

MASTERARBEIT / MASTER'S THESIS

Titel der Masterarbeit / Title of the Master's Thesis

Value-at-Risk Estimation Using Time-Varying Copulas

verfasst von / submitted by

Ana Petrovič

angestrebter akademischer Grad / in partial fulfilment of the
requirements for the degree of

Master of Science (M.Sc.)

Wien, Oktober 2016 / Vienna, October 2016

Studienkennzahl lt. Studienblatt /

Degree programme code as it

appears on the student record sheet:

A 066 920

Studienrichtung lt. Studienblatt /

Degree programme as it appears

on the student record sheet:

Quantitative Economics, Management and Finance

Betreut von / Supervisor:

Prof. Dr. Nikolaus Hautsch

Abstract

One of the most popular and universal measures of financial risk, where the latter significantly influences portfolio valuation and financial decision making, is Value-at-Risk, a measure that quantifies possible portfolio losses due to market movements. Being defined as a quantile of the portfolio return distribution, its estimation requires appropriate specification of marginal distributions of the underlying assets as well as adequate modeling of the dependence between them. The necessary flexibility in this respect is provided by copula functions whose popularity has significantly increased in the last years, leading them to become an elementary tool in modeling multivariate distributions. Value-at-Risk forecasting is enabled by extending the copula theory to the conditional case, an extension that was motivated by empirical evidence of time-varying dependence between financial assets. In this thesis we consider the estimation of Value-at-Risk of a portfolio composed of two exchange rates by modeling the dependence between the exchange rates using a time-varying copula and analyse whether allowing the copula parameter to move according to an evolution equation provides better Value-at-Risk estimates.

Zusammenfassung

Eines der beliebtesten und weltweit gängigsten Maße für Finanzrisiken - und damit für die Bewertung von Portfolios und Investmententscheidungen von großer Bedeutung - ist der Value-at-Risk. Dieses Maß stellt eine Methode zur Quantifizierung von potentiellen Verlusten eines Portfolios auf Grund von Bewegungen im Finanzmarkt dar. Der Value-at-Risk ist als Quantil der Verteilung von Portfoliogewinnen und Portfolioverlusten definiert, wobei die Schätzung einer passenden Modellierung der Marginalverteilung der zugrundeliegenden Finanzinstrumente als auch deren Abhängigkeit untereinander bedarf. Die notwendige Flexibilität in diesem Zusammenhang wird durch Kopula-Funktionen gewährleistet, dessen Popularität über die vergangenen Jahre deutlich zugenommen und sich so zu einem wichtigen Werkzeug zur Modellierung multivariater Verteilungen entwickelt hat. Die Prognose von Value-at-Risk-Kennzahlen wird durch die Erweiterung der Kopula-Theorie auf bedingte Fälle ermöglicht. Diese Erweiterung war motiviert durch den empirischen Nachweis der Zeit-variierten Abhängigkeit von Kapitalanlagen. In dieser Thesis wird der Value-at-Risk eines Portfolios aus zwei Wechselkursen für Währungen durch die Modellierung der Abhängigkeit der Wechselkurse mit Hilfe von Zeit-variierten Kopulas geschätzt. Dabei soll untersucht werden ob die Veränderung der Kopula-Parameter durch eine Entwicklungsfunktion zu einer besseren Schätzung des Value-at-Risks führt.

Contents

1	Introduction	1
2	Value-at-Risk	5
2.1	Value-at-Risk Estimation	7
2.1.1	Historical Simulation	7
2.1.2	Variance Covariance Methods	8
2.1.3	Monte Carlo Simulation Methods	10
2.2	Value-at-Risk and Copulas	12
3	The (Un)conditional Copula	14
3.1	Copula Function	17
3.2	Families of Copulas	19
3.2.1	Elliptical Copulas	20
3.2.2	Archimedean Copulas	21
3.3	Estimation of Copulas	22
4	The Conditional Dependence between the Euro and the Pound	26
4.1	The Models for the Marginal Distributions	28
4.2	Modeling the Dependence between Exchange Rates	32
5	Back-testing	35
5.1	Kupiec's Proportion of Failures Test	36
5.2	Independence Tests	37
5.3	Joint Tests of Unconditional Coverage and Independence	39
5.4	Back-testing results	40
5.4.1	In-sample back-testing	40
5.4.2	Out-of-sample back-testing	45
6	Conclusion	51

List of Tables

1	Some families of Archimedean Copulas	22
2	Summary Statistics	27
3	Results for the Marginal Distributions	29
4	Influence of the "Other" Variable in the Mean and Variance Models .	30
5	p -values of Ljung-Box test	32
6	Results of Anderson Darling Goodnes of Fit Test	32
7	Estimated Parameters of Time-Varying Normal Copula	33
8	Contingency Table for Markov Independence Test	38
9	In-Sample Back-testing Results	41
10	Out-of-Sample Back-testing Results	47

List of Figures

1	Illustration of two distributions with identical marginal distributions and correlation	2
2	Contour plots of Normal and Student's t Copulas	21
3	Contour plots of Gumbel, Clayton and Frank Copulas	22
4	Daily and absolute log-differences of the exchange rates	26
5	Autocorrelation of the log-differences of the exchange rates	28
6	Characteristics of standardized residuals	31
7	Empirical distribution of the transformed series u_t and v_t	31
8	Estimated Parameters of the Normal Copula	34
9	In-Sample Value-at-Risk estimation at $\alpha = 0.10$	42
10	In-Sample Value-at-Risk estimation at $\alpha = 0.05$	43
11	In-Sample Value-at-Risk estimation at $\alpha = 0.01$	44
12	Out-of-Sample Value-at-Risk estimation at $\alpha = 0.10$	48
13	Out-of-Sample Value-at-Risk estimation at $\alpha = 0.05$	49
14	Out-of-Sample Value-at-Risk estimation at $\alpha = 0.01$	50

1 Introduction

Financial risk can be defined as a quantifiable probability of loss or gain that is lower than expected, caused by an event or action that might adversely affect one's ability to achieve given objectives (McNeil et al. (2015)). Risk, being closely connected to uncertainty and randomness, has a strong impact on the valuation of portfolios and financial decision making. When measuring and managing risk, we are in fact modeling the distribution function $F_X(x) = P(X \leq x)$ of a risky asset X , which describes the probability with which the value of that risky asset is smaller than some given number x . In this sense, risk measurement is a statistical issue and the challenge of risk management lies in addressing unexpected and extreme events.

One of the most popular and widely used risk measures is Value-at-Risk. Born in 1993 when the Weatherstone 4.15 report asked for a one-day one-page summary of bank's market risk assessment, it has become an industry-wide standard (McNeil et al. (2015)). Although forecasting Value-at-Risk is of high importance in applied risk management (Berger and Missonig (2014)) and there has been progress in adequate risk modeling and accurate portfolio risk prediction, there is still a strong need to appropriately capture the characteristics of the individual assets in a portfolio and improve the methods of modeling the interdependence between them. This need is directly implied by the definition of Value-at-Risk; determined as a quantile of the return distribution of a portfolio, correct estimation of the former demands adequate modeling of the latter. As the latter depends on the univariate distributions of assets within the portfolio as well as the dependence between them, adequate choice of the marginal distributions and adequate modeling of the dependence structure are imperative for reliable Value-at-Risk estimates.

In financial theory, a central and universal measure of dependence between financial instruments is correlation. Models such as the Capital Asset Pricing Model and the Arbitrage Pricing Theory use correlation in combination with the assumption that the assets' returns are multivariate normally distributed to determine an optimal portfolio. Embrechts et al. (2002) point out however, that correlation is also

one of the most misunderstood concepts. It is often used to describe any kind of dependency between random variables, although this is only true when considering spherical and elliptical distributions. As distributions of random variables we encounter in the world of finance are rarely from this class, the correlation might not always adequately capture the dependence between the underlying assets. Any conclusions based on it might thus be unreliable and could lead to poor decision making.

A simple example demonstrating the importance of appropriate modeling of the dependence structure between random variables is illustrated by Embrechts et al. (2002), where two models with identical marginal distributions and identical correlation have different dependence structure. Examining the point clouds generated by $\text{Gamma}(3, 1)$ marginal distributions and correlation $\rho = 0.5$ depicted in Figure 1, we can determine that although the marginal distributions and correlation in both models are the same, the models do not exhibit the same dependence structure, as extreme losses in the model presented on the right hand side occur together.

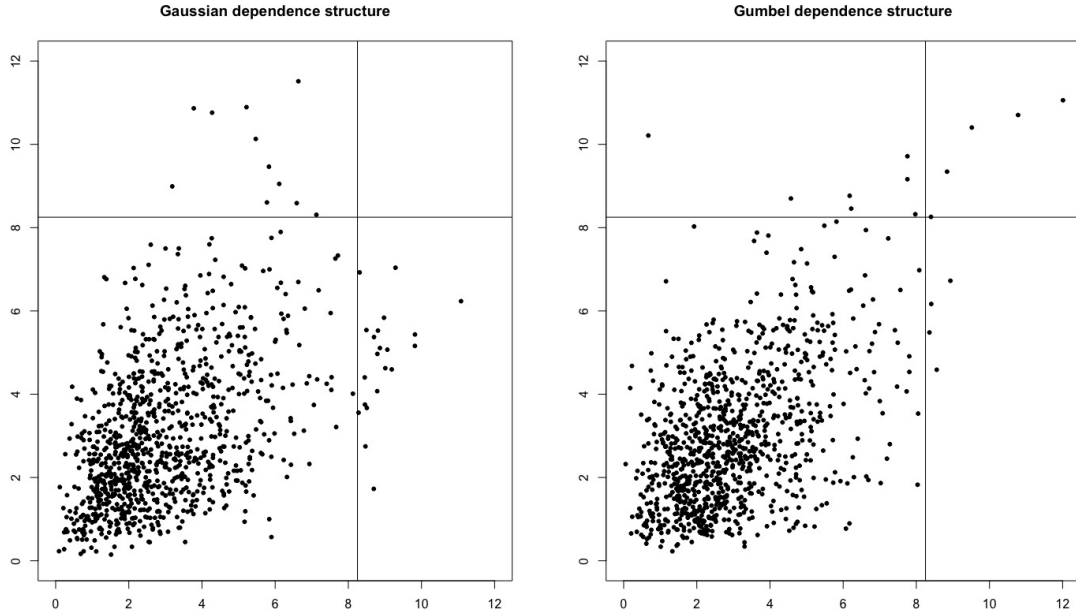


Figure 1: Illustration of two distributions with identical marginal distributions and correlation

Patton (2009) uses a similar approach to illustrate the potential of copula theory, a theory dating back to Sklar (1959) that has received a lot of attention in modeling of univariate and multivariate financial time series in the last decades. Providing plots of a number of bivariate standard normally distributed densities with linear correlation 0.5 and several copula functions describing their dependence structure over a set of parametric copulas, he, similarly as Embrechts et al. (2002), shows that models with the same marginal distributions and correlation may exhibit different dependence structures. In addition, this example indicates the variety of joint densities that may be constructed using copulas that is the largest advantage that copula theory brings the researcher: the flexibility it allows in specifying the model for the joint distribution. In such an application, models for marginal distributions may be specified without any restrictions and separately from the dependence structure that is represented by copula and links the marginal distributions to a joint distribution.

Considering the results provided in the literature exploring how the dependence between financial assets moves with the markets, see for example Andersen et al. (2006) and Bauwens et al. (2006), it seems natural to extend copula theory to the conditional case and allow for time variation. Time variation in copula theory can be included in two different ways. One approach is to consider a time-varying copula parameter that is defined as a function of previous observations, while the copula function remains unchanged through out the sample. On the other hand, one can use a regime switching model, where the time variation is captured in the changing functional form of the copula. Such an extension permits us to use copula functions for modeling of dynamical structures and enables construction of conditional dependence models. Furthermore, in case of Value-at-Risk estimation, a time-varying copula approach is the one allowing to forecast Value-at-Risk. Results of papers such as Jondeau and Rockinger (2006), Ausin and Lopes (2010) and Christoffersen et al. (2012) to list just a few, testify to the value added when considering a time-varying copula model.

This thesis contributes to the existing literature on the use of copulas in financial applications by investigating whether allowing for time variation in copula specifica-

tion leads to improved estimates of Value-at-Risk. We do so by applying the copula approach to the problem of Value-at-Risk estimation of a portfolio of two exchange rates and allowing the copula parameter to vary through time based on an evolution equation, while keeping the functional form of the copula fixed over the entire sample. The conditional copula with time-varying parameters is compared to the plain rolling window approach. Comparing both approaches in an in- and out-of-sample back-tests we determine that while the time-varying approach performs better in the in-sample back-test, it has a lower predictive ability than the plain rolling window approach in case of the out-of-sample back-test. The thesis is structured as follows. Section 2 defines the notion of Value-at-Risk and its estimation methods. Section 3 introduces copula theory and defines both unconditional as well as conditional copula functions. Empirical application of the study starts in Section 4, followed by Section 5 where back-testing results are documented. A conclusion is provided in Section 6.

2 Value-at-Risk

Value-at-Risk is a market risk metric - a measure of possible portfolio losses due to market movements (Linsmeier and Pearson (2000)). It is considered as major determinant of the required risk capital with several attractive features, that quantifies risk in a single, easily understandable number and can be applied to any kind of financial asset, see Alexander (2009) for more details. In addition to Value-at-Risk, two well known risk metrics are volatility and correlation, however, they appropriately summarize uncertainty in the future value of a portfolio only when asset returns are multivariate normally or Student's t distributed.

Value-at-Risk has since its introduction become one of the most used risk management tools. J.P. Morgan's decision in 1994 to make their Riskmetrics database freely available to all market participants, climate created by derivatives disaster such as Procter and Gamble, Kidder Peabody, Orange County, and Barings, as well as the decision of central bank regulators to use Value-at-Risk in calculating a bank's required capital, are just a couple of reasons for the rise of its popularity (Hull and White (1998)).

Alexander (2009) identifies as two basic parameters determining Value-at-Risk the significance level α , and the risk horizon h over which Value-at-Risk is measured. The former is often set by regulators, whereas the risk horizon depends on the liquidity of the assets in the portfolio. For example, banks using internal Value-at-Risk models to assess their market risk capital requirement, are required to measure Value-at-Risk at the 1% significance level over 10 days. In case there are no regulations imposed, the risk horizon should reflect the period of time during which we are exposed to the position, whereas the choice of the significance level depends on the user and the way he wants to interpret the estimated Value-at-Risk. Linsmeier and Pearson (2000) emphasize the importance of these two parameters when interpreting Value-at-Risk estimates. Their absence makes comparison of Value-at-Risk estimates obsolete, as different choices of significance level and risk horizon imply different Value-at-Risk estimates. They estimate, that the loss suffered with 1%

probability can be up to 1.41 times larger than that experienced with 5% probability. The choice of holding period can have an even larger impact in comparison of different VaR estimates.

Linsmeier and Pearson (2000) describe the Value-at-Risk as the loss that is expected to be exceeded with a probability of $\alpha\%$ during the next h -day holding period. Stated differently, Value-at-Risk is the loss that we can expect to be exceeded during $\alpha\%$ of the h -day holding period. Considering a 5% daily Value-at-Risk, we are therefore talking about a loss level that is likely to occur 5% of the time, when the current portfolio is held for 1 day. It also means that we are 95% confident that the Value-at-Risk will not be exceeded when the portfolio is held over 1 day, and similarly, that we expect that a loss as large or larger will occur one day in every 20 days. The estimation of Value-at-Risk given the significance level α and risk horizon h is based on the assumption that we hold our portfolio static over the chosen risk horizon. This assumption might not be reasonable when we are considering longer risk horizons (Linsmeier and Pearson (2000)), and in addition implies that we are assessing the uncertainty of the unrealized profit and loss, $P\&L$, whereas the realized $P\&L$ takes into consideration the adjustment of positions that are made in reality (Alexander (2009)).

The problem of estimating $\alpha\%$ h -day Value-at-Risk at time t is defined as the problem of finding the α quantile $x_{ht,\alpha}$ of the discounted h -day $P\&L$ distribution, such that

$$P(P_{t+h} - P_t < x_{ht,\alpha}) = \alpha,$$

where P_t denotes the value of the portfolio at time t . Then $VaR_{ht,\alpha} = -x_{ht,\alpha}$. We can estimate Value-at-Risk from a $P\&L$ distribution or from the return distribution. The latter approach is preferred, as returns are measured in relative terms and allow comparison at different moments in time, even when price levels are trending. Value-at-Risk estimated from a return distribution is expressed as a percentage of the portfolio's current value. Alexander (2009) serves as a source for the mathematical definition of the Value-at-Risk stated below.

Definition 1. *Let the discounted h -day return on a portfolio be the random variable*

$$X_{ht} = \frac{P_{t+h} - P_t}{P_t}.$$

Then, there exists the α quantile, $x_{ht,\alpha}$, of its distribution

$$P(X_{ht} < x_{ht,\alpha}) = \alpha,$$

such that the current estimate of the $\alpha\%$ h -day Value-at-Risk at time t is

$$VaR_{ht,\alpha} = \begin{cases} -x_{ht,\alpha} & \text{as a percentage of the portfolio value } P_t, \\ -x_{ht,\alpha}P_t & \text{when expressed in value terms.} \end{cases}$$

2.1 Value-at-Risk Estimation

Even though Value-at-Risk is a central concept in risk management as well as a simple measure, it is not easily estimated; considerable effort is needed for construction of the risk factor return distribution. There are three basic approaches to Value-at-Risk estimation, that differ in a variety of aspects, such as their ability to capture the risks of options and option-like instruments, ease of implementation and explanation to senior management, reliability of results, etc., as well as how they construct the risk factor return distribution. Risk manager determines the best model based on the aspects he values the most. For a more detailed review of Value-at-Risk methods see Komunjer et al. (2011).

2.1.1 Historical Simulation

Historical simulation is a simple approach to Value-at-Risk estimation, that does not require a lot of assumptions about the statistical distribution of the underlying assets in the portfolio. To estimate Value-at-Risk using this approach we first collect a sample of historical asset prices or log returns of the underlying assets of the current portfolio and then simulate the discounted $P\&L$ distribution of the portfolio using the weights of the assets. This discounted $P\&L$ distribution is in fact a hypothetical one, as we have estimated it using the current portfolio and realization of actual

past prices of the assets, whereas we did not hold the current portfolio in each of the previous periods.

This approach does not rely on any parametric model assumptions. Nevertheless, an important, and to some extent restricting, is the assumption that the historical distribution of asset prices or log returns is a good approximation of the distribution of future asset prices or log returns. In other words, we assume that the history will repeat itself, which is an assumption that does not necessarily hold. Another disadvantage of this approach is the fact that it requires a large set of historical data in order to be able to estimate Value-at-Risk at high confidence levels. Smaller sets of data might not contain enough observations to estimate quantiles in extreme lower tails. This means that we need to collect data on daily basis for many years, a procedure that is sometimes practically unfeasible and time consuming. If we wanted to estimate Value-at-Risk over longer risk horizons, we would need to scale up the estimated daily Value-at-Risk; scaling of estimated Value-at-Risk is however not a straight forward procedure (Alexander (2009)).

2.1.2 Variance Covariance Methods

Variance covariance approach is the most well known approach of Value-at-Risk estimation, first time fully explained in detail by Morgan (1996) RiskMetrics Technical Document. The approach is based on the assumption that the risk factors' log returns are normally distributed, and that their joint distribution is multivariate normal. This implies that the returns' covariance matrix contains all information about the dependency between the risk factor returns, and we can estimate the Value-at-Risk by estimating the variances and covariances of the underlying assets' log returns, and the sensitivity of the portfolio to them.

To illustrate the approach, let us consider Value-at-Risk estimation of a portfolio with i.i.d. normally distributed returns as described in Alexander (2009). Denoting the return of the portfolio $X \stackrel{\text{i.i.d.}}{\sim} N(\mu, \sigma^2)$, we will derive the formula for the α quantile of the return distribution, x_α , and determine the $\alpha\%$ Value-at-Risk as minus this α quantile. Making use of the mathematical properties of the normal

distribution, we have

$$P(X < x_\alpha) = P\left(\frac{X - \mu}{\sigma} < \frac{x_\alpha - \mu}{\sigma}\right) = P\left(Z < \frac{x_\alpha - \mu}{\sigma}\right), \quad (2.1)$$

where $Z \sim N(0, 1)$. Considering that $P(X < x_\alpha) = \alpha$, we have

$$P\left(Z < \frac{x_\alpha - \mu}{\sigma}\right) = \alpha. \quad (2.2)$$

We know that by definition it holds that $P(Z < \Phi^{-1}(\alpha)) = \alpha$, therefore

$$\frac{x_\alpha - \mu}{\sigma} = \Phi^{-1}(\alpha), \quad (2.3)$$

where Φ is the standard normal distribution function. But as $x_\alpha = -VaR_\alpha$ by definition, and symmetry of the standard normal distribution, $\Phi^{-1}(\alpha) = -\Phi^{-1}(1 - \alpha)$, holds, we can write (2.3) as

$$VaR_{h,\alpha} = \Phi^{-1}(1 - \alpha)\sigma - \mu, \quad (2.4)$$

that is the analytical formula for the Value-at-Risk of a portfolio with i.i.d. normal returns, expressed in percentage terms of the portfolio value. If we would like to obtain the Value-at-Risk in value terms, we multiply the estimated percentage Value-at-Risk with the current value of the portfolio, P_t :

$$VaR_{ht,\alpha} = \left(\Phi^{-1}(1 - \alpha)\sigma_h - \mu_h\right)P_t. \quad (2.5)$$

Alexander (2009) has shown that it is possible to derive a formula for Value-at-Risk also when the risk factor log returns are distributed with Student's t distribution, or a mixture of Student's t and normal distribution. In the latter case, the formula delivers Value-at-Risk in an implicit rather than explicit way, requiring in addition also an application of a numerical method. In addition, Alexander (2009) also states that it is not necessary to assume that each risk factor return follows an i.i.d. process, as it is possible to find a simple scaling rule for linear Value-at-Risk when the risk factor returns are autocorrelated, under the assumption that there is no time-varying volatility. On the other hand, variance covariance approach of Value-at-Risk estimation is not a good choice when our portfolio's returns are not linear functions of portfolio assets' returns (Linsmeier and Pearson (2000), Alexander (2009)).

2.1.3 Monte Carlo Simulation Methods

When using the Monte Carlo Simulation method, Value-at-Risk of a portfolio is estimated by simulating a large number of hypothetical changes in the market factors out of a statistical distribution that is believed to describe the changes in the underlying assets of the portfolio adequately. The choice of this distribution is most of the time based on the changes of underlying market risk factors that have been observed in the past. Each simulation gives a possible value of the portfolio return for the respective time point, and given that we have simulated a large enough number of possible outcomes, we can assume that we have reduced the sampling error in such a way that the simulated distribution of the portfolio values converges to the portfolio's unknown true distribution. This way we can use the simulated distribution to determine the Value-at-Risk of the portfolio.

Linsmeier and Pearson (2000) summarize the Value-at-Risk estimation using the Monte Carlo Simulation in the following steps.

1. Identify the basic market factors of the portfolio, and obtain the formula describing its value.
2. Determine or assume a specific distribution for changes in the basic market factors as well as estimate the parameters of that distribution.
3. Use a pseudo-random generator to generate a large number of hypothetical values of changes in the market factors, that are further on used to calculate hypothetical portfolio values. From each of those, subtract the actual portfolio value at time t and obtain hypothetical daily profits and losses.
4. Order the estimated profits and losses from the largest profit to the largest loss, and determine the Value-at-Risk as the loss that is equaled or exceeded 5% of the time.

Alexander (2009) points out, that, assuming we model the distribution of market risk factors with a normal one, we are considering a normal Monte Carlo Value-at-Risk model, that makes the same assumption as the variance covariance approach.

The two models are essentially the same, they differ only in the way we obtain the evolution of the risk factors; in case of Monte Carlo simulation we simulate it, whereas in the case of variance covariance approach, we obtain it analytically. The variance covariance approach is therefore precise although based on assumptions that are not likely to hold, while the Monte Carlo simulation assuming the normal distribution of log returns might be subject to simulation error. This would imply that any difference that might exist between the two models is caused by the number of simulations used in the Monte Carlo approach not being big enough.

The choice of the distribution underlying the changes in market risk factors is one of the main advantages of the Monte Carlo approach compared to the other two approaches we have discussed. Furthermore, we can use it to estimate Value-at-Risk of a portfolio that is a non-linear function of its underlying assets. To avoid wrong estimates of Value-at-Risk, we should however assume correct pricing models and underlying stochastic processes.

At this point, let us summarize the main advantages and disadvantages, as well as differences between the three Value-at-Risk estimation models we have introduced. Linsmeier and Pearson (2000) provide an extensive overview how the three approaches differ in their ability to capture the risks of options and option-like instruments, how easily they can be implemented, how flexible they are in analyzing the effect of changes in the assumptions, how reliable the results are, etc. The best approach is therefore determined by the user himself, through the dimension of the model he finds the most important.

The main aspect in which the approaches differ one from another is in the assumptions regarding the distribution of the log returns they make. Variance covariance and Monte Carlo approach assume we know the distribution of the risk factors, where in case of variance covariance approach we assume they are normally, Student's t distributed or distributed with a mixture of the two. On the other hand, Monte Carlo simulation allows us to model the log returns with any distribution we deem appropriate. In case of historical approach, we base our estimation on the

empirical distribution, therefore on variations and dependencies between risk factors that have been experienced in the past. The fact that the historical distribution does not make any parametric assumptions on the distribution of the underlying market factors is also a drawback of the approach, as market conditions might have changed through the history. Furthermore, both Monte Carlo and historical approach require a large amount of data. Last but not least, the three approaches differ in the way they consider nonlinear portfolios; only Monte Carlo and historical approach may be used for estimation of Value-at-Risk of a portfolio that is a non-linear function of its underlying assets.

2.2 Value-at-Risk and Copulas

In many applications, log returns have been assumed to follow normal distribution, an assumption that has rarely been proven to be an appropriate one. Longin and Solnik (2001) and Ang and Chen (2002) have shown that the correlation among asset returns is higher during the times when the markets are volatile and when they are experiencing downturns. Proof against the assumption of multivariate normality was also recorded in Erb et al. (1994) and Bae et al. (2003). In addition, Mills et al. (1927) confirms that many economic variables exhibit excess kurtosis, skewness and asymmetric dependence. This implies that univariate distributions of many economic variables are not normal and also deviate from multivariate normality in portfolio settings. Should we fail to correctly specify the distribution of log returns, our Value-at-Risk model will be misspecified and in turn the financial institution too conservative or too risk loving.

Recalling Definition 1, Value-at-Risk is defined as a quantile of the portfolio return distribution. The latter depends on the univariate distributions of the underlying assets in a portfolio as well as on the dependence between them. In terms of modeling the distribution of underlying assets in a portfolio, flexibility is of great importance, as highlighted by Berger (2013) who emphasizes that the dependence model should not pose any restrictions on the distribution of underlying assets in a portfolio. In this sense, the copula approach to Value-at-Risk estimation, providing exactly such

flexibility, is favored over the historical and variance covariance approaches; Monte Carlo simulation approach allows free choice of the distribution of underlying assets and is also in our study combined with the copula approach to estimate Value-at-Risk of the portfolio.

A copula is a function that forms a multivariate distribution function by linking univariate distribution functions without posing any restrictions on them. By definition, as marginal distribution contains all of the univariate information on the individual assets and the joint distribution contains all univariate as well as multivariate information, it contains all information on the dependence between the underlying assets. The theory of copulas is introduced in more detail in next section.

3 The (Un)conditional Copula

Copula theory is a powerful tool in modeling multivariate distributions, as it allows decomposition of any n -dimensional joint distribution function in its n marginal distributions and a copula function - a multivariate distribution function whose margins are uniform on the interval $(0, 1)$. To put it differently, it constructs a multivariate joint distribution function by combining the marginal distributions and the dependence between the variables (Palaro and Hotta (2006)). Copula theory allows the marginal distribution models to be specified separately from the dependence structure represented by copula. Furthermore, marginal distributions may be chosen independently from each other and the choice of them does not put any restrictions on the choice of copula.

The theory of copulas, dating back to Sklar (1959), has seen a large number of applications in various fields. The use of copulas varies from applications in multivariate time series analysis where they characterize the cross-sectional dependence between individual time series, to applications in univariate time series analysis, where they characterize the dependence between a sequence of observations of a scalar time series. In terms of areas of applications, risk management was one of the first fields in economics and finance, where the importance of Value-at-Risk estimation calls for flexible models for modeling the dependence between assets in the portfolio.

One of the first and most influential papers on the use of copulas in finance is that of Embrechts et al. (2002), emphasizing the correct understanding of correlation and consequently the gain of copula representation of dependence for a random vector. Statistical foundation of copula theory is provided by Joe (1997) and Nelsen (2006), whereas the overview of copulas in risk management can be found in McNeil et al. (2015) as well as Denuit et al. (2006) and Alexander (2009). Cherubini et al. (2004) contribute with an introduction based on mathematical finance. In the field of Value-at-Risk estimation, Cherubini and Luciano (2001) use copula theory to evaluate Value-at-Risk of a portfolio of two stock market indices and assume joint distribution functions others than the normal one. Embrechts et al. (2003) use copula theory

to demonstrate the techniques of construction of optimal bounds for risk measures of functions of dependent risks in the case of Value-at-Risk. Embrechts and Höing (2006) use copulas to study the problem of the dependence scenario yielding the worst possible Value-at-Risk at a given confidence level α for dimensions higher than two.

Among others, works of Longin and Solnik (1995), Andersen et al. (2006) and Bauwens et al. (2006), confirming that dependence between time series we encounter in finance is not constant through time, provide the basis for inclusion of time variation into copula theory. Moreover, the topic of correlation and how it moves with the markets has been addressed in a vast number of applications; interdependence of returns and volatilities in stock markets considering a multivariate GARCH framework has been studied in Hamao et al. (1990), Susmel and Engle (1994) as well as Bekaert and Harvey (1995). GARCH setting was expanded to consider time-varying covariances and correlations in Kroner and Ng (1998), Engle and Sheppard (2001), Engle (2002) and Tse and Tsui (2002).

One of the first ones to consider copula theory for conditional case and introduce time variation was Patton (2006b). Studying the dependence between the Deutsche mark and the Japanese yen, he has proven that the correlation between the Deutsche mark-U.S. dollar and the Japanese yen-U.S. dollar exchange rates is higher when Deutsche mark and Japanese yen are depreciating against the U.S. dollar than when they are appreciating. Patton (2006b)'s paper and that of for example Jondeau and Rockinger (2006), Ausin and Lopes (2010), Christoffersen et al. (2012), and Palaro and Hotta (2006), studied models where time variation is introduced by allowing the parameters of a copula to vary through time, in a manner corresponding to a GARCH model for conditional variance, while copula form is fixed through time. Examining the daily returns of stock markets, Jondeau and Rockinger (2006) show that for European market pairs conditional dependency depends on past realizations and is also more extensively affected when returns move in the same rather than opposite direction. Ausin and Lopes (2010) use time-variation in copulas to identify contagion effects during times of financial instability, whereas Christoffersen

et al. (2012) study the dynamic evolution of dependence and tail dependence in developed and emerging markets. Proposing a dynamic asymmetric copula (DAC) model, allowing for asymmetric and dynamic tail dependence, they show that international nonlinear dependence between stock markets has significantly trended up for both markets, developed and emerging. Palaro and Hotta (2006) use the conditional copula to estimate the Value-at-Risk of a portfolio consisting of two stock indices and analyze the performance of conditional copula to the results obtained using traditional approaches to Value-at-Risk estimation. Using different copulas and marginal distributions for their analysis they have concluded that time-varying copulas provide the best results and produced reliable Value-at-Risk estimates.

Applications such as Rodriguez (2007), Chollete (2005), Garcia and Tsafack (2008) and Okimoto (2008) study regime switching models. Rodriguez (2007) investigated the dependence structure of Mexican and Brazilian markets by modeling the degree of dependence using three copulas, Normal, Clayton and Gumbel, considering different periods, and found that the dependence structure of Mexican and Brazilian markets was stronger during the financial crisis of 2008 and intensified after. To answer the question how common and persistent are turbulent periods, Chollete (2005) used a dynamic dependence framework, analyzing a Markov model of copulas. Among other results, he established that turbulent regimes are more common and persistent, and conditional copula is dominated by the unconditional copula. Lee and Long (2005) constructed flexible distributions for the residuals from a multivariate GARCH model using copulas by allowing the GARCH model to capture the time-varying correlation and the copula to capture the dependence that might still remain between conditionally uncorrelated standardized residuals.

In the remainder of this chapter we provide the theoretical foundation of the copula theory based on Palaro and Hotta (2006). We define the (un)conditional copula, introduce some of the existing families of copulas as well as discuss the estimation of copulas.

3.1 Copula Function

The conditional copula function is a distribution function whose marginal distributions are uniform distributions. For simplicity, we consider the bivariate case, however, all concepts presented are easily translated to consider a multivariate application.

Definition 2. *Palaro and Hotta (2006) A 2-dimensional copula function is a function C , whose domain is $[0, 1]^2$ and whose range is $[0, 1]$ with the following properties:*

1. $C(x) = 0$ for all $x \in [0, 1]^2$ when at least one element of x is 0;
2. $C(x_1, 1) = C(1, x_2) = 1$ for all $(x_1, x_2) \in [0, 1]^2$;
3. for all $(a_1, a_2), (b_1, b_2) \in [0, 1]^2$ with $a_1 \leq b_1$ and $a_2 \leq b_2$, we have:

$$V_c([a, b]) = C(a_2, b_2) - C(a_1, b_2) - C(a_2, b_1) + C(a_1, b_1) \geq 0$$

Definition 3. *Palaro and Hotta (2006) The copula function C is a copula for the random vector $\mathbf{X} = (X_1, X_2)^t$, if it is the joint distribution function of the random vector $\mathbf{U} = (U_1, U_2)^t$ where $U_i = F_i(X_i)$, and F_i are the marginal distribution functions of $x_i, i = 1, 2$.*

Theory of copulas dates back to Sklar (1959). The inverse of his Theorem we state below tells us that we may define a valid multivariate distribution by linking any group of n univariate distributions with any copula, and is as such a key result in the theory of copulas. In this sense, Sklar (1959) Theorem enlarges the set of existing multivariate distribution functions and allows construction of more sophisticated models than the much smaller set of available multivariate distributions would imply.

Theorem 1 (Sklar (1959) Theorem for Unconditional Case). *Palaro and Hotta (2006) Let H be a 2-dimensional joint distribution function with marginal distributions F_1, F_2 . Then there exists a copula C such that for all $x \in \mathbb{R}^n$,*

$$H(x_1, x_2) = C(F_1(x_1), F_2(x_2)). \quad (3.6)$$

If F_1, F_2 are continuous, then C is unique; otherwise, C is uniquely determined on $\text{Ran}(F_1) \times \text{Ran}(F_2)$. Conversely, if C is a copula and F_1, F_2 are distribution

functions, then the function H defined by (3.6) is a joint distribution function with margins F_1 and F_2 .

Patton (2006b) extended and proved the validity of the Sklar (1959) theorem for the conditional case. He emphasizes that in this case the conditioning variables, $\mathbf{W}$, on which we condition the marginal models must be the same for both marginal models. Should that not be the case, the joint distribution of random variables X and Y conditioned on variables $\mathbf{W}$, $F_{XY|\mathbf{W}}$, fails to satisfy the conditions to be the joint conditional distribution function.

Theorem 2 (Sklar (1959) Theorem for Conditional Case). *Patton (2006b) Let $F_{X_1|\mathbf{W}}(\cdot|\mathbf{w})$ be the conditional distribution of $X_1|\mathbf{W} = \mathbf{w}$, $F_{X_2|\mathbf{W}}(\cdot|\mathbf{w})$ be the conditional distribution of $X_2|\mathbf{W} = \mathbf{w}$, $F_{X_1X_2|\mathbf{W}}(\cdot|\mathbf{w})$ be the joint conditional distribution of $(X_1, X_2)|\mathbf{W} = \mathbf{w}$, and $\mathcal{W}$ be the support of $\mathbf{W}$. Assume that $F_{X_1|\mathbf{W}}(\cdot|\mathbf{w})$ and $F_{X_2|\mathbf{W}}(\cdot|\mathbf{w})$ are continuous in x_1 and x_2 for all $w \in \mathcal{W}$. Then there exists a unique conditional copula $C(\cdot|w)$ such that*

$$F_{X_1X_2|\mathbf{W}}(x_1, x_2|w) = C(F_{X_1|\mathbf{W}}(x|\mathbf{w}), F_{X_2|\mathbf{W}}(y|\mathbf{w})|w), \quad (3.7)$$

$$\forall (x_1, x_2) \in \mathbb{R} \times \mathbb{R} \quad \text{and each } w \in \mathcal{W}. \quad (3.8)$$

Conversely, if we let $F_{X_1|\mathbf{W}}(\cdot|\mathbf{w})$ be the conditional distribution of $X_1|\mathbf{W} = \mathbf{w}$, $F_{X_2|\mathbf{W}}(\cdot|\mathbf{w})$ be the conditional distribution of $X_2|\mathbf{W} = \mathbf{w}$, and $\{C(\cdot|\mathbf{w})\}$ be a family of conditional copulas that is measurable in $\mathbf{w}$, then the function $F_{X_1X_2|\mathbf{W}}(\cdot|\mathbf{w})$ defined by (3.7) is a conditional bivariate distribution function with conditional marginal distributions $F_{X_1|\mathbf{W}}(\cdot|\mathbf{w})$ and $F_{X_2|\mathbf{W}}(\cdot|\mathbf{w})$.

Each of the marginal distribution functions $F_i, i = 1, 2$ of random variables $X_i, i = 1, 2$ depends only on the parameter ϑ_i , whereas the 2-dimensional copula $\{C_\theta, \theta \in \Omega\}$, where Ω is some parameter space, depends on the unknown vector of parameters θ . We denote the vector of parameters by $\vartheta = (\vartheta_1, \vartheta_2, \theta)$. Considering a sample of size T , $\{(x_{1,t}, x_{2,t})\}_{t=1}^T$, Sklar (1959) Theorem tells us that the parameter vector $\vartheta = (\vartheta_1, \vartheta_2, \theta)$ fully specifies the joint distribution function H (Palaro and Hotta

(2006)):

$$H(x_1, x_2) = C(F_1(x_1; \vartheta_1), F_2(x_2; \vartheta_2); \theta). \quad (3.9)$$

In the conditional case, we use the conditional version of the Sklar (1959) Theorem. Denoting with F_i the conditional distribution of $X_i|\mathbf{W}$, $i = 1, 2$, with H the absolutely continuous conditional distribution of $\mathbf{X}|\mathbf{W}$, $\mathbf{X} = (X_1, X_2)$ and with C the conditional copula function, we have

$$H(x_1, x_2|\mathbf{w}) = C(F_1(x_1|\mathbf{w}), F_2(x_2|\mathbf{w})|\mathbf{w}).$$

We obtain the density function of the joint distribution function H , h , by differentiating (3.9) with respect to all variables:

$$h(x_1, x_2) = c(F_1(x_1), F_2(x_2))f_1(x_1)f_2(x_2). \quad (3.10)$$

Here we denoted the density function of marginal distribution F_i with f_i , and the copula density with c , given by

$$c(u_1, u_2) = \frac{\partial^2 C(u_1, u_2)}{\partial u_1 \partial u_2}.$$

Correspondingly, the density function of the conditional joint distribution function $h(x_1, x_2|\mathbf{w})$ is given by

$$h(x_1, x_2|\mathbf{w}) = c(F_1(x_1|\mathbf{w}), F_2(x_2|\mathbf{w})|\mathbf{w})f_1(x_1|\mathbf{w})f_2(x_2|\mathbf{w}),$$

with $f_i(x_i|\mathbf{w})$ as the conditional density of $X_i|\mathbf{W} = \mathbf{w}$ and $c(u_1, u_2|\mathbf{w})$ conditional copula density function defined as

$$c(u_1, u_2|\mathbf{w}) = \frac{\partial^2 C(u_1, u_2|\mathbf{w})}{\partial u_1 \partial u_2}.$$

3.2 Families of Copulas

There are two popular families of copulas used in financial applications, elliptical and Archimedean copulas.

3.2.1 Elliptical Copulas

Elliptical copulas are derived from multivariate elliptical distributions. The most important copulas in this family are the Gaussian or normal copula, and the Student's t copula.

The Gaussian Copula

The Gaussian copula of a 2-dimensional standard normal distribution, with linear correlation coefficient ρ , is the distribution function of the random vector $(\Phi(X_1), \Phi(X_2))$, where Φ is the univariate standard normal distribution function and $X \sim N_2(0, \rho)$. Hence,

$$C_\rho^{Ga} = P(\Phi(X_1) \leq u_1, \Phi(X_2) \leq u_2) = \Phi_\rho^2(\Phi^{-1}(u_1), \Phi^{-1}(u_2)) \quad (3.11)$$

where Φ_ρ^2 is the distribution function of X .

The Student's t Copula

The Student's t copula of a 2-dimensional standard Student's t distribution with $\nu \geq 0$ degrees of freedom and linear correlation coefficient ρ , is the distribution of the random vector $(t_\nu(X_1), t_\nu(X_2))$, where X has a $t^2(0, \rho, \nu)$ distribution and t_ν is the univariate standard Student's t distribution function. Hence,

$$C_{\nu, \rho}^t = P(t_\nu(X_1) \leq u_1, t_\nu(X_2) \leq u_2) = t_{\nu, \rho}^2(t_\nu^{-1}(u_1), t_\nu^{-1}(u_2))$$

where $t_{\nu, \rho}^2$ is the distribution function of X .

Combining two Student's t distributions F_1 and F_2 with the (same) degrees of freedom ν with a Student's t copula defines the bivariate distribution function $H(x, y) = C(F_1(x), F_2(y))$ that is in fact the standardized bivariate Student's t distribution with $\mu = 0$, degrees of freedom ν and linear correlation coefficient ρ .

To illustrate the different dependence structures implied by different choice of elliptical copulas, we refer to Figure 2 where we have plotted contour plots of Gaussian and Student's t copula, both with standard normal marginal distributions and linear correlation coefficient 0.5.

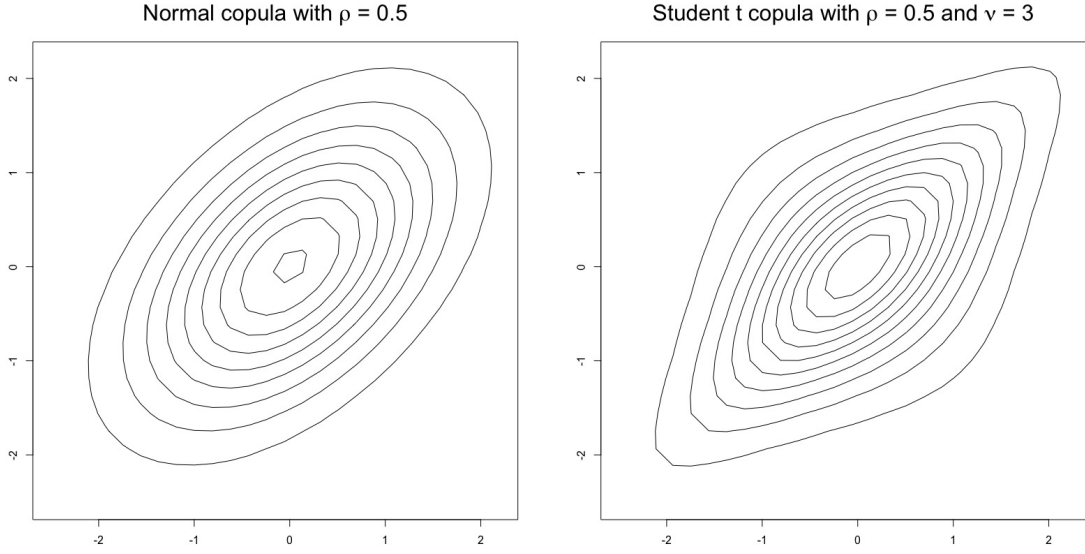


Figure 2: Contour plots of Normal and Student's t Copulas

3.2.2 Archimedean Copulas

Archimedean copulas present an important class of copulas whose members pose many nice properties, can be easily constructed and to which a large number of various families of copulas belongs, and has therefore been applied in many settings. To introduce Archimedean Copulas we make use of Nelsen (2006)'s Lemma 4.1.2; Archimedean copula is defined as

$$C(u, v) = \phi^{[-1]}(\phi(u) + \phi(v)),$$

where $\phi^{[-1]}$ is the pseudo-inverse of ϕ , a continuous decreasing function from $\mathbf{I}$ to $[0, \infty]$ such that $\phi(1) = 0$, referred to as the generator. $\phi^{[-1]}$ and given by

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

In Table 1 we represent just a few of the well known Archimedean copulas, for more families of Archimedean copulas see Nelsen (2006).

As for elliptical copulas, we have plotted the contour plots of the three Archimedean copulas presented in Table 1. Again, the marginal distributions of all three copulas are assumed to be standard normal with correlation coefficient 0.5.

Table 1: Some families of Archimedean Copulas

Family	Copula Function	Generator
Gumbel copula	$uv \exp(-\theta \ln u \ln v)$	$\ln(1 - \theta \ln t)$
Clayton copula	$\left[\max(u^{-\theta} + v^{-\theta} - 1, 0) \right]^{\frac{1}{\theta}}$	$\frac{1}{\theta}(t^{-\theta} - 1)$
Frank copula	$-\frac{1}{\theta} \ln \left(1 + \frac{(e^{-\theta u} - 1)(e^{-\theta v} - 1)}{e^{-\theta} - 1} \right)$	$-\ln \frac{e^{-\theta t} - 1}{e^{-\theta} - 1}$

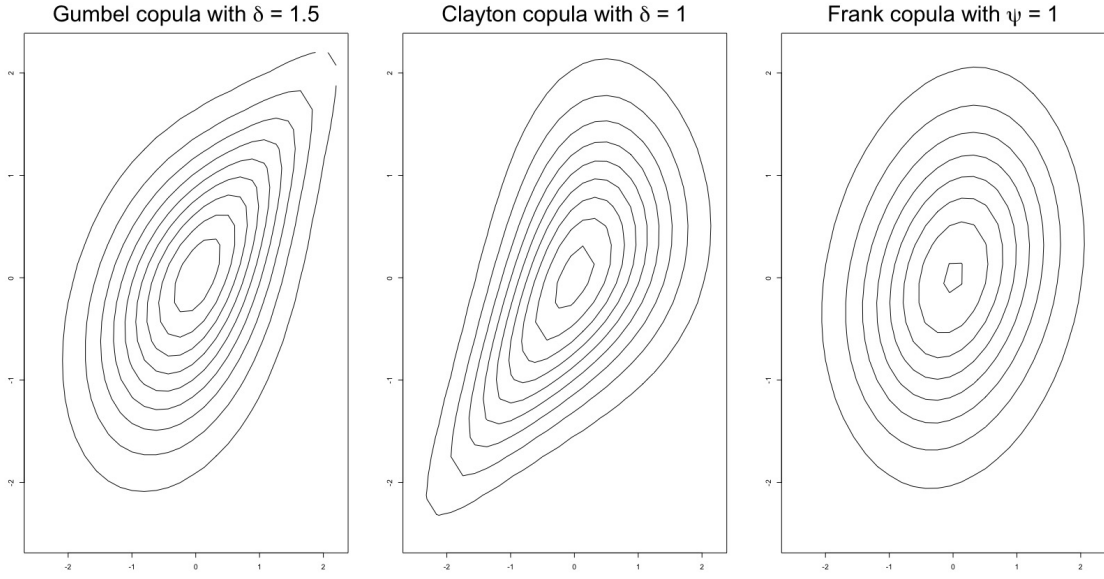


Figure 3: Contour plots of Gumbel, Clayton and Frank Copulas

3.3 Estimation of Copulas

Existing literature has often assumed that in copula-based models the return series of the underlying assets in a portfolio follow an ARMA-GARCH model. These models that have become an important tool in the analysis of time series data in financial applications, are designed to deal with issues common to financial assets, such as for example volatility clustering and heteroskedasticity. By considering such a model for the marginal distributions, we allow for conditional mean and conditional variance.

Denoting the return of a given underlying asset at time point t with r_t , an ARMA(p,q)-

GARCH(r,s) model is defined as

$$r_t = \mu + \sum_{i=1}^p \alpha_i r_{t-i} + \sum_{j=1}^q \beta_j \epsilon_{t-j} \quad (3.12)$$

$$\sigma_t^2 = \omega + \sum_{k=1}^r \delta_k \sigma_{t-k}^2 + \sum_{l=1}^s \gamma_l \epsilon_{t-l}^2 \quad (3.13)$$

$$\epsilon_t = \sigma_t z_t, \quad (3.14)$$

where $z_t \stackrel{iid}{\sim} F_t$ with F_t denoting the conditional distribution of standardized residual, (3.12) specifies the conditional mean and (3.13) conditional variance of the process. Estimating the parameters for the models of the conditional mean and conditional variance, we obtain the estimated standardized residuals $\hat{z}_t$.

There is a number of ways how to estimate copula models, a short overview of which is provided in Patton (2012) that serves as the source for this section. The choice of the method depends on whether the distribution of standardized residuals is parametric or nonparametric. In the parametric case, the distribution of standardized residuals can vary in time as a (parametric) function of $\mathcal{F}_{t-1}$ -measurable variables, or may be constant. In the nonparametric case, the most common approach is that the conditional distribution is constant and estimable via the empirical distribution function.

In case of a fully parametric model, where the (conditional) marginal distributions and the (conditional) copula are all assumed to be known up to a finite dimensional parameter, maximum likelihood is commonly used. In unconditional case, the log-likelihood is given by

$$l(\vartheta) = \sum_{t=1}^T \ln c(F_1(x_{1,t}; \vartheta_1), F_2(x_{2,t}; \vartheta_2); \theta) + \sum_{t=1}^T \sum_{i=1}^2 \ln f_i(x_{i,t}; \vartheta_i) \quad (3.15)$$

and is maximized by its maximum likelihood estimate

$$\hat{\vartheta} = \arg \max_{\vartheta} l(\vartheta).$$

Equivalently, in the conditional case the log-likelihood function is expressed as

$$l(\vartheta) = \sum_{t=1}^T \ln c(F_1(x_{1,t}; \vartheta_1 | \mathbf{w}_t), F_2(x_{2,t}; \vartheta_2 | \mathbf{w}_t); \theta_t | \mathbf{w}_t) + \sum_{t=1}^T \sum_{i=1}^2 \ln f_i(x_{i,t}; \vartheta_i | \mathbf{w}_t),$$

and maximized by its maximum likelihood estimate

$$\hat{\vartheta} = \arg \max_{\vartheta} l(\vartheta).$$

The time-varying copula parameter, denoted with θ_t , varies through time according to a certain evolution equation. Stochastic volatility copula models such as those considered in for example Hafner and Manner (2012) and Manner and Segers (2011) model the time-varying parameter with a latent time series process. On the other hand, ARCH-type copula models employed in Patton (2006b), Jondeau and Rockinger (2006) and Creal et al. (2013) for example define the time-varying parameter by a function of historic observations (Oh and Patton (2016)). More specifically, in the latter case, the evolution equation of the time-varying parameter is defined by a forcing variable that is a function of past observations and might be in some cases, if the parameter has no interpretation, difficult to define.

A large shortcoming of maximum likelihood estimation is the effort the method requires when considering models with large number of parameters or dimensions. If the parameters of the marginal distributions can be separated from each other and from the parameters of the copula, this drawback of the one-stage maximum likelihood can be overcome by estimating the parameters via the multistage maximum likelihood. This approach allows the parameters of the marginal distributions to be estimated using maximum likelihood in the first stage, and the copula parameters to be estimated in the second stage conditional on the estimated parameters of the marginal distributions. Even though multistage maximum likelihood approach is asymptotically less efficient than the one stage maximum likelihood estimation, Joe (2005) and Patton (2006a) have shown that this loss of efficiency is in most cases not extreme.

Estimation of semiparametric models makes use of the attractive aspect of copula theory, namely the decomposition of a joint distribution into the marginal distribution model(s) and the copula model, allowing separate estimation thereof. Semiparametric copula models therefore use a nonparametric model for the marginal distributions and a parametric model for the copula, where the copula parameter

is usually estimated via maximum likelihood and can be referred to as canonical maximum likelihood or pseudo maximum likelihood estimator.

Besides estimation of parametric and semiparametric models, it is also possible to estimate copulas by nonparametric methods, that are based on empirical distributions. These copulas resemble usual multivariate empirical cumulative distribution functions and are highly discontinuous. They prove however to be inefficient in finding a convenient parametric family of copulas by simply visually investigating available data sets (Scaillet and Fermanian (2002)). Fully nonparametric studies of the copula using time series are provided in Scaillet and Fermanian (2002), Fermanian et al. (2004), Sancetta and Satchell (2004) and Ibragimov (2009).

In addition to maximum likelihood estimation, that is the most common approach so far, also other methods have been applied. Genest (1987), Ghoudi and Rémillard (2004), Rémillard (2010) and Genest et al. (2011) consider the method of moments-type estimators, where the parameter of a given family of copulas has a known, invertible mapping to a dependence measure, such as rank correlation or Kendall's tau. The paper of Oh and Patton (2013) applies the simulated method of moments estimations, considering a case where the number of dependence measures might be greater than the number of unknown parameters, as well as an analogous simulation-based method. In his paper, Tsukahara (2005) uses minimum distance estimation, and Bayesian estimation copula methods have been used by Min and Czado (2010), Pitt et al. (2006), Smith et al. (2012) and Rosenberg (2000).

4 The Conditional Dependence between the Euro and the Pound

The theory presented in the previous chapters is applied to the problem of estimating Value-at-Risk of a portfolio of two exchange rates. The dataset consists of daily quotations of Euro-U.S. Dollar and British Pound-U.S. Dollar exchange rates, taken from the database Datastream International, over the period November 30, 2005 to November 30, 2015, yielding 2609 observations. The analysis is performed using the log-differences of daily quotations, to which we refer to as series in the remainder of the thesis.

Figure 4 presents the plots of the series over the observation period. Examining the plots we can see that both return series exhibit high volatility in the period between years 2008 and 2009.

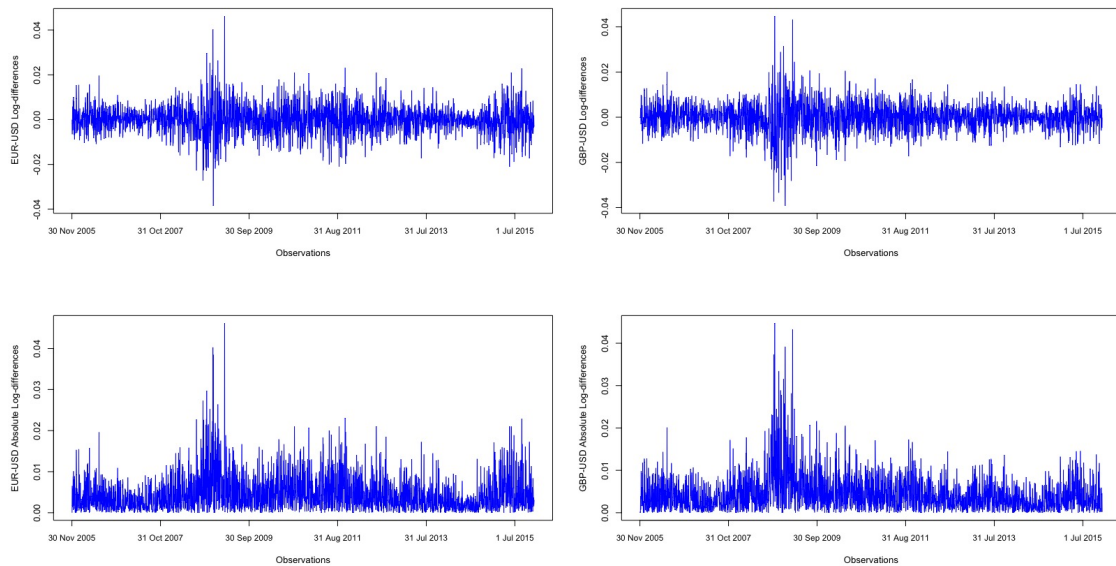


Figure 4: Daily and absolute log-differences of the exchange rates

We present the descriptive statistics of the data, the test statistics of Jarque-Bera and ARCH LM tests, obtained using 100—times the log-differences of the daily exchange rates, and their p -values in brackets in Table 2. Means of both series

are, compared to the standard deviation of each series, relatively small, implying neither exchange rate had a significant trend over the period. The GBP-USD series exhibits slight negative skewness, whereas both series exhibit excess kurtosis. The Jarque-Bera test for normality of the unconditional distribution of each exchange rate rejects the null hypothesis that the series are drawn from a normal population, and the unconditional correlation coefficient between the returns of the two exchange rates indicates high linear dependence. Based on the results of the ARCH LM test, conducted using 10 lags, we can conclude that the residuals of respective series exhibit conditional heteroskedasticity.

Table 2: Summary Statistics

	EUR-USD	GBP-USD
Mean	-0.00	-0.00
Standard Deviation	0.63	0.61
Skewness	0.14	-0.07
Kurtosis	6.44	8.15
Jarque-Bera statistic	1291.51	2882.26
	(0.00)	(0.00)
ARCH LM statistic	197.97	385.69
	(0.00)	(0.00)
Linear correlation	0.64	

The model of the bivariate density of the EUR-USD and GBP-USD exchange rates is specified in two steps. In subsection 3.1 we define the models for the marginal distributions of each exchange rate and in subsection 3.2 we introduce the model for the conditional copula. We estimate the coefficients using multistage maximum likelihood approach discussed in Section 3.3, first estimating the two marginal distribution models and then estimating the copula model.

4.1 The Models for the Marginal Distributions

The models we use for the marginal distributions are assumed to be characterized by AR, st -GARCH models, where st denotes the skewed Student's t distribution. The ARMA-GARCH models have been successfully used to model return series, whereas the choice of skewed Student's t -distribution is justified by papers such as for example Bollerslev (1987), showing that the conditional univariate distribution of daily exchange rate returns can be well described by the Student's t distribution. In Figure 5 we plot the Autocorrelation plots of the log-differences and squared log-differences of both series. The Autocorrelation plot of the log-differences of the EUR-USD series suggests there is no significant correlation to be captured by the ARMA model, while the plot of GBP-USD series suggests autoregressive terms at lags 5 and 11. The results of the ARCH LM test we have performed for both series are confirmed by the Autocorrelation plots of their squared log-differences, showing autocorrelation in the second moments.

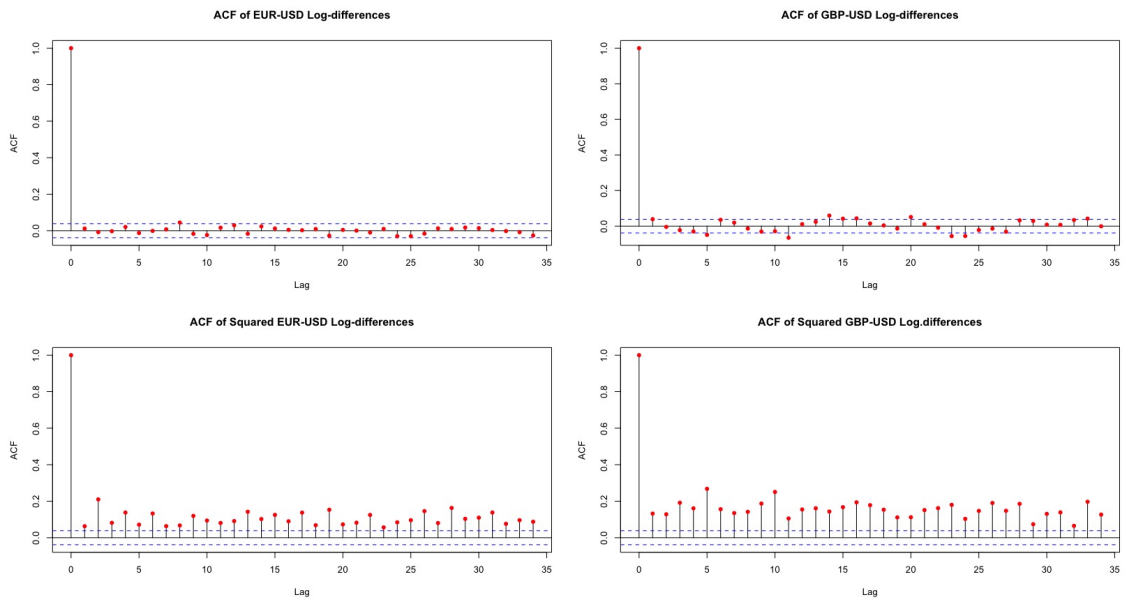


Figure 5: Autocorrelation of the log-differences of the exchange rates

The models for the marginal distributions are determined based on visual examination of the plots of the log-differences of both series as well as the Akaike Information Criterion for the AR process, whereas the order of the GARCH model is motivated

by its frequent use in the existing literature. The marginal distribution of the EUR-USD series follows an AR(0), *st*-GARCH(1,1)

$$X_t = \mu_x + \epsilon_t \quad (4.16)$$

$$\sigma_{x,t}^2 = \omega_x + \beta_x \sigma_{x,t-1}^2 + \alpha_x \epsilon_{t-1}^2 \quad (4.17)$$

$$\epsilon_t = \sigma_{x,t} z_{t,x}, \quad z_{t,x} \stackrel{iid}{\sim} st_{\nu_x, \mu_x}, \quad (4.18)$$

where X_t denotes the exchange rate's log-difference. The marginal distribution of the GBP-USD series is characterized by an AR(5,11), *st*-GARCH(1,1) specification with Y_t as the log-difference of the exchange rate

$$Y_t = \mu_y + \phi_{5y} Y_{t-5} + \phi_{11y} Y_{t-11} + \eta_t \quad (4.19)$$

$$\sigma_{y,t}^2 = \omega_y + \beta_y \sigma_{y,t-1}^2 + \alpha_y \eta_{t-1}^2 \quad (4.20)$$

$$\eta_t = \sigma_{y,t} z_{t,y}, \quad z_{t,y} \stackrel{iid}{\sim} st_{\nu_y, \mu_y}. \quad (4.21)$$

The maximum likelihood estimates of the parameters of the marginal distribution models, with standard errors in parantheses, are presented in Table 3. We notice that the estimated degrees of freedom differ significantly between the two models.

Table 3: Results for the Marginal Distributions

	EUR-USD Margin		GBP-USD Margin	
Constant	-0.00	(0.00)	-0.00	(0.00)
AR(5)	-	-	0.02	(0.02)
AR (11)	-	-	0.05	(0.02)
GARCH constant	0.00	(0.00)	0.00	(0.00)
Lagged variance	0.96	(0.00)	0.96	(0.00)
Lagged e^2	0.03	(0.00)	0.04	(0.00)
Degrees of freedom	7.79	(1.68)	10.79	(2.88)

Patton (2006b) emphasizes the need to test conditional mean and variance models of both exchange rates, defined in (4.16), (4.17), (4.19) and (4.20), for the influ-

ence of variables of the other exchange rate model in each individual case. If no lags of variables of the other exchange rate are significant in the conditional mean and variance models, we can consider only univariate models for the two marginal distribution. As the conditional mean model for the EUR-USD margin does not consider any lag of the series, we only test whether the two significant lags of the GBP-USD margin are important for the conditional mean and variance models of the EUR-USD exchange rate. We test the conditional mean model of EUR-USD exchange rate by regressing the residuals ϵ_t on Y_{t-5} and Y_{t-11} lags of the GBP-USD series and test that all coefficients equal 0. At the same time, we test the conditional variance models by regressing the standardized squared residuals of one series on the lagged squared residuals of the other series and test whether the coefficients are 0. In Table 4 we report the p -values we have obtained performing the described tests. The results confirm it is safe to consider only univariate models for the two marginal distributions.

Table 4: Influence of the "Other" Variable in the Mean and Variance Models

	p -Value
Y_{t-5} and Y_{t-11} in conditional mean model of X_t	0.5603
ϵ_{t-1}^2 in conditional variance model of Y_t	0.3131
η_{t-5} and η_{t-11} in conditional variance model for X_t	0.9053

In order to properly model the conditional copula, our models for the marginal distributions need to be indistinguishable from the true marginal distributions. Should our marginal distributions be misspecified, the probability integral transforms will not be standard uniformly distributed and consequently our copula model will be misspecified (Patton (2006b)). To ensure our copula model is correctly specified, we need to test the marginal distribution models for misspecification.

Comparing the standardized residuals plotted in Figure 6 to plots presented in Figure 4 and 5, we see that the standardized residuals are now approximately independent.

Further on, examining the $q - q$ plots of the transformed series $u_t = \hat{F}_x(x_t | \mathcal{F}_{t-1})$ and $v_t = \hat{F}_y(y_t | \mathcal{F}_{t-1})$, where we denote the estimated marginal distributions of the respective series by $\hat{F}_x$ and $\hat{F}_y$ and the information set available up to time $t - 1$ by $\mathcal{F}_{t-1}$, we can conclude that both series are appropriately specified.

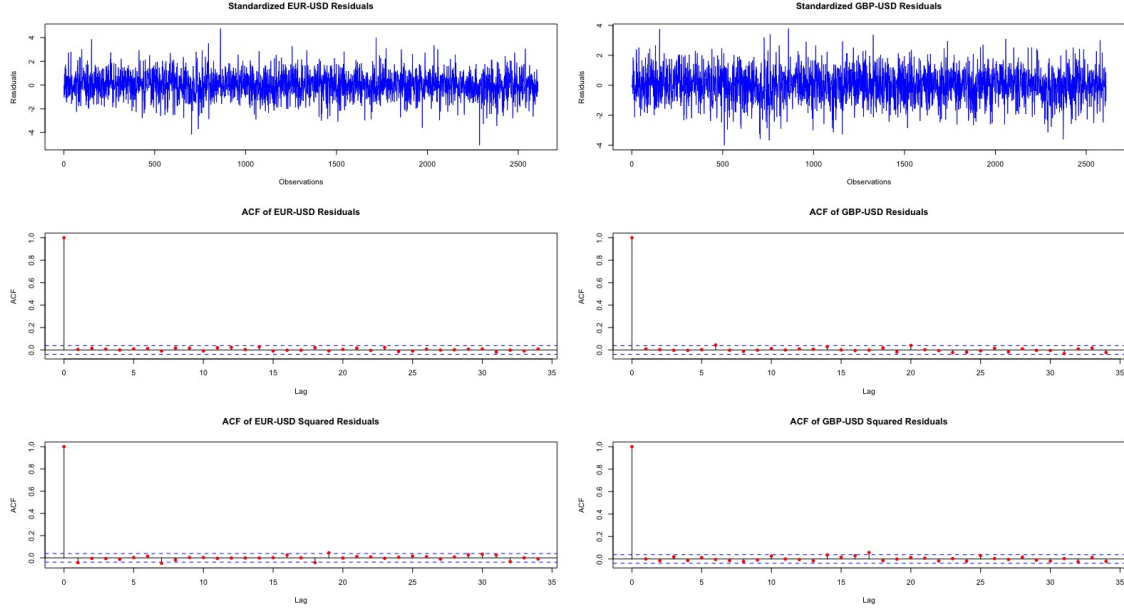


Figure 6: Characteristics of standardized residuals

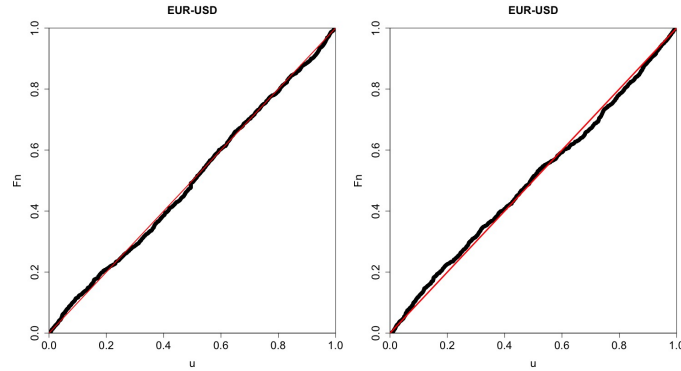


Figure 7: Empirical distribution of the transformed series u_t and v_t

In addition we apply the Ljung-Box test to the residuals and squared residuals of both series and test the null hypothesis of no autocorrelations from lag 1 to 15, 30 and 60 at 5% significance level. The p -values of the tests presented in Table 5 imply that we cannot reject the null hypothesis of no autocorrelation. Lastly, we consider

the Anderson-Darling goodness of fit test and test whether the standardized residual series are a sample of data from the skewed Student's t distribution with estimated parameters. The results of the Anderson Darling Goodness of Fit test are presented in Table 6.

Table 5: p -values of Ljung-Box test

	15 lags	30 lags	60 lags
EUR-USD residuals	0.9141	0.9970	0.9045
GBP-USD residuals	0.8792	0.9262	0.9832
EUR-USD squared residuals	0.5605	0.3489	0.2620
GBP-USD squared residuals	0.7222	0.4666	0.7923

Table 6: Results of Anderson Darling Goodnes of Fit Test

	EUR-USD	GBP-USD
p -value	0.0508	0.0100

Although the Anderson Darling Goodness of Fit test fails to reject the null hypothesis for the GBP-USD series only when considering a confidence level of 1%, we conclude that both models are overall adequately specified.

4.2 Modeling the Dependence between Exchange Rates

We model the dependence between the exchange rates by normal copula with and without time variation. In the time-varying option we assume that the functional form of the copula remains fixed over the sample whereas the parameters are defined by some evolution equation and therefore changing through time.

The time-varying parameter ρ_t of the normal copula is assumed to follow an ARMA(1,10) process with ρ_{t-1} as a regressor that captures the persistence in the dependence parameter, and the mean of the product of the last 10 observations of the transformed

variables $\Phi^{-1}(u_{t-j})$ and $\Phi^{-1}(v_{t-j})$, capturing the variation in dependence. The evolution equation can therefore be written as

$$\rho_t = \Lambda\left(\omega_\rho + \beta_\rho \rho_{t-1} + \alpha \sum_{j=1}^{10} \Phi^{-1}(u_{t-j}) \Phi^{-1}(v_{t-j})\right), \quad (4.22)$$

where $\Lambda(x) \equiv (1 - e^{-x})(1 + e^{-x})^{-1}$ keeps ρ_t in $(-1, 1)$ at all times (Patton (2006b)). We estimate coefficients of (4.22) using maximum likelihood estimation and thus determine the normal copula parameter for each time point t . The maximum likelihood estimates of the parameters of the time-varying copula parameter and respective standard deviations are presented in Table 7.

Table 7: Estimated Parameters of Time-Varying Normal Copula

Constant	-0.11	(0.00)
α	0.12	(0.60)
β	2.40	(0.05)

For Value-at-Risk estimation we assume that the exchange rates have equal weights in the portfolio. In case of the time-varying copula approach we determine the copula for each time point based on the estimated time-varying parameter. To estimate Value-at-Risk of the portfolio we use the Monte Carlo approach as described in Section 2.1.3; from the defined copula with estimated margins we simulate 1000 realizations of the standardized residuals for time point t , scale them back to log returns and based on the resulting values of the portfolio at that time point estimate Value-at-Risk.

Our benchmark model is constructed using the rolling-window approach in the following way. We fix a window of 250 observations that at each step defines the sample we use to estimate the copula parameter, defined as the correlation between the uniform series u_t and v_t we have obtained from the log-returns of both exchange rates. We estimate Value-at-Risk as described above using the Monte Carlo approach, roll the window one observation into the future and repeat the estimation on the new

data sample. With this approach we roll the window over the entire data set, keeping the estimation period constant, and end up with 2359 Value-at-Risk estimates of our portfolio.

In addition to the two approaches described above, we consider for further comparison also the constant copula approach. Value-at-Risk estimation in this case is essentially the same to the time-varying one, with the exception that we determine the copula parameter in the beginning based on the whole data sample and keep it constant through out the observation period. As in case of the benchmark model, the copula parameter is defined as the correlation between the uniform series u_t and v_t . One should note though, that the constant copula can only be used for Value-at-Risk estimation when considering the whole data sample for the estimation of the parameters of the model (only in the in-sample estimation), while it by definition does not allow forecasting. In Figure 8 we plot the estimated parameters of all three approaches.

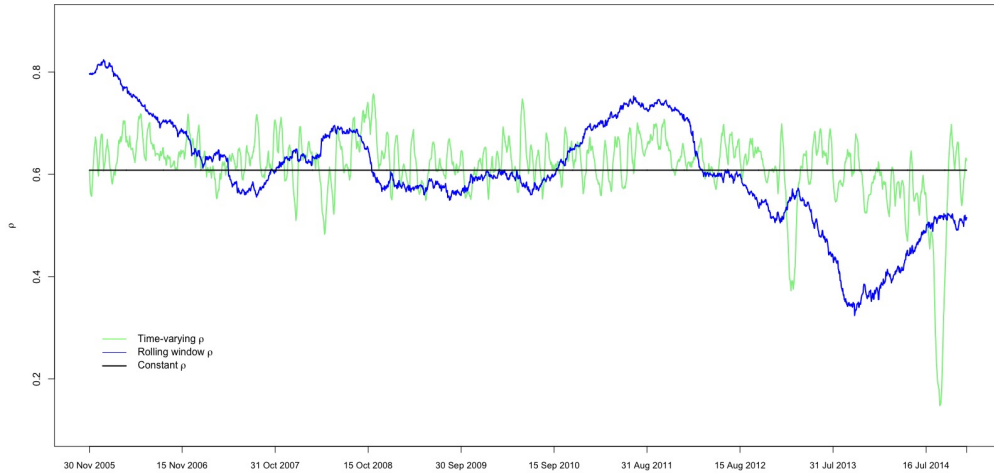


Figure 8: Estimated Parameters of the Normal Copula

5 Back-testing

Value-at-Risk models are useful only if they adequately capture the risk of the portfolio. The accuracy of the Value-at-Risk estimates should therefore always be compared to actual losses, normally done via a statistical procedure referred to as back-testing. We have based this section, providing an overview of commonly used back-test, on Campbell (2006).

A variety of different tests have been proposed, many of which focus on comparing the reported Value-at-Risk to the realized profit or loss. They do so by considering whether the loss on a portfolio exceeds the reported Value-at-Risk, an event that can be described by a hit function

$$I_{t+1}(\alpha) = \begin{cases} 1 & \text{if } x_{t,t+1} \leq VaR_t(\alpha) \\ 0 & \text{if } x_{t,t+1} > VaR_t(\alpha) \end{cases}$$

where $x_{t,t+1}$ denotes the profit or loss on the portfolio between time points t and $t + 1$. Such function describes the history whether or not a loss larger than the reported Value-at-Risk has been realized. Christoffersen (1998) reduced the problem of determining whether a Value-at-Risk model is an accurate one or not to the problem of whether the hit sequence $[I_{t+1}(\alpha)]_{t=1}^{t=T}$ satisfies the unconditional coverage and independence property. A Value-at-Risk measure deemed an accurate one is a measure that produces hit series satisfying both properties (Campbell (2006)).

The unconditional coverage property places a restriction on how often Value-at-Risk violations may occur, whereas the independence property on how the violations may occur. The hit sequence satisfies the unconditional coverage property if the probability of realizing a loss larger than reported Value-at-Risk equals $\alpha\%$. More frequent losses larger than the reported Value-at-Risk imply the Value-at-Risk measure in question systematically understates the risk level of the portfolio. In the opposite way, experiencing too few violations would suggest that the Value-at-Risk measure is too conservative.

On the other hand, the independence property requires that any two elements of

the hit sequence $(I_{t+j}(\alpha), I_{t+k}(\alpha))$ are independent from each other. The history of Value-at-Risk violations must thus not convey any information about whether or not an additional Value-at-Risk violation will occur. If it does, the Value-at-Risk measure is considered an inadequate one, as such clustering of violations suggest that the measure is not as responsive as it should be. If for example $\alpha\%$ Value-at-Risk violation is immediately followed by another violation, the probability of observing a loss in excess of the $\alpha\%$ Value-at-Risk after one has already been observed is 100% and not $\alpha\%$, implying that the reported $\alpha\%$ Value-at-Risk does not accurately reflect the loss that can be expected to be exceeded $\alpha\%$ of time.

The two properties can be combined into a single requirement, that the hit sequence $I_t(\alpha)$ is identically and independently distributed as a Bernoulli random variable with probability α ,

$$I_t(\alpha) \stackrel{i.i.d.}{\sim} B(\alpha).$$

5.1 Kupiec's Proportion of Failures Test

Kupiec (1995) proposed a proportion of failures (POF) test, examining the unconditional coverage property of a Value-at-Risk measure investigating how many times estimated Value-at-Risk is violated over a given period of time. If the number of violations is considerably smaller than the $\alpha\%$ of the sample the model is underestimating the risk, and if the number of violations is considerably larger is it overestimating it.

Denoting the total number of exceptions with $I(\alpha) = \sum_{t=1}^T I_t(\alpha)$ and total number of observations with T , the null hypothesis of POF test can be written as

$$H_0 : \quad \alpha = \hat{\alpha} = \frac{I(\alpha)}{T}.$$

The test statistic is defined as

$$\begin{aligned} POF &= 2 \log \left(\left(\frac{1 - \hat{\alpha}}{1 - \alpha} \right)^{T - I(\alpha)} \left(\frac{\hat{\alpha}}{\alpha} \right)^{I(\alpha)} \right) \\ \hat{\alpha} &= \frac{1}{T} I(\alpha) = \sum_{t=1}^T I_t(\alpha) \end{aligned}$$

and is asymptotically χ^2 distributed with 1 degree of freedom under null hypothesis. The more the proportion of Value-at-Risk violations differs from $\alpha\%$, the more the POF test statistic grows, indicating that the Value-at-Risk measure being tested does not capture the underlying risk of the portfolio accurately.

Campbell (2006) proposes an alternative test of the unconditional coverage property to that of Kupiec (1995), that is based directly on the sample average of the number of Value-at-Risk violations over a given time period, $\hat{\alpha}$. Assuming that the Value-at-Risk measure is accurate, the scaled version of the sample average $\hat{\alpha}$,

$$z = \frac{\sqrt{T}(\hat{\alpha} - \alpha)}{\sqrt{\alpha(1 - \alpha)}}$$

is approximately standard normally distribution, and hypothesis can be tested the same way as in the case of the Kupiec (1995) POF test.

Furthermore, Campbell (2006) reports that the tests of the unconditional coverage property have difficulties in detecting Value-at-Risk measures that systematically under report risk. They exhibit low power in sample sizes consistent with one year, where the size of under reporting can be quite significant. In addition, such tests focus only on the unconditional property of a Value-at-Risk measure and can therefore fail to detect a Value-at-Risk measures that does not fulfill the independence property. A sequence of large dependent losses may have worse consequences for a financial institution than are those caused by large unexpected losses that are experienced more often than expected but spread over time.

5.2 Independence Tests

A popular test of independence property is Christoffersen (1998) Markov test, examining whether the probability of a Value-at-Risk violation depends on whether a Value-at-Risk violation occurred in the previous day. Denoting the number of days when condition j occurred given that the condition i occurred in the previous day with n_{ij} , we can create the contingency table presented below, where $n = n_{00} + n_{01} + n_{10} + n_{11}$.

Table 8: Contingency Table for Markov Independence Test

	$I_{t-1} = 0$	$I_{t-1} = 1$	
$I_t = 0$	n_{00}	n_{10}	$n_{00} + n_{10}$
$I_t = 1$	n_{01}	n_{11}	$n_{01} + n_{11}$
	$n_{00} + n_{01}$	$n_{10} + n_{11}$	n

A Value-at-Risk measure that satisfies the independence property is a measure for which the probability of violating today's Value-at-Risk is not dependent on whether yesterday's Value-at-Risk was violated or not. In other words, the proportions of violations that occur after a previous violation, $I_{t-1} = 1$, and those that occur after a day with no violations, $I_{t-1} = 0$, are the same. The null hypothesis of Markov test is therefore

$$H_0 : \quad \frac{n_{01}}{n_{00} + n_{01}} = \frac{n_{11}}{n_{10} + n_{11}}, \quad (5.23)$$

and the corresponding test statistic can be written as

$$\begin{aligned}
 MT &= -2\ln\left(\frac{(1 - \pi)^{n_{00}+n_{10}} \cdot \pi^{n_{01}+n_{11}}}{(1 - \pi_0)^{n_{00}} \cdot \pi_0^{n_{01}} \cdot (1 - \pi_1)^{n_{10}} \cdot \pi_1^{n_{11}}}\right) \\
 \pi_0 &= \frac{n_{01}}{n_{00} + n_{01}} \\
 \pi_1 &= \frac{n_{11}}{n_{10} + n_{11}} \\
 \pi &= \frac{n_{01} + n_{11}}{n_{00} + n_{01} + n_{10} + n_{11}}.
 \end{aligned}$$

The above specified test statistic is, under null hypothesis, χ^2 distributed with 1 degrees of freedom.

A more recent independence test, proposed by Christoffersen and Pelletier (2004), is based on the fact that for an accurate Value-at-Risk measure the time between Value-at-Risk violations should not exhibit any kind of duration dependence. In more detail, for completely independent Value-at-Risk violations it must hold that the amount of time that has passed between the violations is independent of the time that has passed since the last violation. They prove that the proposed test is

more powerful in detecting Value-at-Risk measure that violates the independence property than the Markov test we have discussed above.

In the world of independence tests, an accurate Value-at-Risk measure will result in a series of independent hits, $[I_t(\alpha)]_{t=1}^{t=T}$. The independence property can however be violated in many ways. For example, in the case of Markov test, a Value-at-Risk measure does not violate the independence property if the probability of violating Value-at-Risk today does not depend on whether the Value-at-Risk was violated yesterday. Nevertheless, it is possible that the probability of violating Value-at-Risk depends on whether it was violated two days or one week ago. It is thus of great importance that any test of the independence property fully describes in which way the independence property might be violated and the alternative hypothesis against which the property is being tested (Campbell (2006)).

5.3 Joint Tests of Unconditional Coverage and Independence

Both Christoffersen (1998) Markov test and the duration test of Christoffersen and Pelletier (2004) can be extended in a way that allows joint testing for independence and unconditional property. The joint Markov test as described in Campbell (2006) examines whether the probability of a Value-at-Risk violation that follows another violation equals to the probability that there was no violation in the previous day, and at the same time examines if either of these probabilities is significantly different than α . In terms of the contingency table, the joint Markov test examines whether the following relation holds:

$$\frac{n_{01}}{n_{00} + n_{01}} = \frac{n_{11}}{n_{10} + n_{11}} = \frac{n_{01} + n_{11}}{n} = \alpha.$$

Campbell (2006) points out that although it might seem that joint tests are preferable to tests of either the unconditional coverage or independence property, they have in fact decreased ability to detect a Value-at-Risk measure that only violates one of the two properties.

5.4 Back-testing results

5.4.1 In-sample back-testing

We first perform an in-sample back-test of the models estimated as described in previous sections. The estimates of Value-at-Risk at different levels of α and realized returns over the whole data sample are presented in Figures 9, 10 and 11. The estimated Value-at-Risk is back-tested to realized returns by performing Kupiec's Proportion of Failure Test and Christoffersen's Markov test discussed in the previous section.

For each of the copula approaches and each α we represent the number of actual exceedances, the test statistic as well as the p -value of both Kupiec's Proportion of Failure test and Christoffersen's Markov test, denoted with LR_{UC} and LR_{CC} , respectively, in Table 9. We cannot obtain the test statistic and consequently the p -value for the Christoffersen's Markov test for $\alpha = 1\%$ for all three copula approaches. The reason lies in the fact that no Value-at-Risk estimation violating the loss experienced at time point t is followed by another Value-at-Risk estimation that violates the experienced loss in the next time point. This means that n_{11} from contingency table defined in Table 8 and consequently π_{11} of the test statistic equals 0, making the test statistic of the Markov test undefined. In terms of the null hypothesis that we are testing this implies that the ratio on the right hand side of the equality in (5.23) equals 1. As n_{00} and n_{01} are both different than 0, the proportion of violations that occur after a previous violation is not the same as the proportion of violations that occur after a day in which no violations occurred and null hypothesis is rejected.

We notice that the observed number of violations is for all copula approaches and all values of α fairly close to the expected number of violations. Should the observed number of violations be much larger than the number of expected ones, the model would overestimate the risk of the portfolio. On the other hand, too few violations imply that the model fails to capture historical information. From the results presented in the Table 9 we thus conclude, that in-sample back-testing the

normal copula with time-varying parameter outperforms the benchmark model as well as the normal copula with constant parameter, as both tests fail to reject their null hypothesis.

Table 9: In-Sample Back-testing Results

	VaR_{0.1}	VaR_{0.05}	VaR_{0.01}
Number of expected violations	235	117	23
Time-varying Normal Copula			
Number of observed violations	257	125	16
LR_{UC} test statistic	2.04	0.44	2.78
LR_{UC} p -value	0.15	0.51	0.10
LR_{CC} test statistic	4.12	0.46	n.a.
LR_{CC} p -value	0.13	0.80	n.a.
Benchmark Model			
Number of observed violations	272	127	15
LR_{UC} test statistic	5.89	0.71	3.63
LR_{UC} p -value	0.02	0.40	0.06
LR_{CC} test statistic	8.12	0.83	n.a.
LR_{CC} p -value	0.02	0.66	n.a.
Constant Normal Copula			
Number of observed violations	256	130	11
LR_{UC} test statistic	1.86	1.26	8.46
LR_{UC} p -value	0.17	0.26	0.00
LR_{CC} test statistic	7.05	1.36	n.a.
LR_{CC} p -value	0.03	0.51	n.a.

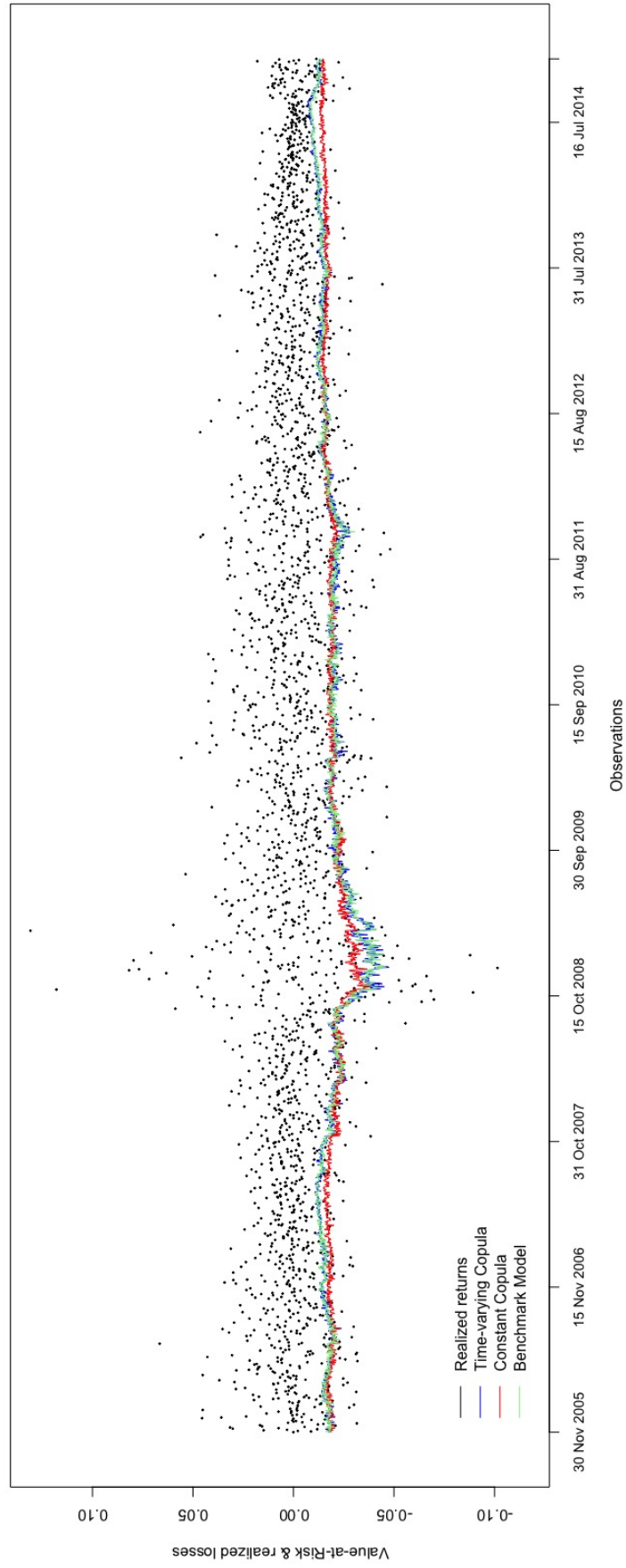


Figure 9: In-Sample Value-at-Risk estimation at $\alpha = 0.10$

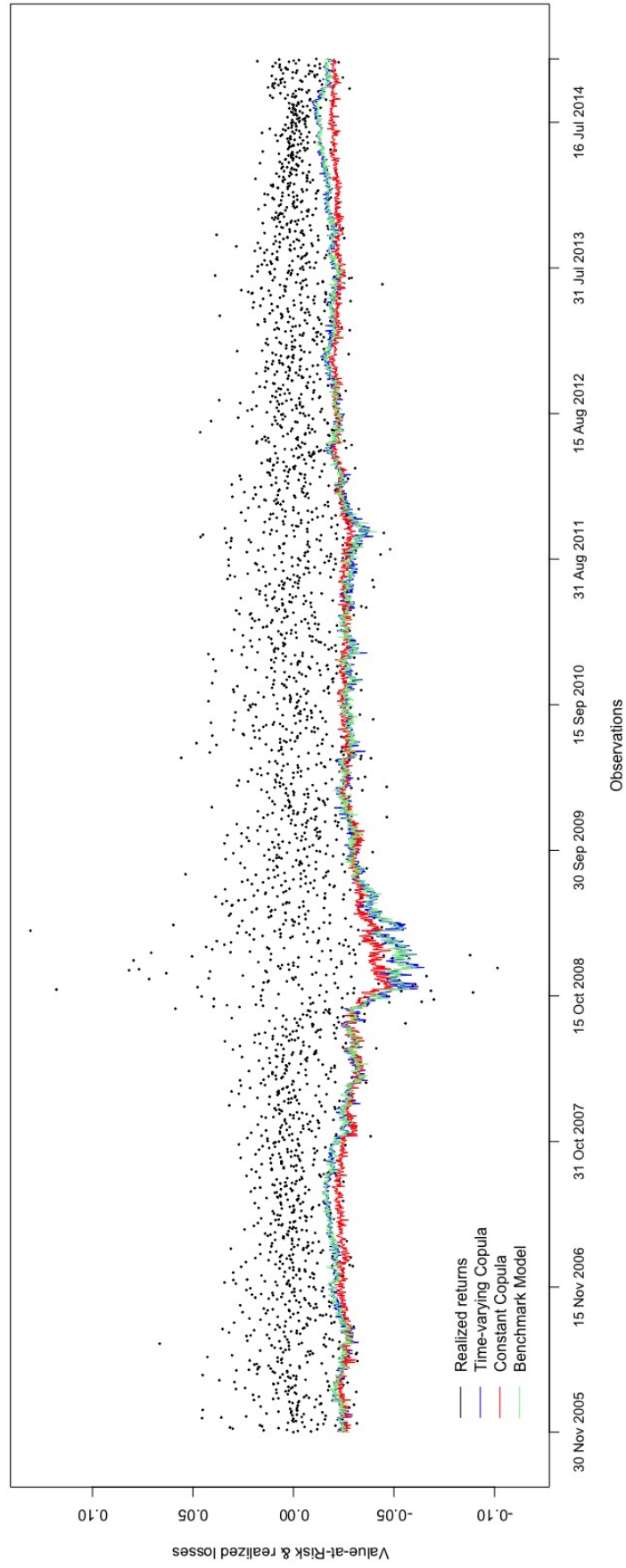


Figure 10: In-Sample Value-at-Risk estimation at $\alpha = 0.05$

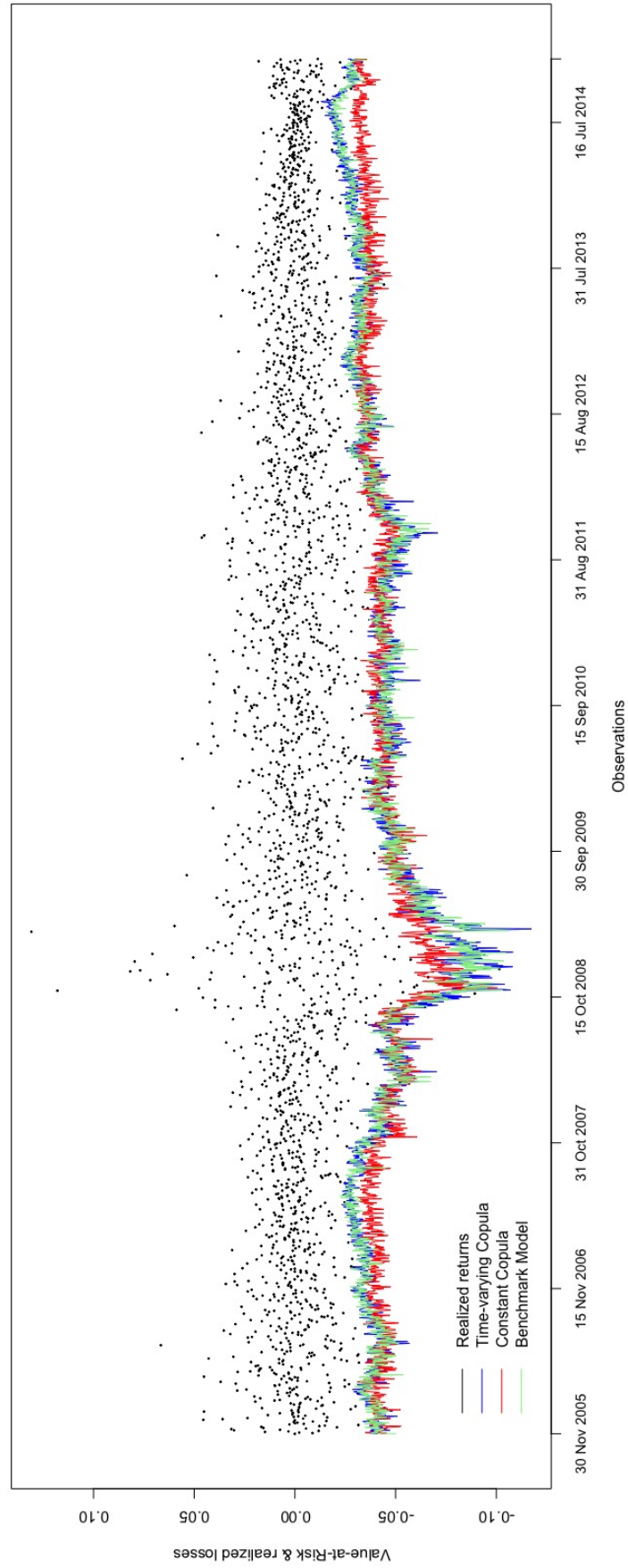


Figure 11: In-Sample Value-at-Risk estimation at $\alpha = 0.01$

5.4.2 Out-of-sample back-testing

To really assess the predictive ability of the Value-at-Risk models and establish whether allowing for time variation in the copula parameter leads to better Value-at-Risk estimates as those obtained considering a simple rolling window model, the Value-at-Risk estimates of both approaches are back-tested to realized returns also via the out-of-sample back-test.

The difference between the in-sample and out-of-sample back-test is that when performing an in-sample back-test, the parameters of marginal distribution models and the copula model as well as Value-at-Risk for respective approaches are estimated based on the data of the whole sample. For time-varying copula approach this means that the parameters of the time-varying copula parameter are estimated based on the whole data sample, while similarly, benchmark copula parameter is estimated on the whole rolling window sample.

When performing an out-of sample back-test however, the copula parameters and values of conditional mean and standard deviation are forecasted based on a sample of data. More specifically, all models are initially estimated on an initial data sample and for every future time point, the copula parameter and the values of conditional mean and conditional standard deviation are forecasted one time point ahead in the future. In the next step, Value-at-Risk is estimated and compared to the realized loss of the portfolio. The window of observations is then moved one observation into the future keeping the length of the data sample, and the models for the marginal distributions and the copula parameters are re-estimated. The window of observations is rolled through out the data sample, repeating the described process at every step, yielding a sequence of one-step ahead Value-at-Risk forecasts that are compared to realized losses of the portfolio.

We have based our decision to use the rolling window approach and re-estimate all models at every step on the works of Palaro and Hotta (2006), Aeppli (2015), Aeppli et al. (2015), Hernández-Lobato et al. (2013), Cerrato et al. (2015) as well as Fengler and Okhrin (2014). We have considered 250 as well as 1000 observations

for the length of the rolling window, while, as out-of-sample back-testing using a window of 250 observation provides similar results as are those obtained using a window of 1000 observations, we present only results of the latter one.

It is necessary to emphasize, that out of the two approaches considered, only time-varying copula approach really allows the copula parameter to be forecasted one time point ahead in the future. In the benchmark model we can only assume that the one-step ahead forecast of the copula parameter is adequately approximated by the copula parameter estimated based on the observations of the previous window of observations.

The back-tests used are the same as in the in-sample back-testing. The results of both Kupiec's Proportion of Failure Test and Christofferson's Markov test for all α s are presented in Table 10. Figures 12, 13 and 14 illustrate the plots of Value-at-Risk forecasts versus the realized losses.

Both tests fail to reject the null hypothesis for both copula approaches and all values of α at confidence level of 5%. At $\alpha = 10\%$ both tests reject their null hypothesis for the time-varying copula approach, whereas in case of the benchmark model only Kupiec's Proportion of Failures test is rejected. As before, the results of the Christoffersen's Markov test are not available for both approaches at $\alpha = 1\%$ as there is no Value-at-Risk estimate violating the actual loss experienced at time point t , followed by another violation at time point $t+1$. Using the same rationale as in the case of in-sample back-testing, this implies that the null hypothesis of the Markov test is rejected. Comparing the results of the two approaches presented in the Table 10, we notice that for $\alpha = 10\%$ and $\alpha = 5\%$ the number of observed violations of the normal copula with time-varying parameter differs more from the number of expected violations than the number of observed violations of the benchmark model and the latter one provides slightly better results. We can thus establish that in case of out-of-sample back-testing, contrary to the results of the in-sample back-testing, benchmark model does a better job estimating the one-day Value-at-Risk at all levels of α .

Table 10: Out-of-Sample Back-testing Results

	VaR_{0.1}	VaR_{0.05}	VaR_{0.01}
Number of expected violations	161	80	16
Time-varying Normal Copula			
Number of observed violations	184	92	16
LR_{UC} test statistic	3.51	1.66	0.00
LR_{UC} p -value	0.06	0.20	0.98
LR_{CC} test statistic	4.92	1.67	n.a.
LR_{CC} p -value	0.09	0.43	n.a.
Benchmark Model			
Number of observed violations	183	90	11
LR_{UC} test statistic	3.21	1.14	1.84
LR_{UC} p -value	0.07	0.29	0.17
LR_{CC} test statistic	3.81	1.34	n.a.
LR_{CC} p -value	0.15	0.51	n.a.

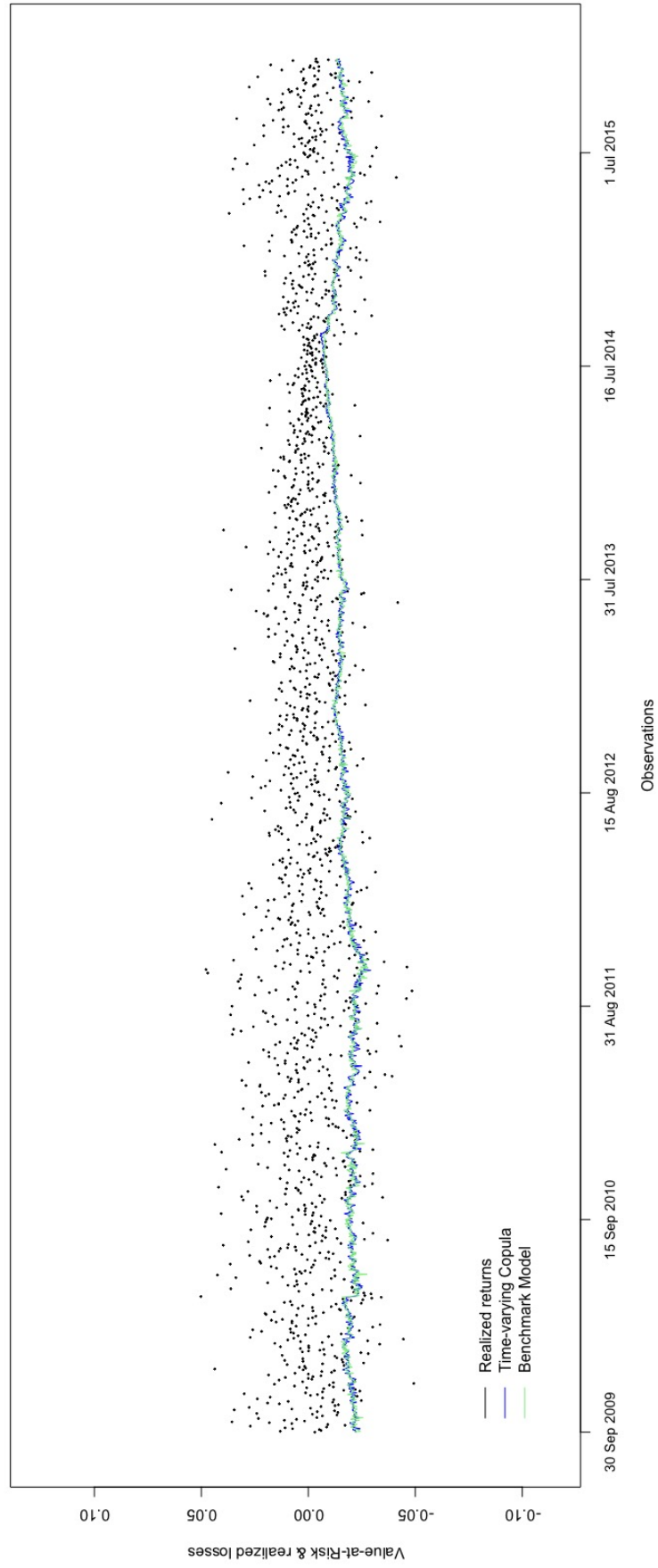


Figure 12: Out-of-Sample Value-at-Risk estimation at $\alpha = 0.10$

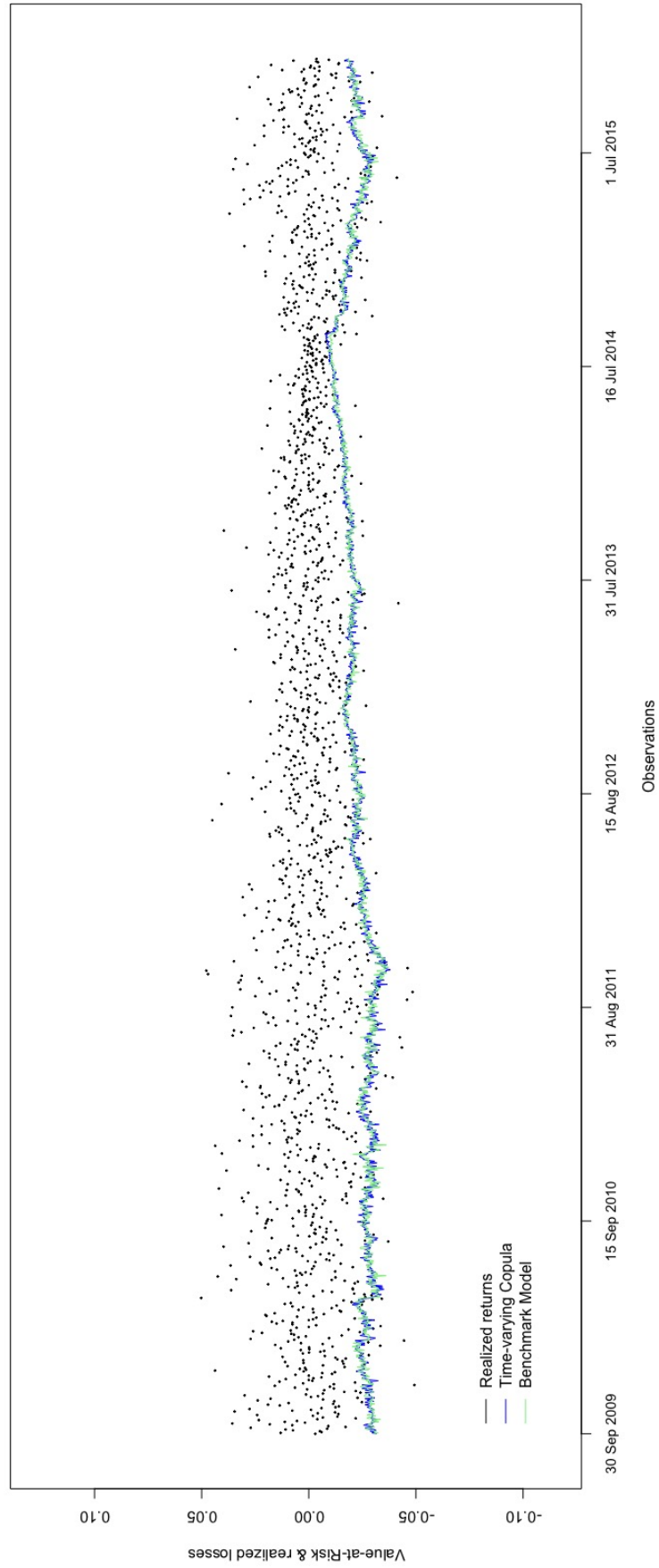


Figure 13: Out-of-Sample Value-at-Risk estimation at $\alpha = 0.05$

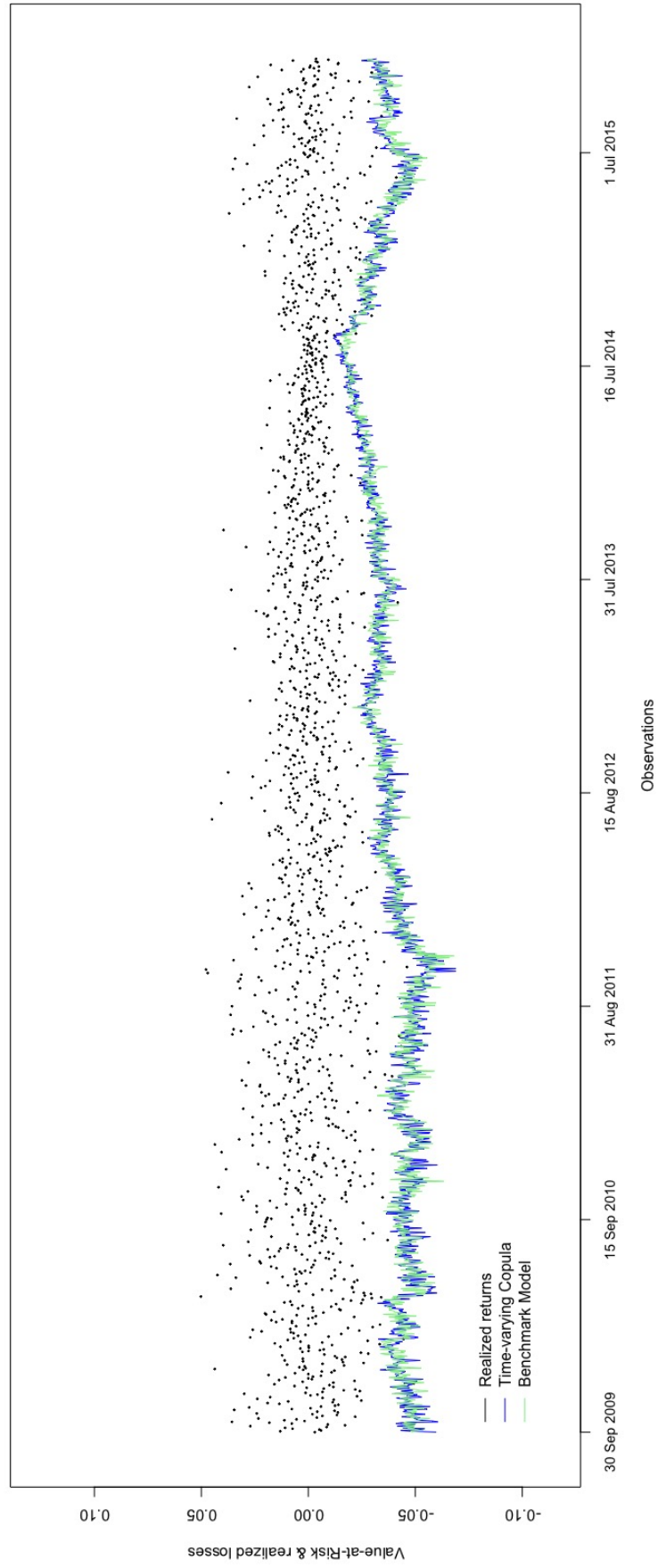


Figure 14: Out-of-Sample Value-at-Risk estimation at $\alpha = 0.01$

6 Conclusion

The popularity of Value-at-Risk in risk management decisions is emphasizing the importance of the accuracy of its estimates. The assumptions that traditional approaches to Value-at-Risk estimation, such as Historical or Variance Covariance approach, make with respect to the distribution of the underlying assets create the need to consider more flexible approaches that allow more liberty in specifying the distributions of the underlying assets and in modeling dependence between them.

In this work we have used the copula theory to estimate Value-at-Risk of a portfolio of two exchange rates. Furthermore, we have allowed for time variation in the copula parameter and compared the results obtained to those resulting from a simple rolling window approach. Based on the analysis of the data we have modeled the marginal distributions of the two exchange rates with AR,*st*-GARCH models while choosing the normal copula for the copula model. In case of the time-varying approach we have assumed the copula parameter follows an evolution equation as defined in Patton (2006b), and estimated the parameters of the model using the multi-stage maximum likelihood estimation. To determine whether consideration of a time-varying copula leads to better Value-at-Risk estimates we have back-tested both approaches using Kupiec's Proportion of Failure test and Christoffersen's Markov test in an in-sample as well as out-of-sample back-test.

Results we have obtained indicate that when estimating the parameters based on the whole data sample, copula approach with a time-varying parameter indeed provides better estimates than the plain rolling window approach. When performing an out-of-sample back-test where we test the predictive power of the models, however, the results of the two back-tests performed show that the benchmark model produces slightly better results than the time-varying copula approach. We therefore conclude that when estimating Value-at-Risk of a portfolio of two exchange rates, especially considering the larger computational effort the time-varying approach requires, allowing for such form of time variation in the copula parameter does not provide better Value-at-Risk estimates than those we obtain using a simple rolling

window approach.

References

- Aeppli, M. D. (2015). *Portfolio risk forecasting-On the predictive power of multivariate dynamic copula models*. Ph. D. thesis, University of St. Gallen.
- Aeppli, M. D., K. Frauendorfer, R. Füss, and F. Paraschiv (2015). Multivariate dynamic copula models: Parameter estimation and forecast evolution. *University of St. Gallen, School of Finance Research Paper* (2015/13).
- Alexander, C. (2009). *Market risk analysis, Value at Risk models*, Volume 4. John Wiley & Sons.
- Andersen, T. G., T. Bollerslev, P. F. Christoffersen, and F. X. Diebold (2006). Volatility and correlation forecasting. *Handbook of economic forecasting* 1, 777–878.
- Ang, A. and J. Chen (2002). Asymmetric correlations of equity portfolios. *Journal of financial Economics* 63(3), 443–494.
- Ausin, M. C. and H. F. Lopes (2010). Time-varying joint distribution through copulas. *Computational Statistics & Data Analysis* 54(11), 2383–2399.
- Bae, K.-H., G. A. Karolyi, and R. M. Stulz (2003). A new approach to measuring financial contagion. *Review of Financial studies* 16(3), 717–763.
- Bauwens, L., S. Laurent, and J. V. Rombouts (2006). Multivariate garch models: a survey. *Journal of applied econometrics* 21(1), 79–109.
- Bekaert, G. and C. R. Harvey (1995). Time-varying world market integration. *The Journal of Finance* 50(2), 403–444.
- Berger, T. (2013). Forecasting value-at-risk using time varying copulas and evt return distributions. *International Economics* 133, 93–106.
- Berger, T. and M. Missonig (2014). Financial crisis, value-at-risk forecasts and the puzzle of dependency modeling. *International Review of Financial Analysis* 33, 33–38.

- Bollerslev, T. (1987). A conditionally heteroskedastic time series model for speculative prices and rates of return. *The review of economics and statistics*, 542–547.
- Campbell, S. D. (2006). A review of backtesting and backtesting procedures. *The Journal of Risk* 9(2), 1.
- Cerrato, M., J. Crosby, M. Kim, and Y. Zhao (2015). Modeling dependence structure and forecasting market risk with dynamic asymmetric copula. *Available at SSRN 2460168*.
- Cherubini, U. and E. Luciano (2001). Value-at-risk trade-off and capital allocation with copulas. *Economic notes* 30(2), 235–256.
- Cherubini, U., E. Luciano, and W. Vecchiato (2004). *Copula methods in finance*. John Wiley & Sons.
- Chollete, L. (2005). Frequent extreme events? a dynamic copula approach. *Unpublished paper, Norwegian School of Economics and Business*.
- Christoffersen, P., V. Errunza, K. Jacobs, and H. Langlois (2012). Is the potential for international diversification disappearing? a dynamic copula approach. *Review of Financial Studies* 25(12), 3711–3751.
- Christoffersen, P. and D. Pelletier (2004). Backtesting value-at-risk: A duration-based approach. *Journal of Financial Econometrics* 2(1), 84–108.
- Christoffersen, P. F. (1998). Evaluating interval forecasts. *International economic review*, 841–862.
- Creal, D., S. J. Koopman, and A. Lucas (2013). Generalized autoregressive score models with applications. *Journal of Applied Econometrics* 28(5), 777–795.
- Denuit, M., J. Dhaene, M. Goovaerts, and R. Kaas (2006). *Actuarial theory for dependent risks: measures, orders and models*. John Wiley & Sons.
- Embrechts, P. and A. Höing (2006). Extreme var scenarios in higher dimensions. *Extremes* 9(3-4), 177–192.

- Embrechts, P., A. Höing, and A. Juri (2003). Using copulae to bound the value-at-risk for functions of dependent risks. *Finance and Stochastics* 7(2), 145–167.
- Embrechts, P., A. McNeil, and D. Straumann (2002). Correlation and dependence in risk management: properties and pitfalls. *Risk management: value at risk and beyond*, 176–223.
- Engle, R. (2002). Dynamic conditional correlation: A simple class of multivariate generalized autoregressive conditional heteroskedasticity models. *Journal of Business & Economic Statistics* 20(3), 339–350.
- Engle, R. F. and K. Sheppard (2001). Theoretical and empirical properties of dynamic conditional correlation multivariate garch. Technical report, National Bureau of Economic Research.
- Erb, C. B., C. R. Harvey, and T. E. Viskanta (1994). Forecasting international equity correlations. *Financial analysts journal* 50(6), 32–45.
- Fengler, M. R. and O. Okhrin (2014). Managing risk with a realized copula parameter. *Computational Statistics & Data Analysis*.
- Fermanian, J.-D., D. Radulovic, M. Wegkamp, et al. (2004). Weak convergence of empirical copula processes. *Bernoulli* 10(5), 847–860.
- Garcia, R. and G. Tsafack (2008). Dependence structure and extreme comovements in international equity and bond markets with portfolio diversification effects. *Working Paper, Risk Asset Management Research Centre*.
- Genest, C. (1987). Frank’s family of bivariate distributions. *Biometrika* 74(3), 549–555.
- Genest, C., J. Nešlehová, and N. Ben Ghorbal (2011). Estimators based on kendall’s tau in multivariate copula models. *Australian & New Zealand Journal of Statistics* 53(2), 157–177.

- Ghoudi, K. and B. Rémillard (2004). Empirical processes based on pseudo-observations. ii. the multivariate case. *Asymptotic methods in stochastics* 44, 381–406.
- Hafner, C. M. and H. Manner (2012). Dynamic stochastic copula models: Estimation, inference and applications. *Journal of Applied Econometrics* 27(2), 269–295.
- Hamao, Y., R. W. Masulis, and V. Ng (1990). Correlations in price changes and volatility across international stock markets. *Review of Financial studies* 3(2), 281–307.
- Hernández-Lobato, J. M., J. R. Lloyd, and D. Hernández-Lobato (2013). Gaussian process conditional copulas with applications to financial time series. In *Advances in Neural Information Processing Systems*, pp. 1736–1744.
- Hull, J. C. and A. D. White (1998). Value at risk when daily changes in market variables are not normally distributed. *The Journal of Derivatives* 5(3), 9–19.
- Ibragimov, R. (2009). Copula-based characterizations for higher order markov processes. *Econometric Theory* 25(03), 819–846.
- Joe, H. (1997). *Multivariate models and multivariate dependence concepts*. CRC Press.
- Joe, H. (2005). Asymptotic efficiency of the two-stage estimation method for copula-based models. *Journal of Multivariate Analysis* 94(2), 401–419.
- Jondeau, E. and M. Rockinger (2006). The copula-garch model of conditional dependencies: An international stock market application. *Journal of international money and finance* 25(5), 827–853.
- Komunjer, I. et al. (2011). Quantile prediction. *Handbook of Economic Forecasting* 2.
- Kroner, K. F. and V. K. Ng (1998). Modeling asymmetric comovements of asset returns. *Review of Financial Studies* 11(4), 817–844.
- Kupiec, P. H. (1995). Techniques for verifying the accuracy of risk measurement models. *The J. of Derivatives* 3(2).

- Lee, T. and X. Long (2005). Copula-based multivariate garch model with uncorrelated dependent errors, university of california. Technical report, Riverside, working paper series.
- Linsmeier, T. J. and N. D. Pearson (2000). Value at risk. *Financial Analysts Journal* 56(2), 47–67.
- Longin, F. and B. Solnik (1995). Is the correlation in international equity returns constant: 1960–1990? *Journal of international money and finance* 14(1), 3–26.
- Longin, F. and B. Solnik (2001). Extreme correlation of international equity markets. *The journal of finance* 56(2), 649–676.
- Manner, H. and J. Segers (2011). Tails of correlation mixtures of elliptical copulas. *Insurance: Mathematics and Economics* 48(1), 153–160.
- McNeil, A. J., R. Frey, and P. Embrechts (2015). *Quantitative risk management: Concepts, techniques and tools*. Princeton university press.
- Mills, F. C. et al. (1927). The behavior of prices. *NBER Books*.
- Min, A. and C. Czado (2010). Bayesian inference for multivariate copulas using pair-copula constructions. *Journal of Financial Econometrics* 8(4), 511–546.
- Morgan, J. (1996). *Riskmetrics: technical document*. Morgan Guaranty Trust Company of New York.
- Nelsen, R. (2006). An introduction to copulas, second edn. springer series in statistics.
- Oh, D. H. and A. J. Patton (2013). Simulated method of moments estimation for copula-based multivariate models. *Journal of the American Statistical Association* 108(502), 689–700.
- Oh, D. H. and A. J. Patton (2016). Time-varying systemic risk: Evidence from adynamic copula model of cds spreads. *Journal of Business & Economic Statistics* (just-accepted), 1–47.

- Okimoto, T. (2008). New evidence of asymmetric dependence structures in international equity markets. *Journal of Financial and Quantitative Analysis* 43(03), 787–815.
- Palaro, H. P. and L. K. Hotta (2006). Using conditional copula to estimate value at risk. *Journal of Data Science* 4, 93–115.
- Patton, A. J. (2006a). Estimation of multivariate models for time series of possibly different lengths. *Journal of applied econometrics* 21(2), 147–173.
- Patton, A. J. (2006b). Modelling asymmetric exchange rate dependence. *International economic review* 47(2), 527–556.
- Patton, A. J. (2009). Copula-based models for financial time series. In *Handbook of financial time series*, pp. 767–785. Springer.
- Patton, A. J. (2012). A review of copula models for economic time series. *Journal of Multivariate Analysis* 110, 4–18.
- Pitt, M., D. Chan, and R. Kohn (2006). Efficient bayesian inference for gaussian copula regression models. *Biometrika* 93(3), 537–554.
- Rémillard, B. (2010). Goodness-of-fit tests for copulas of multivariate time series. *Available at SSRN 1729982*.
- Rodriguez, J. C. (2007). Measuring financial contagion: A copula approach. *Journal of Empirical Finance* 14(3), 401–423.
- Rosenberg, J. V. (2000). Nonparametric pricing of multivariate contingent claims.
- Sancetta, A. and S. Satchell (2004). The bernstein copula and its applications to modeling and approximations of multivariate distributions. *Econometric theory* 20(03), 535–562.
- Scaillet, O. and J.-D. Fermanian (2002). Nonparametric estimation of copulas for time series. *FAME Research paper* (57).

- Sklar, M. (1959). *Fonctions de répartition à n dimensions et leurs marges*. Université Paris 8.
- Smith, M., A. Min, C. Almeida, and C. Czado (2012). Modeling longitudinal data using a pair-copula decomposition of serial dependence. *Journal of the American Statistical Association*.
- Susmel, R. and R. F. Engle (1994). Hourly volatility spillovers between international equity markets. *Journal of International Money and Finance* 13(1), 3–25.
- Tse, Y. K. and A. K. C. Tsui (2002). A multivariate generalized autoregressive conditional heteroscedasticity model with time-varying correlations. *Journal of Business & Economic Statistics* 20(3), 351–362.
- Tsukahara, H. (2005). Semiparametric estimation in copula models. *Canadian Journal of Statistics* 33(3), 357–375.

Ana Petrovič

Contact information +43 699 17380020
Stumpergasse 64/27
1060 Wien
Österreich
ana.petrovic@me.com

Nationality Slovene

Date of birth 7.7.1989

Work experience

Junior Special Exposures Management Process & Project Manager

Raiffiesen Bank International, Am Stadpark 9, 1030 Wien, Österreich
May 2015 onwards

Freelancer at Risk Excellence & Projects (Early Warning System)

Raiffiesen Bank International, Am Stadpark 9, 1030 Wien, Österreich
February 2013 – April 2015

Internship at an actuarial department

Merkur Versicherung, Dunajska 58, 1000 Ljubljana, Slovenia
July 23 2012 – September 21 2012

Assistant at a private firm

Reyan d.o.o., Neubergerjeva 3, 1000 Ljubljana, Slovenia
2007 – 2009

Education

Master of Science in Economics and Business – Quantitative Economics, Management and Finance

October 2012 – October 2016
Faculty of Economics and Business, Universität Wien, Österreich

Level in national or international classification 2nd cycle academic – ISCED 5A

Bachelor of Science in Financial Mathematics

October 1 2009 – July 18 2012
Faculty of Mathematics and Physics, University of Ljubljana, Slovenia

Level in national or international classification 1st cycle academic – ISCED 5A

Matura / high school graduate

	2004 – 2008 Gimnazija Bežigrad, Ljubljana (Bežigrad Grammar School)
Level in national or international classification	ISCED 4A
Languages	
Mother tongue	Slovene
Other languages	English – Full professional proficiency Serbian – Working proficiency German – Limited working proficiency French – Elementary proficiency
Personal skills and competences	
	Reliable, communicative, adaptable, organized, enjoy teamwork, creative.
Computer skills	Excellent competence with Microsoft Office programmes (Word, Excel, PowerPoint) as well as Structured Query Language and iPythonn, good competence with programming languages R and Matlab and document preparation system LaTeX.
	PERSONAL INTERESTS Reading, watching films, cooking. Enjoy nature, meeting new people and learning. Love to travel and experience new cultures.