



UNIVERSITY OF BERGAMO

DEPARTMENT OF MANAGEMENT, ECONOMICS
AND QUANTITATIVE METHODS

Systemic risk measures and contagion models

PhD candidate: Gianluca Farina

Supervisor: Professor Rosella Giacometti

COMPUTATIONAL METHODS FOR FORECASTING
AND DECISIONS IN ECONOMICS AND FINANCE

*Dedicated to my son Gabriele:
the most inquisitive, outspoken,
tenacious researcher I ever met.*

Acknowledgments

The path that lead me here writing these lines has been unusual: it spanned, with interruptions and many twists and turns, more than a decade. It is no surprise then that the list of persons I am very grateful to over these many years is very long as I had a great amount of support along the way.

Starting from my academic career, I was very lucky recently with professor Rosella Giacometti as my supervisor. I would like to thank her for the huge amount of hours she spent reading endless versions of these pages and for her invaluable input, not to mention the emotional support she provided when time pressure was mounting. I am also very grateful to professor Marida Bertocchi for her advice during my PhD years as well as for giving me the opportunity to complete the program. There have been several professors that helped me during my studies: professors Melone, Karatzas, Smirnov have all had huge impact over my choices and passions over the years¹.

My formation was also enriched by my working experiences, both in Italy and in the City of London. I had the privilege to work with many distinguished professionals, some of whom have become close friends. In particular I would like to thank Giuseppe Cinquemani as he "has always believed in me from the very start" (his own words) and few of the old quants at UBS: Bernard Raynaud, Mark Spittle, Georgeos Tsouderos, Bertrand Duret have shared with me not only their invaluable practical knowledge but also many stormy days.

There have been many friends who helped me keeping the right balance in life. I am

¹*There is also a professor of mathematics at University of Naples, Engineering faculty, that has played a key role in my decision of becoming a mathematician rather than an engineer; I would like to thank him even if I don't know his name and he had no chance to know me as I was sitting in a room of 300 students while he was talking about white flies.*

very grateful to have met all of them: Maria and Vitto, Giuppe and Benny, Mario, Mario and Pasqualino, from my days in Caserta. Claudio and Franz for the endless Risk (not the journal!) nights in Milan, Penny and Ken from my days in New York, Toby and James (and the Bath's gang), Lorenzo and Alessandro, Farhad and Aurora, Chris and Hugo and the Gadesky's Saturday group for their friendship on and especially off the British fencing pistes.

My parents have always helped me in countless ways. Without their example and their gentle guidance I would have not become the person I am today. I am also very grateful to my brother Michele for his continuous encouragement.

Last, but not least, I wish to express my gratitude to Claudia. When she was working at the computer center during my first PhD year, she helped me printing a document; somehow, she never stopped supporting me ever since.

Contents

Thesis structure	11
Introduction	13
I Framework and tools	15
1 General setting	17
1.1 The switch from micro to macro prudential approach	18
1.2 The problem of the definition	18
1.3 Possible causes	22
1.3.1 Direct	22
1.3.2 Indirect	23
1.4 Regulators response	24
2 Systemic risk measures	27
2.1 Introduction	28
2.2 Characteristics	29
2.3 Literature review	31
2.3.1 Measures based on portfolio return/losses analysis	32
2.3.2 Network models	34
2.3.3 Econometric indicators	37
2.3.4 Measures based on multivariate default distribution	39
2.4 Conclusive remarks	44
3 Tools	47
3.1 A basic introduction to copulas	48
3.1.1 Definitions	48

3.1.2	Sklar's theorem	50
3.1.3	Archimedean copulas	51
3.1.4	Kendall's τ	54
3.1.5	Calculating joint survival/default probabilities	54
3.2	The one factor Gaussian base correlation approach	57
3.2.1	Introduction	57
3.2.2	Description	57
3.2.3	CDS and CDO pricing	58
3.2.4	Base correlation approach	62
3.2.5	Further readings	63
II	Articles	65
4	CIMDO <i>posterior</i> stability study	67
4.1	Introduction	68
4.2	CIMDO description	68
4.2.1	The distance function	69
4.2.2	The constraints	69
4.2.3	Mathematical formulation and solution	70
4.3	The independence issue	72
4.3.1	Extended independence result	73
4.3.2	Consequences	73
4.4	<i>Posterior</i> stability study	74
4.4.1	The base case	75
4.4.2	Family issues	75
4.4.3	Parameter choice	79
4.4.4	Default probabilities	84
4.5	Conclusive remarks	87
5	A model of infectious defaults with immunization	91
5.1	Introduction	92
5.2	A new model	97
5.3	Theoretical results	99
5.3.1	Marginal distribution	99
5.3.2	Joint default/survival probability	101

5.3.3	Portfolio loss distribution	104
5.4	An application to CDO Pricing	108
5.4.1	A restricted version of the model	109
5.4.2	The portfolio	110
5.4.3	Operational ranges	110
5.4.4	A calibration exercise	114
5.4.5	Single name deltas	115
5.5	Conclusive remarks	118
6	A new measure of systemic risk in contagion models	121
6.1	Introduction	122
6.2	Contagion Losses Ratio measures	122
6.2.1	A thought experiment	122
6.2.2	CLR and CoCLR	125
6.3	Calculating CLR and CoCLR in practice	126
6.3.1	Davis and Lo	127
6.3.2	Sakata, Hisakado and Mori	127
6.3.3	Contagion model introduced in chapter 5	129
6.4	A fictional banking system	133
6.4.1	System specification	133
6.4.2	Systemic risk measures in the base case	134
6.4.3	Conditioned CLR	136
6.4.4	Sensitivity results	139
6.5	Conclusive remarks	143
7	Conclusions	145
III	Appendices	147
A	Selected proofs	149
A.1	CIMDO extended independence theorem	149
A.2	Contagion model - Joint default/survival probability	152
A.3	Contagion model - Recursive formulae for Θ_A and Λ_A	154

B Computational considerations	155
B.1 CIMDO methodology	155
B.1.1 Factorising out μ	157
B.2 Loss distribution in the contagion model introduced in chapter 5	159
B.2.1 Performances	161
C Portfolio spreads	163
List of figures	166
List of tables	169
References	170

Thesis structure

The work is structured over seven chapters divided into two parts.

Part I

Formed by chapters 1 to 3, part I introduces the framework and some of the tools used in the rest of the work.

In particular, the first chapter analyzes some the primary aspects of systemic risk, among which its definition(s) and possible causes. A brief overview of the actions adopted by governments and regulators worldwide to counter balance recent financial turmoils complete this introductory chapter.

Chapter 2 contains a review of the systemic risk measures that have been advanced in recent years. After analyzing some of the features that systemic risk measures can exhibit, we propose a novel categorization of the literature by grouping together works that use similar modeling frameworks. Four broad categories have been identified: measures based on portfolio return/losses analysis, network models, econometric indicators and measures based on multivariate default distribution.

Chapter 3 is formed by two separate components: a review of copulas (with special focus on Archimedean ones) and some background material on the one factor Gaussian model. The chapter is meant as a quick reference in order to facilitate the reading of the second part of the thesis and can be skipped by readers familiar with these topics.

Part II

The second part of the work (chapters 4 to 6) groups few original results and represents the bulk of the thesis.

Chapter 4 is centered around the CIMDO methodology, a framework heavily used for systemic risk purposes. We present theoretical results as well as an applied stability study aimed at identifying the strength of the relationship between the inputs and the output of the methodology.

Chapter 5 approaches systemic risk from a different angle, the modeling one, as we propose a new contagion model. Theoretical results are presented regarding both marginal and joint default distributions together with a recursive algorithm for the calculation of the portfolio loss distribution. An application to the problem of CDO pricing is presented.

In chapter 6 we suggest two new systemic risk measures in the context of contagion models based on attributing portfolio losses to either idiosyncratic or to infection-driven events. In order to test the ability of these measures to offer a fresh perspective when compared to other established methodologies, we present an application to a fictional banking system.

Conclusion and appendix

Chapter 7 concludes the work while in appendix there are few proofs of theoretical results presented, some considerations on the numerical aspects of the work performed and the portfolio data used in chapter 5.

Chapter independence

We tried to arrange the material so that each chapter could be read independently from the others. At the same time, we wanted to minimize any possible repetition. The final result is therefore a compromise between these two opposite wishes. In particular, we suggest to read chapter 6 after chapters 2 and 5. Section 3.1 could instead be useful for chapter 4. Similarly, section 3.2 is suggested before reading chapter 5.

Introduction

The main theme of this thesis is systemic risk measurement intended as the set of methodologies that can be used to assess systemic risk of a given financial network, often belonging to the banking sector². The concept under scrutiny, systemic risk, is actually a very elusive one: several possible definitions have been proposed over the years but a global consensus is yet to be achieved. At the same time, the debates on its possible causes and on which of the many recent crisis should be considered as a systemic event are on full swing.

Among all this fuzziness, there is probably only one aspect of systemic risk over which a worldwide agreement has indeed been reached: its importance in today's financial markets. There is no national or international supervisory agency that doesn't list monitoring and assessing systemic risk as one of its top priorities. For many years it has been evident that the traditional risk-based approach of measuring the health of a system by only looking at firm's level indicators is inadequate: the micro-prudential approach needs to be supplemented with a new macro-prudential one that focuses on the system as a whole and on the interconnections of its members.

Another one of the very few certainties surrounding systemic risk is the belief that quantitative tools will be essential for this type of augmented regulatory approach to succeed. We can identify two fields of research related to measuring systemic risk that rely on highly quantitative components. The first tries to answer the question of what to measure: it involves uncovering some of the system's frailties and identifying clear indicators of systemic risk, hopefully capable of forecasting when a particularly set of critical conditions might occur in the future. The second line of research, instead, answers the question of how to compute the above indicators: it deals with the modeling

²The central role played by banks in systemic risk studies is reflected by the terminology used in many works on the subject where the term *bank* is usually preferred to the generic *firm*.

framework behind the measures as well as the required tools to interpret the data and covers a wide spectrum of sophistication.

The two approaches are actually very entangled to each other: a measure very often relies on the underlying modeling framework to provide the necessary estimates/forecasts³. Conversely, a modeling approach often suggests a 'natural' point of view on systemic issues: for example, when financial systems are modeled as networks, it is quite reasonable to assess systemic risk via network fragility measures.

In the present work we will present original contributions to both lines of research. In fact, in chapter 6 we will introduce a new systemic risk measure in the context of contagion models that belongs, therefore, to the first branch of literature. In chapter 5, instead, we will present a new methodology for modeling default contagion that is therefore linked to the second field of research mentioned earlier.

Publications extracted from the thesis

Two separate articles extracted from this thesis have been submitted to international journals for publication: in particular, an extract of chapter 4 has been submitted to the *Journal of Financial Services Research* while an extract from chapter 5 has been submitted to the *Journal of Banking and Finance*. At the time of writing, we are considering submitting an extract of chapter 6.

Notation caveat

When working with a Merton style model, we will indicate with X_i the i^{th} firm assets and with K_i its default threshold. There is no consensus in literature regarding which side of the assets distribution should be considered the defaulted one⁴. We adapted our notation between different sections in order to facilitate the comparison with the literature on the subject.

³There have already been examples of works that suggest different modeling routes to compute essentially similar risk measures: for example, Brownlees and Engle (2012) shows an alternative way to compute a measure defined by Acharya, Pedersen, Philippon, and Richardson (2010)

⁴For example, most of the works on the Gaussian model assume $P\{i \text{ defaults}\} = P\{X_i \leq K_i\}$ while works in the systemic risk context often assume the opposite: $P\{i \text{ defaults}\} = P\{X_i \geq K_i\}$.

Part I

Framework and tools

Chapter 1

General setting

In this chapter we will briefly introduce some of the most important open problems regarding systemic risk, starting with its definition. We will then analyze possible causes and conclude with a brief historical review of recent regulators efforts in order to give an idea of the importance of the subject.

1.1 The switch from micro to macro prudential approach

Table 1.1 lists just few of the financial crisis and economic turmoils occurred in recent years that, in our view, can be considered as highly representative of systemic events¹. The events listed (together with many others we omitted) had at least one positive consequence: they taught regulators worldwide to rethink the rationale of banking regulation. The traditional micro-prudential approach has been proven inadequate in today's highly interconnected financial markets; the main reason is that only ensuring the soundness of individual firms, as in Basel I and Basel II frameworks, fails to assess the behavior of the system as a whole. Conversely, all of the reported events have a strong, if not predominant, systemic component and show how a different perspective needs to complement the old one.

A macro-prudential approach has indeed been advocated from multiple and prestigious voices² and monitoring systemic risk is today on the agenda of virtually every financial supervisory authority. As noted by Huang, Zhou, and Zhu (2012)

[This new] perspective has become an overwhelming theme in the policy deliberations among legislative committees, bank regulators, and academic researchers.

The debate regarding systemic risk is very far from being an academic exercise as the Financial Stability Board (2009) stated

financial institutions should be subject to requirements commensurate with the risks they pose to the financial system

1.2 The problem of the definition

By looking again at table 1.1, we could rephrase the famous starting lines of Tolstoy's *Anna Karenina* and state that bullish markets are very similar but every financial crisis

¹We did not intend to provide an exhaustive list of systemic events; more detailed information can be found, for example, in the works of Bordo, Eichengreen, Klingebiel, and Martinez-Peria (2001) and Laeven and Valencia (2010).

²Interesting readings at regard are Crockett (2000), Knight (2006), Borio (2003, 2009), Brunnermeier, Goodhart, Persaud, Crockett, and Shin (2009) and Clement (2010).

1997	Asian financial crisis	The crisis started when the Thai currency collapsed. As the crisis spread, most of Southeast Asia and Japan saw slumping currencies, devalued stock markets and other asset prices. See Radelet and Sachs (2000) for more details.
1998	LTCM default	Considerable hedge funds losses that spilled over to the trading floors of both commercial and investment banks. See Rubin, Greenspan, Levitt, and Born (1999) for more information.
2008	Sub-prime mortgage crisis and liquidity crunch	Started from special investments vehicles, it affected financial markets worldwide and lead to the demise of Bear Stearns and Lehman Brothers and to the bail out of AIG. See Brunnermeier (2009), Adrian and Shin (2010) and Williams (2010) for more details.
2008	Icelandic banking crisis	Short-term liquidity and depositors runs were the major causes of the collapse of most of the country's banking system. Kindleberger and Aliber (2011) provide some background on the crisis.
2010	Flash crash	On the 6th of May, the Dow Jones index experienced a sudden downward jump of around 9% in less than 30 minutes. The simultaneous interaction of many high frequency trading algorithms and the high uncertainty of the markets caused by the ongoing Euro debt crisis were later blamed (see Securities, Commission, et al. (2010)).

Table 1.1: Selected systemic events.

occurs in its own way. Indeed there is very little in common between the causes of the crisis listed earlier but there is wide consensus that each represented a systemic shock. This highlights one of the main problems related to measuring systemic risk: the difficulty in finding a universal definition.

Several definitions have been proposed over the recent past; we will review the most illustrative ones starting with the one provided by the G-10 Working Group (2001) that identifies systemic risk as:

the risk that an event will trigger a loss of economic value or confidence in
[...] a substantial portion of the financial system that is serious enough to
[...] have significant adverse effects on the real economy

Another possible definition can be found in the survey of De Bandt and Hartmann (2000):

A systemic crisis can be defined as a systemic event that affects a considerable number of financial institutions or markets in a strong sense, thereby severely impairing the general well-functioning of the financial system. While the special character of banks plays a major role, we stress that systemic risk goes beyond the traditional view of single banks' vulnerability to depositor runs. At the heart of the concept is the notion of contagion, a particularly strong propagation of failures from one institution, market or system to another.

Billio, Getmansky, Lo, and Pelizzon (2012) instead propose the following essential definition:

any set of circumstances that threatens the stability of or public confidence
in the financial system

Adrian and Brunnermeier (2011) center the definition they use on the intermediation:

the risk that the intermediation capacity of the entire financial system is
impaired, with potentially adverse consequences for the supply of credit to
the real economy

Yet another definition (Daula (2010)) puts financial intermediaries as its heart as it threats systemic risk as

financial system instability, potentially catastrophic, caused or exacerbated by idiosyncratic events or conditions in financial intermediaries

A more elaborated definition is instead given by Schwarcz (2008)

the risk that an economic shock such as market or institutional failure triggers (through a panic or otherwise) either the failure of a chain of markets or institutions or a chain of significant losses to financial institutions, resulting in increases in the cost of capital or decreases in its availability, often evidenced by substantial financial-market price volatility

We could have extracted more definitions from published papers and official documents alike but we believe that the ones mentioned should suffice in showing both the variety found in literature and the concepts considered. In the following list we will try to distillate a set of common features underlying most of the definitions proposed:

Structure

Most of the definitions measure systemic risk as the capability of some initial shock to affect the well-functioning or the stability of the financial system. The form of the initial shock is often left unspecified while in some cases it is assumed to be either an idiosyncratic failure of some financial institution or a market driven signal. It is also interesting to note how the equation 'less stable $\Rightarrow$ impaired' is implicitly applied to financial systems by many sources³.

Effect on real economy

Many definitions connote systemic events in the financial world by their ability to affect the real economy, usually via a deterioration of the banking sector that leads to increase of the cost of money or to evaporating credit supplies.

Confidence

Another common feature of systemic events is, according to some of the above definitions, their impact on public confidence. The impact on the market perception (and

³Many systems in nature are extremely dynamic (i.e. they are very far from being stable) and yet optimally functioning. In many cases, it is this very characteristic that makes natural systems resilient to outside changes.

the price depression that usually follows) is sometimes assumed to be more important than the actual losses caused by the initial shock.

Contagion

Many definitions stress the importance of contagion mechanisms to connote systemic events although the nature of the contagion is usually left unspecified. It is worth stressing that the infection effects taken into consideration are not only limited to internal contagion between the same sector or country but both cross-sectors and cross-countries infection channels are usually considered.

1.3 Possible causes

The same fuzzy boundaries encountered in the definition(s) of systemic risk can be found when considering the possible causes for it as several candidates have been proposed. We can summarize the possible sources of systemic risk in two broad categories: direct and indirect ones.

1.3.1 Direct

The direct sources of systemic risk are represented by contractual obligations that link together the components of the system. Among these sources we can list the interbank deposit market that makes banks highly dependent on each other for their short-term funding (see for example Rochet and Tirole (1996) for a discussion in merit). Syndicated loans are another way to create strong contractual links between the participants of the loan. Finally, there is the issue of counter party exposures that affects not only the banking sector but the entire financial world.

In theory, it is possible to estimate these interconnections for any subset of banks and therefore assess the topology of the entire system. In practice, there are several obstacles that make this estimations a prohibitive task, among which the necessity to have full undisclosed access to every bank's books in order to uncover looping mechanisms and the huge number of possible connections as the networks involved are composed of many hundreds of nodes.

1.3.2 Indirect

The category of indirect sources is both wider and less defined than the direct one. It contains all sorts of inter-dependencies between market participants, including effects that lie outside of the financial world. Among the main ones, we can cite the following sources:

Depositors runs This is the classical example of indirect systemic event⁴. It can be described as the case of a bank suffering a sudden rush of withdrawals by depositors. Given the low ratio of liquidity versus deposits, banks are usually unable to sustain similar cash flights and the run leads to the targeted bank's default; this in turn can initiate a cascade effect on other banks (due mostly to the interbank deposit market). See Pedersen (2009) and references therein for more information on this highly studied mechanism while a recent example is the case of Northern Rock in 2007.

Common factors/markets exposures Also known as systematic risk, the common exposures of financial firms to the same risk sources ultimately increases the likelihood of firms co-movements. Under the hypothesis of several market theories, this type of risk can be minimized through diversification but never completely eliminated.

Asset bubbles Financial imbalances can take several months to build up, usually in a low volatility environment, only to unravel in an unexpected and rapid period of high volatility. This is known as the 'volatility paradox' as default clustering is realized in stormy environments but actually prepared by calm periods.

Fire sales Fire sales can be sources of systemic risk in two ways. The first is when a firm suffers considerable losses and is forced to unwind its positions by selling assets at distressed (fire sale) prices. This might force other market participants to revalue their books suffering market-to-market losses, possibly initiating a vicious circle. Similar mechanisms might be triggered by haircut spirals (see Brunnermeier and Pedersen (2009)) and margin requirements (see Gârleanu and Pedersen (2007)). The second way is when companies assume that abnormally cheap prices will be permanent and fail to protect their positions against increases

⁴Indeed the very term systemic event was first used to denote this type of situation.

in the prices. See for example Allen and Gale (2009) and Acharya and Yorulmazer (2007).

Information contagion This broad category groups many different herding behaviors, both from investors and financial firms alike. The base ingredient is usually high market volatility created or inflated by panic and deterioration of market perception about critical conditions. Initially, the term 'information contagion' was only applied to depositors' runs in the banking context; the 2010 Flash crash example shows how the spectrum of possible situation that can arise because of information contagion is actually much wider.

OTC derivatives The evolutions of financial products and their trading over the counter (OTC) has created another source of systemic risk. It is rather illustrative the motivation adduced by US authorities for the bail out of AIG in 2008⁵:

The board determined that, in current circumstances, the disorderly failure of AIG could add to already significant levels of financial market fragility and lead to substantially higher borrowing costs, reduced household wealth, and materially weaker economic performance.

The US authorities concern about market fragility was mostly driven by the credit default swap exposures to AIG and the difficulty in assessing their actual size and ramifications.

Please note how, even with full access to confidential information that is usually granted to financial regulators, indirect effects are much harder to quantify than direct ones and approaches capable of fully including their impact are very rare.

1.4 Regulators response

In the autumn of 2008 the global financial crises reached its peak. In the following months, regulators worldwide intensified their efforts to stabilize banking sectors and, among other initiatives, they performed a series of stress tests exercises. One in particular, known as SCAP and dedicated to the US banking industry, has almost become a

⁵On 16 September 2008 the US authorities announced that they would take the unprecedented step of offering emergency financial support to AIG, a large insurance conglomerate.

benchmark among researchers on systemic risk: many papers proposing new measures design the data sets used in their applications to coincide with the SCAP ones in order to compare the behavior of the new measure with SCAP results. It is worth then briefly describing the SCAP test and similar stress tests performed by supervisory authorities in the UK and in the EU zone in recent years⁶. The global discussions on the supervisory roles during distress times are very interesting: further readings we would suggest are G-20 Working Group (2009), Tarullo (2010) and Foglia (2008).

US

In February 2009, the US government announced a two-month long stress test to be performed on the 19 largest bank holding companies with the aim of assessing capital requirements (see Federal Reserve (2009b) for details). The initiative, conducted by the Federal Reserve, was known as the Supervisory Capital Assessment Program, SCAP for brevity. The test was based on two different scenarios, a baseline one and a more pessimistic one, that differ on macroeconomics forecasts (including real GDP, unemployment rate and house prices index) for the next two years. The goal was to establish how much additional liquidity buffer each company needed in order to face worsening economical conditions.

The results of the exercise were published in May 2009: 10 companies were required to raise a total of \$74.6 billion in capital. Although there were some critics, the test was overall considered reliable and the response of the markets were positive. More details about this exercise and other successive refinements can be found in Federal Reserve (2009b,a, 2011, 2012, 2013).

UK

The Financial Services Authority (FSA) has also been very active during the crisis and set guidelines for stress testing that are similar to the SCAP ones and that are based on the following three components (see Financial Services Authority (2008)):

- Firms own stress testing implemented in order to assess in a robust and effective way the firm's capital and liquidity requirements in stressed conditions.

⁶Similar exercises were conducted also in South-East Asia: a Japanese study of 2012 conducted by the International Monetary Fund can be found in Das, Blancher, and Arslanalp (2012) while the China Banking Regulatory Commission (CBRC), instead, held in 2012 a Large Bank Supervisory Work Conference.

- Supervisory stress testing of specific high impact firms performed on a regular basis.
- Simultaneous system-wide stress testing undertaken by firms using a common scenario for financial stability purposes.

An addition proposed by the Financial Services Authority (2009) is the concept of reverse-stress testing, i.e. the idea that companies should design their stress testing environment starting from the risks they see more likely to cause their default. For more information on the FSA approach, Turner et al. (2009) is an ideal reading.

It is also worth mentioning the approach of Bank of England, best embodied by Aikman et al. (2009); their RAMSI (Risk Assessment Model for Systemic Institutions) framework is built on a modular approach where different macroeconomic shocks and scenarios can be fed to different balance sheet based models. Feedback effects (via fire sales) and liquidity risk are also included.

EU

The EU European authorities faced even a tougher scenario than the US and UK ones for at least two reasons. The first was the problem of working on an (imperfectly) integrated banking system but with national based regulations. The second was the sovereign debt crises that increased markets instabilities. In this delicate environment, European Union-wide banking stress tests were conducted by the Committee of European Banking Supervisors every year since 2009 as mandated by the Council of the European Union via the ECOFIN. See European Banking Authority (2009, 2010, 2011) and the report from De Larosière (2009) for more info while Cardinali and Nordmark (2011) provides some insights on the informative power of such tests.

The European Central Bank (ECB), similarly to the RAMSI effort we saw from BoE, published several theoretical studies on the subject of systemic risk among which European Central Bank (2009, 2010).

Chapter 2

Systemic risk measures

The chapter is split in two main sections: the first describes some features useful in order to compare systemic risk measures. In the second part of the chapter we present a literature review organized according to the modeling framework underlying the proposed indicators. Given the amplitude of the corpus of resources on the topic, our review is clearly partial; in particular we focused more on the measures we extensively relied upon in the rest of the thesis. More detailed reviews on the subject are the works by Markeloff, Warner, and Wollin (2011, 2012), De Bandt and Hartmann (2000) and Malz (2013).

2.1 Introduction

The literature on systemic risk measures has quickly grown over the past decade, partly in response to recent financial crises and the need for risk measurement approaches that take into consideration the system as a whole. Given the ambiguity of the definition(s) and the ample range of possible causes, it is no surprise that a wide range of risk measures have been suggested over the last decade. Figure 2.1 shows the number of papers with the words 'systemic risk' in the title per year according to Google scholar service and provides indirect evidence on the importance of the topic.

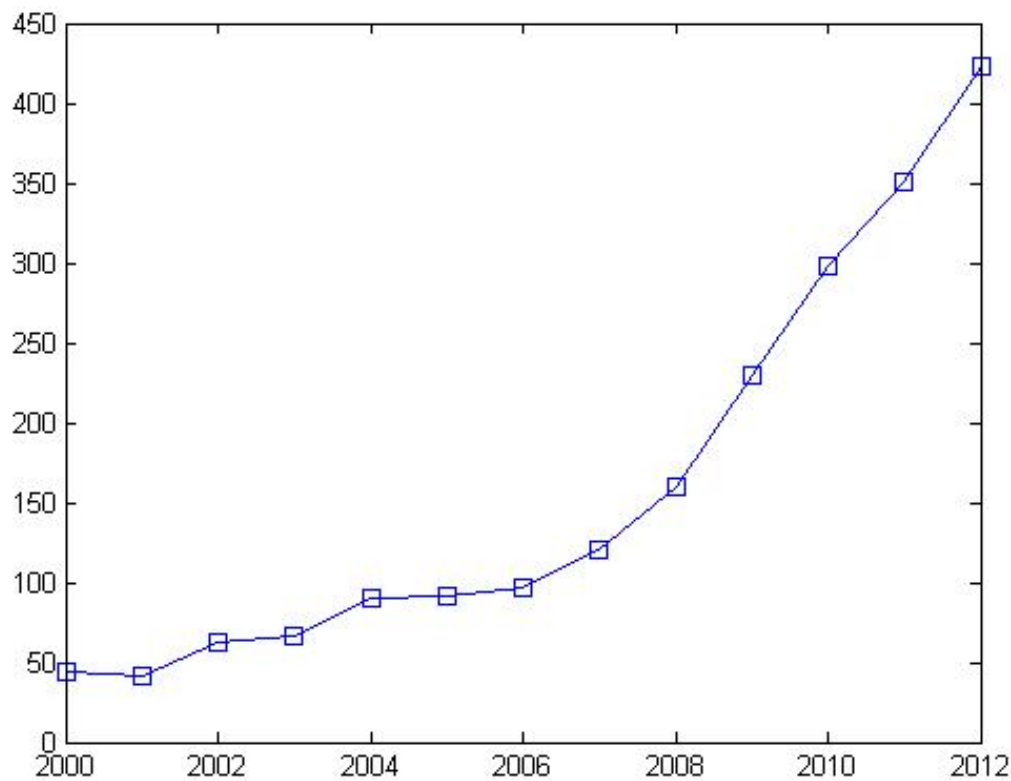


Figure 2.1: Number of papers with the words 'systemic risk' in the title per year according to Google scholar service.

2.2 Characteristics

There are many criteria that can be used in order to differentiate between the many measures proposed. Some of the most important ones are reported in the following subsections.

System vs. individual measures

Most measures are designed to assess the systemic risk level for the entire system under consideration. There are, however, measures that can give more granular information by providing insights on how much a given entity is linked to the rest of the system. In particular, there are two complementary approaches that can be considered. Some indicators try to assess how much a given entity contributes to generate or sustain a global crisis; other measures instead aim at estimating how much a specific company is exposed to a global crisis.

Finally, there are measures that can give information on any given pair of entities in the system, providing estimates for the potential impact of the default of firm i on another company j . In this case, it is useful to distinguish between symmetric and asymmetric measures: the latter can reach more flexibility by allowing firm's i default to impact firm j in a different way than the impact of j on i .

Data

Another important way to look at measures is to consider the type of data they are built upon. Some indicators rely on statistical/econometric tools, usually applied to historical series regarding economic fundamentals. Other measures are based on balance sheet information, accounting variables or confidential data regarding the internal state of the system.¹ The main advantage of this type of data is, of course, their granularity, especially in the case of supervisory studies.

Finally, some studies rely on market observables (the most common being CDS prices, bonds yields, options data). There are several advantages of using this type of data that come from market efficiency hypothesis; in particular, market prices (see Tarashev, Borio, and Tsatsaronis (2009) for a detailed discussion):

¹Many applied studies on national banking systems have been carried out by financial regulators; often such agencies have access to private data among which: capital structure, asset composition, banks lending portfolios.

- summarize the considered opinion of market participants based on the information at their disposal.
- reflect market participants views of all potential sources of risk.
- are easily available on a timely basis.
- are forward looking, i.e. they encompass market participants expectations about future events.

On the negative side, they are more exposed to model uncertainty problem (prices require a model in order to be interpreted) and might not be available for all types of institutions.

Additivity

As noted by Brunnermeier, Goodhart, Persaud, Crockett, and Shin (2009), systemic risk measure should be able to identify both of the following types of risk sources:

Individually systemic institutions This source of risk concerns those financial institutions that are so interconnected and large that they can cause negative risk spillover effects on the rest of the system².

Herding behavior This source of risk deals with institutions that behave as part of an herd. Adrian and Brunnermeier (2011) appropriately summarized this effect noting that

...a group of 100 institutions that act like clones can be as precarious and dangerous to the system as the large merged identity...

In some cases, a system-wide measure can be decomposed into single components, i.e. the systemic risk of the entire system is the sum of the systemic risks of each component. This additive property is extremely useful when seen from regulators point of view: although systemic risk should be measured globally, the ability to allocate such measures to the individual firms is central in order to design efficient and transparent Pigouvian taxing mechanisms.

Tarashev, Borio, and Tsatsaronis (2009, 2010) propose a very elegant approach to the

²Often two acronyms are found in the press: TBTF stands for too big to fail while TITF stands for too interconnected to fail.

problem of attributing a system wide measure to each component based on a game's theory tool³. In cooperative games, the attribution problem has been tackled in order to measure the importance of each player in the game. The attribution methodology proposed by Shapley (see Shapley (1952)) gives each player a value equal to the average of the player's marginal contributions to the value created by all possible subsets of players. In other words, it calculates the marginal increase in the output that the player can bring to each possible coalition. The concept is naturally translated to risk attribution by Tarashev, Borio, and Tsatsaronis obtaining an attribution methodology that benefits from many of the theoretical advantages of the Shapley value, among which additivity and fairness (in the sense that the incremental risk created by the interaction of two institutions is split equally between them). The Shapley value can be calculated starting from the characteristic function of the game⁴; this methodology shows also a linear property as the Shapley value of the sum of two characteristic functions is the sum of the Shapley values calculated on each function separately. This allows the attribution methodology to be applied to any linear combination of systemic risk measures.

2.3 Literature review

Works proposing systemic risk measures usually consist of two steps; the first is to identify what to measure while the second answers the question of how to measure it. Not only the second step is as important as the first one (there is little value of a risk measure that cannot be computed in any way), but often it is the driving force behind the choice of the modeling frameworks used. We have hence organized the literature review by grouping systemic risk measures according to the tools that have been suggested to compute them. This categorization allows for a fairer and more direct comparison between indicators that rely on similar modeling assumptions.

In particular we identified four main categories:

- Measures based on portfolio return/losses analysis.

³Other approaches to the same problem are the ones by Lütkebohmert and Gordy (2007) and Koyluoglu and Stoker (2002).

⁴This is the function θ that assigns an output value to any coalition, i.e. to any possible subset of the given system. Not every function can be used as characteristic function but there are few properties it needs to satisfy, in particular it should be super-additive ($\theta(S \cup T) \geq \theta(S) + \theta(T)$ for $S \cap T = \emptyset$), monotonic ($\theta(S) \geq \theta(T)$ for $S \supseteq T$) and such that $\theta(\emptyset) = 0$.

- Network models.
- Econometric indicators.
- Measures based on multivariate default distribution.

Please note that the field has yet to reach a full maturity and therefore there is still plenty of room for unorthodox approaches that, although completely legitimate, are difficult to classify under common umbrellas. We have listed some of them in closing section 2.4.

2.3.1 Measures based on portfolio return/losses analysis

CoVar

Probably the best known branch of the literature is based on considering financial systems as portfolios of assets; systemic risk is hence measured as extreme and/or conditional scenarios. The most well known measure in this set is the Δ CoVaR proposed by Adrian and Brunnermeier (2011) and extended recently to a multivariate version by Cao (2012). Its emphasis is on analyzing the returns of the entire system conditioned on a given entity being in distress. In a slightly different version, it can also measure how a particular firm is exposed to a global crisis. Let X^i be the return of a financial firm and X^S the return of the system, while let Z, Y be generic variables. The starting point of the *CoVaR* approach is the quantity $CoVaR_q(Z|C)$, implicitly defined via the following equation⁵

$$P\{Z \leq CoVaR_q(Z|C)|C\} = q$$

where C is an event. In particular, the authors define the $\Delta CoVaR_q(Y, Z)$ as the difference between the $CoVaR_q(Y|Z \text{ in distress})$ and $CoVaR_q(Y|Z \text{ at normal state})$. By setting $(Y, Z) = (X^S, X^i)$ or $(Y, Z) = (X^i, X^S)$ we get two complementary measures:

- $\Delta CoVaR_q(X^S, X^i)$. The emphasis is on analyzing the entire system return conditioned on i^{th} entity state; this measure helps attributing a global distress state to a particular firm.

⁵This definition mimics the well-known $VaR_q(Z)$ definition, also an implicit one:

$$P\{Z \leq VaR_q(Z)\} = q$$

- Exposure $\Delta CoVaR_q(X^i, X^S)$. We now study the entity i^{th} conditioned on entire system state and measure of much firm i^{th} is exposed to a global crisis.

In their study Adrian and Brunnermeier (2011) also showed how the *CoVaR* gives additional info regarding riskiness of firms when compared against the *VaR* on its own by applying it to a selections of global firms with data for 2006, Q4. Finally, they provide several methods for estimating it, the simplest one being quantile regression on the returns distributions.

The recent extension of the *CoVaR* introduced by Cao (2012) lead to an indicator called Multi-*CoVaR*, where the conditional expectations are taken assuming that multiple firms are in distress. At the same time, the author defined distress and normal state in a more sophisticated way.

MES, SES and SRISK

Acharya, Pedersen, Philippon, and Richardson (2010) split the system return into its constituents and define the marginal expected shortfall MES, a risk measure based on the expected shortfall of the entire portfolio. In formulas, they assumed $X^S = \sum w_i \cdot X_i$, where the weights w_i are assumed constant in time. The measure proposed, the marginal expected shortfall *MES*, is based on the expected shortfall *ES*:

$$ES = \sum w_i \cdot E(X^i | X^S \leq VaR^S)$$

$$MES_i = \frac{\partial ES}{\partial w_i} = E(X^i | X^S \leq VaR^S)$$

In the same work, the authors then extend this approach by considering a two-step time model, $t \in \{0, 1\}$. They assume that each bank has an initial balance sheet constraint of

$$w_0^i + b^i = a^i$$

where a^i are bank i assets, b^i represents its debt and w_0^i is the value of its equity at time $t = 0$. When the time moves from 0 to 1, the only quantity that changes is the equity value that becomes w_1^i . The authors then define the systemic expected shortfall of bank i *SES*^{*i*} as

$$SES^i = E[z \cdot a^i - w_1^i | W_1 < z \cdot A]$$

where W_1 is the system aggregated equity value at time $t = 1$, A represents the system assets and z is a minimum ratio equity/assets during normal times (leverage).

The work by Brownlees and Engle (2012) can be considered an extension of the previous one as it is also based on the *MES*. The authors derive a measure called *SRISK* that is similar to the *SES*; the main difference relies in the modeling framework suggested to calculate it. Both the market and the individual firm returns are modeled via correlated stochastic processes⁶. The authors provide expressions for *MES* that can be used to calibrate the model while the *SRISK* indicator is obtained via numerical simulations.

DIP

The distress insurance premium DIP (Huang, Zhou, and Zhu (2012)) is based on the system losses $L = \sum_i L_i$ rather than returns:

$$DIP = E[L|L \geq L_m]$$

where L_m is a threshold for systemic distress. They too provide a way to allocate the system risk down to its constituents by taking derivatives with respect to L_i :

$$\frac{\partial DIP}{\partial L_i} = E[L_i|L \geq L_m]$$

The above indicator is highly dependent on the model required to estimate L ; the authors show one possible application by estimating the marginal default probabilities via CDS information and then calculating the correlation between banks from equity data.

2.3.2 Network models

Another branch of literature models financial entities as connected nodes of directed graphs; systemic risk is then measured by studying the network fragility or robustness. In particular, indicators as the number of links in and/or out of a given node, the average length of paths connecting random or given nodes and various clustering coefficients have all been used to measure systemic risk.

⁶The correlation is estimated via the dynamic conditional correlation (DCC) tool (see Engle (2002) and Engle (2009)). The volatilities are instead modeled via a threshold autoregressive conditional heteroskedasticity approach (TARCH, see Rabemananjara and Zakoïan (1993))

Granger causality

Billio, Getmansky, Lo, and Pelizzon (2012) build a network using Granger causality⁷ (both in its linear and non-linear version) on indices and equities leading to interesting results on the interaction between the shadow and real banking systems. The authors apply this concept to weight the level of inter-connection between the top 25 entities of each of four sectors considered: banks, brokers, insurance companies and hedge funds. They also performed a sector based analysis by analyzing the Granger causality between the four sectors and defining a systemic risk indicator as the ratio of the observed connections over the total number of possible connections. The conclusions they reached are very interesting as their study shows that the so-called shadow banking system (hedge funds) cannot be held responsible for spreading infective shocks to the wider economy but rather the opposite (i.e. hedge funds returns were affected by shocks emerging from banks, insurance companies and brokers).

SSR

Another interesting study is the one by Markose, Giansante, Gatkowski, and Shaghaghi (2010) that uses CRT (credit risk transfers products, mostly CDSs and credit notes) exposures data for 26 US banks in 2008, Q4. They showed that the network of CDS exposures was particularly prone to contagion by looking at the nodes distribution versus simulated graphs with similar connectivity⁸. For each entity, they also introduced the Systemic Risk Ratio (SSR) as the fraction of the losses caused by the entity's default over the total size of the system.

Banking networks

The work of Nier, Yang, Yorulmazer, and Alentorn (2007) aims at understanding how the structure of networks affects system stability. The authors simulate networks of banks via random graphs⁹ and then apply perturbations to single banks. One of the most interesting results of this study is the effect of both capitalization and network

⁷Given two series, X and Y , X is said to Granger cause Y when statistical analysis shows that the past history of X has some predictive power over values of Y above the past history of Y itself.

⁸They also calculated the May-Wigner stability measure, an indicator of stability of networks suggested by May (1972) (for more details and examples, see May (2001)).

⁹Given a set of N nodes, a random graph is usually built assigning randomly links between the nodes with a given probability p . The parameter p is called the Erdos-Renyi probability after the name of the authors who first explored the properties of such graphs (see Erdos and Renyi (1959)).

connectivity on the system stability.

Espinosa-Vega and Solé (2011) perform several scenario analysis on European banking systems and consider three types of scenarios¹⁰:

1. The first scenario is driven by simple credit shock, i.e. a single bank failure on some of its inter-banking obligations.
2. The second type of scenario includes liquidity effects together with interbank defaults. Liquidity squeeze is assumed to limit the ability of banks to raise new capital by selling assets, forcing the bank into a fire sale.
3. Finally, the third set of analysis adds another level of complexity by including contingent claims that cause counter-party risks.

The authors perform the three above types of scenarios by assuming the default on the entire banking system of a given European country and then analyze the domino effects on other national banking systems. The study is based on Bank for International Settlements (BIS) data on cross-country bilateral exposures.

Similar is the work of Hale (2012) but the author goes one step deeper in terms of granularity of the data: the network built in fact is based on interbank lending data from a vast database of international syndicated bank loans¹¹ and can therefore drill down to the individual institution (as opposed to grouping them by country).

Cross border studies

Haldane (2009) instead uses a study on cross-border networks to analyze the nodes degree distribution (the degree of a node is simply the number of connections passing through it) concluding that the vast majority of nodes have a low number of links but few nodes (called super spreaders) have a huge number of connections. This finding supports the assert that financial systems are becoming robust yet fragile: robust as eliminating randomly one node there is good chance that the system will be mostly unaffected (as most of the nodes have a small number of links) but at the same time fragile as a big shock is caused by the elimination of a super spreader.

Another study that is based on cross-border banking data is the one by Minoiu and Reyes (2013). In this study, that covers 184 countries for a period spanning over three

¹⁰A rare example of a network framework incorporating indirect systemic risk features.

¹¹Almost 8000 institutions are included from 141 different countries.

decades (1978-2010), the authors divide the network into two parts: the core (consisting of 15 countries with advanced economies) and the periphery. One of the interesting effects uncovered is the relative size of the capital's flows between the two parts: the volumes of lending internal to the core is ten times bigger than the flows from the core to the periphery. The implications for systemic risk are similar to the ones seen for Haldane (2009).

2.3.3 Econometric indicators

There are several studies that propose econometric or statistical indicators as warning signals for systemic risk crises. We will review just the most representative works in this category and add a list of further readings for this vast branch of the literature.

Co-Risk

In the Co-Risk approach (Chan-Lau, Espinosa, Giesecke, and Solé (2009)) the authors use pairwise quantile regression on CDS data:

$$CDS_i = \beta_j CDS_j + \sum_k \beta_k Y_k$$

where the vector of macro-economical variables Y captures several effects, including

- **Economy wide default risk:** captured by the difference between the daily returns on the S&P 500 index and the three-month US treasury rate.
- **Business cycle:** accounted for via the slope of the US yield curve.
- **Interbank default:** measured as the one-year LIBOR spread over the one-year constant maturity US treasury yield.
- **Liquidity:** the yield spread between the three-month collateral repo rate and the three-month US treasury rate.
- **Risk appetite:** approximated via the implied volatility index VIX.

Strong linear dependency is a symptom of a concentrated system, one that is more at risk of systemic events.

PCA analysis

The principal component analysis (PCA) methodology is a statistical tool based on eigenvalue decomposition of the covariance matrix of time series. The aim is to attribute the total variance of the system to each eigenvector; a system where few eigenvectors are responsible for most of the variance is seen as carrying an high degree of systemic risk.

We have already mentioned the paper by Billio, Getmansky, Lo, and Pelizzon (2012) in relationship to network approaches. In the same paper a PCA analysis is conducted on the same data set. The study is aimed at measuring systemic risk by looking at how much system variance can be explained by few factors.

The work by Kritzman, Li, Page, and Rigobon (2010) also relies on PCA techniques and a new measure, called absorption ratio (AR), is introduced. The AR is defined as

$$AR = \frac{\sum_{i=1}^5 \sigma_i^2}{\sum_{j=1}^M \sigma_j^2}$$

where M is the covariance matrix rank and $\sigma_1, \dots, \sigma_M$ are its eigenvalues in order of decreasing magnitude. The sum on the numerator is limited to the top 5 components.

Dynamic factor

Schwaab, Koopman, and Lucas (2011) is a very good example of an approach based on a sophisticated econometric framework. The mixed-measurement dynamic factor model (MM-DFM, Koopman, Lucas, and Schwaab (2010)) is an econometric tool where factors include macro, regional and industry specific frailties and it is calibrated via simulated maximum likelihood. The authors relied on a considerable credit data set (covering more than 12.000 firms) in order to study the behavior of three broad geographical areas: the US, the EU and the remaining countries. They also proposed a set of indicators to either measure the current level of systemic risk (thermometers) or to forecast impending crisis (crystal balls).

Further readings

- Warning indicators for banking crises from statistical analysis of historical correlations: Borio and Drehmann (2009), Borio and Lowe (2004), Drehmann, Borio, and Tsatsaronis (2011).

- Bubbles and system fragility: Alessi and Detken (2009), Tymoigne (2011).

2.3.4 Measures based on multivariate default distribution

Another line of research uses multivariate distribution of defaults and then simply counts defaults either unconditionally or based on some extreme event.

DiDe

Segoviano and Goodhart (2009) define the Distress Dependency (DiDe) matrix as a tool to monitor the impact of one firm's distress on each other component of the system. Its (i, j) entry represents the probability that bank i experiences distress conditioned on firm j default:

$$DiDe_{\{i,j\}} = \text{probability } i \text{ defaults given } j \text{ default}$$

It is also possible to compare entities by looking at entire column or row: the information contained in row i will give the exposure of firm i to the rest of the system. A single column j will instead show the impact of default of entity j on the rest of the system.

JPoD

The joint probability of default has been suggested by Segoviano and Goodhart (2009). It simply monitors the joint probability that all the firms in the portfolio experience distress at the same time. In formulas, we can write it as ¹²:

$$JPoD = P\{X_1 \geq K_1, \dots, X_n \geq K_n\}$$

The measure is very intuitive and quite simple to calculate under many modeling frameworks used to retrieve the multivariate defaults distribution.

Radev (2012b) extended this measure into a conditional version called ΔCoJPoD (an approach similar to the already mentioned ΔCoVaR by Adrian and Brunnermeier (2011)). The starting point is the definition of the CoJPoD_h as the probability that

¹²In terms of the joint probability density function of the assets $r(\mathbf{x})$ (where $\mathbf{x} = x_1, \dots, x_n$) we can write it as:

$$JPoD = \int_{K_1}^{\infty} \dots \int_{K_n}^{\infty} r(\mathbf{x}) d\mathbf{x}$$

the entire system is in distress given that the firm h has defaulted:

$$CoJPoD_h = \frac{JPoD}{P\{X_h \geq K_h\}} \quad (2.1)$$

The author then suggests to compare the risk of the system when entity h is included and defaults, to the situation where entity h is excluded via the following:

$$\Delta CoJPoD_h = CoJPoD_h - JPoD_h$$

where $JPoD_h$ represents the $JPoD$ measure on the system without name h . The measure can be viewed hence as the difference between the effects of default on systemic fragility when the system is dependent or independent of the given entity h .

BSI

The banking stability index (also known simply as Stability Index in some papers) was initially developed as a conditional expectation of default probability measure by Huang (1992) and then applied in the systemic risk context by Segoviano and Goodhart (2009). It is the expected number of defaults given that at least one firm has defaulted; the lower limit of the measure is 1 (implying a very weak banking linkage). It can be computed via the following formula:

$$BSI = \frac{\sum_{i=1}^n P\{X_i \geq K_i\}}{1 - P\{X_1 < K_1, \dots, X_n < K_n\}}$$

SFM

Introduced by Radev (2012a) and based on a work by Avesani and Garcia Pascual (2006), the systemic fragility measure SFM represents the probability of having at least two defaults in the system in the same time window:

$$SFM = P\{X_i \geq K_i, X_j \geq K_j, i \neq j\}$$

PAO

The Probability of at least One Default, instead, is another conditional measure of the importance of a given entity j :

$$PAO_j = \text{probability at least another entity defaults given } j \text{ default}$$

This name specific measure can be extended into a system wide indicator:

$$\text{system } PAO = \text{probability at least another entity defaults given one default occurred}$$

It can be calculated as:

$$\text{system } PAO = \frac{P\{N^d \geq 2\}}{P\{N^d \geq 1\}}$$

where we used the symbol N^d for the variable that counts the defaults occurred in the system. It is very similar to the *BSI* measure (indeed their calculations have the same denominator) with one crucial difference: the *BSI* counts the defaults occurred distinguishing between scenarios that have different number of defaults. The *system PAO* instead only measures the probability mass of observing 2 or more defaults given that at least one occurred. The difference, in spirit, is similar to the distinction between *VaR* and expected shortfall. In the reminder of the thesis we will simply use the form *PAO* instead of *system PAO* when there is no ambiguity.

Radev (2012a) introduces a generalization of the PAO measure, the probability that at least $N-1$ entities default give a particular firm defaults (PNED)¹³:

$$PNED^n(i) = P\{\text{at least } n \text{ entities default} \mid \text{entity } i \text{ defaulted}\}$$

A pictorial comparison

In this section we will provide a pictorial comparison on four measures we heavily relied upon in the rest of the thesis (in particular in chapter 4): *JPoD*, *PAO*, *BSI* and *SFM*. The risk measures mentioned above sample the distributional space $\mathbb{R}^n$ across different regions. For example, while the *JPoD* focuses on the extreme corner in which every entity defaults, the other three take in considerations integrals of the joint density over a wider area. In order to gain some intuition on the generic case, we can take a closer

¹³The PNED notation is ours; Radev (2012a) uses PNBD for banking systems and PNSD for sovereigns.

look in 2 and 3-dimensions by plotting the area explored by the four measures. Even if the *BSI* and the *PAO* use the same areas/volumes of the probability space, we decided to plot them separately for consistency. Figure 2.2 refers to the bi-dimensional case¹⁴. We can notice how the *SFM* (based on the probability of having 2 defaults) coincides with the *JPoD*. As the *BSI* is a conditional measure, its value can be obtained as the ratio of 2 probabilities: the numerator is represented by the sum of the areas of the 2 rectangles (yellow and pale green) in the upper central plot while the denominator is represented by the red rectangle in the bottom central plot. Please note that the denominator is equal to 1 minus the area of the union of the areas summed up in the numerator. Similar considerations can be made for the *PAO*.

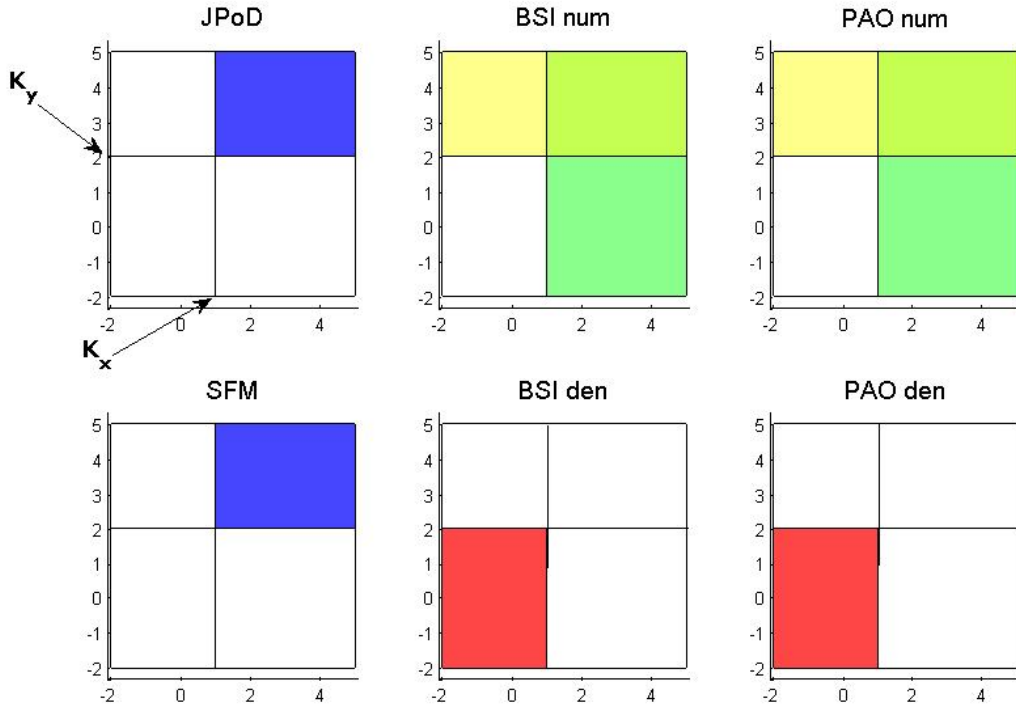


Figure 2.2: Comparison of the area exploited by different risk measures when $n = 2$. For the first plot, on each axis (representing the values of a firm's asset variable), the default threshold K is highlighted.

Moving to the 3-dimensional case, we can now appreciate the difference between *JPoD* and *SFM* as the latter is based on a wider volume. The *BSI*, again, is more

¹⁴Of course where we displayed rectangles the reader should instead imagine infinite areas reaching $\pm\infty$ according to the region under consideration.

complex as it is based on the ratio of the sum of the volumes of the 3 parallelepipeds (yellow, pale green and cyan) in the bottom central plot over the area of the red parallelepiped in the bottom central plot. This complexity means that is not always easy to predict the change in the BSI measure when moving to a distribution with stronger tail dependency. Again, note how *BSI*'s final value is driven by the subtle difference between the volume of the union of the 3 parallelepipeds (denominator) versus the sum of their volumes (numerator). Of course, a similar consideration can be made for the *PAO*.

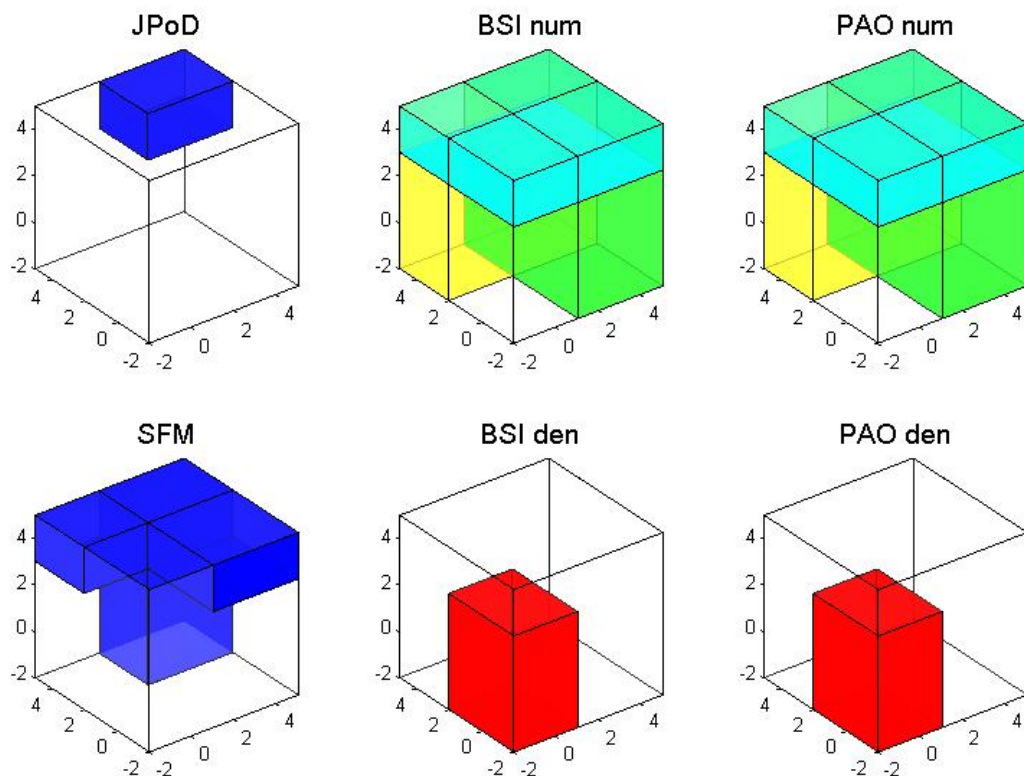


Figure 2.3: Comparison of the area exploited by different risk measures when $n = 3$. The values at which each cube is defined represent the default thresholds for each of the 3 names under consideration.

2.4 Conclusive remarks

We will conclude this chapter devoted to literature review by looking at two methodologies that we excluded from the categorization presented so far. The main reasons for their exclusion is that we deemed their application to systemic risk not sufficiently developed¹⁵. By including these approaches in this final section, we hope readers will be able to appreciate the heterogeneity of the literature and gain some insights of possible directions that future research on systemic risk might explore.

Agent based model

Thurner (2011) introduce an agent-based model where there is only one type of asset traded by three kinds of investors also known as agents:

- **Noise traders:** collectively represented by a unique instance, traders actually drive the asset price.
- **Hedge funds:** there are several hedge funds that differ from each other by their amount of aggressiveness measured by their leverage. Hedge funds react to the asset price by trying to exploit any mispricing.
- **Institutional investors:** as in the traders case, they are represented by a unique instance. The role of the institutional investor is to allocate its wealth among the hedge funds according to their performances.

It is really interesting to observe that even this simple setup can generate price bubbles and default clustering among the hedge funds (default clustering is an important aspect of systemic risk). In particular, the simulations show the following extremely realistic pattern: first, highly leveraged hedge funds build huge positions by exploiting small price fluctuations in times of high volatility; then, when the noise traders start to sell the asset, hedge funds start rushing to liquidate their asset positions; this in turn push the price further down causing a fire sale spiral that lead to hedge funds defaulting.

¹⁵This is definitely the case for the work of Thurner (2011); regarding the CCA field, some of its applications have been fully explored but we believe that the one by Gray and Jobst (2010) is extremely promising and deserves even more attention from researchers.

Contingent claim analysis

Contingent Claim Analysis (CCA) is a generalization of the option pricing theory by Black and Scholes (1973). CCA is based on three assumptions:

1. Values of liabilities are derived from the values of the assets
2. Liabilities can be ranked in term of seniority
3. Asset values follow stochastic processes

There are several examples of application of CCA to systemic risk¹⁶; we will focus on a recent one by Gray and Jobst (2010). The idea behind their application of CCA to systemic risk is to estimate the percentage of the system losses expected to be covered by government's intervention. The methodology is based on the usual Merton (1974) framework where owners of corporate equity in leveraged firms are seen as call option holders. Following the usual Black-Scholes argument and the arbitrage-free assumption, spreads on corporate bonds can be used to extract firm expected losses. The authors then use CDS data to calculate the market estimate of the percentage of the total firm losses that should be covered by the government. Market expectations regarding government's liabilities in a crisis are linked to the systemic risk taken on by governments.

¹⁶A part from the one we will discuss in this section, readers might find interesting Gray and Malone (2008) and Gray, Merton, and Bodie (2007).

Chapter 3

Tools

In this chapter we will briefly introduce few topics that are important for the rest of the work; in particular, we will explore Archimedean copulas and the one factor Gaussian model in some details.

This chapter is meant as a quick reference that readers unfamiliar with the above subjects could use in order to more easily approach the rest of the work. It is not, by any means, an attempt to provide complete and exhaustive reports on the state of the art of the topics included; we limited the exposition to the bare minimum and we only reported those aspects of the fields we deemed necessary for the reminder of the exposition. At the beginning of each section, interested readers will find a list of references to more comprehensive reviews.

3.1 A basic introduction to copulas

In this section we will introduce those aspects and concepts of the copula's theory that we think are useful for understanding the rest of the work. We will also include the description of the algorithm we developed in order to calculate the joint default/survival probability for Archimedean copulas when the dimensionality of the problem is small. Nelsen (1999), McNeil, Frey, and Embrechts (2010) and Mai and Scherer (2012) are all excellent references on the subject.

3.1.1 Definitions

Given two n -dimensional vectors $a = (a_1, \dots, a_n)$ and $b = (b_1, \dots, b_n)$ (where each a_i and $b_i \geq a_i$ are in $\overline{\mathfrak{R}}$), we will denote with $[a, b]$ the *hyper-rectangle*, i.e. the subset of $\overline{\mathfrak{R}}^n$ $[a_1, b_1] \times \dots \times [a_n, b_n]$. The vertices of the hyper-rectangle are the vectors $(c_1, \dots, c_n)$ where each c_i can be either a_i or b_i ; let's denote with $H(B)$ the set of all possible vertices of a given hyper-rectangle B . Let's denote with f a real function whose domain $Dom(f)$ is a subset of $\overline{\mathfrak{R}}^n$; for every hyper-rectangle $B \subseteq Dom(f)$, we can then define its **f -measure** $V_f(B)$ ¹ as

$$V_f(B) = \sum_{c \in H(B)} sgn(c) \cdot f(c) \quad (3.1)$$

where the sgn function is defined as

$$sgn(c) = \begin{cases} +1 & \text{if } c_k = a_k \text{ for an even number of } k \\ -1 & \text{if } c_k = a_k \text{ for an odd number of } k \end{cases} \quad (3.2)$$

For example, if $B = [a_1, b_1] \times [a_2, b_2]$, we have

$$V_f(B) = f(a_1, a_2) - f(a_1, b_2) - f(b_1, a_2) + f(b_1, b_2) \quad (3.3)$$

The f -measure of a subset can also be defined in terms of the first order differential operator Δ :

$$V_f(B) = \Delta_{a_1}^{b_1} \dots \Delta_{a_n}^{b_n} f(x) \quad (3.4)$$

¹Some authors refer to $V_f(B)$ as the f -volume of B .

where

$$\Delta_{a_k}^{b_k} f(x) = f(x_1, \dots, x_{k-1}, b_k, x_{k+1}, \dots, x_n) - f(x_1, \dots, x_{k-1}, a_k, x_{k+1}, \dots, x_n) \quad (3.5)$$

An n -dimensional real function f is **n -increasing** if $V_f(B) \geq 0$ for every hyper-rectangle B inside the domain of f . It's important to note that being n -increasing does not mean that f is non-decreasing in every argument; in general the two concepts are unrelated to each other.

Assume $\text{Dom}(f) = S_1 \times \dots \times S_n$ and that each S_k has a minimum and a maximum element, a_k and b_k respectively. Then f is **grounded** if $F(t) = 0 \forall t \in \text{Dom}(f)$ such that $t_k = a_k$ for at least one k . We then say that f has **margins** f_k if

$$f(b_1, \dots, b_{k-1}, x, b_{k+1}, \dots, b_n) = f_k(x), \quad \forall x \in S_k \quad (3.6)$$

where $f_k : \text{Dom}(f_k) \supseteq S_k \rightarrow \Re$ for every k .

We can now give the definition of copula².

An n -dimensional **copula** is a function C with the following properties:

1. $\text{Dom}(C) = [0, 1]^n$
2. C is grounded and n -increasing
3. C has margins C_k that satisfy $C_k(u) = u, \quad \forall u \in S_k$

Another possible characterization of a copula is that of a function C that satisfies the following conditions:

1. $\text{Dom}(C) = [0, 1]^n$
2. $C(u) = 0$ if $u_k = 0$ for at least one k
3. $C(1, \dots, 1, u_k, 1, \dots, 1) = u_k, \forall u_k \in [0, 1]$
4. $\forall a, b \in [0, 1]^n$ such that $a \leq b$, then $V_C([a, b]) \geq 0$

²A caveat is actually due at this point; we decided to introduce directly the concept of a copula and voluntarily skipped the definition of sub-copulas. The difference between the two concepts, although fundamental for the proof of many properties of copulas, is of no interest for the rest of this work.

An interesting property of copulas is that they are uniformly continuous in their domain, a consequence of the following proposition (see Nelsen (1999, theorem 2.10.7))

Proposition 1 *For every u and v in $[0, 1]^n$, we have*

$$|C(u) - C(v)| \leq \sum_{k=1}^n |u_k - v_k| \quad (3.7)$$

We can find boundaries on the possible values that a copula C can take; indeed we have

Proposition 2 *Given a copula C , for every $u \in [0, 1]^n$, we have*

$$W^n(u) \leq C(u) \leq M^n(u) \quad (3.8)$$

where

$$W^n(u) = \max(0, 1 - n + \sum_{k=1}^n u_k)$$

$$M^n(u) = \min(u_1, \dots, u_n)$$

The functions W^n and M^n are known as Fréchet-Hoeffding boundaries. While M^n is always a copula (corresponding to the co-monotone case), W^n is a copula only when $n = 2$ (counter-monotonic case). Another important copula is the independent one

$$\Pi^n(u) = u_1 \cdot \dots \cdot u_n$$

3.1.2 Sklar's theorem

Sklar's theorem (Sklar (1959), Sklar (1996)) provides a crucial link between copulas and multivariate probability cumulative functions.

Consider $X_1, \dots, X_n$ random variables and assume that their marginal distribution functions are denoted with $F_1, \dots, F_n$, while let F denote their multivariate cumulative distribution function:

$$P\{X_i \leq x_i\} = F_i(x_i) \quad (3.9)$$

$$P\{X_1 \leq x_1, \dots, X_n \leq x_n\} = F(x_1, \dots, x_n) \quad (3.10)$$

Theorem 3 *Let $X_1, \dots, X_n$ be random variables whose marginal distribution functions are denoted with $F_1, \dots, F_n$ and let F denote their multivariate cumulative distribution*

function. Then there exists an n -dimensional copula C such that

$$F(x_1, \dots, x_n) = C[F_1(x_1), \dots, F_n(x_n)] \quad (3.11)$$

Moreover, if every F_i is continuous then C is unique; otherwise, C is uniquely determined on $\text{Ran}(F_1) \times \dots \times \text{Ran}(F_n)$.

Conversely, for every n -dimensional copula C , the function G defined as

$$G(x_1, \dots, x_n) := C[F_1(x_1), \dots, F_n(x_n)] \quad (3.12)$$

is a multivariate cumulative distribution function of the random vector X with margins equal to F_i .

A quasi inverse of a cumulative function G is any function $G^{(-1)}(t)$ for which

$$G[G^{(-1)}(t)] = t, \forall t \in \text{Ran}(G)$$

Indeed $G^{(-1)}(t)$ is defined as:

$$\begin{aligned} & \{ \text{any } x \in \bar{\mathfrak{R}} \text{ such that } G(x) = t \} && \text{if } t \in \text{Ran}(G) \\ & \inf \{ x | G(x) \geq t \} && \text{if } t \notin \text{Ran}(G) \end{aligned}$$

Of course, when G is invertible, $G^{(-1)} \equiv G^{-1}$ and the quasi-inverse is unique.

Equipped with this definition, the following equation holds as a corollary of Sklar's theorem

$$C[u_1, \dots, u_n] = F[F_1^{(-1)}(u_1), \dots, F_n^{(-1)}(u_n)] \quad (3.13)$$

where $u_i = F_i(x_i)$.

3.1.3 Archimedean copulas

Archimedean copulas are among the most applied copulas in finance; the main reasons for their widespread usage are the facility with which they can be constructed and the availability of efficient algorithms for numerical simulations. Before giving a formal definition, we will introduce a couple of new concepts.

Let ϕ be a continuous, strictly decreasing function from $[0, 1]$ to $[0, \infty]$ such that

$\phi(1) = 0$. The **pseudo-inverse** of ϕ is the function $\phi^{[-1]} : [0, \infty] \rightarrow [0, 1]$ defined as

$$\phi^{[-1]}(t) = \begin{cases} \phi^{-1}(t) & t \in [0, \phi(0)] \\ 0 & t > \phi(0) \end{cases} \quad (3.14)$$

Note how $\phi^{[-1]}$ is continuous, non-increasing on $[0, \infty]$ and strictly decreasing in $[0, \phi(0)]$; when $\phi(0) = \infty$, $\phi^{[-1]} \equiv \phi^{-1}$. Given continuous, strictly decreasing function ϕ from $[0, 1]$ to $[0, \infty]$ such that $\phi(1) = 0$, let consider the following multivariate function from $[0, 1]^n$ to $[0, 1]$

$$C_\phi(u) = \phi^{[-1]}[\phi(u_1) + \cdots + \phi(u_n)] \quad (3.15)$$

The function C_ϕ has the right shape to be a copula and copulas that can be written in the form of equation (3.15) are known as **Archimedean copulas** and the function ϕ is the (additive) **generator** of the copula. The question we need to answer is which properties ϕ needs to exhibit for C_ϕ to be a copula? The answer depends on the dimensionality n of the problem. When $n = 2$, the following theorem (Alsina, Frank, and Schweizer (2003)) provides the answer

Theorem 4 C_ϕ is a two-dimensional copula if and only if ϕ is convex.

In higher dimensions, we need an additional concept. A function $f(t)$ is **completely monotonic** on the real interval H if it is continuous, infinitely derivable and with derivatives that alternate in sign:

$$(-1)^k \frac{d^k}{dt^k} f(t) \geq 0, \forall t \in \check{H} \quad (3.16)$$

where we used $\check{H}$ to represent the interior points of H .

We can now state the result (Kimberling (1974), Schweizer and Sklar (1983)) for the generic case $n \geq 3$:

Theorem 5 C_ϕ is an n -dimensional copula if and only if ϕ^{-1} is completely monotonic on $[0, \infty)$.

Selected examples

The variety of families belonging to the Archimedean class is yet another reason of the success of this type of copulas in applications. In this section we will focus on the three one-parameter families of Archimedean copulas we used in the rest of the work: Frank,

Gumbel and Clayton. In the following we will indicate with θ the single parameter that characterizes the generator function, denoted by ϕ_θ . Finally, we will indicate with C_θ the corresponding copula function.

Frank

$$\phi_\theta(t) = -\ln \left(\frac{e^{-\theta \cdot t} - 1}{e^{-\theta} - 1} \right) \quad (3.17)$$

This family first appeared in Frank (1979); both Genest (1987) and Nelsen (1986) provide analysis on some statistical properties of this family. The parameter θ can vary in the range $(0, \infty)$ and note how the family spreads from independence ($C_\theta \rightarrow \Pi$ when $\theta \rightarrow 0$) to complete monotonicity ($C_\theta \rightarrow M$ when $\theta \rightarrow \infty$). When $n = 2$, the parameter θ can also take negative values; moreover, C_θ approaches W when $\theta \rightarrow -\infty$, making the Frank family comprehensive³.

Gumbel

$$\phi_\theta(t) = [-\ln(t)]^\theta \quad (3.18)$$

Also known as Gumbel-Hougaard, this family was introduced in Gumbel (1960) and also described in Hougaard (1986). The parameter θ can vary in the range $[1, \infty)$; we have $C_1 = \Pi$ while C_θ approaches M when $\theta \rightarrow \infty$.

Clayton

$$\phi_\theta(t) = \frac{t^{-\theta} - 1}{\theta} \quad (3.19)$$

This family, introduced by Clayton (1978), is also known under other names among which Pareto family (Hutchinson and Lai (1990)) and generalized Cook and Johnson family (Genest and Mackay (1986)). The parameter can vary in $(0, \infty)$ for a generic dimensionality and this family is, like the Frank's one, comprehensive when $n = 2$ as we have $C_0 = \Pi$, $C_\infty = M$ and $C_{-1} = W$ (the last equality holds only for $n = 2$).

Table 3.1, extracted from Nelsen (1999, table 4.1), summarizes the above descriptions.

³A one-parametric copula family is said to be comprehensive if includes as special cases M , W and Π .

Copula Family	$\phi(\theta)$	$\theta \in$	Limiting and special cases
Frank	$-\ln \left(\frac{e^{-\theta \cdot t} - 1}{e^{-\theta} - 1} \right)$	$(-\infty, +\infty) \setminus \{0\}$	$C_{-\infty}^* = W, C_0 = \Pi, C_\infty = M$
Gumbel-Hougaard	$[-\ln]^\theta$	$[1, \infty)$	$C_0 = \Pi, C_\infty = M$
Clayton	$\frac{t^{-\theta} - 1}{\theta}$	$[-1, \infty) \setminus \{0\}$	$C_{-1}^* = W, C_0 = \Pi, C_\infty = M$

Table 3.1: Archimedean copula families. * only valid for $n = 2$.

3.1.4 Kendall's τ

Kendall's τ is a measure of dependency of joint distributions. In order to give a formal definition of Kendall's τ , let define the concept of concordant/discordant pairs.

Given two pairs of observations (x_i, y_i) and (x_j, y_j) from two variables X and Y , we say they are **concordant** if $(x_i - x_j) \cdot (y_i - y_j) > 0$, **discordant** if $(x_i - x_j) \cdot (y_i - y_j) < 0$.

We can now define Kendall's τ in the bivariate case:

$$\tau = \frac{c - d}{c + d}$$

where c is the number of concordant pairs while d is the number of discordant ones. In the Archimedean case, we have a very useful expression that links Kendall's τ to the copula's generator:

$$\tau = 1 + 4 \cdot \int_0^1 \frac{\phi(t)}{\phi'(t)} dt \quad (3.20)$$

Equation (3.20) can be used either directly for calculating τ or indirectly to calibrate copulas to the desired level of dependency.

3.1.5 Calculating joint survival/default probabilities

Some of the most common systemic risk measures proposed in literature⁴ can be calculated from quantities like

$$P\{X_1 \diamond_1 K_1, \dots, X_n \diamond_n K_n\} \quad (3.21)$$

where $\diamond_i$ can be either $\leq$ or $>$ and we assume that the joint distribution is driven by an Archimedean copula. We will show an efficient algorithm that calculates (3.21); it consists of two steps:

⁴Especially the ones seen in section 2.3.4.

1. Split calculation of probabilities of the type (3.21) into terms containing only the $\leq$ sign:

$$P\{X_1 \leq G_1, \dots, X_n \leq G_n\}$$

where G_i is either K_i or ∞ .

2. Using Sklar's theorem, we then calculate

$$P\{X_1 \leq G_1, \dots, X_n \leq G_n\} = C[F_1(x_1), \dots, F_n(x_n)]$$

where C is the copula function and $F_i(\cdot)$ is the marginal cumulative density function for variable X_i and $F_i(x_i) = P\{X_i \leq K_i\}$ if $G_i = K_i$, or $F_i(x_i) = 1$ if $G_i = \infty$

Before showing the sketch of the algorithm for a generic value of n , let look at a little example in 3 dimensions. Suppose we want to calculate $\hat{P} := P(X > x, Y \leq y, Z > z)$; we can also write $\hat{P}$ as

$$\hat{P} = [P(X > x, Y \leq y)] - [P(X > x, Y \leq y, Z \leq z)].$$

Applying the same principle again to each term in the previous equation we get

$$\hat{P} = [P(Y \leq y) - P(X \leq x, Y \leq y)] - [P(Y \leq y, Z \leq z) - P(X \leq x, Y \leq y, Z \leq z)]$$

Note how we can rewrite the first term to make it contain all 3 variables

$$P(Y \leq y) = P(X \leq \infty, Y \leq y, Z \leq \infty)$$

and similarly for every term missing any of the variables. Let $u = F_X(x)$, $v = F_Y(y)$ and $w = F_Z(z)$ ⁵; the right hand side can now be calculated in copulas terms as

$$\hat{P} = C(1, v, 1) - C(u, v, 1) - C(1, v, w) + C(u, v, w)$$

We can now give the description of the generic algorithm in its two steps:

1. **Recursive step:** if there is at least a $>$ sign among the $\diamonds$, let j be the first index for which we have a $>$ sign, use the following recursive structure until all

⁵For every real-valued variable V , we have $F_V(\infty) = 1$.

$>$ have been eliminated:

$$\begin{aligned} P\{X_i \leq K_i, X_j > K_j, X_h \diamond_h K_h\} = \\ P\{X_i \leq K_i, X_j \leq \infty, X_h \diamond_h K_h\} - \\ P\{X_i \leq K_i, X_j \leq K_j, X_h \diamond_h K_h\} \end{aligned}$$

$$\forall i < j, \forall h > j$$

2. **Call to C :** when every $\diamond$ is a $\leq$, simply call the copula function

$$P\{X_1 \leq G_1, \dots, X_n \leq G_n\} = C[F_1(x_1), \dots, F_n(x_n)]$$

In the case of Archimedean copulas, the right hand side of the last expression can be written as

$$C[F_1(K_1), \dots, F_n(K_n)] = \phi^{[-1]}(\phi(u_1) + \dots + \phi(u_n))$$

where $u_i = F_i(K_i)$.

A caveat

An important caveat is actually due at this point: the above algorithm suffers from the problem of loss of precision due to cancellations effects because of the many subtractions between terms of similar magnitude. Although stable for the sizes we handled in the present work, a more detailed analysis of this problem will be required before using it for medium sized portfolios.

3.2 The one factor Gaussian base correlation approach

In this section we will review the one factor Gaussian model with a particular attention to the base correlation approach. A part from the original description of the model by Li (2000), readers might be also interested in Burtschell, Gregory, and Laurent (2009) and also Longstaff and Rajan (2008) for comparisons of CDO pricing models and Finger (2004) and Eberlein, Frey, and Von Hammerstein (2008) for a more generic overview of pricing CDO products.

3.2.1 Introduction

The one factor Gaussian model (OFG in short) was the de-facto standard modeling tool for pricing CDO products before recent global financial crisis; even now, it is still widely used as lingua franca and represents the baseline for many models currently used in practice. Its success derives from its simplicity, the parsimony of the parameters needed⁶ and its tractability.

3.2.2 Description

The main underlying assumptions of the OFG model are:

$$\begin{aligned} A_i &= \sqrt{1-\rho} \cdot V_i + \sqrt{\rho} \cdot V, \quad \rho \in [0, 1] \\ V_i, V &\approx N(0, 1) \\ P\{i \text{ defaults}\} &= P\{A_i \leq K_i\} \end{aligned} \tag{3.22}$$

So defaults are linked to the event of the latent variable A_i (often assumed to represent assets values) going below a threshold K_i . The A_i are then modeled as sum of two normal distributed quantities reflecting idiosyncratic effects and the common state of the world, modeled via the variable V .

From the first and second equation of (3.22), we easily get that $A_i \approx N(0, 1)$ making the calibration of the default thresholds K_i straightforward once the marginal probability of default, PD_i , is given:

$$K_i = \Phi^{-1}(PD_i)$$

⁶On top of the single name info, only a single correlation number is needed in its most basic form.

where Φ is the cumulative density function of the standard normal distribution $N(0, 1)$. It also worth noting that

$$E[A_i \cdot A_j] = E[\rho V^2] = \rho \quad \forall i \neq j$$

From (3.22) also derives the conditional independence property, i.e. the probability of defaults of any two names i and j conditioned on a given value of the common factor are independent. In particular, the equation for the conditional probability of default of name i is given by

$$P\{A_i \leq K_i | V = v\} = P\left\{V_i \leq \frac{K_i - \sqrt{\rho} \cdot v}{\sqrt{1 - \rho}}\right\} = \Phi\left(\frac{K_i - \sqrt{\rho} \cdot v}{\sqrt{1 - \rho}}\right) \quad (3.23)$$

The above equation is central for the applicability of the model as it suggests a strategy to obtain portfolio loss distribution:

- integrate numerically over the common factor V ;
- for every value of the common factor, use equation (3.23) in order to obtain the conditional default probabilities;
- exploit the independence of the probabilities in the previous point to get the conditional loss distribution. This can be done in many ways, including recursive algorithm as the one in Andersen, Sidenius, and Basu (2003) or using fast Fourier transforms.

Figure 3.1 shows, in the top part, the shape of $P\{A_i \leq K_i | V = v\}$ when the correlation ρ is varied for two values of the common factor. The bottom part, instead, shows the shape of the conditional probability when the common factor varies for several levels of correlation ρ .

3.2.3 CDS and CDO pricing

In this section we will give basic information regarding CDO products. The best starting point is actually a simpler product: the credit default swap (CDS) is essentially a form of insurance against the default of a specific firm, called the reference entity. In this contract the protection buyer pays a premium at fixed schedule until the maturity of the

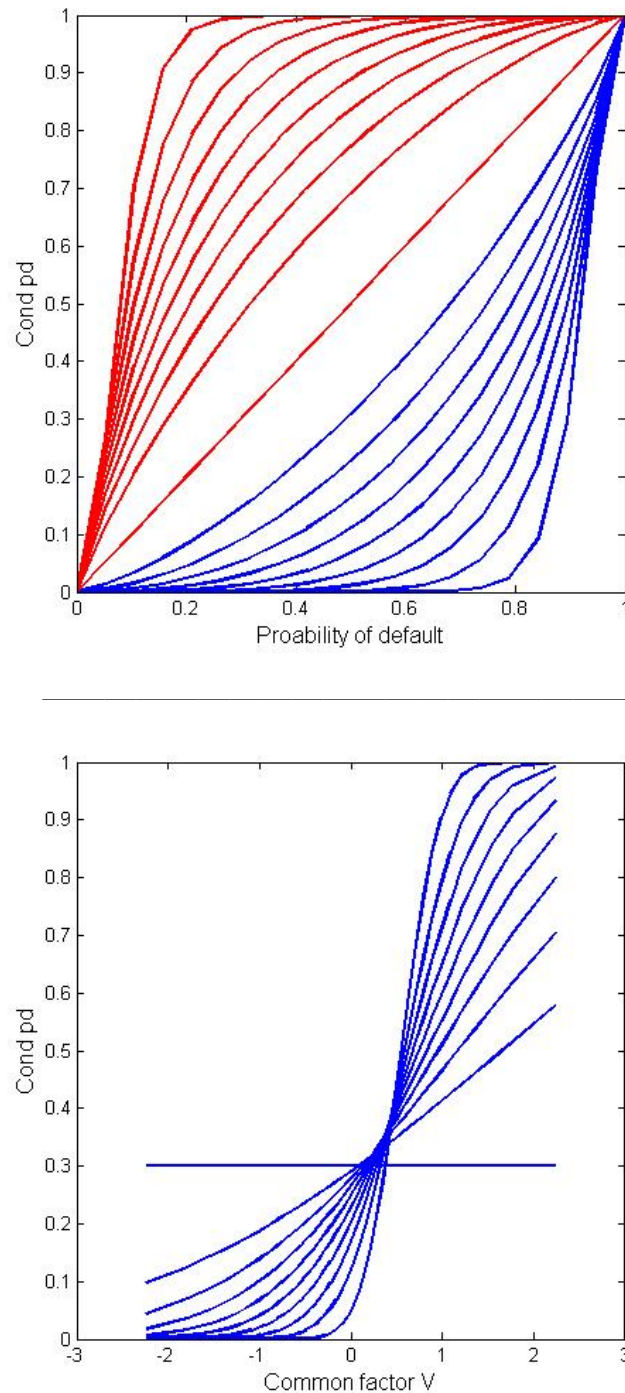


Figure 3.1: **Top plot:** conditional (y -axis) vs. marginal (x -axis) probability of default. The blues lines correspond to $V = -1.5$ while the red ones to $V = 1.5$. Lines are obtained with increasing values of correlation ρ (low for lines closer to the diagonal, high for external ones).

Bottom plot: Conditional probability of default vs. the value of the common factor V on the x -axis. The gradient of the curves increases as the value of the correlation ρ is increased. The marginal probability of default used is 30%.

contract or a default event of the reference entity⁷. Upon default, the protection seller compensates the buyer by either paying the face value in exchange of the defaulted bond (physical delivery) or by paying a cash amount (cash settlement). A CDS trade has hence two legs; the coupon leg represents the expected value of the payments made by the protection buyer while the offsetting leg is the expected value of the default payment made by the protection seller. In formulas

$$DfltLeg = N \cdot (1 - rec) \cdot \int_0^M -\frac{\partial s(t)}{\partial t} \cdot v(t) dt \quad (3.24)$$

$$BPV = N \cdot \sum_i dcf(t_i, t_{i+1}) \cdot v(t_{i+1}) \cdot s(t_{i+1}) \quad (3.25)$$

$$CpnLeg = \text{premium} \cdot BPV \quad (3.26)$$

$$V_{CDS} = CpnLeg - DfltLeg \quad (3.27)$$

where

$v(t)$ is the discount factor at time t

$s(t)$ is the probability that the reference entity is still alive by time t

$dcf(t, s)$ is the day count fraction between the coupon dates t and s

rec is the expected recovery rate

and N is the notional of the trade while M is the maturity. Please note that the term $-\frac{\partial s(t)}{\partial t}$ represents the probability of default happening at time t ⁸. The par spread of the trade is calculated as the fair value of the cpn , i.e. the level of cpn that makes the value of trade zero:

$$\text{ParSpread} = \frac{DfltLeg}{BPV} \quad (3.30)$$

⁷It is **a** default event rather than **the** default event as there are several conditions that might trigger a protection claim under the most common specification of the contract. These include failure to repay a debt, filing for bankruptcy but even debt or company restructuring.

⁸This is based on the most common reduced form model by Jarrow and Turnbull (1995) where the default time τ is modeled as the first event of a Poisson counting process. In this case we have

$$P\{\tau < t + \delta t | \tau \geq t\} = \lambda(t) \delta t \quad (3.28)$$

$$s(t) = e^{-\int_0^t \lambda(t) dt} \quad (3.29)$$

where $\lambda(t)$ is the hazard rate. See Jarrow and Turnbull (1995) or Lando (2009) (chapter 4) for details.

Par spreads are among the most common quoting mechanisms in the CDS world.

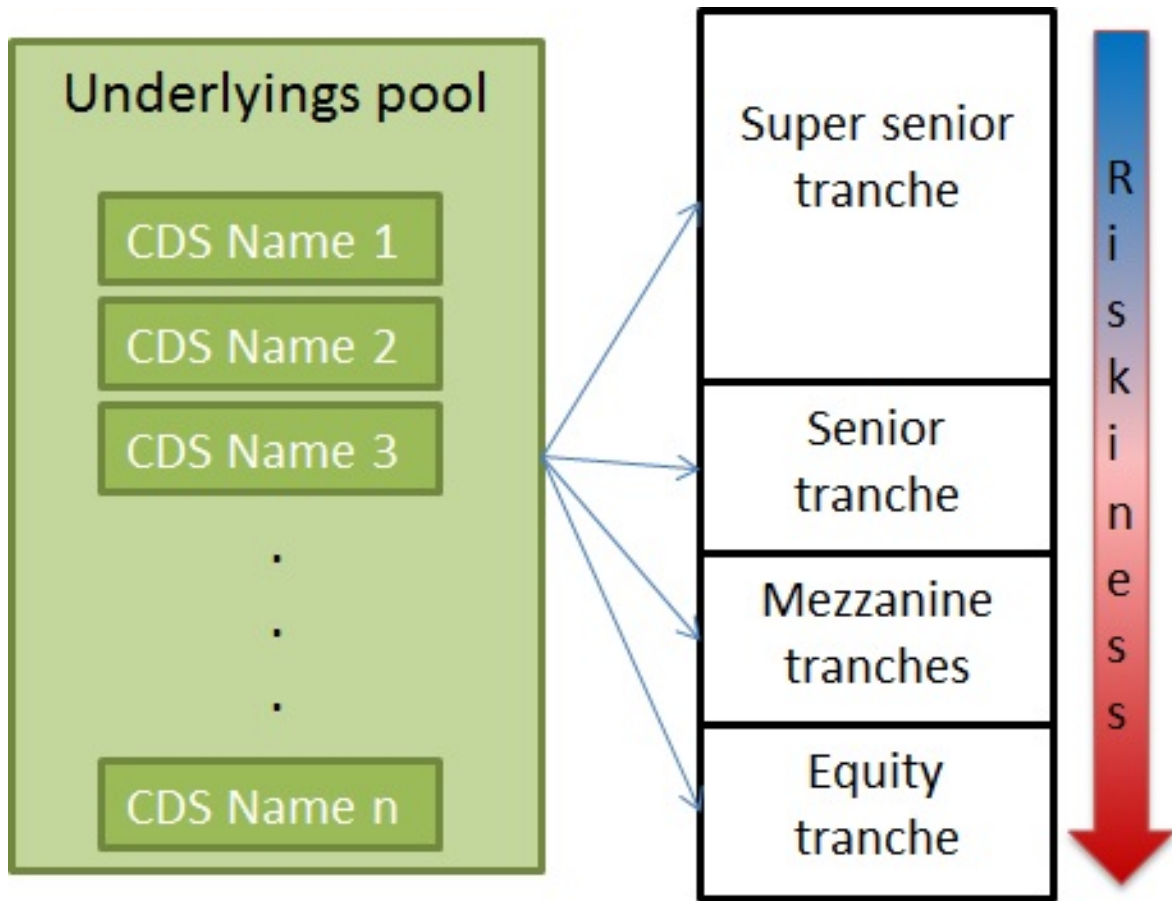


Figure 3.2: Synthetic CDO structure

CDSs represent the basic building blocks of more complex structured products for at least two reasons. The first is that CDSs are heavily used to calibrate complex models and hedge structured products that rely on them. Secondly, the structure of basket and portfolio credit trades can be better understood by looking at similar features in the single name world. For example, let's look at CDO products, in particular synthetic ones⁹. A synthetic CDO is a complex structured product in which a pool of underlying entities is grouped together. The CDO is then sliced into tranches characterized by an attachment and detachment point: these determine the subordination of the tranche. The portfolio of underlying entities will suffer losses due to defaults. If the losses suffered are above the attachment of a given tranche, the notional of the tranche itself is eroded,

⁹Synthetic CDOs use CDSs instruments as opposed to cash CDOs that use mortgages and bonds.

even to the point of wiping out the entire tranche when the portfolio losses pass the detachment level. Junior tranches (the ones with low attachment and detachment points) will suffer higher losses while very senior ones (high attach and detach) are more protected as all the tranches below them have to be eroded before they start suffering losses. This explain why junior tranches pay premium that are much higher than the equivalent for senior tranches: they have to compensate for riskier loss profile. Figure 3.2 provides a visual description of the structure just summarized.

The pricing formulas are very similar to the CDS case with the difference that instead of a fixed notional, we now have to use the expected value of tranche notional. Where we used s for CDS products, for a CDO tranche with attachment and detachment (att, det) , we will instead use the following quantity:

$$S(att, det, t) = (det - att) \cdot P\{L(t) \leq att\} + \int_{att}^{det} (det - x) \cdot P\{L(t) = x\} dx \quad (3.31)$$

where $L(t)$ is the portfolio losses at time t and $S(att, det, t)$ can be interpreted as the expected notional left in the tranche at time t . The above definition makes the pricing formulas for a CDO tranche look very familiar:

$$DfltLeg = \int_0^M -\frac{\partial S(att, det, t)}{\partial t} \cdot v(t) dt \quad (3.32)$$

$$BPV = \sum_i dcf(t_i, t_{i+1}) \cdot v(t_{i+1}) \cdot S(att, det, t_{i+1}) \quad (3.33)$$

$$CpnLeg = cpn \cdot BPV, \quad ParSpread = \frac{DfltLeg}{BPV} \quad (3.34)$$

$$V_{CDO} = CpnLeg - DfltLeg \quad (3.35)$$

3.2.4 Base correlation approach

The main reason behind the success of the OFG model as reference for the CDO market lies in the concept of base correlation; the OFG ρ is in fact used as a communication tool among practitioners, in a similar way to what happens on option trading desk with Black-Scholes σ .

Base correlation is built following the following recipe. Every tranche is priced as a difference of two equity tranches, i.e. tranches with attachment set at zero. For example, the 3-6 tranche is valued as a 0-6 tranche minus a 0-3 one with similar maturity

and coupon structure. Then, one correlation ρ is calibrated for each seniority to match the market quotes and it is known as base correlation (as opposed to the compound correlation, i.e. the single ρ that matches a given tranche priced in isolation and not as difference of equities). This calibration strategy is based on the monotonicity of par spreads for equity tranches with respect to the ρ parameter. There is then one single correlation for every seniority traded, causing a correlation skew with respect to seniority that is similar, in principle, to the Black-Scholes volatility smile observed with respect to strikes in the option world.

Of course, this is also source of problems as the different specifications of the OFG are not necessarily consistent with each other, leaving the door open to mispricing; while the liquidity of the markets makes arbitrages between standard tranches less likely, the same guarantee is not given for non-standard ones for which interpolation and extrapolation techniques are required¹⁰.

Figure 3.3 gives an idea of the model behavior when increasing correlation for a simple homogeneous portfolio of 40 names. The x -axis (on the bottom right) follows correlation while the y -axis (bottom left) follows number of defaults. By increasing correlation we can move from an independent bell-shaped distribution on the left (dark blue) to an almost bimodal distribution on the right (dark red) where only two scenarios are actually contemplated: either every name survive or the entire portfolio defaults.

3.2.5 Further readings

Several extensions of the OFG model have been proposed in recent years with the aim of overcoming some of its limitations. A very successful branch of literature has focused on making the recovery rate dependent on the common factor in order to increase losses in the tail (two good examples are Amraoui, Cousot, Hitier, and Laurent (2012) and Andersen and Sidenius (2004)). Other authors have instead used different distributions for the common factor¹¹ including gamma distribution (Joshi and Stacey (2006)), the normal inverse (Kalemanova, Schmid, and Werner (2007)) and several attempts with t-Student (for example, Vrins (2009)). Others have instead worked on modifying the

¹⁰The main issue is that simple interpolation in the parameter space can generate inconsistent loss distributions, even if there are no arbitrages in the starting nodes of the interpolation; see Parcell (2007) for more details.

¹¹From a theoretical point of view, this type of approach is straightforward; the difficulty is how to keep a good degree of tractability.

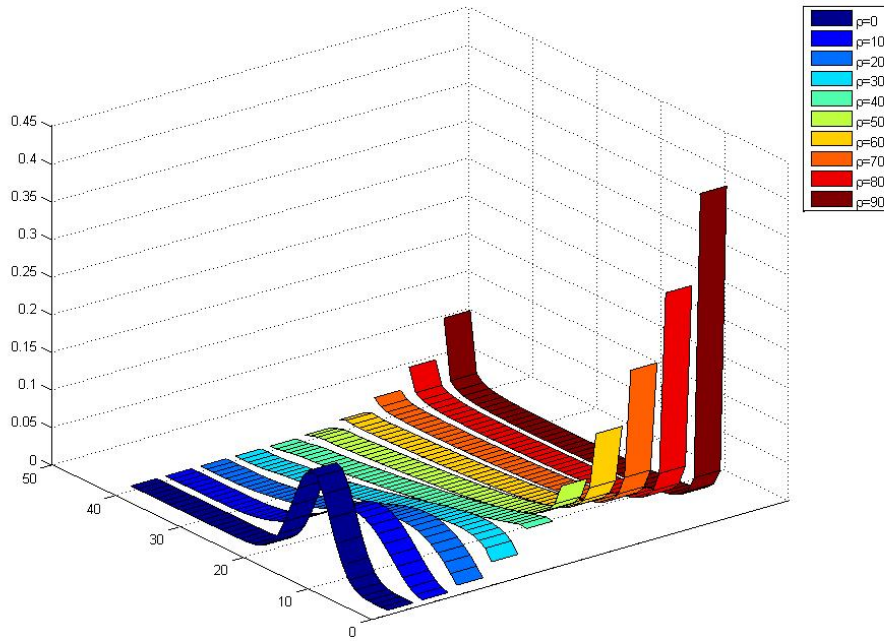


Figure 3.3: Loss distribution vs. ρ for an homogeneous portfolio of 40 names.

underlying copula: Crane and Van Der Hoek (2008), Hofert and Scherer (2011), Hull and White (2010) and Wang, Rachev, and Fabozzi (2009) are all good resources in merit.

There are, of course, several works that significantly depart from the OFG scheme. A very elegant branch of the literature deals with Markov approaches (for example, Bielecki, Cousin, Crépey, and Herbertsson (2012), Di Graziano and Rogers (2009) and Laurent, Cousin, and Fermanian (2011)). Some papers try to minimize the modeling framework to the bare minimum by implying information directly from CDO prices as done by Cont, Deguest, and Kan (2010) and Walker (2009), among others. Finally, interesting readings are the papers from Frey and Backhaus (2010), Errais, Giesecke, and Goldberg (2010) and Willemann (2007).

Part II

Articles

Chapter 4

CIMDO *posterior* stability study

Following recent crisis, great attention has been given in the recent past to the problem of measuring systemic risk in financial systems. Among the many tools proposed, the CIMDO methodology stands out as a flexible and elegant approach to the problem of specifying a multivariate distribution starting from marginal info and it has been widely used by financial regulators worldwide. In particular, its Bayesian spirit suggests that the CIMDO *posterior* can lessen the model uncertainty problem, a feature partially confuted by a recent theoretical result proved by Radev and extended in this chapter. The aim of this chapter is to analyze the relationship between the *prior* and the *posterior* distributions by studying their performance on few systemic risk measures. Our results suggest that the choice of the *prior* is critical when using the CIMDO approach for measuring systemic risk and should serve as first guidance to any users of such methodology.

4.1 Introduction

Segoviano (2006) introduced the CIMDO¹ methodology designed to give an answer to the problem of finding a multivariate distribution to model the joint behavior of a given system when only marginal/incomplete statistics are available. We will see the details of the methodology in the following sections but we can anticipate two of its advantages. The first is its simplicity and elegance, qualities that help explain why many financial regulators have listed the CIMDO approach among their regulatory tools². The second reason for the methodology appeal is its Bayesian flavor: one can hope that, due to the Bayesian spirit of the methodology, the CIMDO approach is capable of limiting the impact of the choice of the parametric family, ultimately reducing the model uncertainty.

Since the first work on the subject, an increasing number of papers using and adapting this methodology has appeared, mostly concerning the problem of measuring systemic risk in financial systems. Segoviano and Goodhart (2009), applied the methodology to this end and quickly many regulators worldwide enlisted the CIMDO approach among their tool sets used to assess the stability of national banking systems³. The model has also been used in Jin and De Simone (2013) as a building block toward a dynamic t-copula model.

4.2 CIMDO description

The task of finding a multivariate distribution starting only from marginal info is, in general, an under-identified mathematical problem that has an infinite number of solutions. In order to reduce its degree of freedom, usually additional structure is enforced on the solution, either by picking a parametric distribution or by specifying a dependency structure, for example embracing the copula approach. These steps turns the original problem into a well-identified one, reducing its complexity and usually increasing its tractability at the cost of adding spurious information when none was actually

¹CIMDO stands for Consistent Information Multivariate Density Optimizing.

²The audience of financial stability reports, for example, is extremely wide and diversified in terms of sophistication; methodologies that are too complex and/or seen as black boxes are usually discharged in favor of simpler approaches.

³Just as examples, Lee, Ryu, and Tsomocos (2012) used it for Korean banks while Johansson (2011) applied it to the Swedish case. These are just but two examples of national agencies embracing the CIMDO approach. In addition, many recent official publications of the IMF (International Monetary Fund) also provide CIMDO based results.

observed.

This is where the CIMDO framework offers some improvements over the standard approaches. It constructs a multivariate density (the *posterior*) that will be as close as possible to the starting *prior* but also satisfying the observed statistics: this is achieved by solving a constrained optimization problem where the shape of the objective function used minimizes the difference between the two distributions while the constraints represent the marginal statistics observed.

In the next sections we will review the single components of the optimization problem used to generate CIMDO's *posterior* multivariate distribution p .

4.2.1 The distance function

The difference between the *prior* and the *posterior* is mathematically measured via the Kullback cross-entropy approach in which the available information is used to infer the unknown distribution. Let p and q be 2 real-valued functions whose domains lie in $\mathbb{R}^n$. In our particular case, they will be the *prior* (q) and *posterior* (p) distributions. Let define the cross-entropy objective function as

$$C(p, q) = \int_{\mathbb{R}^n} p(\mathbf{x}) \cdot \ln \left[\frac{p(\mathbf{x})}{q(\mathbf{x})} \right] \mathbf{d}\mathbf{x} \quad (4.1)$$

Equation (4.2.1) falls short of being a distance function as it is not symmetric. Indeed, The cross entropy function is technically known as a quasi-metric satisfying 3 of the 4 axioms of a metric:

$$C(p, q) \geq 0$$

$$C(p, q) = 0 \Leftrightarrow p = q$$

$$C(p, q) \leq C(p, z) + C(z, q)$$

but not the fourth one as $C(p, q) \neq C(q, p)$. It can be interpreted as the difference between the 2 distributions p and q .

4.2.2 The constraints

Once we have introduced the function $C(p, q)$, the only piece needed for the CIMDO methodology is the marginal information that comes via two separate inputs:

Current default probabilities PoD_i

These represent the individual default probability of entity i . The framework is completely agnostic regarding how the PoD are calibrated.

Historical default thresholds K_i

These represent, in a Merton style model (Merton (1974)), the historical default threshold of entity i , i.e. quantities such that

$$\text{Historical average of } P\{i \text{ defaults}\} = \int_{K_i}^{\infty} \phi_i(x) dx$$

where $\phi_i(x)$ is the marginal probability density function of variable X_i . The average is taken over a medium time interval (usually 1 year).

These inputs are combined into a set of constraints (one for each entity) requiring the *posterior* density p to match the $PoDs$ over the default regions specified by the K s:

$$\int_{\mathbb{R}} \cdots \int_{K_m}^{\infty} \cdots \int_{\mathbb{R}} p(\mathbf{x}) d\mathbf{x} = \int_{\mathbb{R}^n} p(\mathbf{x}) \cdot \chi_d^m(\mathbf{x}) d\mathbf{x} = PoD_m, \quad m = 1, \dots, n$$

where the function χ_d^i is defined as

$$\chi_d^i(\mathbf{x}) = \begin{cases} 1 & x_i \geq K_i \\ 0 & \text{otherwise} \end{cases}$$

Finally, the last constraint needed is a technical one to make sure that the *posterior* p is a properly formed probability density function:

$$\int_{\mathbb{R}^n} p(\mathbf{x}) d\mathbf{x} = 1$$

4.2.3 Mathematical formulation and solution

Putting all together, the solution is then recovered by the following minimization problem:

$$\text{Minimize: } C(p, q) \quad (4.2)$$

$$\int_{\mathbb{R}^n} p(\mathbf{x}) d\mathbf{x} = 1 \quad (4.3)$$

$$\int_{\mathbb{R}^n} p(\mathbf{x}) \cdot \chi_d^m(\mathbf{x}) d\mathbf{x} = PoD_m, \quad m = 1, \dots, n \quad (4.4)$$

Using the method of Lagrange multipliers, we need to minimize the function

$$L(p, \lambda, \mu) = C(p, q) + \sum_{m=1}^n \lambda_m \left[\int_{\mathbb{R}^n} p(\mathbf{x}) \cdot \chi_d^m(\mathbf{x}) d\mathbf{x} - PoD_m \right] + \mu \left[\int_{\mathbb{R}^n} p(\mathbf{x}) d\mathbf{x} - 1 \right]$$

The solution is then found by imposing $\nabla L(p, \lambda, \mu) = 0$, i.e. solving the following system:

$$\begin{cases} \frac{\partial L}{\partial p} = 0 \\ \frac{\partial L}{\partial \lambda} = 0 \\ \frac{\partial L}{\partial \mu} = 0 \end{cases} \quad (4.5)$$

We can go one step further thanks to the particular shape of $C(p, q)$; noting that $\frac{\partial C(p, q)}{\partial p(\mathbf{x})} = \ln \left(\frac{p(\mathbf{x})}{q(\mathbf{x})} \right) + 1$, the first part of the previous system of equations, $\frac{\partial L}{\partial p(\mathbf{x})} = 0$, becomes

$$\ln \left[\frac{p(\mathbf{x})}{q(\mathbf{x})} \right] + 1 + \mu + \sum_{m=1}^n \lambda_m \cdot \chi_d^m(\mathbf{x}) = 0$$

From which the final solution for p has to be of the form

$$\hat{p}(\mathbf{x}) = q(\mathbf{x}) \cdot \exp \left[-1 - \hat{\mu} - \sum_{m=1}^n \hat{\lambda}_m \chi_d^m(\mathbf{x}) \right] \quad (4.6)$$

where $\hat{\mu}$ and $\hat{\lambda}_m$ represents the optimal values for the correspondent parameters.

Regarding equation (4.6), it is worth noting that p can be factorized as the product of the *prior* q and the function

$$f(\mathbf{x}, \hat{\mu}, \hat{\lambda}) = \exp \left[-1 - \hat{\mu} - \sum_{m=1}^n \hat{\lambda}_m \chi_d^m(\mathbf{x}) \right]$$

Let a survival/default event be any of the following events:

$$\{x_1 \diamond K_1, \dots, x_G \diamond K_G\}$$

where $\diamond$ can be either $>$ or $\leq$ ⁴. We can then note that $f(\mathbf{x}, \hat{\mu}, \hat{\lambda})$ is constant inside such default/survival regions; this feature not only allows for faster calculations (please see the appendix B.1 for a detailed study on this issue) but it is also crucial for the proof of our independence result in the next section.

4.3 The independence issue

One of the most appealing features of the CIMDO methodology is the fact that it lessens the dependency from the choice of the *prior*. As the *posterior* will be fitted on top of it, one can think that the choice of the distributional form and of the parametrization of the *prior* becomes less important.

This belief is somehow mitigated by a theoretical stability analysis performed by Segoviano (2006): he showed that the choice of the entropy function leads to more stable *posteriors* with respect to infinitesimal changes to the *prior* than other metric functions. In more details, let E be an alternative distance function between p and q defined as:

$$E_r(p, q) = \int_{\mathbb{R}^n} |p(\mathbf{x}) - q(\mathbf{x})|^r d\mathbf{x}$$

Segoviano showed that

$$\left| \frac{\partial C(p, q)}{\partial q} \right| \leq \left| \frac{\partial E_r(p, q)}{\partial q} \right|$$

when $r \geq 2$. in other words, the objective function used in the minimization problem in order to find the *posterior* is more stable with respect to perturbations of the *prior* q if we use $C(p, q)$ rather than $E_r(p, q)$, for $r \geq 2$. Thanks to the minimization step, a bigger stability of the objective function is naturally translated into a more stable *posterior*. This result also means that the choice of the *prior* q is likely to dictate the shape of the *posterior* p as the latter cannot depart from q freely (or at least, it can depart from it but less than if it was obtained via E_r instead).

Indeed, in a recent paper, Radev (2012a) proved that the link between q and p is even

⁴For simplicity we will be working with absolute continuous probability measures, freeing us from worrying to much about the equal sign.

stronger in some cases due to the following result:

Theorem 1 *If the prior is normal with the identity matrix as variance/covariance matrix, then default/survival events are independent under the CIMDO posterior p .*

4.3.1 Extended independence result

Theorem 1 casts some doubts over the ability of the CIMDO framework to ease the model uncertainty; even more so as we have extended Radev's result by eliminating the normal assumption and shown that a double implication actually holds (please refer to the appendix A.1 for the proof):

Theorem 2 *If default/survival events are independent under the prior distribution q , they will also be independent under the CIMDO posterior p :*

$$\begin{aligned} P_q \{x_1 \diamond K_1, \dots, x_G \diamond K_G\} &= \prod_{i=1}^G P_q \{x_i \diamond K_i\} \\ &\Leftrightarrow \\ P_p \{x_1 \diamond K_1, \dots, x_G \diamond K_G\} &= \prod_{i=1}^G P_p \{x_i \diamond K_i\} \end{aligned} \tag{4.7}$$

This result makes the link between q and p even stronger as the independence of one is transferred to the other for default/survival regions no matter which distributional assumptions are taken.

4.3.2 Consequences

The results showed in this section uncover a strong (and we suspect mostly unknown) link between the *prior* and the *posterior*. The ability of the CIMDO to build robust *posterior* even in the presence of misrepresentations of the *prior* is clearly weakened. It is legitimate to question which other behaviors the *posterior* inherits from the *prior* and the other CIMDO inputs, an aspect that will be discussed in details in the rest of the chapter.

There is also another source of concern regarding possible flaws in the application of the CIMDO methodology to systemic risk measurement. For example, if we assume that the *prior* chosen threats default/survival events as independent, it is easy to see that the *JPoD* measured under the *posterior* distribution will simply be:

$$JPoD(p) = \prod_i PoD_i$$

In this case, the entire machinery used to retrieve the *posterior* p adds no information at all on systemic risk, as we can obtain it simply from the set of inputs PoD_i ⁵. This is an example of a bad combination of choice of *prior* and choice of risk measure. Please note that there is nothing inherently wrong with the *JPoD* measure: similar results can be obtained for most risk measures based on default/survival events. Similarly, there is nothing wrong with assuming independence when choosing the *prior*; it is only the combination of the two that can lead to unexpected results.

4.4 *Posterior* stability study

The aim of the rest of the chapter is to study the effect that each CIMDO input has on the final *posterior* with a special focus on systemic risk measurement. The recipe we will follow in order to study the impact on the *posteriors* is the following:

1. Build a small fictional yet realistic specification of the model⁶.
2. Perturb its components obtaining new modified *posteriors*.
3. Compute the BSI, SFM, JPoD and PAO systemic risk measures implied by both the different *posteriors* generated by our perturbations and the original *priors*.
4. Use the results in the previous point to compare the effects of the perturbations.

We will show the results divided in three parts: the effects of changes to the *prior*'s parametric assumptions, the impact of changes of the parametrization of q and the dependency of p from the marginal info PoD and K .

⁵For the curious reader, in this case also the *prior* q adds no extra info as we have

$$JPoD(q) = \prod_i \text{Hist}PoD_i$$

where $\text{Hist}PoD_i := \text{Historical average of } P\{i \text{ defaults}\}$.

The above holds just (reasonably) assuming that the *prior* and the marginal probability density functions agree on the probability mass assigned to individual default events:

$$\int_{\mathbb{R}} \cdots \int_{K_i}^{\infty} \cdots \int_{\mathbb{R}} q(\mathbf{x}) d\mathbf{x} = \int_{K_i}^{\infty} \phi_i(x) dx$$

⁶The decision to use a fictional specification of the model rather than relying on market data allowed us to retain full control over the feature we wanted to test.

4.4.1 The base case

The starting point of our analysis will be a small portfolio of 10 names. Figure 4.1 shows the default probabilities PoD versus the historical default thresholds K for the base case. We explicitly built the base portfolio to cover a wide section of the probability

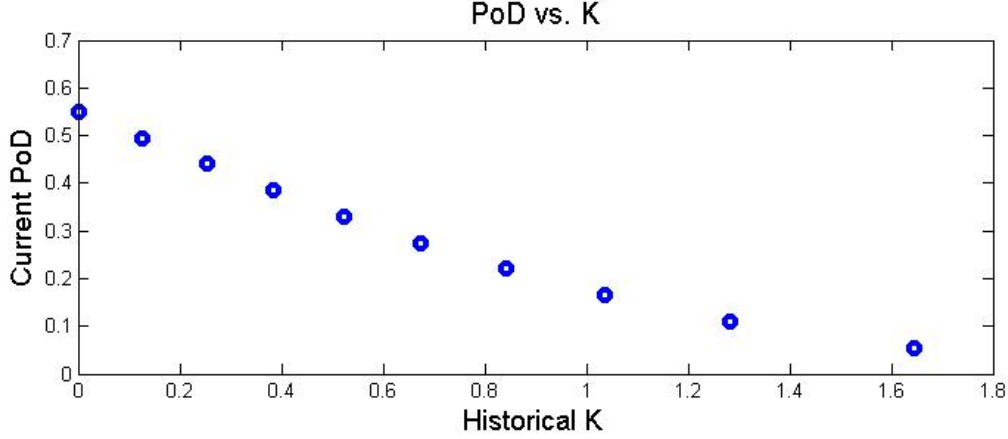


Figure 4.1: Current PoD vs. the historical default thresholds K .

space, going from safe names to risky ones and we set the current default rates to be above the historical levels⁷ for every name in the portfolio.

4.4.2 Family issues

In this section we will start by studying the impact of changing the *prior* q by varying its parametric shape. In the next section we will complete the analysis of the changes of q by focusing on changes of the dependency structure keeping fixed the parametric shape.

We used two different sets of distributions for this investigation: a set of elliptical distributions and a set of Archimedean copulas. In both cases, we measured the changes between p and q with respect to the four systemic risk measure mentioned earlier (BSI, SFM, JPoD and PAO). There are two reasons we choose these particular sets; firstly, they are among the most used distributions in financial modeling. The second reason is

⁷Please note that in the special case where $HistPoD_i \equiv PoD_i, \forall i$, i.e.

$$K_i = \Phi_i^{-1}(1 - PoD_i)$$

the CIMDO solution becomes trivial as the *posterior* will be identical to the *prior* ($p \equiv q$).

the tractability of the *prior* when integrated as the CIMDO methodology can be quite demanding in terms of computational costs.

Figure 4.2 shows the change in the four measures used when the *prior* is either a t-Student (with $\nu = 3$ and $\nu = 5$) or a standard normal distribution. In all cases, the variance/covariance matrix is such that the implied pairwise correlation $\rho = 0.3$. Calculations for this set was based on an adaptation of a numerical integration routine by Genz and Bretz (2002) to compute integrals on the *prior*⁸.

Few remarks on figure 4.2:

- The *JPoD* shows bigger dependency on the *prior* than the other three measures that are almost indifferent on the model chosen.
- In all cases, we can see how important is the choice of the *prior* q as its risk measure is always very close to the one derived from the *posterior* p .

For the second set of distributions, we explored Archimedean copulas (please see section 3.1 for details on copulas). Remember that copulas belonging to this class can be described by a single generator function ϕ that depends, for the families we considered, by a single parameter θ . In order to make a fair comparison, we decided to calibrate the copula parameters to provide a similar dependence structure as measured by the Kendall's τ . In particular, we used equation (3.20) to calibrate the generator's parameters in order to match a desired value for τ . In the study, we used Frank, Gumbel (also known as Gumbel-Hougaard) and Clayton copulas.

Copula Family	$\phi(\theta)$	ϕ/ϕ'	τ	θ
Clayton	$\frac{t^{-\theta}-1}{\theta}$	$\frac{t^{\theta+1}-t}{\theta}$	$\frac{\theta}{\theta+2}$	0.8571
Gumbel-Hougaard	$[-\ln]^\theta$	$\frac{t \cdot \ln}{\theta}$	$\frac{\theta-1}{\theta}$	1.4286
Frank	$-\ln \left(\frac{e^{-\theta \cdot t}-1}{e^{-\theta t}-1} \right)$	$-\frac{\phi}{\theta} \cdot \frac{e^{-\theta \cdot t}-1}{e^{-\theta \cdot t}}$	*	2.9174

Table 4.1: Copula's parameter calibrated to have $\tau = 0.3$.

* For the Frank family, the reported value of θ was found numerically by integrating ϕ/ϕ' over $[0, 1]$ according to (3.20).

Figure 4.3 shows the change in the four measures when the *prior* is one of the Archimedean copulas introduced earlier.

The remarks on figure 4.3 are similar to the ones for figure 4.2:

⁸Please see appendix B.1 for more details.

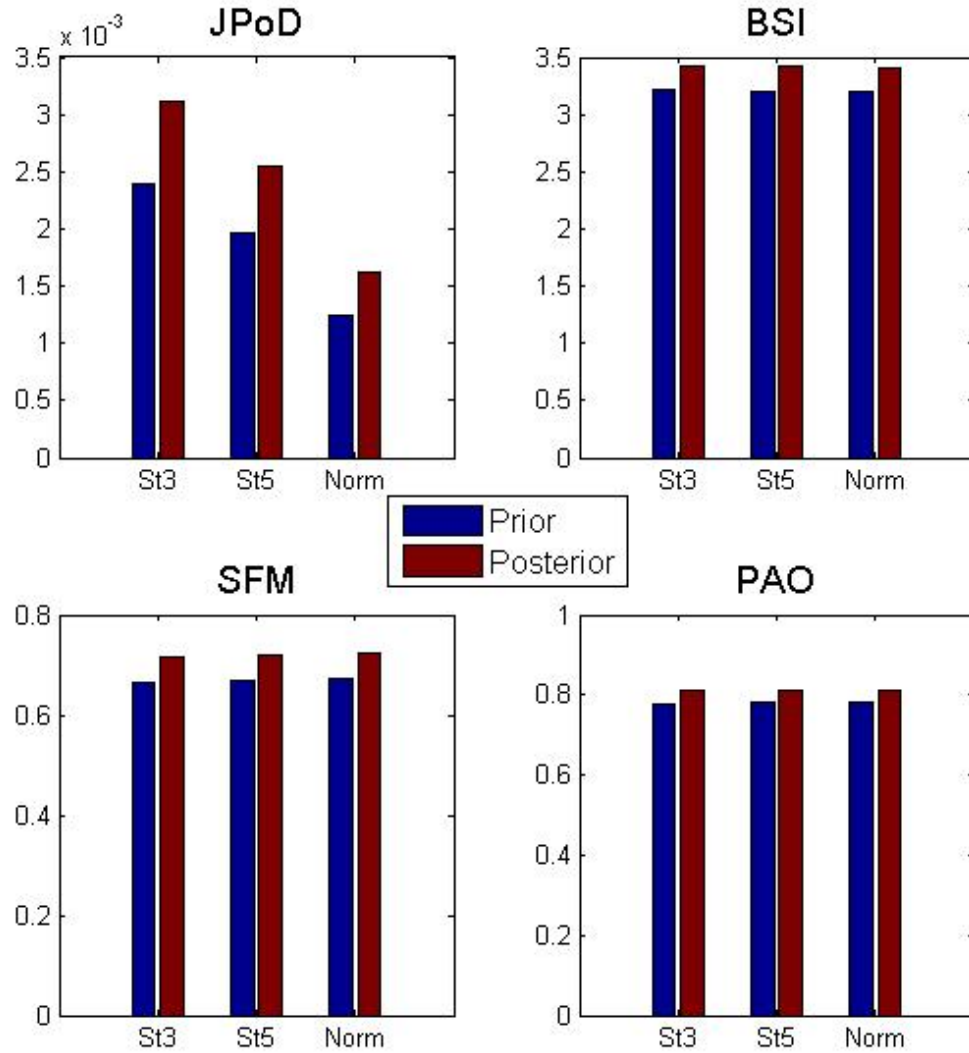


Figure 4.2: Comparison of two t-Student (one with $\nu = 3$, the other with $\nu = 5$) and a normal distribution. The variance/covariance matrix implies $\rho = 30\%$. Each box shows a risk measure computed on both the *prior* and the relative *posterior*.

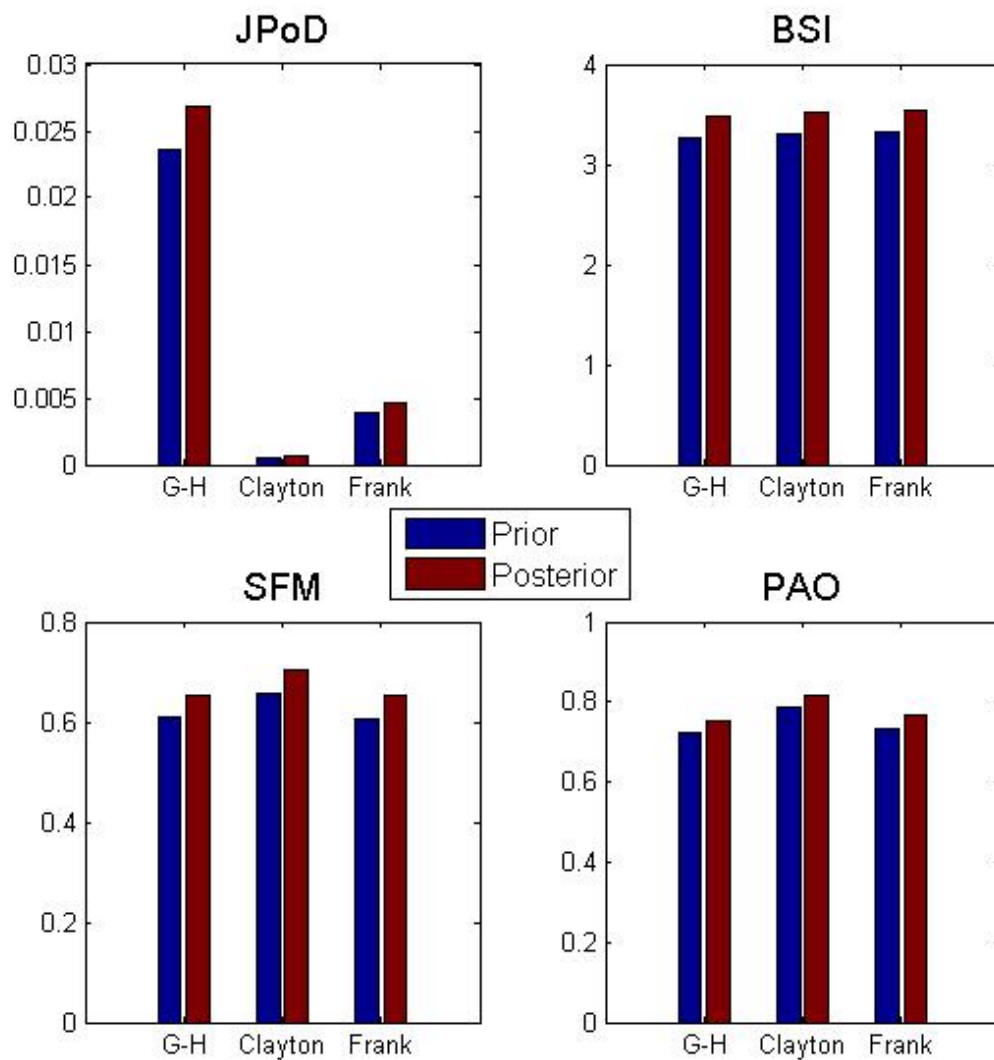


Figure 4.3: Comparison of 3 Archimedean copulas calibrated to exhibit Kendall's $\tau = 0.3$. Each box shows a risk measure computed on both the *prior* and the *relative posterior*.

- The *JPoD* shows much bigger dependency on the *prior* than the other three measures, especially compared to the *BSI* and the *PAO*; the *SFM* instead shows some limited variation.
- Again, the choice of the *prior* q is critical.

4.4.3 Parameter choice

After studying the choice of the parametric family, we focused on the issue of the impact of the model calibration. We found that the dimensionality n plays a big role when changing the parametrization and decided therefore to alter slightly our approach. Instead of using the base portfolio as in the previous and in the next sections, we considered a universe of names each with $HistPoD = 15\%$ and $PoD = 16.5\%$ (10% higher than the $HistPoD$). We then considered homogeneous portfolios of size ranging from 2 to 15 included⁹. We performed several studies with different parametric distributions but the qualitative results were similar and therefore we report only one set of findings from a Clayton copula. We changed both the dependency structure of the *prior* and the size of the portfolio n and studied the effects on the systemic risk measures. For every surface shown in figures 4.4, 4.5, 4.6 and 4.7, the y axis (on the right) is the dimension n while the x axis (on the left) follows the pairwise Kendall's τ (from 0 to 0.9); the parameter θ is varied according to the formula $\theta = \frac{2\tau}{1-\tau}$. If we indicate with $H(r)$ the value of a specific risk measure (either *JPoD*, *BSI*, *SFM* or *PAO*) under the distribution r , then the different surfaces reported for each measure are:

- The top row simply reports the measures under the *posterior* ($H(p)$, on the left) and the *prior* ($H(q)$, on the right).
- The bottom left graph shows the difference between the *posterior* and the *prior*: $H(p) - H(q)$
- The bottom right graph is similar to the previous one with the difference that what is reported is the relative change between the two distributions: $\frac{H(p)-H(q)}{H(q)}$

Few features are worth noting:

⁹The complexity of the CIMDO algorithm is exponential in n and medium-big sized portfolios are not tractable in this context.

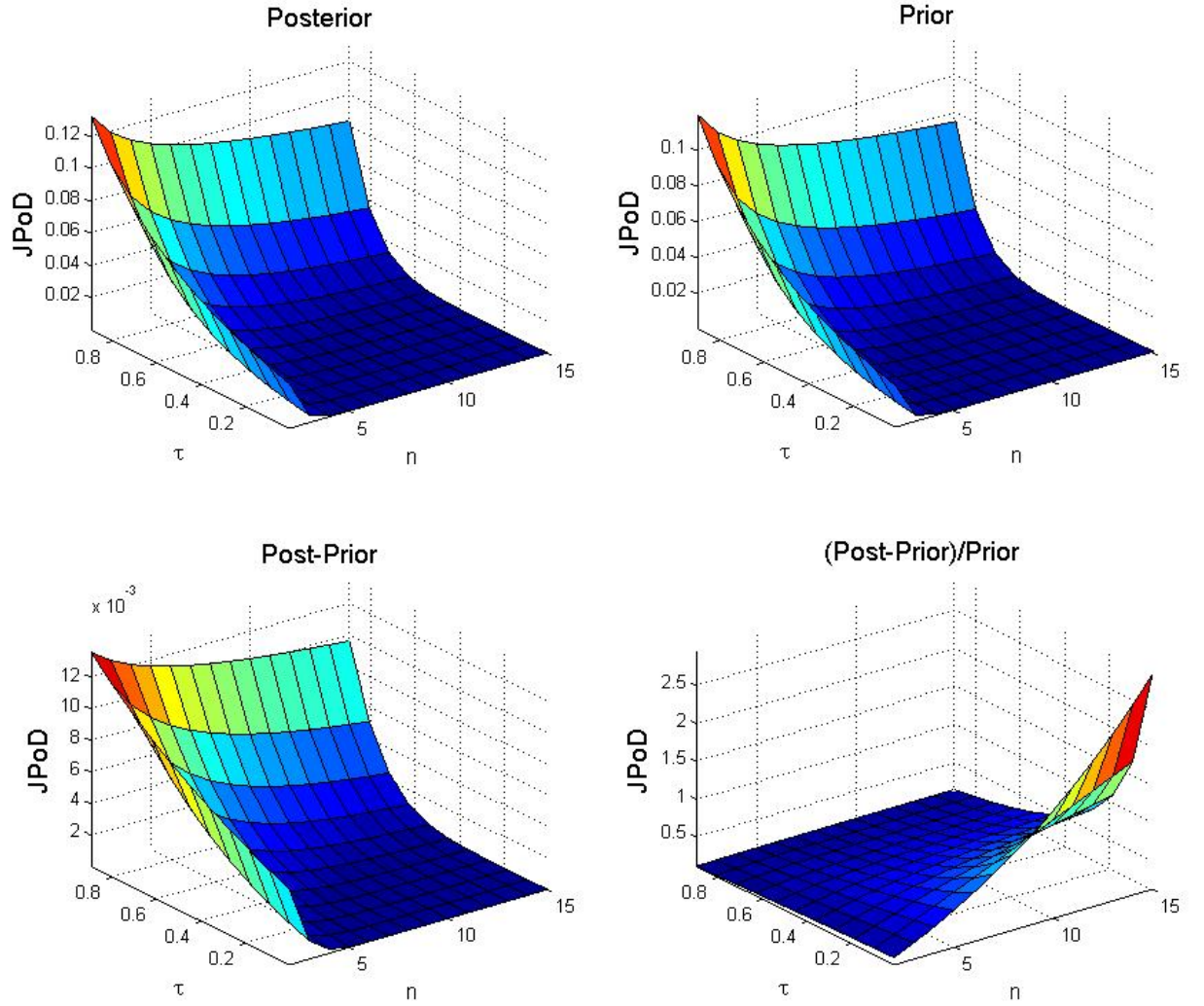


Figure 4.4: *Prior* and *posterior* $JPoD$ risk measure vs. τ and dimension n for a Clayton copula with $\theta = \frac{2\tau}{1-\tau}$.

The top row simply reports the measures under the *posterior* ($JPoD(p)$, on the left) and the *prior* ($JPoD(q)$, on the right). The bottom left graph shows the difference between the *posterior* and the *prior*, $JPoD(p) - JPoD(q)$, while the bottom right graph reports the relative change between the two distributions, $\frac{JPoD(p) - JPoD(q)}{JPoD(q)}$.

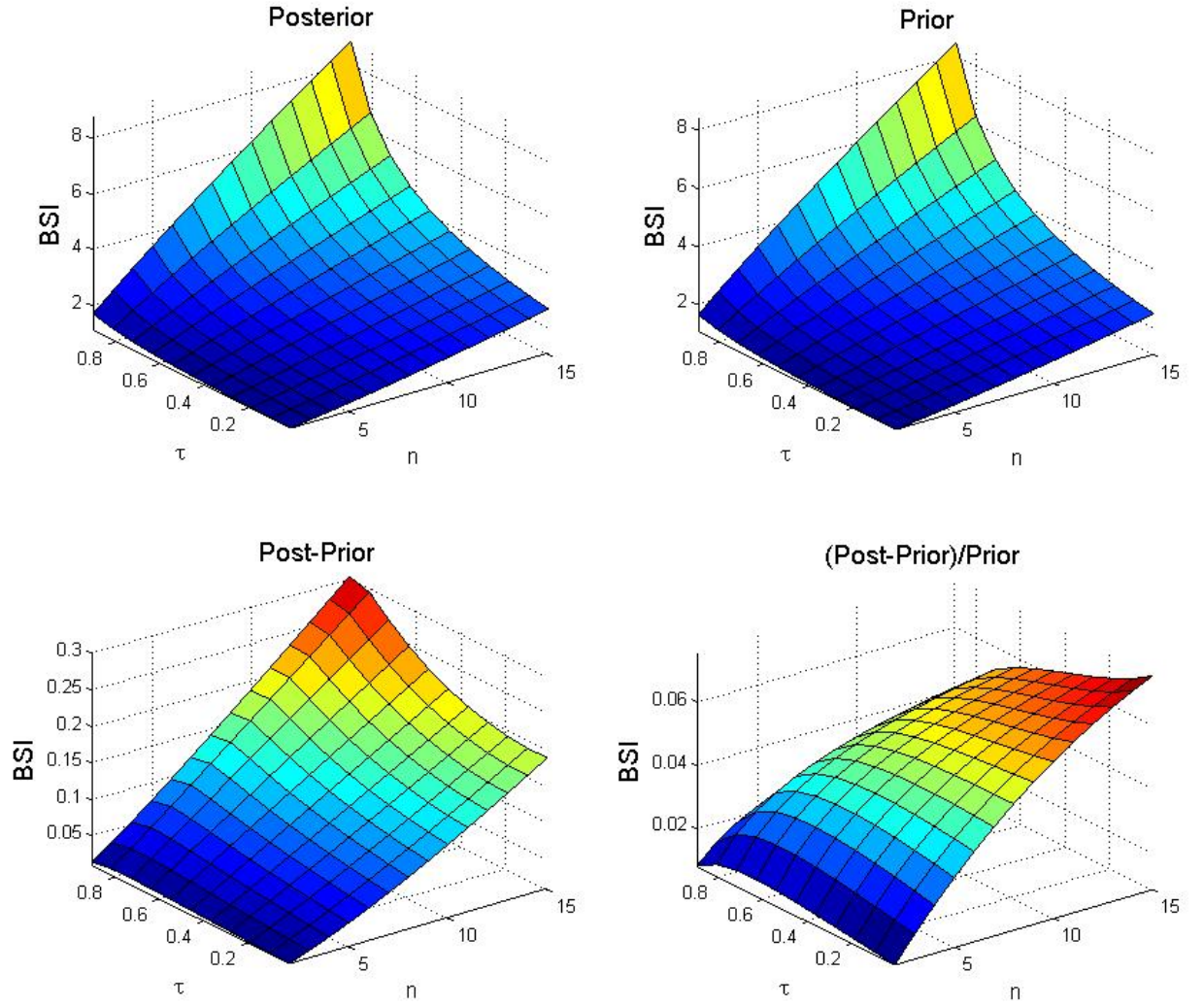


Figure 4.5: *Prior* and *posterior* BSI risk measure vs. τ and dimension n for a Clayton copula with $\theta = \frac{2\tau}{1-\tau}$.

The top row simply reports the measures under the *posterior* ($BSI(p)$, on the left) and the *prior* ($BSI(q)$, on the right). The bottom left graph shows the difference between the *posterior* and the *prior*, $BSI(p) - BSI(q)$, while the bottom right graph reports the relative change between the two distributions, $\frac{BSI(p) - BSI(q)}{BSI(q)}$.

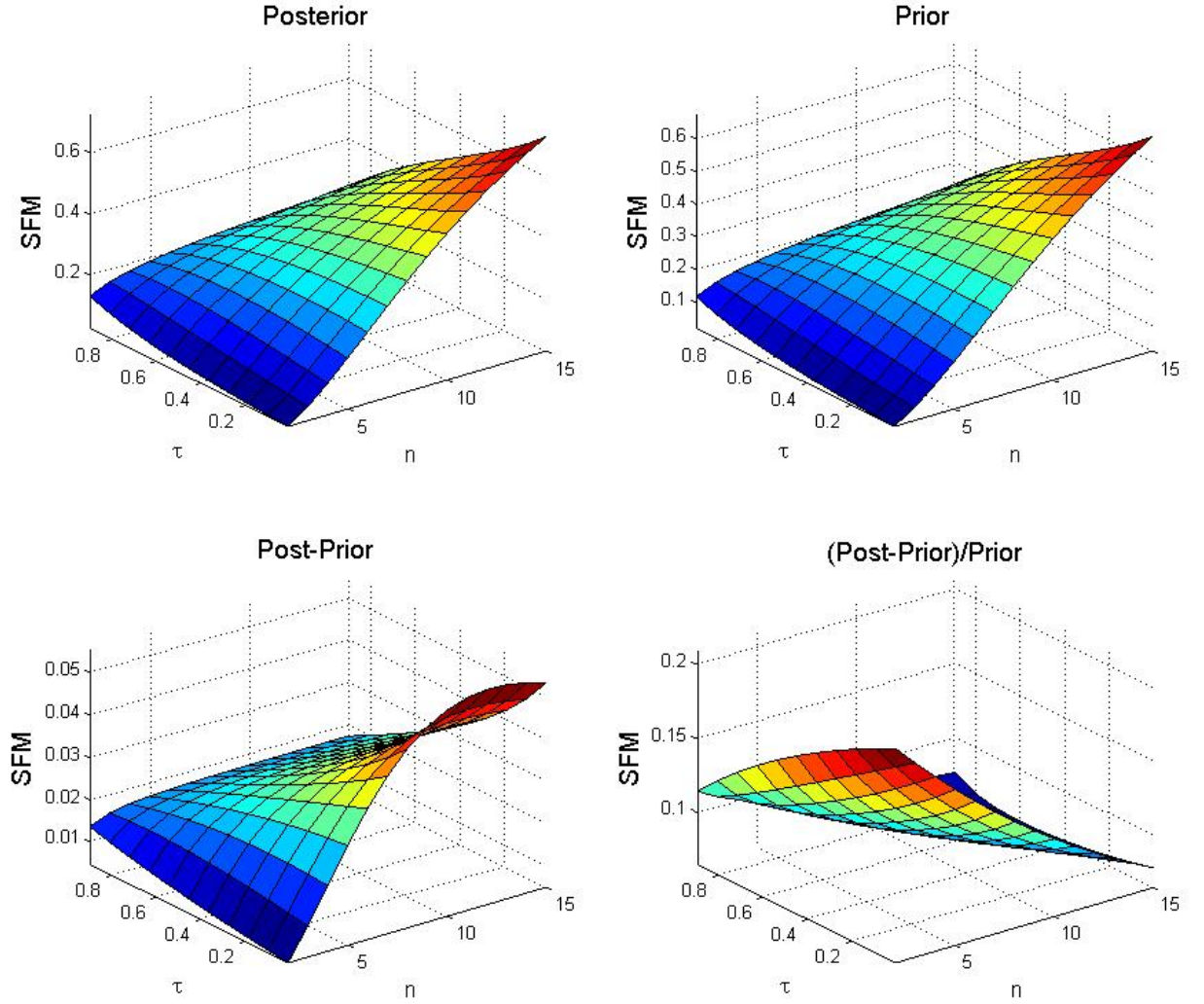


Figure 4.6: *Prior* and *posterior* SFM risk measure vs. τ and dimension n for a Clayton copula with $\theta = \frac{2\tau}{1-\tau}$.

The top row simply reports the measures under the *posterior* ($SFM(p)$, on the left) and the *prior* ($SFM(q)$, on the right). The bottom left graph shows the difference between the *posterior* and the *prior*, $SFM(p) - SFM(q)$, while the bottom right graph reports the relative change between the two distributions, $\frac{SFM(p) - SFM(q)}{SFM(q)}$.

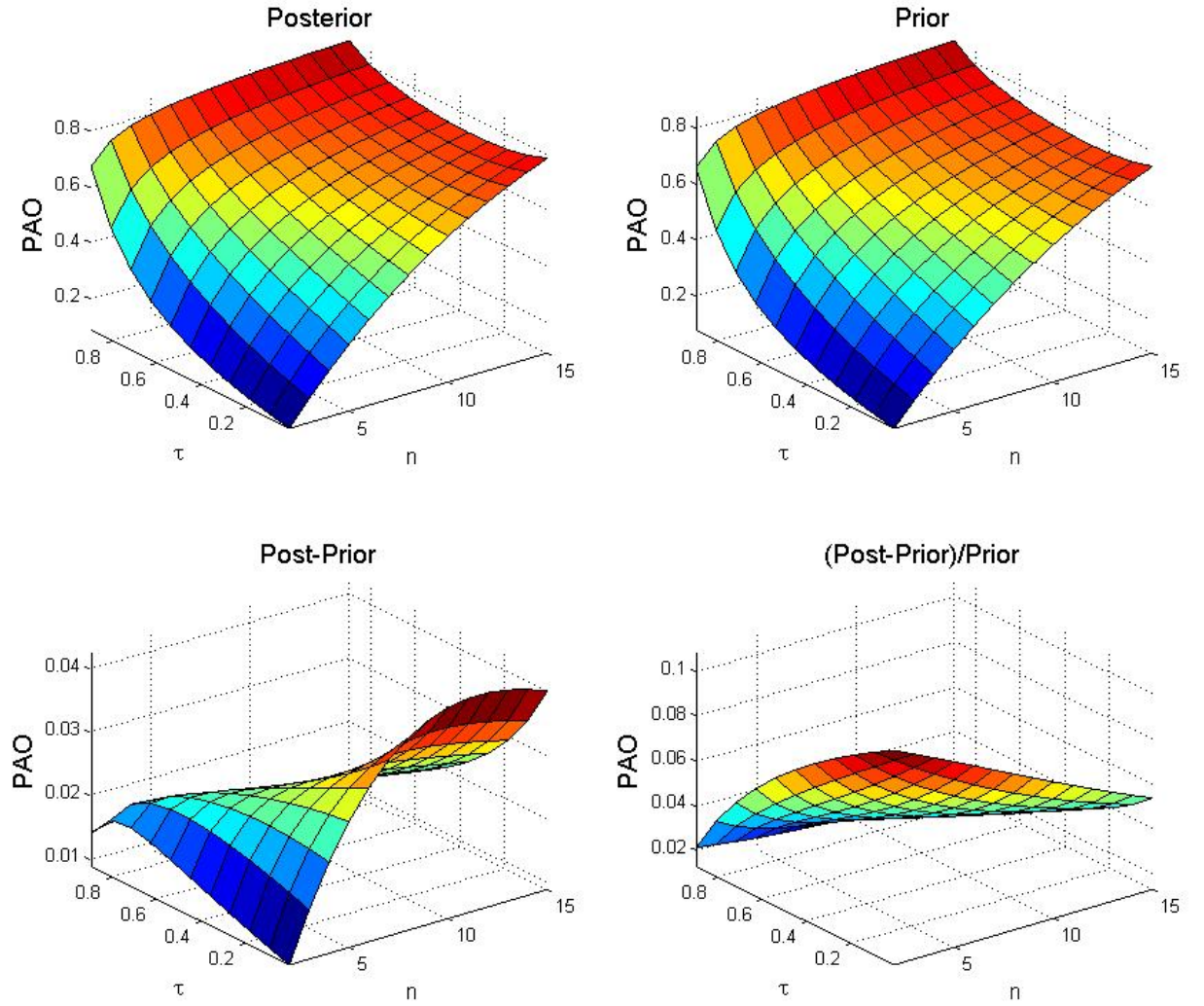


Figure 4.7: *Prior* and *posterior* PAO risk measure vs. τ and dimension n for a Clayton copula with $\theta = \frac{2\tau}{1-\tau}$.

The top row simply reports the measures under the *posterior* ($PAO(p)$, on the left) and the *prior* ($PAO(q)$, on the right). The bottom left graph shows the difference between the *posterior* and the *prior*, $PAO(p) - PAO(q)$, while the bottom right graph reports the relative change between the two distributions, $\frac{PAO(p) - PAO(q)}{PAO(q)}$.

- By looking at the surfaces in the first rows, we can see how the four risk measures differ in the way they react to changes in the parameter τ . On an intuitive level, we would expect all four risk indicators to increase when the parameter is increased no matter what dimensionality we are working with. This is in fact the case for the *JPoD*, the *BSI* and the *PAO*. The *SFM* measure, instead, shows a peculiar behavior as it decreases as the dependency is heightened for sizes $n \geq 4$.
- Looking still at the first row for each plot and focusing on the slope with respect to τ , we can notice that, while the *JPoD* shows a strong degree of change when moving τ for all the values of n considered, this is not the case for the *BSI* and the *PAO*. For the first, we can see how it is almost indifferent to τ for low levels of n ; the opposite is true for the *PAO*, as it shows flat gradients when n is high.
- Another difference between the *SFM* and the other three measures can be noticed when comparing the gradient of the surfaces across the dimension and the parameter direction. Unlike in the other measures, for the *SFM* the variation due to the parameter is comparable, in magnitude, to the one due to the dimensionality. We observed the same feature even after we tried to normalize the dimensionality effect by dividing $SFM(p)$ by the independent case (i.e. instead of using $SFM(p)$ directly, we considered $SFM(p)/SFM(\pi)$, where π is the density of the independent distribution).
- By looking at bottom row, instead, we can see how the magnitude of the difference between $H(p)$ and $H(q)$ is usually small (less than 5%); even when there are peaks where the relative change seems big (for example in the *JPoD* case for high dimensionality and low values of the parameter), we can see that they correspond to zones where both the value $H(q)$ and $H(p)$ are very close to zero, artificially bumping the relative change. This means that the values of the risk measures under the *posterior* distribution are never too far away from the corresponding ones from the *prior*. This result is in line with what we have already found in the previous case: the choice of q (and its parametrization in the specific) is crucial for the measurement of systemic risk using the CIMDO methodology.

4.4.4 Default probabilities

There are two inputs to the CIMDO methodology that are not directly part of the *prior*: the marginal probabilities of default *PoD* and the default thresholds K . In both

cases we will start from a *prior* based on a Clayton copula whose θ has been calibrated to imply a Kendall $\tau = 0.3$. We will also assume that the marginals follow standard normal distributions.

We started from the same base portfolio as in section 4.4.2 and then applied a multiplicative bump to the *PoDs* so that the perturbed ones $\tilde{PoD}$ will be given by

$$\tilde{PoD}_i = \omega \cdot PoD_i$$

where ω is the multiplicative factor.

In order to compare directly the impact of changes in *PoDs* and changes in *Ks*, instead of perturbing directly the default thresholds we decided to apply the same multiplicative factor to the historical default probabilities and then recalibrate the *K* accordingly:

$$\tilde{K}_i = 1 - \Phi_i^{-1}(\omega \cdot HistPoD_i)$$

Figure 4.8 shows the change in the systemic risk measures when varying these two inputs: the black line represents the effects of the changes to the marginals *PoD* that will be matched by the *posterior* thanks to the constraints (4.4) in the optimization problem; the light green line, instead, represents the effects of the changes to the historical probabilities of defaults *HistPoD* that in turn will be used to calculate new default thresholds *K*. The *x* axis shows the multiplicative factor ω .

It is interesting to note the following points:

- The lines cross when $\omega = 1$ as in this case we are back to the base portfolio.
- For *BSI*, *SFM* and *PAO*, the changes due to *PoD* are in absolute terms more substantial than the changes due to *HistPoD*. Table 4.2 shows the absolute ratio of the variations due to *PoD* over the equivalent statistic for *K*. This is not true for the *JPoD* measure for which the two effects are very similar, both in absolute terms and direction.
- All four measures increase when the marginals are increased. This is intuitive as it is quite natural to expect that the joint default probabilities will raise proportionally to the marginal probabilities.

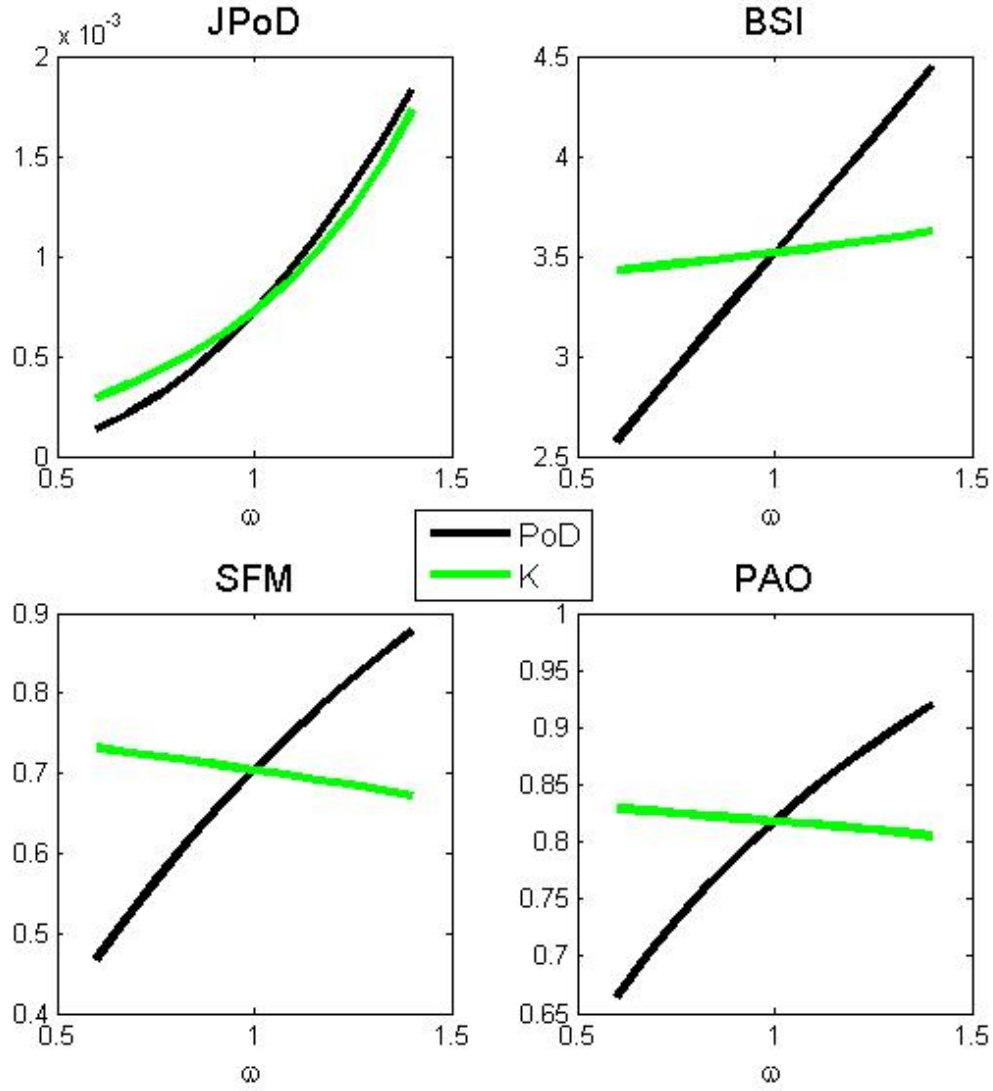


Figure 4.8: Change in the *posterior* systemic risk measures vs. changes in either the *PoD* or the *K*.

Measure	Variation for <i>PoD</i>	Variation for <i>HistPoD</i>	Absolute ratio
JPoD	0.0017	0.0014	1.1791
BSI	1.8665	0.2008	9.2964
SFM	0.4085	-0.0603	6.7707
PAO	0.2551	-0.0238	10.7045

Table 4.2: Variations due to either *PoD* or *HistPoD*.

- Regarding the changes due to *HistPoD*, please note that we have

$$K_i = \Phi_i^{-1}(1 - \text{HistPoD}_i)$$

where Φ_i is the marginal cumulative density function for the X_i variable. The K s will then decrease when the *HistPoD*s increase. This means that the left side of the x axis corresponds to higher values of K while the right side to lower values. The four measures react in different ways to changes in *HistPoD*:

- The *JPoD* shows a similar effect to the one observed for changes due to *PoD*; increasing *HistPoD* (i.e. decreasing K) causes *JPoD* to increase significantly.
- The *BSI* shows a similar direction as the one observed in the *JPoD* case but the effect is much smaller. Both the numerator and denominator involved in its calculations react in the same way and therefore the final effect on the measure is reduced.
- Both *SFM* and *PAO* increase with K although this effect is small and dimension dependent. In particular, the slope $\frac{\partial SFM}{\partial \text{HistPoD}}$ is slightly positive for very small values of n and then decreases steadily as n increases becoming quickly negative. Figure 4.9 shows the change of *SFM* (left) and *PAO* (right) when the dimensionality and the default thresholds are changed.

4.5 Conclusive remarks

We tried to analyze the question of the ability of the CIMDO methodology to ease the model uncertainty problem in the context of measuring systemic risk. The question is not trivial as theoretical results from Radev and ourselves show that the wrong choice of the *prior* can, in the worst case, completely void any systemic risk analysis performed on the obtained *posterior*. We therefore set ourselves to the task of studying the relationship between the input and the output distributions using few well-known systemic risk measures as proving grounds. The results obtained show that the decisions surrounding the *prior* are critical, both in terms of the choice of the distribution family and in terms of parameters calibration. We also tested the stability of the *posterior* with respect to changes to the other inputs of the CIMDO approach.

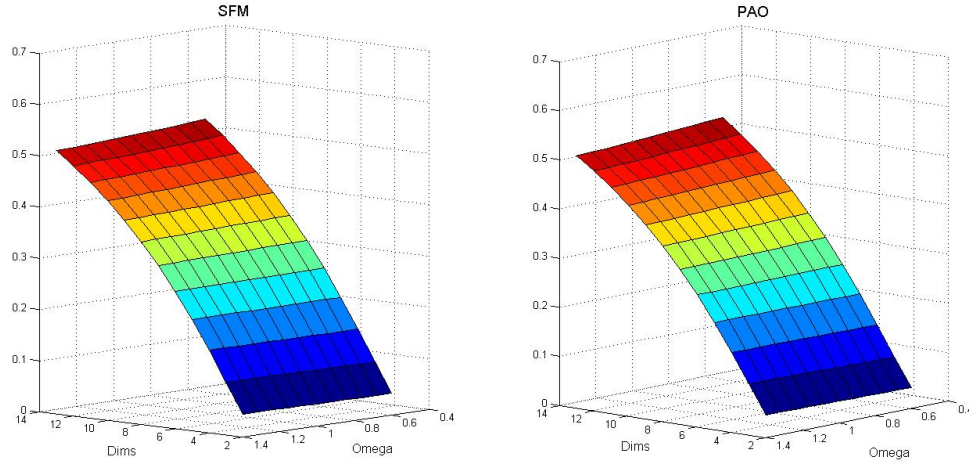


Figure 4.9: Change in the *posterior SFM* and *PAO* risk measures vs. changes in the dimensionality n and the default thresholds K .

In order to give some guidance to any practitioner willing to use the CIMDO approach, table 4.3 summarizes the order of magnitude of changes on the four systemic risk measures when moving some of the inputs of the methodology¹⁰.

Measure	PoD	$HistPoD$	τ	Copula
JPoD	0.0017	<u>0.0014</u>	<u>0.0428</u>	0.0261
BSI	1.8665	0.2008	<u>2.3506</u>	<u>0.0534</u>
SFM	<u>0.4085</u>	0.0603	0.3673	<u>0.0502</u>
PAO	<u>0.2551</u>	<u>0.0238</u>	0.0815	0.0635

Table 4.3: Comparison of variations in the *posterior* measures. For every measure, the biggest and the smallest values obtained are highlighted with a top or bottom line, respectively.

In general, each measure reacts in different ways to the changes in the inputs (see figure 4.10) but we can note the following:

- The two most important factors in terms of ability to affect the systemic risk observed through the CIMDO *posterior* are the model parametrization (represented by the value of τ in this example) and the marginal probabilities of default (PoD).

¹⁰All the results reported in the table have been obtained starting from the base portfolio described earlier.

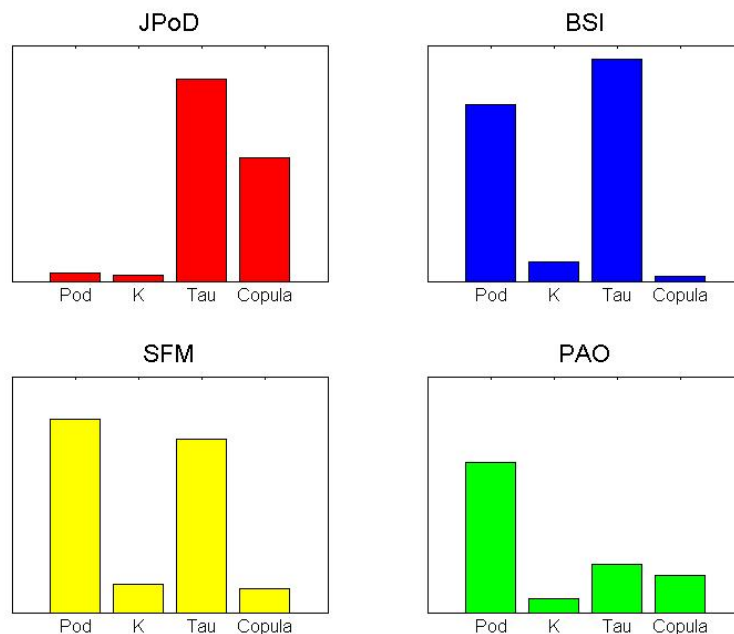


Figure 4.10: Ranges of values reached on few systemic risk measures bby varying the CIMDO inputs, both marginals (Pod and K) and joint (τ and copula type).

For every measure, they are the first and second most important factor (but for the $JPoD$, where the choice of the copula also plays a significant role).

- The choice of the copula seems less important when compared to the other inputs analyzed¹¹. The same can be said for the choice of the historical default threshold.

¹¹This could be due to the limited number of alternative copulas used. A more detailed and extended analysis could be the subject of further research in this area.

Chapter 5

A model of infectious defaults with immunization

This chapter introduces a new model that takes inspiration from the approach of Davis and Lo (2001). Some homogeneous assumptions have been relaxed while the contagion mechanism proposed is the result of two independent components: an infection attempt generated by defaulting firms and a failed defense from healthy ones. Theoretical results on both marginal and joint default distribution are presented. We also provide an efficient recursive algorithm for the portfolio loss distribution similar, in spirit, to the one commonly used for CID (conditionally independent) models. A version of the model with a simplified parameter structure is then applied to the problem of pricing (and hedging) CDO instruments and its performance compared to the standard one factor Gaussian model.

5.1 Introduction

One of the main problems in credit modeling is default clustering: it has been observed, especially during recessions, that defaults are not uniformly spaced in time but rather tend to concentrate over small periods of time. Even the most used approaches to introduce dependence among defaults (via exposure to common factors, by correlating the default intensity processes and by direct application of copula methods) struggle to replicate this observed pattern. One valid alternative is to add asymmetric dependency structure to the modeling framework; these effects, usually referred to as contagion or infection, not only have the ability to increase the probability of observing extreme losses in the portfolio but can also account for default clustering¹. The fact that they can also replicate the economical interdependence between firms is usually considered a plus.

From a pure technical point of view, adding these type of effects introduces a looping mechanism that makes calibration more problematic: the probability of default of each name can impact and is impacted by the probability of default of the others. Several attempts have been proposed in order to resolve the looping issue when adding such effects.

Jarrow and Yu (2001) suggest to separate the firms into two groups: primary names (i.e. belonging to the first group) can only default idiosyncratically while names in the second group can also default because of infection starting from the first set of names. Their work generalizes reduced-form models by making default intensities depend on counterpart default. The primary/secondary separation has been applied by other authors too, for example Rösch and Winterfeldt (2008): they start from a one-factor model where the number of defaults of primary names can affect the default probabilities of secondary names. In the usual tradition of CID (conditionally independent) models, the primary names latent variable is given by

$$R_i = \sqrt{\rho} \cdot F + \sqrt{1 - \rho} \cdot U_i$$

¹See Azizpour, Giesecke, et al. for evidence about default clustering and the role of contagion. See also Allen and Gale (2000) for a survey of the mathematical techniques used in modeling financial contagion.

where F is the common factor and U_i are idiosyncratic components. For secondary names instead we have

$$R_j = \sqrt{\rho} \cdot F + \sqrt{1 - \rho} \cdot U_j - \beta \cdot \frac{D^P}{N^P}$$

where N^P is the number of firms in the primary group, D^P of which are infecting firms. The parameter β governs the contagion effect ($\beta = 0$ is the no-contagion case). Conditioning on F , the default probabilities of primary names become independent and therefore one can easily obtain the conditional distribution of D^P ; conditioning also on this, the probability of default of names in the secondary group become independent and the final conditional portfolio loss distribution can be obtained. Integrating over the common factor (externally) and the distribution of D^P (internally), one obtains the portfolio loss distribution. Numerical results are shown for homogeneous portfolios only, although we suspect that a mixture of numerical integration and convolution techniques (for example, fast Fourier transform) could be effective.

Another paper starting from classical factor models but adding contagion mechanism is the one by Neu and Kühn (2004). Firms can have either mutually supportive or competitive relationships between them; a default will hence decrease the default probability of competitors but increase the same quantity for firms that had positive inflows from the name in distress (for example suppliers/clients).

Egloff, Leippold, and Vanini (2007) instead add micro-structural dependencies via a directed weighted graph and show how even well diversified portfolios carry significant credit risk when such inter-dependencies are accounted for. The open problem in their approach (as well as in other network based models) is the calibration of the weights that form the neural network.

Yu (2007) works with an intensity based model where the default intensities are driven also by past defaults history, in addition to exogenous factors. A special case of this model is the copula approach presented by Schönbucher and Schubert (2001).

Two other valuable readings are the works by Giesecke and Weber (2004) and by Frey and Backhaus (2010); the first adds contagion processes to standard common factor approaches while the second applies a Markov-chain model with default contagion to the problem of dynamically hedge CDO products.

Our approach belongs to a branch of the literature started by Davis and Lo (2001). In the original formulation of the model, Davis and Lo propose two possible causes of

default: a name can default either directly (*idiosyncratic factor*) or because infected by defaults of other entities in the portfolio (*contagion factor*). Note, however, that only names that default directly (as opposed by contagion) can infect other names: no domino effects are taken in consideration. Let Z_i represents the default variable for name i (i.e. $Z_i = 1$ iif name i defaults) in a portfolio of n entities; the authors model Z_i via the following decomposition

$$Z_i = X_i + (1 - X_i) \left[1 - \prod_{j \neq i} i(1 - X_j Y_{i,j}) \right] \quad (5.1)$$

where X_i and $Y_{i,j}, i \neq j$ are $n + n(n-1)$ independent Bernoulli variables ($i, j = 1, \dots, n$). From the latter equation it is quite clear how the two mechanisms mentioned earlier are implemented. In particular, the terms of the form $X_i \cdot Y_{i,j}$ explain how the contagion is spread across names: each variable $Y_{i,j}$ measures the probability that name i , once defaulted, infects name j . In order to get closed-form solutions, the authors impose very strict assumptions on the X, Y variables, namely:

$$P(X_i = 1) = p$$

$$P(Y_{i,j} = 1) = q$$

They are now requiring that all the names in the portfolio have the same probability of default and, in addition, the same probability of infecting other names. With these assumptions they are able to get their main results; let $G_n = Z_1 + \dots + Z_n$ be the number of defaults and let $C_j^h = h!/j!(h-j)!$. We then have:

$$P[G_n = k] = C_k^n \cdot \alpha(p, q, n, k)$$

$$\alpha(p, q, n, k) = p^k (1-p)^{n-k} (1-q)^{k(n-k)} + \sum_{i=1}^{k-1} C_i^k p^i (1-p)^{n-i} [1 - (1-q)^i]^{k-i} (1-q)^{i(n-k)}$$

They also provide equations for the expected value and the variance of G_n :

$$E[G_n] = n \cdot [1 - (1-p) \cdot (1-pq)^{n-1}]$$

$$var[G_n] = E[G_n] - (E[G_n])^2 + n(n-1) \cdot \beta(p, q, n)$$

where

$$\beta(p, q, n) = p^2 + 2p(1-p)[1 - (1-q)(1-pq)^{\bar{n}}] + (1-p)^2[1 - 2(1-pq)^{\bar{n}} + (1-2pq + pq^2)^{\bar{n}}]$$

and $\bar{n} = n - 2$.

Please note also that all the results they obtained are for the distribution of the number of defaults in the portfolio and not about its loss distribution, i.e. how much losses each default scenario implies. In order to link the two, the additional assumption that every name has the same loss given default (LGD) is required.

Other authors have used Davis and Lo's original approach as starting point of more sophisticated models. Sakata, Hisakado, and Mori (2007) extended Davis and Lo's original model by assuming that idiosyncratic defaults might in fact be avoided with the help of non-defaulted names. In their model there is hence not only a (negative) contagion that causes more defaults but also a positive one (called recovery spillage) that prevents entities from defaulting. Mathematically, the model has a shape that is similar to equation (5.1) but extended to show an elegant default/survival symmetry:

$$Z_i = X_i \cdot \prod_{j \neq i} [1 - Y'_{i,j}(1 - X_j)] + (1 - X_i) \left[1 - \prod_{j \neq i} (1 - X_j Y_{i,j}) \right] \quad (5.2)$$

The main results are similar to the ones obtained in Davis and Lo (2001) in terms of both complexity and assumptions required (homogeneous portfolio). The paper is completed by an interesting discussion on the limit of the loss distribution when the size of the portfolio $n \rightarrow \infty$. The authors also try to fit their model to the loss distribution implied from iTraxx-CJ with mixed results.

Another very interesting work is the one by Cousin, Dorobantu, and Rullière (2013) that extends the original approach in many ways. First, the authors use a multiple time step model where defaults happening at the previous time interval can still cause contagion later on. In addition, they consider the case where more than one infection is needed to cause a default by contagion and, from a theoretical point of view, they relax several of the original paper assumptions. Unfortunately, the results for the most generic specification of their model involve quite heavy calculations that are computationally very demanding. In fact, the numerical applications shown are based on

assumptions that are in line with Davis and Lo (2001) in term of complexity².

We can notice how from one side we have very generic extensions of the model that are cumbersome to use³ and from the other side extremely fast versions that can only deal with unrealistic instances of the model due to the severe restrictions imposed on the starting assumptions. We suspect that neither one of the two cases is actually usable in real life and that this has been the main reason this type of model has not been applied to a wider set of problems. Ideally, one would like to keep the contagion mechanism from the generic form of the model but benefit from the speed of the restricted versions: the model we are presenting in this chapter is a compromise that moves toward that ideal solution. From the computational point of view, portfolio losses can be calculated via a recursive algorithm with reasonable costs. From the modeling side, we can specify different default probabilities as well as different infection rates. We also added an immunization mechanism that healthy names can use to protect themselves against infections. The only concession we had to make with respect to the original model is the following: instead of having several different contagion shocks when a name defaults (one for each name in the portfolio), we now model the event that the defaulting name might infect the entire system. We had to exclude the possibility that one defaulting firm infects some firms but not others in favor of the more draconian scenario in which it either tries to infect all names or none. If from one side this difference makes our model *not* an extension of Davis and Lo's one, from the other it is crucial in obtaining efficient loss distribution algorithms.

The rest of the chapter is structured as follows: section 2 introduces our new model and provides explanations and examples regarding the parameters choice. Section 3 presents some theoretical results and describes the algorithm for the portfolio loss distribution. In section 4 a version of the model with a simplified parameter structure is applied to the problem of pricing (and hedging) CDOs and its performances compared against the standard one factor Gaussian copula model. Section 5 concludes.

²In particular, the authors require that the Bernoulli variables used are exchangeable.

³Montecarlo based with usually big dimensionality: in the original approach the number of possible scenarios is an impressive 2^n .

5.2 A new model

The following equations show our new approach:

$$Z_i = X_i + (1 - X_i) \cdot (1 - U_i) \cdot \left[1 - \prod_{i \neq j} (1 - X_j \cdot V_j) \right] \quad (5.3)$$

$$P[X_i = 1] = p_i, \quad P[U_i = 1] = u_i, \quad P[V_i = 1] = v_i \quad (5.4)$$

where we have postulated the existence of $3N$ mutually independent Bernoulli variables $X_1, \dots, X_n, V_1, \dots, V_n, U_1, \dots, U_n$. As in the Davis and Lo paper, we also have the two mechanisms for default seen earlier, i.e. idiosyncratic defaults and contagion. In particular, name i infects name j only if two independent conditions are satisfied:

1. **Infection attempt from i :** name i defaults idiosyncratically and attempt to spread the infection to *all* other names. This component is driven by the V variables.
2. **Failed defense from j :** name j fails to defend itself from *every* possible infection. This component is driven by the U variables.

Note the independence assumption among the building blocks of (5.3); it will be crucial when we will prove theoretical results in later sections.

Another assumption of the original approach we find limiting is the request that all names share the same loss given default (LGD). In reality, LGDs do vary by name and we have tried to include this aspect of reality in our modeling efforts. In order to work with discrete losses, we will assume that the LGDs can take the following form

$$LGD_i = D \cdot \begin{cases} d_i & \text{idiosyncratic default} \\ c_i & \text{infection} \end{cases} \quad (5.5)$$

where each d_i and c_i are integers representing units of losses and D is instead the minimum loss amount we can represent in our model⁴.

⁴This strategy to discretize the LGDs is used, for example, in Andersen, Sidenius, and Basu (2003). To give an example, assume all the names in the portfolio have LGDs in the set $\{33\%, 40\%, 44\%\}$; if we choose D to be 10%, the above names will be associated with $\{3, 4, 4\}$ units of losses, respectively. The approximated LGDs will hence be $\{30\%, 40\%, 40\%\}$, leading to rounding errors of magnitude $\{3\%, 0\%, 4\%\}$. Note how a finer choice of D , for example $D = 5\%$, leads to the approximated LGDs

The set of parameters $(p_i, u_i, v_i, d_i, c_i)$ for a given name i defines its behavior in the model. In particular:

High values of v represent names that, upon default, are extremely infective. This can be used for pivotal names that are regarded as critical for the well being of the entire system.

Low values of u should be used for names that have weak defenses or, differently put, are strongly dependent on the health of the rest of the system.

Meaning of p This parameter controls only the probability of idiosyncratic default, not the final probability of default for the name. It is important to understand that the final probability of default is the result not only of p but also of the rate of infections from the other names as well as its own ability to gain immunization via u .

Use of d_i and c_i By varying (d_i, c_i) we can model very different behavior between normal defaults and systemic scenarios. For example, increasing c with respect to d will make the contagion effects stronger, putting more emphasis (and risk) on systemic events rather than on the "normal" idiosyncratic defaults. Note that we can even set $d = c = 0$ to add names whose only function is to bring bad news to the rest to the system, prompting an infection.

Before moving to the properties of the model it is worth noting a final point. The independence assumption between the building blocks should be seen as a tractable way of creating dependency. We use a specific combinations of independent variables to generate dependent ones, in exactly the same way it is done, for example, in factor models (there the idiosyncratic and common factor variables play the role of our X , U and V).

$\{35\%, 40\%, 45\%\}$ that show smaller rounding errors ($\{2\%, 0\%, 1\%\}$), at the cost of dealing with higher units of losses ($\{7, 8, 9\}$). Higher units of losses translates in heavier calculation burdens as the dimension of the vector used for the portfolio loss distribution increases. This instance of the common compromise between accuracy and speed makes the choice of D in real life applications a problem more complicated than one suspects. Parcell (2006) offers a very effective strategy to minimize the rounding errors.

5.3 Theoretical results

In this section we will provide some useful results: at first, we will explore single name default probability under the model assumptions as well as the expected LGD. Then we will present results on the joint default/survival probabilities. Finally, we will show how to efficiently calculate the portfolio loss distribution. A few necessary tools and additional notational shortcuts will be introduced along the way.

Proposition 6 *Let $A \subseteq \{1 \dots, n\}$; the probability that at least one name in A spreads an infection is given by the quantity I_A defined as*

$$I_A := \left[1 - \prod_{j \in A} (1 - p_j \cdot v_j) \right] \quad (5.6)$$

Proof:. The probability that an infection starts from inside A is given by

$$P\{X_j \cdot V_j = 1 \text{ for at least one } j \in A\} = 1 - P\{X_j \cdot V_j = 0, \forall j \in A\} \quad (5.7)$$

We can now use the independence assumption on both X and V to get

$$\begin{aligned} P\{X_j \cdot V_j = 0, \forall j \in A\} &= \\ \prod_{j \in A} P\{X_j \cdot V_j = 0\} &= \\ \prod_{j \in A} (1 - P\{X_j \cdot V_j = 1\}) &= \\ \prod_{j \in A} (1 - p_j \cdot v_j) & \end{aligned} \quad (5.8)$$

where the last passage is obtained thanks to the independence of X_j and V_j . ■

5.3.1 Marginal distribution

We will now focus on single name properties. Let $\bar{A}$ be the complement of set A . The first result gives the formula for the probability of default of single names⁵:

Proposition 7 *Let $\tilde{p}_i := P\{Z_i = 1\}$. We have:*

$$\tilde{p}_i = p_i + (1 - p_i) \cdot (1 - u_i) \cdot I_{\bar{i}} \quad (5.9)$$

⁵This result is especially useful for calibration; we can assume that $\tilde{p}$ will be available (for example via calibration from CDS data) and use equation (5.9) to match some of the other parameters.

Proof. Let remember the form of variable Z_i as expressed by equation (5.3)

$$Z_i = X_i + (1 - X_i) \cdot (1 - U_i) \cdot \left[1 - \prod_{j \neq i} (1 - X_j \cdot V_j) \right] \quad (5.10)$$

We have that

$$P\{Z_i = 1\} = P\{Z_i = 1|X_i = 1\} \cdot P\{X_i = 1\} + P\{Z_i = 1|X_i = 0\} \cdot P\{X_i = 0\} \quad (5.11)$$

It is easy to see that

$$P\{Z_i = 1|X_i = 1\} = 1 \quad (5.12)$$

and

$$P\{X_i = 1\} = 1 - P\{X_i = 0\} = p_i \quad (5.13)$$

so that the only part left to calculate is $P\{Z_i = 1|X_i = 0\}$ for which we have

$$P\{Z_i = 1|X_i = 0\} = P\left\{(1 - U_i) \cdot \left[1 - \prod_{j \neq i} (1 - X_j \cdot V_j) \right] = 1\right\} \quad (5.14)$$

Both variables in the last expression can only take binary values (0, 1) so their product can only be 1 if they both take value 1. This implies that

$$P\left\{(1 - U_i) \cdot \left[1 - \prod_{j \neq i} (1 - X_j \cdot V_j) \right] = 1\right\} = P\left\{(1 - U_i) = 1, \left[1 - \prod_{j \neq i} (1 - X_j \cdot V_j) \right] = 1\right\} \quad (5.15)$$

We can use the independence assumption between the various building blocks to split the right hand side as

$$P\{(1 - U_i) = 1\} \cdot P\left\{\left[1 - \prod_{j \neq i} (1 - X_j \cdot V_j) \right] = 1\right\} \quad (5.16)$$

end hence, thanks to proposition 6

$$P\{Z_i = 1|X_i = 0\} = (1 - u_i) \cdot I_i^- \quad (5.17)$$

that concludes the proof. ■

An intuitive explanation can be obtained by observing that name i can default in two ways: idiosyncratically (with probability p_i) or by contagion if it survives $(1 - p_i)$, fails to defend itself $(1 - u_i)$ and an external infection is active. The independence of the building blocks and the previous proposition then are sufficient to prove the assert. So name i has a p_i probability of idiosyncratic default and a $(\tilde{p}_i - p_i)$ probability of being infected. We can use the above result to obtain the formula for the expected value of LGD_i :

$$LGD_i = D \cdot [p_i \cdot d_i + (\tilde{p}_i - p_i) \cdot c_i] \quad (5.18)$$

5.3.2 Joint default/survival probability

In this subsection we will show results regarding joint survival/default events. The notation might seem complicated at first so we will add plenty of examples. Let introduce ps , us , bd and id as the vectors defined by

$$\begin{aligned} ps_i &= (1 - p_i) \cdot (1 - u_i) && \text{(non immune) partial survival} \\ us_i &= (1 - p_i) \cdot u_i && \text{(full) unconditional survival} \\ bd_i &= p_i \cdot (1 - v_i) && \text{(non infective) benign default} \\ id_i &= p_i \cdot v_i && \text{infective default} \end{aligned}$$

They will be useful for simplifying formulas later on. Let A be a subset of $\{1, \dots, n\}$ and let also $m_A = |A|$. Given an integer $h \leq m_A$ and two vectors x and y , let introduce the following operator

$$\Lambda_A^h(x, y) = \sum x_{i_1} \cdots x_{i_h} \cdot y_{j_1} \cdots y_{j_{m_A-h}}$$

where the summation is intended over all possible choices of indices $i_1, \dots, i_h$ and $j_1, \dots, j_{m_A-h} \in A$. It is the sum of all possible products of factors of m_A elements of the 2 vectors x, y with exactly h elements taken from x and the rest from y .

Similarly, but now for 3 vectors (x, y, z) rather than 2, we can define

$$\Theta_A^{h,k}(x, y, z) = \sum x_{i_1} \cdots x_{i_h} \cdot y_{j_1} \cdots y_{j_k} \cdot z_{r_1} \cdots z_{r_{m_A-h-k}}$$

It is the sum of all possible products of factors of m_A elements of the 3 vectors x, y, z with exactly h elements of x and k of y ⁶.

To clarify further, let's look at a simple example; assume $A = \{1, 2, 3\}$, we have

$$\Lambda_A^1(x, y) = x_1y_2y_3 + y_1x_2y_3 + y_1y_2x_3$$

as the 3 terms on the right hand side represent all the possible ways we can take products of elements of x and y with only one element coming from x . Similarly:

$$\begin{aligned}\Lambda_A^0(x, y) &= y_1y_2y_3 \\ \Lambda_A^2(x, y) &= x_1x_2y_3 + x_1y_2x_3 + y_1x_2x_3 \\ \Lambda_A^3(x, y) &= x_1x_2x_3\end{aligned}$$

$$\Theta_A^{1,1}(x, y, z) = x_1y_2z_3 + x_1z_2y_3 + y_1x_2z_3 + z_1x_2y_3 + y_1z_2x_3 + z_1y_2x_3$$

where the last equation shows an example applied to Θ instead.

Finally, let's use the following shortcut

$$\Pi_A(x) = \prod_{i \in A} x_i$$

The following result allows to calculate the probability of every possible combination of survival/default (please refer to the appendix A.2 for the proof):

Proposition 8 *Let A and B be two non empty subsets of $\{1, \dots, n\}$ with $A \cap B = \emptyset$. Define with C the set of all the names that are neither in A nor in B : $C = \overline{A \cup B}$ and let finally $P(A, B)$ represent the probability of all names in A default while all names in B survive, i.e.*

$$P(A, B) = P(Z_i = 1, Z_j = 0, \forall i \in A, \forall j \in B)$$

Then, we have

$$P(A, B) = P_1 + P_2 + P_3 \tag{5.20}$$

⁶Note that we have

$$\Lambda_A^h(x, y) = \Theta_A^{h, m_A - h}(x, y, \cdot) \tag{5.19}$$

so that we can consider Λ_A as a special case of Θ_A .

with

$$\begin{aligned}
P_1 &= \Pi_A(bd) \cdot (1 - I_C) \cdot \Pi_B(1 - p) \\
P_2 &= \left[\sum_{h=1}^{m_A} \sum_{k=0}^{m_A-h} \Theta_A^{h,k}(id, ps, bd) \right] \cdot (1 - I_C) \cdot \Pi_B(us) \\
P_3 &= \left[\sum_{h=0}^{m_A} \Lambda_A^h(p, ps) \right] \cdot I_C \cdot \Pi_B(us)
\end{aligned} \tag{5.21}$$

The three components reflect different cases; in particular, the first corresponds to an infection-free world where names in A default on their own in non-infective way and names in B survive without needing necessarily to immunize themselves. In the second component the infection is internal to A and names in B need full immunization. The same is true for the last case where the infection is now starting in C and names in A default either idiosyncratically (in an infective or benign way) or by contagion.

When $B = \bar{A}$, equation (5.20) gives the probability of exclusive defaults of the entire subset A :

$$P(A, \bar{A}) = \Pi_A(bd) \cdot \Pi_{\bar{A}}(1 - p) + \left[\sum_{h=1}^{m_A} \sum_{k=0}^{m_A-h} \Theta_A^{h,k}(id, ps, bd) \right] \cdot \Pi_{\bar{A}}(us) \tag{5.22}$$

Similarly, when $B = \emptyset$, we obtain the probability that at least every name of A defaults:

$$P(A, \emptyset) = \left\{ \Pi_A(bd) + \left[\sum_{h=1}^{m_A} \sum_{k=0}^{m_A-h} \Theta_A^{h,k}(id, ps, bd) \right] \right\} \cdot (1 - I_{\bar{A}}) + \left[\sum_{h=0}^{m_A} \Lambda_A^h(p, ps) \right] \cdot I_{\bar{A}} \tag{5.23}$$

Calculating Θ_A and Λ_A

Equations (5.20)-(5.22)-(5.23) are very interesting but they would be of no practical use if we cannot provide an efficient way of calculating Θ_A and Λ_A . Luckily there is a relationship linking Θ_A and Λ_A to Θ_W and Λ_W where W is a strict subset of A (please refer to the appendix A.3 for the proof):

Proposition 9 *Let A be a subset of $\{1, \dots, n\}$ containing at least 2 elements and let $t \in A$; let also $W = A/t$ be the subset of names of A excluding t . The following identities hold:*

$$\Lambda_A^h(x, y) = x_t \cdot \Lambda_W^{h-1}(x, y) + y_t \cdot \Lambda_W^h(x, y) \tag{5.24}$$

$$\Theta_A^{h,k}(x, y, z) = x_t \cdot \Theta_W^{h-1,k}(x, y, z) + y_t \cdot \Theta_W^{h,k-1}(x, y, z) + z_t \cdot \Theta_W^{h,k}(x, y, z) \tag{5.25}$$

As an example, let again $A = \{1, 2, 3\}$; we have already seen that

$$\Lambda_A^1(ps, bd) = ps_1 bd_2 bd_3 + bd_1 ps_2 bd_3 + bd_1 bd_2 ps_3.$$

Using the right hand side of (5.24) instead we would have

$$\Lambda_A^1(ps, bd) = ps_1 \cdot \Lambda_{A/1}^0(ps, bd) + bd_1 \cdot \Lambda_{A/1}^1(ps, bd)$$

and the identity is satisfied as

$$\Lambda_{A/1}^0(ps, bd) = \Lambda_{\{2,3\}}^0(ps, bd) = bd_2 bd_3$$

and

$$\Lambda_{A/1}^1(ps, bd) = \Lambda_{\{2,3\}}^1(ps, bd) = ps_2 bd_3 + bd_2 ps_3.$$

We can recursively use (5.24) and (5.25) until we reach (on the right hand side) the case $W = \{i\}$ for which we have⁷:

$$\Lambda_{\{i\}}^h(ps, bd) = 0, \text{ if } h > 1 \text{ or } h < 0$$

$$\Lambda_{\{i\}}^1(ps, bd) = ps_i, \quad \Lambda_{\{i\}}^0(ps, bd) = bd_i$$

$$\Theta_{\{i\}}^{h,k}(id, ps, bd) = 0, \text{ if } h + k > 1 \text{ or } \min(h, k) < 0$$

$$\Theta_{\{i\}}^{1,0}(id, ps, bd) = id_i, \quad \Theta_{\{i\}}^{0,1}(id, ps, bd) = ps_i, \quad \Theta_{\{i\}}^{0,0}(id, ps, bd) = bd_i$$

In the homogeneous case we can get simplified expressions for Θ and Λ :

$$\Theta_A^{h,k}(x, y, z) = C_h^{m_A} \cdot C_k^{m_A-h} \cdot x^h \cdot y^k \cdot z^{m_A-h-k}$$

$$\Lambda_A^h(x, y) = C_h^{m_A} \cdot x^h \cdot y^{m_A-h}$$

5.3.3 Portfolio loss distribution

Let L_n represents the unit of losses of a portfolio with n names; for a given integer h , we will show an efficient algorithm for calculating $P\{L_n = h\}$ similar, in spirit, to the one presented by Andersen, Sidenius, and Basu (2003) for conditionally independent

⁷These results come straight from the definitions of Θ and Λ .

models. In such modeling framework, one can obtain conditional default probabilities that are independent from each other by conditioning on the value of a common factor. Once the conditional default probabilities are obtained, one needs only to compute their convolution to get the portfolio loss distribution (under the chosen value of the common factor). Integrating numerically over the common factor leads to the unconditional portfolio loss distribution.

Andersen, Sidenius, and Basu showed that an efficient way of performing the convolution of the independent conditioned default probabilities is to construct the portfolio loss distribution by adding each name one by one via the following recursive relationship

$$P_{n+1}(h) = P_n(h) \cdot (1 - p_j) + P_n(h - d_j) \cdot p_j \quad (5.26)$$

where p_j is the conditional probability of default of name j , d_j are the units of losses associated with it while $P_n(h) = P\{L_n = h\}$, i.e. the probability of having h units of losses in a portfolio of n names.

The intuition behind (5.26) is the following. There are two ways in which we can obtain h units of losses when adding an extra name: either we already reached h losses with the previous names and the new one survived (first component) or the name defaulted adding d_j units of losses to the $h - d_j$ ones already reached before adding j (second component).

In our model, we can exploit the independence of the building blocks to obtain a similar recursive algorithm that constructs the portfolio loss distribution by adding an extra name; this will be the aim of the rest of this section, in which, for ease of explanation, we will omit the subscript j for the name we are adding.

Let start by defining L_n^C as the the units of losses due to contagion events and $L_n^I = L_n - L_n^C$ for losses due to idiosyncratic effects. When adding an extra name to the calculations, there are two situations we need to face: either we are in a *contaminated* world, i.e. there has already been an infectious default, or we are in an *infection-free* case. In the latter case we need to consider the scenario of the first infective default that triggers potential losses, i.e. losses that have not been realized yet but have accumulated as the result of previous names surviving without getting full immunization. Let L_n^R represent the number of units of such potential losses and note that they will become L_n^C units of realized losses if an infection appears. Let define two quantities that will

take the role played by $P_n(h)$ in (5.26):

$$\begin{aligned}\alpha_n(h, k) &:= P\{L_n^I = h, L_n^C = 0, L_n^R = k, \text{ no infection active}\} \\ \beta_n(h, k) &:= P\{L_n^I = h, L_n^C = k, L_n^R = 0, \text{ at least one infection active}\}\end{aligned}\quad (5.27)$$

So $\alpha_n(h, k)$ represent the probability of realizing h units of losses in an uncontaminated world of n names, in which there are also k units of losses at risk should an infection appear. On the other hand, $\beta_n(h, k)$ represent the probability of realizing $h + k$ units of losses (in a contaminated universe of n names), of which h are due to idiosyncratic defaults and k are due to pure infection. The above quantities are sufficient to describe the distribution of L_n :

$$P\{L_n = h\} = \sum_k \beta_n(k, h - k) + \sum_k \alpha_n(h, k) \quad (5.28)$$

The following system of equations can be used to calculate the $\alpha_n(\cdot, \cdot)$ and $\beta_n(\cdot, \cdot)$:

Proposition 10 *We have*

$$\begin{aligned}\alpha_{n+1}(h, k) &= (1 - p) \cdot u \cdot \alpha_n(h, k) + \\ &\quad (1 - p) \cdot (1 - u) \cdot \alpha_n(h, k - c) + \\ &\quad p \cdot (1 - v) \cdot \alpha_n(h - d, k) \\ \beta_{n+1}(h, k) &= (1 - p) \cdot u \cdot \beta_n(h, k) + p \cdot \beta_n(h - d, k) + \\ &\quad (1 - p) \cdot (1 - u) \cdot \beta_n(h, k - c) + \\ &\quad p \cdot v \cdot \alpha_n(h - d, k)\end{aligned}\quad (5.29)$$

with the following boundary conditions:

$$\begin{aligned}\alpha_0(0, 0) &= 1 \\ \alpha_0(i, j) &= 0 \quad \forall (i, j) \neq (0, 0) \\ \beta_0(i, j) &= 0 \quad \forall i, j\end{aligned}\quad (5.30)$$

Proof:. In order to obtain a set of equations for $\alpha_n(\cdot, \cdot)$, consider that there are 3 ways of reaching $\alpha_{n+1}(h, k)$ starting from $\alpha_n(h, k)$ and adding a new name:

1. Full survival

$$(1 - p) \cdot u \cdot \alpha_n(h, k) \quad (5.31)$$

The name survives $(1 - p)$ and protects itself from future aggressions (u) . No losses are realized neither potential ones added.

2. Partial survival

$$(1 - p) \cdot (1 - u) \cdot \alpha_n(h, k - c) \quad (5.32)$$

The name survives $(1 - p)$ but fails to protect itself against future aggressions $(1 - u)$. Its c units of losses are at risk should an infection spread.

3. Non-infectious default

$$p \cdot (1 - v) \cdot \alpha_n(h - d, k) \quad (5.33)$$

The name defaults directly (p) but it is not trying to start an infection $(1 - v)$.

Similarly, there are 4 ways of reaching $\beta_{n+1}(h, k)$:

1. Full survival

$$(1 - p) \cdot u \cdot \beta_n(h, k) \quad (5.34)$$

The name survives $(1 - p)$ and protects itself against the current and future infections (u) .

2. Default by contagion

$$(1 - p) \cdot (1 - u) \cdot \beta_n(h, k - c) \quad (5.35)$$

The name survives $(1 - p)$ but fails to protect itself against the existing infection $(1 - u)$.

3. Direct default

$$p \cdot \beta_n(h - d, k) \quad (5.36)$$

The name defaults (p) and in this case we don't need to consider separately the cases in which it spreads or not the infection as we are already in an infected world.

4. First infection

$$p \cdot v \cdot \alpha_n(h - d, k) \quad (5.37)$$

The name defaults (p) and spreads the contagion (v) in a previously uncontaminated world causing the k units of potential losses to become real ones.

Putting together the previous equations, we get system (5.29). ■

Figure 5.1 provides a visual description of the relationships between the $\alpha(\cdot, \cdot)$ and the $\beta(\cdot, \cdot)$ on the (h, k) space.

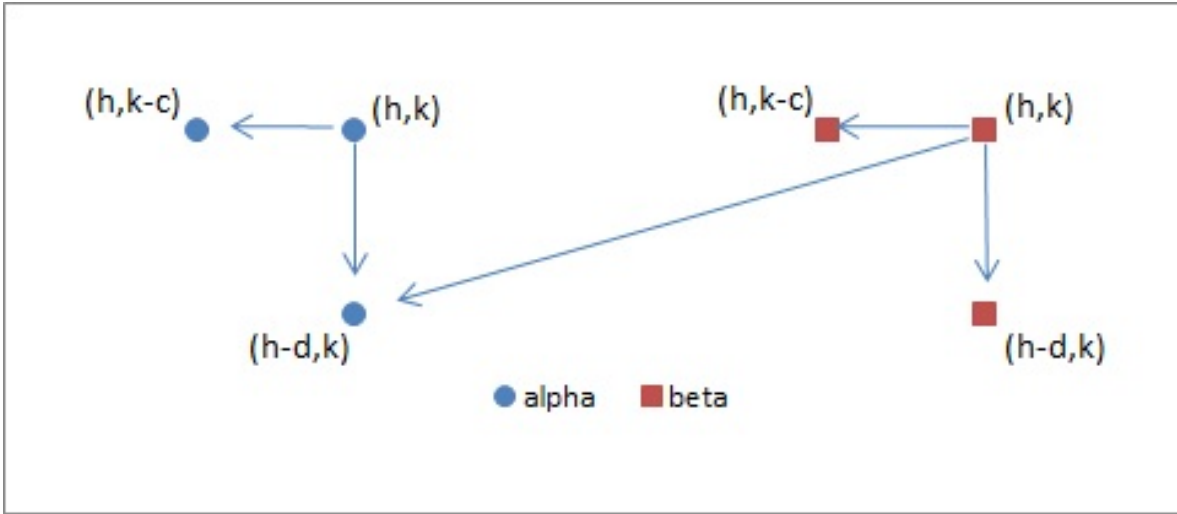


Figure 5.1: Recursive links between the $\alpha(\cdot, \cdot)$ and the $\beta(\cdot, \cdot)$: the vertical axis represents h changes while the horizontal one follows k .

We hope that is clear now why we had to abandon the approach of letting each name infect independently or not each other name in the portfolio. Instead of considering only the two cases of contaminated or infection-free world, we would have had to "remember" this type of information for every name already added, making the calculations rapidly unusable as we would have needed 2^n states by the last step.

The system of equations (5.29)-(5.30) is easily implemented via a recursive algorithm for both $\alpha_n(\cdot, \cdot)$ and $\beta_n(\cdot, \cdot)$. It is therefore quite efficient to get the portfolio loss distribution $P\{L_n = h\}$ for real-life values of n . In the appendix B.2, there is a description of the algorithm and a performance comparison with the recursive algorithm presented in Andersen, Sidenius, and Basu (2003).

5.4 An application to CDO Pricing

In this section we will test the new model by pricing CDO contracts. Handling portfolio credit products requires the ability to model and calculate the entire portfolio loss

distribution at several time steps. This application will hence show the tractability of the system of equations (5.29)-(5.30).

5.4.1 A restricted version of the model

The model described so far is very generic and highly tractable. One issue a potential user would face is the calibration of the many parameters involved: $5 \cdot n$ parameters (for each name, we need p_i, v_i, u_i, d_i, c_i). It is usually difficult to obtain a stable calibration in models with an high number of interconnected parameters if data is scarce. In this section we will therefore specify a particular strategy to reduce the degree of freedom of the model. We will assume the following shape for p , v and u :

$$\begin{aligned} p_i &= (1 - \omega) \cdot \tilde{p}_i \\ v_i &= \mu \cdot (1 - \tilde{p}_i) \\ u_i &= 1 - \frac{\tilde{p}_i - p_i}{(1 - p_i) \cdot I_i} \end{aligned} \tag{5.38}$$

Some considerations

- ω is a factor that controls how much probability of default comes from idiosyncratic effects versus contagion ones. In particular, $\omega = 0$ represents the no-contagion case, while an higher value for ω causes most of the losses to derive from contagion events.
- The v_i are in an inverse relationship with $\tilde{p}_i$; this reflects the fact that healthier firms have a bigger impact in case of idiosyncratic default than riskier ones; the market is expecting default of risky firms (hence the high probability of default) and therefore the shock when the event finally happen is minor. The strength of the effect is controlled by μ .
- The u_i are chosen according to (5.9) in order to satisfy marginal constraints. It is worth noting that we get the calibration to the marginal info embedded in our multivariate model almost by construction.

Regarding the parametrization of the LGDs, we decided to simply set⁸

$$d_i = c_i = \text{round}(LGD_i/D), \forall i \tag{5.40}$$

⁸An alternative is to specify c_i (assuming d_i is given, for example via $d_i = \text{round}(LGD_i/D)$) as in the previous case) in order satisfy (5.18) as closely as possible (the solution might not be exact given

In this case the LGDs become independent on the mechanism that caused the default (idiosyncratic or contagion).

5.4.2 The portfolio

We had some difficulties in retrieving the actual quotes for CDO tranches on the main indices (ITraxx and CDX) as well as their portfolio composition. We hence decided to use a realistic yet artificial set of quotes as portfolio composition, whose details can be found in the appendix C. The dimensionality of the portfolio is set to $n = 100$.

5.4.3 Operational ranges

We compared our new approach against the standard one factor Gaussian (OFG) copula model; readers unfamiliar with it might refer to the section 3.2 for more info. In order to make a fair comparison between the two models, we fixed μ^9 and let only ω vary, resulting in an ω skew. Figure 5.2 shows the range of possible spreads for an equity tranche (0 – 3%) that the two models can reach as ω and ρ vary¹⁰. The model behavior with respect to equity tranche pricing is critical for calibration as every other tranche is priced as difference of equities. From this point of view, we can happily notice that the spread is a monotonic function of ω , making calibration unique (when a solution exists) and the slope between the two models is very similar. It is also worth noting that the new model can reach a wider range of spreads compared to OFG, meaning that it can still be calibrated on prices where the standard OFG would instead fail.

A similar analysis can be done for other tranches with the difference that we now have two inputs to move: the parameter at the detachment and the one at the attachment¹¹. Instead of a single line, we will have now a surface of possible spreads; for

that c_i has to be an integer and (5.18) is specified in the continuous):

$$c_i = \text{round} \left(\frac{\tilde{p}_i \cdot LGD_i / D - p_i \cdot d_i}{\tilde{p}_i - p_i} \right) \quad (5.39)$$

A heuristic version of this strategy is to allow values for c_i to vary only in the set $\{d_i - 1, d_i, d_i + 1\}$.

⁹We set $\mu = 0.01$. The value might seem quite low at first sight but remember that we only need one name to infect the rest of the system; even small values of μ (and therefore v) cause significant contagion when we are dealing with portfolios of medium-big size.

¹⁰We will use the acronym MCM to refer to the contagion model in plots and tables.

¹¹In a real calibration exercise, the value of the parameter for the attachment would have already been calibrated to match the previous tranche.

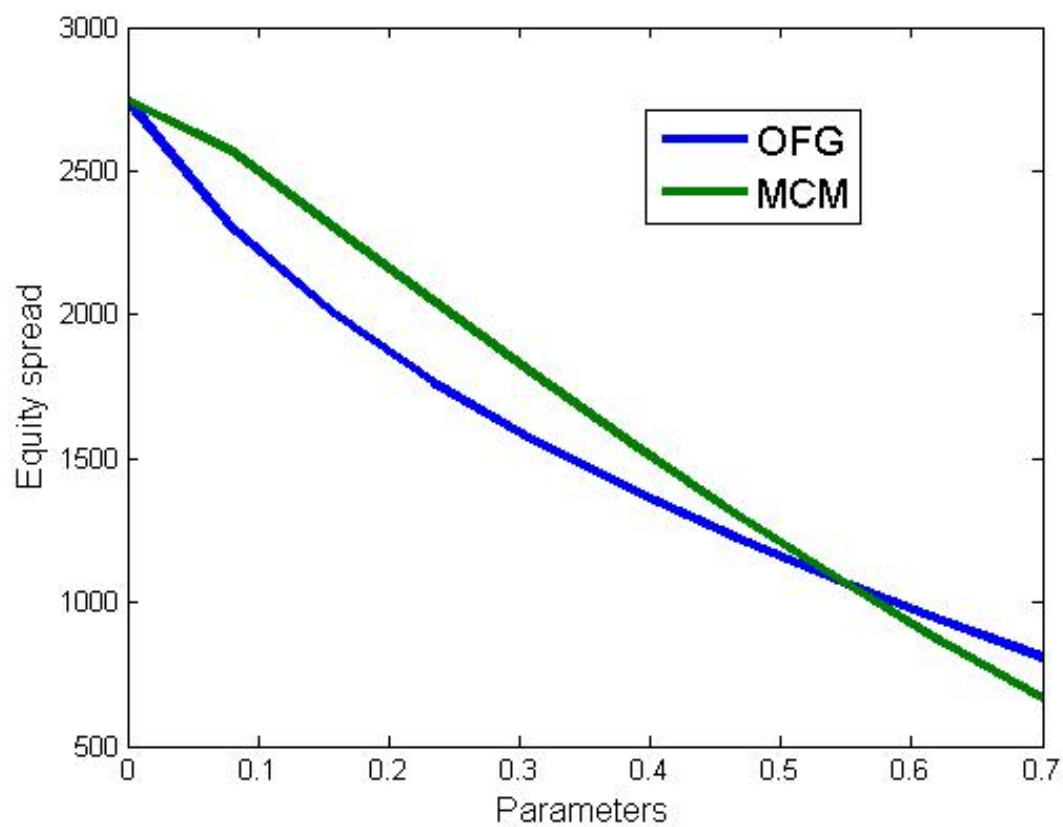


Figure 5.2: Price ranges: CDO spread for 0-3% tranche when varying the value for ω (contagion model) and ρ (OFG).

example, figure 5.4 shows such surfaces for a 6% to 9% tranche. Figures 5.3, 5.5, 5.6 deal with the other tranches.

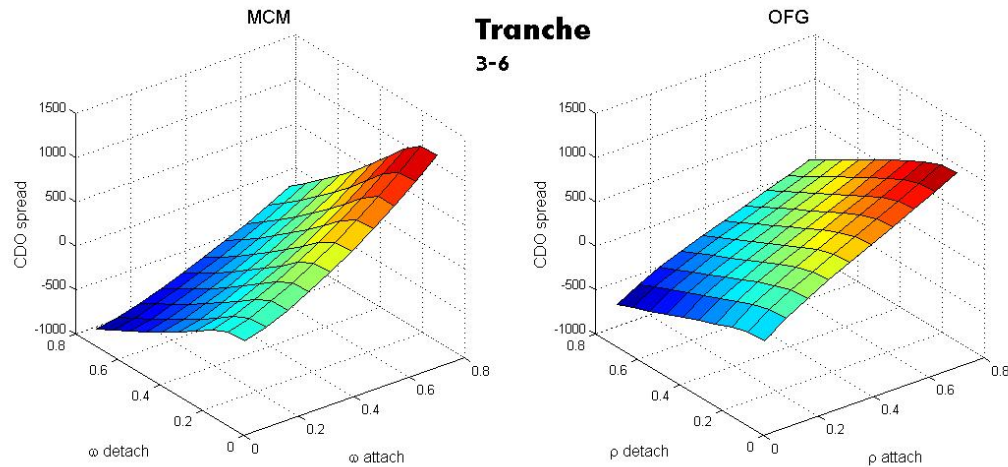


Figure 5.3: Price ranges: CDO spread for 3-6% tranche when varying the values for ω (contagion model) and ρ (OFG).

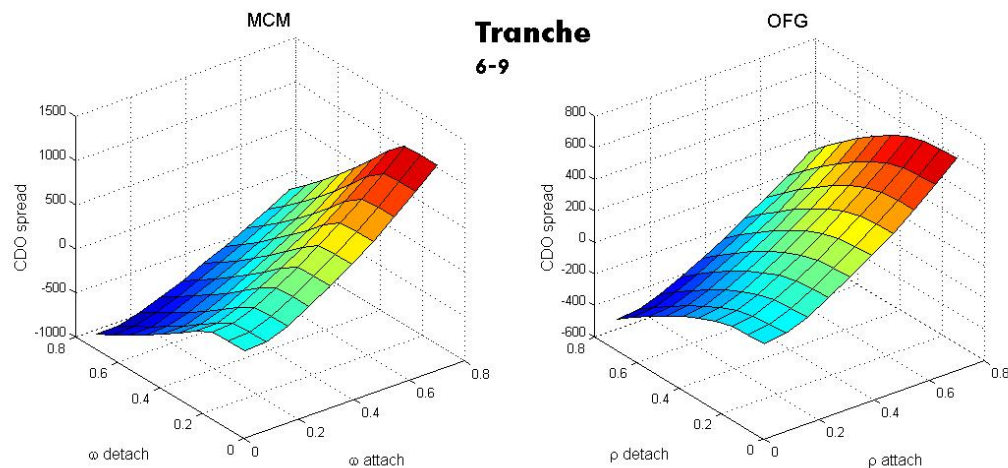


Figure 5.4: Price ranges: CDO spread for 6-9% tranche when varying the values for ω (contagion model) and ρ (OFG).

Table 5.1 shows the minimum and maximum spreads theoretically possible under the two models. Note that the negative spreads are just a realization of the shortcomings of the base approach when we let the two parameters move freely (the model used to price the attachment equity and the one used for the detachment one are not mathematically

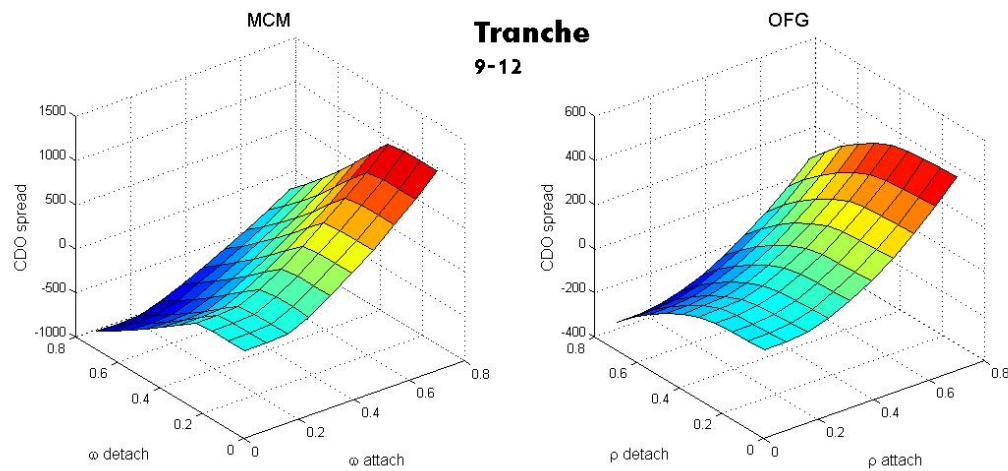


Figure 5.5: Price ranges: CDO spread for 9-12% tranche when varying the values for ω (contagion model) and ρ (OFG).

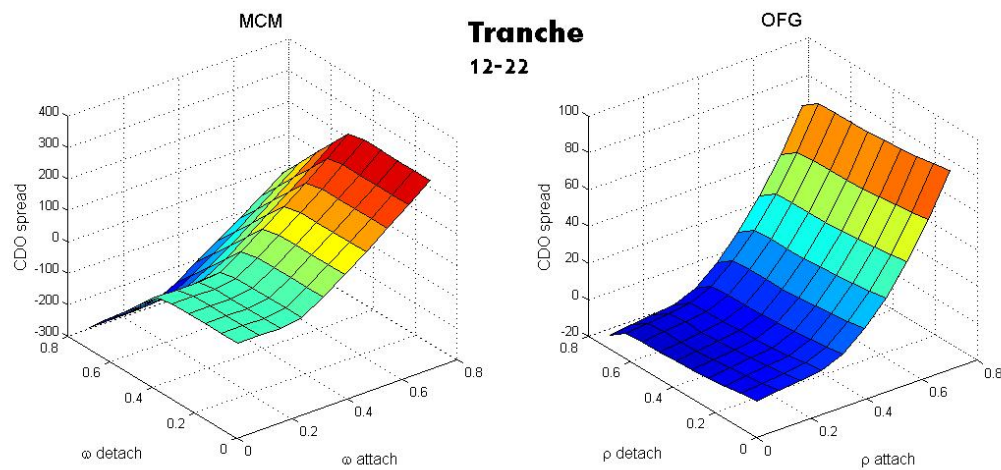


Figure 5.6: Price ranges: CDO spread for 12-22% tranche when varying the values for ω (contagion model) and ρ (OFG).

consistent with each other). Note also that the spreads that can be achieved with the contagion model are far higher than those reachable with the OFG one.

Tranche %	Min OFG	Max OFG	Min MCM	Max MCM
0-3	712.41	2,276.98	565.14	2,281.60
3-6	-519.27	1,206.86	-790.74	1,403.82
6-9	-407.55	746.53	-817.83	1,331.02
9-12	-275.11	478.60	-788.65	1,260.01
12-22	-12.23	88.18	-229.84	303.13

Table 5.1: Spread theoretical ranges for OFG and contagion model.

5.4.4 A calibration exercise

Because of the difficulties of accessing actual quotes, we decided to test calibration abilities of the model via a different route: we started from a realistic correlation skew for a OFG model and used it to generate quotes for five tranches (0%-3%, 3%-6%, 6%-9%, 9%-12% and 12%-22%); we then calibrated the contagion model based on these quotes and observed how severe was the skew in the ω space.

Tranche %	$\rho\%$	Quote	$\omega\%$
0-3	27.00	1,666.33	34.95
3-6	38.00	295.40	26.60
6-9	45.00	136.96	25.89
9-12	52.00	68.60	28.19
12-22	66.00	41.95	34.46

Table 5.2: OFG ρ and quotes versus calibrated ω and quotes ($\mu = 0.01$, maturity 5Y). All the marks were successfully calibrated.

Table 5.2 provide details on the calibration exercise. The error is extremely small but the most positive aspect, in our opinion, is the fact that the smile, shown in figure 5.8, is much less pronounced for the contagion model than for the OFG one. A flatter ω skew translates into a less arbitrage prone modeling framework.

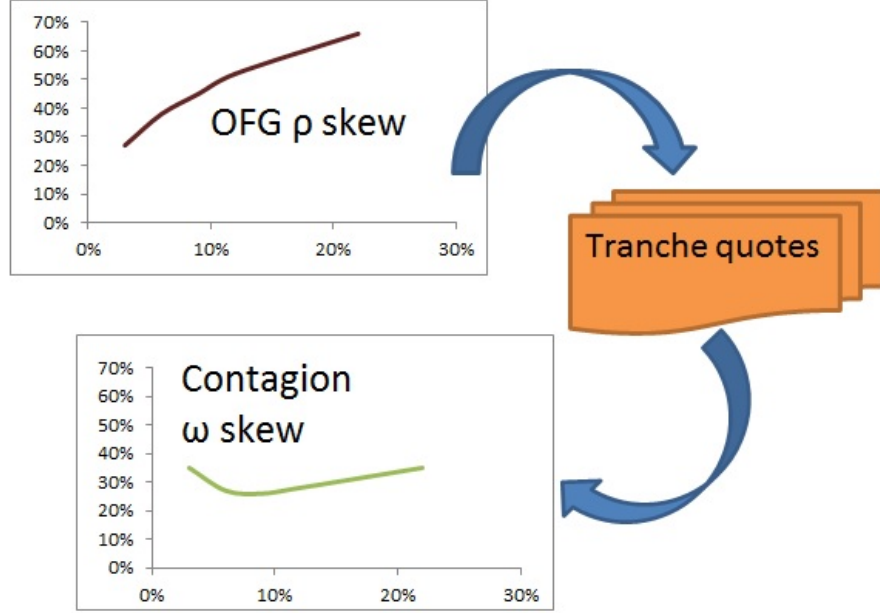


Figure 5.7: Strategy for the calibration exercise performed.

5.4.5 Single name deltas

Finally, we looked at tranche deltas with respect to single name quotes; this is a crucial information in order to assess the behavior of the model with respect to hedging. While the internal calculations are usually done in terms of single name default intensities, we preferred to show deltas to the actual CDS quotes from which the default intensities were calibrated as these will be the instruments used in real hedging strategies for structured credit products¹². Lets indicate with $q_i(T_j)$ the CDS quote for name i and maturity T_j ; what we are interested in is the following quantity

$$\Delta_i := \frac{\partial V_{CDO} [q_i(T_j)]}{\partial q_i(T_j)} \Big|_{\omega \text{ or } \rho} \quad (5.41)$$

where the derivative is taken by keeping fixed the dependence parameter (i.e. ω for the contagion model and ρ for the OFG one) and we stressed with the notation the dependency of V_{CDO} from the quantity we are studying ($q_i(T_j)$). We have actually

¹²In reality, structured credit dealers also use the full index for hedging purposes; this omission in our experiment does not impact the applicability of our results in any way.

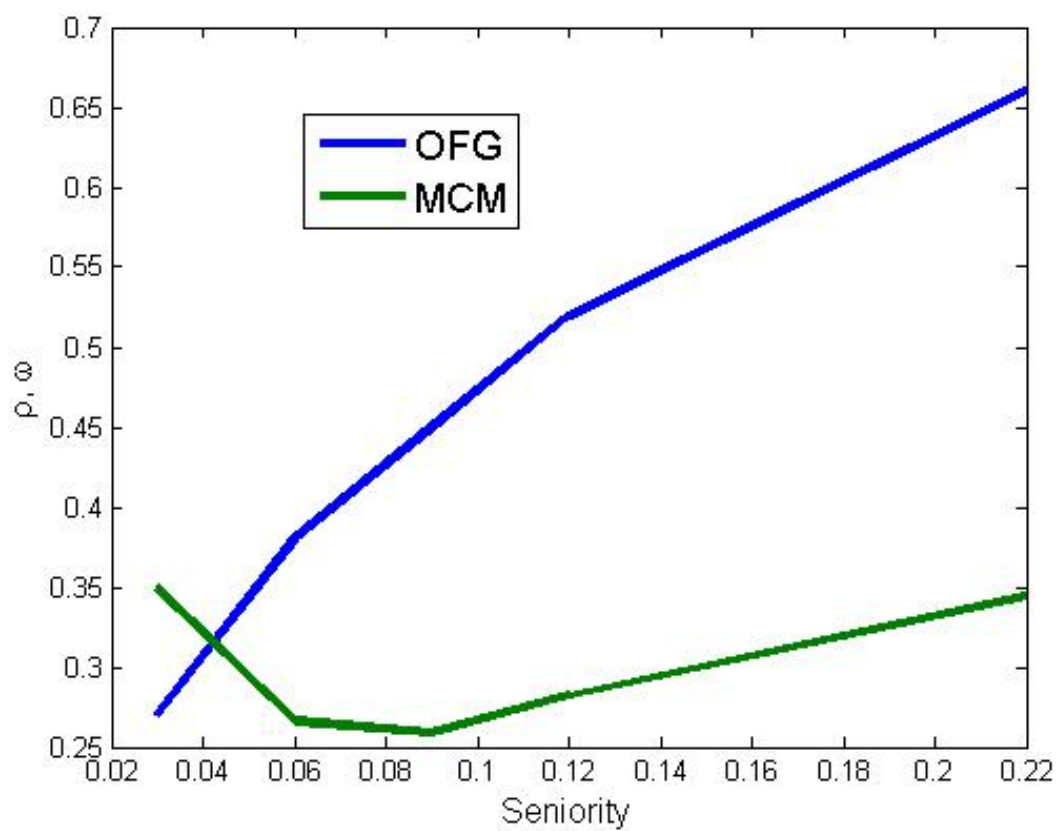


Figure 5.8: Smiles: parameter skew when both the OFG and the contagion models are calibrated to the same set of quotes.

calculated a numerical approximation of Δ_i ¹³:

$$\widetilde{\Delta}_i = \frac{V_{CDO} [q_i(T_j) + q_i(T_j) \cdot \delta] - V_{CDO} [q_i(T_j)]}{q_i(T_j) \cdot \delta} \quad (5.42)$$

A part from convexity issues, $\widetilde{\Delta}_i$ is usually a good proxy for Δ_i .

We analyzed deltas for the five tranches usually quoted on the market. We repeated the experiment for three values of the parameters and for both the OFG and the contagion model. Results are presented in figure 5.9. We can make some considerations by comparing the two sets of plots:

Magnitude The two models assign different amount of risk to different tranches. For example, the contagion model shows more sensitivity to single name changes for equity tranches than the OFG one. For most of the other tranches, the reverse is true. Smaller hedges (roughly 40% of their OFG equivalent) are suggested by the contagion model for mezzanine and senior tranches and bigger ones (25%) for equity ones. For example, while hedging a 0-3% tranche in the OFG model with $\rho = 50\%$ would require an hedge proportional¹⁴ to 105 for a name with a CDS spread of 100 basis points, the equivalent quantity for the contagion model with $\omega = 50\%$ is 190 (see the top left plots of each set). When instead we take a 6%-9% tranche (top right plots for each set) the situation is reversed as the required OFG hedge (68) is much bigger than the contagion one (14).

Steepness Lets look at the risk profile of the equity tranche for the OFG model. One can notice that the higher the correlation ρ , the steeper is the curve with respect to the riskiness of the underlying names. The equity tranche has always very little exposure to safe names (the ones with low CDS spreads): the reason is that it will probably be already wiped out by the time a default of a safe name occurs. The opposite is true for super senior tranches where risky names are assumed to default and most of the exposure comes from safe names (see top left and bottom right plots for the OFG model).

The same is not true in the contagion model: although it is still the case that

¹³In the results we presented we built the curves with only one quote (at 5Y maturity) and set the multiplicative bump value to $\delta = 5\%$.

¹⁴The numbers shown are numerical approximation of mathematical derivatives; the actual hedging notionals depend on the trade notional and the hedging instruments notionals.

safe names are less likely to default than riskier ones, their infection rate is much higher. This indirect effect is comparable to the idiosyncratic one and indeed the risk for an equity tranche is higher for safe names than for risky ones (see top left plot of the contagion set).

Convexity The effect highlighted on the previous point with regard to the steepness of the equity risk profile applies to the other tranches too. It makes their risk profiles rather flat while the OFG model exhibits very tranche-specific behavior. In general, the risk profile in the contagion model has very little convexity, i.e. every name is equally likely to impact the tranche, making static hedging strategies more effective in the contagion setting.

Dependency on the parameter The shape of the risk profiles for the OFG model is highly sensitive to the value of ρ . For example, consider the 9%-12% tranche; its risk goes from an almost straight line for low-medium values of ρ to a bell shape for high correlations. The equivalent change in shape for the contagion model is much less pronounced; moving ω usually affects the risk profile overall level but not its shape¹⁵. This is due to the way the ω parameter is applied uniformly across the spectrum of names: an increase in its value is translated in an almost parallel shift in the risk numbers.

5.5 Conclusive remarks

In this chapter we presented a new model that uses Davis and Lo (2001) as a starting point. Unlike other extensions of such model, the one introduced here can achieve reasonable performances with heterogeneous portfolios as we provided both theoretical and practical results for the efficient computation of the portfolio loss distribution.

We then introduced a restricted version of the model and applied it to the problem of pricing CDO products. We showed that the range of possible prices obtainable by the contagion model is comparable (if not wider) to the one obtained with the one factor Gaussian model. We also performed a calibration exercise and showed that the skew in the ω space is less pronounced than the one in the equivalent Gaussian ρ space, resulting in a more robust modeling framework. Finally, we studied the tranche risks with

¹⁵An exception is the super senior tranche 12%-22% for which a change in the sign of the slope can be observed when increasing ω from low to medium values.

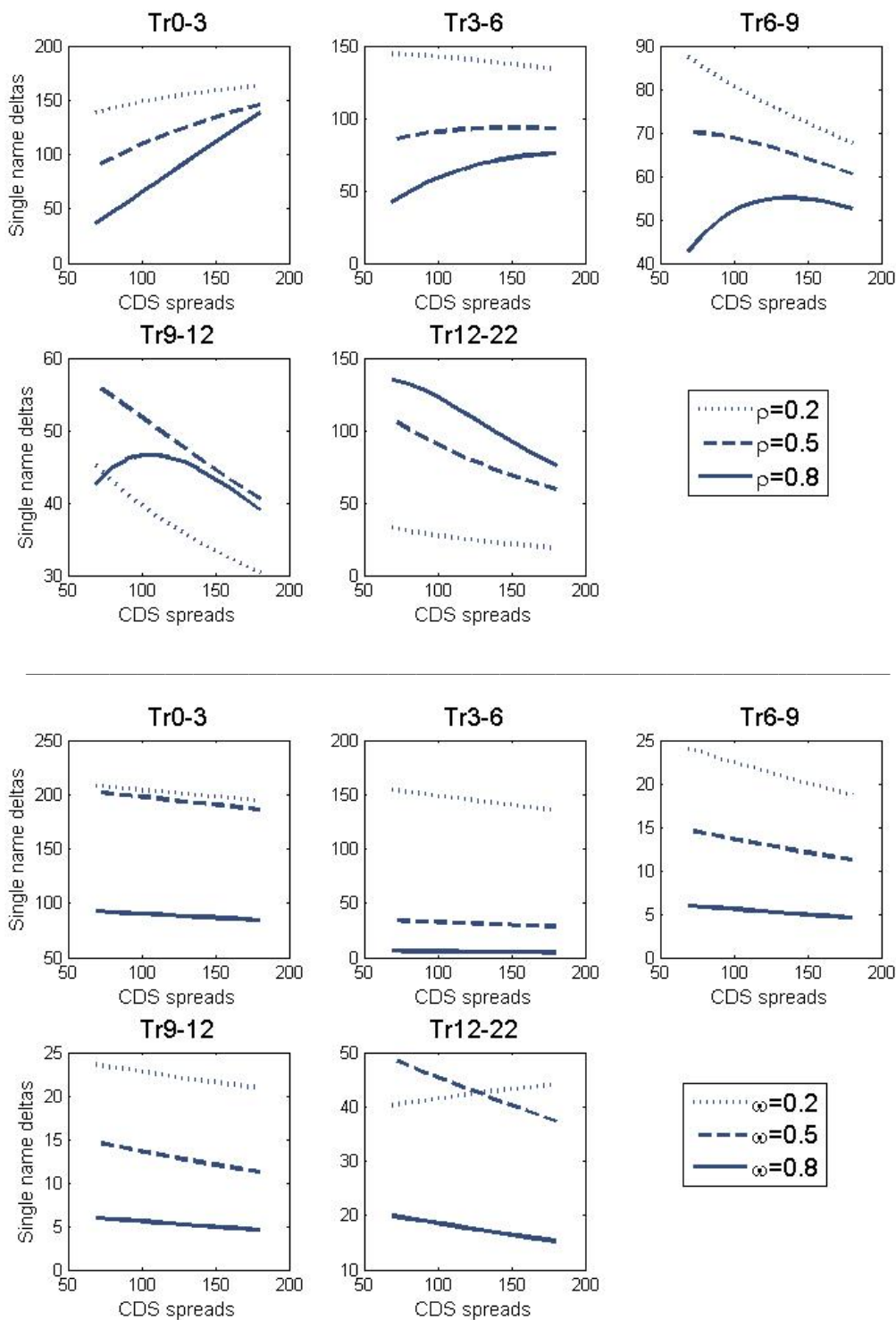


Figure 5.9: Single name deltas: $\tilde{\Delta}$ for the OFG model (above) and the contagion model (below). Each box shows risk to single names of a particular tranche (with notional set to 1 million and coupon set to 1%). The x axis shows the CDS quotes for the individual names in the portfolio: risky names are on the right, safe ones on the left.

respect to the credit worthiness of the underlying names, concluding that the contagion model shows less convexity than the one factor Gaussian one.

Chapter 6

A new measure of systemic risk in contagion models

In this chapter we suggest a new measure, called Contagion Loss Ratio, for assessing systemic risk in contagion models. It is computed as the percentage of the portfolio expected loss that is caused by infection. In order to demonstrate the applicability of the new measure, we show how to calculate it for three contagion models: Davis and Lo (2001), Sakata, Hisakado, and Mori (2007) and the model we introduced in chapter 5. Finally an application to an artificial banking system is presented.

6.1 Introduction

In this chapter we introduce a new systemic risk measure. In chapter 2 we adopted a classification of the existing literature on the subject based on the modeling frameworks used; under that point of view, the measure we introduce belongs to the branch of literature that considers financial systems as portfolios of assets as our work is based on the class of contagion models pioneered by Davis and Lo (2001). There have been many attempts to introduce asymmetric dependency structure in credit risk modeling with the double aim of increasing the probability of large losses (fat tail problem) and explaining default clustering¹. We would like to cite few interesting readings with regard to adding contagion mechanism to the standard formalism of credit modeling. Jarrow and Yu (2001), Rösch and Winterfeldt (2008), Neu and Kühn (2004), Egloff, Leippold, and Vanini (2007), Yu (2007) (that generalizes Schönbucher and Schubert (2001)), Giesecke and Weber (2004) and Frey and Backhaus (2010) are but a few of the principal works on this topic.

In the next section we will introduce the new indicator. In section 3 we will show how to calculate it for the models presented in Davis and Lo (2001), Sakata, Hisakado, and Mori (2007) and for the model introduced in chapter 5. Section 4 presents an application to a fictional banking system in order to show the behavior of the new measure with respect to changes in the model parameters. Section 5 concludes.

6.2 Contagion Losses Ratio measures

6.2.1 A thought experiment

Before formally introducing the new risk measure, let's make a thought experiment. Suppose, we have three systems each composed of four entities. In the first system, all names are independent, have a default probability of 50% each and a LGD of 1\$. The second system is identical to the first with the difference that the first three names have a probability of default of just 1% while the last has an LGD of 4\$ and a probability of default of 49.25%. In the third system every name has a LGD = 1\$ (like in the first system) but it has a different dependency structure: there are three firms that will

¹See the work from Azizpour, Giesecke, et al. for evidence about default clustering and the role of contagion.

never default on their own but they are linked to the survival of the last name that has a 50% probability of default (figure 6.1 shows this dependency structure). Table 6.1 summarizes the three systems

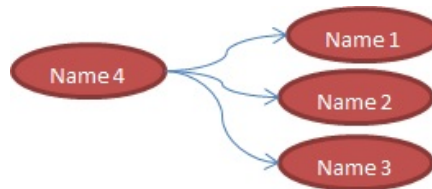


Figure 6.1: Dependency structure of system C.

	System A		System B		System C	
	IPD	LGD	IPD	LGD	IPD	LGD
First 3 names	50%	1\$	1%	1\$	0%	1\$
Last name	50%	1\$	49.25%	4\$	50%	1\$
Dep structure	Independent		Independent		Survival hierarchy	
Expected loss	2\$		2\$		2\$	

Table 6.1: Summary of the three systems of the thought experiment. IPD represents the idiosyncratic probability of default.

All three portfolios have an expected loss of 2\$ but they reach it in different ways; figure 6.2 shows the loss distribution for the three systems. What can we say about systemic risk of the three systems? The first observation might be obvious but it is worth mentioning it here anyway: the portfolio total expected loss gives information on how risky a portfolio is but says very little about how much systemic risk it carries. System A is probably the best we can hope for in terms of systemic risk: every default adds the minimum amount of losses possible and they are all independent from each other and yet shows the same expected loss as the other two systems. System B and C, on the other hand, achieve their expected loss via a single catastrophic event. Although the loss distributions of system B and C are almost identical, there is a subtle difference on how the catastrophic event occurs. In both systems the total expected loss is basically caused by a single default but while in the first this is the default of a huge (in relative terms) firm whose LGD is four times the size of any other name, in system C this default has nothing spectacular in it, if considered on its own; it is the chain of

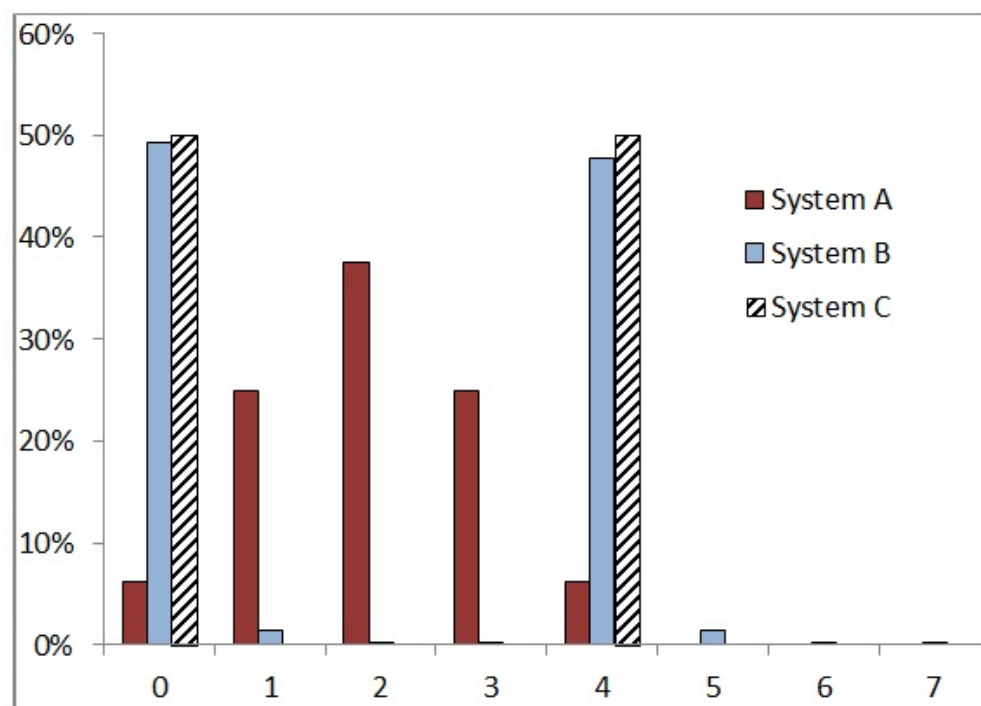


Figure 6.2: Loss distribution for the three systems.

defaults that it triggers that cause the damage. The final result might be the same, but we consider system C much more problematic to deal with from the regulators point of view; a quick look at the relative sizes in system B would raise warnings (the size of the last name is bigger than the sum of the rest of the system) while a detailed knowledge of the dependency structure of system C is needed in order to uncover its risks.

It is very rare to have the full information about the interactions of every firm for the complicated systems we deal with in real life, especially for today's highly interconnected financial systems. What we do have, often, is a way to model the system loss distribution but that information alone does not resolve the issue either.

6.2.2 CLR and CoCLR

What truly differentiates the two examples B and C above is the role played by losses resulting from contagion events: in system C they constitute the bulk of the total expected losses, while they are completely absent from system B. We can exploit the ability of most contagion models to attribute losses to either infective events or standard idiosyncratic defaults to measure the degree of systemic risk of the portfolio. The indicator we propose is then the contagion losses ratio, CLR, and it is the ratio of infection-driven losses versus the total portfolio expected loss. The higher the CLR, the more likely is that the system will experience losses due to contagion events².

CLR measures the contribution of contagion losses for the entire loss distribution. Sometimes it is more useful to focus on the tails of such distribution, in particular the right one, i.e. high losses scenarios. We can tailor the CLR to look specifically in that area. This naturally leads to a new version of the indicator, the conditional CLR or CoCLR. Its definition is given by

$$CoCLR = \frac{E[L_n^C | L_n \geq T]}{E[L_n | L_n \geq T]}$$

where L_n^C represent losses due to contagion and L_n the total losses, both for a system of n names. Therefore, CoCLR is computed as the ratio of two conditional expectations.

²Just to complete the thought experiment, the CLR for systems A and B is 0 while the same indicator for system C scores a staggering 75% (the default of the last name does not count among the contagion losses because it is an idiosyncratic one).

The CoCLR is a generalization of the CLR as $CLR = CoCLR$ when $T = 0$.

6.3 Calculating CLR and CoCLR in practice

The risk measure we propose is model dependent as it requires the ability to attribute losses to contagion or idiosyncratic effects. This is normally a reasonable task for such models; in the next subsections we will show how to calculate the CLR measure for three contagion models (Davis and Lo (2001), Sakata, Hisakado, and Mori (2007) and the one introduced in chapter 5). The base idea is to get a formula for the probability of having h units of idiosyncratic losses and k units of contagion-driven ones. We will indicate with $\gamma_n(h, k)$ such probabilities for a portfolio of n names and with L_n the total losses of the portfolio. In the three models considered, there are only two ways for defaulting: either idiosyncratically or via contagion. We can therefore write

$$P\{L_n = z\} = \sum_h \gamma_n(h, z - h) \quad (6.1)$$

as we are sure to cover all possible combinations that can cause z units of losses. The CLR can then be written as

$$CLR = \frac{\sum_k [k \cdot \sum_h \gamma_n(h, k)]}{\sum_z z \cdot P\{L_n = z\}} = \frac{\sum_k [k \cdot \sum_h \gamma_n(h, k)]}{\sum_z z \cdot [\sum_h \gamma_n(h, z - h)]} \quad (6.2)$$

Thanks to equation (6.2), for every model we will only need to show how to calculate $\gamma_n(h, k)$ in order to be able to calculate the CLR.

We can also get an expression for CoCLR in terms of $\gamma_n(\cdot, \cdot)$; consider that

$$E[L_n | L_n \geq T] = \sum_z z \cdot P\{L_n = z | L_n \geq T\} = \frac{\sum_{z \geq T} z \cdot P\{L_n = z\}}{P\{L_n \geq T\}}$$

hence

$$E[L_n | L_n \geq T] = \frac{\sum_{z \geq T} z \cdot \sum_h \gamma_n(h, z - h)}{P\{L_n \geq T\}} \quad (6.3)$$

On the other hand

$$E[L_n^C | L_n \geq T] = \sum_k k \cdot P\{L_n^C = k | L_n \geq T\} = \frac{\sum_k k \cdot P\{L_n^C = k, L_n \geq T\}}{P\{L_n \geq T\}}$$

from which

$$E[L_n^C | L_n \geq T] = \frac{\sum_k k \cdot \sum_{h \geq T-k} \gamma_n(h, k)}{P\{L_n \geq T\}} \quad (6.4)$$

Taking the ratio of equations (6.4) and (6.3), we get the desired expression

$$CoCLR = \frac{\sum_k [k \cdot \sum_{h \geq T-k} \gamma_n(h, k)]}{\sum_{z \geq T} z \cdot [\sum_h \gamma_n(h, z-h)]} \quad (6.5)$$

6.3.1 Davis and Lo

The main result presented in Davis and Lo (2001) is the formula for $P\{L_n = h\}$ when the portfolio is homogeneous (also in terms of LGDs):

$$P[L_n = z] = C_z^n \cdot \sum_{i=1}^z C_i^h p^i (1-p)^{n-i} [1 - (1-q)^i]^{h-i} (1-q)^{i(n-z)} \quad (6.6)$$

where $C_j^i = i!/j!(i-j)!$ and p and q are the model's parameters. In particular:

- p is the idiosyncratic probability of default of each name.
- q is the pairwise infection probability, i.e. the probability that a defaulting name infects another firm.

Every term in the summation on the right hand side of equation (6.6) covers the scenario of having i idiosyncratic defaults and $z-i$ infection-driven ones. It is easy then to see that $\gamma_n(h, k)$ is given by

$$\gamma_n(h, k) = C_{h+k}^n \cdot C_h^{h+k} \cdot p^h (1-p)^{n-h} [1 - (1-q)^h]^k (1-q)^{h[n-(h+k)]}$$

6.3.2 Sakata, Hisakado and Mori

The model by Sakata, Hisakado, and Mori (2007) is an extension of the one by Davis and Lo that allows for positive recovery spillages to flow from healthy firms to distressed ones. The usual contagion effect from defaulting names to healthy ones is also present. Their main result is a closed form formula for $P\{L_n = h\}$ and, again, is obtained for

the homogenous case:

$$P[L_n = z] = C_z^n \cdot \sum_{i=1}^z \sum_{m=0}^{n-z} \left\{ \begin{array}{l} C_i^z C_m^{n-z} p^{n-z-m+i} (1-p)^{z-i+m} \cdot \\ (1-q')^{i(z+m-i)} (1-q)^{m(n-z-m+i)} \cdot \\ [1 - (1-q)^{n-z-m+i}]^{z-i} \cdot \\ [1 - (1-q')^{z+m-i}]^{n-(z+m)} \end{array} \right\} \quad (6.7)$$

where

- p is the idiosyncratic probability of default of each name.
- q is the pairwise infection probability, i.e. the probability that a defaulting name infects another firm.
- q' is the pairwise recovery probability, i.e. the probability that an healthy name recovers a distressed firm.

The terms in the double summation of equation (6.7) represent the scenario where z names default (of which i idiosyncratically and $z-i$ because of contagion) and $n-k$ survive (of which m on their own and the rest thanks to the support of other names). With a little bit of rearrangements, we can get

$$\gamma_n(h, k) = \sum_m \left\{ \begin{array}{l} C_{h+k}^m C_h^{h+k} C_m^{n-(h+k)} p^{n-k-m} \cdot \\ (1-p)^{k+m} (1-q)^{m(n-k-m)} \cdot \\ (1-q')^{h(k+m)} [1 - (1-q)^{n-k-m}]^k \cdot \\ [1 - (1-q')^{k+m}]^{n-(h+k+m)} \end{array} \right\} \quad (6.8)$$

Figure 6.3 shows the expected loss and three systemic risk indicators for a small homogenous portfolio with 15 names modeled with Davis and Lo model. Similarly, figures 6.4 and 6.5 show the same information for the Sakata, Hisakado and Mori approach for two levels of the parameter q' . Together with the total portfolio losses EL and the contagion losses ratio CLR, we also added two other measures; the first is a loss based measure called the distress insurance premium DIP³ introduced by Huang, Zhou,

³The DIP is calculated as

$$E[L_n | L_n \geq L_T]$$

where L_T is a threshold level that should represent the event of the system being in a distressed state; for this particular example, we set $L_T = 11$, i.e. system loosing around a third of its potential total loss.

and Zhu (2012). The second is a default counting one, the banking stability index BSI⁴ by Segoviano and Goodhart (2009).

Every surface shows the change in one measure when p (on the bottom left axis) and q (on the bottom right axis) vary while q' is fixed at 10% (6.4) or 20% (6.5) for the Sakata, Hisakado and Mori model. For both models, the change in the EL is significant in both the p and q directions suggesting an equal increase in the portfolio riskiness. Both the DIP and BSI measures follow a similar pattern, with the DIP giving slightly more importance to changes in q than to changes in p ⁵. The CLR instead shows the biggest variation along the q direction and is instead decreasing with p : this is intuitive as increasing p will cause increases in idiosyncratic losses at a faster pace than the increase in the contagion ones (each idiosyncratic loss carries the ability to cause infection-driven defaults). So the ratio between contagion and total losses decreases with p . A system where more losses come from idiosyncratic events has less systemic risk than one where contagion losses prevail, assuming the total portfolio loss is the same. Both the BSI and the DIP fail to make this distinction and rely on portfolio riskiness as the main signal of systemic risk.

Finally, we can notice how the addition of the recovery spillage mechanism in the Sakata, Hisakado and Mori approach lower the EL but does not change the qualitative behaviors of the three measures.

6.3.3 Contagion model introduced in chapter 5

Another example of calculation of CLR can be seen with the model we introduced in chapter 5. We can write $\gamma_n(h, k)$ in terms of $\alpha(\cdot, \cdot)$ and $\beta(\cdot, \cdot)$ defined previously (equation (5.27) on page 106) as

$$\gamma_n(h, k) = \beta_n(h, k) + \delta_{k,0} \sum_i \alpha_n(h, i)$$

⁴The formula to calculate the BSI is given by

$$BSI = \frac{\sum_{i=1}^n P\{X_i \text{ defaults}\}}{P\{\text{At least one default}\}}$$

Remember that the BSI is based on default counting and not on actual losses.

⁵We can note how the slope along the q direction is bigger, in absolute terms, than the slope along the p one:

$$\left| \frac{\partial DIP}{\partial q} \right| \geq \left| \frac{\partial DIP}{\partial p} \right|$$

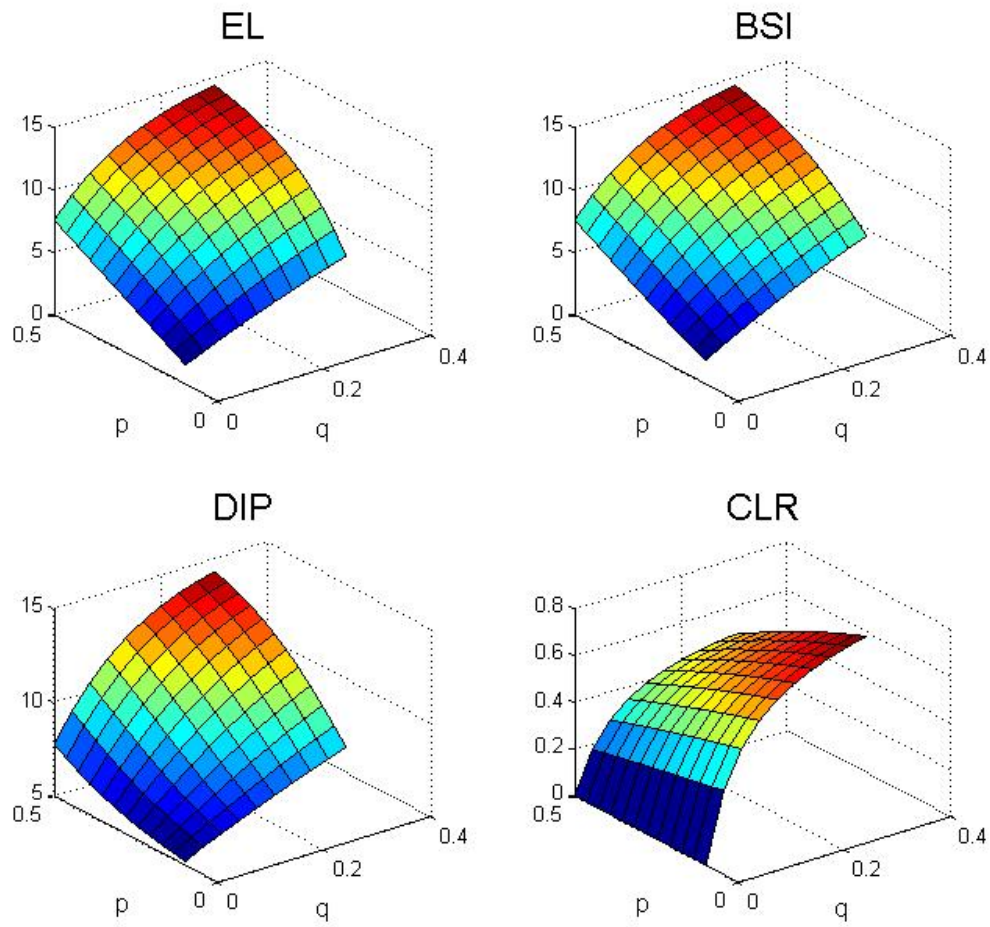


Figure 6.3: Changes in the EL (upper left), BSI (upper right), DIP (bottom left) and CLR (bottom right) when p and q vary for a portfolio with 15 names.

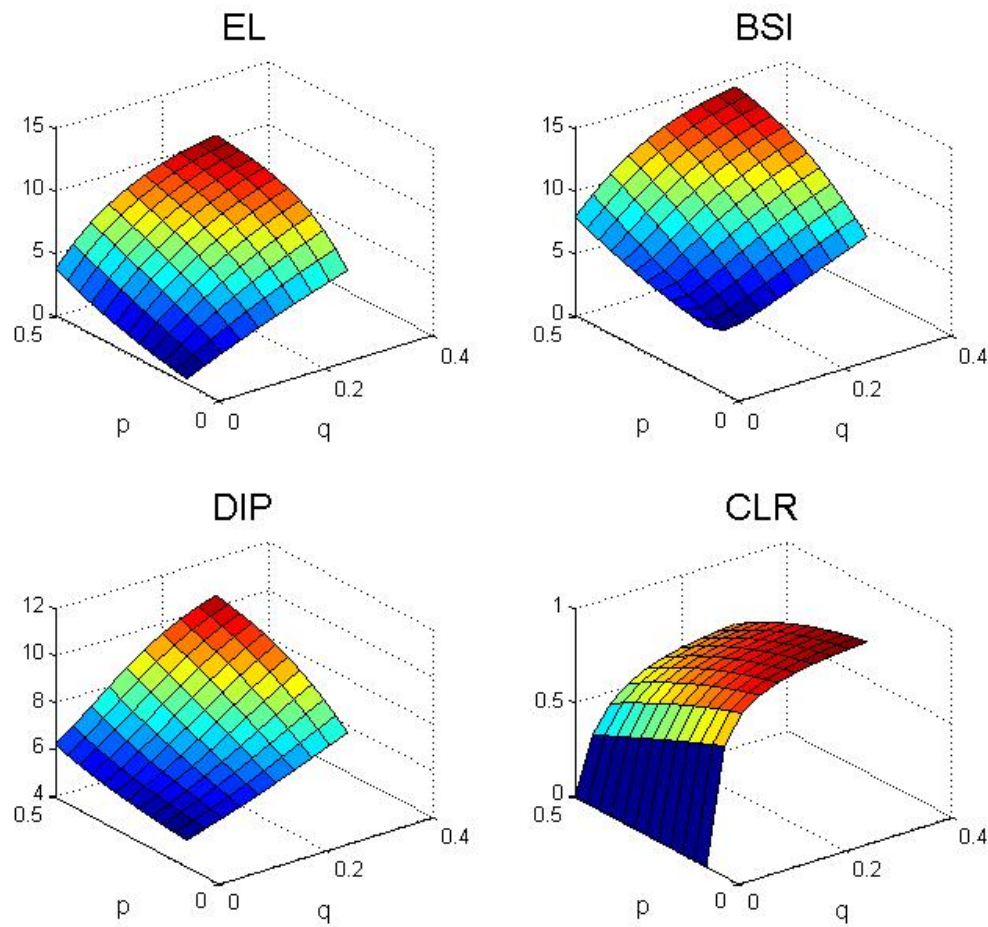


Figure 6.4: Changes in the EL (upper left), BSI (upper right), DIP (bottom left) and CLR (bottom right) when p and q vary for a portfolio with 15 names with $q' = 10\%$.

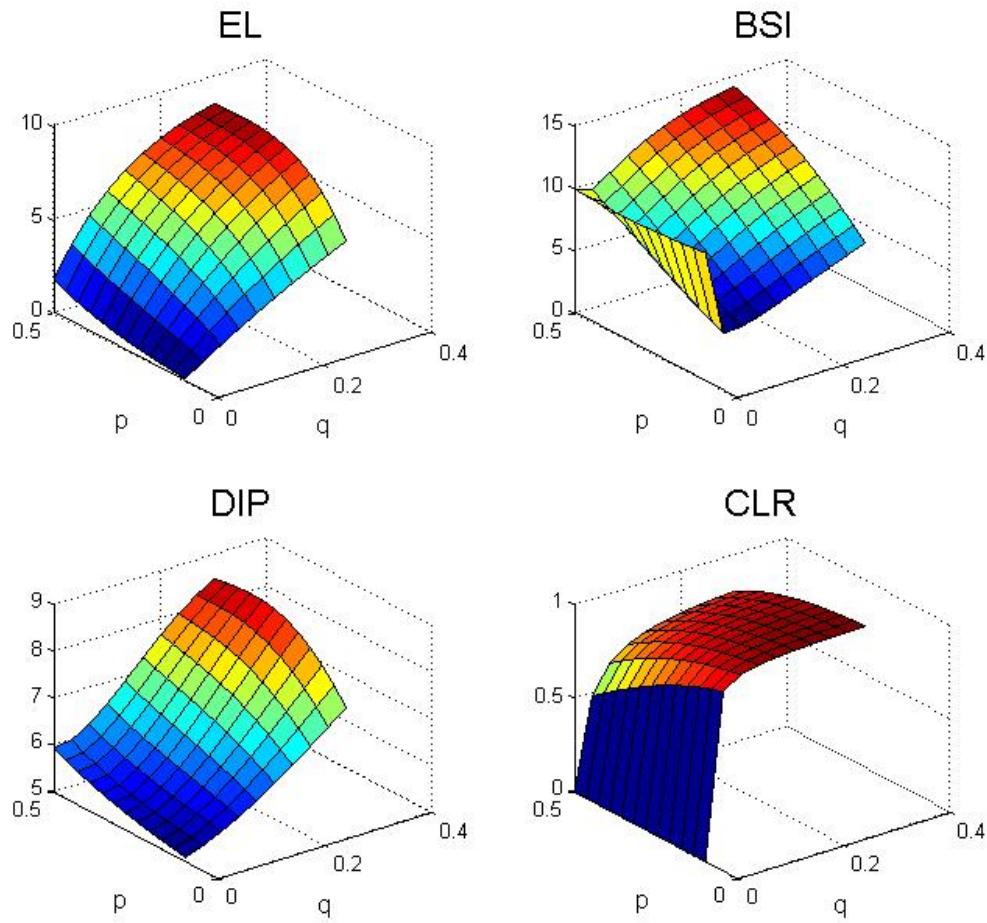


Figure 6.5: Changes in the EL (upper left), BSI (upper right), DIP (bottom left) and CLR (bottom right) when p and q vary for a portfolio with 15 names with $q' = 20\%$.

where $\delta_{i,j}$ is the Kronecker's symbol whose value is 0 if $i \neq j$ or 1 otherwise.

In the next section, we will exploit this modeling context's ability to handle heterogeneous portfolios in order to test the CLR in a more complex situation.

6.4 A fictional banking system

In this section we will build a fictional banking system and use it to compare the CLR against two other measures, the DIP and the BSI, encountered earlier. The modeling framework we will use in this exercise is the one introduced in chapter 5 where five parameters are needed for each name i in the portfolio modeled:

- The parameter p_i governs the idiosyncratic probability of default.
- The parameter u_i describes the ability of name i to defend itself against infections.
- The parameter v_i controls the probability that i , upon defaults, attempts to infect the rest of the system.
- The parameters d_i and c_i model the LGD for name i in idiosyncratic and contagion-driven default event respectively.

6.4.1 System specification

The universe we will consider in this application is designed to describe an interconnected network of firms representing the national banking system of a small country. We will model four different types of banks:

- The first type of entities consists of very safe banks, whose main activity is old-fashioned lending. Their size is big (and therefore their losses upon default will also be significant) but the probability of idiosyncratic default is small. Since they provide most of the liquidity of this fictional system, the probability of contagion they carry is quite high.
- The second group is represented by very aggressive, investment bank style institutes. Their size is considerable but smaller than those in the first group and their business is quite risky. The market partially anticipates their default and hence the infection rate they show is small.

- The third group represents smaller, regional banks. Their main feature is to depend on the well being of the system. They don't have the strength to sustain any upset in the financial system they rely upon, hence their defense parameter u is zero.
- Finally, we included a central bank. Its default should be read as the possibility of extremely negative news for the entire system but causes no direct losses (the LGD is set to 0). We exploited the flexibility of the modeling environment that allows to add entities solely as carriers of shocks to the system but that add no direct financial losses to the portfolio.

We will consider only one central bank but three entities for each of the remaining groups, for a total of 10 banks. Table 6.2 provides details on the parametrization chosen.

Type (num)	p	u	v	d	c
Safe (3)	0.1	0.8	0.9	5	6
Risky (3)	0.7	0.4	0.1	3	4
Local (3)	0.2	0	0.01	1	2
Central (1)	0.01	1	1	0	0

Table 6.2: An artificial banking system. Both d and c represent units of losses.

6.4.2 Systemic risk measures in the base case

Before analyzing systemic risk issues, let's take a look at the base configuration. Table 6.3 reports the probability of default of every type of entity, split into idiosyncratic and contagion components. The last column shows the ratio of contagion versus total probability of default. We can notice how regional banks are the most affected by contagion: they more than double their default probability due to infection threats. Risky banks instead have already a significant default probability and the increase caused by contagion is marginal. Finally, safe banks show a medium increase in their riskiness when contagion is considered.

Table 6.4, instead, provides information on the various indicators in the current setting as well as the case where no contagion is allowed. The BSI for the starting setting is 4.40, meaning that we expect to see 4.40 defaults given that at least one default has

Type	Total	Idiosyncratic	Contagion	Contagion/Total %
Safe	16.20	10.00	6.20	38.28
Risky	76.46	70.00	6.46	8.44
Local	52.19	20.00	32.19	61.68
Central	1.00	1.00	0	0

Table 6.3: Individual default probabilities: total, idiosyncratic and contagion component and relative importance of the latter over the total. All numbers shown are percentages.

Indicator	No contagion	With contagion	Percentage increase
EL	8.29	12.22	9.58%
BSI	3.34	4.40	6.85%
DIP	12.26	18.25	18.76%
CLR	0	0.31	50%

Table 6.4: Systemic risk indicators for the banking system (contagion on/off).

If we indicate with c the contagion value and with a the one obtained without contagion, the relative change is calculated as $\frac{c-a}{0.5 \cdot (c+a)}$ to deal with the zero CLR case.

occurred. In absence of contagion, we would instead expect 3.04 defaults conditioned on the same event. The increase is moderate.

The DIP scores 18.25 for the initial setting (12.26 if no contagion was allowed), representing the expected loss of the portfolio, once it has exceeded 10 units of losses. The increase due to contagion in this case is more pronounced, causing an addition 6 units of losses to be added to the tally.

Finally, the CLR is 31% (it would be zero without infection), i.e. 31% of the expected loss can be attributed to contagion events. The measure is designed to look for infection-driven losses so it is natural that the change in this case is maximum.

Both the CLR and the CoCLR are based on the ability to disentangle from the portfolio loss distribution the information regarding the causes of such losses, i.e. contagion or idiosyncratic. We can achieve that in virtue of equation (6.1)

$$P\{L_n = z\} = \sum_h \gamma_n(h, z - h)$$

that splits $P\{L_n = z\}$ into the individual $\gamma_n(\cdot, \cdot)$ components. Figure 6.6 shows both the loss distribution (top plot) and the set of the $\gamma_n(h, k)$ (bottom plot) on a 2-dimensional surface. The loss distribution in the top plot can be recovered by summing diagonally across the γ_n plot. Regarding the γ_n plot, it is interesting to note how we can split the default region in two areas: the first shows three peaks along the no-contagion direction ($k = 0$ on the right axis) corresponding to the three possible configurations of defaults of the old-style big banks (only one default, two defaults or three defaults). The other area is more complex: it is located in the middle of the surface and contains all the possible combinations of idiosyncratic (but infective) defaults and losses due to contagion. Very little is happening outside these two areas.

6.4.3 Conditioned CLR

Next we present in figure 6.7 an example of the behavior of the CoCLR calculated on the fictional banking system. We also added the DIP measure for comparison⁶. The horizontal axis follows the tail level T while the dark blue line shows the value of the corresponding CoCLR (left vertical axis) where the light green line shows the

⁶The DIP measure also requires to specify a tail parameter T to determine when the system is considered to be in distress.

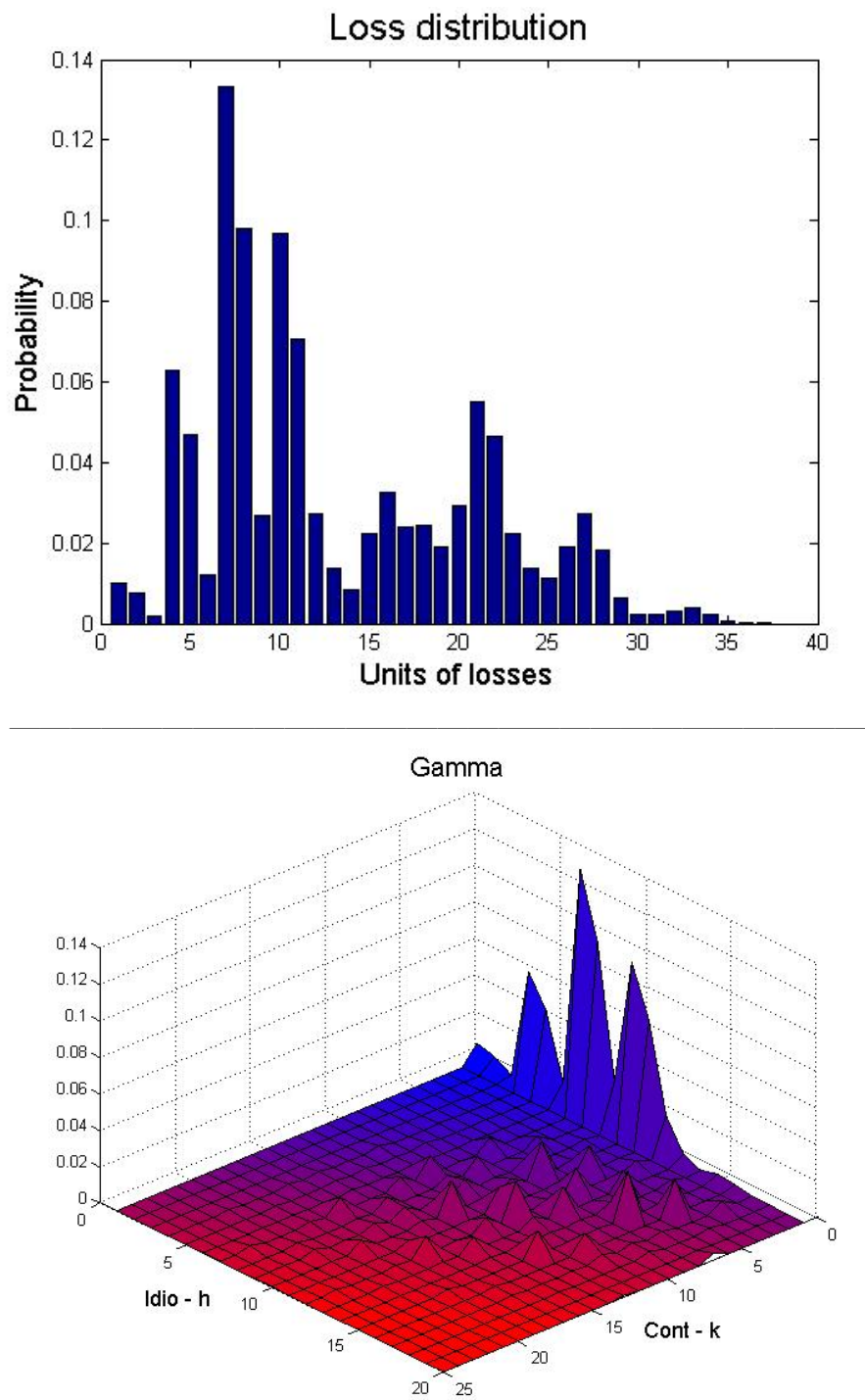


Figure 6.6: Top plot: loss distribution. Bottom plot: distribution of $\gamma_n(h, k)$. Each point represents the probability of having h units of losses from idiosyncratic effects and k units coming from contagion. The region plotted has been cut to cover the most significant losses (more than 95% of probability mass included). The color goes from blue for small total losses to red for high losses.

DIP values (right vertical axis). Please note that the two measures are not directly comparable in terms of units (the CoCLR is expressed in percentages while the DIP in conditional expected losses) but the juxtaposition is useful in terms of the relative slope of the two curves. They are both generally increasing with T but we can notice how the CoCLR shows some drops; they can be explained as zones in the loss distribution where losses are due mainly to idiosyncratic effects rather than to contagion ones. In this case, the denominator of formula (6.5) grows more rapidly than the numerator causing the CoCLR measure to decrease. We can also note how the CoCLR measure goes to 1 signaling that the most extreme loss scenarios (i.e. central bank defaulting) can only be reached via contagion⁷.

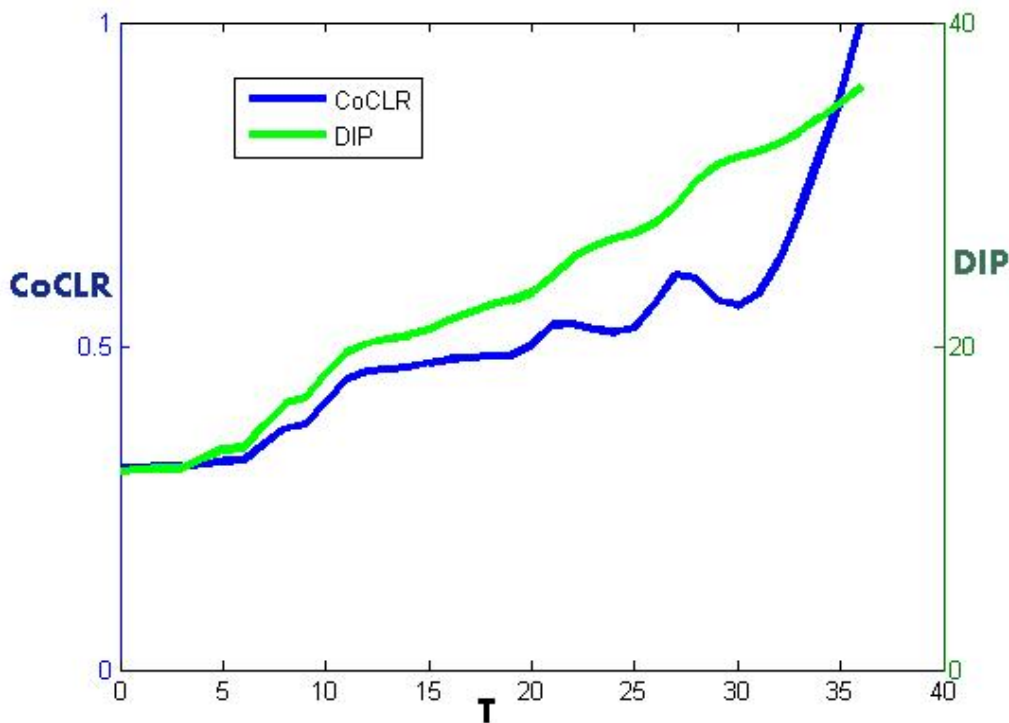


Figure 6.7: CoCLR and DIP vs. T for the banking system.

⁷Remember that the LGD for the central bank was set to zero; in a scenario where the central bank defaults, all the subsequent losses would come by contagion.

6.4.4 Sensitivity results

Finally we performed a sensitivity analysis of the three measures considered so far by moving the parameters p , u and v for one name of each of the four types. Results are presented in figure 6.8: it shows changes to the EL, BSI, DIP and CLR when moving the model parameters. In each box there are four lines, one for each bank type.

A part from being useful in order to show model/measure behaviors, these plots

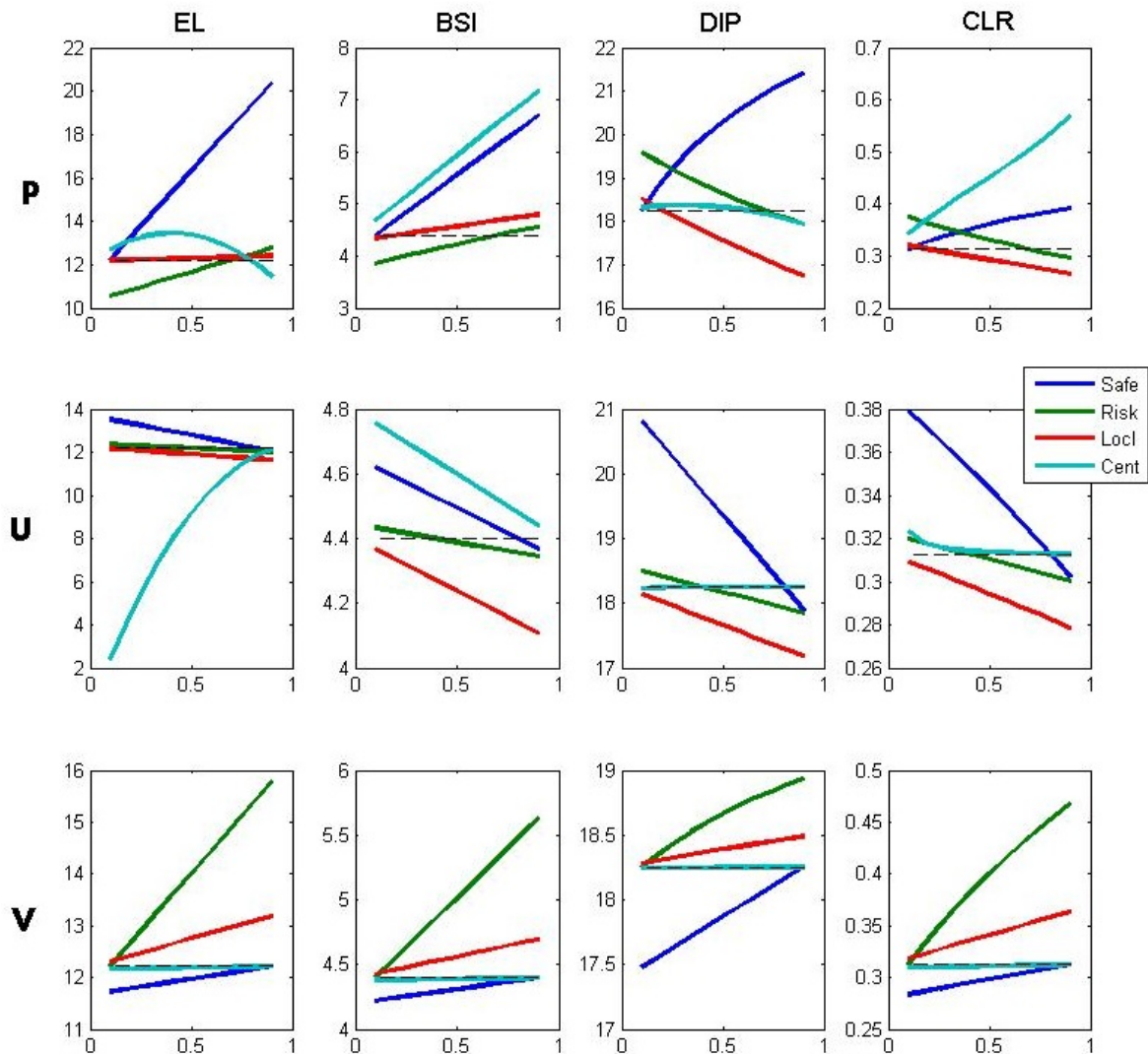


Figure 6.8: Change in the risk measures when varying the parameters p (first row), u (second row) and v (last row) for one name of each category.

can also serve as guidance for regulators' interventions. What would an hypothetical regulator do in order to reduce the systemic risk of this fictional banking system?

This provocative question has many possible answers, each of which based on many assumptions. Firstly we need to assume that he/she will trust the representation of reality as depicted by the model used. Secondly, he/she needs to decide what systemic risk is (an elusive task) and which measure he/she will rely upon for his/her judgment. Depending on the answer to the last two questions, several courses of action are possible.

The least controversial one would be to move the system toward independence by reducing the level of negative inter dependencies between the banks. Under the current model assumptions, this can be achieved in two ways:

Reducing v (third row) This effect is present for all the systemic risk measures considered. Decreasing v for any bank will ultimately reduce the overall level of losses due to contagion; the effect is more marked for those type of banks (risky and local) that are more likely to default idiosyncratically and therefore more likely to be the source of contagion in the system.

Increasing u (second row) This effect is also present for all the systemic risk measures considered. Again, the ultimate consequence is to reduce the amount of losses due to contagion but this time the impact is maximal for safe banks; the reason is that the LGD in case of default due to contagion is maximum for this type of banks and a change in u will therefore have a bigger impact on the overall risk measure considered.

Weakening the channels through which infection can spread is surely a safe, non controversial way of reducing systemic risk but it is also a very difficult route to take. The number of possible connections a potential regulator has to monitor is huge (2^n); in addition, these links in real life take an extremely wide set of forms, going from direct economical connections to informational/perception driven ones.

It is no surprise then that the most common approach regulators have taken in the past is to work directly on the single banks, ultimately affecting their default probability, for example by imposing capital ratios. In our modeling framework, this translates in studying the impact of changes of p (top row). It could be intuitive to assume that by decreasing the probability of idiosyncratic default of any participant of the system, the overall level of systemic risk should decrease. By looking at the BSI plot vs. p (second plot from the left) this is indeed the case, as the measure simply counts the conditioned probability of default.

Looking at the other two measures, instead, reveals a more complicated pattern; if we focus on the DIP (third plot on the first row) we can see that we have banks whose increased probability of default actually benefits the entire system (risky and local banks). By increasing p for bank A , two effects happen:

- the probability of defaults caused by the bank A increase, overall increasing the amount of losses of the system
- the probability of bank A to default by contagion decreases. As the LGD in case of default by contagion is higher than in the idiosyncratic case, the overall effect is to decrease the amount of losses of the system

The two above effects pull in different directions; for extremely infective banks (like safe ones) the first one dominates and the overall system is in worse shape if bank A probability of idiosyncratic default increases. On the other hand, banks that are not very infective show an opposite behavior. The BSI is a measure too simple (it just counts defaults) to capture these type of dynamics as it fails to attribute the appropriate importance to different defaults like the DIP does.

The CLR is also sophisticated enough to show patterns similar to the ones exhibited by the DIP. The main difference between the two measures can be seen by looking at the central bank⁸. By increasing p for the central bank we are actually increasing the probability of a doomsday scenario; the DIP measure is a conditional measure and can be calculated as a ratio: in this case both the numerator and the denominator would increase similarly making the DIP almost indifferent to this parameter.

The situation for the CLR is very different (fourth plot, top row): the central bank is by far the most infective actor in the system and the expected losses due to contagion are extremely sensitive to its parameters.

Finally, we can note how the magnitude of the changes due to u and v is smaller than the one due to p meaning that the overall impact of moving p is greater than the equivalent for u and v . Table 6.5 provides additional info regarding the ranges obtainable when varying the model parameters on each measure⁹.

⁸Remember that we have $LGD = 0$ for this type of bank as it models the probability of having extremely negative news for the system.

⁹The percentage changes are calculated as $\frac{Max-Min}{Base\ value}$.

Measure	Type	P			U			V		
		Min	Max	% Δ	Min	Max	% Δ	Min	Max	% Δ
BSI	Safe	4.40	6.70	52.38	4.37	4.62	5.70	4.22	4.40	4.07
	Risky	3.87	4.57	15.93	4.34	4.43	1.98	4.40	5.63	27.92
	Local	4.34	4.79	10.20	4.11	4.37	5.91	4.43	4.70	6.10
	Central	4.68	7.17	56.60	4.44	4.76	7.23	4.38	4.40	0.40
DIP	Safe	18.25	21.39	17.19	17.89	20.81	15.98	17.48	18.25	4.24
	Risky	17.94	19.58	9.01	17.85	18.50	3.54	18.25	18.93	3.71
	Local	16.74	18.51	9.71	17.17	18.13	5.27	18.28	18.49	1.14
	Central	17.93	18.37	2.41	18.23	18.25	0.13	18.25	18.25	0.04
CLR	Safe	0.31	0.39	25.25	0.30	0.38	24.57	0.28	0.31	9.35
	Risky	0.30	0.37	25.20	0.30	0.32	6.23	0.31	0.47	49.64
	Local	0.27	0.32	17.09	0.28	0.31	9.77	0.32	0.36	14.38
	Central	0.34	0.56	68.30	0.31	0.32	3.20	0.31	0.31	0.89

Table 6.5: Banking system study: BSI, DIP and CLR ranges info for changes in p , v and u .

6.5 Conclusive remarks

One issue potential users of the new measure have to face is the problem of model calibration. The CLR relies on detailed information coming from the model, more than most of the other systemic risk measures proposed so far; it is therefore heavily exposed to the problem of model uncertainty. Contagion models aim at capturing the complex interactions between firms and the set of information they require (in their most generic forms) is very vast. As a consequence, calibrating contagion models under generic assumptions has always been a prohibitive task¹⁰. The solution most often found in literature is to reduce their complexity by forcing homogeneous assumptions on the system modeled. Nevertheless, we suspect that the capability of calibrating the model underlying the CLR measure will always be a demanding challenge.

Several strategies have been used in literature for network models that could be effective in this case too, depending on the data available. Billio, Getmansky, Lo, and Pelizzon (2012) calibrated a network model using CDS data, an approach that could be used effectively in both the sovereign and corporate context. In the context of banking systems, instead, regulatory data can be used to estimate the degree of infectivity of the system; the studies from Hale (2012) and Espinosa-Vega and Solé (2011) are two good examples based on interbank lending data while Cont, Moussa, and Santos (2011) were able to use a unique set of mutual exposure for the Brazilian banking case. Finally, we showed an alternative route in chapter 5 where a calibration example based on CDO data was presented on a reduced version of the model.

This chapter introduced a new systemic risk measure, the CLR, that can be applied in the context of contagion models and relies on an information set that is different from the one used by other measures. Although model dependent, we showed how to easily calculate it for a variety of models. We applied it in a controlled experimental environment represented by a small artificial banking system. The results we provided with respect to the behavior of the new measure show that the CLR adds a fresh perspective on systemic risk in the context of contagion models.

¹⁰The exceptions are, of course, the cases where the system consists of few names and/or the full information about their interaction is available.

Chapter 7

Conclusions

Systemic risk measurement is a fascinating field of research for many reasons. First of all, it is a relatively recent area where there are no established theories researchers have to adhere to. Indeed the variety of approaches contributes to form a very creative environment where methodologies and ideas coming from others fields are welcome. We already stressed in the introductory chapters the importance of the topic: macro-prudential approaches represents one of the pillars of today's regulators strategies in order to tackle banking instability. As a consequence, there is an unwritten balance required between sophisticated quantitative tools and the need to be understandable from a very wide audience. This adds another level of complexity for academics working on this topic.

This thesis contributes to such rich and challenging environment in at least four ways:

- The literature review presented in chapter 2 has been organized according to a novel categorization scheme: we hope that it could represent a useful map for researchers navigating through this vast and growing corpus of resources.
- The theoretical result we provided in the context of the CIMDO methodology extends significantly previous ones. At the same time, we presented a detailed study of the CIMDO *posterior* stability with respect to the inputs of the methodology that should serve as a guidance for practitioners. To the best of our knowledge, this is the first attempt to test in some details the limits of the CIMDO framework.
- We introduced a new systemic risk framework in the context of infection models

that is intuitive and yet capable of showing behaviors that are qualitative different from other more established measures. The new measure has many positive properties and we showed recipes to calculate it for several contagion models.

- Finally, we introduced an original contagion model in chapter 5; unlike other models of the same type, it exhibits good balance between tractability and flexibility. We also provided a rich mixture of pure theoretical results and practical algorithms. Finally, its versatility can be inferred from both the CDO application shown in chapter 5 and its use in the final part of chapter 6.

Part III

Appendices

Appendix A

Selected proofs

In this section we will report the proof of some of the results obtained in this work.

A.1 CIMDO extended independence theorem

We will show here the proof of theorem 2 by proving equation (4.7) that we rewrite here

$$\begin{aligned} P_q \{x_1 \diamond K_1, \dots, x_G \diamond K_G\} &= \prod_{i=1}^G P_q \{x_i \diamond K_i\} \\ &\Leftrightarrow \\ P_p \{x_1 \diamond K_1, \dots, x_G \diamond K_G\} &= \prod_{i=1}^G P_p \{x_i \diamond K_i\} \end{aligned} \tag{A.1}$$

Let also rewrite from page 71 equation (4.6) for convenience

$$\hat{p}(\mathbf{x}) = q(\mathbf{x}) \cdot \exp \left[-1 - \hat{\mu} - \sum_{m=1}^n \hat{\lambda}_m \chi_d^m(\mathbf{x}) \right]$$

Proof. We will prove the $\Rightarrow$ part of (A.1) by induction on n ; if $n = 1$, the assert is obvious. We will then assume it true for n and prove it for $n + 1$. In order to simplify notation, let

$$z = x_1, \dots, x_n, \quad y = x_{n+1}$$

Thanks to the assumption on the *prior*, we can write

$$q(z, y) = q_Z(z) \cdot q_Y(y)$$

where q_Z and q_Y are the marginals of, respectively $Z = (X_1, \dots, X_n)$ and $Y = X_n$. Let also define¹:

$$\begin{aligned}\bar{\mu} &= e^{-1-\mu} \\ f_Z(z) &= e^{-\sum_{m=1}^n \lambda_m \chi_d^m} \\ f_Y(y) &= e^{-\lambda_{n+1} \chi_d^{n+1}}\end{aligned}$$

so that, thanks to (4.6), we can factorize $p(z, y)$ as

$$p(z, y) = \bar{\mu} \cdot f_Z(z) q_Z(z) \cdot f_Y(y) q_Y(y)$$

Because of the inductive hypothesis, we only need to prove that

$$P_p\{x_1 \diamond K_1, \dots, x_n \diamond K_n, x_{n+1} \diamond K_{n+1}\} = P_p\{x_1 \diamond K_1, \dots, x_n \diamond K_n\} \cdot P_p\{x_{n+1} \diamond K_{n+1}\} \quad (\text{A.2})$$

The first probability of the right hand side in (A.2) is given by the following integral

$$\begin{aligned}P_p\{x_1 \diamond K_1, \dots, x_n \diamond K_n\} &= \int_{\mathfrak{R}} dy \int_A p(z, y) dz \\ &= \bar{\mu} \int_{\mathfrak{R}} f_Y(y) q_Y(y) dy \cdot \int_A f_Z(z) q_Z(z) dz\end{aligned}$$

where $\mathfrak{R}^n \supseteq A := \prod_{i=1}^n (l_i, u_i)$ with $(l_i, u_i) = (K_i, \infty)$ if $x_i > K_i$, $(l_i, u_i) = (-\infty, K_i]$ otherwise.

Similarly, the second term of the right hand side in (A.2) is

$$\begin{aligned}P_p\{x_{n+1} \diamond K_{n+1}\} &= \int_{\mathfrak{R}^n} dz \int_B p(z, y) dy \\ &= \bar{\mu} \int_{\mathfrak{R}^n} f_Z(z) q_Z(z) dz \cdot \int_B f_Y(y) q_Y(y) dy\end{aligned}$$

where $\mathfrak{R} \supseteq B := (l_{n+1}, u_{n+1})$. The product of the two terms is then

$$\left(\bar{\mu} \int_{\mathfrak{R}} f_Y(y) q_Y(y) dy \cdot \int_A f_Z(z) q_Z(z) dz \right) \cdot \left(\bar{\mu} \int_{\mathfrak{R}^n} f_Z(z) q_Z(z) dz \cdot \int_B f_Y(y) q_Y(y) dy \right)$$

¹As a reminder on the notation

$$\chi_{[K_i, \infty)}^i \equiv \chi_d^i$$

Rearranging terms leads to

$$\left(\bar{\mu} \int_{\mathfrak{R}} f_Y(y) q_Y(y) dy \cdot \int_{\mathfrak{R}^n} f_Z(z) q_Z(z) dz \right) \cdot \left(\bar{\mu} \int_A f_Z(z) q_Z(z) dz \cdot \int_B f_Y(y) q_Y(y) dy \right) \quad (\text{A.3})$$

Thanks to the last constraint in the optimization problem (4.3), we have:

$$\bar{\mu} \int_{\mathfrak{R}} f_Y(y) q_Y(y) dy \cdot \int_{\mathfrak{R}^n} f_Z(z) q_Z(z) dz = \int_{\mathfrak{R}^{n+1}} p(z, y) dz dy \equiv 1$$

and hence (A.3) simplifies to

$$\begin{aligned} & \bar{\mu} \int_A f_Z(z) q_Z(z) dz \cdot \int_B f_Y(y) q_Y(y) dy = \\ & \int_{A \times B} p(z, y) dz dy = \\ & P_p \{x_1 \diamond K_1, \dots, x_n \diamond K_n, x_{n+1} \diamond K_{n+1}\} \end{aligned}$$

proving (A.2) and the first part of the theorem.

The proof of the $\Leftarrow$ part of (A.1) is obtained by switching the role of p and q ; by redefining $\bar{\mu}$, $f_Z(z)$ and $f_Y(y)$ as the inverse of the previous definitions,

$$\begin{aligned} \bar{\mu} &= e^{+1+\mu} \\ f_Z(z) &= e^{+\sum_{m=1}^n \lambda_m \chi_d^m} \\ f_Y(y) &= e^{+\lambda_{n+1} \chi_d^{n+1}} \end{aligned}$$

we will now have²

$$\begin{aligned} p(z, y) &= p_Z(z) \cdot p_Y(y) \\ q(z, y) &= \bar{\mu} \cdot f_Z(z) p_Z(z) \cdot f_Y(y) p_Y(y) \end{aligned}$$

and the demonstration follows as above.

■

²The first equation by the independence assumption on p , the second by construction of p obtained in (4.6).

A.2 Contagion model - Joint default/survival probability

Let's remember the result we want to prove (see section 5.3.2, page 101, for details). Let $P(A, B)$ represent the probability of all names in A default while all names in B survive, i.e.

$$P(A, B) = P(Z_i = 1, Z_j = 0, \forall i \in A, \forall j \in B)$$

We have:

$$P(A, B) = P_1 + P_2 + P_3 \quad (\text{A.4})$$

with

$$P_1 = \Pi_A(bd) \cdot (1 - I_C) \cdot \Pi_B(1 - p)$$

$$P_2 = \left[\sum_{h=1}^{m_A} \sum_{k=0}^{m_A-h} \Theta_A^{h,k}(id, ps, bd) \right] \cdot (1 - I_C) \cdot \Pi_B(us) \quad (\text{A.5})$$

$$P_3 = \left[\sum_{h=0}^{m_A} \Lambda_A^h(p, ps) \right] \cdot I_C \cdot \Pi_B(us)$$

Proof. The three cases we need to identify are 1) infection free world, 2) infection starting inside A but not in C and 3) infection starting in C (and possibly but not necessarily in A). Lets see each of them in detail.

1. This is the simplest case. Names in A can only default idiosyncratically in non infective way with probability $\Pi_A(bd)$. Names in B don't need to get immunization and only need to avoid idiosyncratic default. The total probability of this case is hence

$$P_1 = \Pi_A(bd) \cdot (1 - I_C) \cdot \Pi_B(1 - p)$$

2. We are still requiring no infection coming from C (with an associated probability of $1 - I_C$) but now we need to consider contagion effects coming from names in A . Given that at least one of these names is required to default in an infective way, the rest can either be infected (partial survival) or default (both with or without spreading of the infection). Names in B now cannot accept anything less than full, unconditional survival. The total tally is

$$P_2 = \left[\sum_{h=1}^{m_A} \sum_{k=0}^{m_A-h} \Theta_A^{h,k}(id, ps, bd) \right] \cdot (1 - I_C) \cdot \Pi_B(us)$$

3. In this last case, the infection starts from C with probability of I_C . Names in B still need full survival ($\Pi_B(us)$) and names in A can either default (in both ways, infective or not) or default by infection. This leads to the last equation

$$P_3 = \left[\sum_{h=0}^{m_A} \Lambda_A^h(p, ps) \right] \cdot I_C \cdot \Pi_B(us)$$

■

A.3 Contagion model - Recursive formulae for Θ_A and Λ_A

We will prove the result showed in section 5.3.2, i.e.

$$\Lambda_A^h(x, y) = x_t \cdot \Lambda_W^{h-1}(x, y) + y_t \cdot \Lambda_W^h(x, y) \quad (\text{A.6})$$

$$\Theta_A^{h,k}(x, y, z) = x_t \cdot \Theta_W^{h-1,k}(x, y, z) + y_t \cdot \Theta_W^{h,k-1}(x, y, z) + z_t \cdot \Theta_W^{h,k}(x, y, z) \quad (\text{A.7})$$

Proof. We will only show proof of equation (A.7) as the one for equation (A.6) is analogous. Let s be one of the terms that form $\Theta_A^{h,k}(x, y, z)$; there are 3 possible cases for s regarding the value in position t : it contains either x_t , or y_t or z_t . In the first case, we can write

$$s = x_t \cdot g_1$$

where g_1 is a combination of $h - 1$ values taken from x (but not in the t position), k values from y (again not in the t position) and therefore is part of $\Theta_W^{h-1,k}(x, y, z)$. We can repeat the same reasoning for all other terms that contain x_t to get the first part of the right hand side: $x_t \cdot \Theta_W^{h-1,k}(x, y, z)$.

The second and third cases are similar. In the second case, i.e. s contains y_t , we can write

$$s = y_t \cdot g_2$$

and we know that g_2 will have h values from x , $k - 1$ from y and therefore will belong to $\Theta_W^{h,k-1}(x, y, z)$; repeating the reasoning for all terms of the form $s = y_t \cdot g_2$ we conclude the second part of the right hand side: $y_t \cdot \Theta_W^{h,k-1}(x, y, z)$.

Finally, when

$$s = z_t \cdot g_3$$

we have that g_3 is part of the summation of $\Theta_W^{h,k}(x, y, z)$ leading to the last part of the right hand side: $z_t \cdot \Theta_W^{h,k}(x, y, z)$. ■

Appendix B

Computational considerations

B.1 CIMDO methodology

The computational burden of the CIMDO methodology increases considerably with the size of the portfolio. Going from n to $n + 1$ names, not only we will need to calculate an extra integral for the marginal constraint on the new name, but all the integrals involved will be in $n + 1$ dimensions rather than n -dimensional ones. And they need to be calculated while solving a $n + 1$ optimization problem. Given the shape of the *posterior*, it is unlikely that analytical formulas will be of any use and quadrature methods become necessary. Multi dimensional numerical integration is usually a tricky business, and almost always a slow one, not one you want to rely upon when optimizing. Luckily there is an alternative: we can exploit the shape of the *posterior* (4.6) as it can be decomposed into the product of the *prior* and a piecewise constant function that depends only on the parameters and the default boundaries. Such decomposition allows us to speed up the calculations of the constraints. In particular, every marginal constraint integral can be decomposed into 2^{n-1} products of integrals on the *prior* and simple functions involving only the parameters. The last non-leaking constraint can be similarly split into 2^n components. For example, in 3 dimensions (let $\mathbf{d} = dx dy dz$):

$$\int_{\mathbb{R}} \int_{K_y}^{\infty} \int_{\mathbb{R}} p(x, y, z) \mathbf{d} = e^{-1-\mu-\lambda_y} \cdot \left[\begin{array}{l} e^{-\lambda_x-\lambda_z} \int_{K_x}^{\infty} \int_{K_y}^{\infty} \int_{K_z}^{\infty} q(x, y, z) \mathbf{d} \\ + e^{-\lambda_z} \int_{-\infty}^{K_x} \int_{K_y}^{\infty} \int_{K_z}^{\infty} q(x, y, z) \mathbf{d} \\ + e^{-\lambda_x} \int_{K_x}^{\infty} \int_{K_y}^{\infty} \int_{-\infty}^{K_z} q(x, y, z) \mathbf{d} \\ + \int_{-\infty}^{K_x} \int_{K_y}^{\infty} \int_{-\infty}^{K_z} q(x, y, z) \mathbf{d} \end{array} \right] \quad (\text{B.1})$$

Please note that the integrals on the *prior* need only to be calculated once, before the optimization takes place. Additionally, for some choices of the *prior* we might have closed form or quasi analytical techniques available to considerably speed up the calculations. The cost is that we will now need 2^n n -dimensional integrals on the *prior*, plus a little bit of housekeeping to store and use the decomposition.

The following figures show the computational times (in seconds) needed in order to

- Calculate the 2^n building block integrals on the *prior* - blue line
- Solve for $\hat{\mu}$ and $\hat{\lambda}$ to find the *posterior* - green line

The red line shows the total of the two. In particular, figure B.1 uses a numerical routine optimized for t-student distributions. On the other hand, figure B.2 relies on the methodology described earlier (section 3.1.5 on page 54) for Archimedean copulas. Note how in this latter case the calculations of the integrals over the *prior* takes roughly half the time while the solver performance is similar between the 2 examples.

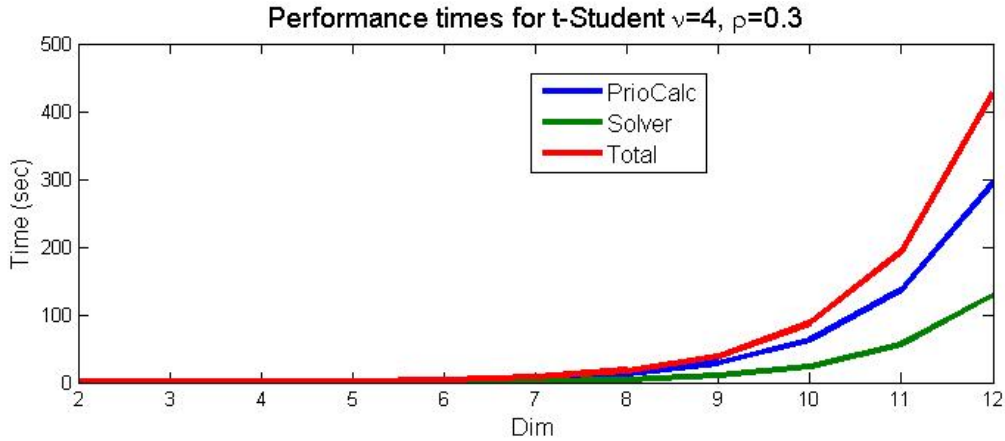


Figure B.1: Time (in sec) necessary to find the CIMDO *posterior* vs. dimensionality for t-Student distribution with $\nu = 4$ and $\rho = 0.3$.

Given the exponential complexity, this approach is not really feasible for large values of n (in our implementation, the cases $n \leq 15$ performed in a reasonable time).

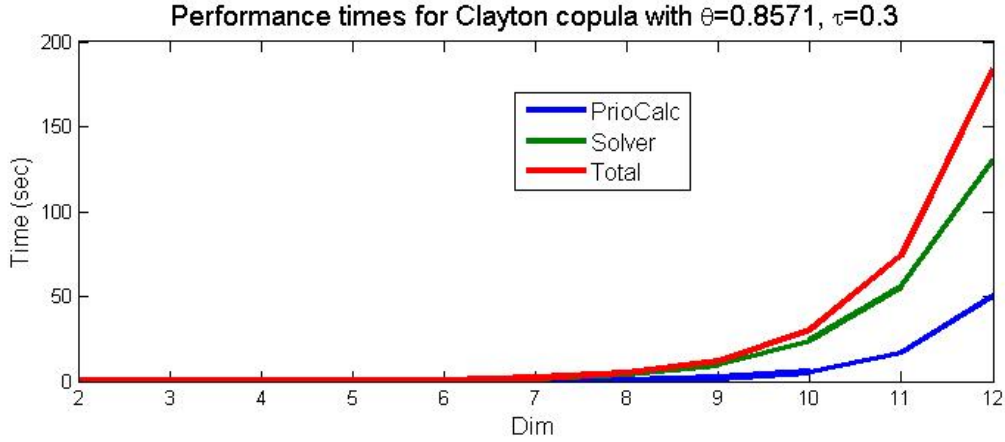


Figure B.2: Time (in sec) necessary to find the CIMDO *posterior* vs. dimensionality for a Clayton copula with $\tau = 0.3$.

B.1.1 Factorising out μ

Given the shape of the solution (4.6), the system of equations to be solved in order to recover the optimal parameters becomes:

$$e^{-1-\mu} \int_{\mathbb{R}^n} g(\mathbf{x}) \cdot e^{[\sum_{i=1}^n -\lambda_i \cdot \chi_d^i(\mathbf{x})]} \mathbf{d}\mathbf{x} = 1$$

$$e^{-1-\mu} \int_{\mathbb{R}^n} g(\mathbf{x}) \cdot e^{[\sum_{i=1}^n -\lambda_i \cdot \chi_d^i(\mathbf{x})]} \cdot \chi_d^j(\mathbf{x}) \mathbf{d}\mathbf{x} = PoD_t^{X_j}$$

where $j = 1, \dots, n$. We can notice how the term $e^{-1-\mu}$ can be factorized out; for example, from the first equation we can get

$$e^{1+\mu} = \int_{\mathbb{R}^n} g(\mathbf{x}) \cdot e^{[\sum_{i=1}^n -\lambda_i \cdot \chi_d^i(\mathbf{x})]} \mathbf{d}\mathbf{x} \quad (\text{B.2})$$

The latter equation is useful in two ways; first, we can substitute it in the rest of the system obtaining n equations in the n unknowns $\lambda_1, \dots, \lambda_n$:

$$\int_{\mathbb{R}^n} g(\mathbf{x}) \cdot e^{[\sum_{i=1}^n -\lambda_i \cdot \chi_d^i(\mathbf{x})]} \cdot \chi_d^j(\mathbf{x}) \mathbf{d}\mathbf{x} = PoD_t^{X_j} \cdot \int_{\mathbb{R}^n} g(\mathbf{x}) \cdot e^{[\sum_{i=1}^n -\lambda_i \cdot \chi_d^i(\mathbf{x})]} \mathbf{d}\mathbf{x}$$

or

$$\int_{\mathbb{R}^n} g(\mathbf{x}) \cdot e^{[\sum_{i=1}^n -\lambda_i \cdot \chi_d^i(\mathbf{x})]} \cdot \left\{ \chi_d^j(\mathbf{x}) - PoD_t^{X_j} \right\} \mathbf{d}\mathbf{x} = 0$$

The second way to use (B.2), once the n -dimensional sub-system is solved, is simply

to determine μ as:

$$\mu = \ln \left\{ \int_{\mathfrak{R}^n} e^{\sum_{i=1}^n -\lambda_i \cdot \chi_d^i(\mathbf{x})} g(\mathbf{x}) d\mathbf{x} \right\} - 1$$

.

B.2 Loss distribution in the contagion model introduced in chapter 5

The structure of the system of equations (5.29)-(5.30) (page 106) naturally suggests a recursive algorithm for the portfolio loss distribution $P\{L_n = h\}$ that is tractable even for real-life values of n . The required inputs are the vectors of p , u , v , d and c as defined in section 5.2 (page 97).

The algorithm can be split into three steps:

1. **Step 1.** Initialize the boundary conditions.
2. **Step 2.** Update the $\beta(\cdot, \cdot)$. This can be done in place but must be done before changing the $\alpha(\cdot, \cdot)$ (as the $\beta(\cdot, \cdot)$ depend on them).
3. **Step 3.** Update the $\alpha(\cdot, \cdot)$. Again, this can be done in place.

The following frame shows the algorithm written in Matlab code:

```

1  % Boundary conditions – numB represents the number of loss buckets ...
   (for example numB = n if the LGDs are the same)
2  bt = zeros(numB+1, numB+1); % beta
3  ap = zeros(numB+1, numB+1); ap(1,1) = 1; % alpha
4  MaxLoss = 0; %maximum level of losses reacheable so far
5
6  tfs = (1-p) .* u; % full survival
7  tps = (1-p) .* (1-u); % partial survival
8  tid = p .* v; % infective default
9  tni = p .* (1-v); % non-infective default
10
11 for i=1:n % recursive loop over the number of names n
12     MaxLoss = MaxLoss + max(d(i), c(i));
13     for h = MaxLoss:-1:0
14         % updating betas first (as they depend on the alphas)
15         for k = MaxLoss:-1:0
16             bt(h+1, k+1) = bt(h+1, k+1) * tfs(i);
17             if k ≥ c(i); bt(h+1, k+1) = bt(h+1, k+1) + tps(i) * ...
                 bt(h+1, k+1 - c(i)); end;
18             if h ≥ d(i)
19                 bt(h+1, k+1) = bt(h+1, k+1) + p(i) * bt(h+1 - d(i), ...
                     k+1);
20                 bt(h+1, k+1) = bt(h+1, k+1) + tid(i) * ap(h+1 - ...
                     d(i), k+1);
21             end
22         end
23
24         % updating alphas now
25         for k = MaxLoss:-1:c(i)
26             ap(h+1, k+1) = ap(h+1, k+1) * tfs(i);
27             if k ≥ c(i); ap(h+1, k+1) = ap(h+1, k+1) + tps(i) * ...
                 ap(h+1, k+1 - c(i)); end;
28             if h ≥ d(i); ap(h+1, k+1) = ap(h+1, k+1) + tni(i) * ...
                 ap(h+1 - d(i), k+1); end;
29         end
30     end
31 end

```

B.2.1 Performances

The performances of the algorithm described in the previous section are directly comparable to the ones shown by the standard one factor Gaussian model embedded with the algorithm presented in Andersen, Sidenius, and Basu (2003), with the difference that while the latter is linear in n , our method shows quadratic complexity.

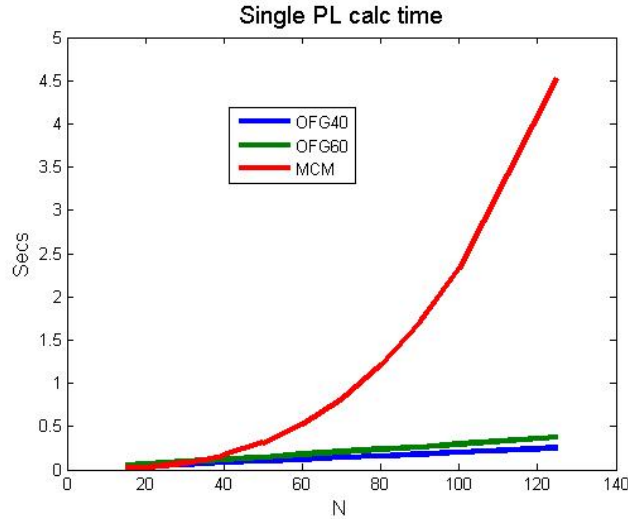


Figure B.3: Times (in seconds) necessary to calculate portfolio loss distribution for a 0 – 5% tranche at $t = 5$ when varying the portfolio dimension n .

Figure B.3 provides a comparison of performance times between the model introduced here (MCM) and two instances of the standard one factor Gaussian model (OFG) that differ for the number of points used to sample the common factor distribution¹. Note also that, even in our Matlab prototyped implementation, the contagion model takes less than 5 seconds to calculate the portfolio loss distribution for real-life sized portfolios.

Finally, note also that is possible to get accurate risk calculations by inverting the above algorithm to remove a given name and then calculate derivatives of (5.29), completely analogous to what is done for Andersen, Sidenius, and Basu (2003).

¹Please note that in no way we tried to optimize our Matlab code to make usable for actual applications of either model types; readers should be wary that by changing the implementations and using proper programming languages, the times shown here should both decrease. This said, we think that the relative costs of the two models should stay roughly the same as the ones presented here.

Appendix C

Portfolio spreads

In this section we will report few additional statistics on the set of quotes used for section 5.4.2. Table C.1 contains the min, max, mean and standard deviation of the CDS par quotes used for the portfolio. Every individual survival curve is built starting from 4 CDS quotes (for 3, 5, 7 and 10 years maturities respectively) whose distribution has been generated to mimic real life portfolios¹. Each name is assumed to have a 40% recovery and have the same weight inside the portfolio.

Maturity	3Y	5Y	7Y	10Y
Min	59.2367	69.0244	81.0120	93.1496
Max	115.4230	180.2951	243.2890	315.2737
Mean	87.9749	114.7784	148.4421	190.1074
Std Dev	16.8945	28.0898	38.8019	55.4140

Table C.1: Portfolio single name CDS par spreads statistics. The reported numbers are in basis points.

Figure C.1 reports the distribution of the quotes for each maturity.

Table C.2 instead reports the full list of quotes used.

¹The results shown here are based on a portfolio built to replicate the riskiness and dispersion of the ITraxx portfolio.

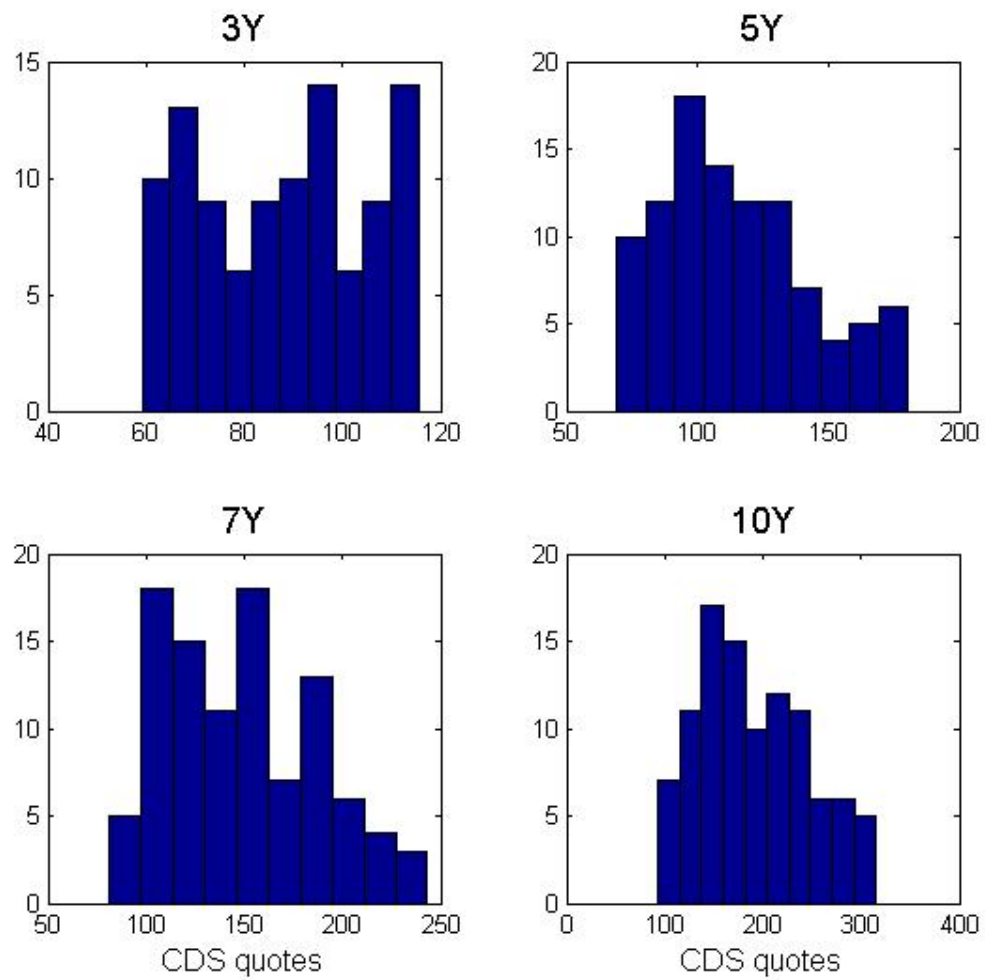


Figure C.1: CDS quotes distributions

Name	3Y	5Y	7Y	10Y	Name	3Y	5Y	7Y	10Y	Name	3Y	5Y	7Y	10Y
N.1	106	163	176	272	N.34	69	80	87	94	N.67	64	74	89	125
N.2	95	101	118	156	N.35	109	147	195	212	N.68	66	95	101	141
N.3	114	180	197	312	N.36	108	148	180	235	N.69	87	128	183	282
N.4	114	147	218	237	N.37	82	85	98	105	N.70	110	132	188	210
N.5	83	129	190	299	N.38	69	79	99	102	N.71	60	87	113	146
N.6	97	99	149	232	N.39	111	174	225	291	N.72	111	152	208	315
N.7	98	142	206	254	N.40	78	120	147	157	N.73	105	142	157	180
N.8	97	106	152	154	N.41	104	128	147	182	N.74	110	112	145	159
N.9	74	77	81	121	N.42	64	69	108	170	N.75	115	165	214	275
N.10	99	118	185	189	N.43	92	95	109	132	N.76	62	87	89	93
N.11	84	103	151	222	N.44	106	107	110	121	N.77	89	94	140	209
N.12	69	90	114	158	N.45	96	138	192	244	N.78	100	110	153	200
N.13	100	145	169	238	N.46	90	106	154	171	N.79	115	160	237	301
N.14	97	106	114	147	N.47	98	109	133	183	N.80	84	125	131	142
N.15	114	138	186	211	N.48	104	109	170	249	N.81	68	85	127	188
N.16	102	118	154	218	N.49	87	109	139	164	N.82	62	77	101	126
N.17	110	174	231	250	N.50	88	115	171	253	N.83	97	133	156	197
N.18	67	77	116	134	N.51	96	118	175	231	N.84	59	94	104	110
N.19	106	121	189	229	N.52	79	123	188	250	N.85	80	90	116	139
N.20	70	80	110	141	N.53	95	128	144	170	N.86	114	177	182	263
N.21	79	118	160	212	N.54	86	98	147	164	N.87	74	93	123	193
N.22	112	131	191	277	N.55	72	79	90	113	N.88	83	131	155	221
N.23	81	108	113	117	N.56	76	119	150	166	N.89	97	129	183	256
N.24	89	131	204	220	N.57	111	176	223	238	N.90	69	74	118	131
N.25	92	117	118	142	N.58	73	91	124	144	N.91	60	81	123	173
N.26	68	100	119	156	N.59	93	133	151	162	N.92	69	85	108	172
N.27	68	93	107	149	N.60	76	90	113	147	N.93	67	102	142	174
N.28	99	143	181	190	N.61	63	73	109	110	N.94	69	87	113	121
N.29	72	111	121	181	N.62	113	162	209	282	N.95	93	105	130	175
N.30	90	143	150	190	N.63	72	92	145	193	N.96	73	86	117	136
N.31	65	102	102	150	N.64	89	101	131	180	N.97	106	169	243	293
N.32	106	161	169	210	N.65	98	121	148	236	N.98	92	98	152	232
N.33	73	109	137	212	N.66	61	93	143	212	N.99	106	123	166	169
N.100	83	99	108	120										

Table C.2: Portfolio single name CDS par spreads. The reported numbers are in basis points.

List of Figures

2.1	Systemic risk popularity	28
2.2	Comparison of the area exploited by different risk measures when $n = 2$	42
2.3	Comparison of the area exploited by different risk measures when $n = 3$	43
3.1	OFG Conditional probability of default.	59
3.2	Synthetic CDO structure	61
3.3	OFG loss distribution vs. ρ	64
4.1	Current <i>PoD</i> vs. the historical default thresholds K	75
4.2	Comparison in the elliptical family	77
4.3	Comparison among the Archimedean copulas	78
4.4	<i>JPoD</i> risk measure on q and p vs. τ and n for a Clayton copula	80
4.5	<i>BSI</i> risk measure on q and p vs. τ and n for a Clayton copula	81
4.6	<i>SFM</i> risk measure on q and p vs. τ and n for a Clayton copula	82
4.7	<i>PAO</i> risk measure on q and p vs. τ and n for a Clayton copula	83
4.8	Change in the <i>posterior</i> systemic risk measures vs. changes in either the <i>PoD</i> or the K	86
4.9	Change in the <i>posterior</i> <i>SFM</i> and <i>PAO</i> when n and K change.	88
4.10	CIMDO inputs relative importance.	89
5.1	Recursive links between the $\alpha(\cdot, \cdot)$ and the $\beta(\cdot, \cdot)$	108
5.2	Price ranges 0-3%	111
5.3	Price ranges 3-6%	112
5.4	Price ranges 6-9%	112
5.5	Price ranges 9-12%	113
5.6	Price ranges 12-22%	113
5.7	Strategy for the calibration exercise performed.	115

5.8	Parameter skews	116
5.9	Single name deltas	119
6.1	Dependency structure of system C.	123
6.2	Loss distribution for the three systems.	124
6.3	Davis&Lo contagion models: sensitivity to p and q	130
6.4	Sakata et. al. contagion model: sensitivity to p and q when $q' = 10\%$	131
6.5	Sakata et. al. contagion model: sensitivity to p and q when $q' = 20\%$	132
6.6	γ and loss distribution	137
6.7	CoCLR and DIP vs. T for the banking system.	138
6.8	Risk measure sensitivity with respect to p , u and v	139
B.1	CIMDO performance for t-Student	156
B.2	CIMDO performance for Clayton copula	157
B.3	Performance comparison: contagion model vs. OFG when n varies	161
C.1	CDS quotes distributions	164

List of Tables

1.1	Selected systemic events.	19
3.1	Archimedean copula families	54
4.1	Copulas parameters	76
4.2	Variations due to either <i>PoD</i> or <i>HistPoD</i>	86
4.3	Comparison of variations in the <i>posterior</i> measures	88
5.1	Spread theoretical ranges for OFG and contagion model.	114
5.2	Contagion model calibration	114
6.1	Thought experiment: systems info	123
6.2	Artificial banking system: info	134
6.3	Individual default probabilities	135
6.4	Systemic risk indicators for the banking system (contagion on/off) . . .	135
6.5	Banking system study: BSI, DIP and CLR ranges info for changes in p , v and u	142
C.1	Portfolio single name CDS par spreads statistics	163
C.2	Portfolio single name CDS	165

Bibliography

- Viral Acharya and Tanju Yorulmazer. Too many to fail. an analysis of time-inconsistency in bank closure policies. *Journal of Financial Intermediation*, 16(1): 1–31, 2007.
- Viral Acharya, Lasse Pedersen, Thomas Philippon, and Matthew Richardson. Measuring systemic risk. *NYU Working Paper*, 2010.
- Tobias Adrian and Markus Brunnermeier. Covar. Technical report, National Bureau of Economic Research, Inc, 2011.
- Tobias Adrian and Hyun Shin. The changing nature of financial intermediation and the financial crisis of 2007–2009. *Annual Review of Economics*, 2:603–618, 2010.
- David Aikman, Piergiorgio Alessandri, Bruno Eklund, Prasanna Gai, Sujit Kapadia, Elizabeth Martin, Nada Mora, Gabriel Sterne, and Matthew Willison. Funding liquidity risk in a quantitative model of systemic stability. *Bank of England, working paper No. 372*, 2009.
- Lucia Alessi and Carsten Detken. 'real time'early warning indicators for costly asset price boom/bust cycles: a role for global liquidity. 2009.
- Franklin Allen and Douglas Gale. Financial contagion. *Journal of political economy*, 108(1):1–33, 2000.
- Franklin Allen and Douglas Gale. *Understanding financial crises*. Oxford University Press, 2009.
- Claudi Alsina, Maurice Frank, and Bert Schweizer. Problems on associative functions. *Aequationes Mathematicae*, 66(1-2):128–140, 2003.

- Salah Amraoui, Laurent Cousot, Sébastien Hitier, and Jean-Paul Laurent. Pricing cdfs with state-dependent stochastic recovery rates. *Quantitative Finance*, 12(8): 1219–1240, 2012.
- Leif Andersen and Jakob Sidenius. Extensions to the gaussian copula: Random recovery and random factor loadings. *Journal of Credit Risk Volume*, 1(1):05, 2004.
- Leif Andersen, Jakob Sidenius, and Susanta Basu. All your hedges in one basket. *Risk Magazine*, 16(11):67–72, 2003.
- Renzo Avesani and Antonio Garcia Pascual. A new risk indicator and stress testing tool: A multifactor nth-to-default cds basket. 2006.
- Shariar Azizpour, Kay Giesecke, et al. Self-exciting corporate defaults: contagion vs. frailty.
- Tomasz Bielecki, Areski Cousin, Stéphane Crépey, and Alexander Herbertsson. Dynamic hedging of portfolio credit risk in a markov copula model. *Journal of Optimization Theory and Applications*, pages 1–13, 2012.
- Monica Billio, Mila Getmansky, Andrew Lo, and Lorian Pelizzon. Econometric measures of connectedness and systemic risk in the finance and insurance sectors. *Journal of Financial Economics*, 104(3):535–559, 2012.
- Fischer Black and Myron Scholes. The pricing of options and corporate liabilities. *The journal of political economy*, pages 637–654, 1973.
- Michael Bordo, Barry Eichengreen, Daniela Klingebiel, and Maria Soledad Martinez-Peria. Financial crises: lessons from the last 120 years. *Economic Policy*, 16(32): 51–82, 2001.
- Claudio Borio. Towards a macroprudential framework for financial supervision and regulation? *CESifo Economic Studies*, 49(2):181–215, 2003.
- Claudio Borio. Implementing the macroprudential approach to financial regulation and supervision. 2009.
- Claudio Borio and Mathias Drehmann. Towards an operational framework for financial stability:” fuzzy” measurement and its consequences. *Documentos de Trabajo (Banco Central de Chile)*, (544):1, 2009.

- Claudio Borio and Philip Lowe. Securing sustainable price stability: should credit come back from the wilderness? 2004.
- Christian Brownlees and Robert Engle. Volatility, correlation and tails for systemic risk measurement. *Available at SSRN 1611229*, 2012.
- Markus Brunnermeier. Deciphering the liquidity and credit crunch 2007-2008. *The Journal of economic perspectives*, 23(1):77–100, 2009.
- Markus Brunnermeier and Lasse Heje Pedersen. Market liquidity and funding liquidity. *Review of Financial studies*, 22(6):2201–2238, 2009.
- Markus Brunnermeier, Charles Goodhart, Avinash Persaud, Andrew Crockett, and Hyun Shin. The fundamental principles of financial regulation. 2009.
- Xavier Burtschell, Jon Gregory, and Jean-Paul Laurent. A comparative analysis of cdo pricing models. *The Journal of Derivatives*, 16(4):9–37, 2009.
- Zhili Cao. Multi-covar and shapley value: A systemic risk measure. Technical report, 2012.
- Andrew Cardinali and Jakob Nordmark. How informative are bank stress tests. *Bank Opacity in the European Union. Lund University*, 2011.
- Jorge Chan-Lau, Marco Espinosa, Kay Giesecke, and Juan Solé. Assessing the systemic implications of financial linkages. *IMF Global Financial Stability Report*, 2, 2009.
- David Clayton. A model for association in bivariate life tables and its application in epidemiological studies of familial tendency in chronic disease incidence. *Biometrika*, 65(1):141–151, 1978.
- Piet Clement. The term ‘macroprudential’: origins and evolution. *BIS Quarterly Review*, March, 2010.
- Rama Cont, Romain Deguest, and Yu Hang Kan. Default intensities implied by cdo spreads: inversion formula and model calibration. *SIAM Journal on Financial Mathematics*, 1(1):555–585, 2010.
- Rama Cont, Amal Moussa, and Edson Santos. Network structure and systemic risk in banking systems. *Edson Bastos e, Network Structure and Systemic Risk in Banking Systems (December 1, 2010)*, 2011.

- Areski Cousin, Diana Dorobantu, and Didier Rullière. An extension of davis and lo's contagion model. *Quantitative Finance*, 13(3):407–420, 2013.
- Glenis Crane and John Van Der Hoek. Using distortions of copulas to price synthetic cdo's. *Insurance: Mathematics and Economics*, 42(3):903–908, 2008.
- Andrew Crockett. Marrying the micro-and macro-prudential dimensions of financial stability. *BIS speeches*, 21, 2000.
- Udaibir Das, Nicholas Blancher, and Serkan Arslanalp. Japan: Financial sector stability assessment update. *IMF Global Financial Stability Report*, 2012.
- Tom Daula. Systemic risk: Relevance, risk management challenges and open questions, 2010.
- Mark Davis and Violet Lo. Infectious defaults. *Quantitative Finance*, 1(4):382–387, 2001.
- Olivier De Bandt and Philipp Hartmann. Systemic risk: A survey. 2000.
- Jacques De Larosière. *The High-Level Group on financial supervision in the EU: report*. European Commission, 2009.
- Giuseppe Di Graziano and Leonard Rogers. A dynamic approach to the modeling of correlation credit derivatives using markov chains. *International Journal of Theoretical and Applied Finance*, 12(01):45–62, 2009.
- Mathias Drehmann, Claudio Borio, and Konstantinos Tsatsaronis. Anchoring counter-cyclical capital buffers: the role of credit aggregates. 2011.
- Ernst Eberlein, Rüdiger Frey, and Ernst August Von Hammerstein. *Advanced credit portfolio modeling and CDO pricing*. Springer, 2008.
- Daniel Egloff, Markus Leippold, and Paolo Vanini. A simple model of credit contagion. *Journal of Banking & Finance*, 31(8):2475–2492, 2007.
- Robert Engle. Dynamic conditional correlation: A simple class of multivariate generalized autoregressive conditional heteroskedasticity models. *Journal of Business & Economic Statistics*, 20(3):339–350, 2002.

- Robert Engle. *Anticipating correlations: a new paradigm for risk management*. Princeton University Press, 2009.
- Paul Erdos and Alfred Renyi. On random graphs. *Publ. Math. Debrecen*, 6:290–297, 1959.
- Eymen Errais, Kay Giesecke, and Lisa Goldberg. Affine point processes and portfolio credit risk. *SIAM Journal on Financial Mathematics*, 1(1):642–665, 2010.
- Marco A Espinosa-Vega and Juan Solé. Cross-border financial surveillance: a network perspective. *Journal of Financial Economic Policy*, 3(3):182–205, 2011.
- European Banking Authority. Ceb’s press release on the results on the eu-wide stress testing exercise. 2009.
- European Banking Authority. Aggregate outcome of the 2010 eu wide stress test exercise coordinated by ceb in cooperation with the ecb. 2010.
- European Banking Authority. Eu-wide stress testing aggregate report. 2011.
- European Central Bank. Recent advances in modelling (sic) systemic risk using network analysis. 2009.
- European Central Bank. New quantitative measures of systemic risk. 2010.
- Federal Reserve. The supervisory capital assessment program: Overview of results. 2009a.
- Federal Reserve. The supervisory capital assessment program: Design and implementation. 2009b.
- Federal Reserve. 2011 comprehensive capital assessment review (ccar). 2011.
- Federal Reserve. 2012 comprehensive capital assessment review (ccar). 2012.
- Federal Reserve. 2013 comprehensive capital assessment review (ccar). 2013.
- Financial Services Authority. Stress and scenario testing. 2008.
- Financial Services Authority. Stress and scenario testing. feedback on cp08/24 and final rules. 2009.

- Financial Stability Board. Report of the financial stability forum on addressing procyclicality in the financial system. Financial Stability Forum, 2009.
- Christopher Finger. Issues in the pricing of synthetic cdos. *The Journal of Credit Risk*, 1(1):113–124, 2004.
- Antonella Foglia. Stress testing credit risk: a survey of authorities’ approaches. *Bank of Italy Occasional Paper*, (37), 2008.
- Maurice Frank. On the simultaneous associativity of $f(x, y)$ and $x+y-f(x, y)$. *Aequationes Mathematicae*, 19(1):194–226, 1979.
- Rüdiger Frey and Jochen Backhaus. Dynamic hedging of synthetic cdo tranches with spread risk and default contagion. *Journal of Economic Dynamics and Control*, 34(4):710–724, 2010.
- G-10 Working Group. Consolidation in the financial sector, 2001.
- G-20 Working Group. Enhancing sound regulation and strengthening transparency: final report., 2009.
- Nicolae Gârleanu and Lasse Heje Pedersen. Liquidity and risk management. *The American Economic Review*, 97(2):193–197, 2007.
- Christian Genest. Frank’s family of bivariate distributions. *Biometrika*, 74(3):549–555, 1987.
- Christian Genest and Jock Mackay. The joy of copulas: Bivariate distributions with uniform marginals. *The American Statistician*, 40(4):280–283, 1986.
- Alan Genz and Frank Bretz. Comparison of methods for the computation of multivariate t probabilities. *Journal of Computational and Graphical Statistics*, 11(4):950–971, 2002.
- Kay Giesecke and Stefan Weber. Cyclical correlations, credit contagion, and portfolio losses. *Journal of Banking & Finance*, 28(12):3009–3036, 2004.
- Dale Gray and Andreas Jobst. Systemic cca-a model approach to systemic risk. In *Deutsche Bundesbank/Technische Universität Dresden Conference: Beyond the Financial Crisis: Systemic Risk, Spillovers and Regulation, Dresden*. Citeseer, 2010.

- Dale Gray and Samuel Malone. *Macrofinancial risk analysis*, volume 433. Wiley. com, 2008.
- Dale Gray, Robert Merton, and Zvi Bodie. New framework for measuring and managing macrofinancial risk and financial stability. Technical report, National Bureau of Economic Research, 2007.
- Emil Gumbel. Distributions des valeurs extrêmes en plusieurs dimensions. *Publ. Inst. Statist. Univ. Paris*, 9:171–173, 1960.
- AG Haldane. Rethinking the financial network, 2009.
- Galina Hale. Bank relationships, business cycles, and financial crises. *Journal of International Economics*, 88(2):312–325, 2012.
- Marius Hofert and Matthias Scherer. Cdo pricing with nested archimedean copulas. *Quantitative Finance*, 11(5):775–787, 2011.
- Philip Hougaard. A class of multivariate failure time distributions. *Biometrika*, 73(3):671–678, 1986.
- Xin Huang, Hao Zhou, and Haibin Zhu. Systemic risk contributions. *Journal of financial services research*, 42(1-2):55–83, 2012.
- John Hull and Alan White. An improved implied copula model and its application to the valuation of bespoke cdo tranches. *Journal of investment management*, 8(3):11, 2010.
- Timothy Hutchinson and Chin Lai. *Continuous bivariate distributions, emphasising applications*. Rumsby Scientific Publishing Adelaide, 1990.
- Robert Jarrow and Stuart Turnbull. Pricing derivatives on financial securities subject to credit risk. *The journal of finance*, 50(1):53–85, 1995.
- Robert A Jarrow and Fan Yu. Counterparty risk and the pricing of defaultable securities. *the Journal of Finance*, 56(5):1765–1799, 2001.
- Xisong Jin and Francisco De Simone. Banking systemic vulnerabilities: A tail-risk dynamic cindo approach. Technical report, 2013.

- Martin Johansson. A systemic risk indicator for the swedish banking system. pages 1–8, 2011.
- Mark Joshi and Alan Stacey. *Intensity gamma: a new approach to pricing portfolio credit derivatives*. Centre for Actuarial Studies, Department of Economics, the University of Melbourne, 2006.
- Anna Kalemanova, Bernd Schmid, and Ralf Werner. The normal inverse gaussian distribution for synthetic cdo pricing. *The journal of derivatives*, 14(3):80–94, 2007.
- Clark Kimberling. A probabilistic interpretation of complete monotonicity. *Aequationes Mathematicae*, 10(2):152–164, 1974.
- Charles Kindleberger and Robert Aliber. *Manias, panics and crashes: a history of financial crises*. Palgrave Macmillan, 2011.
- Malcolm Knight. Marrying the micro and macroprudential dimensions of financial stability: six years on. In *speech delivered at the 14th International Conference of Banking Supervisors, BIS Speeches, October*, volume 119, 2006.
- Siem Jan Koopman, Andre Lucas, and Bernd Schwaab. Macro, frailty, and contagion effects in defaults: Lessons from the 2008 credit crisis. 2010.
- Ugur Koyluoglu and Jim Stoker. Honour your contribution. *Risk*, 15(4):90–94, 2002.
- Mark Kritzman, Yuanzhen Li, Sebastien Page, and Roberto Rigobon. Principal components as a measure of systemic risk. *SSRN eLibrary*, 2010.
- Luc Laeven and Fabian Valencia. *Resolution of Banking Crises: The Good the Bad and the Ugly*, volume 10. International Monetary Fund, 2010.
- David Lando. *Credit risk modeling: theory and applications*. Princeton University Press, 2009.
- Jean-Paul Laurent, Areski Cousin, and Jean-David Fermanian. Hedging default risks of cdos in markovian contagion models. *Quantitative Finance*, 11(12):1773–1791, 2011.
- Jong Han Lee, Jaemin Ryu, and Dimitrios Tsomocos. Measures of systemic risk and financial fragility in korea. *Annals of Finance*, pages 1–30, 2012.

- David Li. On default correlation: A copula function approach. *Journal of Fixed Income*, March, 2000.
- Francis Longstaff and Arvind Rajan. An empirical analysis of the pricing of collateralized debt obligations. *The Journal of Finance*, 63(2):529–563, 2008.
- Eva Lütkebohmert and Michael Gordy. Granularity adjustment for basel ii. Technical report, Discussion Paper, Series 2: Banking and Financial Supervision, 2007.
- Jan-Frederik Mai and Matthias Scherer. *Simulating Copulas: Stochastic Models, Sampling Algorithms, and Applications*, volume 4. World Scientific, 2012.
- Allan Malz. Risk-neutral systemic risk indicators. *Available at SSRN 2241567*, 2013.
- Richard Markeloff, Geoffrey Warner, and Elizabeth Wollin. Modeling systemic risk to the financial system. 2011.
- Richard Markeloff, Geoffrey Warner, and Elizabeth Wollin. Modeling systemic risk to the financial system. a review of additional literature. 2012.
- Sheri Markose, Simone Giansante, Mateusz Gatkowski, and Ali Shaghaghi. Too interconnected to fail: Financial contagion and systemic risk in network model of cds and other credit enhancement obligations of us banks. 2010.
- Robert May. Will a large complex system be stable? *Nature*, 238:413–414, 1972.
- Robert May. *Stability and complexity in model ecosystems*, volume 6. Princeton University Press, 2001.
- Alexander McNeil, Rüdiger Frey, and Paul Embrechts. *Quantitative risk management: concepts, techniques, and tools*. Princeton university press, 2010.
- Robert Merton. On the pricing of corporate debt: The risk structure of interest rates. *The Journal of Finance*, 29(2):449–470, 1974.
- Camelia Minoiu and Javier Reyes. A network analysis of global banking: 1978–2010. *Journal of Financial Stability*, 2013.
- Roger Nelsen. Properties of a one-parameter family of bivariate distributions with specified marginals. *Communications in statistics-Theory and methods*, 15(11):3277–3285, 1986.

- Roger Nelsen. *An introduction to copulas*. Springer, 1999.
- Peter Neu and Reimer Kühn. Credit risk enhancement in a network of interdependent firms. *Physica A: Statistical Mechanics and its Applications*, 342(3):639–655, 2004.
- Erlend Nier, Jing Yang, Tanju Yorulmazer, and Amadeo Alentorn. Network models and financial stability. *Journal of Economic Dynamics and Control*, 31(6):2033–2060, 2007.
- Ed Parcell. Loss unit interpolation in the collateralised debt obligation pricing model, 2006.
- Ed Parcell. Wiping the smile off your base (correlation curve). 2007.
- Lasse Heje Pedersen. When everyone runs for the exit. *International Journal of Central Banking*, 2009.
- Roger Rabemananjara and Jean-Michel Zakoïan. Threshold arch models and asymmetries in volatility. *Journal of Applied Econometrics*, 8(1):31–49, 1993.
- Steven Radelet and Jeffrey Sachs. The onset of the east asian financial crisis. In *Currency crises*, pages 105–162. University of Chicago Press, 2000.
- Deyan Radev. Assessing systemic fragility-a probabilistic perspective. *Available at SSRN 2090242*, 2012a.
- Deyan Radev. Systemic risk and sovereign debt in the euro area. 2012b.
- Jean-Charles Rochet and Jean Tirole. Interbank lending and systemic risk. *Journal of Money, credit and Banking*, 28(4):733–762, 1996.
- Daniel Rösch and Birker Winterfeldt. Estimating credit contagion in a standard factor model. *Risk*, pages 78–82, 2008.
- Robert. Rubin, Alan Greenspan, Arthur Levitt, and Brooksley Born. Hedge funds, leverage, and the lessons of long-term capital management. 1999.
- Ayaka Sakata, Masato Hisakado, and Shintaro Mori. Infectious default model with recovery and continuous limits. *Journal of the Physical Society of Japan*, 76(5):054801, 2007.

- Philipp Schönbucher and Dirk Schubert. Copula-dependent default risk in intensity models. In *Working paper, Department of Statistics, Bonn University*. Citeseer, 2001.
- Bernd Schwaab, Siem Jan Koopman, and André Lucas. Systemic risk diagnostics: coincident indicators and early warning signals. 2011.
- Steven Schwarcz. Systemic risk. In *American Law & Economics Association Annual Meetings*, page 20. bepress, 2008.
- Berthold Schweizer and Abe Sklar. *Probabilistic metric spaces*. DoverPublications. com, 1983.
- Securities, Exchange Commission, et al. Findings regarding the market events of may 6, 2010. *Report of the Staffs of the CFTC and SEC to the Joint Advisory Committee on Emerging Regulatory Issues*, 2010.
- Miguel Segoviano. Consistent information multivariate density optimizing methodology. 2006.
- Miguel Segoviano and Charles Goodhart. Banking stability measures. *IMF Working Papers*, pages 1–54, 2009.
- Lloyd Shapley. A value for n-person games. Technical report, DTIC Document, 1952.
- Abe Sklar. *Fonctions de répartition à n dimensions et leurs marges*. Université Paris 8, 1959.
- Abe Sklar. Random variables, distribution functions, and copulas: a personal look backward and forward. *Lecture notes-monograph series*, pages 1–14, 1996.
- Nikola Tarashev, Claudio Borio, and Kostas Tsatsaronis. The systemic importance of financial institutions. *BIS Quarterly Review*, 75:87, 2009.
- Nikola Tarashev, Claudio Borio, and Kostas Tsatsaronis. Attributing systemic risk to individual institutions. Technical report, Bank for International Settlements, 2010.
- Daniel Tarullo. Lessons from the crisis stress tests. In *Speech at the Federal Reserve Board International Research Forum on Monetary Policy. Washington, DC March*, volume 26, 2010.

- Stefan Thurner. Systemic financial risk: agent based models to understand the leverage cycle on national scales and its consequences. *Multi-disciplinary issues international futures program*. Jan. 14th, 2011.
- Adair Turner et al. *The Turner Review: a regulatory response to the global banking crisis*, volume 7. Financial Services Authority London, 2009.
- Eric Tymoigne. Measuring macroprudential risk: Financial fragility indexes. 2011.
- Frederic Vrans. Double-t copula pricing of structured credit products: practical aspects of a trustworthy implementation. *Journal of Credit Risk*, 5(3):91–109, 2009.
- Michael Walker. The static hedging of cdo tranche correlation risk. *International Journal of Computer Mathematics*, 86(6):940–954, 2009.
- Dezhong Wang, Svetlozar Rachev, and Frank J Fabozzi. Pricing of credit default index swap tranches with one-factor heavy-tailed copula models. *Journal of Empirical Finance*, 16(2):201–215, 2009.
- Søren Willemann. Fitting the cdo correlation skew: a tractable structural jump-diffusion model. *Journal of Credit Risk*, 3(1):63–90, 2007.
- Mark Williams. Uncontrolled risk: Lessons of lehman brothers and how systemic risk can still bring down the world financial system. 2010.
- Fan Yu. Correlated defaults in intensity-based models. *Mathematical Finance*, 17(2):155–173, 2007.