are strictly increasing on $\text{Ran}X$ and $\text{Ran}Y$ respectively, then $C_{\alpha(X)\beta(Y)} = C_{XY}$. This means for the strictly increasing transformations of X and Y , C_{XY} does not change.*

Proof. Firstly, let F_1, G_1, F_2 and G_2 be the distribution functions of the random variables; $X, Y, \alpha(X)$ and $\beta(Y)$ respectively. Since α and β are strictly increasing we have,

$$F_2 = P(\alpha(X) \leq x) = P(X \leq \alpha^{-1}(x)) = F_1(\alpha^{-1}(x))$$

$$G_2 = P(\beta(Y) \leq y) = P(Y \leq \beta^{-1}(y)) = G_1(\beta^{-1}(y)) \quad .$$

$$\forall x, y \in \overline{\mathbf{R}}$$

$$\begin{aligned}
C_{\alpha(X)\beta(Y)}(F_2(x), G_2(y)) &= P(\alpha(X) \leq x, \beta(Y) \leq y) \\
&= P(X \leq \alpha^{-1}(x), Y \leq \beta^{-1}(y)) \\
&= C_{XY}(F_1(\alpha^{-1}(x)), G_1(\beta^{-1}(y))) \\
&= C_{XY}(F_2(x), G_2(y))
\end{aligned}$$

Continuity of X and Y leads $\text{Ran}F_2 = \text{Ran}G_2 = I$, thus we have $C_{\alpha(X)\beta(Y)} = C_{XY}$ on I^2 . $\square$

2.2.3 Survival Copula

A population representation of the lifetime as random variables is a common study area in many applications. Basically, a survival function searches the probability of living or surviving for objects through a given time x . The function is also known as the reliability function. It is denoted as $\bar{F}(x) = P(X > x) = 1 - F(x)$, here $F(x)$ is again the distribution function with random variable X . For the pair of random variables (X, Y) with the joint distribution function H , the joint survival function is written as $\bar{H}(x, y) = P(X > x, Y > y)$. Furthermore, marginals of the joint survival function $\bar{H}$ are calculated by $\bar{H}(x, -\infty) = \bar{F}(x)$ and $\bar{H}(-\infty, y) = \bar{G}(y)$ where the distributions $\bar{F}$ and $\bar{G}$ are the univariate survival functions.

After this point, one can think of the idea that if there is a connection between the joint survival functions and their univariate marginal survival functions in the same manner in Sklar's theorem. One can observe the link between the survival functions and copulas easily. Let C be a copula with the random variables X and Y , then

$$\begin{aligned}
\bar{H}(x, y) &= P(X > x, Y > y) = 1 - F(x) - G(y) + H(x, y) \\
&= \bar{F}(x) + \bar{G}(y) - 1 + C(F(x), G(y)) \\
&= \bar{F}(x) + \bar{G}(y) - 1 + C(1 - \bar{F}(x), 1 - \bar{G}(y))
\end{aligned}$$

Now let us define function $\hat{C}$ from I^2 to I such that

$$\hat{C}(u, v) = u + v - 1 + C(1 - u, 1 - v) \quad (2.10)$$

This yields

$$\overline{H}(x, y) = \hat{C}(\overline{F}(x), \overline{G}(y)) \quad (2.11)$$

$\hat{C}$ is the *survival copula* of X and Y . Beware of confusing the survival copula $\hat{C}$ with the joint survival function $\overline{C}$ for two uniformly distributed random variables with the joint distribution function copula C . Hence, $\overline{C}$ is given by $\overline{C} = P(U > u, V > v) = 1 - u - v + C(u, v) = \hat{C}(1 - u, 1 - v)$.

Example 2.2.3. *A bivariate Pareto distribution with random variables X, Y whose joint survival function is*

$$\overline{H}_\theta(x, y) = \begin{cases} (1 + x + y)^{-\theta}, & x \geq 0, y \geq 0, \\ (1 + x)^{-\theta}, & x \geq 0, y < 0, \\ (1 + y)^{-\theta}, & x < 0, y \geq 0, \\ 1, & x < 0, y < 0 \end{cases}$$

where the parameter is defined as $\theta > 0$. Therefore we can deduce the marginal survival functions as

$$\overline{F}_\theta(x) = \begin{cases} (1 + x)^{-\theta}, & x \geq 0, \\ 1, & x < 0 \end{cases} \quad \overline{G}_\theta(y) = \begin{cases} (1 + y)^{-\theta}, & y \geq 0, \\ 1, & y < 0 \end{cases}$$

We obtain the marginal survivals with the similar structure. In order to construct the survival copula applying the survival version of Corollary 2.2.1 is very straightforward. Firstly by inverting the marginal survival functions, we get $\overline{F}_\theta^{-1}(u) = u^{\frac{1}{\theta}} - 1$ and $\overline{G}_\theta^{-1}(v) = v^{\frac{1}{\theta}} - 1$. After putting them into the joint survival function and the survival

copula yields

$$\hat{C}_\theta(u, v) = (u^{\frac{1}{\theta}} + v^{\frac{1}{\theta}} - 1)^{-\theta}$$

Example 2.2.4. Recall the copula, Gumbel's bivariate distribution in Example 2.2.1 which is

$$C_\theta = u + v - 1 + (1 - u)(1 - v)e^{-\theta \ln(1-u)\ln(1-v)},$$

By simply using the result (2.10) we have the survival copula as

$$\begin{aligned}\hat{C}_\theta(u, v) &= u + v - 1 + 1 - u + 1 - v - 1 + uve^{-\theta \ln u \ln v} \\ &= uve^{-\theta \ln u \ln v}.\end{aligned}$$

CHAPTER THREE

FAMILIES OF COPULA

In literature, there exist lots of copula types and variety of methods to construct them. In this study, we focus on the two famous families namely, *Archimedean* and *elliptical*. For each section first we start by with information about the families they belong and then examine them.

3.1 Archimedean Copula

This prominent family of copulas and their applications are first presented by Genest & MacKay (1986), Genest & Mackay (1986). Due to their potential in applications, they are often preferred. For example, they are used for the most of the actuarial studies. Mainly the reasons for their preference are; they can be constructed rather easily, many of the copula families are included in this class, and they have many useful properties.

First, we begin with the general ideas for this family of copulas. Recall that $\forall x, y \in \overline{\mathbf{R}}$, the continuous random variables X and Y are independent if $H(x, y) = F(x)G(y)$. Describing the joint distribution by factoring it into product of two functions as this is the only case of doing it. However, there are some situations which a function of H can be factored into a product of functions F and G .

As a result for a function of λ , we can describe a joint distribution function as $\lambda(H(x, y)) = \lambda(F(x))\lambda(G(y))$. And surely function λ must take positive values in the interval $(0, 1)$. Furthermore by defining φ as; $\varphi(t) = -\ln\lambda(t)$, we can alter the structure of equality as $\varphi(H(x, y)) = \varphi(F(x)) + \varphi(G(y))$ i.e. rewritten as product of the marginals. When the copula is added to the equality, we get

$$\varphi(C(u, v)) = \varphi(u) + \varphi(v) \quad (3.1)$$

For example, recall the survival copula we found in Example 2.2.4, the function which

satisfies the form in (3.1) is $\varphi(t) = t^{-1/\theta} - 1$. Now we focus on the form, $C(u, v) = \varphi^{[-1]}(\varphi(u) + \varphi(v))$ to acknowledge how to describe copulas in terms of function φ .

Definition 3.1.1. Let we have a continuous and strictly decreasing function $\varphi : I \rightarrow [0, \infty)$, which satisfies the condition; $\varphi(1) = 0$. The *pseudo-inverse* of the function φ is, $\varphi^{[-1]} : [0, \infty) \rightarrow I$ given as

$$\varphi^{[-1]}(t) = \begin{cases} \varphi^{-1}(t), & 0 \leq t \leq \varphi(0), \\ 0, & \varphi(0) \leq t \leq \infty \end{cases} \quad (3.2)$$

Clearly, this definition provides the copula to lay in the interval I when it is written in the form of $\varphi^{-1}(t)$. Furthermore, the properties of $\varphi^{-1}(t)$ are

1. Continuous and nonincreasing on $[0, \infty)$,
2. Strictly decreasing on $[0, \varphi(1)]$,
3. $\varphi^{[-1]}(\varphi(u)) = u$ where $u \in I$,

With the third property we can immediately define

$$\varphi(\varphi^{[-1]}(t)) = \begin{cases} t, & 0 \leq t \leq \varphi(0), \\ \varphi(0), & \varphi(0) \leq t \leq \infty \end{cases}$$

4. Observe that if $\varphi(0) = \infty$, then $\varphi^{[-1]} = \varphi^{-1}$.

Lemma 3.1.1. Let the function φ has the features given in the Definition 3.1.1, and $C : I \rightarrow I^2$ with

$$C(u, v) = \varphi^{[-1]}(\varphi(u) + \varphi(v)) . \quad (3.3)$$

then C has the conditions defined in (2.2), (2.3) and (2.4).

Proof. By plugging the zeros to the function, we get $C(u, 0) = \varphi^{[-1]}(\varphi(u) + \varphi(0))$ and $C(0, v) = \varphi^{[-1]}(\varphi(0) + \varphi(1)) = 0$ by the definition of $\varphi^{[-1]}$. Also we have,

$$C(u, 1) = \varphi^{[-1]}(\varphi(u) + \varphi(1)) = \varphi^{[-1]}(\varphi(u)) = u \text{ and } C(1, v) = \varphi^{[-1]}(\varphi(1) + \varphi(v)) = \varphi^{[-1]}(\varphi(v)) = v. \quad \square$$

Lemma 3.1.2. *Let the functions of φ , $\varphi^{[-1]}$ and C have the features which are given in the Lemma 3.1.1, we say C is 2-increasing if and only if $\forall v \in I$ and $u_1 \leq u_2$*

$$C(u_2, v) - C(u_1, v) \leq u_2 - u_1 \quad (3.4)$$

Proof. In order to prove the lemma, first observe that 3.4 equals to $V_C = ([u_1, u_2] \times [v, 1]) \geq 0$, and if C is 2-increasing, form holds. Now for the other way around, assume that C satisfies the 3.4. Say we have v_1, v_2 from I where $v_1 \leq v_2$. Since $\varphi, \varphi^{[-1]}$ are continuous, C is continuous. As a result there exists a $t \in I$ so that $C(t, v_2) = v_1$, equivalently $\varphi(v_2) + \varphi(t) = \varphi(v_1)$. Then we have

$$\begin{aligned} C(u_2, v_1) - C(u_1, v_1) &= \varphi^{[-1]}(\varphi(u_2) + \varphi(v_1)) - \varphi^{[-1]}(\varphi(u_1) + \varphi(v_1)) \\ &= \varphi^{[-1]}(\varphi(u_2) + \varphi(v_2) + \varphi(t)) - \varphi^{[-1]}(\varphi(u_1) + \varphi(v_2) + \varphi(t)) \\ &= C(C(u_2, v_2), t) - C(C(u_1, v_2), t) \\ &\leq C(u_2, v_2) - C(u_1, v_2) \end{aligned}$$

Thus C is 2-increasing. $\square$

Theorem 3.1.1. *Let φ and $\varphi^{[-1]}$ be the functions defined as in the Definition 3.1.1, the function $C : I \rightarrow I^2$ which has the form (3.3) is a copula if and only if φ is convex.*

Proof. See Alsina et al. (2003). $\square$

Thus we covered the features of the function φ that it should have to determine a copula. In that case, the copula classes which can be written in the form of (3.3) are called *Archimedean copulas*. Moreover, the function φ which is used to construct a particular copula is called as the *generator* of that copula. Whenever $\varphi(0) = \infty$ holds, the function φ is called as the *strict generator* of the copula. And a copula constructed by a strict generator is known as the *strict Archimedean copula*.

Example 3.1.1. Let we have a generator such that

$$\varphi(t) = -\ln t \text{ where } t \in [0, 1]$$

Observe, since $\varphi(0) = \infty$, the generator is strict. It means that $\varphi^{[-1]}(t) = \varphi^{-1}(t) = e^{-t}$. By using the generator according to (3.1.3), we get $C(u, v) = \exp(-((- \ln u) + (- \ln v))) = uv = \Pi(u, v)$. So the product copula is indeed a strict Archimedean copula.

Example 3.1.2. Now say we have the generator $\varphi_\theta(t) = \ln(1 - \theta \ln t)$ where θ lies in the interval $(0, 1]$. Again we can see that $\varphi_\theta(t)$ is strict because $\varphi_\theta(0) = \infty$. Then $\varphi^{-1}(t) = \exp((1 - e^t)/\theta)$. To generate a copula C_θ based on the φ_θ , we apply

$$\begin{aligned} C(u, v) &= \varphi^{[-1]}(\varphi(u) + \varphi(v)) \\ &= \exp\left(\frac{1 - \exp[\ln((1 - \theta \ln u)(1 - \theta \ln v))]}{\theta}\right) \\ &= \exp\left(\frac{1 - (1 - \theta \ln u)(1 - \theta \ln v)}{\theta}\right) \\ &= \exp(\ln u + \ln v - \theta \ln u \ln v) \\ &= uv \exp(-\theta \ln u \ln v) . \end{aligned}$$

In fact, by the given generator φ , we again build the copula we have found out in Example 2.2.4, which is the survival copula of the Gumbel's bivariate distribution.

3.1.1 Gumbel Copula

This type of copula was first studied by Gumbel (1960), such families are strict and have the form

$$C_\theta(u, v) = \exp\left(-[(- \ln u)^\theta + (- \ln v)^\theta]^{1/\theta}\right) \quad (3.5)$$

where $\theta \in [1, \infty)$ and $\varphi_\theta = (- \ln t)^\theta$.

3.1.2 Frank Copula

The first appearance of this family did not involve a statistical context. Frank (1979) found this type of function as a solution for a functional equation. The family is studied as the notion of copula by Genest (1987). Such families are strict and have the form

$$C_\theta(u, v) = -\frac{1}{\theta} \ln \left(1 + \frac{(e^{-\theta u} - 1)(e^{-\theta v} - 1)}{e^{-\theta} - 1} \right) \quad (3.6)$$

$$\text{where } \theta \in (-\infty, \infty) - \{0\} \text{ and } \varphi_\theta = -\ln \frac{e^{-\theta t} - 1}{e^{-\theta} - 1} .$$

3.1.3 Clayton Copula

This family of copulas has first appeared in the work of Clayton (1978). Such families are strict for $\theta \geq 0$ and have the form.

$$C_\theta(u, v) = \left(\max(u^{-\theta} + v^{-\theta} - 1, 0) \right)^{-1/\theta} \quad (3.7)$$

$$\text{where } \theta \in [-1, \infty) - \{0\} \text{ and } \varphi_\theta = \frac{1}{\theta} (t^{-\theta} - 1) .$$

3.2 Elliptical Copula

Definition 3.2.1. Let X be an n -dimensional random vector, $\mu \in \mathbb{R}^n$ and let $\Sigma \in \mathbb{R}^{n \times n}$ be a symmetric positive semi-definite matrix. We say that X has an elliptical distribution if the characteristic function $\varphi_{X-\mu}$ of $X - \mu$ can be defined as $\varphi_{X-\mu}(t) = \Psi(t^T \Sigma t)$ where $\psi : \mathbb{R}^{+*} \rightarrow \mathbb{R}$. Elliptical distribution is denoted as $X \sim (\mu, \sigma, \psi)$ and ψ is called characteristic generator.

Theorem 3.2.1. $X \sim (\mu, \Sigma, \psi)$ with $\text{rank}(\Sigma) = k$ if and only if there exists a random variable $Y \geq 0$ which is independent of U , uniformly distributed k -dimensional random vector on unit hyper-sphere $\{z \in \mathbb{R}^k | z^T z = 1\}$, and an $n \times k$ matrix A where

$AA^t = \Sigma$ so that

$$X \stackrel{d}{=} \mu + YAU.$$

Proof. See Fang et al. (1990). □

In the light of Definition 3.2 and Theorem 3.2, the normal (Gaussian) and t copula belong to the elliptical family. For more details, address McNeil et al. (2003).

3.2.1 Normal Copula

The normal copula gives the information for the dependency structure generated by the bivariate normal distribution. By applying the transformation (2.7) to the random variables X, Y which are distributed over standard normal distribution; $N(0, 1)$, we have the normal copula via covariance θ as

$$C_\theta(u, v) = H_\theta(\Phi^{-1}(u), \Phi^{-1}(v)) \quad (3.8)$$

where Φ is univariate standard normal,

H_θ bivariate standard normal distribution and $\theta \in (-1, 1)$

Thus we have

$$C_\theta(u, v) = \frac{1}{2\pi\sqrt{1-\theta^2}} \int_{-\infty}^{\Phi^{-1}(u)} \int_{-\infty}^{\Phi^{-1}(v)} \exp\left(\frac{-(s^2 - 2\theta st + t^2)}{2(1-\theta^2)}\right) ds dt \quad (3.9)$$

also, the density of normal copula is given by

$$c_\theta(u, v) = \frac{1}{2\pi\sqrt{1-\theta^2}} \exp\left(\frac{-(s^2 - 2\theta st + t^2)}{2(1-\theta^2)}\right) \frac{d}{du} \Phi^{-1}(u) \frac{d}{dv} \Phi^{-1}(v) \quad (3.10)$$

3.2.2 Student's t copula

With the same procedure as the normal copula, we can adopt the t copula from bivariate Student's t distribution. The univariate t distribution is given by

$$t_v(x) = \int_{-\infty}^x \frac{\Gamma(\frac{v+1}{2})}{\sqrt{\pi v} \Gamma(\frac{v}{2})} \left(1 + \frac{s^2}{v}\right)^{-\frac{v+1}{2}} ds \quad (3.11)$$

where v is degree of freedom, and Γ is Euler function. Moreover, the bivariate t distribution $t_{p,v}$ via Pearson's correlation coefficient p , is given as

$$t_{p,v}(\mathbf{x}) = \int_{-\infty}^{x_1} \int_{-\infty}^{x_2} \frac{1}{2\pi\sqrt{1-p^2}} \left(1 + \frac{s^2 + t^2 - 2pst}{v(1-p^2)}\right)^{-\frac{v+2}{2}} ds dt, \quad (3.12)$$

then the bivariate Student's t copula is governed by the bivariate function

$$C_{p,v}(u, v) = t_{p,v}(t_v^{-1}(u), t_v^{-1}(v)) \quad (3.13)$$

$$= \int_{-\infty}^{t_v^{-1}(u)} \int_{-\infty}^{t_v^{-1}(v)} \frac{1}{2\pi\sqrt{1-p^2}} \left(1 + \frac{s^2 + t^2 - 2pst}{v(1-p^2)}\right)^{-\frac{v+2}{2}} ds dt \quad (3.14)$$

The density of Student's copula is given by

$$c_{p,v}(u, v) = p^{-\frac{1}{2}} \frac{\Gamma(\frac{v+2}{2})\Gamma(\frac{v}{2})}{\left(\Gamma(\frac{v+1}{2})\right)^2} \frac{\left(1 + \frac{\zeta_1^2 + \zeta_2^2 - 2p\zeta_1\zeta_2}{v(1-p^2)}\right)^{-\frac{v+2}{2}}}{\prod_{i=1}^2 \left(1 + \frac{\zeta_i^2}{v}\right)^{-\frac{v+2}{2}}}, \quad (3.15)$$

where $\zeta_1 = t_v^{-1}(u)$, $\zeta_2 = t_v^{-1}(v)$.

3.3 Visualization Of Copula Types

This section is devoted to observe the shapes of the mentioned copulas. Main observations are based on the different joint behaviours and tail characteristics of their densities. Parameters of copulas are determined by the two Kendall's tau measures, 0.3 and 0.7 (address Section 4 and Section 5). The used visualizations styles are 3D graphs and contour plots.

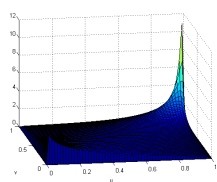


Figure 3.1 Gumbel density via $\rho = 0.3$

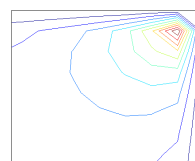


Figure 3.2 Gumbel level curves of density via $\rho = 0.3$

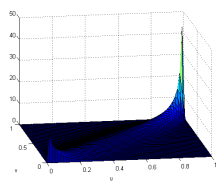


Figure 3.3 Gumbel density via $\rho = 0.7$

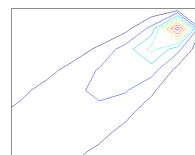


Figure 3.4 Gumbel level curves of density via $\rho = 0.7$

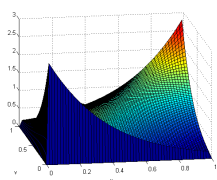


Figure 3.5 Frank density via $\rho = 0.3$

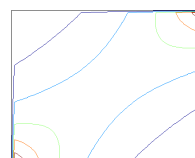


Figure 3.6 Frank level curves of density via $\rho = 0.3$

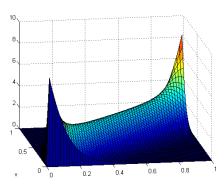


Figure 3.7 Frank density via $\rho = 0.7$

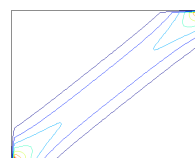


Figure 3.8 Frank level curves of density via $\rho = 0.7$

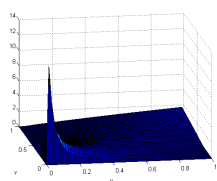


Figure 3.9 Clayton density via $\rho = 0.3$

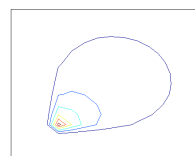


Figure 3.10 Clayton level curves of density via $\rho = 0.3$

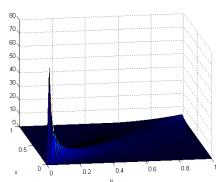


Figure 3.11 Clayton density via $\rho = 0.7$

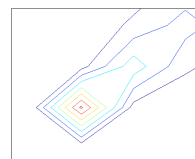


Figure 3.12 Clayton level curves of density via $\rho = 0.7$

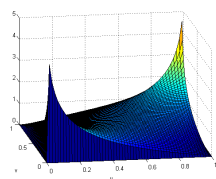


Figure 3.13 Normal density via $\rho = 0.3$

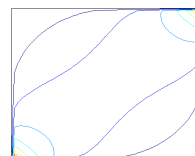


Figure 3.14 Normal level curves of density via $\rho = 0.3$

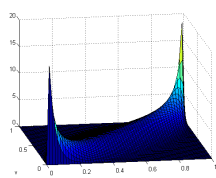


Figure 3.15 Normal density via $\rho = 0.7$

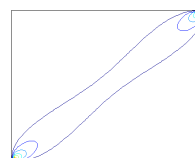


Figure 3.16 Normal level curves of density via $\rho = 0.7$

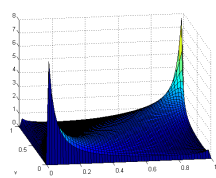


Figure 3.17 Student's t density via $\rho = 0.3$ and d.o.f.=4

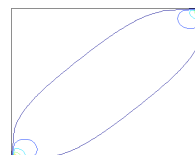


Figure 3.18 Student's t level curves of density via $\rho = 0.3$ and d.o.f.=4

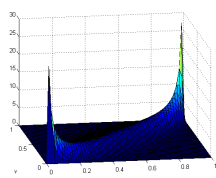


Figure 3.19 Student's t density via $\rho = 0.7$ and d.o.f.=4

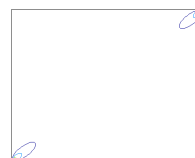


Figure 3.20 Student's t level curves of density via $\rho = 0.7$ and d.o.f.=4

Remark 3.3.1. Gumbel copula shows upper tail tendency, Frank behaves symmetrical on the tails and Clayton captures a lower tail behaviour. These Archimedean copulas get relatively thinner with the high dependence level.

If the degrees of freedom tend to be very large, the Student's t copula diverges to the normal copula. But for the limited range of values, these two elliptical family of copulas show very different characteristics. One can easily notice, Student's t copula captures more information about the tails than the normal copula.



CHAPTER FOUR

DEPENDENCE

4.1 Dependence Structures

Investigation of the dependence structure between random variables is one of the most endeavored areas in the history of probability and statistics. Random variables to be searched for their association could belong to many areas, such as DNA sequencing, prey-predator dynamics, calculating risks of financial portfolios, optimization, forecasting natural disasters, renewable energy and much more... Adapting the elements of a particular concept and study for the dependency structure between random variables is a very core idea of analyzing a dataset.

There exist very different formats to measure the dependency relation. Definitions of dependency measurements are too important, while some are appropriate for a studied subject, others might not get valid outcomes. One must approach the data carefully with the measurements and the assumptions. Many of the properties and measures for dependency are *scale-invariant* (the term used by Höffding (1940), Hoeffding (1941)) which means they do not change under the strictly increasing transformations of random variables. In Theorem 2.2.6, we validated that random variables of a copula remain invariant under such transformations. Thus, copula random variables can express the scale-invariant feature of properties and measures.

In this section, the role of copulas in the measure of association is investigated. First, the well known scale-invariant measure of associations is introduced. They are namely the population versions of *Kendall's tau* and *Spearman's Rho*. Both of the measures nourish from the idea of *concordance*. Before beginning with the form of concordance, note that 'correlation coefficient' is the term that belongs to the linear dependence between random variables so to use the measurement of dependency in a modern sense we use the term 'measure of association'.

4.2 Concordance Measures

At first glance, one can define the concordance as the agreement of the large or small tendency of individual random variables in a random pair. Or we call the pair of random variables as concordant when one of them tends to be large (small), other one tends to be large (small) too. Now let us redefine it mathematically, let (x_i, y_i) and (x_j, y_j) be two arbitrarily chosen observations from a vector (X, Y) which consists continuous random variables. Pairs of (x_i, y_i) and (x_j, y_j) are said to be concordant if $x_i < x_j$ and $y_i < y_j$, or if $x_i > x_j$ and $y_i > y_j$. In the same manner, pairs are said to be discordant if $x_i < x_j$ and $y_i > y_j$, or if $x_i > x_j$ and $y_i < y_j$. Alternatively, we can reformulate the concordance as; (x_i, y_i) and (x_j, y_j) are concordant if $(x_i - x_j)(y_i - y_j) > 0$ and discordant if $(x_i - x_j)(y_i - y_j) < 0$.

4.2.1 Kendall's tau

Kendall's tau is a widely known measure for association between random variables. The sample version of this measurement was defined by the concordance of random variables Kruskal (1958). Let we have $\{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\}$ as a random sample via number of n observations from a vector (X, Y) of continuous random variables. There exist number of $\binom{n}{2}$ distinct pairs of (x_i, y_i) and (x_j, y_j) that we can crosscheck whether pairs of random variables concordant or discordant. Let c be the number of concordant pairs and d be the number of discordant pairs then, for the sample, Kendall's tau is given as

$$t = \frac{c - d}{c + d} = \frac{c - d}{\binom{n}{2}} . \quad (4.1)$$

To see the Kendall's tau in another view, observe that it is the probability of concordant pair minus probability of discordant pair of (x_i, y_i) and (x_j, y_j) which are chosen randomly from the sample. The population format of the Kendall's tau for a vector (X, Y) which consists of continuous random variables with joint distribution function H has a similar fashion. Let we have two independent and identically

distributed random vectors (X_1, Y_1) and (X_2, Y_2) where each has the joint distribution function H . Then Kendall's tau for the population is again defined as the probability of concordance minus probability of discordance as,

$$\tau = \tau_{X,Y} = P((X_1 - X_2)(Y_1 - Y_2) > 0) - P((X_1 - X_2)(Y_1 - Y_2) < 0) . \quad (4.2)$$

Before showing how the copula involves in the matters of concordance and measure of association, let us define Q as the 'concordance function' which is the probabilities of concordance and discordance between the vectors (X_1, Y_1) and (X_2, Y_2) with the continuous random variables and different joint distributions H_1 and H_2 but with the same marginals of F and G . Then with the following theorem, we can deduce that concordance function only has connection with the distributions of (X_1, Y_1) and (X_2, Y_2) by their copulas. And by that theorem, one can see how practical to find measure of association with the help of copula.

Theorem 4.2.1. *Let (X_1, Y_1) and (X_2, Y_2) be independent vectors made up of continuous random variables with joint distributions H_1 and H_2 respectively. With same marginals namely, F via X_1 and X_2 and G via Y_1 and Y_2 . Then let C_1 and C_2 be the copulas for the vectors (X_1, Y_1) and (X_2, Y_2) respectively where $H_1(x, y) = C_1(F(x), G(y))$ and $H_2(x, y) = C_2(F(x), G(y))$. The concordance function Q is given by*

$$Q = P((X_1 - X_2)(Y_1 - Y_2) > 0) - P((X_1 - X_2)(Y_1 - Y_2) < 0) , \quad (4.3)$$

then we have the result

$$Q = Q(C_1, C_2) = 4 \iint_{I^2} C_2(u, v) dC_1(u, v) - 1 . \quad (4.4)$$

Proof. We need to rearrange the concordance function first. Due to the continuity of random variables we can write,

$$P((X_1 - X_2)(Y_1 - Y_2) < 0) = 1 - P((X_1 - X_2)(Y_1 - Y_2) > 0) ,$$

then we have

$$Q = 2P((X_1 - X_2)(Y_1 - Y_2) > 0) - 1 . \quad (4.5)$$

Observe the following revision,

$$P((X_1 - X_2)(Y_1 - Y_2) > 0) = P((X_1 > X_2)(Y_1 > Y_2)) + P((X_1 < X_2)(Y_1 < Y_2)) \quad (4.6)$$

Since both products can be positive or negative, we can calculate the probabilities by integrating one of the distributions of (X_1, Y_1) and (X_2, Y_2) . Let us take the (X_1, Y_1) .

$$\begin{aligned} P(X_1 > X_2, Y_1 > Y_2) &= P(X_2 < X_1, Y_2 < Y_1) \\ &= \iint_{R^2} P(X_2 < x, Y_2 < y) dC_1(F(x), G(y)) \\ &= \iint_{R^2} C_2(F(x), G(y)) dC_1(F(x), G(y)) \end{aligned}$$

By applying the probability transforms that is $U = F(x)$ and $V = G(y)$

$$P(X_1 > X_2, Y_1 > Y_2) = \iint_{I^2} C_2(u, v) dC_1(u, v)$$

We deal with the another part in the same manner,

$$\begin{aligned} P(X_1 < X_2, Y_1 < Y_2) &= \iint_{R^2} P(X_2 > x, Y_2 > y) dC_1(F(x), G(y)) \\ &= \iint_{R^2} (1 - F(x) - G(y) + C_2(F(x), G(y))) dC_1(F(x), G(y)) \\ &= \iint_{I^2} (1 - u - v + C_2(u, v)) dC_1(u, v) \end{aligned}$$

The copula distribution C_1 has the pair (U, V) whose random variables distribute

uniformly. Thus $E(U) = E(V) = 1/2$. By using this fact, we can get

$$\begin{aligned}
P(X_1 < X_2, Y_1 < Y_2) &= \iint_{I^2} dC_1(u, v) - \iint_{I^2} u dC_1(u, v) - \iint_{I^2} v dC_1(u, v) + \\
&\quad + \iint_{I^2} C_2(u, v) dC_1(u, v) \\
&= 1 - \frac{1}{2} - \frac{1}{2} + \iint_{I^2} C_2(u, v) dC_1(u, v) \\
&= \iint_{I^2} C_2(u, v) dC_1(u, v)
\end{aligned}$$

Therefore

$$P((X_1 - X_2)(Y_1 - Y_2) > 0) = 2 \iint_{I^2} C_2(u, v) dC_1(u, v)$$

Finally by plugging into (4.5)

$$Q = 4 \iint_{I^2} C_2(u, v) dC_1(u, v) - 1$$

□

Corollary 4.2.1. *Let we have C_1 and C_2 as copulas and Q as the concordance function. Then, $Q(C_1, C_2) = Q(C_2, C_1)$ i.e. concordance function is symmetric.*

Proof. By Theorem 4.2.1 we know that $Q(C_1, C_2) = 4 \iint_{I^2} C_2(u, v) dC_1(u, v) - 1$. Recall the part (4.2.6), if one takes the integral according to (X_2, Y_2) for the following procedure the result follow as

$$P((X_1 - X_2)(Y_1 - Y_2) > 0) = 2 \iint_{I^2} C_1(u, v) dC_2(u, v) .$$

Again by plugging into (4.5)

$$Q = 4 \iint_{I^2} C_1(u, v) dC_2(u, v) - 1 = 4 \iint_{I^2} C_2(u, v) dC_1(u, v) - 1 .$$

Hence $Q(C_2, C_1) = Q(C_1, C_2)$.

□

Theorem 4.2.2. *Let X and Y be continuous random variables with copula C . Kendall's*

τ is given by

$$\tau_{X,Y} = \tau_C = Q(C, C) = 4 \iint_{I^2} C(u, v) dC(u, v) - 1 . \quad (4.7)$$

Observe that the integral can be seen as the expected value of the function $C(U, V)$ of uniform random variables U and V which lie on $(0, 1)$ and has the distribution function C as

$$\tau_C = 4E(C(U, V)) - 1 . \quad (4.8)$$

Also note that, when the copula C belongs to a parametric family of copulas, which is denoted such as C_θ or $C_{\alpha,\beta}$, we can easily denote the Kendall's tau as τ_θ and $\tau_{\alpha,\beta}$ respectively.

Remark 4.2.1. Since the concordance function Q is the difference of probabilities, $Q(C, C) \in [-1, 1]$.

Theorem 4.2.3. Let C_1 and C_2 be copulas, then the following equality holds

$$\iint_{I^2} C_1(u, v) dC_2(u, v) = \frac{1}{2} - \iint_{I^2} \frac{\partial C_1(u, v)}{\partial u} \frac{\partial C_2(u, v)}{\partial v} du dv$$

Proof. Absolute continuity of copulas leads us to use the integration by parts

$$\iint_{I^2} C_1(u, v) dC_2(u, v) = \int_0^1 \int_0^1 C_1(u, v) \frac{\partial C_2(u, v)}{\partial u \partial v} du dv \quad (4.9)$$

Let $s = C_1(u, v)$ then $ds = \frac{\partial C_1(u, v)}{\partial u} du$ and $dt = \frac{\partial C_2(u, v)}{\partial u \partial v} du$ then $t = \frac{\partial C_2(u, v)}{\partial v}$.

By evaluating the inner part of the integral we have

$$\begin{aligned} & \int_0^1 C_1(u, v) \frac{\partial C_2(u, v)}{\partial u \partial v} du \\ &= C_1(u, v) \frac{\partial C_2(u, v)}{\partial v} \Big|_0^1 - \int_0^1 \frac{\partial C_1(u, v)}{\partial u} \frac{\partial C_2(u, v)}{\partial v} du \\ &= v - \int_0^1 \frac{\partial C_1(u, v)}{\partial u} \frac{\partial C_2(u, v)}{\partial v} du \end{aligned}$$

By plugging the result, we have the inner part and by taking the definite integral of v

from 0 to 1, the desired equality follows as

$$\iint_{I^2} C_1(u, v) dC_2(u, v) = \frac{1}{2} - \iint_{I^2} \frac{\partial C_1(u, v)}{\partial u} \frac{\partial C_2(u, v)}{\partial v} dudv$$

Additionally, essence of this proof is approximating C_1 and C_2 by sequences of absolutely continuous copulas. For details, see Li et al. (2002). $\square$

Determining the Kendall's tau for the member of Archimedean family of copulas is even easier. Tau can be calculated directly by the generator of a copula. This is one of the advantages that Archimedean family provides we mentioned in the Section 3. Because clearly dealing with the one variable function -the generator- is simpler than dealing with bivariate function -the copula-. Following corollary gives the evaluation of Kendall's tau for the Archimedean families.

Corollary 4.2.2. *Let X, Y be random variables with Archimedean copula C generated by φ . The Kendall's tau, τ_C for X and Y is evaluated as*

$$\tau_C = 1 + 4 \int_0^1 \frac{\varphi(t)}{\varphi'(t)} dt . \quad (4.10)$$

Proof. See, (Nelsen 2006). $\square$

Example 4.2.1. *Let we have Gumbel copula, for $\theta \geq 0$ by the Corollary 4.2, the Kendall's tau can be evaluated as*

$$\frac{\varphi_\theta(t)}{\varphi'_\theta(t)} = \frac{(-\ln t)^\theta}{\theta(-\ln t)^{\theta-1} \frac{1}{t}} = \frac{t \ln t}{\theta}$$

$$\begin{aligned} \tau_\theta &= 1 + 4 \int_0^1 \frac{t \ln t}{\theta} dt \\ &= \frac{\theta - 1}{\theta} . \end{aligned}$$

4.3 Properties of Dependence

On dependence studies, the most emphasized property is obviously the independence of random variables. Basically, two continuous random variables are independent if their joint distribution can be described as the product of its marginal distributions. If we have an independence condition between X and Y , it means that in the set of all joint distribution H has a particular place of its own as a subset. This subset can be described by the independence copula Π . One can describe many of the dependence properties by the help of copula, same as describing independence with Π . In this section, we investigate the properties of copulas that give the structure for another type of dependence.

First we examine very basic dependence structures namely positive and negative dependence. One can think of it as the positive tendency or negative tendency that is similar the notion of concordance. Strictly speaking, properties of positive dependence imply, random variables have the tendency of getting large or small values together, and negative dependence imply the one variable tends to get large while the other tends to get small.

4.3.1 Quadrant Dependence

Definition 4.3.1. Let X and Y be random variables, as Lehmann (1966) proposed X and Y said to be *positively quadrant dependent* (PQD) if $\forall (x, y) \in \mathbf{R}^2$

$$P(X \leq x, Y \leq y) \geq P(X \leq x)P(Y \leq y) \quad (4.11)$$

or

$$P(X > x, Y > y) \geq P(X > x)P(Y > y) \quad (4.12)$$

In the case of 4.11 or 4.12, we denote the positive quadrant dependence as $PQD(X, Y)$. Further, by reversing the orders we get the conditions for *negative quadrant dependence*

which is denoted as $NQD(X, Y)$ in the same manner.

When X and Y are continuous random variables with joint distribution H which has copula C and marginal distributions F and G , PQD is determined by

$$H(x, y) \geq F(x)G(y) \quad \forall (x, y) \in \mathbf{R}^2 \quad (4.13)$$

Thus

$$C(u, v) \geq uv \quad \forall (u, v) \in I^2 \quad (4.14)$$

So if the continuous random variables are PQD, their joint distribution H and their copula C are also PQD.

4.3.2 Tail Monotonicity

Positive dependence defined in 4.3.1 can be rewritten as

$$P(Y \leq y | X \leq x) \geq P(Y \leq y)$$

or

$$P(Y \leq y | X \leq x) \geq P(Y \leq y | X \leq \infty)$$

To understand this expression, let X and Y be the lifetimes of two components, then the probability above stands for the probability of Y having a lesser lifetime increases, given that X has a lesser lifetime. Such characteristic of right tails are formalized by the definition below by Esary et al. (1972).

Definition 4.3.2. For the random variables X and Y , we have the following properties;

1. If $P(Y \leq y | X \leq x)$ is a nonincreasing function belong to $x \forall y$, then Y is said to be *left tail decreasing* in X , and denoted as $LTD(Y|X)$.
2. If $P(X \leq x | Y \leq y)$ is a nonincreasing function belong to $y \forall x$, then X is said to be *left tail decreasing* in Y , and denoted as $LTD(X|Y)$.

3. If $P(Y > y|X > x)$ is a nondecreasing function belong to $x \forall y$, then Y is said to be *right tail increasing* in X , and denoted as $RTI(Y|X)$.

4. If $P(X > x|Y > y)$ is a nondecreasing function belong to $y \forall x$, then X is said to be *right tail increasing* in Y , and denoted as $RTI(X|Y)$.

Next definition has the adaptation of tail monotonicity to copula.

Definition 4.3.3. For the continuous random variables X and Y with copula C , we have the following properties;

1. $\exists v \in I$ such that $\frac{C(u,v)}{u}$ is nonincreasing in u if and only if $LTD(Y|X)$.
2. $\exists u \in I$ such that $\frac{C(u,v)}{v}$ is nonincreasing in v if and only if $LTD(X|Y)$.
3. $\exists v \in I$ such that $\frac{1-u-v+C(u,v)}{1-u}$ is nondecreasing in u if and only if $RTI(Y|X)$.
4. $\exists u \in I$ such that $\frac{1-u-v+C(u,v)}{1-v}$ is nondecreasing in v if and only if $RTI(X|Y)$.

Thus for the continuous random variables, tail monotonicity is observed as a property of the copula. Moreover, proof is clear since marginals are nondecreasing.

4.3.3 Tail Dependence

Tail dependence is just an another approach used to observe comovement of random variables, involving largeness or smallness. But in this concept, we examine the relationship among the extreme observations whose dynamics might be critical in sensitive researches, for example in finance. For a model which has heavy tails, (for example, time series) without tail analysis it is exposed to the unexpected outcomes in the future. So inference about the tails holds a great importance.

Definition 4.3.4. Let random variables (X, Y) have joint distribution H with marginals

F_X and G_Y , then we have

$$\lambda_U = \lim_{x \rightarrow 1^-} (G(y) > t | F(x) > t) \quad (4.15)$$

as the *upper tail dependence coefficient* (upper TDC). If $\lambda_U > 0$, the pair (X, Y) is said to be upper tail dependent, if $\lambda_U = 0$ then it is upper tail independent. Also, similarly, we have

$$\lambda_L = \lim_{x \rightarrow 0^+} (G(y) \leq t | F(x) \leq t) \quad (4.16)$$

as the lower tail dependence coefficient (lower TDC) in the same manner. We can interpret those dependencies as the probability of one marginal overlapping high (low) threshold given that the other one overlaps high (low) threshold. Copula concept can be adapted to the tail dependence Sklar (1959). For a copula C of the joint distribution H . TDCs can be computed as

$$\lambda_U = \lim_{x \rightarrow 1^-} \frac{1 - 2t + C(t, t)}{1 - t} \quad \lambda_L = \lim_{x \rightarrow 0^+} \frac{C(t, t)}{t} \quad (4.17)$$

Note that copula's invariance property still holds for TDC, thus strictly increasing transformations of marginals does not change the TDC.

Example 4.3.1. Consider Gumbel copula in 3.5. We can determine its TDCs by following computations.

$$\begin{aligned} C(t, t) &= \exp(-[2^{1/\theta}(-\ln t)]) \\ &= t^{2^{1/\theta}} \end{aligned}$$

$$\begin{aligned} \lambda_U &= \lim_{x \rightarrow 1^-} \frac{1 - 2t + t^{2^{1/\theta}}}{1 - t} \\ &= \lim_{x \rightarrow 1^-} \frac{-2 + 2^{1/\theta} t^{2^{1/\theta}-1}}{-1} \\ &= 2 - 2^{1/\theta} \end{aligned}$$

and

$$\begin{aligned}\lambda_L &= \lim_{x \rightarrow 0^+} \frac{t^{2^{1/\theta}}}{t} \\ &= \lim_{x \rightarrow 0^+} \frac{2^{1/\theta} t^{2^{1/\theta}-1}}{1} \\ &= 0\end{aligned}$$

Example 4.3.2. Now consider Normal and t copula. Elliptical copulas have equal TDCs i.e. $\lambda_U = \lambda_L$. For them, TDC is denoted by λ . It has been shown that for Normal copula $\lambda = 0$ and for t copula, we have (Embrechts 2003)

$$\lambda = 2 - 2t_{v+1} \left(\sqrt{v+1} \frac{\sqrt{1-p}}{\sqrt{1+p}} \right) \quad (4.18)$$

Remark 4.3.1. From the examples of TDC, one can deduce that Gumbel copula is a right tailed distribution and t copula has symmetric tails. Those tail behaviours and other mentioned copula's tail behaviours can easily be seen in the Visualization of Copula Types Section.

4.3.3.1 Non-parametric Approach to Tail Dependence

Ferreira (2013) states that usage of parametric estimators contains risk for a model, because they may conduce to misinterpretations on the dependence structure. By using non-parametric methods, one can avoid this situation, but generally the cost is dealing with larger variance. Moreover non-parametric methods are very useful alternatives if there exist doubts for the structure of the studied data. First, we introduce another description of TDC.

$$\lambda = \lim_{t \rightarrow \infty} P(F(X) > 1 - 1/t \mid G(Y) > 1 - 1/t),$$

With the same fashion, we can adapt copula into this description by

$$\lambda = 2 - \lim_{t \rightarrow \infty} t P(F(X) > 1 - 1/t \text{ or } G(Y) > 1 - 1/t) \quad (4.19)$$

$$= 2 - \lim_{t \rightarrow \infty} t(1 - C(1 - 1/t, 1 - 1/t)). \quad (4.20)$$

Changing the (4.20) by empirical parts, Huang (1992) stated the estimator

$$\hat{\lambda}^{(H)} = 2 - \frac{1}{k_n} \sum_{i=1}^n 1\left(\hat{F}(X_i) > 1 - \frac{k_n}{n} \text{ or } \hat{G}(Y_i) > 1 - \frac{k_n}{n}\right) \quad (4.21)$$

where $\hat{F}$ and $\hat{G}$ are empirical distributions,

$$\hat{F}(x) = \frac{1}{n+1} \sum_{i=1}^n 1(X_i \leq x)$$

$$\hat{G}(y) = \frac{1}{n+1} \sum_{i=1}^n 1(Y_i \leq y)$$

Here $k_n = k$ is a parameter which satisfies $k \rightarrow \infty$ and $\frac{k}{n} \rightarrow 0$, as $n \rightarrow \infty$, where $k \in \{1, \dots, n\}$. Assigning of a k value that satisfies an optimum result is a hard endeavor. Because small values create big variances, on the other hand large values cause powerful bias.

Similarly for the LTD the nonparametric measure of Huang can be used as

$$\hat{\lambda}^{(H)} = \frac{1}{k_n} \sum_{i=1}^n 1\left(\hat{F}(X_i) \leq \frac{k_n}{n}, \hat{G}(Y_i) \leq \frac{k_n}{n}\right) \quad (4.22)$$

CHAPTER FIVE

METHODS OF PARAMETER ESTIMATION AND GOODNESS OF FIT TESTING

5.1 Parameter Estimation

We emphasize that copula provides a simple expression to tackling multivariate analysis. Classic statistical methods are not applicable for most of the multivariate inferences. Solely, the theory of maximum likelihood estimator (MLE) can be used. Other than there are alternative approaches to deal with high scale computations to calculate MLE's. Such techniques consist of nonparametric inference and simulation methods.

In this section, we investigate the usage of the statistical inference methods for theory of copula. The first method is the *exact maximum likelihood*. This process requires numerical approximations for the objective function of the estimation. Because of the hard computational scale of multivariate nature, computations have mixed derivative structure. In addition, we examine more compact forms to estimate parametrically.

Also if one has the sufficient information for a particular data, nonparametric estimations can be used with high performance while they provide simplicity for computations. Although we focus on the bivariate case note that especially for the inferences involve high dimensions, efforts to find the best combination of marginal and copula are challenging, and one can lose track easily. Non-parametric methods for marginals and copula are usable for this kind of situations. In general, such methods have the advantage of expressing the copula by the dataset itself, so we touch on the non-parametric ways of estimation too.

Then we continue the chapter with the goodness of fit testing which is a very rooted phase in the matter of statistical modeling. In this part the importance of testing a data fit is emphasized and some methods of goodness of fit testing is explained.

5.2 Parametric Ways

5.2.1 Exact Maximum Likelihood Estimation

First step recall form below as a representation for the derivative of a joint distribution function

$$f(x, y) = c(F(x), G(y))f(x)f(y) \quad , \quad (5.1)$$

where

$$c(F(x), G(y)) = \frac{\partial C(F(x), G(y))}{\partial F(x) \partial G(y)} \quad .$$

Observe that here f and g are the densities of the univariate marginal distributions and c is the density of the copula C . Later on, let $\{x_i, y_i\}_{i=1}^n$ be the sample data matrix then the log-likelihood function is defined as

$$l(\boldsymbol{\theta}) = \sum_{i=1}^n \ln c(F(x_i), G(y_i)) + \sum_{i=1}^n [\ln f(x_i) + \ln g(y_i)] \quad . \quad (5.2)$$

Here $\boldsymbol{\theta}$ is the parameter set of marginals and the copula. Thus by the maximization of the log-likelihood function, we adopt the estimations for marginals and the copula as

$$\hat{\boldsymbol{\theta}}_{MLE} = \max_{\boldsymbol{\theta} \in \Theta} l(\boldsymbol{\theta}) \quad (5.3)$$

Remark 5.2.1. Note that for the theory of asymptotic maximum likelihood the regularity conditions are assumed by Serfling (2009). The conditions hold for both copula and its marginals. These mentioned regularity conditions are that; maximum likelihood estimator exists, and it is consistent and asymptotically efficient. Additionally, asymptotically normal property holds

$$\sqrt{n}(\hat{\boldsymbol{\theta}}_{MLE} - \boldsymbol{\theta}_0) \rightarrow N(0, \mathcal{I}^{-1}(\boldsymbol{\theta}_0)) \quad (5.4)$$

where $\mathcal{I}$ denotes the Fisher's information matrix and $\boldsymbol{\theta}_0$ is the exact value of estimation.

5.2.2 IFM Method

This method is the slightly modified version of the exact maximum likelihood estimation. The preceding way of estimation requires parameter estimation of marginal distributions and parameters of the copula which represent the dependence structure. Naturally, such estimation process can get complicated especially if high dimensions are involved. Basically, the log-likelihood function is made up of both density of copula and its parameters and, marginals and all of their parameters. This observation is the starting point for the method *inference for the margins (IFM)* and it is proposed by Joe & Xu (1996) to have a computational ease. IFM process has two steps for the parameter estimation.

1. Firstly, marginal parameters denoted by θ_1 are estimated via estimation of univariate marginal distributions

$$\hat{\theta}_1 = \text{ArgMax}_{\theta_1} \sum_{i=1}^n [\ln f(x_i; \theta_1) + \ln g(y_i; \theta_1)] \quad (5.5)$$

2. Secondly, with the estimated $\hat{\theta}_1$ the copula parameter denoted by θ_2 can be estimated as

$$\hat{\theta}_2 = \text{ArgMax}_{\theta_2} \sum_{i=1}^n \ln c(F(x_i), G(y_i); \theta_2, \hat{\theta}_1) \quad (5.6)$$

Hence, the IFM estimator is denoted by

$$\hat{\theta}_{IFM} = (\hat{\theta}_1, \hat{\theta}_2)' \quad (5.7)$$

Furthermore, the log-likelihood function l , l_j denotes the j th marginals' log-likelihood function and l_c is the log-likelihood function of copula. By that, the equation for the IFM estimation is given as

$$\left(\frac{\partial l_1}{\partial \theta_{11}}, \frac{\partial l_2}{\partial \theta_{12}}, \frac{\partial l_c}{\partial \theta_2} \right) = \mathbf{0}' \quad (5.8)$$

in contrast, the MLE is calculated by

$$\left(\frac{\partial l}{\partial \theta_{11}}, \frac{\partial l}{\partial \theta_{12}}, \frac{\partial l}{\partial \theta_2} \right) = \mathbf{0}' \quad (5.9)$$

The estimators calculated by the (5.8) and (5.9) are not necessarily equivalent, and in order to get an exact MLE, the IFM clearly provides a great start.

5.3 Non-Parametric Ways

5.3.1 Maximum Pseudolikelihood Estimator

It is also possible to estimate the copula parameters without determining marginal distributions. This approach based on the ranks and interchanging the marginal distributions by their empirical counterparts. So, to work with this method, transformation of sample data $\{X_i, Y_i\}_{i=1}^n$ to their ranked observations $\{R_i, S_i\}_{i=1}^n$ is obtained. The procedure for the estimation of copula parameters is given below.

1. Describe the marginal distributions by their empirical distributions, which denoted as $\hat{F}(x_i)$ and $\hat{G}(y_i)$;

$$\hat{F}(x) = \frac{1}{n+1} \sum_{i=1}^n 1(X_i \leq x) ,$$

$$\hat{G}(y) = \frac{1}{n+1} \sum_{i=1}^n 1(Y_i \leq y) .$$

2. By maximizing the rank-based log-likelihood argument the estimation of the parameter is obtained.

$$l(\boldsymbol{\theta}) = \sum_{i=1}^n \ln\left(c\left(\frac{R_i}{n+1}, \frac{S_i}{n+1}\right)\right) , \quad (5.10)$$

$$\frac{\partial}{\partial \boldsymbol{\theta}} l(\boldsymbol{\theta}) = \sum_{i=1}^n \frac{\frac{\partial}{\partial \boldsymbol{\theta}} c\left(\frac{R_i}{n+1}, \frac{S_i}{n+1}\right)}{c\left(\frac{R_i}{n+1}, \frac{S_i}{n+1}\right)} = 0 . \quad (5.11)$$

5.3.2 Estimation of Parameter by Kendall's Tau

Clearly, we have a direct relationship between Kendall's tau and the parameter of a particular distribution. With this relation, it is possible to estimate a parameter which depends on the estimated sample version of Kendall's tau. Consider the FGM distribution which has the Kendall's tau measure as

$$\tau = \frac{2\theta}{9}$$

The empirical form of the sample version of (4.2.1) given as

$$\tau_n = \frac{c_n - d_n}{\binom{n}{2}} \quad (5.12)$$

With the estimated sample value of τ , we can estimate the parameter θ as

$$\hat{\theta} = \frac{9\tau_n}{2}$$

Note that the empirical form of τ again based on ranks which make this adaptation a nonparametric one.

5.4 A Fundamental Problem of Modeling

Instinctively, a researcher tends to make assumptions for a model of a particular data as the first move in an analysis and suspect that it can be described by a particular probabilistic family. These initial observations can be based on the historical experience of research belongs to a particular area or visual representation of the dataset. Although such observations give an idea for characteristics of data, it is not enough for a final decision. In order to conduct a realistic approach to the fit condition of data, the goodness of fit (GoF) procedures must be used. GoF also provides feedback to see that if initial assumptions were true. Let F be the

distribution function and vessel for observations in a dataset, then we are testing

$$H_0 = F = F_0$$

vs

$$H_1 = F \neq F_0 .$$

Here F_0 is the chosen distribution. Another way to describe the test is that

$$H_0 = F \in \mathfrak{F}$$

vs

$$H_1 = F \notin \mathfrak{F} ,$$

where $\mathfrak{F}$ is family of distributions denoted as $\mathfrak{F} = \{F_{\theta}, \theta, \Theta\}$.

The tests that can be conducted for any distribution known as a universal procedure, which means that during the process, the properties of studied distribution or family are not considered. Here, we focus on the universal processes.

The goodness of fit test with a universal procedure is based on the construction of empirical distributions. In another words, it can be shown as $\mathbb{F}_n = \sqrt{n}(F_n - F_0)$ Fermanian (2013). Some famous tests follow that approach are; Kolmogorov-Smirnov (KS), Anderson-Darling (AD), Cramer von Mises (CvM).

The classical GoF can be easily modified for copulas as

$$H_0 = C = C_0$$

vs

$$H_1 = C \neq C_0$$

or

$$H_0 = C \in \mathfrak{C}$$

vs

$$H_1 = C \neq \mathfrak{C}$$

where $\mathfrak{C}$ is studied copula family; $\mathfrak{C} = \{C_\theta, \theta \in \Theta\}$. In the same manner, empirical copulas can be adapted to GoF procedure. Empirical copulas are first defined in the studies of Deheuvels (1979), Deheuvels (1981). For $u, v \in I^2$ empirical copulas are given as

$$C_n(u, v) = H_n(F_n^{-1}(u), G_n^{-1}(v)) \quad (5.13)$$

or

$$\overline{C}_n(u, v) = \frac{1}{n+1} \sum_{i=1}^n 1(F_n(x_i) \leq u, G_n(y_i) \leq v) \quad (5.14)$$

Easily one can conclude that $\|C_n - \overline{C}_n\|_\infty \leq 2n^{-1}$ Fermanian et al. (2004). Asymptotically one will get the same result from GoF regardless of working with C_n or $\overline{C}_n$.

5.4.1 Kolmogorov-Smirnov Test Statistic

The well known test statistic Kolmogorov-Smirnov is established by the core idea of measuring distance between empirical copula $\overline{C}_n(u, v)$ and theoretical copula C_θ where θ is a convenient estimator representing the distribution. The test statistic is given by

$$T_{KS} = \sup_{u, v \in I^2} |\sqrt{n}(\overline{C}_n(u, v) - C_\theta)(u, v)| \quad (5.15)$$

5.4.2 Cramer von Mises Test Statistic

Procedure of CvM test statistic is again a natural GoF approach regarding measuring the distance between empirical and theoretical copula. The testing is done by using the ranked observations. According to the authors Genest, Rémillard, and Beaudoin 5.14 is the most objective benchmark for GoF copula considering that it is a nonparametric

procedure. The test statistic for CvM is given by Genest et al. (2008)

$$T_{CvM} = n \int_{I^2} \left[\bar{C}_n\left(\frac{R_i}{n+1}, \frac{S_i}{n+1}\right) - C_\theta\left(\frac{R_i}{n+1}, \frac{S_i}{n+1}\right) \right]^2 d\left[\bar{C}_n\left(\frac{R_i}{n+1}, \frac{S_i}{n+1}\right)\right] \quad (5.16)$$

$$= \sum_{i=1}^n \left[\bar{C}_n\left(\frac{R_i}{n+1}, \frac{S_i}{n+1}\right) - C_\theta\left(\frac{R_i}{n+1}, \frac{S_i}{n+1}\right) \right]^2. \quad (5.17)$$

5.4.2.1 Bootstrap Procedure

Bootstrap techniques help to conduct non-parametric GoF tests. After calculating the test statistic, we need to compare it with values of the studied probability family. For this goal, samples are generated randomly which depend on the parameter of the original data. For each sample, test statistics are calculated, then among them, one is chosen for comparison. The value is chosen according to the procedure of a particular bootstrap. Finally, p-value for the test can be determined again by the procedure's description. Note that, each calculated test statistic is named as the bootstrapped test statistics. We use the bootstrap procedure with Cramer von Mises with following steps.

- Test statistic is calculated as in (5.17),
- M samples are generated based on the distribution with originally estimated parameter, and we calculate test statistic for each sample with the formula above. Let's call them S_j 's. Note that for all the samples, their parameters are estimated and used in individual calculations,
- Arrange all S_j from small to large and choose the term which is, $S_{\lfloor (1-\alpha)N \rfloor}$. Computed value provides a comparison for the test statistics,
- Finally p-value is determined by $\frac{1}{N} \#\{j : S_j \geq CM_n\}$.

CHAPTER SIX

MANAGING AND MEASURING RISKS

6.1 Overview

In the world of the economy, typical move of an organization or a businessman notice a promising area for investment and took the opportunity. Carrying out some serious investment could have some effective results later on. The volatile nature of prices brings out an unknown nature for the success of an investment in the future. If changes in the financial dynamics go unpredictably, they can cause a loss whose amount depends on the magnitude of the investment. Moreover, here the loss we refer could be millions that are resulting from a sudden change in financial flow. This kind of possibilities brings the need of *risk* awareness. During the decision phase of business without evaluating the risk factors surely one should not press the button to start it. Risk analysis is the vital aspect of the decision process in financial affairs.

Making concrete factors of risk is one of the main efforts in the statistical discipline. By concrete, we refer to measuring and managing risk. There exist many methods to succeed that. Before continuing with the details first let us establish an understanding of the concept of risk.

In literature, the risk is generally defined as the mischance of having severe consequences or possible negative outcomes of an event. The risk goes hand in hand with the terms; uncertainty, decisions, consequences and events. In this paper, we cover the management of financial risks. We can explain the financially risky situations as chances that might keep an investor or organization from achieving their strategic goals. Alternatively, with another point of view, we can say the financial risk is the possibility of loss or gaining less amount of returns than expected. Hence risk is highly connected with the sense of randomness. In 1933 A. N. Kolmogorov brought a sensible definition Kolmogorov (1933) for the matters of randomness and probability. That is the well known axiomatic probability system (for details see

Kolmogorov 1933.)

To give an example to a fundamental risk model, let the random variable X tells us the value of risk, moreover, let X is governed by the distribution function $F(x) = P(X \leq x)$ which considers the probability by the end of a period. This probability measures the value of risk which we called X less than or equal to a fixed value of x . If the financial risk management of a particular market has more than one risky situations (i.e. it is mostly the case) random vectors of $(X_1, X_2, \dots, X_n)$ is used to model the risk.

For the most common financial risk that encountered in the industry, we can immediately give *market risk* as an example. The market risk involves the chances of change in the value of a financial asset. Such changes in the value occurs due to variations in the bottom line elements referred as exchange rates, commodity prices, stock and bond prices, etc.

As we mentioned before there exist a variety of approaches to measure and manage risk. To make it clear with a simple process, suppose our portfolio has the number of n investments with weights $w_1, w_2, \dots, w_n$. Hence the volatilities in the value of portfolio given as

$$X = \sum_{i=1}^n w_i X_i$$

where X_i represents the change in the i th investment and observation for the portfolio is restricted within a specified period.

To achieve measuring the risk factors of financial activity, finding an appropriate distribution for the portfolio is crucial. If there are multiple investments, in this case, a careful preparation of for joint model is needed. Hence tackling with measurement of a financial risk is primarily based on the statistical theory. Ingredients of risk analysis consist of the history of data, given portfolio structure, statistically estimating the distribution accurately where such distributions describe the change in the values. Observation of changes on prices points out a well-known data structure namely the time series. Indeed they have volatile changes and such characteristics cause risky environments. In the next parts, we continue with the elements of financial time series

and focus on how to model volatility of financial data. Then, we present the crucial process for risk analysis namely *GARCH* (generalized autoregressive conditionally heteroscedastic). And in the next part, we investigate the notions of approaching the risk issue. First, we start with the primary purposes that risk measurement is used, then continue with the ways of general approaches to risk, finally sum up with the phenomenons for the risk assessment, namely Value-at-Risk (VaR) and Conditional Value-at-Risk (CoVaR), and backtesting methods to test the adequacy of VaR results.

6.2 Financial Time Series

Financial time series are set of empirical investigations. Inferences of such series are mainly applied to analyze series of risk changes such as log-returns on equities, indexes, commodity prices and exchange rates. Long time investigations of financial time series led some econometrical facts arose based on this experience. These facts for the time series usually still hold for longer time intervals(weekly or monthly returns) or for shorter time intervals (intra-daily returns). The facts are given as McNeil et al. (2015)

1. Return series have a little serial correlation. However, they are not iid,
2. Absolute or squared return series has a heavy serial correlation,
3. Conditional expected returns are nearly zero,
4. Volatility seems to change by time,
5. Return series have heavy tails,
6. Extreme returns found in clusters.

Remark 6.2.1. Note that for analysis of time series the data is converted into the logarithmic returns. Say $\{S_t\}_{t=0}^n$ are index rate series then daily log-returns are

computed by

$$X_t = \ln\left(\frac{S_t}{S_{t-1}}\right), \quad t = 1, \dots, n$$

We continue with the basic definitions of time series analysis which are used in the modeling risk of return series.

A time series model is a family of random variables denoted by $\{X_t\}_{t \in \mathbb{Z}}$. Moreover, it is a stochastic process with one risk factor.

Definition 6.2.1. (strict stationary) The time series $\{X_t\}_{t \in \mathbb{Z}}$ is said to be strictly stationary if

$$(X_{t_1}, X_{t_2}, \dots, X_{t_n}) \stackrel{d}{=} (X_{t_1+k}, X_{t_2+k}, \dots, X_{t_n+k}) ,$$

$$\forall t_1, t_2, \dots, t_n, k \in \mathbb{Z} \text{ and } \forall n \in \mathbb{N} .$$

Definition 6.2.2. (covariance stationary) The time series $\{X_t\}_{t \in \mathbb{Z}}$ is said to be covariance stationary if there exist first moments and they satisfy

$$\mu(t) = \mu \quad t \in \mathbb{Z} ,$$

$$\gamma(t, s) = \gamma(t + k, s + k) \quad t, s, k \in \mathbb{Z} .$$

Remark 6.2.2. Preceding stationary definitions were given for the understanding of the behaviors of time series. Change on the mean, variance or covariances are systematically inconsistent with stationary.

Before defining autocorrelation function note that, covariance stationary means that $\gamma(t - s, 0) = \gamma(t, s) = \gamma(s, t) = \gamma(s - t, 0)$ holds $\forall s, t$. Thus covariance between X_t and X_s solely based on the separation of $|s - t|$, this situation is known as *lag*. Therefore, autocovariance function to describe covariance stationary process is written as

$$\gamma(h) = \gamma(h, 0) \quad \forall h \in \mathbb{Z}$$

Also, $\gamma(0) = \text{Var}(X_t) \quad \forall t$

Definition 6.2.3. (autocorrelation function) Covariance stationary process $\{X_t\}_{t \in \mathbb{Z}}$ has

the autocorrelation function (ACF) $\phi(h)$ with

$$\phi(h) = \phi(X_h, X_0) = \frac{\gamma(h)}{\gamma(0)} \quad \forall h \in \mathbb{Z} .$$

Here autocorrelation or serial correlation refer to $\phi(h)$ at lag h .

Desired time series models are build up by *white noise processes* which are *stationary processes without serial correlation*.

Definition 6.2.4. (White Noise) $\{X_t\}_{t \in \mathbb{Z}}$ is said to be a white noise process if it is covariance stationary determined by the autocorrelation function

$$\phi(h) = \begin{cases} 1, & h = 0, \\ 0, & h \neq 0 . \end{cases}$$

A centered white noise process is denoted as $WN(0, \sigma^2)$ where $\sigma^2 = Var(X_t)$.

Definition 6.2.5. (Strict White Noise) $\{X_t\}_{t \in \mathbb{Z}}$ is said to be strict white noise process if its random variables iid with finite variance. And centered strict white noise process is denoted as $SWN(0, \sigma^2)$ in the same manner.

Remark 6.2.3. Note that, SWN is the simplest noise process. Before going into GARCH which is an another noise process, we need to define the martingale-difference. First, assume time series $\{X_t\}_{t \in \mathbb{Z}}$ is filtered to acquire $\{\mathcal{F}_t\}_{t \in \mathbb{Z}}$ which denotes the information growth over time. Also, filtered data contains the information of past and present values of the time series $\{X_s\}_{s \leq t}$. The history up to a horizon t is represented by $\mathcal{F}_t = \sigma(\{X_s : s \leq t\})$. Such filtrations are also known as *natural filtration*.

Definition 6.2.6. (martingale-difference) Time series $\{X_t\}_{t \in \mathbb{Z}}$ is called as a martingale-difference sequence which is adapted to the filtration $\{\mathcal{F}_t\}_{t \in \mathbb{Z}}$ if for $E|X_t| < \infty$ X_t is $\mathcal{F}_t$ -measurable(adapted) and

$$E(X_t | \mathcal{F}_{t-1}) = 0 \quad \forall t \in \mathbb{Z} .$$

Clearly unconditional mean is zero i.e.

$$E(X_t) = E(E(X_t|\mathcal{F}_{t-1})) = 0 \quad \forall t \in \mathbb{Z}.$$

Also if $\forall t \ E(X_t^2) < \infty$ then for autocovariances, we have

$$\begin{aligned} \gamma(t, s) &= E(X_t, X_s) \\ &= \begin{cases} E(E(X_t X_s | \mathcal{F}_{s-1})) = E(X_t E(X_s | \mathcal{F}_{s-1})) = 0, & t < s, \\ E(E(X_t X_s | \mathcal{F}_{t-1})) = E(X_s E(X_t | \mathcal{F}_{t-1})) = 0 & t > s \end{cases} \end{aligned}$$

Therefore a martingale-difference sequence with finite variance has zero mean and zero covariance. If, for each t , the variance is constant, then we have a white noise process.

6.2.1 GARCH Models

For modeling, the risk factors of return series GARCH models are of the essence. Because this process uses the past variance to forecast the variance state in the future. With an alternative view, it captures the homoscedastic nature of volatile changes in time series and leads us to accurate results. One can see it also as a preparation of financial data for a risk management. Preference of GARCH model depends on two important aspects. Firstly, it considers general risk factors of time series rather than lost in detailed risk analysis work of single series hold a bigger significance. A simple GARCH process indeed follows that idea. Secondly, GARCH process provides an excellent starting point regarding quantitative work for a financial data and makes a reasonable interpretation for possible volatility. GARCH process is defined by Bollerslev (1986).

6.2.1.1 GARCH Process

Definition 6.2.7. Let $\{Z_t\}_{t \in \mathbb{Z}}$ be a $SWN(0, 1)$. The $\{X_t\}_{t \in \mathbb{Z}}$ process is said to be a GARCH(p,q) one if it is strictly stationary and if $\forall t \in \mathbb{Z}, \exists \{\sigma_t\}_{t \in \mathbb{Z}}$ (strictly positive

values sequence) following holds

$$X_t = \sigma_t Z_t, \quad \sigma_t^2 = \alpha_0 + \sum_{i=1}^q \alpha_i X_{t-i}^2 + \sum_{j=1}^p \beta_j \sigma_{t-j}^2, \quad (6.1)$$

where

$$\alpha_0 > 0 \quad \alpha_i \geq 0 \quad i = 1, \dots, p \quad \beta_j \geq 0 \quad j = 1, \dots, q.$$

Here the squared volatility σ_t^2 depends on the past.

Mostly, low order GARCH models are preferred. High volatilities are stubborn in this model, because if $|X_{t-1}|$ or σ_{t-1} is large $|X_t|$ can be large too. From 6.1 model of GARCH(1,1) corresponds to

$$\sigma_t^2 = \alpha_0 + (\alpha_1 Z_{t-1}^2 + \beta) \sigma_{t-1}^2. \quad (6.2)$$

6.2.1.2 EGARCH Process

EGARCH process is slightly different version of classical GARCH and originally introduced by Nelson (1991). The version we consider here varies from the original one and there are many different approaches too. The EGARCH process with the leverage effect is defined by

$$\log \sigma_t^2 = \alpha_0 + \sum_{i=1}^q \alpha_i \left[\left(\frac{|X_{t-i}|}{\sigma_{t-i}} \right) - E \left(\frac{|X_{t-i}|}{\sigma_{t-i}} \right) \right] + \sum_{i=1}^q L_i \left(\frac{X_{t-i}}{\sigma_{t-i}} \right) + \sum_{j=1}^p \beta_j \log \sigma_{t-j}^2, \quad (6.3)$$

where

$$E(|Z_{t-j}|) = E\left(\frac{X_{t-j}}{\sigma_{t-j}}\right) = \sqrt{\frac{v-2}{\pi}} \frac{\Gamma(\frac{v-1}{2})}{\Gamma(\frac{v}{2})},$$

for the Student's t distribution with $v > 2$.

Note that EGARCH is different, because it has not nonnegativity conditions for the parameters like the standard GARCH has. Also, volatility clustering is permitted in GARCH by the α_i and β_j combinations, but in EGARCH β_j captures the volatility.

Moreover, in this process Z_t plays a forcing role for both conditional variance and the error. Finally, to comment on leverage effect note that, it is a principal in theory of economics which states, information about market should affect the volatility asymmetrically. Simply put, it considers the tendency that a decrease in equity value of a company contributes the volatility.

6.3 Foremost Goals of Risk Measurement

- *Determining the capital amount* that an institution needs to have is crucial since such accumulation will be a compensation for unexpected losses in the future and prevent a regulator from worry about the solvency of the firm. Or similarly, to plan an investors margin requirements for an organized investment move, determining the risks is essential.
- Measurement of risk is a *tool of management*. Measures are used to draw a limit for the risk level that is logical to take. To give an example, bank traders mostly try not to exceed a bound of daily 95% VaR.
- Compensation for risks of an insurance company is provided by *insurance premiums*. The amount of the compensation is determined by the risk of insured claims.

McNeil et al. (2015)

6.4 Ways of Measuring Risks

We will categorize the basic approaches to financial risk measurement into four parts.

6.4.1 Notional Amount Approach

This quantitative method for analyzing the risks of a portfolio is the oldest one. In this method, the summation of notional values of each security implies risk factors of a portfolio. Here every notional value can be weighted by a risk factor which defines the risk level of an asset. The advantage of the approach is its simplicity. It is criticized for that; the method does not consider the difference between long and short positions and time does not interpret how to diversify overall risk.

6.4.2 Factor-Sensitivity Measures

By this method, the change of a portfolio can be determined by the previously investigated change that a risk factor causes. These measures help to explain the robustness of a portfolio which makes them useful. However, they lack measuring the overall risk.

6.4.3 Risk Measures Based on Loss Distributions

These techniques are a modern way to describe risk. For a decided horizon t they help to analyze conditional and unconditional loss distribution of a portfolio. Trying to analyze by this approach requires a solid calibration of a distribution which will describe the loss. Sometimes this creates tough problems, especially when the estimation of loss distribution is needed for large portfolios. But when problems overcome the model, it gives accurate results for risk of a portfolio since its realistic approach. Other than that, in this method we are relying on the past experiences of the market changes, this angle is naturally limited since dynamics of governing financial markets are constantly changing. To use this approach effectively, loss distribution should be supported with hypothetical scenarios of risk. Furthermore model should have the characteristics which cover the expectation of an investor. For example, modeling of volatilities along with the statistical computation is very significant to

achieve dependable results.

6.4.4 Scenario-Based Risk Measures

This method simply deals with measuring of the portfolio risk by using the information of possible risk scenarios in the future. The risk is determined as the maximum possible loss that might be the result when all scenarios involve. Here, the advantage of the approach is that it is very explanatory and useful for the portfolios can be effected by small risk factors. And the downside is accuracy of scenarios which added to model and deciding their weights.

6.5 Value-at-Risk

In this part, we give the most preferred risk measurement technique, Value-at-Risk (VaR). Say we have a portfolio which includes risky assets and a chosen time horizon t . The distribution function indicating the return can be defined as $F_R(l) = P(R \leq r)$. Now, we want to interpret F_R statistically to measure how risky is our portfolio over the horizon t . We can define VaR as; under a fixed probability, it is the threshold for loss which is unlikely to be overlapped.

Definition 6.5.1. (VaR) For a determined confidence level $\alpha \in (0, 1)$, VaR is calculated as the loss that exceeded with probability (α) or will not be exceeded with probability $1 - \alpha$. Mathematically speaking, it is,

$$VaR_\alpha = \sup\{r \in \mathbf{R} : F_R(r) = P(R \leq r) \leq \alpha\} . \quad (6.4)$$

So we can see that, VaR is the α quantile of return distribution. Common α usages are, 0.05 and 0.01. By the definition of VaR, it does not consider the information of loss for the probability below α . This situation is taken into account as a drawback for the method. Moreover, if there are two different assets in a portfolio, one needs to consider two return distributions belong to them and their weights. Let R_1 and R_2 be the random

variables for two different assets, then their weights on the portfolio can be described by $R = bR_1 + (1 - b)R_2$ where $b \in (0, 1)$. As a result, VaR can be computed with the specified weights on the two assets for a fixed time horizon t with the same manner in 6.4.

$$VaR_\alpha = \sup\{l \in \mathbf{R} : P(bR_1 + (1 - b)R_2 \leq r) \leq \alpha\} \quad (6.5)$$

6.6 Backtesting

After conducting a risk analysis by applying VaR, the whole analysis has not been done yet. Because there is an important last bit of the analysis which explains whether the calculated VaR is reflecting the market characteristics realistically. In other words, we can test how accurate and adequate our risk estimates are. Methods to test the adequacy of a VaR model called as backtesting. Brown indicated that “VaR is only as good as its backtest. When someone shows me a VaR number, I don’t ask how it is computed, I ask to see the backtest.” to emphasize the role of backtesting in VaR analysis. Brown (2008)

In backtesting methods, the motivation is comparing the real data of profits and losses with the estimated VaR results. If an estimation of VaR is made based on the $1 - \alpha$, then the expectation is that probability of VaR violation is close to the given probability level $1 - \alpha$. To be specific let the chosen confidence level is 0.99, then in an appropriate VaR model exception should occur on one day out of 100 days. A very basic approach to applying backtesting is checking within a time horizon whether the number of exceptions are synchronous with the confidence level. But there is more to that, to mention adequacy of VaR model also should have independent exceptions or not clustered exceptions. Clustered exceptions are the sign of inaccuracy in the modeling of market volatilities. The following procedures of backtesting pursue the mentioned expectancies from a VaR model.

6.6.1 POF-Test

The proportion of failures (POF) test is a well-known one to investigate if the exceptions in the VaR model are reasonable for the determined confidence level. It is proposed by Kupiec (1995). Let T be the observation size, x be the number of VaR violations and c be the confidence level, then for the test, we have the null hypothesis as Dowd (2006)

$$H_0 : p = \hat{p} = \frac{x}{T} .$$

The hypothesis helps to control if the proportion of failures close to the chosen interval level. The test statistics for POF carry out with likelihood-ratio (LR) test,

$$POF = -2\log\left(\frac{(1-p)^{T-x}p^x}{(1-\frac{x}{T})^{T-x}(\frac{x}{T})^x}\right) \quad (6.6)$$

POF is asymptotically chi-squared distributed with the degree of freedom one if the model is appropriate. The LR value is bigger than the critical value of chi-square distribution means that we reject the null hypothesis which suggests the model is not adequate. With more information provided in a dataset, one can reject a false model much less difficulty Jorion (1997). But POF test is not sufficient for a complete backtest since it does not consider the relationship between exceptions, but only numbers of them. For this reason, there is a possibility that it can fail to reject a model which has clustered exceptions.

6.6.2 Mixed Kupiec Test

To investigate the dependency between VaR violations, this test provides an effective way. Because it is based on an exhaustive analysis for dependency and this feature is important, since an occurrence of a violation can be originated from a violation which belongs five days ago for example. Haas (2001) offered the mixed Kupiec test aiming to capture the dependence with a more general view. For every exception, test statistic

yields

$$LR_i = -2\log\left(\frac{(p(1-p))^{v_i-1}}{(\frac{1}{v_i})(1-\frac{1}{v_i})^{v_i-1}}\right) \quad (6.7)$$

Here v_i denotes the time passed between exceptions i and $i-1$. By considering each exception in the test statistic, the null hypothesis established as that the exceptions are independent events. By combining all exception, the test statistic for the independence is given by

$$IND = \sum_{i=2}^n -2\log\left(\frac{(p(1-p))^{v_i-1}}{(\frac{1}{v_i})(1-\frac{1}{v_i})^{v_i-1}}\right) - 2\log\left(\frac{(p(1-p))^{v-1}}{(\frac{1}{v})(1-\frac{1}{v})^{v-1}}\right) \quad (6.8)$$

Where v is the time until the first exception and n is the number of total exceptions. Similarly, the test statistic has chi-squared distribution with n degrees of freedom. With the combination of POF test, one can adopt a mixed test which considers both exception rate and independence. This version of test statistic yields

$$MIX = POF + IND \quad (6.9)$$

The test statistic MIX has a chi-squared distribution with $n+1$ degrees of freedom. If the test statistic is greater than the critical value, we reject the null hypothesis, and it suggests that the considered VaR model is not adequate regarding exceptions and independence.

6.7 Conditional Value at Risk

Considering the risk participation effect of outside markets is crucial to understand risk status of interested portfolio. Because, spillovers from an highly influential market may lead very dramatic results. Hence, notion of systemic risk analysis is

emerged. Classic VaR method is lacking in capturing systemic risk, because it measures risk of an isolated economic structure. As a systemic risk measure CoVaR is very famous. Basically, it is modified version of traditional VaR method with conditional part. CoVaR is designed to analyse the VaR of asset(institution) i , when asset(institution) j is distressed. Hence, CoVaR provides interpretations for risk participations of outside interference. CoVaR measure was firstly proposed by Brunnermeier & Adrian (2011) and then generalized by Girardi & Ergün (2013). These two approaches differ in the setup in the conditional part. Let R_i and R_j be random variables represent two return series and we are interested in how R_i is affected when R_j is in trouble. Brunnermeier & Adrian defined the condition as $R_j = VaR_j$ that is portfolio of j exactly at its VaR threshold. Girardi & Ergün used a different way and defined condition as $R_j \leq VaR_j$ and it refers that j can be beyond its VaR. Latter approach provides more detailed view (effects of the events about the tails) in the analysis of systemic risk. Additionally, Mainik & Schaanning (2014) showed that CoVaR and dependence parameter are monotonic in the second model, so that as association increases CoVaR also increases. In contrast, when returns are equal to their VaR, CoVaR decreases at high dependence level. This inconsistency of the first model could lead very problematic decisions since CoVaR fails to detect risk in case of high association.

According to the Brunnermeier & Adrian, CoVaR can be obtained similar to the quantile function of VaR as

$$P(R_i \leq CoVaR_{\alpha, \beta}^- \mid R_j = VaR_{\beta}) = \alpha \quad (6.10)$$

where β is determined as the confidence level for the quantile financial condition i.e., $P(R_j \leq VaR_{\beta}) = \beta$.

Later CoVaR equation was generalized by Girardi & Ergün as

$$\begin{aligned}
P(R_i \leq CoVaR_{\alpha,\beta}^{\leq} \mid R_j \leq VaR_{\beta}) &= \alpha \\
\frac{P(R_i \leq CoVaR_{\alpha,\beta}^{\leq}, R_j \leq VaR_{\beta})}{P(R_j \leq VaR_{\beta})} &= \alpha \\
P(R_i \leq CoVaR_{\alpha,\beta}^{\leq}, R_j \leq VaR_{\beta}) &= \alpha\beta
\end{aligned} \tag{6.11}$$

Models are referred as $CoVaR_{\alpha}^{\leq}$ and $CoVaR_{\alpha}^{\leq}$.

After calculating the CoVaR, in order to realize the risk contribution of asset j , another CoVaR must be calculated which considers R_j , when it has stable dynamics. In other words, financial condition needs to have a benchmark state. We set the benchmark state as median of the R_j . In this case $\beta = 0.5$ and in the light of this, the risk contribution of j then can be calculated as

$$\Delta CoVaR_{\alpha,\beta} = CoVaR_{\alpha,\beta} - CoVaR_{\alpha,0.5} \tag{6.12}$$

Remark 1. Let R_i and R_j have the joint distribution H and marginals F_i and F_j respectively. H has the copula distribution $C(u, v)$ which is differentiable with respect to u , it has a strictly positive density function such that,

$$d(u, v) := \frac{\partial C(u, v)}{\partial u}$$

Functions C and d are strictly increasing and invertible with respect to v for each fixed $u \in [0, 1]$. For the equations (6.10) and (6.11) we have the following theorems.

Theorem 6.7.1. *Based on the assumptions in Remark 1, for given levels of $\alpha, \beta \in (0, 1)$ equation (6.10) and (6.11) can be calculated as*

$$CoVaR_{\alpha,\beta}^{\leq} = F_i^{-1}(d^{-1}(\alpha, \beta)) \tag{6.13}$$

$$CoVaR_{\alpha,\beta}^{\leq} = F_i^{-1}(C^{-1}(\alpha, \beta)). \tag{6.14}$$

These adopted formulas are generalized way of the implicit CoVaR equations, and they are quite handy for calculations. Both of them based on the transformation $v = F_i(CoVaR_{\alpha,\beta})$ and $u = F_j(VaR_\beta)$ and involves the calculation of the variable v , with the equations $d^{-1}(\alpha, \beta) = v$ and $C^{-1}(\alpha, \beta) = v$. Hence if we define confidence level as $\tilde{\alpha} = d^{-1}(\alpha, \beta)$ or $\tilde{\alpha} = C^{-1}(\alpha, \beta)$, it is clear that CoVaR is nothing but VaR calculation with level $\tilde{\alpha}$. Moreover, by the Theorem 6.7.1. risk contribution 6.12 can be redefined as

$$\Delta CoVaR_{\alpha,\beta}^- = F_i^{-1}(d^{-1}(\alpha, \beta)) - F_i^{-1}(d^{-1}(\alpha, 0.5)) \quad (6.15)$$

$$\Delta CoVaR_{\alpha,\beta}^\leq = F_i^{-1}(C^{-1}(\alpha, \beta)) - F_i^{-1}(C^{-1}(\alpha, 0.5)) \quad (6.16)$$

CHAPTER SEVEN

APPLICATION

Sudden changes in oil prices is a very encountered phenomenon throughout the history of the world economy. Many of the literature found a significant effect of oil price shocks on the economies around the world. See, Hamilton (1983), Hamilton (2003), Kilian & Park (2009). The past decade had unprecedented swings in oil prices. In particular, for the five year bit, we noted such different dynamics. Factors that contributed this situation include the slow pacing in global growth and increased supply of oil and gas. The mentioned fluctuations in oil prices have significant effects on the economies around the world. Investigating the changes in a market of a country gives an idea about how that particular economy affected by external and internal factors and investment possibilities in the future. The structure of core financial elements such as oil and stock indices is effective in terms of acceleration in a country's economy. For the worst case scenario, price shocks in oil could lead a financial crisis globally. Since these are the well-known dynamics in oil markets, there exist many studies which are dedicated to investigating the effects of oil price shocks. Some of these studies examined the association between oil prices and stock indices. Mohanty et al. (2010) analyzed the relation between oil prices and stock returns of oil and gas for Central and Eastern European (CEE) countries by using linear factor pricing model, Mohanty et al. (2010). Aloui et al. (2013) used the time varying copula approach to investigate the dependence structure between Brent crude oil price and stock markets for CEE economies. Gulpinar and Katata (2014) studied on the modeling oil and gas supply disruption risks using extreme value copula.

7.1 Data and Empirical Results

In this study, we focus on the stock indices of various emerging markets. We explicitly analyze the relationship between changes in oil prices and stock market returns for emerging markets. These stock indices consist of South Korea, Taiwan,

Mexico, China, Brazil and Turkey. Our goal is to investigate how Brent oil and emerging markets are affected by each other and compute risk measures for these two assets. The relation between oil prices and stock market returns for emerging countries has not been examined in the concept of copula approximation. We are motivated by the fact that emerging markets have very different characteristic than developed ones. For example, the demand growth in emerging economies can lead to dramatic swings in oil prices during times of sluggish production. Also, markets are chosen from different regions to capture a wide angle for the study.

The data set has the information for daily prices. It includes closing prices of Europe Brent as dollars per barrel and closing prices of six stock indices. These indices are; Korea Composite Stock Price Index (KOSPI), the Taiwan Stock Exchange Corporation Weighted Index (TAIEX), the Indice de Precios y Cotizaciones (IPC, Mexican Index), the Shanghai Stock Exchange Composite Index (SSEC), the Bolsa de Valores de Estado de Sao Paulo Index (IBOVESPA), the Borsa Istanbul (BIST 100). These markets have the status, ‘emerging’ according to the current financial sources. The period of the whole data is from 4 January 2010 to 18 December 2015. The data is collected from investing.com which is a trusted financial portal. Plot of each index closings and Brent oil closings are shown in Figure 7.1. In this study, we split our financial sample into two parts namely in-sample and out-sample. For in-sample data ranges from 4 January 2010 to 10 July 2014 with 927 observations and for out-sample, from 11 July 2014 to 18 December 2015 with 300 observations. We built our analysis onto forecasting each daily risk factor through 300 days by using each information of out-sample data.

In particular, for the past five years of Brent crude oil prices, we examine a steady growth from the start of 2010. Later on, prices followed stable ups and downs close to \$100, reaching its peak in 2012 with \$125. At the end of 2014, Brent oil prices started a gradual decreasing and fallen below \$50 in January 2015. That level of decrease was examined for the first time since May 2009. The descriptive statistics for closing prices of six emerging markets and Brent oil is given in the Table 7.1.

Before starting the analysis, we need to filter our data since in financial time series

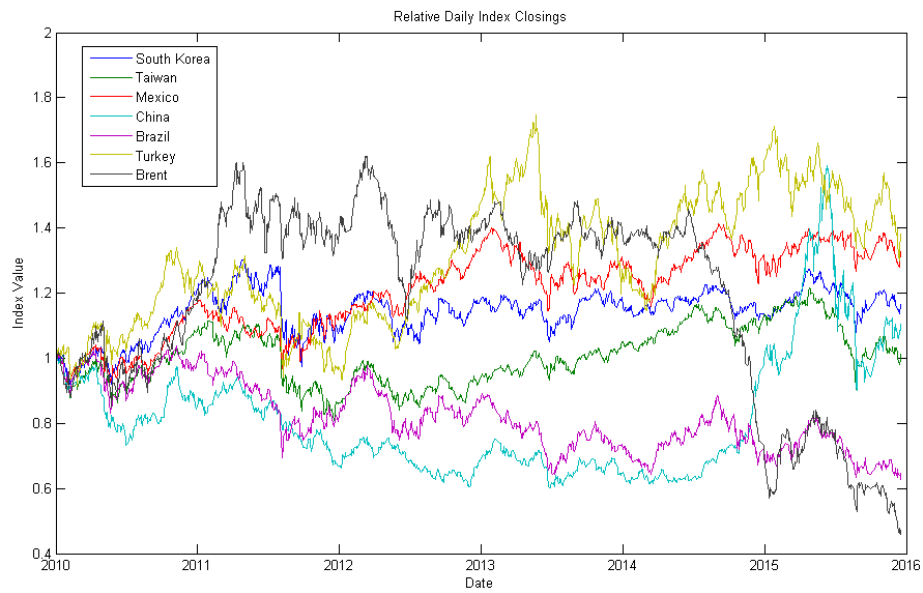
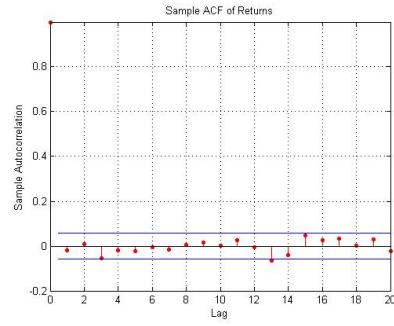


Figure 7.1 Daily closings through 2010-2015

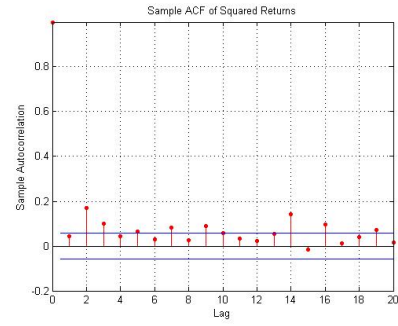
Table 7.1 Descriptive Statistics

Variable	Mean	StDev	Min	Median	Max	Skewness	Kurtosis
South Korea	1938.8	123.3	1552.8	1961.5	2208.4	-0.74	0.32
Taiwan	8291.2	730.2	6633.3	8259.5	9973.1	0.12	-0.88
Mexico	39542	4192	30368	40414	46313	-0.40	-0.97
China	2646.2	619.2	1950	2429.3	5166.4	1.48	2.29
Brazil	57442	7170	43911	56283	72996	0.27	-0.96
Turkey	69698	10478	48739	69037	93179	0.03	-1.11
Brent	93.714	22.941	36.290	104.665	128.140	-0.82	-0.56

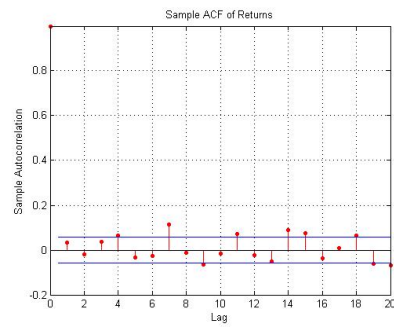
we encounter heteroscedasticity and some autocorrelation, hence we can not say they are i.i.d. observations as we emphasize in the section Financial Time Series in Chapter Six. Dynamics of returns for the whole sample can be visualized with autocorrelation graphs in Figure 7.2. Especially, for the squared returns we see big exceeding volatilities.



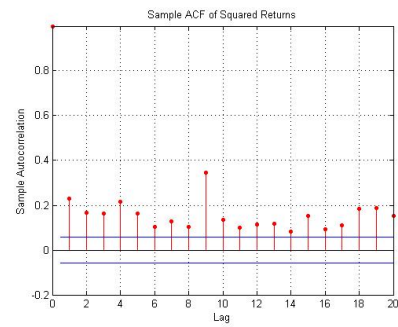
(a) Brazil & Brent



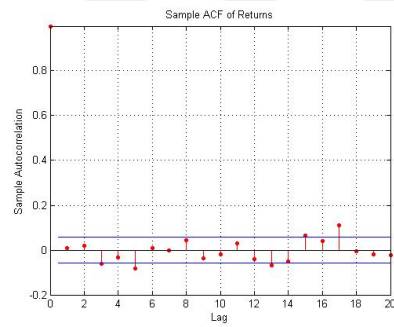
(b) Brazil & Brent Squared



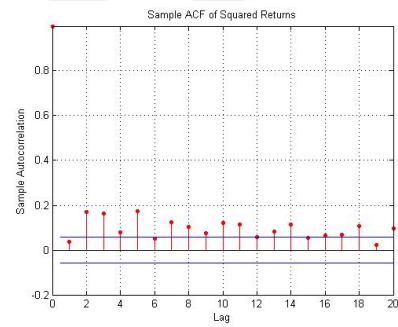
(c) China & Brent



(d) China & Brent Squared

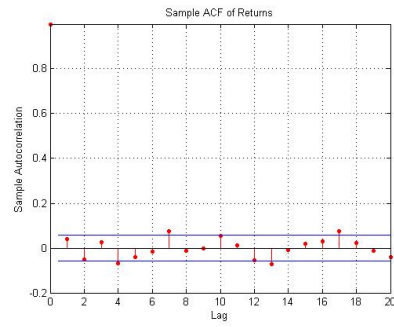


(e) Mexico & Brent

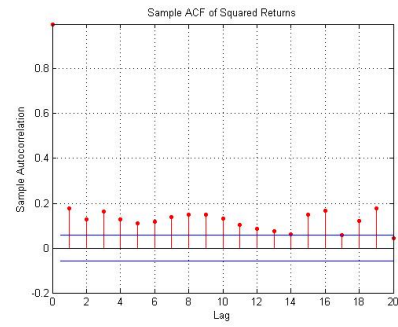


(f) Mexico & Brent Squared

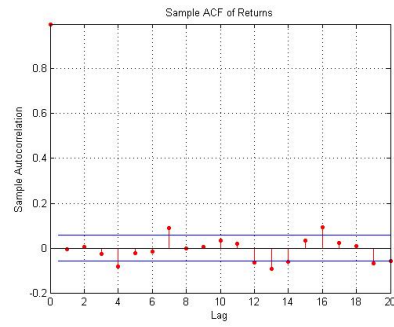
Figure 7.2 Autocorrelation graphics of returns and squared returns of datasets



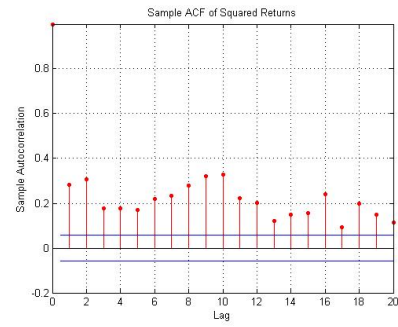
(g) Taiwan & Brent



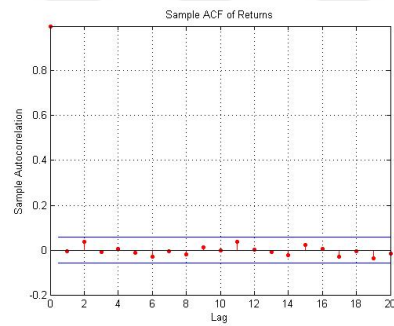
(h) Taiwan & Brent Squared



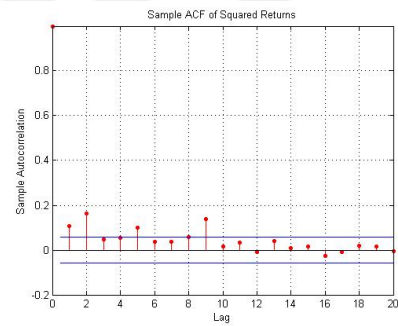
(i) S. Korea & Brent



(j) S. Korea & Brent Squared



(k) Turkey & Brent



(l) Turkey & Brent Squared

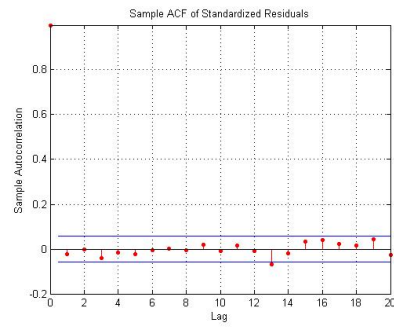
Figure 7.2 continues

Hence we conduct GARCH(1,1) process with using Student's t distribution innovations for marginals. Then we standardize the filtered data by dividing each return to its conditional standard deviation. Parameter results of EGARCH(1,1) is given in Table 7.2. The parameters are significant hence we have the proper filtered model with stationary variance. Additionally, after the process we see a homoscedastic structure of returns as in the autocorrelation graphics of standardized residuals in Figure 7.3.

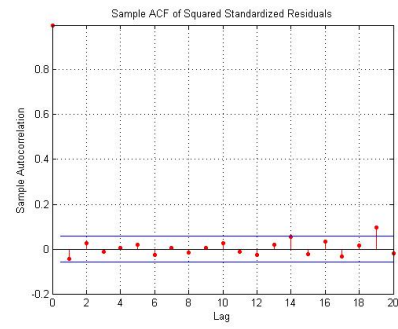
Table 7.2 Estimation of EGARCH(1,1) parameters (standard errors)

Parameter	SKorea	Taiwan	Mexico	China	Brazil	Turkey	Brent
C	-0.0005 (0.0002)	-0.0008 (0.0002)	-0.0006 (0.0002)	-0.0001 (0.0003)	6.15×10^{-6} (0.0003)	-0.0012 (0.0003)	0.0001 (0.0004)
K	-0.0844 (0.0457)	0.0044 (0.0346)	-0.1146 (0.0708)	-0.1046 (0.0569)	-0.1037 (0.0623)	-0.3934 (0.195)	1.0594e-06 (0.0281)
EGARCH(1)	0.9907 (0.0050)	1.0000 (0.0038)	0.9874 (0.0076)	0.9875 (0.0067)	0.9876 (0.0074)	0.9524 (0.0234)	0.9967 (0.0034)
ARCH(1)	0.1069 (0.0244)	0.0927 (0.0218)	0.1052 (0.0269)	0.1184 (0.0292)	0.0742 (0.0202)	0.1747 (0.0404)	0.0909 (0.0221)
Leverage	-0.1193 (0.0199)	-0.1080 (0.0151)	-0.1182 (0.0195)	-0.0129 (0.0179)	-0.0853 (0.0158)	-0.0748 (0.0253)	-0.0545 (0.0106)
DoF	7.6788 (1.6777)	6.1742 (1.1873)	6.2907 (1.3109)	4.7631 (0.8055)	9.7381 (2.9267)	6.0000 (1.1877)	6.3609 (1.3671)

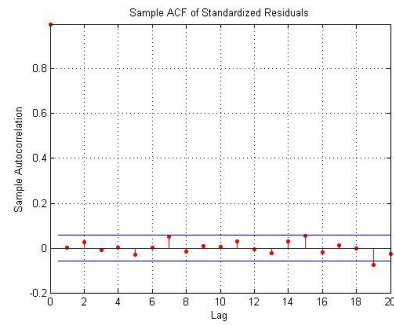
We check if our financial data fit to normal distribution by using Kolmogorov-Smirnov test, results are given in Figure 7.4. As an expected outcome, our time series data does not distribute normally. In fact by applying the goodness of fit test, we acknowledge that the data needs a good representation for its fat tails. The pairwise association is measured by Kendall's tau, and results are obtained in Table 7.3. The largest Kendall's tau is between South Korea and Taiwan with 0.4835. Outcome reasonable since these two far east countries have a great economic collaboration and their trading has an increasing trend. Second, largest is Brazil and Mexico with 0.3767. Similarly, that is possible because they have a good economic relationship. Furthermore, Brazil and Mexico have higher association with Brent oil with 0.1584 and 0.1563, respectively. It seems a reasonable result since they are making more oil exporting comparing with the other countries in our study.



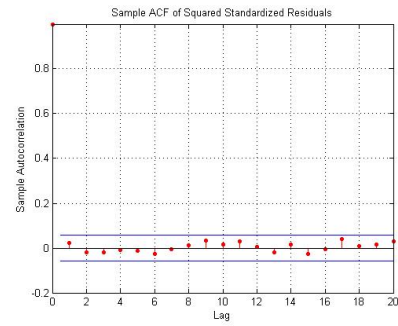
(a) Brazil & Brent



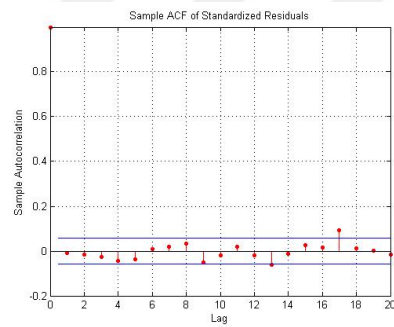
(b) Brazil & Brent Squared



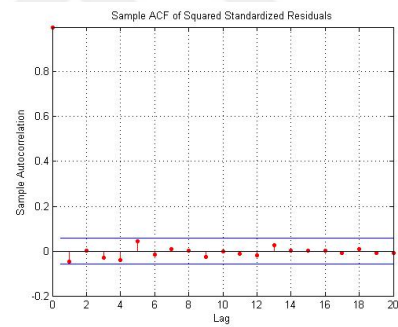
(c) China & Brent



(d) China & Brent Squared

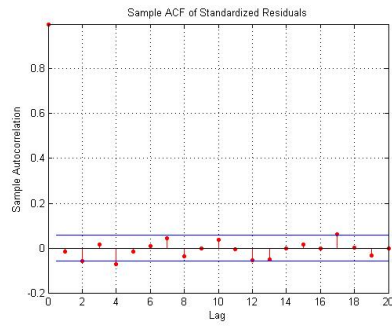


(e) Mexico & Brent

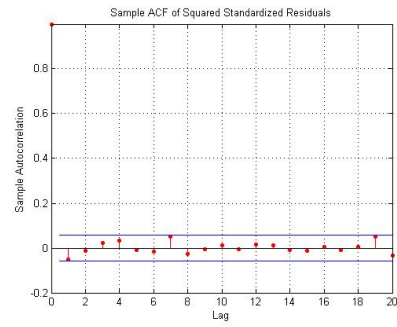


(f) Mexico & Brent Squared

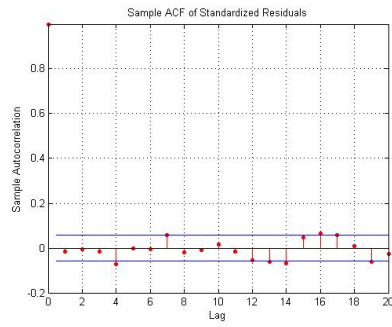
Figure 7.3 Autocorrelation graphics after EGARCH(1,1) process



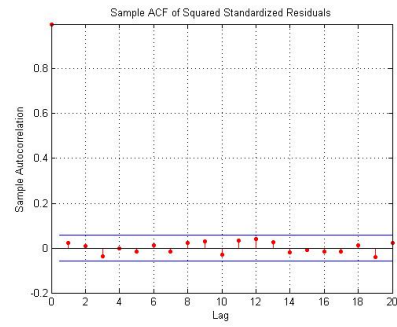
(g) Taiwan & Brent



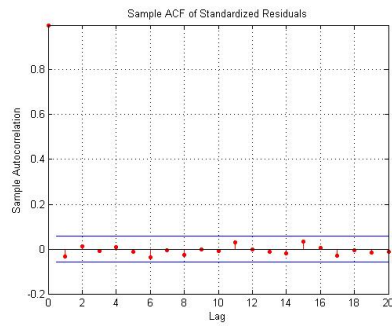
(h) Taiwan & Brent Squared



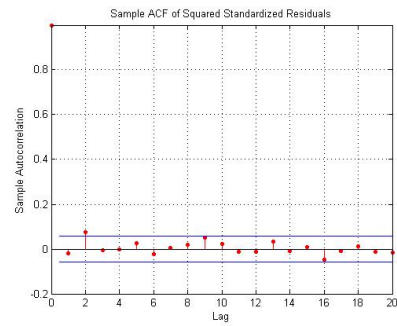
(i) S. Korea & Brent



(j) S. Korea & Brent Squared

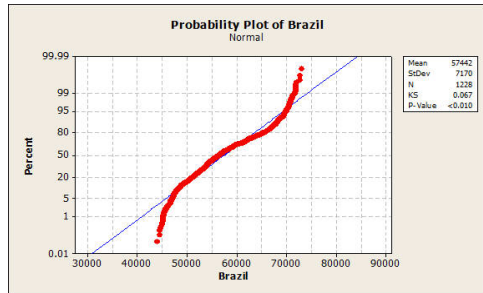


(k) Turkey & Brent

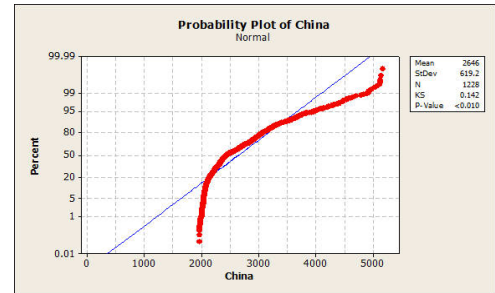


(l) Turkey & Brent Squared

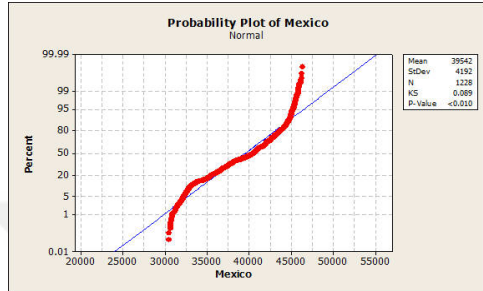
Figure 7.3 continues



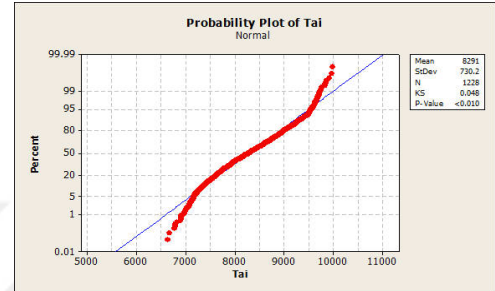
Brazil & Brent



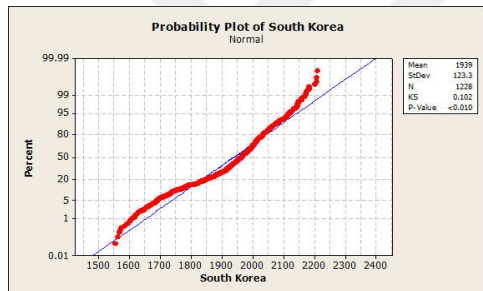
China & Brent



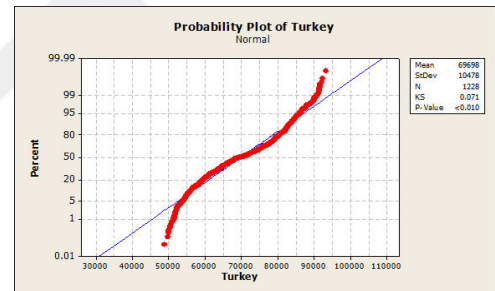
Mexico & Brent



Taiwan & Brent



S. Korea & Brent



Turkey & Brent

Figure 7.4 Probability plots of each datasets for Kolmogorov-Smirnov Test

Table 7.3 Kendall's tau values of the financial dataset

	South Korea	Taiwan	Mexico	China	Brazil	Turkey	Brent
South Korea	1.0000	0.4835	0.1748	0.1961	0.1608	0.1596	0.1361
Taiwan	0.4835	1.0000	0.1700	0.2184	0.1551	0.1461	0.1225
Mexico	0.1748	0.1700	1.0000	0.1121	0.3767	0.2273	0.1563
China	0.1961	0.2184	0.1121	1.0000	0.1089	0.0762	0.0910
Brazil	0.1608	0.1551	0.3767	0.1089	1.0000	0.2170	0.1584
Turkey	0.1596	0.1461	0.2273	0.0762	0.2170	1.0000	0.0962
Brent	0.1361	0.1225	0.1563	0.0910	0.1584	0.0962	1.0000

Parameter estimations are computed by using IFM method and results are given in Table 7.4. For all datasets, maximized log likelihood is the highest for t copula. In addition, in Taiwan & Brent dataset Clayton copula is relatively close to maximized log likelihood of t copula and for Turkey & Brent we see a large log likelihood value among others. Also, Gumbel in Brazil & Brent and Normal in South Korea & Brent is relatively high.

Cramer-von-Mises method is used for GoF testing as the mentioned steps in the Bootstrap procedure. We test at 95 % confidence level and with $M = 10000$ trial samples. Within the calculations, IFM is used for parameter estimation. The p-values for the CvM are displayed in Table 7.5. Each bivariate dataset has relatively high p-value among it that shows sharper clues. Brent vs. Taiwan is more fitting to Clayton copula with p-value 0.367. And all the other datasets are more fitting to Student's t copula due to their high p-values. Especially for Brent vs. Turkey shows a great fit to t copula with 0.9695. In the light of these, we use Clayton for one pair and Student's t copula for others in our study.

Table 7.4 Parameter Estimations (Standard Errors) [Maximized log likelihood]

	Gumbel	Frank	Clayton	Normal	T
South Korea & Brent	1.139 (0.022) [25.7]	1.307 (0.174) [27.77]	0.2711 (0.037) [28.67]	0.2215 (0.025) [30.28]	0.223 (0.029) [34.34]
Taiwan & Brent	1.125 (0.022) [22.43]	1.183 (1.171) [22.66]	0.2643 (0.039) [28.77]	0.2084 (0.026) [26.72]	0.205 (0.03) [30.27]
Mexico & Brent	1.196 (0.025) [50.8]	1.509 (0.174) [36.27]	0.3261 (0.044) [42.2]	0.275 (0.025) [47.39]	0.2608 (0.032) [56.66]
China & Brent	1.078 (0.02) [8.095]	0.8325 (0.1704) [11.39]	0.1674 (0.036) [12.49]	0.1355 (0.025) [11.14]	0.1408 (0.03) [16.2]
Brazil & Brent	1.187 (0.024) [44.71]	1.489 (0.174) [35.31]	0.3054 (0.042) [36.9]	0.2662 (0.024) [44.28]	0.2572 (0.03) [52.1]
Turkey & Brent	1.11 (0.02) [18.18]	0.9383 (0.163) [13.8]	0.2087 (0.035) [18.93]	0.1542 (0.023) [14.47]	0.1563 (0.032) [38.38]

Table 7.5 p-values (Test statistics) of Cramer-von Mises Test

	Gumbel	Frank	Clayton	Normal	T
Brent vs. SKorea	0.0586 (0.0025)	0.0286 (0.0638)	0.034 (0.084)	0.0208 (0.3241)	0.0187 (0.3912)
Brent vs. Taiwan	0.0842 (0.0002)	0.0498 (0.0015)	0.0216 (0.367)	0.0432 (0.009)	0.039 (0.0118)
Brent vs. Mexico	0.04340 (0.0165)	0.0529 (0.0009)	0.0646 (0.0033)	0.0403 (0.0139)	0.0270 (0.10)
Brent vs. China	0.0350 (0.0608)	0.0203 (0.2702)	0.0222 (0.3604)	0.0169 (0.5696)	0.0149 (0.6623)
Brent vs. Brazil	0.0280 (0.1215)	0.0493 (0.0019)	0.0807 (0.0005)	0.0329 (0.0471)	0.0244 (0.1527)
Brent vs. Turkey	0.0259 (0.1848)	0.0325 (0.0294)	0.040 (0.0421)	0.0252 (0.1789)	0.0098 (0.9695)

Before adding the tail analysis in the study, first we embark on a classical VaR computation covering a monthly risk measure. For this risk estimation, we determine the tail fraction for generalized Pareto distribution as 0.1 to fit the standardized residuals and use proportion of failure(POF) test for backtesting procedure. Please address the Value-at-Risk section below for details. In the backtesting phase, we compare the VaR values with simulated observations of returns. Results for the first VaR application of the study are displayed in Table 7.6. Each VaR estimate indicates that, at the end of the month the losses of portfolio will not exceed that particular estimated value with 90% probability. As one can see, the values are not that apart from each other. Brent & Brazil and Brent & China have relatively risky portfolios compared to the other ones. Furthermore, all models pass from POF test. This suggests for all VaR estimations, proportion of violations are tolerable hence models are adequate.

Table 7.6 Simulated VaR estimates for fifty-fifty Brent and Index assets with 90% confidence and p-values for the VaR models determined by POF test

	VaR Values	p-values of POF
Brent & SKorea	-0.104	0.9405
Brent & Taiwan	-0.0989	0.6533
Brent & Mexico	-0.1042	0.5467
Brent & China	-0.1170	0.4549
Brent & Brazil	-0.1286	0.7858
Brent & Turkey	-0.1075	0.5467

7.2 Tail Dependence Analysis

Tail behaviors between parametric and non-parametric approaches are investigated in this section. Tail dependence coefficients of six bivariate portfolios are calculated parametrically by the resulting equations in Example 4.3.1. and Example 4.3.2. and Huang upper tail dependence is used as a non-parametric measure.

To calculate the daily changes on TDC, the in-sample is shifted one day forward to a day from out-sample, after each estimation of TDC. For the datasets which are fitted to t copula, nonparametric UTD is also added since t copula has symmetric tails. Results of TDC is given in Figure 7.5. For Brazil(IBOVESPA) Index & Brent nonparametric estimations are distant to parametric ones. Nonparametric UTD is much closer and overall, estimations behave similar. In China(SSEC) Index & Brent nonparametric LTD starts up high then continues in a steady manner, nonparametric UTD is first going steady then get much higher in the last bit. Nonparametric estimations distant to parametric ones. Nonparametric UTD and parametric tail estimations are very close and behave the same in Mexico(IPC) Index & Brent until between the days 100-150 and nonparametric LTD is generally higher than the other ones. For Taiwan(TAIEX) & Brent which fitted to Clayton copula, share same trends in nonparametric and parametric estimations through the period. For Turkey(BIST 100) & Brent, non-parametric LTD has same trends as parametric ones but nonparametric UTD has big leaps after day 100. Overall, there are considerable differences between parametric and nonparametric changes, nonparametric estimations seem rather stagnant than the parametric ones. Generally, parametric tail estimations are more volatile and mostly has lesser tail coefficients.

7.3 Value-at-Risk

VaR calculations for equally weighted emerging market & Brent portfolios are conducted by using both parametric and nonparametric approach. Following

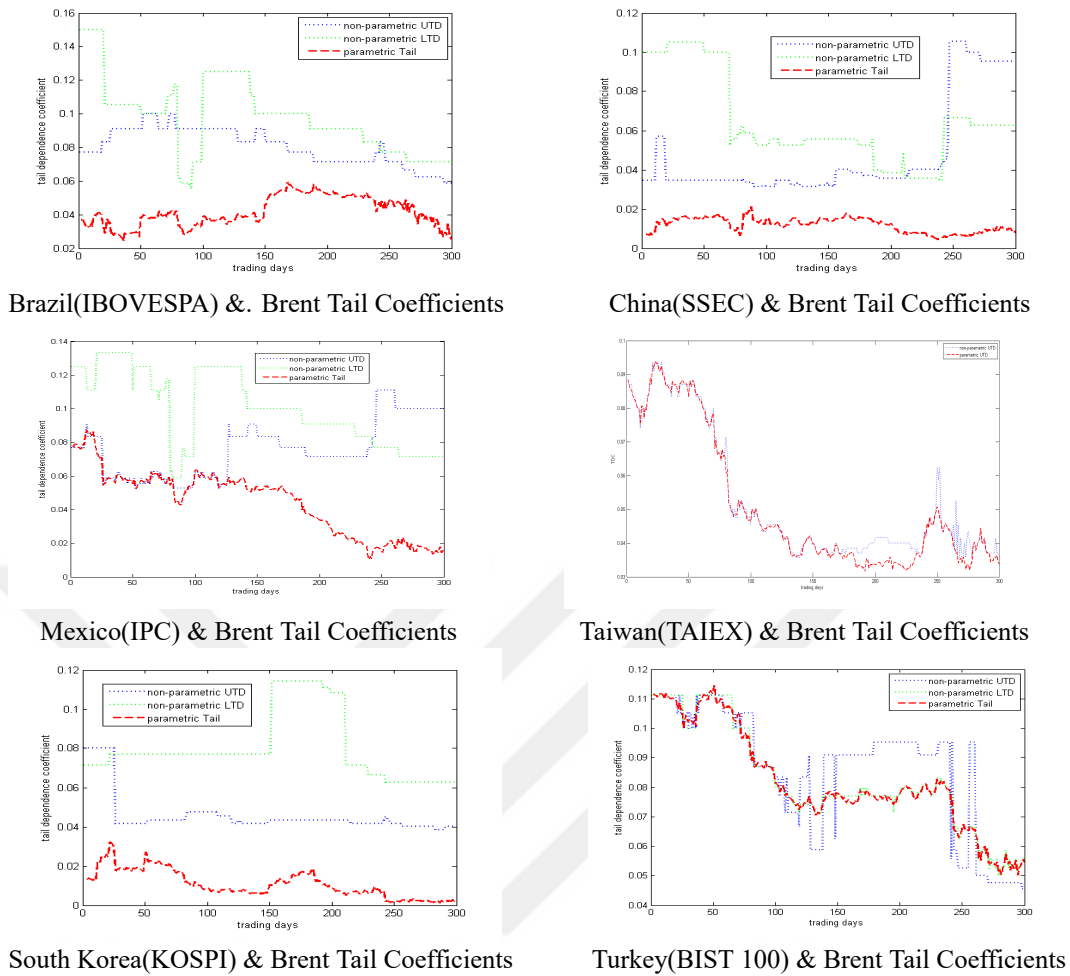


Figure 7.5 Visuals of parametric and non parametric tail coefficients

out-sample forecasting algorithm is used for 300 days of VaR estimations. The algorithm is similar which used in Siburg et al. (2015).

- First we re-arrange 300 return samples by updating in-sample with a day from out-sample. Later on, EGARCH(1,1) procedure is applied to each sample with Student's t innovations and standardized residuals are obtained.
- To find the percentage of extreme observations in each sample Huang nonparametric estimation is used to obtain lower tail (4.25), and the k values are obtained.
- To capture the tail information about the tails standardized residuals are fitted to generalized Pareto distribution by the tails with determined k values as tail

fractions. And they are fitted to empirical distribution by the center. Hence marginals are constructed.

- For the parametric part, marginals are transformed into uniform data and their parameters are estimated by IFM. For nonparametric part, parameters are estimated by using the Huang's LTD estimator.
- Number of 2000 trials or samples are generated by the determined copula distribution and parameter. By using the information of marginals, inverses of those samples are obtained. In other words, first copula then standardized residuals are simulated.
- By using the last information on the standardized residuals and returns, conditional standard deviations, returns and obtained inverses of samples, simulations of returns are created by GARCH simulation in MATLAB software.
- Finally VaR is calculated for each equally weighted simulated returns with for 0.05 quantile.

Results of VaR simulations are visualized by including actual returns of out-samples in Figure 7.6 and Table 7.7. At first glance, we examine similar trends and little differences between parametric and nonparametric calibration of VaR estimates. Generally, all of VaR estimations keep a cautious tendency below the actual returns with just small amount of sharp cuts. Graphics and descriptive table suggest nonparametric VaR estimates seem to capture the stress more for Brazil & Brent, China & Brent and South Korea & Brent. Other portfolios seem capture stress more for parametric modeling, especially for Brent & Turkey, parametric model is superior as relatively big difference between standard deviations suggest. But there is a necessity to investigate risk analysis in detail. Because during forecasting for large time periods, overlooking some bits of important information is quite possible and skipping these information may lead to unexpected losses. Hence we examine the performance of VaR models in detail by conducting backtesting. The results of Mix

Kupiec Test (which has the test statistic given in (6.9)) are displayed in Table 7.8. According to the p-values, all VaR models came out adequate except the parametric model of South Korea & Brent due to irremissible underestimation of risk. Nonparametric calibration outperforms parametric ones in four cases due to their higher p-values. In contrast for Brent & Taiwan, it seems both models have an indistinguishable closeness and for Brent & Turkey parametric model is superior due to strong fit of t copula as GoF test showed this evidence before.

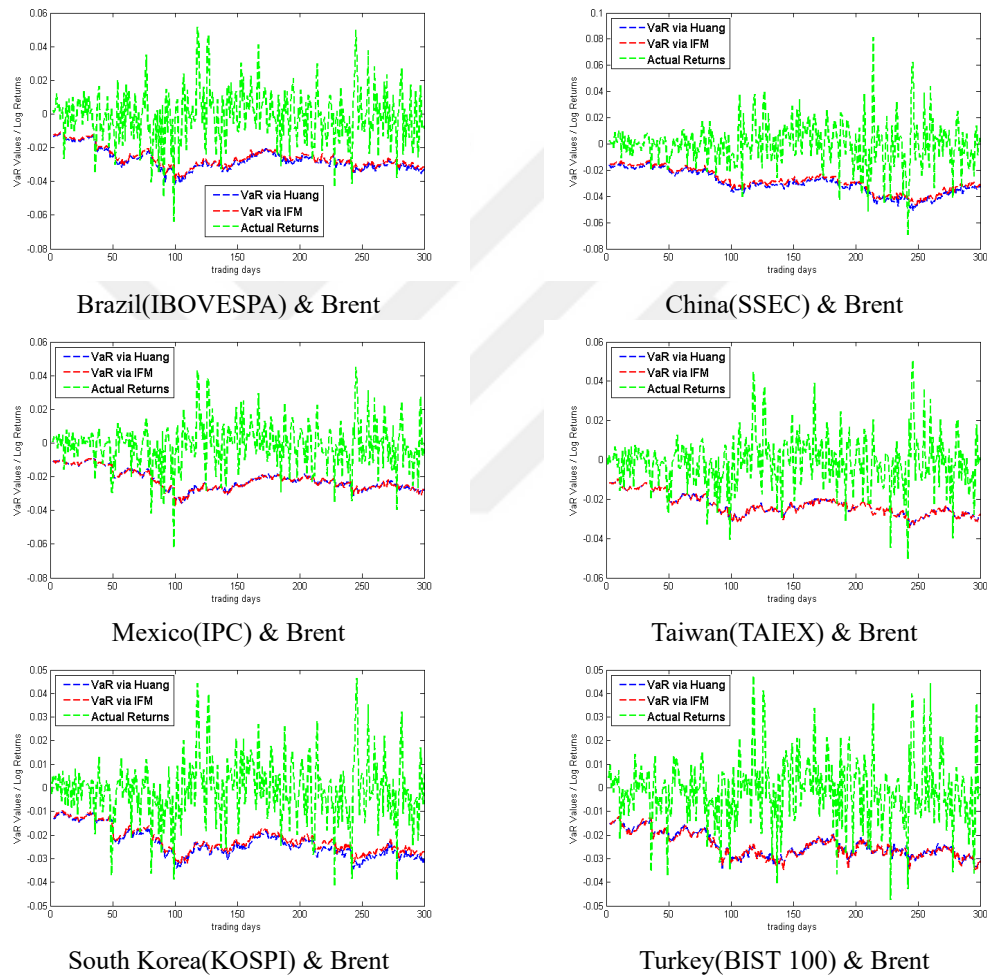


Figure 7.6 Parametric and nonparametric estimations of VaR for 300 trading days

7.4 Conditional Value-at-Risk

In the next step of our application, we determine financial condition as distressed Brent market and analyse its risk participation to the emerging markets. We use both

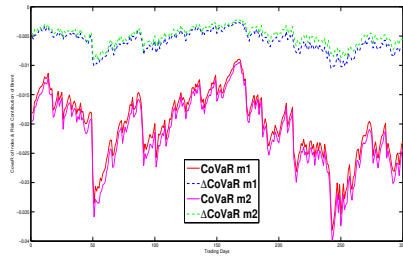
Table 7.7 Descriptive Statistics for VaR Estimations

	Method	Min	Mean	Max	Std. Dev.
Brent & SKorea	Para.	-3.2448	-2.1976	-0.9806	0.5484
	Nonpara.	-3.4418	-2.3333	-1.0236	0.5933
Brent & Taiwan	Para.	-3.4430	-2.3054	-1.0065	0.5428
	Nonpara.	-3.4288	-2.2975	-1.0017	0.5409
Brent & Mexico	Para.	-3.7896	-2.2185	-0.9211	0.6042
	Nonpara.	-3.7676	-2.2130	-0.9325	0.6016
Brent & China	Para.	-4.6817	-2.8257	-1.2805	0.8567
	Nonpara.	-5.1456	-3.0369	-1.3994	0.8801
Brent & Brazil	Para.	-3.9758	-2.5996	-1.0983	0.5932
	Nonpara.	-4.1593	-2.7288	-1.1780	0.6175
Brent & Turkey	Para.	-3.5150	-2.4638	-1.2375	0.5190
	Nonpara.	-3.4195	-2.4445	-1.2838	0.5103

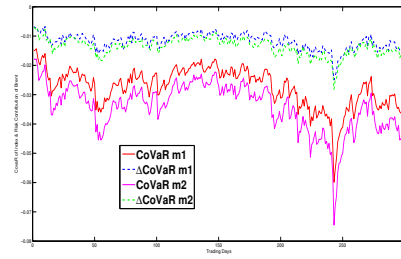
Table 7.8 Backtesting results with 0.05 significance

	VaR via parametric model	VaR via non-parametric model
Brent & SKorea	0.0341	0.1788
Brent & Taiwan	0.3643	0.3643
Brent & Mexico	0.1303	0.3455
Brent & China	0.3502	0.3814
Brent & Brazil	0.2124	0.2503
Brent & Turkey	0.5168	0.2183

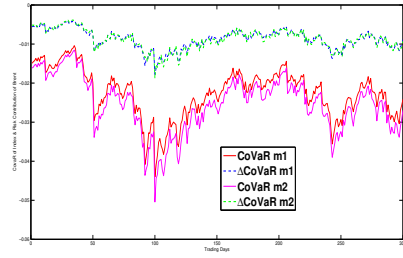
models that proposed by Brunnermeier & Adrian and Girardi & Ergün to calculate CoVaR to see the investigate difference between those approaches and clearly visualize the results when a system relevant financial condition exceeds the threshold. The CoVaR computation framework is similar to the VaR estimation algorithm explained before. The difference is that parametric estimations are computed for all of the simulated returns and by using those numerical calculations are carried out to obtain quantiles $\tilde{\alpha}_1$, $\tilde{\alpha}_2$ for each day. The descriptive outcomes and graphs of $\Delta CoVaR$ values can be examined in Table 7.9 and Figure 7.7. Moreover violations for all CoVaR models are displayed in Table 7.10. According to the table, China is the most affected market from the unhealthy Brent oil returns. Among the other emerging markets, China has more oil activity in its market, hence this result is expected. On the other hand, Brazil and Turkey are exposed to a considerable affect from Brent. Figures show that second model of CovaR is able to capture more risk, which suggests first model underestimates the risk. This observation also can be seen in the trend of risk participation as well. Lastly, we consider the table of violations and joint violations. Joint violations are the days which we encounter CoVaR violations within the days Brent already has VaR exceedance. In the table for violations, we see that there is no evidence first model outperforms the second one. And our CoVaR models are able to capture violations despite the short time period of 34 days.



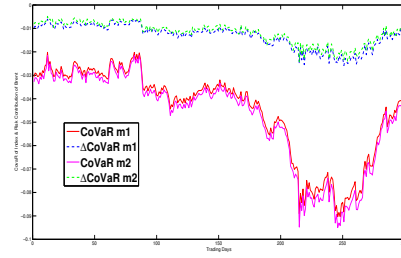
Brazil(IBOVESPA)



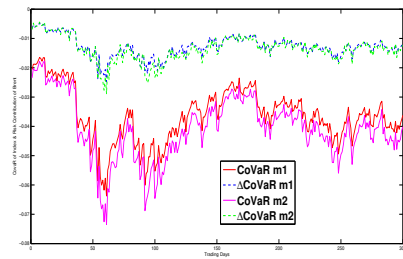
China(SSEC)



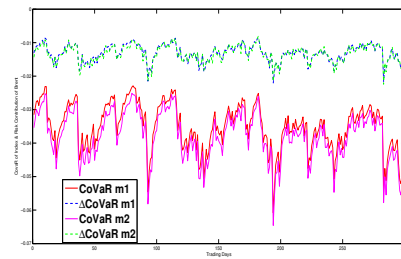
Mexico(IPC)



Taiwan(TAIEX)



South Korea(KOSPI)



Turkey(BIST 100)

Figure 7.7 CoVaR estimations of through 300 trading days with $\alpha = 0.05$ and $\beta = 0.05$, y-axis being CoVaR values and risk contribution and x-axis being trading days

Table 7.9 Descriptive Statistics for $\Delta CoVaR$ in Percentage

	Method	Min	Mean	Max	Std. Dev.
Brent & SKorea	$\Delta CoVaR_{0.05,0.05}^=$	-2.7954	-1.4361	-0.6274	0.4063
	$\Delta CoVaR_{0.05,0.05}^{\leq}$	-3.2042	-1.6569	-0.7220	0.4720
Brent & Taiwan	$\Delta CoVaR_{0.05,0.05}^=$	-3.4956	-1.5689	-0.7414	0.4132
	$\Delta CoVaR_{0.05,0.05}^{\leq}$	-4.6329	-2.0631	-0.9997	0.5396
Brent & Mexico	$\Delta CoVaR_{0.05,0.05}^=$	-2.6280	-1.3914	-0.6297	0.3821
	$\Delta CoVaR_{0.05,0.05}^{\leq}$	-3.1683	-1.6473	-0.7448	0.4525
Brent & China	$\Delta CoVaR_{0.05,0.05}^=$	-6.7313	-3.2432	-1.3503	1.4408
	$\Delta CoVaR_{0.05,0.05}^{\leq}$	-7.3171	-3.6223	-1.5541	1.5512
Brent & Brazil	$\Delta CoVaR_{0.05,0.05}^=$	-3.9308	-2.3575	-1.1144	0.5781
	$\Delta CoVaR_{0.05,0.05}^{\leq}$	-4.5571	-2.7262	-1.3086	0.6664
Brent & Turkey	$\Delta CoVaR_{0.05,0.05}^=$	-3.8759	-2.1951	-1.3742	0.4505
	$\Delta CoVaR_{0.05,0.05}^{\leq}$	-4.3071	-2.4401	-1.5503	0.4920

Table 7.10 Individual and joint violation numbers where joint violations investigated among the VaR violations of Brent which has number of 34 trading days

		# Vio.	# Joint Vio.
Korea	$CoVaR_{\alpha,\beta}^=$	4	2
	$CoVaR_{\alpha,\beta}^{\leq}$	4	2
Taiwan	$CoVaR_{\alpha,\beta}^=$	4	2
	$CoVaR_{\alpha,\beta}^{\leq}$	0	0
Mexico	$CoVaR_{\alpha,\beta}^=$	6	4
	$CoVaR_{\alpha,\beta}^{\leq}$	3	2
China	$CoVaR_{\alpha,\beta}^=$	8	1
	$CoVaR_{\alpha,\beta}^{\leq}$	8	1
Brazil	$CoVaR_{\alpha,\beta}^=$	6	3
	$CoVaR_{\alpha,\beta}^{\leq}$	5	2
Turkey	$CoVaR_{\alpha,\beta}^=$	6	1
	$CoVaR_{\alpha,\beta}^{\leq}$	5	1

CHAPTER EIGHT

CONCLUSION

The theory of copula has many applications in the financial fields due its specialty of capturing details and being practical. In this paper, we are motivated by the useful features of copula distributions and the idea of combining copula first with the well known risk measure Value-at-Risk and then more efficient measure Conditional Value-at-Risk for a sensitive and realistic risk assessment. Also, we focus using copula distributions on analyzing the comovements and risk factors between emerging markets and Brent oil prices and attempt to forecast risk percents for 300 trading days. Because these assets are very sensitive to changes and constitute an important place in the future of the world economy. Emerging markets have different characteristics than the developed ones and a change in world economy might have a big impact on them. On the other hand, fluctuations of Brent oil prices have a big potential to shake the balances in the world economy. Thus, it is very fitting to use copula distributions as a vessel for these assets to capture their extreme behaviours. So, we form a bivariate dataset of emerging market and Brent oil for 1228 days of trading days where all days are matching. In addition to that, we benefit from the tail information. To do so, we use Huang tail dependence coefficient to find the tail percentages of datasets and use it to fit our standardized residuals to generalized Pareto distributions. Because generalized Pareto distribution is efficient to capture the information about the tails.

In this paper, our main focuses are the investigation of the risk environment between emerging market and Brent market and compare the performances of risk models. In order to apply our aims, first we analyze the datasets as univariate and validated that marginals fit to Student's t-distribution. Then to adopt a stationary variance for bivariate datasets, EGARCH(1,1) process is used and filtrations are standardized by dividing each of them to their conditional standard deviation. Proper copula distributions for financial datasets are determined by the Cramer-von-Mises test statistic with bootstrap method. Six bivariate portfolios are fitted to the copulas

with parametric (IFM estimations) and non-parametric calibration (Using Huang tail estimates for estimation of parameters). Daily VaR forecasts for 300 trading days are carried out by using the information of in-sample and out-sample. With 95% confidence, daily VaR estimations of six portfolios show a moderate risk which suggest the chosen emerging markets have similar patterns. We can say that six portfolios do not have a big risky situation. On the other hand, the most risky markets are Brent & China and Brent & Brazil with their relatively higher mean and standard deviation as the descriptive study of VaR forecasting shows. According to the results of Mixed Kupiec test, most of the models carried out by nonparametric tail estimation more adequate since they are not underestimate risk which can lead more violations. In fact, parametric model of South Korea & Brent could not pass the test. In general, other ones resulted adequate in terms of expected violation number and independence between days which have violation. Additionally, models of Brent & Taiwan has very close estimations of risk such that it did not reflect on the backtesting. The only model which can be interpreted better parametrically is Brent & Turkey, because of strong fit. In the final step we carry out two CoVaR model calculations to see the impact of Brent oil to the emerging markets. Overall, Brent has a relatively stable risk participation on the emerging markets. However, China show some extreme values with large standard deviation. This suggests that this market can show unstable behaviours while Brent is being distressed. Also, there is no huge difference between the CoVaR models which suggests, the condition $R_j \leq VaR_j$ does not lead to extreme CoVaR values. Regardless, using the generalized CoVaR model is healthy in terms of precision and consistency in risk model.

REFERENCES

- Aloui, R., Hammoudeh, S., & Nguyen, D. K. (2013). A time-varying copula approach to oil and stock market dependence: The case of transition economies. *Energy Economics*, 39, 208–221.
- Alsina, C., Frank, M., & Schweizer, B. (2003). Problems on associative functions. *Aequationes mathematicae*, 66(1-2), 128–140.
- Ang, A., & Chen, J. (2002). Asymmetric correlations of equity portfolios. *Journal of financial Economics*, 63(3), 443–494.
- Bollerslev, T. (1986). Generalized autoregressive conditional heteroskedasticity. *Journal of econometrics*, 31(3), 307–327.
- Brown, A. (2008). Private profits and socialized risk—counterpoint: Capital inadequacy. *Global Association of Risk Professionals*. June/July, 8.
- Brunnermeier, M., & Adrian, T. (2011). Covar. *Federal Reserve Bank of New York Staff Reports*, 348.
- Clayton, D. G. (1978). A model for association in bivariate life tables and its application in epidemiological studies of familial tendency in chronic disease incidence. *Biometrika*, 65(1), 141–151.
- Deheuvels, P. (1979). La fonction de dépendance empirique et ses propriétés. académie royale de belgique. *Bulletin de la Classe des Sciences*, 65(5), 274–292.
- Deheuvels, P. (1981). A kolmogorov-smirnov type test for independence and multivariate samples. *Revue roumaine de mathématiques pures et appliquées*, 26(2), 213–226.
- Dowd, K. (2006). Retrospective assessment of value at risk. *Risk Management: A Modern Perspective, Amsterdam et al*, 183–203.
- Embrechts, P., McNeil, A., & Straumann, D. (2002). Correlation and dependence in risk management: properties and pitfalls. *Risk management: value at risk and beyond*, 176223.

- Esary, J. D., Proschan, F. et al. (1972). Relationships among some concepts of bivariate dependence. *The Annals of Mathematical Statistics*, 43(2), 651–655.
- Fang, K.-T., Kotz, S., & Ng, K. W. (1990). *Symmetric multivariate and related distributions*. Chapman and Hall.
- Fermanian, J.-D. (2013). An overview of the goodness-of-fit test problem for copulas. In *Copulae in Mathematical and Quantitative Finance* (61–89). Springer.
- Fermanian, J.-D., Radulovic, D., Wegkamp, M. et al. (2004). Weak convergence of empirical copula processes. *Bernoulli*, 10(5), 847–860.
- Ferreira, M. S. (2013). Nonparametric estimation of the tail-dependence coefficient. *RevStat*, 11(1), 1–16.
- Fisher, N. I. (1997). Copulas. *Encyclopedia of statistical sciences*, 159–163.
- Frank, M. J. (1979). On the simultaneous associativity of (x, y) and $x+y - F(x, y)$. *Aequationes mathematicae*, 19(1), 194–226.
- Genest, C. (1987). Frank's family of bivariate distributions. *Biometrika*, 74(3), 549–555.
- Genest, C., & MacKay, J. (1986). The joy of copulas: Bivariate distributions with uniform marginals. *The American Statistician*, 40(4), 280–283.
- Genest, C., & Mackay, R. J. (1986). Copules archimédiennes et familles de lois bidimensionnelles dont les marges sont données. *Canadian Journal of Statistics*, 14(2), 145–159.
- Genest, C., Remillard, B. et al. (2008). Validity of the parametric bootstrap for goodness-of-fit testing in semiparametric models. In *Annales de l'Institut Henri Poincaré, Probabilités et Statistiques*, volume 44, (1096–1127). Institut Henri Poincaré.
- Girardi, G., & Ergün, A. T. (2013). Systemic risk measurement: Multivariate garch estimation of covar. *Journal of Banking & Finance*, 37(8), 3169–3180.

- Gumbel, E. J. (1960). Bivariate exponential distributions. *Journal of the American Statistical Association*, 55(292), 698–707.
- Haas, M. (2001). New methods in backtesting. *Financial Engineering Research Center, Bonn*.
- Hamilton, J. D. (1983). Oil and the macroeconomy since world war ii. *Journal of political economy*, 91(2), 228–248.
- Hamilton, J. D. (2003). What is an oil shock? *Journal of econometrics*, 113(2), 363–398.
- Hoeffding, W. (1941). Masstabinvariante korrelationsmasse für diskontinuierliche verteilungen. *Archiv für mathematische Wirtschafts-und Sozialforschung*, 7, 49–70.
- Höfdding, W. (1940). Masstabinvariante korrelationstheorie. *Schriften des Mathematischen Instituts und Instituts für Angewandte Mathematik der Universität Berlin*, 5, 181–233.
- Huang, X. (1992). Statistics of bivariate extreme values: Statistiek van bivariate extreme waarden.
- Joe, H., & Xu, J. J. (1996). The estimation method of inference functions for margins for multivariate models.
- Jorion, P. (1997). *Value at risk: the new benchmark for controlling market risk*. Irwin Professional Pub.
- Kilian, L., & Park, C. (2009). The impact of oil price shocks on the us stock market. *International Economic Review*, 50(4), 1267–1287.
- Kolmogorov, A. (1933). Grundbegriffe der wahrscheinlichkeitsrechnung (engl. transl. foundations of probability, chelsea (1956)).
- Kruskal, W. H. (1958). Ordinal measures of association. *Journal of the American Statistical Association*, 53(284), 814–861.

- Kupiec, P. H. (1995). Techniques for verifying the accuracy of risk measurement models. *The Journal of Derivatives*, 3(2), 73–84.
- Lehmann, E. L. (1966). Some concepts of dependence. *The Annals of Mathematical Statistics*, 1137–1153.
- Li, D. X. (2000). On default correlation: A copula function approach. *The Journal of Fixed Income*, 9(4), 43–54.
- Li, X., Mikusinski, P., & Taylor, M. (2002). Some integration-by-parts formulas involving 2-copulas. *Distributions with Given Marginals and Statistical Modelling*, Kluwer, Dordrecht, 153–159.
- Longin, F., & Solnik, B. (2001). Extreme correlation of international equity markets. *The Journal of Finance*, 56(2), 649–676.
- Mainik, G., & Schaanning, E. (2014). On dependence consistency of covar and some other systemic risk measures. *Statistics & Risk Modeling*, 31(1), 49–77.
- McNeil, A. J., Frey, R., & Embrechts, P. (2015). *Quantitative risk management: Concepts, techniques and tools*. Princeton university press.
- McNeil, F., Embrechts, P., Lindskog, A., & Rachev, S. (2003). Modelling dependence with copulas and applications to risk management. *Handbook of Heavy Tailed Distributions in Finance*, edited by Rachev ST, Elsevier/North-Holland, Amsterdam, 329–384.
- Mohanty, S., Nandha, M., & Bota, G. (2010). Oil shocks and stock returns: The case of the central and eastern european (cee) oil and gas sectors. *Emerging Markets Review*, 11(4), 358–372.
- Nelsen, R. (2006). An introduction to copulas, 2nd. New York: SpringerScience Business Media.
- Nelson, D. B. (1991). Conditional heteroskedasticity in asset returns: A new approach. *Econometrica: Journal of the Econometric Society*, 347–370.

Serfling, R. J. (2009). *Approximation theorems of mathematical statistics*, volume 162. John Wiley & Sons.

Siburg, K. F., Stoimenov, P., & Weiß, G. N. (2015). Forecasting portfolio-value-at-risk with nonparametric lower tail dependence estimates. *Journal of Banking & Finance*, 54, 129–140.

Sklar, M. (1959). Fonctions de repartition an dimensions et leurs marges. *Publ. Inst. Statist. Univ. Paris*, 8, 229–231.

