

Imperial College London

Finance

DOCTORATE DISSERTATION IN

September 2021

DATE

Essays on Systemic Risk

THESIS TITLE

Oren Schneorson

Author

Professor Franklin Allen

Primary supervisor

Professor David Miles

Secondary supervisor

Imperial College Business School

Department



D.Phil. in Finance
September 2021

D.Phil. Thesis

Essays on Systemic Risk

D.Phil. Candidate: Oren Schneerson

Advisor: Prof. Franklin Allen
Prof. David Miles

Statement of Originality and Copyrights

I hereby declare that this thesis is my own work, and that all credits to other works, images, charts, or any other material that this thesis relies upon is appropriately referenced and credited. The copyright of this thesis rests with the author. Unless otherwise indicated, its contents are licensed under a Creative Commons Attribution-Non Commercial 4.0 International Licence (CC BY-NC). Under this licence, you may copy and redistribute the material in any medium or format. You may also create and distribute modified versions of the work. This is on the condition that: you credit the author and do not use it, or any derivative works, for a commercial purpose. When reusing or sharing this work, ensure you make the licence terms clear to others by naming the licence and linking to the licence text. Where a work has been adapted, you should indicate that the work has been changed and describe those changes. Please seek permission from the copyright holder for uses of this work that are not included in this licence or permitted under UK Copyright Law.

*To my family without whose unwavering support and dedication to education this would
never have happened.*

Acknowledgements

I would like thank my supervisors Professors Franklin Allen and David Miles for their support, encouragement and direction. Special thanks to Professor Alex Michaelides who acted as my discussant, and to all Imperial College's staff from whom I sought advice. Thanks also to my good friends Dr Patrick Moran, Patrick Schnedier and Dr Can Gao for many hours of discussion, consultation and proof reading. I thank Imperial College that funded my position.

Abstract

The theme of this dissertation is economic stability. In the first chapter I show that interbank connections also have a positive effect on financial stability, via the bank-run channel. This is because they provide risk-sharing benefits to short-term creditors (outsiders), thus incentivizing them to ‘run on the bank’ less frequently. I argue that there is a trade-off between risk-sharing and cascading losses, both of which have an effect on stability. My novelty is the link I make between long-term welfare and financial stability.

In view of the risk from the domino effect, policy makers sought to reduce credit exposures between Systemically Important Financial Institutions (SIFIs). One such policy was an amendment to the U.S. bankruptcy code in 2005-6. which expanded the exclusion of derivative positions from mandatory stay. Since this policy has two opposing effects, the question whether it was beneficial for financial stability is empirical. In the second chapter of this dissertation I ask exactly that question. I describe the challenges to identification, and document what had actually happened to financial stability in the four years period before the financial crisis.

The final chapter deals with the question: how is it that economic outcomes can be highly volatile, while changes in economic fundamentals seem to be rather small? I develop a novel amplification mechanism that is closely related to the existence of multiple equilibria. I apply techniques from the global games literature in a macroeconomic model, showing that an economy with nominal wage rigidity exhibits a knife-edge property with erratic fluctuations around a threshold.

Keywords: Systemic Risk, Banking, Macroeconomics, Applied Microeconomics

JEL Classification: G21, D84, D50

Contents

List of Figures	xiii
List of Tables	xv
Presentation	1
1 Interbank Credit Exposures and Financial Stability	3
1.1 Introduction	3
1.2 Model	8
1.2.1 Physical Environment	8
1.2.2 Agents and Preferences	9
1.2.3 Information Structure	10
1.2.4 Institutional Environment	10
1.3 Equilibrium	12
1.4 Interbank Connectivity	15
1.5 Comparative Statics	19
1.5.1 Calibration: Preferences, Technology and Information	19
1.5.2 Calibration: Banking Sector	20
1.5.3 Calibration: Target	22
1.5.4 Simulation	23
1.5.5 Results	23
1.6 Discussion	25
2 Do Netting Rules Reduce Systemic Risk?	31
2.1 Introduction	31
2.1.1 Numerical Example	36
2.1.2 Related Literature	38
2.2 Methodology	41

2.2.1	Probability of Default	44
2.2.2	Bank Run Index	45
2.2.3	Distance to Default	45
2.3	Data	46
2.3.1	CDS ParSpread and the Risk Neutral Probability of Default	47
2.3.2	Distance to Default	51
2.4	Results	51
2.5	Discussion	54
3	Almost Purely Endogenous	57
3.1	Introduction	58
3.2	Model: Complete Information, Multiple Equilibria	65
3.3	Model: Incomplete Information, Unique Equilibrium	70
3.4	Almost Purely Endogenous	74
3.5	Infinite Horizon OLG	79
3.6	Discussion	82
	Conclusions	85
	References	87
	Appendix A	95
A.1	Proof of Proposition 1	95
A.1.1	Definitions and Preliminary Results	95
A.1.2	Unique Threshold-Equilibrium	98
A.2	Proof of Proposition 2	99
A.2.1	Implicit Function Theorem, Conditions	101
A.3	Technical Appendix	101
A.3.1	Region I	103
A.3.2	Region II	104
A.3.3	Region III	105
A.3.4	Simulation: Graphs	109
	Appendix B	111
B.1	Equity Information	111
B.2	Term Structure of the Risk Free Rate	112
B.3	Matching Procedure	112
B.4	Charts	121

Appendix C	129
C.1 Proof of Lemma 1	129
C.2 Proof of Proposition 3	130
C.3 Proof of Proposition 4	131
C.3.1 Numerical Example	137
C.4 Optimal Asset Prices, Derivation	137
C.5 Relaxing Wage Rigidity	140
C.5.1 Fraction of Workers have Sticky Wages	140
C.5.2 Short-run Sticky Wage	141
C.6 Additional Points for Discussion	143
C.6.1 Nominal Wage Rigidity	143
C.6.2 Neutrality of Money	145
C.6.3 Price Index Indeterminacy	146
C.6.4 Trading Mechanism	147
C.7 Related Literature, Macroeconomics	147

List of Figures

1.1	Global OTC derivative markets	7
1.2	Proportion of early withdrawals $n(\theta)$	14
1.3	State-space partition (regions)	16
1.4	S&P 500 10-year average returns and SPF forecasts	22
1.5	Threshold strategy (θ^*), varying interbank connectivity (B); $\gamma_1 = 0$	28
1.6	Heatmap of optimal interbank connectivity	29
1.7	Threshold strategy (θ^*), varying interbank connectivity (B); $\gamma_2 = 0$	30
2.1	Systemic- β : The Correlation between fundamentals and default risk	34
2.2	Interbank credit exposures and the domino effect	37
2.3	1-year risk neutral PoD (index, North America)	48
2.4	Bank run index (weighted average): North America	50
2.5	Distance to default (index)	52
2.6	First Stage Regressions: All sectors – All entities	53
2.7	Chow tests, fixed window: All sectors	55
2.8	Chow tests, moving window: All sectors	55
3.1	Graphical example	61
3.1	The labour market with nominal wage rigidity, comparative statics	68
3.1	Potential equilibria with negligible variance of $Q(\theta)$	75
A.1	Comparing integration methods: expected differential utility $a\Delta$	109
B.1	T-years risk neutral PoD (indices), by location	121
B.2	1-year PoD, large banks and non-large banks	122
B.3	Bank run index, by location (weighted average)	122
B.4	Bank run index, large banks and non-large banks	123
B.5	Distance to default, by location (index)	123
B.6	Distance to default, large banks and non-large banks	124

B.7	First stage regressions: Systemic- β , Financials vs Non-financials	124
B.8	First stage regressions: Systemic- β , Large banks vs Other entities	124
B.9	First stage regressions: R^2 , Large banks vs Other entities	125
B.10	First stage regressions: R^2 , Financials vs Non-financials	125
B.11	Chow tests, fixed window: Financials	125
B.12	Chow tests, fixed window: Non-Financials	126
B.13	Chow tests, fixed window: Large banks	126
B.14	Chow tests, fixed window: Other than large banks	126
B.15	Chow tests, moving window ($y = 1$): Financials vs Non-financials	127
B.16	Chow tests, moving window ($y = 1$): Large banks vs Other entities	127
C.1	Graphical example, varying η	138

List of Tables

1.1	State-space partition (Regions)	17
1.2	Calibration: Preferences, Technology and Information	20
1.3	Calibration: Banking Sector	21
2.1	Summary statistics	47
2.2	Testing parallel trends (dependent variable: $\hat{\beta}_{jt}$)	54
A.1	State-space partition: boundaries	110
B.1	List of participating companies (excerpt), Markit Ticker to Compustat gvkey	113
B.2	List of large banks	114

Presentation

The overarching theme that ties this dissertation together is economic stability. I began research in this area during my Master's degree, working on contagion in network models. I studied how shocks are transmitted via the financial network, and how governments can effectively arrest this transmission by allocating bailouts to Systemically Important Financial Institutions (SIFIs). My insights from working on network models brought me to my first research idea, which is also the first chapter in this dissertation. In short, I learned that banks with large interbank liabilities transmit a large proportion of their losses to other banks. The flip-side of this is that banks who are not highly connected are less likely to transmit shocks. Rather, shocks are 'absorbed' by 'outside creditors' who do not play an active role in affecting stability. Therefore, from the perspective of network models, less connection between banks is always good for stability.

In the first chapter of this dissertation I show that interbank connections also have a positive effect on financial stability, via the bank-run channel. This is because they provide risk-sharing benefits to short-term creditors (outsiders), thus incentivizing them to 'run on the bank' less frequently. I argue that there is a trade-off between risk-sharing and cascading losses, both of which have an effect on stability. My novelty is the link I make between long-term welfare and financial stability.

Policy makers have long argued that interbank credit risk undermines the stability of the financial sector, because it gives rise to cascading losses (domino-effect contagion). In view of the risk from the domino effect, they sought to reduce credit exposures between Systemically Important Financial Institutions (SIFIs). One such policy was an amendment to the U.S. bankruptcy code in 2005-6. which expanded the exclusion of derivative positions from mandatory stay. The first chapter of this dissertation questions the effectiveness of this policy by showing that, in equilibrium, it may have adverse effects on stability because it reduced the risk sharing benefits of interbank credit exposures. Since this policy has two opposing effects, the question whether it was beneficial for financial stability is empirical. In the second chapter of this dissertation I ask exactly that question. I describe the challenges

to identification, and document what had actually happened to financial stability in the four years period before the financial crisis.

The final chapter of this dissertation deals with a deeper question in economics and finance relating to stability, namely how is it that economic outcomes can be highly volatile, while changes in economic fundamentals seem to be rather small? My starting point is that market structure / institutions could have ‘amplification mechanisms’, through which innovations have increased impact. In this chapter I offer a novel amplification mechanism that is closely related to the existence of multiple equilibria. I apply techniques from the global games literature in a macroeconomic model, showing that an economy with nominal wage rigidity exhibits a knife-edge property with erratic fluctuations around a threshold.

1

Interbank Credit Exposures and Financial Stability

This chapter investigates how interbank credit exposures affect financial stability. Policy makers often see such exposures as undermining stability by exacerbating cascading losses through the financial system. I develop a model that features a trade-off between cascading losses and risk-sharing. In contrast to previous studies I find that reducing interbank connectivity may destabilize the financial system via the bank-run channel. This is because it decreases the risk-sharing benefits of interbank connectivity. A bank-run model features two islands that are connected via a long term debt claim. Varying the size of this claim (interbank connectivity), I study how the decision to ‘run on the bank’ is affected. I run a simulation of the model, calibrated to the U.S. banking system between 1997-2007. I find that large bankruptcy costs are required to trump the risk-sharing benefits of interbank credit exposures.

1.1 Introduction

Should financial institutions be treated more favorably, in the context of a financial crisis, to other companies or individuals? This question has been the topic of active debate in recent years following the events of 2007-8. In practice, regulators the world over seem to answer it with a definite yes: from bail-outs to favorable treatment in bankruptcy, financial institutions are deemed systemically important and have earned an acronym to that effect (SIFI).¹

Policy makers have argued that giving seniority to SIFIs could stabilize the financial system (Mengle, 2010). This is because it reduces interbank credit exposures, insulating the

¹Systemically Important Financial Institutions.

financial sector from cascading losses (domino-effect; Bliss and Kaufmann, 2004; Duffie and Skeel, 2012). However, giving seniority to banks implies that other creditors recover less in bankruptcy. Among those are short-term creditors whose panic may destabilize the financial system (Roe, 2013).² This chapter investigates how interbank credit exposures affect financial stability. It studies the trade-off between domino-effect contagion and risk-sharing as it affects short-term creditors' decision to 'run on the bank'.

Studying this trade-off in equilibrium requires a model in which both the domino-effect and risk-sharing are present, and where non-bank creditors can respond to changes in the financial network. I propose a two-period bank-run model in which banks have two types of creditors: depositors and other banks. Depositors may decide whether to withdraw their deposit in period 1 or wait until long-term investment bears fruit in period 2. If enough of them withdraw in period 1, this decision to 'run on the bank' is individually optimal even though fundamentals could be high enough for it to be socially sub-optimal. Following Goldstein and Pauzner (2005, GP), I assume depositors don't have common knowledge about the economy's fundamentals; instead, each agent observes a private signal about fundamentals.³ This allows me to evaluate existing policy in equilibrium, because it resolves the problem of multiple self-fulfilling equilibria typical of Diamond and Dybvig (1983) bank-run models.

The first contribution of this chapter is the finding that reducing interbank credit exposures may in fact decrease the stability of the financial system. In absence of bankruptcy costs, connecting the banks is always beneficial because it allows depositors to share idiosyncratic risk. This finding contrasts with previous studies on domino-effect contagion, who do not model agents' endogenous response to changes in the financial network (Acemoglu et al., 2015; Gai and Kapadia, 2010).

The second contribution is to generate predictions as to which effect is stronger, depending on the magnitude of bankruptcy costs and the degree of risk-aversion. If bankruptcy costs are positive, the risk-sharing channel conflicts with domino-effect contagion, as the latter reduces long-term value when the financial sector is highly interconnected. I simulate the model, calibrated to the U.S. banking sector between 1997-2007, and find that if bankruptcy costs are below 20%, the optimal policy is to have interbank connectivity in excess of 60% of banks' balance sheets. Conversely, if bankruptcy costs are as high as 50% netting is optimal. These results suggest that bank seniority could be counter-productive if bankruptcy costs are not too high.

²"If derivatives and repo counterparties bear less risk, as they do, due to the Bankruptcy Code's favoritism, then other creditors that are poorly prioritized bear more risk and thus have more incentive for market discipline."

³The noisy signal could also be which interpreted as agents' opinion about the soundness of the economy.

The size of bankruptcy costs is important because it determines the strength of the domino-effect. In the literature on domino-effect contagion bankruptcy costs are often calibrated as high as 50-100%, referencing recovery rates on corporate bonds of companies that entered bankruptcy. However, it is problematic to ascribe low recovery rates to bankruptcy costs. The Lehman Brothers case serves as a good example. There, legal and administration costs amounted to about 2.5% of total claims approved by the courts (roughly \$9 billion). Yet, general creditors recovered only 31% of their claims, which is significantly lower than the banking sector historical average.

This chapter contributes to a number of areas in economics and finance. First, it is related to the literature on domino-effect contagion discussed above. My model is not the first that contains an endogenous response to instability. Erol and Vohra (2018) study endogenous network formation in face to the contagion risk. In their model, stability is a function of agents' choices, but is limited only to the domino-effect. The value of the network is at odds with its stability: the higher is the value of a link (financial or other) – the more agents want to connect – the larger is the impact of the domino-effect, thus decreasing stability. In contrast, the cause of instability in my model is the interplay between exogenous and endogenous risk. Agents' choice of whether to run on the bank is directly affected by long term value: the higher is this value – the more stable is the banking sector. Interbank credit exposures augment this value because they allow agents to share idiosyncratic risk.

It also contributes to the area of banking by putting forth a novel mechanism by which risk-sharing increases financial stability. Other works in this vein modelled risk-sharing as occurring between banks themselves, e.g. via credit lines that smooth liquidity shocks (Allen and Gale, 2000; Ladley, 2013). The mechanism in my model is different to theirs: interbank connectivity is only effective when banks are already insolvent; risk-sharing occurs at the level of depositors, not banks.

Another contribution to the banking literature is the application of the seminal work of GP to the study of public policy. Other works applied GP to study the effect of government guarantees (Allen et al., 2018) and bank heterogeneity in banks' asset holdings (Goldstein et al., 2020) on financial stability. Both are applications using a similar framework but studying separate questions. Liu (2016) studies the joint occurrence of interbank credit freezes and bank-runs. His focus is on liquidity, while I study the interplay between liquidity and solvency. In order to do this I introduced one crucial modification that can be easily overlooked. In the original GP model, all illiquid assets can be used to generate period 1 value by applying the usual discount for early liquidation. In contrast, I allow for liquidation up to a constraint based on the portfolio choice of the bank. This allows me to keep at least some value for period 2, without which there is never any role for bankruptcy proceedings.

In the literature on insolvency, my work is related to Matta and Perotti (2016), who study how risk arising from premature liquidation of illiquid assets can contribute to the probability of a bank-run – and the way in which mandatory⁴ stay helps in dealing with this problem. In contrast, this chapter focuses on priority rules that define the boundaries of the estate upon which mandatory stay is imposed. Bolton and Oehmke (2015) study how seniority rules affect banks' risk-taking behaviour. Their focus is the effect of seniority on investment efficiency and welfare, whereas mine is financial stability.

Seniority of interbank liabilities is achieved in practice by excepting Qualified Financial Contracts (QFCs, mostly derivatives) from mandatory stay. It has the effect of reducing interbank connectivity, as it allows SIFIs to net mutual assets and liabilities with an insolvent bank.⁵ This would not be allowed ordinarily. Figure 1.1 illustrates the trend in market value of OTC derivatives, which grew rapidly in the run-up to the 2007-8 crisis. In December 2008, netting reduced banks' credit risk by a staggering \$30 trillion, about 85% of their global gross market value (Figure 1.1). Netting is one of the two most prominent exceptions to mandatory stay (Wood, 2007).

Insolvency netting was gradually introduced to the U.S. Bankruptcy Code between 1978-2006 in a series of amendments often referred to as 'safe harbors' (Mooney Jr, 2014).⁶ This legislation was put forth "with congressional intent in creating... safe harbors to promote the stability and efficiency of financial markets" (Chapman, 2016). Other countries treat insolvency netting in different ways, deviating from the guidelines set up in the Basel Accords.⁷ In the EU, similar rules exist, exempting QFC-like transactions from stays (Perotti et al., 2010). However, U.S. law is likely to be relevant in most advanced economies, as banks may choose the legal system which governs their Master Netting Agreement.⁸ Duffie and Skeel (2012) discuss the arguments for and against netting, with particular analysis for repurchase agreements (repos) and central clearing counterparties (CCPs). One of the main

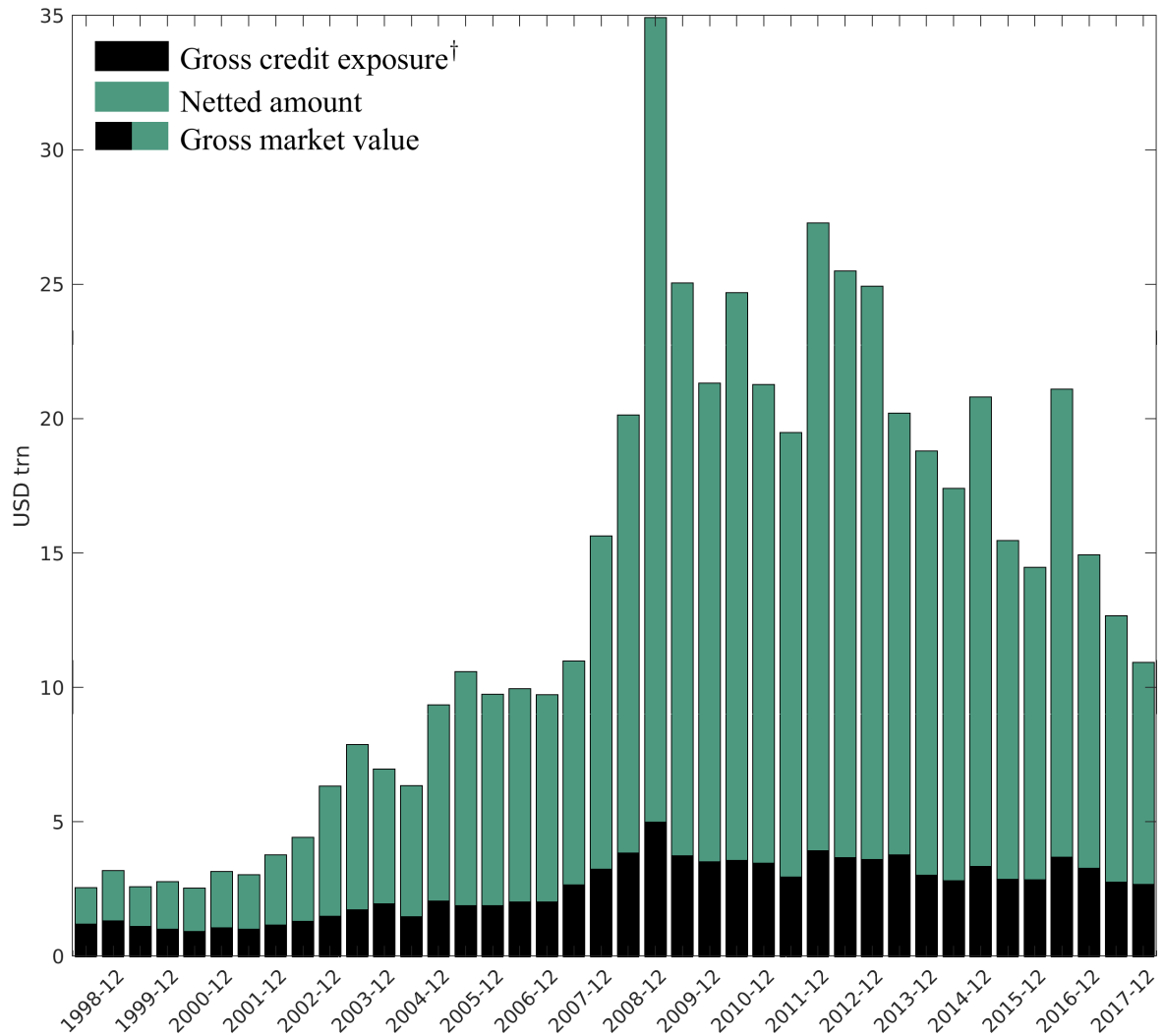
⁴Mandatory Stay: an injunction imposed on the creditors of a bankrupt company which disallows them from independently pursuing the collection of their claims.

⁵The reason it is pertinent particularly to SIFIs is that it requires some form of mutual claims, which typically arise from their role as dealers in OTC derivative markets and as liquidity providers (Bliss and Kaufmann, 2004).

⁶In 1978, safe harbors were first introduced to the Bankruptcy Code in order to enhance commodity market stability. Hence, a series of amendments passed by Congress expanded the type of institutions and contracts that could benefit from the protection of safe harbors, and consequently, netting became more and more robust. In the U.K., netting is an old legal practice, dating back as early as the 18th century, and is not limited to financial contracts.

⁷The Basel Accords set a minimum reporting requirement for two securities to be eligible for set-off, that countries can only make more strict, i.e. less favourable to netting.

⁸From conversations I had with practitioners, they estimated that in the majority of cases the law governing set-off would be either American or English – mainly due to the advantages related to certainty attributed to common law – the difficulty of a judge to set a precedent under a common-law system, in contrast to the ability of a judge to interpret the law in her own way under a code-of-laws system.

Figure 1.1 Global OTC derivative markets

† Gross credit exposure is the amount of credit risk left before collateralization and after netting; according to this number firms post initial margins. Source: Bank of International Settlements.

arguments they raise for excluding QFCs from mandatory stay is that it resolves at least some of the uncertainty around the default of SIFIs. They do not discuss the risk sharing channel that this chapter studies. In relation to repo, their main concern is about a fire-sales of illiquid collateral.

The rest of the chapter is organized as follows. The next section sets up the model and section 1.3 describes the equilibrium. In 1.4 I explicitly analyze how an interbank connection affects the equilibrium and state the main result of the chapter: giving priority to banks may increase the probability of runs. In 1.5 I describe how I calibrate the model, and present simulation results. Section 1.6 concludes with a discussion of the implications of this chapter for policy. All proofs are provided in the appendix.

1.2 Model

This section spells out a bank-run model with idiosyncratic risk based on Goldstein and Pauzner (2005, GP). In essence, bank-run models are a coordination game: there are situations in which if agents' believed that other agents will not run – they would not run themselves. The main problem these models deal with is the existence of a “bad” equilibrium (a bank-run), which is Pareto-dominated – alongside a “good” equilibrium, which is Pareto-dominating. In the absence of equilibrium uniqueness, bank-run models do not provide direction as to which equilibrium is selected, and thus policy analysis is precluded.⁹

GP apply methods from the theory on global games to Diamond and Dybvig's seminal bank-run model. By modelling the structure of the economy's fundamentals and agents' information, their model pins down a unique ex-ante probability of bank-runs. Moreover, one of their forceful points is to show that bank-runs are typically an equilibrium phenomenon even in the second-best case (equilibrium that is constrained by agents' private information about their type). In their model, bank-runs occur even in states of the world where fundamentals are strong enough so that in absence of the coordination problem, there would not be a run on the bank. In this sense bank-runs are a self-fulfilling prophecy.¹⁰

1.2.1 Physical Environment

An economy with two islands, a single consumption good and two assets (short and long) exists for three periods: 0, 1 and 2. For the moment consider only an individual island; the islands are identical ex-ante. Island k is inhabited by a continuum of depositors and a representative bank. All agents have access to a risk free asset (with a return normalized to 0) and a risky asset that takes 2 period to mature; if allowed to mature, it yields θ_k . Total return on the risk-free and risky asset is given by $R(\theta_k)$,

$$R(\theta_k) = x + (1 - x)\theta_k \quad (1.1)$$

⁹Multiplicity of equilibria in coordination games can be thought of as an artefact of extreme and implausible assumptions about common knowledge. These assumptions are intended to simplify analysis, i.e. they are not the result of an underlying reason internal to the logic of the model.

¹⁰The literature on self-fulfilling prophecies and coordination games dates back at least to Aumann (1976) who establishes the notion of common knowledge for game-theoretic models. Models with multiple equilibria were studied in banking (Chari and Jagannathan, 1988), monetary policy (Benhabib and Schmitt-Grohé, 2001) and macroeconomics (Farmer and Benhabib, 1994). Deviations of the common knowledge assumption gave rise to a vast literature that this review cannot hope to span. Notable references are Carlsson and Van Damme (1993); Monderer and Samet (1989); Morris and Shin (1998); Rubinstein (1989).

where x is the investment in the short asset. Returns are a function of a state variable $\theta_k \in \Theta \subseteq \mathbb{R}$ (economic fundamentals); R is continuous and monotonically increasing in θ_k , so high values of the state variable imply good news to investors. The fundamentals in island k are composed of a common factor θ and a mean zero idiosyncratic factor ν_k ,

$$\theta_k = \theta + \nu_k \quad (1.2)$$

$$\theta \sim f_\theta \quad \mathbb{E}[\theta_k] = \mu_\theta \quad (\nu_k, \nu_{-k}) \sim f_\nu \quad \mathbb{E}[\nu_k] = 0$$

Let $\boldsymbol{\theta} = (\theta, \nu_k, \nu_{-k})$ be the state of nature, and its probability density function f_θ . Since the islands are ex-ante identical, f_ν is symmetric: $f_\nu(\nu, \nu') = f_\nu(\nu', \nu)$, which also implies $\text{Var}(\nu_k) = \text{Var}(\nu_{-k})$. I assume that the correlation between ν_k and ν_{-k} is not perfectly positive,

$$\frac{\text{Cov}(\nu_k, \nu_{-k})}{\text{Var}(\nu)} \neq 1 \quad (1.3)$$

A proportion up to x can be liquidated from the total portfolio at no cost to yield value in period 1. In case $x < n \cdot c$, the bank incurs an illiquidity cost of size γ_1 . The scrap value $(1-x)R(1-\gamma_1)$ is only yielded in period 2.¹¹

Our main interest is banks' seniority in insolvency procedures. In order to have a meaningful distinction between insolvency and illiquidity, we must have some value also in period 2. Limiting the amount of the long asset that can be liquidated aides in making this distinction, as it prevents a bank-run from totally ravaging the bank's assets. This is one of the differences between my model and GP.

1.2.2 Agents and Preferences

The economy is populated by a unit continuum of depositors, and one bank per island. There are two types of depositors: impatient (with proportion λ) and patient (with proportion $1 - \lambda$). Impatient depositors consume in period 1 and patient in period 2. Depositors don't know their type in period 0, which they observe only in period 1. Moreover, their type is their private information and is unverifiable, so that patient depositors can always disguise themselves as impatient. Ex-ante expected utility of an agent is given by

$$\mathbb{E}u = \int_{\boldsymbol{\theta}} \left[n q u(c_1) + (1 - n + (n - \lambda)(1 - q)) u(c_2) \right] f_{\boldsymbol{\theta}}(\boldsymbol{\theta}) d\boldsymbol{\theta} \quad (1.4)$$

¹¹The assumption that the portfolio is liquidated proportionally is unreal, but it simplifies the proof of a unique equilibrium.

where θ is the state of nature, $n = n(\theta)$ is the proportion of depositors who withdraw in period 1, $c_t = c_t(\theta)$ is consumption in period t and $q = q(\theta)$ is the probability of consumption in period 1. In case an impatient depositor couldn't get her deposit, she consumes zero. If consumption in period 2 is positive, then the probability of consuming is 1; I therefore omit the probability of consumption in period 2 (and consequently its t subscript). Rather than designating another symbol for the interest rate I used a shorthand to the following notation, omitting the time subscript for consumption in period 1 and simply denoting it as $c_1 = c$. If an agent consumes in period 1, it will always consume a fixed amount, which is the also the gross return on bank deposits. The utility function u is a monotonically increasing ($u' > 0$) and has decreasing marginal gain ($u'' < 0$).

I assume that banks expect zero profits, thus maximizing agents' welfare. Banks face a portfolio allocation problem, investing x in the risk-free asset that yields R_f , and $1 - x$ in the risky asset with expected return of μ_θ . In all, they choose (c, x) in order to maximize depositors' period 0 expected utility (eq. 1.4).

1.2.3 Information Structure

Recall that returns on island k depend on a common factor θ and an idiosyncratic factor ν_k (eq. 1.2). The idiosyncratic factor (ν_k, ν_{-k}) is observed only in period 2. The common factor θ is observed imperfectly already in period 1. Agent i observes her type, and a private noisy signal θ_i

$$\theta_i = \theta + \varepsilon_i \quad (1.5)$$

$$\varepsilon \sim U[-\epsilon, \epsilon] \quad \varepsilon \perp \theta$$

where ε_i is the realization of the noise in agent i 's signal; ε is uniformly distributed with range 2ϵ around the real value of θ . The rest of the details about agents preferences and the physical environment as well as the institutional environment, are all common knowledge. No information is revealed publicly, and for simplicity assume that there is no communication possible between agents.¹²

1.2.4 Institutional Environment

The rationale for a banking sector is risk-sharing of individual liquidity risk – i.e. being patient or impatient. In period 0, agents can deposit their funds at their regional bank, which

¹²Even if it was possible to communicate, agents would not have had proper incentive to tell the truth about their signal, thus an unverifiable signal is enough.

promises a convertible debt contract with gross interest rate c if the deposit is withdrawn in period 1, and otherwise a debt-plus-equity claim on the assets of the bank in period 2.¹³¹⁴

Denote by D the size of the period 2 debt claim¹⁵ and by A_k the external assets of bank k from its own portfolio, (i.e. unrelated to bank $-k$)

$$A_k = (1 - \min[n \cdot c, x])R_k \quad (1.6)$$

In period 0 the bank invests x in the risk-free, and $1 - x$ in the risky asset; period 1 value is capped at x . In period 1 it liquidates proportionally from its total portfolio to pay for first period withdrawals.¹⁶ Impatient depositors always withdraw their deposit in period 1. Patient depositors may decide whether to withdraw in period 1 or 2.

Definition 1 (Liquidity). *A bank is liquid if the total amount of period 1 claims is lower than liquid assets, $n \cdot c \leq x$; otherwise it is deemed illiquid. When a bank is illiquid, the asset is not allowed to mature – i.e. it is liquidated yielding a scrap value of $(1 - \gamma_1)A_k$ in period 2. Denote by $\tilde{\gamma}_1 = \tilde{\gamma}_1(n)$ the value of illiquidity costs depending on the level of withdrawals,*

$$\tilde{\gamma}_1(n) = \begin{cases} 0 & x \geq n \cdot c \\ \gamma_1 & x < n \cdot c \end{cases} \quad (1.7)$$

Moreover, when a bank is liquid period 1 claims are paid with certainty, whereas when it is illiquid period 1 claims are paid on first-come-first serve basis. The probability of consumption in period 1 $q = q(n)$ is

$$q(n) = \begin{cases} 1 & x \geq n \cdot c \\ \frac{x}{n \cdot c} & x < n \cdot c \end{cases} \quad (1.8)$$

Definition 2 (Solvency). *Bank k is solvent if the total level of assets in period 2 exceeds the total level of debt claims, $A_k(1 - \tilde{\gamma}) \geq D$; otherwise it is deemed insolvent, and all general*

¹³Observe that in absence of any other period 2 debt claims, the debt-plus-equity claim is equivalent to a simple equity claim. Demandable debt is a standard assumption in the banking literature for at least two reasons: (1) it is the optimal contract to overcome moral hazard (Calomiris and Kahn, 1991); (2) it is a constrained-optimal risk-sharing contract between agents in period 0 facing asymmetric information in period 1: “banks can be viewed as providing insurance that allows agents to consume when they need to most. Our simple model shows that asymmetric information lies at the root of liquidity demand.” (Diamond and Dybvig, 1983)

¹⁴Convertible debt is a risk-sharing mechanism. It is not optimal ex-ante because it fixes the interest rate where market clearing would have it vary. However, it does address the asymmetric information problem, thus allowing agents to share liquidity risk (Diamond and Dybvig, 1983). If agents could coordinate on the good equilibrium (no run), then it is the constrained optimal contract.

¹⁵For the sake of brevity, I abstract from optimal capital structure.

¹⁶This assumption ensures that when agents run on the bank, returns in period 2 are lower.

creditors are paid pro-rata. If a bank is insolvent in period 2, its estate is liquidated yielding a scrap value of $(1 - \gamma_2)(1 - \tilde{\gamma}_1)A_k$. Denote by $\tilde{\gamma}_2 = \tilde{\gamma}_2(\boldsymbol{\theta})$ the value of illiquidity costs depending on the level of withdrawals,

$$\tilde{\gamma}_2(\boldsymbol{\theta}) = \begin{cases} 0 & A_k(1 - \tilde{\gamma}_1) \geq D \\ \gamma_2 & A_k(1 - \tilde{\gamma}_1) < D \end{cases} \quad (1.9)$$

Illiquidity costs γ_1 and bankruptcy costs γ_2 affect period 2 value. The distinction between them is crucial, since illiquidity costs do not contribute to the domino effect. This is because they are not *caused* by the insolvency of another bank.

1.3 Equilibrium

In this section I describe the equilibrium of the model. In order to focus on this chapter's contribution, I analyze the equilibrium for a given a choice of (x, c) . The first best (FB) resource allocation can be achieved when there is no asymmetric information, and is given by:

$$\mathbb{E}_\theta \left[\frac{u'(1 - \lambda c^{FB})}{u'(c^{FB})} R(\theta) \right] = 1 \quad (1.10)$$

This allocation is only available if agents' types are not private information; otherwise bank-runs (patient depositors withdrawing early) may disrupt this optimality. Bank-runs are a result of a failure of patient agents to coordinate on the 'good equilibrium'.

Definition 3 (Bank-runs). *A 'good equilibrium' obtains if all patient depositors withdraw late whenever fundamentals are high enough to support the optimality of withdrawing late. Bank-runs occur in states in which, although fundamentals are high enough to support the good equilibrium, nevertheless there is at least some early withdrawal by patient depositors and a Pareto sub-optimal equilibrium obtains.*

Impatient agents always withdraw their funds in period 1. However, patient depositors need to choose whether to withdraw in period 1 or wait until period 2. Their decision is based upon all available information. Crucially, agents use their signal in two ways: (a) to infer information about exogenous variables: returns R_k ; (b) to infer information about endogenous variables: the proportion of agents who withdraw early, n .

Let $v = v(\theta)$ be the patient agent differential utility between withdrawing late and early at state θ . The relevant welfare evaluation made by a patient agent is captured by $\Delta = \mathbb{E}[v | \theta_i]$, the expected utility difference between withdrawing late to withdrawing early.

Definition 4 (Patient agents' utility differential). *Let $\Delta(\cdot)$ be patient agents' utility differential between withdrawing early to withdrawing late:*

$$\Delta(\theta_i) = \iiint_{\theta} v(\theta) f_{\theta}(\theta | \theta_i) d\nu_k d\nu_{-k} d\theta \quad (1.11)$$

$$\begin{aligned} v(\theta) &= u(c_2) - [qu(c_1) + (1 - q)u(c_2)] \\ &= q[u(c_2) - u(c_1)] \end{aligned} \quad (1.12)$$

where $q = q(\theta) = \min[\frac{x}{n \cdot c}, 1]$ (see eq. 1.8)

$\Delta(\theta_i)$ indicates the expected value of v , conditional on observing signal θ_i . v depends on the proportion of withdrawals n , the short-term interest rate c , liquid assets x , returns R and the capital structure in period 2. In keeping with the global games method, I assume there exist some values of fundamentals in which patient agents' binary decision (withdraw early or late) does not depend on other agents' actions:

Assumption 1 (Extreme locales of fundamentals). *Two extreme locales of fundamentals are assumed to exist in which patient agents' actions are independent of their beliefs concerning other agents' actions.*

a. *Lower dominance locale: $\theta \in [0, \underline{\theta}(c, x)]$, where $\underline{\theta}$ is defined by the relation*

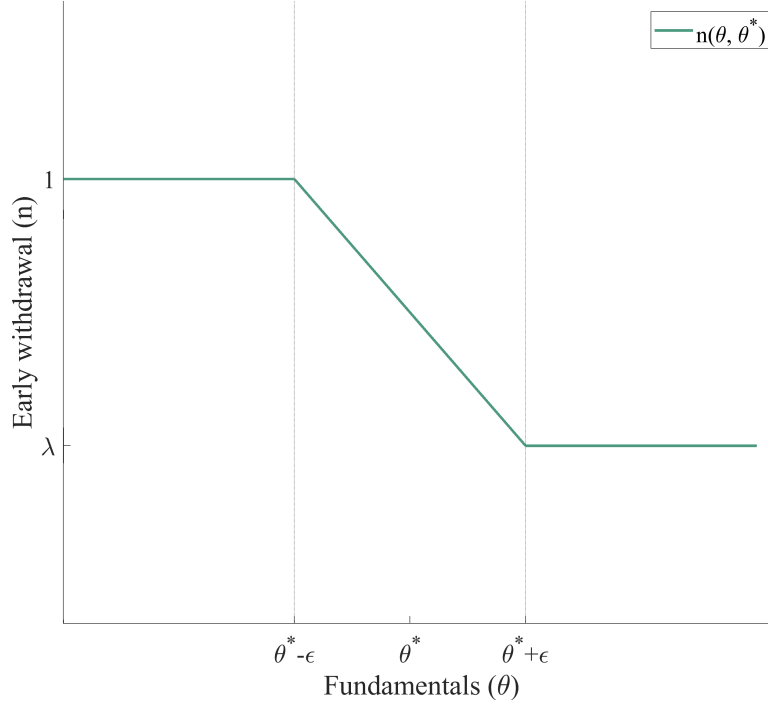
$$c = \frac{1 - \lambda c}{1 - \lambda} R(\underline{\theta}) \quad (1.13)$$

b. *Upper dominance locale: $\theta \in [\bar{\theta}(c, x), 1]$, where*

$$\frac{(1 - x)R(\bar{\theta})(1 - \gamma_1)(1 - \gamma_2)}{1 - \lambda} \geq c \quad (1.14)$$

Furthermore, $\bar{\theta}$ is assumed to satisfy $\bar{\theta} < 1 - 2\epsilon$.

The interpretation of 1 is as follows: a. if fundamentals are low enough so that the consumption in period 1 is equated to that of period 2 even in the best of states (no run) – then a patient depositor's strictly dominant strategy is to withdraw early regardless of other agents' actions; b. if fundamentals are high enough so that consumption in period 2 is higher

Figure 1.2 Proportion of early withdrawals $n(\theta)$ 

than in period 1 even in the worst of states – then a patient depositor’s strategy is to withdraw late regardless of other agents’ actions.

Let $s(\theta_i)$ be patient depositor i ’s (mixed) strategy – the probability with which she withdraws early. It is equal to 1 (0) whenever $\Delta(\theta_i, \cdot) < 0$ (> 0). If in equilibrium s^* is increasing from 0 to 1 at some θ^* and is 1 for $\theta_i < \theta^*$ and 0 for $\theta_i > \theta^*$, I call this a *threshold equilibrium* (θ^* is a threshold value of the signal below which she withdraws early even if she is patient).

Proposition 1 (Unique threshold-equilibrium). *There exists a unique equilibrium. Patient agents withdraw early whenever they observe a signal $\theta_i < \theta^*$, and withdraw late otherwise.*

All proofs are provided in the appendix. In a threshold equilibrium the proportion of depositors withdrawing in period 1 in each island n :

$$n(\theta, \theta^*) = \begin{cases} \lambda & \text{if } \theta \in [\theta^* + \epsilon, 1] \\ \lambda + (1 - \lambda) \left(\frac{1}{2} + \frac{\theta^* - \theta}{2\epsilon} \right) & \text{if } \theta \in [\theta^* - \epsilon, \theta^* + \epsilon] \\ 1 & \text{if } \theta \in [0, \theta^* - \epsilon] \end{cases} \quad (1.15)$$

Note that since $\varepsilon_i \sim U(-\epsilon, \epsilon)$, then all signals are within the bounds of $\pm\epsilon$ from the common factor θ , and in a given threshold-equilibrium (θ, θ^*) : $n(\theta, \theta^*)$ is degenerate. For values of $\theta \geq \theta^* + \epsilon$ no patient agent will withdraw early, while for $\theta \leq \theta^* - \epsilon$ all patient agents

withdraw early. In between these two values the proportion of depositors withdrawing early is linearly decreasing in θ with slope $\frac{1-\lambda}{2\epsilon}$ (see Figure 1.2).¹⁷

There are three main differences between my model and that of GP: an additional island, asset returns and bankruptcy costs. First, in GP's model there is a single island, which precludes any heterogeneity in period 2. Second, GP have a fixed return with varying probability, while in my model returns vary with θ_k , but are certain for a given θ_k . Third, my model has bankruptcy/illiquidity costs whereas GP's model doesn't. In both models there are no costs to liquidating the long asset up to a proportion x , at which point first period claims are paid on a first-come-first-served basis. However, in GP $x = 1$ so there are no costs at all. Conversely, in my model liquidation of the long asset is cost-free up to a proportion $x \leq 1$. Not all assets are liquidated to satisfy first period withdrawals, which means there is always some value left to agents who did not run or didn't manage to withdraw their deposit in time.

This concludes the setup of the model. In sum, I spelled out a variation of an otherwise standard bank-run model, originally put forward by Diamond and Dybvig (1983). Crises are a result of depositors inability to coordinate on the good equilibrium. Equilibrium multiplicity is resolved by relaxing the assumption of common knowledge about economic fundamentals (and therefore other agents' behaviour). Depositors observe noisy signals about fundamentals, and use a threshold strategy in equilibrium – as in Goldstein and Pauzner (GP).

1.4 Interbank Connectivity

Section 1.2 provided the general setup of the model, with the aim of giving rise to a framework in which banks have two types creditors: depositors and banks. In the last section I described the equilibrium. In this section I consider an interconnected financial system: banks who have mutual assets and liabilities.

Banks have a debt claim of size B on one another, which matures in period 2. Recall that the size of non-bank liabilities is D , and denote by $\psi = \frac{B}{B+D}$ banks' recovery rate. If at least one bank is insolvent, the interbank connection implies a transfer from the strong to the weak bank. Let $\tilde{A}_k$ be the value of bank k 's external assets in period 2 including bankruptcy costs

$$\tilde{A}_k = A_k \times \Gamma \tag{1.16}$$

$$\Gamma = \overbrace{(1 - \tilde{\gamma}_1)}^{\text{Illiquidity costs}} \times \overbrace{(1 - \tilde{\gamma}_2)}^{\text{Bankruptcy costs}}$$

¹⁷Linearity is due to uniform distribution of the signal.

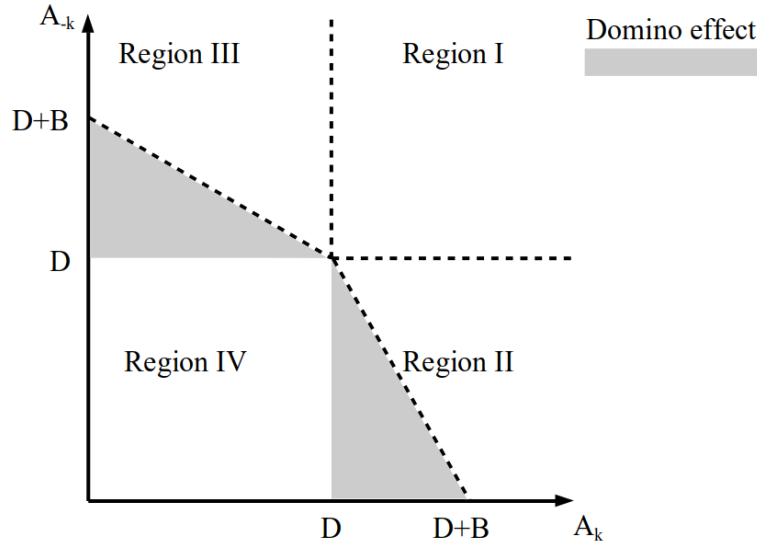


Figure 1.3 State-space partition (regions)

where the definition of $\tilde{\gamma}_2$ is altered in order to reflect the interbank connection

$$\tilde{\gamma}_2(\theta) = \begin{cases} 0 & A_k(1 - \tilde{\gamma}_1) + \min[\psi(\tilde{A}_{-k} + B), B] \geq D + B \\ \gamma_2 & (1 - \tilde{\gamma}_1) + \min[\psi(\tilde{A}_{-k} + B), B] < D + B \end{cases} \quad (1.17)$$

Connecting the banks with mutual assets and liabilities directly alters consumption in period 2 in two ways: on the one hand it effects a payment from the island with high fundamentals to the one with low fundamentals if at least one of them is bankrupt (risk sharing); on the other hand it transmits bankruptcy/illiquidity costs between islands (domino effect contagion).¹⁸

There are four possible cases, which I call regions, depending on which bank is solvent or not: (region I) both banks solvent; (region II) k solvent, $-k$ insolvent; (region III) k insolvent, $-k$ solvent; (region IV) both banks insolvent. Table 1.1 summarizes all possibilities, and provides the conditions for being in each region depending on (A_k, A_{-k}) . Figure 1.3 depicts those in 2D space; the horizontal (vertical) axis is the realization of period 2 assets of bank k ($-k$). Region IV, where both banks are insolvent, includes the southwestern quadrant – plus two congruent triangles (gray) due to the domino effect.

Equations 1.18-1.21 state consumption in period 2 in island k for each possibility. If both banks are solvent (region I), consumption is the same as it would have been in absence of

¹⁸By setting agents noisy signals around the common factor, I preclude island heterogeneity in period 1, thus illiquidity costs will play no role.

Region	Solvency		Condition
	Bank k	Bank $-k$	
I	Yes	Yes	$A_j(1 - \tilde{\gamma}_1) \geq D \forall j \in \{k, -k\}$
II	Yes	No	$A_{-k}(1 - \tilde{\gamma}_1) < D \wedge$ $A_k(1 - \tilde{\gamma}_1) \geq D + B(1 - \psi) - \psi \tilde{A}_{-k}$
III	No	Yes	$A_k(1 - \tilde{\gamma}_1) < D \wedge$ $\tilde{A}_{-k} \geq D + B(1 - \psi) - \psi \tilde{A}_k$
IVa	No	No	$A_j(1 - \tilde{\gamma}_1) < D \forall j \in \{k, -k\}$
IVb	No*	No	$A_{-k}(1 - \tilde{\gamma}_1) < D \wedge$ $D \leq A_k(1 - \tilde{\gamma}_1) < D + B(1 - \psi) - \psi \tilde{A}_{-k}$
IVc	No	No*	$A_k(1 - \tilde{\gamma}_1) < D \wedge$ $D \leq \tilde{A}_{-k} < D + B(1 - \psi) - \psi \tilde{A}_k$

Table 1.1 State-space partition (Regions)

interbank connection. Consumption in island k depends only the fundamentals in of bank k ,

$$c_2(\theta | \theta \in \Theta_1) = \frac{\tilde{A}_k}{1 - \min[n, x/c]} \quad (1.18)$$

In case bank k is solvent and bank $-k$ is not (region II), consumption in island k depends on the level of fundamentals for both banks,

$$c_2(\theta | \theta \in \Theta_2) = \frac{\tilde{A}_k - B + \psi[B + \tilde{A}_{-k}]}{1 - \min[n, x/c]} \quad (1.19)$$

If bank $-k$ is solvent but bank k is not (region III), consumption in island k is higher than it would absent of an interbank connection, but does not depend on the fundamentals of $-k$.

Depositors in island k consume a proportion $1 - \psi$ of total assets available to them,

$$c_2(\theta|\theta \in \Theta_3) = \frac{(1 - \psi)(\tilde{A}_k + B)}{1 - \min[n, x/c]} \quad (1.20)$$

Finally, when both banks are insolvent (region IV), depositors in island k consume a proportion $\frac{1}{\psi}$ from bank k 's external assets, and $\frac{\psi}{1+\psi}$ from bank $-k$'s external assets,

$$c_2(\theta|\theta \in \Theta_4) = \frac{\frac{1}{1+\psi}\tilde{A}_k + \frac{\psi}{1+\psi}\tilde{A}_{-k}}{1 - \min[n, x/c]} \quad (1.21)$$

Connecting the banks changes expected utility from consuming in period 2. As before, utility in island k is (strictly) increasing in the fundamentals of bank k , though this time it is also (weakly) increasing in the fundamentals of bank $-k$. Bank interconnectivity may affect the probability of consumption in period 1 only via θ^* , i.e. through agent endogenous choice of whether to run on the bank. Otherwise, it does not affect period 1 value.

Corollary 1. *There exists a unique threshold equilibrium even if banks are interconnected.*

What effect might varying B have on the incentive of patient agents to run on the bank? On the one hand, we can see that bank inter-connectivity has an insurance effect. To the extent that agents are risk averse, they would prefer a more equal consumption across islands. On the other hand, this insurance may come at a cost. By pushing both banks over the cliff-edge of bankruptcy costs instead of just one, connecting the islands may destroy value in some states (the domino-effect).

Proposition 2 (Insolvency netting). *If $\gamma_2 = 0$ (no bankruptcy costs due to insolvency), optimal interbank connectivity implies merging the banks ($B \rightarrow \infty$). If $\gamma_2 > 0$ the effect of interbank connectivity on stability is ambiguous*

The reason interbank connectivity unambiguously decreases the probability of bank-runs if $\gamma_2 = 0$, is that bankruptcy costs are incurred only via period 1 illiquidity. The domino effect is absent due to the assumption that agents observe a noisy signal of the *common factor* θ . In this case connecting the banks can only increase expected welfare in period 2 due to the insurance effect. This could be seen in Fig. 1.3, since the (period 1) expectation of (A_k, A_{-k}) lies on the 45 degree line, with a symmetric pdf. If $\gamma_2 > 0$ there is a trade-off between insurance and bankruptcy costs. Insolvency netting has the effect of reducing interbank connectivity. In no case does insolvency netting increase financial stability unambiguously, as is suggested by the networks literature.

1.5 Comparative Statics

In the previous Section I demonstrated how higher interbank connectivity might enhance the stability of the financial sector via the bank-run channel. Proposition 2 shows that in absence of bankruptcy costs in period 2 ($\gamma_2 = 0$), interbank connectivity is always beneficial to stability due to its insurance effect. If, however, $\gamma_2 > 0$ then interbank connectivity may also undermine stability due to the domino effect. The prediction of the model is thus ambiguous. In this section I present simulation results for a set of calibrated parameters, with the aim of resolving this ambiguity in prediction. I shall demonstrate that for a range of ‘plausible’ parameters, one may obtain either result depending on the social cost of bankruptcy and the magnitude of risk aversion.

1.5.1 Calibration: Preferences, Technology and Information

I make assumptions about the utility, technology and the distribution of θ and ν_k . The utility function exhibits constant relative risk aversion:

$$u(c_t) = 1 - \exp(-\alpha c_t)$$

where the CARA parameter α ranges between 0.1 to 4 (benchmark case: $\alpha = 1$). The values of λ are chosen so to be consistent with observed levels of deposit rates and liquid asset holdings, $x > \lambda c$. Returns on total assets are linear in θ :

$$R_k = R(\theta_k, x) = xR_f + (1-x)\theta_k \quad \theta_k = \theta + \nu_k$$

where R_f is the gross risk free rate and θ_k the return on the long asset. I calibrate R_f using 3-months T-Bill rates which averaged 3.62% p.a. between 1997-2007. The common component in returns is independent from the idiosyncratic components, who are also themselves independent and distributed normally.

$$\begin{pmatrix} \nu_k \\ \nu_{-k} \end{pmatrix} = N \left(\mathbf{0}, \sigma_\nu^2 \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \right)$$

I calibrate σ_ν using the volatility of 3-months cumulative returns on the S&P 500 (roughly 7.4%, or 16% annually for iid processes). Deviation from common knowledge (ϵ) is matched to half the interquantile range of SPF¹⁹ forecasts, 1.5% per annum, or about 0.37% over three months.

¹⁹Survey of Professional Forecasters, see further details below.

Parameter Reference	Parameter	Value(s)	Data Reference	Value(s)
CARA	α	1.0		
Impatient depositors	λ	0.05		
Risk free rate	R_f	1.018	3-months TBill	0.88%
Return S.D. (long asset)**	σ_ν	0.074	S.D. of S&P 500 3-months cumulative return	7.4%
Signal noise	ϵ	0.0037	SPF interquantile range	0.74%

Table 1.2 Calibration: Preferences, Technology and Information

1.5.2 Calibration: Banking Sector

Short-term deposit rates are calibrated based on financial-sector commercial paper rates with maturity 1-3 months (source: Federal Reserve H.15 Selected Interest Rates).²⁰ Liquid assets x is calibrated based on the ratio of high quality liquid assets (HQLA) to total assets. For a sample of large U.S. bank holding companies in mid-2007, values ranged between 3% to 9% (Yankov et al., 2020). Non-bank debt claims D are calibrated to match the leverage ratio for a sample of the largest U.S. banks between 1997-2007: Leverage ratio = $1/(1 - D) = 25$.²¹ Realistic values for interbank mutual claims can be gauged via data on derivative assets and liabilities of U.S. bank holding companies from the Fed (Annual Report of Holding Companies - FR Y-6). Gross market values vary between 30%-150% of banks' balance sheets.²²

Finally, we need to input values for illiquidity and bankruptcy costs (γ_1 and γ_2 respectively). Since in my model γ_1 doesn't contribute to the domino effect, I shall set it to zero. All costs will be incurred via γ_2 , so that the domino effect may have its largest efficacy.

²⁰In order to match the model frequency, the annual figure needs adjustment to a 3-months frequency, $(1 + \text{annualized rate})^{1/4}$.

²¹A higher value of D strengthens the domino effect, since it implies that bankruptcy costs are borne more often.

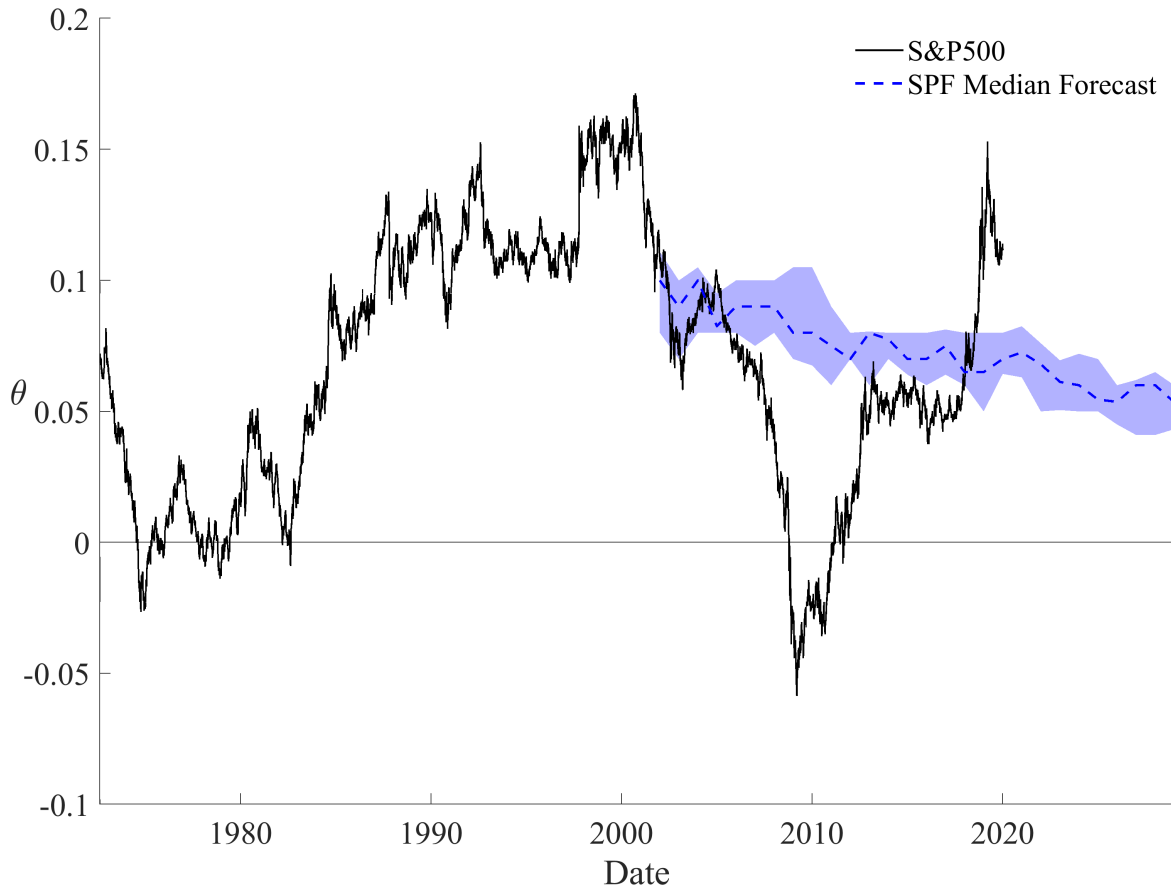
²²Note that derivative trading assets are presented by large banks as an off-balance sheet item, which is why it can be higher than the balance sheet.

Parameter Reference	Parameter	Value(s)	Data Reference	Value(s)
Short-term bank rate (short asset; net)	c	1.01	Commercial paper – financials (3-months)	0.96%
Liquid assets	x	0.06	HQLA	3% - 9%
Non-bank debt claims	D	0.96	Leverage ratio	25
Interbank mutual claims	B	0 - 1.5	Gross market value of derivatives	30-150%
Illiquidity costs	γ_1	0		
Bankruptcy costs	γ_2	0-0.015		

Table 1.3 Calibration: Banking Sector

γ_2 determines whether interbank connectivity enhances stability or undermines it, since it controls the magnitude of the domino effect. Due to the sensitivity of the results to this parameter, I use it in comparative statics rather than attempt to provide a point estimate.

In order to map the bankruptcy costs parameter γ_2 into the economic magnitude of costs in practice, one also has to take account of the probability of default. The unconditional probability of default is not matched well by the model: if the unconditional expectation of quarterly returns is 2.4% and the standard deviation 7.4%, a leverage ratio of 25 imply that banks in the model default around 3% of the time. In reality, investment grade corporate bonds default about 0.5% of the time. This then means γ_2 has an increased impact compared to its real-world counterpart, implying a conservative approach to estimating optimal interbank exposures.

Figure 1.4 S&P 500 10-year average returns and SPF forecasts

1.5.3 Calibration: Target

The main output of the model is the threshold level of expected returns at which bank-runs occur. This is also the variable my calibration aims to match.²³ It is important to establish a benchmark probability of bank-runs that the model would target, i.e. a level of θ^* that would be in accord with historical experience.

Due to the rarity of bank-runs, I refrain from matching a hard target. Rather, I will match a wide area below expected returns, using the median forecast of average 10-year returns on the S&P 500 Index (available from the Survey of Professional Forecasters, SPF).²⁴ This proxies for investors' expectations about the value of long term investment. The steep decline in the S&P 500 during the 2007-8 financial crisis shows that investors were concerned about the long term prospects of the economy as a whole. Equity markets also provide a key public

²³In order to deduce from this output a frequency at which bank-runs occur, one would have to make a stand on the unconditional distribution of agents' expectations.

²⁴The Federal Reserve Bank of Philadelphia surveys a panel of professional forecasters on variables of interest on a quarterly basis. Since 1992-Q1, they began collecting forecasts of average 10-year returns on the S&P 500, collected on an annual basis. I use the median forecast 2008-Q1 as a target θ^* for the model to match.

signal by which investors measure ‘fundamentals’ in practice. Figure 1.4 plots the realized 10-year average returns on the S&P 500 between 1972-2020, along with the SPF median forecast (shifted by 10 years to match realized returns). Panelists’ expectations are much less volatile than realized returns.²⁵ Between 1992-2020, they dropped 450 b.p. from 10% to 5.5% per annum. I define a benchmark case where the time between periods is 3 months, thus I target threshold strategies that are lower than $(1 + .1)^{1/4} - 1 \approx 2.5\%$.

1.5.4 Simulation

In this subsection I outline my simulation methodology. I rely on numerical techniques to get calibration results. Since the model has no closed form solutions, it requires a fixed-point algorithm to locate θ^* , the threshold signal below which patient agents run on the bank. Locating this threshold strategy involves evaluating patient agents’ expected differential utility Δ (equation 1.11).

The model is set up under minimal assumptions about utility and technology. In principle one could work with any odd function for these objects, so long as it satisfies generic conditions described in section 1.2. Simply evaluate the triple integral for a candidate equilibrium θ_0^* ,

$$\int_{\theta=\theta_0^*-\epsilon}^{\theta_0^*+\epsilon} \int_{\theta_k} \int_{\theta_{-k}} v(\cdot) f(\theta_k, \theta_{-k} | \theta) d\theta_k d\theta_{-k} d\theta$$

insert the result to the fixed point algorithm that will suggest the next candidate θ_1^* ; then repeat. The process is concluded after T steps when $\Delta(\theta_T^*) \approx 0$.

It could be computationally costly to evaluate Δ in this way, especially due to multiple uncertainty sources (idiosyncratic risk plus uncertainty about the common component). The combination of CARA utility and Normal Distribution addresses this issue, because it yields partially closed form solutions.²⁶ Further details are provided in the technical appendix.

1.5.5 Results

Figure 1.5 graphs the equilibrium threshold strategy θ^* as it varies with interbank connectivity (B) when bankruptcy costs are: (a) zero; (b) small (0.55% for 6 months, $\gamma_2 = 0.0055 - 10\%$ over 10 years); (c) intermediate (0.9% for 6 month, $\gamma_2 = 0.009 - 17\%$ over 10 years); (d)

²⁵Because of this I opt to match survey data. Asset prices are too volatile to be a good estimate of investors’ long-term growth, and would result in an difficulty to generate a realistic threshold strategy.

²⁶Essentially reducing the triple integral to components of a double / single integrals. This speeds up each evaluation from about 18 seconds to 1 second. Figure A.1 demonstrates that the two algorithms’ outputs are approximately the same.

large (1.45% for 6 month, $\gamma_2 = 0.0145 - 26\%$ over 10 years). In line with Proposition 2, in absence of bankruptcy costs (Panel A), agents run on the bank less often (θ^* lower) when the banks are more highly interconnected.

For example, when the banks are not connected, agents run on the bank when they expect long-term returns to be 4.34% per annum; when mutual claims are 150% of the balance sheet size, this number drops by 17 b.p. to 4.17%. Compare this to intermediate bankruptcy costs in Panel C. When interbank connectivity is 30% of the balance sheet size, agents run on the bank when they expect long-term returns to be 5.26% per annum; when mutual claims are 150%, this number increases by 3 b.p. to 5.29%. In this case the optimal policy is for an intermediate value of interbank connectivity.

We can see that bankruptcy costs have two effects on the probability of bank-runs. First, all else equal, higher bankruptcy costs mean agents run on the bank more often. This makes sense because if agents incur higher bankruptcy costs, a higher reward is necessary to dissuade them from running on the bank. Second, bankruptcy costs strengthen the domino effect, so that higher interbank connectivity is less beneficial. The former increases the level of the curves in Figure Fig. 1.5; the latter shifts their minima to the left.

The heatmap in Figure 1.6 graphs the optimal level of interbank connectivity, varying risk aversion and bankruptcy costs. A coordinate on this chart is essentially the interbank connectivity (B) which attains the minimum threshold θ^* for a given level of bankruptcy costs (γ_2 , horizontal axis) and risk aversion (α , vertical axis). The more green is the color, the higher is the optimal level of interbank connectivity. The latter varies between 0-150% of banks' balance sheets. On the one hand, the higher is risk aversion the higher is the optimal interbank connectivity. This makes sense because higher risk aversion implies high risk sharing benefits from interbank connectivity. On the other hand, higher bankruptcy costs imply a lower optimal interbank connectivity. Higher costs due to bankruptcy imply a stronger domino-effect. In many of the cases, netting is not the optimal policy.

It is important to establish the economic magnitude of these effects on financial stability. A change of 3-17 b.p. in θ^* may seem small at first glance, especially considering the high volatility of equity markets. However, although it is true that long-term growth varies widely, the expectation of it varies much less (see Figure 1.4). The magnitude of changes should be compared to the latter. The interquantile range of forecasts of long term growth is 300 basis points.

Overall the model does a good job at matching a reasonable level of the threshold strategy. Values of θ^* ranging between 4.2%-6% p.a. (when $\gamma_1 = 0$) imply a spread between the threshold strategy and the deposit rate of 30-300 b.p., depending on the level of γ_2 and risk aversion. The size of this spread can be increased by assuming higher values of illiquidity

costs. For example, assuming illiquidity costs of $\gamma_1 = 0.0155$ and no bankruptcy costs $\gamma_2 = 0$ implies roughly 300 b.p. increase in θ^* (see Figure 1.7). The effect of connecting the banks agrees with the model's predictions: in absence of bankruptcy costs, the optimal policy is to completely merge the banks – so as to achieve maximum risk sharing. The optimal level of interbank exposures decreases as we increase bankruptcy costs. Illiquidity costs γ_1 have no impact on the domino-effect because banks are only heterogeneous in period 2; γ_1 would contribute toward the domino-effect if banks were heterogeneous in period 1.

1.6 Discussion

This chapter analyzes the effects of interbank connectivity on financial stability, with reference to a concrete policy – insolvency netting. Netting reduces the size of interbank connectivity by allowing banks to settle mutual assets/liabilities on a net basis. The rationale for this is to stabilize the financial system by preventing the spillover of bankruptcy costs in the event of a crisis (domino-effect contagion). However, this is done by concentrating the losses at a particular group of creditors, undoing the risk-sharing effect that an interconnected financial system provides.

To study the effect of interbank credit exposures in equilibrium, I propose a two period bank-run model with bank heterogeneity in period 2. Bank inter-connectivity has two opposing effects on stability: a domino-effect and a risk-sharing effect. In absence of bankruptcy costs ($\gamma_2 = 0$), I find that insolvency netting *decreases stability*; in case there are costs also due to insolvency ($\gamma_2 > 0$), the effect of netting is ambiguous.

These results stand in stark contrast to previous research on contagion in financial markets. Most of the work in this field neglect to take account of the endogenous response of rational agents to the limiting of financial inter-connectivity. The welfare effect from limiting inter-connectivity has no impact on financial stability in those models. In this chapter, the financial network serves both as insurance to participants in financial markets, and as a shock transmission mechanism. Financial stability is achieved by maximizing investors long term (period 2) value. Whether insolvency netting does this is questionable, as it brings about two opposing effects.

In order to decide which of the two effects is stronger, one must compare the premium that a patient agent would be willing to pay for such insurance – and then compare it with the expected destruction of value due to domino-effect contagion. Recovery rates in financial sector bankruptcies could be as low as 50% (Denison et al., 2019), suggesting rather high bankruptcy costs. Such bankruptcy costs are likely to trump any insurance benefits that an interconnected financial network provides.

It is unclear, however, to what extent low recovery rates should be attributed to bankruptcy costs rather than low fundamentals. A bankrupt company may have made bad investment decisions, and as a result the value of its assets is marked down when this information reaches the market. Consider the case of Lehman Brother for example. On the one hand, general creditors recovered only about 31% of their claims,²⁷ a low number even compared to historical experience. On the other hand, only about 2.5% (roughly \$9 billion) are estimated to be the expenses due to the liquidation process (Denison et al., 2019).

It took more than a decade to liquidate Lehman's estate. It seems unlikely that these losses are the result of a prolonged fire-sale. They could in part be due to the loss of Lehman's franchise value, and/or difficulty of asset valuation in a market with severe asymmetric information. But even the latter is not a pure bankruptcy loss, insofar as the loans underlying many of Lehman's structured finance suffered only modest losses. What matters for systemic risk is the social cost of bankruptcy, not simple transfers. The debate on whether these costs should be chiefly attributed to the bankruptcy event (as in Ball, 2018) or to bad investment decisions (hidden from public view by Lehman's accounting gimmicks, Valukas, 2010) is still ongoing.

The questions raised in this chapter are relevant for policy makers around the world. Most countries set up priority rules for the financial system based on incomplete view of the consequence of this legislation. Two instances of this are the introduction of Master Netting agreements in 2005, and the development of central counterparty clearing houses (CCPs). Both of these had the effect of increasing the eligibility of contracts to insolvency netting, and by extension to reducing credit exposures between SIFIs. This chapter challenges the view that prompted recent changes in U.S. legislation – that distributional effects of netting have no impact on systemic stability. American law is relevant to other advanced economies, as SIFIs choose the jurisdiction which governs their master netting agreements.

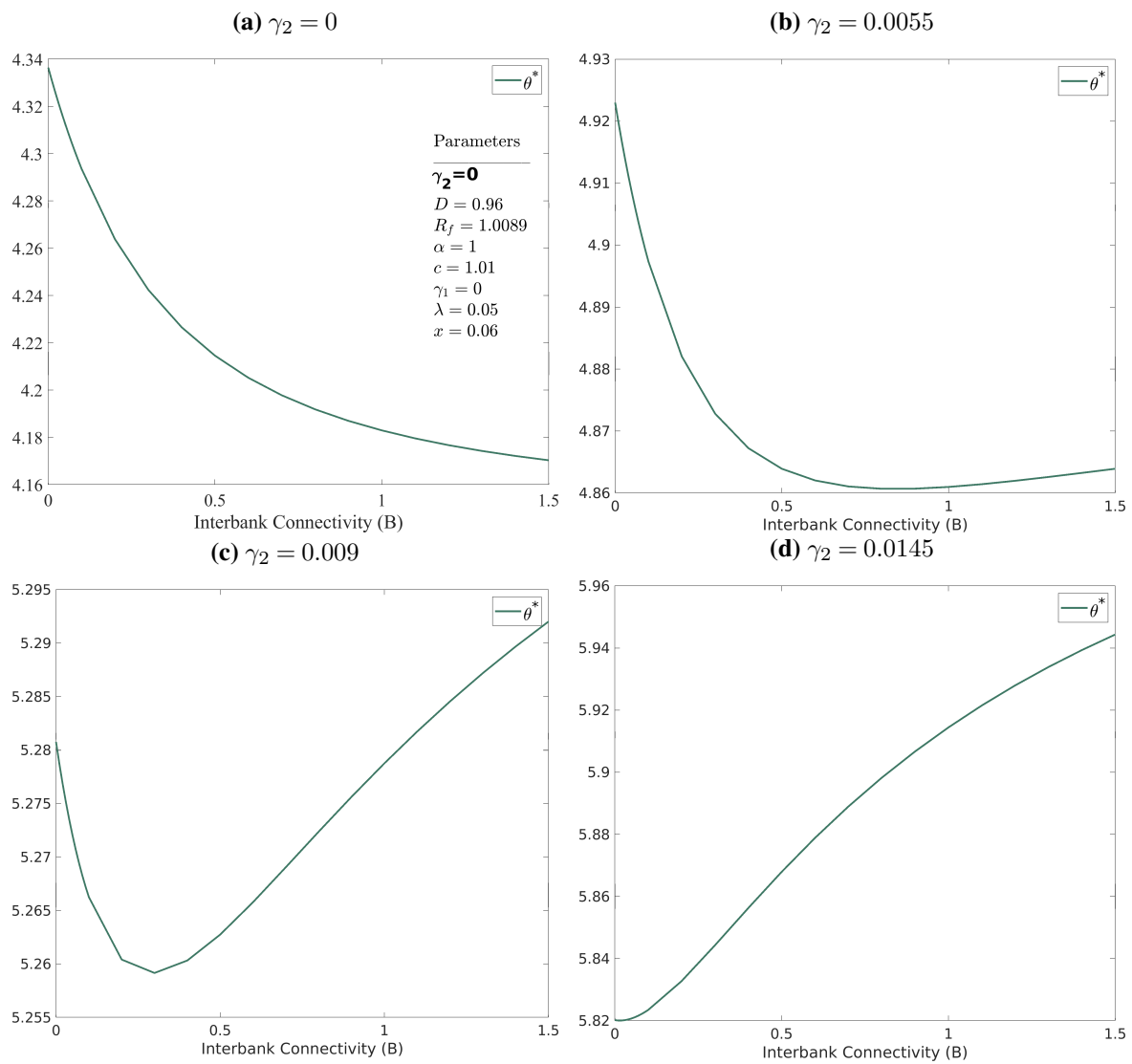
From a legal perspective, netting is in conflict with a central principle in bankruptcy law – mandatory stay. By disallowing creditors to collect their assets, mandatory stay aims to allow an orderly liquidation of the estate in order to dispense to creditors maximal value on pro-rata basis. Netting is a form of exclusion from mandatory stay: creditors who are allowed to net are in effect able collect their assets by not paying their liabilities.²⁸

²⁷21% if this figure is discounted by corporate bond yield.

²⁸There are two main categories of claims on the bankrupt that are excluded from the estate: claims with right to set-off and secured claims. The main difference between the two categories is that behind the exclusion of a secured claim from mandatory stay stands a firm legal principle, namely security defeats stay – *since the debtor essentially granted the creditor ownership of the security*. Netting negates mandatory stay directly, since there is no transfer of ownership. Moreover, since there is no ownership transfer, such contracts escape negative pledges – meaning third party creditors cannot challenge them in courts (Wood, 2007).

Rather than being based on a legal principle, the rationale for the right to net is based on two consequentialist arguments. First, in absence of the right to set-off, OTC dealer markets may suffer a reduction in intermediation services.²⁹ Second, if a crisis occurs – the inability to set-off will contribute to domino-effect contagion. My work adds another consequentialist argument to the latter, demonstrating that netting of interbank liabilities may in fact decrease financial stability due to the bank-run channel.

²⁹Because taking large positive and negative positions in the trade book would be more risky and thus more costly.

Figure 1.5 Threshold strategy (θ^*), varying interbank connectivity (B); $\gamma_1 = 0$ 

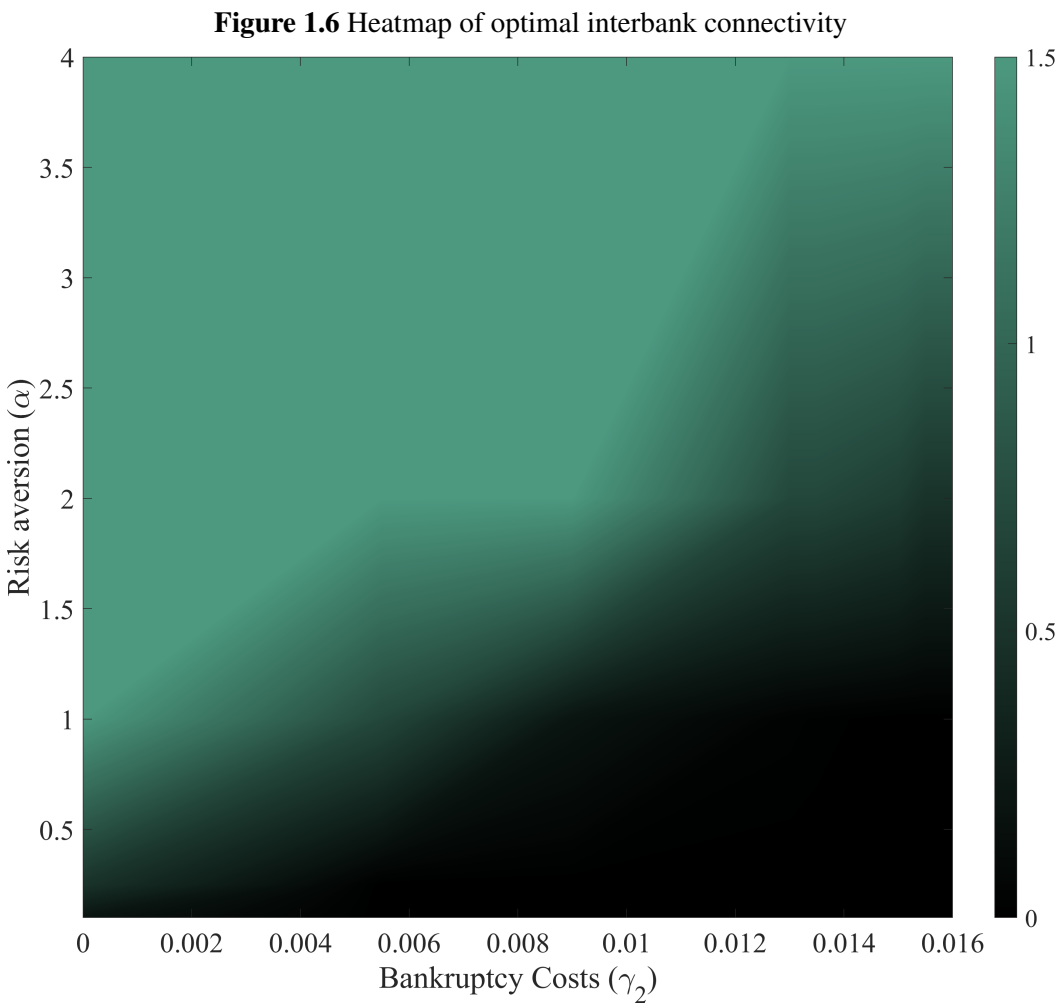
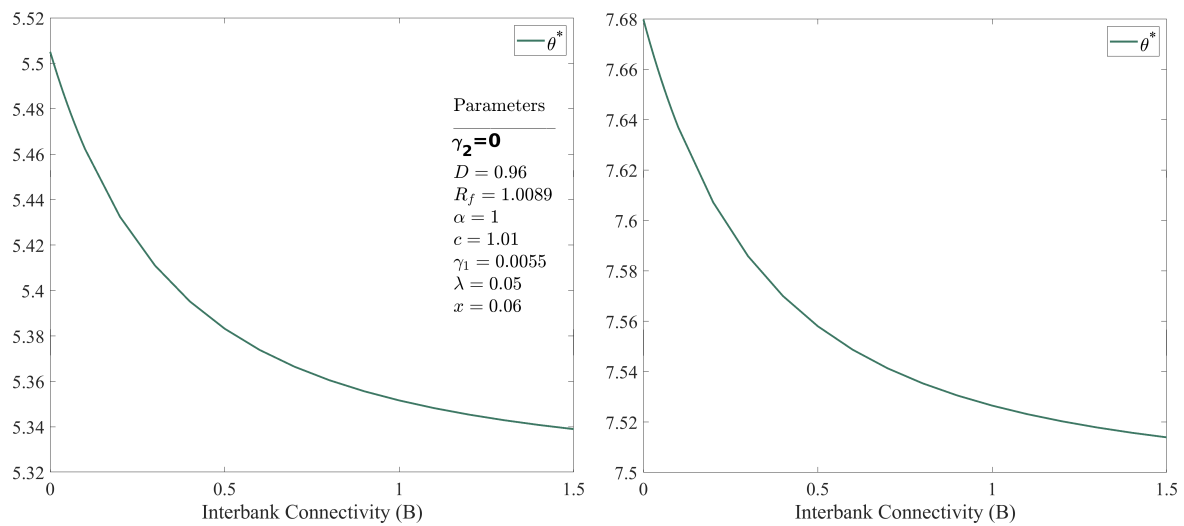


Figure 1.7 Threshold strategy (θ^*), varying interbank connectivity (B); $\gamma_2 = 0$ **(a)** $\gamma_1 = 0.0055$ **(b)** $\gamma_1 = 0.0155$  θ^* is presented in annualized percent

2

Do Netting Rules Reduce Systemic Risk?

This chapter brings to view new evidence about the effect of netting rules on financial stability. It develops a new methodology for evaluating changes in systemic risk, making use of recent developments in banking theory by Goldstein and Pauzner (2005). Using this methodology, I document the co-movement of firms' fundamentals (defined by Merton's Distance to Default) and default risk measures based on CDS prices. I find no support that netting rules reduced systemic risk – which was their main aim. Comparing American to European companies, I show that a two years trend in weakening correlations stopped for American firms shortly before a legislation that changed the U.S. bankruptcy code was enacted – but not so for European and British firms. Furthermore, I find that the financial sector's pre-financial crisis correlations between fundamentals and default risk were either completely absent, or were small and of the opposite sign to the predictions of economic theory and elementary economic intuition. This was the case both in Europe and in North America. In contrast, non-financial firms in both economic regions exhibited strong and statistically significant correlations that are consistent with economic theory. This contrast between the financial sector and other sectors suggests that market participants priced-in implicit government guarantees to the financial sector before the financial crisis of 2007-8.

2.1 Introduction

When a company is bankrupt, an injunction (called Mandatory Stay) is imposed on the creditors of the delinquent which disallows them from independently pursuing the collection of their claims. This is done in order to avoid the uncoordinated scramble of creditors to

increase their own recovery rate. Assets are typically distributed to the creditors on a pro rata basis. For example, if a company owes \$100 to two creditors – \$50 to each – but it only has \$80 on hand, the creditors would be entitled to \$40 each. If instead one creditor was owed \$60 and the other \$40, then the first would be entitled to \$48 and the second to \$32. This simple rule is referred to as *pari passu* (on equal footing), meaning that every creditor is treated without preference to every other creditor. It is a central tenant of bankruptcy law in all advanced economies, and has long been accepted as the principle for the division of value between creditors acting collectively against a common debtor, both in and outside of court (Baird, 2017). However, contemporary bankruptcy law (in the U.S. and elsewhere) treats systemically important financial institutions (SIFIs) more favourably than other companies, by allowing them to ‘net’ certain mutual assets and liabilities with an insolvent company.

Netting is a form of exclusion from the *pari passu* rule. In the special case where one of the creditors is also a debtor of the bankrupt, netting allows her to pay herself by not paying the bankrupt. This is especially prevalent for large banks due to their activity as security dealers, giving rise to mutual credit exposures between them. The rationale behind allowing netting is to reduce this credit risk, so as to lessen the social cost of cascading losses during a financial crisis (Chapman, 2016). In practice this is accomplished by excluding qualified financial contracts (QFCs) from the mandatory stay on a delinquent’s assets. Netting rules were first introduced into U.S. legislation in 1978, and have since evolved substantially.¹ Until the early 2000’s, eligibility for netting was smaller than today. In order to net an asset and a liability with a delinquent counterparty, the products being netted had to be of similar underlying and maturity. But this changed in 2005 when Master Netting agreements became fully effective, and there was no more need for products to have similar attributes in order to net them (Mooney Jr, 2014). The practice of netting has two direct effects: (a) an effect on financial stability due to contagion and risk sharing; (b) an effect on trade in derivatives due to reduced marginal cost of intermediation (see Bolton and Oehmke, 2015).

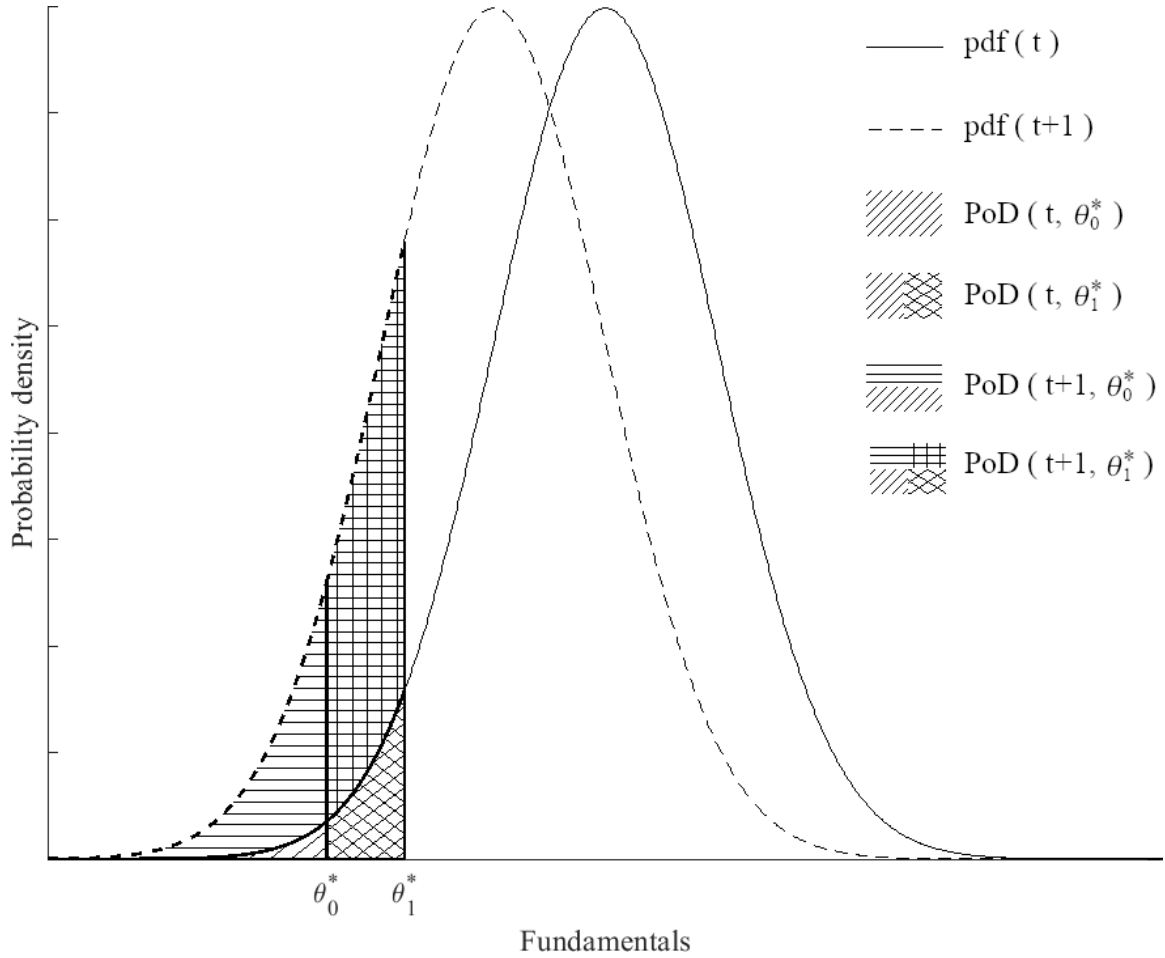
The aim of this chapter is to empirically evaluate (a): do netting rules reduce systemic risk? To this purpose I investigate a panel of financial and non-financial companies, documenting how a measure of financial crash risk changed before and after the legislation was enacted: the Bank Run Index (BRI, Veronesi and Zingales, 2010). Although the measure was originally conceived to capture risks in the banking sector, in fact nothing prevents us from computing it for non-financial companies as well. Calling it a *Bank Run Index* is somewhat of a misnomer, as what it actually captures is the relative importance of short term risk. See the Discussion (section 2.5) for a further deliberation of this assertion. I use

¹Historically, English common law allowed netting since the 18th century, while French Napoleonic Code have always disallowed it. In Germanic legal traditions, netting is allowed only when stringent conditions are fulfilled about the type of securities being netted. Most other countries follow one of these legal traditions.

the BRI in conjunction with a company's Distance to Default (Merton, 1974) in order to capture firms' fundamentals. The novelty in my approach to measuring systemic risk is the connection I make between the two: a company's default risk and its fundamentals. An increased correlation between the two indicates a heightened systemic risk. I shall refer to this correlation as systemic- β . Unlike the simple measures used to evaluate market participants' concern of an imminent crisis, such as the Probability of Default (PoD) and BRI, using the correlation between fundamentals and these measures could reveal structural changes to the endogenous threshold at which investors refuse to roll over a company's debt. This approach is inspired by Goldstein and Pauzner (2005), who reconcile the two canonical positions in the debate about the origin of banking crises, namely whether they are driven by real shocks (Gorton, 1988) or coordination failures (Kindleberger and Aliber, 1978).

The basic idea runs as follows. A bank run occurs when creditors lose confidence in a bank's ability to repay its liabilities the next time round, therefore refusing to roll over their claims this time round. The loss of confidence depends on both: (a) the objective circumstances (the fundamentals of the bank); (b) the subjective evaluation of each creditor, regarding the actions of other creditors in equilibrium. The subjective circumstances are important because of the resemblance of rolling over debt to the game of musical chairs: one never wants to be the most optimistic forecaster of the time the music stops. The objective circumstances matter for two reasons: directly, indicating financial health; indirectly, for their function as a coordination device. Thus, a loss of confidence would typically occur over a short period of time, in which multiple measures of financial health would decline steeply. As a consequence, we would expect to observe a negative correlation between public signals of financial health (such as market capitalization) and the price of insurance against default – that is, would expect systemic- β to be negative.

Figure 2.1 depicts the theoretical relationship between the PoD and a company's fundamentals. The full (dashed) line stands for the probability density function of the distribution of fundamentals, conditional on time t ($t + 1$) information. The area below it filled with diagonal hatching is the probability of default, assuming creditors lose confidence in the company when its fundamentals fall below θ_0^* . Should this threshold increase to θ_1^* , the PoD would be given by the former area and in addition the area filled with crossed diagonal hatching. Time $t + 1$ PoDs are given analogously with additional horizontal and crossed hatching. It is readily clear that the higher is the threshold θ^* , the larger is the change in PoD for a given change in fundamentals. Thus, if we can observe proxies of both fundamentals and of PoD, in principle we may detect changes to θ^* by comparing the absolute change in PoD for a given change in fundamentals.

Figure 2.1 Systemic- β : The Correlation between fundamentals and default risk

What is the difference, then, between previous studies on the measurement of systemic risk and the one I propose here? One of the leading approaches in the literature was provided by Adrian and Brunnermeier (2016), who analyze the cross-section covariance of tail risk (i.e. conditional covariance, conditioning on negative tail events).² In contrast, my approach concerns the (unconditional) covariance between fundamentals and PoD *within a single company*. Furthermore, my approach is informed by an economic theory about the origins of correlated tail risks as an endogenous phenomenon, while previous studies are not. For example, Acharya et al. (2017) and Glasserman and Young (2015) are based on the domino-effect literature, which views tail risk as a result of mechanic contagion. Most empirical studies analyze systemic risk without a direct reference to the theory, instead opting to define it ad hoc in terms of correlated tail risks. Remaining agnostic about the origin of these correlations leaves us with the attempt to make inference at the shape of the conditional joint

²It is different to the traditional asset pricing view of risk, that emphasizes the unconditional covariance of returns with the market (and/or other factors).

distribution, using data that is typically conditioned on a different state. Put differently, the assumption is that we can learn about what will happen to asset prices during a crisis from a non-crisis sample. My approach does not require this. Using it we may not be able to identify the exact threshold of fundamentals below which creditors of a company will no longer roll over their claims, but we are able to detect changes in this threshold insofar as we observe variation in both fundamentals and PoD (or their proxies).

Turning to the empirical results of the chapter, my contribution is threefold. First, I document the co-evolution of the Merton model distance to default (DD) – the measure I use to proxy fundamentals – with CDS implied risk-neutral PoD and its derivative, the BRI. I do so for a large sample of public companies both in Europe and in North America, showing that: (1) both PoD (an absolute default risk measure) and BRI (measuring the relative importance of short-term default risk) decreased significantly between 2003-2007; (2) both PoD and BRI of large banks were lower than other companies before the financial crisis, reversing only in mid-2007; (3) both PoD and the BRI of American companies were on average higher as compared to European companies before the financial crisis; (4) this gap persistently narrowed over the years before the financial crisis.

Second, I find that the financial sector's pre-crisis systemic- β (defined via the correlation between DD to BRI) were either completely absent, or were small and of the opposite sign to the predictions of economic theory and elementary economic intuition. This was the case in Europe and in North America, and in both regions it was driven mainly by large banks rather than other financial companies. In contrast, non-financial firms in both continents exhibited strong and statistically significant systemic- β that are consistent with economic intuition. This contrast between the financial sector and other sectors supports the view that market participants priced in implicit government guarantees to the financial sector before the financial crisis of 2007-8, as in Kelly et al. (2016). The latter study financial crash insurance in options markets to learn about implicit government guarantees to the financial sector. Their findings demonstrate the difficulty to learn about systemic risk from market prices, because of market participants' expectation of implicit government guarantees, clouding absolute measures of risk. My findings serve as an example of how this clouding effect obscures the measurement of systemic risk. Nevertheless, not all sectors enjoy implicit guarantees, thus at least some inference can be made with relative confidence.

Lastly, comparing American and non-American companies, I find no support to the claim that netting rules reduced systemic risk – which was their main goal. Comparing American to European companies, I show that a two years trend in weakening systemic- β slowed markedly for American non-financial firms exactly at the time when the legislation was enacted, but didn't for European and British firms. This was the case although fundamentals

stayed roughly at the same level they were prior to the change to U.S. legislation. Instead of observing an acceleration in the weakening of correlations between fundamentals and default risk measures, we observed the exact opposite. Financial companies did not exhibit any notable change in trends. These findings do not show that the expansion of netting rules *caused* the shift in the trend, because the nature of the policy was that it affected the whole market, making it difficult to identify a proper control group. However, insofar as we accept the premise that American companies are more highly affected by changes to U.S. legislation, this finding is at least partial evidence that the policy did not attain its declared aims.

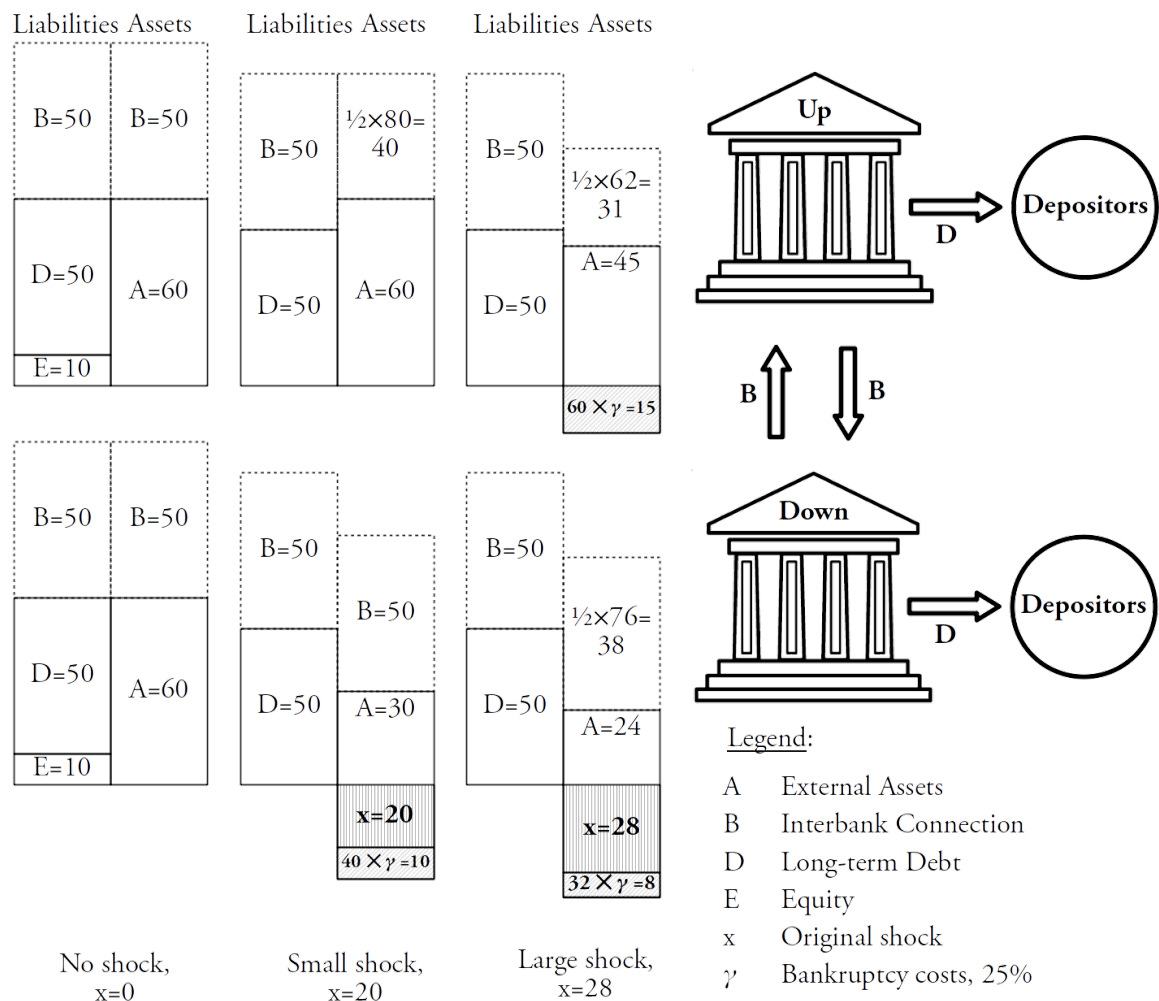
The rest of the chapter is organized as follows. In the next subsection I provide a numerical example of how interbank credit exposures could give rise to a social cost due to the domino effect. The subsection after that reviews the related literature. In section 2.2 I describe the methodology by which I propose to measure changes to systemic risk. Section 2.3 describes the data, providing summary statistics and briefly discussing the time-series of relevant variables. Section 2.4 contains the empirical results of the chapter, and section 2.5 concludes.

2.1.1 Numerical Example

Figure 2.2 illustrates how the domino effect is manifested through the propagation of shocks between banks. The rectangles in the diagram depict the balance sheets of two banks, Up and Down, who are connected via an interbank mutual claim worth \$50 (denoted by B). At an initial stage (left panel, no shock, $x=0$) they are both identical. In addition to B, each bank owns \$60 worth of external assets (denoted by A), making the total size of their balance sheet \$110. On the liability side, they have the interbank claim as well as another claim by depositors worth \$50 (denoted by D). At this initial stage, the equity claim on each of those banks is worth \$10 (denoted by E).

Should bank Down suffer a \$20 loss on its external assets (middle panel, small shock, $x=20$), it would no longer be able to honor its liabilities in full to its creditors – its own depositors and bank Up. Let us assume that on bankruptcy a bank loses an additional 25% of the value of its external (non-interbank) assets, augmenting bank Down's impairment to \$30. Provided that bank Up pays its liabilities to bank Down in full, the latter would have \$80 in total assets to be distributed equally between bank Up and its depositors, in accord with the *pari passu* rule. This then means that bank Up would suffer a loss of \$10 on its interbank assets, eliminating its equity though escaping default. Because bank Up is not bankrupt, it has to honor its liabilities to bank Down in full, even though bank Down does not pay in full.

In this case there is no domino effect, because bank Up is not bankrupt. If however, bank Down would be hit by a larger shock destroying \$28 worth of its external assets (right panel,

Figure 2.2 Interbank credit exposures and the domino effect

large shock, $x=28$) – and added to this an additional \$8 due to bankruptcy – its external assets would be worth \$24, leaving it unable to meet its liabilities once more. In contrast to the former case (small shock), this time bank Up would default as well, meaning it will lose \$15 due to bankruptcy costs. All in all, bank Up would be able to distribute \$76 – \$38 to each creditor. Bank Down, who suffered the original shock, would only be able to pay \$72 – \$31 to each creditor. In this example, the loss of \$15 due to bankruptcy is the social cost of the domino effect, because it is a loss that is due to the banks being connected via credit exposures.

Eligibility for netting of %40 of interbank debt would have prevented this. It reduces B to 30, meaning bank Down would have \$54 to distribute to creditors, \$20.25 to bank Up and \$33.75 to its own depositors. Bank Up would lose \$9.75, a loss that would be absorbed by its equity holders – bank Up would remain afloat, avoiding the costs associated with bankruptcy.

2.1.2 Related Literature

This chapter is related to four main strands of literature in economics and finance: measurement of systemic risk, default risk, financial intermediation and financial stability. Since the latter topics are extensively covered in the former chapter of this dissertation, to avoid duplication this section is limited to the first two.

Measuring ‘systemic risk’ is not a straightforward task. There are multiple ways by which academics differentiate this term from the more well known ‘systematic risk’ (Allen and Carletti, 2013), and measurement will depend on the exact phenomenon one wishes to capture. The term ‘systemic risk’ preceded the financial crisis, but after 2008 there was a boost to research on the topic since it was difficult to reconcile the highly correlated tail risks we observed during the 2007-8 crisis with traditional asset pricing theory. Much of the literature emphasized high correlations and inter-dependencies, especially in tail risk, as an operational definition in measuring systemic risk – while remaining agnostic about its origin. Goodhart et al. (2009) use Consistent Information Multivariate Density to infer the non-linear relationship in banks’ distress, using equity and CDS prices. Borio and Drehmann (2009) examine leading indicators using equities, bonds and property prices. De Jonghe (2010) applies extreme value analysis to equities, assessing the tail risk of a company conditional on an index equities experiencing a bad draw. Billio et al. (2012) document that the financial system’s asset returns have become more correlated using PCA and Granger-causality network. Like the current study, Huang et al. (2009) use equity and CDS prices to study systemic risk. They propose to measure systemic risk as the theoretical premium that would have to be paid for protecting against large defaults. Acharya et al. (2017) propose to measure systemic risk via bank under-capitalization. Since bankruptcy is costly, not holding sufficient capital imposes an externality on other banks. This approach is rooted in stress testing contagion models, emphasizing the practicality of the measurement for the purpose of central bank supervision.

Adrian and Brunnermeier (2016) revamp the practical Value-at-Risk (VaR) measure so that it takes into consideration the co-variation in risk of individual units in the financial system; conditional on a subset of institutions coming under strain, it asks: what is the risk to the rest of the system? The authors propose a forward looking measure of systemic risk (forward- $\Delta CoVaR$, an early warning signal) based on quantile regressions. Interestingly, this measure is negatively correlated with its future realizations, making it consistent with the notion that low volatility environments are a breeding ground for systemic risk. This stylized “fact” can be readily impressed upon us by observing the low volatility in the CDS market just two years prior to the financial crisis (see Figure 2.3) – which is coincidentally also the period after the enactment of insolvency netting in the United States. That CDS markets are

less volatile after the law changed makes it less likely to reject the hypothesis of this chapter, that a given change in fundamentals brings about a lower change in the expected probability of default. This fact in itself is not sufficient for rejection though, because the volatility of fundamentals can also change. In the final analysis, the notion of CoVar puts the emphasis on the magnitude of co-variation of bad events across separate financial entities, but for the same security (e.g. equity). In contrast, this chapter emphasizes the magnitude of co-variation across two different securities (stocks and bonds) within the same entity. The significance of analyzing multiple securities for a single entity is that in this way one may analyze the relationship between different types of risk, keeping constant unobserved properties of the firm.

Another branch of literature related to the current chapter deals with the estimation of default risk. Modern default analysis, one that applies rigorous statistical methods for the study of bankruptcy risk, dates back to Altman (1968) and Ohlson (1980). These studies use the actual occurrence of bankruptcy to assess the prediction power of self-reported financial ratios of profitability, liquidity, and solvency. This approach has a range of problems inherent to using this type of data, such low frequency, timeliness (being backward looking) and even concerns about reliability of self-reported information.³

A fundamentally different approach to studying default risk is to extend the textbook analysis of equity- to bond-markets (see e.g. Campbell, 1987; Fama and French, 1989; Fama and Schwert, 1977; Keim and Stambaugh, 1986). It is different from the former in that it relies on market information to detect investors' expected returns – in contrast to Altman / Ohlson where the actual occurrence of bankruptcy is used to define risk. My thesis shares the approach of using market expectations to define risk, and in doing so escapes these problems. The downside of using market based measures is that default risk has to be disjoined in some way from other types of risk, because default is only one part of a more complex risk-reward compound; this is a challenge my thesis is also facing. The way I handle this is by using multiple securities carving up the same underlying cash flow.

There are, however, two problems with extending the commonly applied risk-reward analysis from equity to bond markets which my thesis does not share. First, bond markets are much less liquid than equity markets, exacerbating noisy measurement issues. In contrast, I use data from CDS markets that are more liquid than bond markets (Longstaff et al., 2005; Zhu, 2006). Second, textbook equity analysis relies on historical returns to learn about market expectations. Although equity/bond prices are forward looking, returns are not. Alternative methods have been developed in recent years, using options and other derivatives

³Recent applications of this approach include Dichev (1998) and Griffin and Lemmon (2002), who utilize accounting based measures as well as market information from equity markets.

to learn more directly about market expectations (Martin and Gao, 2021; Ross, 2015). My thesis draws upon these developments, using CDS prices to back out expected probability of default.⁴

Yet another popular way to analyze default risk with market based information is through the yield spread between bonds with low and high credit ratings (Dichev and Piotroski, 2001; Hand et al., 1992; Holthausen and Leftwich, 1986). Using this approach, it has been shown that bond downgrades are typically followed by negative equity returns, that default risk varies with the business cycle, and is related to macroeconomic factors (Denis and Denis, 1995).⁵ This approach recognizes that default risk is a complex phenomenon that can be difficult to capture with parsimonious models. Rather than attempt to construct a credit model, it relies on professionals who's job it is to do just that. Studies using this approach typically assume that all assets within a rating category share the same default risk. This level of aggregation is bound to lose information contained in the cross section.⁶ It also assumes that credit downgrades/upgrades capture changes in default probability without a lag (Kealhofer et al., 1998). In addition to the latter critique of the yield spread methodology, there are also empirical problems with it. Some studies find that the information contained in the default-spread is largely unrelated to default risk (see e.g. Elton et al., 2001). Doubts have also been cast on whether credit ratings contain any information above and beyond what is already contained in market prices.

The Merton (1974) model is the classical reference in this regard, providing a 'distance-to-default' (DD) measure by combining equity market and balance sheet information. My thesis makes considerable use of the Merton model. Since its inception, a large literature used it to analyze default risk. Vassalou and Xing (2004) examine how distance to default is related to equity returns. They find that measures based on a Merton model contain different information from aggregate default spreads. Benos and Papanastasopoulos (2007) showed that using distance to default, we can predict credit rating changes. Duffie et al. (2007) show that the probabilities of default from a Merton model exhibit considerable predictive power. There are also papers that are critical of the Merton model. Du and Suo (2004) and Bharath and Shumway (2008) find that distance to default is not a sufficient statistic for credit quality, as is suggested by the theory underlying it.⁷ For the purpose of the current study, a strong

⁴Strictly speaking the expectation is under the risk-neutral measure, since the discount factor I use does not assume any form for utility function of investors (see Borovička et al., 2016).

⁵There are, however, other papers that find that default risk is idiosyncratic (Asquith et al., 1994; Opler and Titman, 1994), which seems to be in contradiction to the latter.

⁶Also, default risk is defined via the historical average of default, and is thus backward looking.

⁷This could reflect the fact that some of assumptions made by the model are rather strong, e.g. that the company's debt claims are all of equal maturity. Indeed, practitioners have long developed databases to address this issue, e.g. the model used in practice by Moody's relaxes this assumption.

correlation between DD and PoD/BRI (systemic- β) would make it easier to detect structural changes, yet it is not a necessary condition for testing the hypothesis the study sets out to assess. It is sufficient to have *some* level correlation, enough so that the changes in it can be statistically detected given the sample size. The advantage of using the DD is that it inherits the relatively high frequency with which information is processed by equity markets, since it is essentially an adjustment of equity market information for bond-holders. Moreover, it is a parsimonious model, and is well rooted in asset pricing theory.

2.2 Methodology

In this section I explain the hypothesis this chapter sets out to assess, namely that reducing credit risk between large banks reduces systemic risk. I flesh out a formal test by which I propose to do this, describing the challenges to identification and how I propose to sidestep those challenges until a proper identification strategy is conceived. Finally, I explain in detail how I construct each of the objects I use to in the formal test.

I define systemic risk in terms a threshold of ‘fundamentals’ at which the endogenous response of agents bring about a financial crisis. A policy is beneficial for financial stability if it is able to push down this threshold, i.e. that systemic events occur at worse fundamentals. In order to test whether a policy brought about an improvement to systemic risk, one has to analyse the co-movement of firms’ fundamentals and firms’ probability of default (PoD) – before and after the policy (exogenous treatment) is introduced. As an operational definition of systemic risk I take the average within-firm correlation between fundamentals and credit risk, and term this correlation systemic- β . Typically we would expect systemic- β to be negative because as the financial prospects of a firm decrease, its credit risk would rise. A more strongly negative systemic- β could indicate either that fundamentals worsened, or that the threshold of fundamentals at which creditors refuse to roll over their claims has increased. Thus if we are able to control for firms’ fundamentals, we could in principle observe changes in this threshold. The rest of this section describes how I propose to measure systemic- β .

I shall adopt widely used measures of credit risk and fundamentals. Let $\text{PoD}_{it}(T)$ be the expected probability that firm i defaults in the next T periods, conditional on time t information. If a policy is to reduce systemic risk, we should – all other things equal – observe a lower PoD after the policy was introduced. This effect should be immediate, potentially preceding the actual change in policy by a short amount of time. The main variable we need to keep “equal” in this comparison is firms’ fundamentals. I propose to measure it via the distance to default of a company (DD, Merton, 1974). Distance to default is a volatility-adjusted measure of leverage (Duffie et al., 2007); thus it captures two essential properties

of fundamentals: (a) the market value of the company; (b) the uncertainty about this value. There are two main factors driving PoD: the long term value of assets (fundamentals); and the short term risk of a systemic event. The market value of a company is the best estimate of the former; the near- term structure of the company's CDS best captures the latter.⁸

A derivative of PoD constructed to capture the relative importance of short term risk is the Bank Run Index (read further down this section for additional information on how its constructed). At a first stage I measure the average systemic- β of a group of companies by regressing BRI_{it} on the distance to default DD_{it}

$$BRI_{it} = \alpha + \beta DD_{it} + \varepsilon_{it} \quad (2.1)$$

This captures the average correlation between fundamentals and short-term default risk at the firm level. As explained above, ex-ante we would expect β to be negative, though potentially dynamic because of changing macro conditions. I run panel regressions of equation 2.1 for North American and European firms separately, using an estimation window of 90-calendar days – continuously shifting it in time from 2001 to 2010 in order to capture the dynamics of this average correlation in each group. Let $\hat{\beta}_{jt}$ be the estimate of β in equation 2.1 for group j where the estimation window ends on date t . One simple way to test whether the two groups had parallel trends before the intervention is to regress these estimates $\hat{\beta}_{jt}$ on a constant, and a time trend for each group,

$$\hat{\beta}_{jt} = \gamma_0^j + \gamma_1^j \times t + \eta_{jt} \quad \text{for } t \in \text{PRE}(20\text{-Apr-}2005, 2) \quad (2.2)$$

for a sample before the intervention: $\text{PRE}(20\text{-Apr-}2005, 2) = \{t \mid t < 20\text{-Apr-}2005 \wedge t \geq 20\text{-Apr-}2003\}$. The parallel trends hypothesis is that $\gamma_1^j = \gamma_1^{-j}$. In section 2.4 I show that the parallel trends assumption is violated. Until a proper identification strategy is conceived, I propose to test whether each of the groups experienced a structural break on 20 April 2005. I run multiple Chow tests for each group separately, regressing the estimates $\hat{\beta}_{jt}$ on a constant and a time trend for three samples: FULL, before date τ (PRE) and after after date τ (POST)

$$\hat{\beta}_{jt} = \gamma_0^j + \gamma_1^j \times t + \eta_{jt} \quad (2.3)$$

⁸This is not to say equity is not affected by systemic risk, nor that near-term CDS prices are not affected by long-term risk. There is however, the question of the magnitude of influence of each type of risk on each security. If we think that systemic events occur at rather low fundamental value, the loss to equity holders from this type of risk is also low. Similarly, strong fundamentals do imply lower CDS prices, but since the main beneficiaries are residual-claim holders, we would expect near-term insurance prices to be less sensitive to changes in the long term value of assets.

Formally, let τ be any potential date we want to test for a potential structural break in the trend, let $\text{FULL}(\tau, y) = \{t \mid |t - \tau| < y\}$ be a set of dates y years about τ ; $\text{PRE}(\tau, y) = \{t \mid t \in \text{FULL}(\tau, y) \wedge t < \tau\}$ and $\text{POST}(\tau, y) = \{t \mid t \in \text{FULL}(\tau, y) \wedge t > \tau\}$ split FULL into two equal periods at τ . For each date τ I compute the Chow test-statistic $C^j(\tau, y)$ is given by

$$C^j(\tau, y) = \frac{(\text{SSR}_{\text{FULL}}^j - \text{SSR}_{\text{PRE}}^j - \text{SSR}_{\text{POST}}^j)/2}{(\text{SSR}_{\text{PRE}} + \text{SSR}_{\text{POST}})/(2y - 4)} \quad (2.4)$$

where SSR_m^j is the sum of squared residuals for group j and the sample dates $m \in \{\text{FULL}, \text{PRE}, \text{POST}\}$. If the errors are normally distributed, $C^j(\tau, y)$ follows an F -distribution with $2y - 4$ degrees of freedom. I run those tests for pairs (τ, y) such that $y \in \{\frac{1}{4}, 1\}$ and dates between January 2003 to July 2007. Finally, I fix the four year window around 20-Apr-2005, running Chow tests at each date within it. If there was a structural break, we would expect a high Chow test-statistic for the North American sample, on (or shortly before) 20-Apr-2005 – but not so in Europe.

Moreover, since we are trying to assess not only that there was *a* structural break, but also that this break *reduced* systemic risk, we have one additional testable implication. If investors' concerns about an imminent crisis fell over time, we would expect the coefficient γ_1 to be positive. Given these expectations, under the null hypothesis that the effect of the policy was to reduce systemic risk, we would expect $\gamma_{1,\text{POST}} > \gamma_{1,\text{PRE}}$ for companies in the U.S., and not so (or to a lesser extent) for companies in Europe. These tests do not provide the ultimate evidence that an intervention caused any potential structural break. However, insofar as the intervention achieved its desired effect, we would at least expect a tendency in that direction.

Ideally, we would be able to identify a control group which is unaffected by the policy intervention, and a treatment group which is fully affected by it and to which companies are randomly assigned. These ideal conditions are difficult to meet in any 'natural experiment', but especially so for insolvency netting because it is plausible that all companies are to some extent affected by this change to U.S. bankruptcy law. This challenge can potentially be overcome by recognizing that not all companies should be affected by the policy in the same way. For example, we might expect the effect to be different for financial companies, especially those who are active in derivative markets – since the policy directly benefited them. Furthermore, we may expect companies in the country where legislation changed to be more highly affected by it. This recognition does not fully resolve the challenge of identification since there is still the problem of non-random selection and the violation of the parallel trends assumption. This thesis does not have an adequate fix for this problem, thus

claiming causality is precluded. Keeping these reservations in mind, I maintain that it still worth the while to empirically examine the effects of the policy.

I run the regressions and tests above for groups of companies differing in two main attributes, comparing geography (North America and Europe) and further breaking down each group by company type (financial / non-financial, large banks / other entities). In the subsections below I describe how I construct each of the measures used in these regressions: BRI, PoD and DD.

2.2.1 Probability of Default

I follow Veronesi and Zingales (2010) to compute firms' survival probabilities $\mathcal{Q}(T)$, using daily CDS prices from Markit:

$$\mathcal{Q}(T) = e^{-\int_{u=0}^T p(u) du} \quad (2.5)$$

where p is a vector of default intensities which has to be estimated from the term structure of CDS. The estimation uses the following approximation

$$\text{ParSpread}(T) = (1 - \delta) \int_{\tau=0}^T p(\tau) DF(\tau, r, p) d\tau \quad (2.6)$$

$$DF(\tau, r, p) = \frac{e^{-\int_{u=0}^{\tau} [r(u) + p(u)] du}}{\int_{v=0}^T e^{-\int_{u=0}^v [r(u) + p(u)] du} dv} \quad (2.7)$$

where T is the Tenor of the CDS contract, r is the vector of the (expected) risk-free rates, DF is the discount factor and δ the expected recovery rate. I estimate the discretized⁹ version of p numerically, using a least squares approach: finding the vector p which minimizes the sum of squared residuals of the model-based par spread (as computed from equation 2.6) and observed par spreads. I use LIBOR and swap data from FRED to estimate the term structure of expected risk free rates (see Appendix B.2 for further details), and Markit's RealRecovery to as input for δ .¹⁰ Finally, the probability of default, PoD is given by

$$\text{PoD}(T) = 1 - \mathcal{Q}(T) \quad (2.8)$$

⁹The time steps I use for discretization are: (3M, 6M, 9M, 1Y, 2Y, 3Y, 4Y, 5Y, 7Y, 10Y, 15Y, 20Y, 30Y) where M stands for months and Y for years.

¹⁰In case of missing information, I use the value 0.4, which is standard assumption in the literature.

2.2.2 Bank Run Index

The Bank Run Index, $BRI(T)$, is a measure of short term risk proposed by Veronesi and Zingales (2010). It utilizes the term structure of CDS to gauge at the gap between the probability of default before year T and the conditional probability of default before year $T + 1$, $\mathcal{P}(T + 1)$ – conditioning on default not occurring before year T . I follow Veronesi and Zingales and take $T = 1$ year. In “regular” times, the probability of default is an increasing function of T . However, if a company is facing a heightened risk that its creditors won’t roll-over their claims, then we should observe an inversion of this relation. Conditional on surviving the short-term, we would expect the probability of default to decrease. To compute the conditional probability of default, $\mathcal{P}(T + 1)$, we use Bayes’ rule

$$\mathcal{P}(T + 1) = \frac{\mathcal{Q}(T) - \mathcal{Q}(T + 1)}{\mathcal{Q}(T)} \quad (2.9)$$

Then $BRI(T)$ is defined via

$$BRI(T) = PoD(T) - \mathcal{P}(T + 1) \quad (2.10)$$

thus in times of concern about a company’s short-term prospects, $BRI(T)$ should be positive. Otherwise it would be negative or close to zero.

2.2.3 Distance to Default

Following Vassalou and Xing (2004) I compute distance to default (DD) using a Black-Scholes-Merton (BSM) model. Assuming the market value of the firm’s total assets V follows a geometric Brownian motion:

$$\frac{dV}{V} = \mu dt + \sigma dW \quad (2.11)$$

where μ, σ are respectively V ’s instantaneous drift and volatility, and W is a standard Wiener process. Then, the the distance to default is given by

$$DD = \frac{\ln(V/F) + \mu - \frac{1}{2}\sigma^2}{\sigma\sqrt{T}} \quad (2.12)$$

where F is the face value of the firm’s debt, and T its maturity. Since V is not directly observable, one has to estimate it using market data on equity and balance sheet data on its liabilities. Assuming the firm’s capital structure is composed from a simple debt claim with

face value F and maturity T , and a residual claim of value E , the BSM model asserts that the total value of the firm's assets V is related to its equity value E via the relation:

$$E = V\Phi(d_1) - Fe^{-rT}\Phi(d_2) \quad (2.13)$$

where r is the risk free rate, Φ is the CDF of a standard Normal distribution, and d_1, d_2 are given by:

$$d_1 = \frac{\ln(V/F) + \left(r + \frac{1}{2}\sigma^2\right)T}{\sigma\sqrt{T}} \quad d_2 = d_1 - \sigma\sqrt{T} \quad (2.14)$$

The output from equation 2.13 is V , which then allows us to compute σ and μ to use as inputs in equation 2.12. As inputs, equation 2.13 takes the firm's market capitalization, the face value of its liabilities, the risk free rate and the standard deviation σ . Following Bharath and Shumway (2008), I adopt an iterative procedure whereby V is computed for each time period t solving equation 2.13 numerically. At the end of each iteration, σ is computed as the standard deviation of changes in $\ln(V)$ in the preceding year to t , which is then used as an input for the next iteration. At the initial iteration, I take σ as the standard deviation of equity returns. The iteration process concludes when σ converges up to desired approximation (I used $1e^{-4}$).

As a proxy of F I used two different variables: total liabilities (from Compustat, lttq); and a composition of short and long-term liabilities. I follow Bharath and Shumway (2008), who combine short-term liabilities with 1/2 of long term liabilities. For non-financial companies I used Compustat's lttq for short-term liabilities, and lltq for long term liabilities. For banks, I used Debt in Current Liabilities (dlcq) and Total Deposits (dptcq). However, I found that Compustat had limited data for non-American companies (most have data only from 2007). Bloomberg was unhelpful in that regard because it also had data starting only 2007. This meant that using a composition of short and long term debt would lose many European financial companies. For this reason I chose the Merton DD based on total liabilities as a benchmark. Results did not change qualitatively for those companies that did have the data. To get market capitalization (E) at daily frequency I used Compustat data on firms' equity, see Appendix B.1 for further details.

2.3 Data

I obtained data on firms' CDS prices, equities and balance sheets from three sources: Markit, Compustat and Bloomberg. In order to use these data-sets together I matched unique

Table 2.1 Summary statistics

Variable	Mean	Std_dev	Min	Pctile_25	Median	Pctile_75	Max
BRI	0.119	0.180	-0.523	0.018	0.051	0.149	0.982
PoD(6M)	0.114	0.164	0.001	0.016	0.046	0.133	0.871
PoD(1Y)	0.146	0.177	0.001	0.031	0.072	0.184	0.982
PoD(5Y)	0.260	0.202	0.001	0.102	0.201	0.359	1.000
DD	3.264	2.069	0.001	1.648	2.948	4.566	11.652
E (\$mm)	19635	36422	4	3225	7827	18865	530,426
F (\$mm)	60,781	218,453	1	3,072	7,918	22,951	3,672,476
V (\$mm)	74.776	205,222	7	7,143	17,752	47,867	2,970,370
μ_V	0.172	0.169	-0.741	0.090	0.141	0.217	4.411
σ_V	0.284	0.161	0.016	0.192	0.246	0.331	2.901

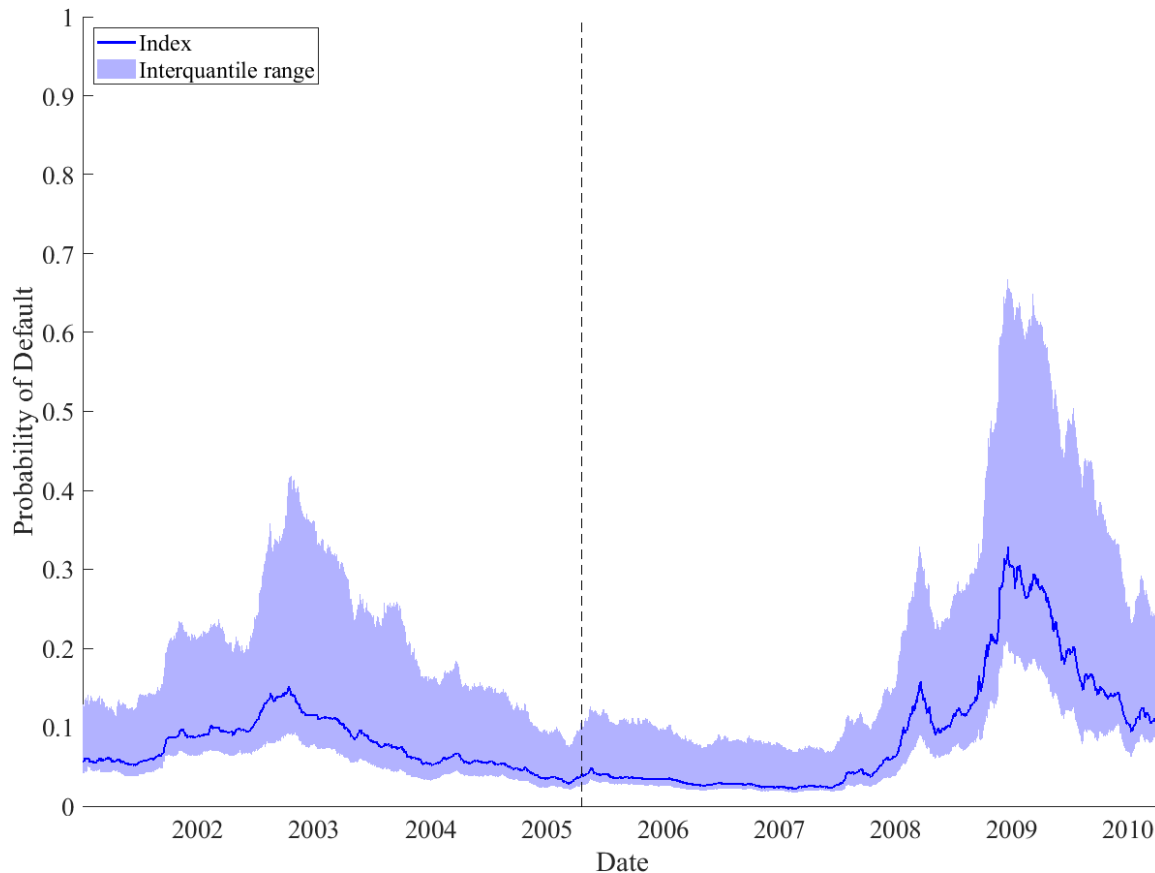
identifiers in each dataset. The matching process was largely manual. A full account of the matching process is given in Appendix B.3. Whenever I refer to a weighted average / index of a variable, I use firms' market capitalization for weighting. The index of the Bank Run Index (BRI) is a weighted average of BRI. The period we are particularly interested in is a limited window around April 2005, when a significant amendment to U.S. bankruptcy code was enacted. As a benchmark I propose to consider two years on either side of this date, to avoid the effect of the financial crisis on the results.

The sample consists of 768 North American and 409 European public companies, with overall 3,823,295 time-firm observations within the estimation window. Table 2.1 provides summary statistics for the full sample. The average market value of a company in the sample is \$18.1 bn, and its median is \$7.2 bn; the largest observation recorded is of Citigroup on 26 Dec 2006, with \$459.1 bn. It should come as no surprise that the companies in the sample are large, as small companies typically don't have CDS written on their debt obligations. In the subsections below I provide further details on the data sources and describe the time-series of three main variables: risk neutral survival probability, bank-run index (BRI) and distance to default (DD).

2.3.1 CDS ParSpread and the Risk Neutral Probability of Default

A CDS par-spread is the premium paid for insurance against default of a company, and can thus be used to back out market expectations about its probability of default.¹¹ I obtained data on this variable from Markit, using the universe of single-name CDS securities at a daily

¹¹It is denoted in percent, so that the insuree pays the premium times the notional amount insured. In the event of default by the company named in the contract (as defined in the contract's document clause), the insurer pays the insuree the notional amount.

Figure 2.3 1-year risk neutral PoD (index, North America)

frequency. Data is available starting in 2001, just a few years after the commencement of commercial use in this type of derivative.¹²

CDS contracts vary across four main attributes: tier, currency, document clause and tenor. For the estimation I used the most common contracts in terms of tier (SNRFOR) and currency (USD and EUR). I then computed default intensities for each available document clause at the ticker level, utilizing all available tenors (see section 2.2.1 above for further information). As a benchmark I take contracts denominated in USD and modified-restructuring (MR) document clause. Since not all companies have this particular combination, I estimated the average effect of currency and document clause on par spreads, and then used this correction for any company that lacked it. I used all available tenors to estimate each firm's default intensities, linearly completing the term structure of par spreads and finding default intensities that best match this term structure.

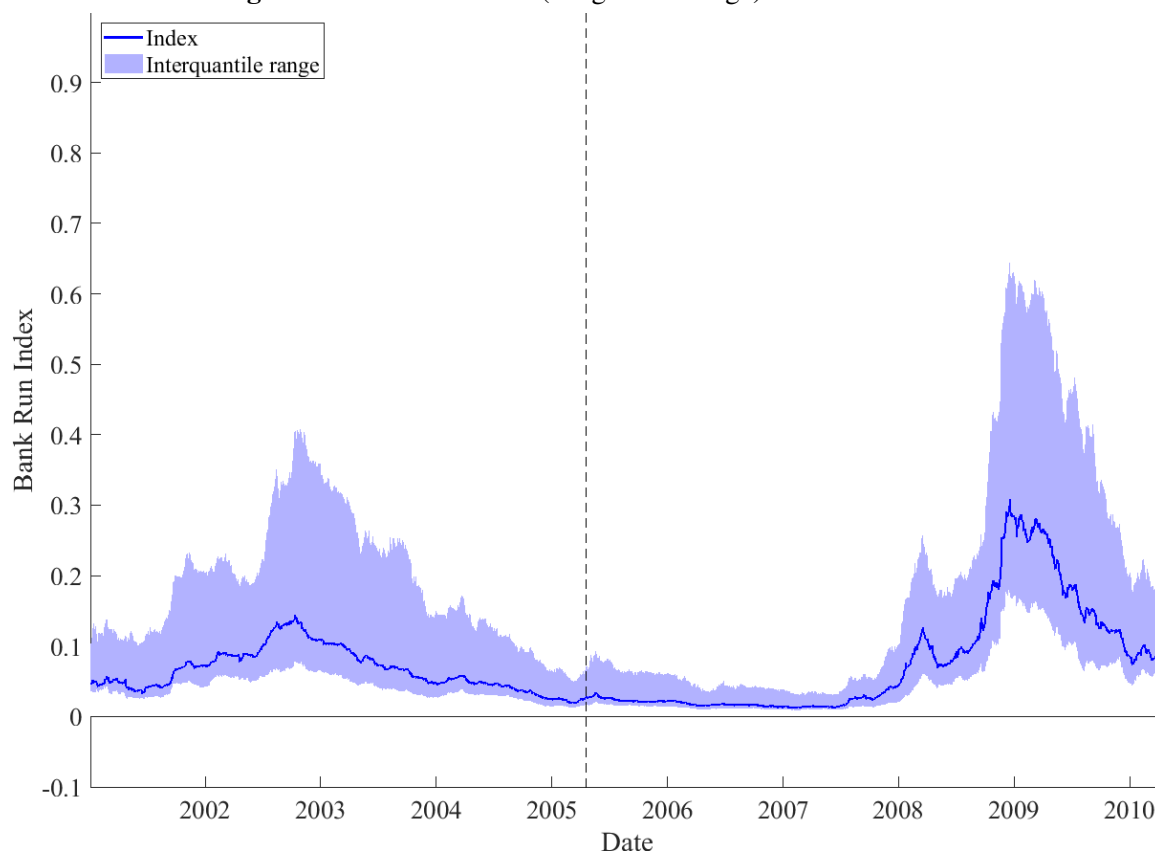
I opt to describe the data in terms of (risk neutral) PoDs rather than par spreads, because it is a more standardized variable, facilitating cross-section comparison. PoDs are strongly

¹²1997 is generally acknowledged as the year of inception of the broad CDS market (Tett, 2009).

counter-cyclical, and generally have a strong common component that reacts to market-wide events/trends (especially during market stress). The weighted average of 1-Year PoD for all North American companies in the sample (see Figure 2.3) varied between 0.04 and 0.38 in the decade 2001-2010, averaging 0.18. Its interquantile range widened in times of stress and took about 6-10 quarters to return to pre-stress levels. Moreover, throughout the decade the index was located at the lower end of the distribution, reflecting the fact that smaller companies tend to have higher PoD – which is consistent with a common explanation for Fama-French SMB factor standing in for default risk. Over the period between 2001-2007, CDS prices dropped across almost all companies around the world, except for several months before August 2002 (when American Airlines declared a technical bankruptcy). Accordingly, the implied risk neutral PoD increased. In 2005 CDS prices stabilized at the lowest pre-crisis levels, which can be seen in Figure 2.3 by observing a relatively low PoD.

Although the patterns described above generally apply for other advanced economies, some geographical differences persisted in the sample. In terms of 1-year PoD, European and British companies maintained a gap below American ones, that narrowed over time until they almost converged before the crisis (Figure B.1). A similar picture arises when limiting observations to large banks. Interestingly, this pattern does not repeat further down the term-structure. The gap in 5-year PoD remains more or less the same in the four years window around April 2005, and only narrows at the onset of the financial crisis. In terms of the difference between large banks and other entities (Figure B.2), a similar picture arises both in the U.S. and in Europe, where large banks were considered safer entities up until the financial crisis, albeit the gap between those groups narrowed slowly over time.

Finally, about two month before the legislation was enacted, we can observe a market-wide shift upward in PoD, both in Europe and in North America, for large banks and for other entities. Increments are observed both in short-term and long term insurance prices, while the BRI increased slightly. Marked spikes are observed in the BRI of the largest global banks in this short period. It would behoove us to refrain from interpreting this as the effect of insolvency netting, as other events could also be driving it. For example, in 2005, General Motors entered into difficulties, experiencing credit rating downgrades both in May and December. To the extent investors expected this and thought it could cause problems in credit markets, a market-wide increase in insurance prices would seem like a natural reaction, potentially unrelated to the structural change this chapter studies. What is pointedly clear is that investors' concerns of short-term risk increased, though only for a short period of about three months.

Figure 2.4 Bank run index (weighted average): North America

Bank Run Index

The Bank Run Index (BRI) proposed by Veronesi and Zingales (2010) is a function of risk neutral survival probabilities. The weighted average BRI for American companies (Figure 2.4) is strongly counter-cyclical; like survival probabilities, its interquantile range expands at times of stress and narrows back as credit conditions ease. A high BRI signals that investors are anxious about the ability of the company to repay its debt in the short-term. Although the word ‘bank’ in it might suggest it is not appropriate for non-bank companies, in fact nothing prevents us from computing this measure for those companies too. Non-bank companies also need to regularly replace old with new debt, albeit at lower frequency. Until mid-2007, companies that are not large banks had a higher BRI than large banks (Figure B.4), both in Europe and in the US; after that large U.S. banks had a higher BRI, whereas in Europe the indices of the two groups almost coincide. Comparing a weighted average of BRI by location (Figure B.3), until 2005 there was a gap between European/British and American companies, that narrowed over time until the three indices almost coincide by mid 2007.

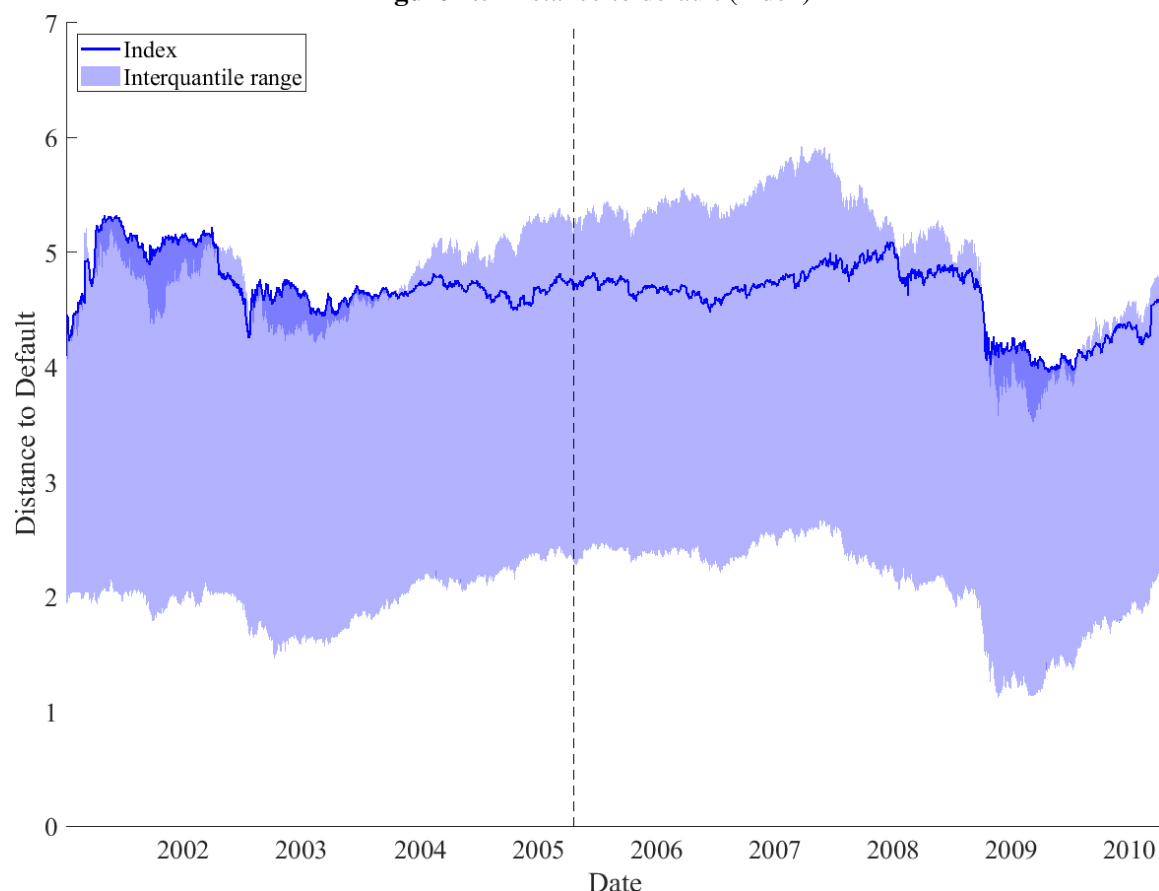
2.3.2 Distance to Default

A DD-index for all companies in the U.S. (Figure 2.5) averaged 5.79 between January 2001 and April 2007, falling to 4.61 in the last quarter of 2008. Its interquantile range was stable at roughly 2.6 over the whole period between 2001-2010. As before with PoD, the index is situated at the high end of the distribution, at times exceeding the interquantile range (indicated by a darker blue color). Since the weighting is by market capitalization, this points to the fact that larger companies are also further away from default in the Merton sense. However, this is not true across the board. In the financial sector, the index is situated in the centre or slightly lower than the median of the distribution – reflecting the fact that financial companies tend to be highly leveraged, thus tending to be closer to default. Here too, there are substantial differences across regions and sectors. Companies in the U.S. were significantly more distant from default compared to European and British companies (Figure B.5), which is the opposite of the pattern in survival probabilities. This reversal supports the claim that distance to default is not a sufficient statistic for default risk (see Bharath and Shumway, 2008; Du and Suo, 2004). Large banks were considerably closer to default compared to other companies both in the U.S. and in Europe, though in the U.S. this difference was more notable (Figure B.6).

2.4 Results

All the results presented in this section are for the BRI as the dependent variable. Although the charts below show the evolution of estimates / statistics for the full period between 2001-2010, in the text I only review the estimation window. Charts depicting the first-stage regression estimates are presented in Figures 2.6 and B.7-B.10 in the Appendix. Each data point at date t represents the output from a regression of equation 2.1 spanning 90 calendar days ending on date t ; 95%-confidence bounds are coloured around the estimates.

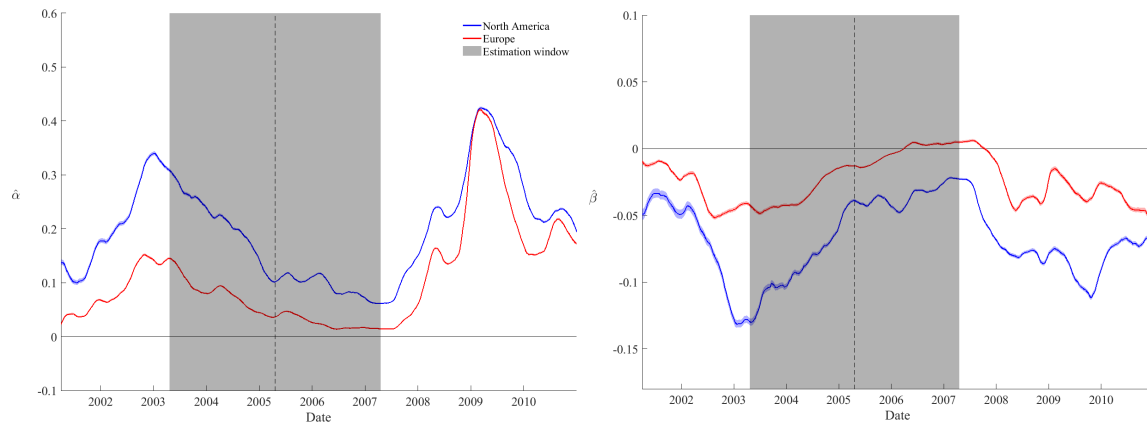
Figure 2.6 depicts the estimates of both α (Panel (a), constant) and β (Panel (b), DD) for all companies in North America (blue) and Europe (red). Estimates of α varied between 0.08 to 0.3 for North American companies, and 0.02 to 0.14 for European companies. North American companies had a higher BRI on average, though this gap narrowed slowly over time. Estimates of β were generally negative for the non-financial sector, and were either positive or very close to zero for financial companies (Figure B.7). The latter was driven mainly by large banks. An estimate of -0.08 (North American non-financial companies on 20 April 2005) indicates that a 10% increase in DD lowered the BRI by 80 basis points.

Figure 2.5 Distance to default (index)

Estimates for non-financial companies exhibited an upward trend between 2003-2007, both in Europe and in North America. This trend was stronger in North America until 2005, thence it slowed to about the same rate of increase as in Europe. The R^2 for these regressions ranged between 0.1 to 0.25 for North America, and between 0 and 0.33 in Europe (Figure B.9b). In contrast, estimates of β for financial companies hovered around zero in both economic regions, exhibiting a notably lower R^2 (Figure B.9a). In the earlier period of the estimation window estimates were even positive (Figure B.8a), contrary to elementary economic intuition. To get a positive correlation between a company's BRI and its DD means that when the company's equity price increased, investors were *more* concerned about its short-term risk of default. Together with a low R^2 coefficient, it suggests that – as far as the financial sector was concerned – there was little or no connection between a company's market value and its default risk.

One way to interpret this finding is that investors priced in implicit government guarantees to the financial sector, thereby disconnecting CDS prices from the real underlying risk of default. This would also be consistent with the fact that financial sector PoD, as implied by

Figure 2.6 First Stage Regressions: All sectors – All entities
 (a) Constant (α) (b) Systemic- β



its CDS par spreads, was lower than in other sectors in the period up to mid-2007 (Figure B.2), even though the financial sector was much more highly leveraged (reflected in its low distance to default, Figure B.6). In contrast to bond-holders, stockholders have little to gain from such a guarantee, so it is likely that equity markets were not affected by it in the same way. This finding echoes Kelly et al. (2016), who presented evidence of mispricing of aggregate tail risks using deep ‘out of the money’ options.

I tested whether the two groups (North America / Europe) had parallel trends in these average correlations, by regressing equation 2.2 for each sub-group by company type; results are reported in Table 2.2. We can see that the parallel trends assumption is rejected at the 5% significance level for the unconditional sample (all sectors), and at the 1% significance level for non-financial companies, for large banks and for entities other than large banks. The only group the assumption cannot be rejected for is financial companies, but this group exhibited almost no trend at all.

To formally provide evidence regarding a potential structural break in the trend, I ran Chow tests for each group in two ways. First, I used a moving window with a range $y \in \{\frac{1}{4}, 1\}$ years on each side of date τ , varying τ from January 2003 to July 2007. The results from these tests are presented in Figures 2.7 and B.11-B.14 in the Appendix. For non-financial companies we can observe that the Chow test-statistic reached a local maximum at, or several weeks before, April 2005 (both for $y = \frac{1}{4}$ and for $y = 1$). For large banks and for financial companies, such a peak is completely absent around those dates.

I ran further Chow tests for a fixed window (and varying sizes of PRE / POST samples), within an estimation window of one year before and after April 2005. The results of these tests are presented in Figures 2.8 and B.15-B.16. The fixed-window Chow tests complement the above results, indicating that a structural break occurred in the trend of correlations for

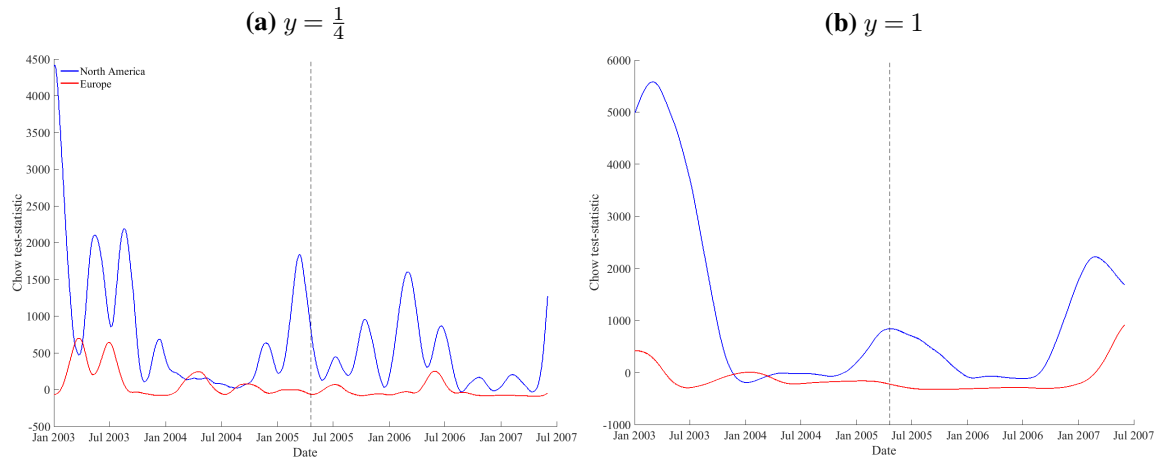
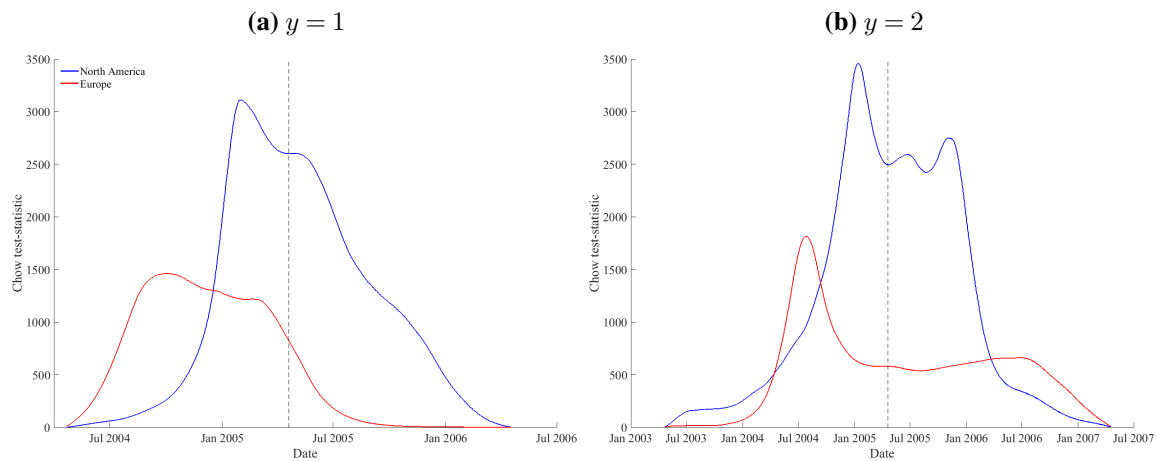
Table 2.2 Testing parallel trends (dependent variable: $\hat{\beta}_{jt}$)

Covariate / Statistic	All sectors	Financials	Non- financials	Large banks	Other entities
Constant (γ_0) (N.A.)	−0.1394*** (0.4129e − 6)	−0.0035*** (0.1130e − 6)	−0.2116*** (0.0716e − 5)	0.0317*** (0.0970e − 5)	−0.1546*** (0.4675e − 6)
Constant (γ_0) (Europe)	−0.0889*** (0.6107e − 6)	−0.0022*** (0.1672e − 6)	−0.1160*** (0.1059e − 5)	0.0191*** (0.1435e − 5)	−0.1015*** (0.4675e − 6)
Time trend (γ_1) (common)	0.0177** (0.0340e − 6)	0.0004** (0.0093e − 6)	0.0214*** (0.0059e − 5)	−0.0045*** (0.0080e − 5)	0.0197*** (0.0385e − 6)
Time trend (γ_1) (N.A.)	0.0029** (0.0139e − 6)	0.0001 (0.0038e − 6)	0.0085*** (0.0024e − 5)	−0.0018*** (0.0033e − 5)	0.0029*** (0.0157e − 6)
R^2	0.9311	0.0286	0.9529	0.2961	0.9333
T	1461	1461	1461	1461	1461

North American companies, shortly before the legislation was passed. Figures for European companies do not indicate such a break around April 2005.

2.5 Discussion

This chapter presents new evidence about the efficacy of netting rules in reducing systemic risk. It develops a new way to assess changes in systemic risk by measuring firm-level correlations between fundamentals and default risk. I call this correlation systemic- β . The measurement is operationalized by defining fundamentals via distance to default (DD, Merton, 1974) and default risk via the Bank Run Index (BRI, Veronesi and Zingales, 2010). I find that firm's systemic- β weakened (less negative) both in Europe and in North America between 2003-2007, mainly driven by non-financial companies. The systemic- β for financial companies were small in both economic regions, at times they were even positive – a fact going against elementary economic intuition. This is so because it suggests that, as far as the financial sector was concerned, before the 2007-8 crisis there was little to no connection between information about the equity value and the credit risk of companies. Echoing Kelly et al. (2016), I offer one interpretation to this puzzling finding, namely that market participants priced government guarantees into CDS prices of the financial sector. Consistent

Figure 2.7 Chow tests, fixed window: All sectors**Figure 2.8** Chow tests, moving window: All sectors

with the theory, firms' systemic- β strengthened (more negative) across the board as the crisis unfolded starting mid-2007.

Comparing the trends between North America and Europe, I bring evidence that the systemic- β for non-financial companies in North America experienced a structural break at or before April 2005 – the month in which a major change to the U.S. bankruptcy code was enacted. I find no evidence that such a structural break occurred in the European trend. This suggests that the change in American legislation affected companies in North America in a stronger way. As to financial companies, no visible effect can be ascertained. If indeed market participants priced in government guarantees to the financial sector's default insurance, this would obscure our ability to infer information from financial companies' CDS prices. By extension, it would make it difficult to determine what effect the new legislation had on the financial sector using my methodology.

If correlations between fundamentals and short term default risk (systemic- β) increased for non-financial companies, does that really cause a problem to systemic risk? I acknowledge that financial firms are commonly the ones we worry about when think of a coordination problem. However, this is a matter of severity, not of principle. The coordination problem in the financial sector is more acute because of the relatively low duration of its liabilities. Yet, in theory it applies to any company using debt finance.

It may be that by favouring some companies over others, insolvency netting decreased the threshold of fundamentals at which creditors refuse to roll-over their claims – for those companies that benefit from it. Since the maturity structure of their liabilities is typically more conducive to a coordination problem, the total effect might be good for systemic risk. However, this could not be corroborated by the data using my methodology. It may just as well be that the foregone risk-sharing benefits from a well connected financial system increased systemic risk also in the financial sector.

Now, the fact that we distinctly observe an adverse effect for companies that do not benefit directly from netting could suggest this was true also for financial companies that are excluded my sample. The companies in my sample are large, larger than the average company that is not included in it. This is because my methodology is limited to companies that are both publicly traded and have a CDS contract written on their bonds. Smaller financial companies (such as regional banks and/or money market funds) with a similar liability structure are unlikely to benefit directly from the new legislation. Indeed, many of those companies are the ones that provide short term finance to dealer banks, or are otherwise connected to them indirectly by lending to companies that are financing dealer banks. This is because many countries prohibit deposit taking institutions from investing in risky securities directly, thus compelling them to lend any excess funds they obtain from deposits to those who are allowed to do so. Even without any regulatory pressure, the fact that investing in certain asset classes requires expertise may underlie the connection between small and large banks. Thus, even if it was the case that the new legislation increased the stability of some companies at the expense of others, it might be that, overall, systemic risk was not reduced.

3

Almost Purely Endogenous

[T]he simple microeconomic idea of efficient levels of employment does not seem to be able to explain the large observed fluctuations in the level of employment over the business cycle... Neither the public's preferences about work and consumption nor the productive technology shift suddenly, as far as can be determined. Economists may acknowledge that people work harder when there is more work to do, but the macro efficiency principle does not explain why there is sometimes distinctly less work to do in the whole U.S. economy.

Robert Hall, 1980

This chapter studies a stylized macroeconomic model in which a coordination problem between producers gives rise to non-negligible fluctuations in output, and consequently asset prices – even when exogenous risk is assumed to be negligible. In an otherwise standard setting I assume a downward nominal wage rigidity, and show that this implies a multiplicity of equilibria that can be Pareto ranked. Following Frankel et al. (2003), I go on to resolve equilibrium multiplicity by using a global games approach. I then assume that exogenous variation is negligible and show that nevertheless non-negligible endogenous fluctuations in output persist in equilibrium. Finally, I extend the model to an infinite horizon OLG, and study asset pricing implications.

3.1 Introduction

In May 2007, on the eve of the financial crisis, Federal Reserve Chairman Ben Bernanke said he believed mortgage defaults would not seriously harm the economy (Bernanke et al., 2007). Two months later he maintained this line, estimating losses associated with subprime borrowing at \$50-100 billion – a small fraction of the \$56.2 trillion in U.S. household net worth (Kaiser, 2007).¹ And yet, many economists believe today that sub-prime mortgage defaults were a central cause of the financial crisis that ensued. How could such a small ‘shock’ have had such dramatic consequences?

The contrast between small shocks to their large consequences motivated a perception of risk in which market institutions are understood to have ‘amplification mechanisms’, through which innovations have increased impact (e.g. Acemoglu et al., 2012; Allen and Gale, 2004; Bernanke et al., 1994; Brunnermeier and Sannikov, 2014; Holmstrom and Tirole, 1997; Kiyotaki and Moore, 1997). Within this understanding of risk, variation in economic variables is explained via institutions / structures of the economy, rather than directly via external shocks. This falls within the purview of the substantial and developing area of ‘systemic risk’. Yet, it still remains the case that the systemic risk literature and its predecessors rely on exogenous shocks to be of not too small a size, as the amplification mechanisms put forth hitherto all have a finite signal-to-noise ratio. This immediately raises the question of how much amplification actually occurs? This is difficult to resolve empirically (Cochrane, 1994; Kocherlakota et al., 2000).²

This chapter develops a stylized general equilibrium model in which variation in fundamentals (payoff-relevant exogenous variables) can be of negligible magnitude, but nevertheless non-negligible aggregate fluctuations in output, employment and asset prices persist in equilibrium. The cause of these fluctuations is a coordination problem between monopolistic producers: since the demand for each firm’s product depends on the aggregate purchasing power of consumers (effective demand), which in turn depends on the amount of labour firms wish to employ at that level of effective demand, firms may fail to coordinate on an equilibrium in which effective demand equals the first / second best levels. In my model, this intuition is formalized in a one-period simultaneous action game, populated by agents (producers) who hold rational expectations.

Three main frictions play a part in my model: monopolistic competition, nominal wage rigidity and lack of common knowledge. In section 3.2 I show that in an economy with

¹A year later Bernanke revised this figure upwards to \$300 billion, but this includes European institutions as well (Bernanke, 2008).

²Empirical evidence of amplification is provided in the credit constraint literature, e.g. in Adrian et al. (2011, 2014, 2016); Enders and Kollmann (2010); Gabaix et al. (2007); He et al. (2017); Siriwardane (2015).

the first two frictions, there exist a continuum of equilibria: output depends on producers' expectations about other producers' intensity of production. When producers expect others to produce at a high intensity, it is optimal for all producers to do so – *and vice-versa*. Output declines can be driven by sunspots: equilibrium is a self fulfilling prophecy.

The first two frictions, and the resulting continuum of equilibria, are all in line with Keynes's General Theory (Keynes, 2018). However, while the New Keynesian approach has long incorporated monopolistic competition into its models, it is common in that literature to replace Keynes's nominal wage rigidity with nominal price rigidity (Krause and Lubik, 2007). Although price rigidities could theoretically give rise to multiple equilibria, this is not a robust phenomenon (Blanchard and Kiyotaki, 1987).³

The concept of self-fulfilling equilibria was studied extensively since the 1980's in the sun-spots literature (Azariadis, 1981; Diamond, 1982; Woodford, 1986), where it was hypothesized that business cycle fluctuations could arise from self-fulfilling beliefs about endogenous variables.⁴ The main issue with equilibrium indeterminacy in an economic model is that it is impossible to use it to provide any comparative statics and/or give policy evaluations. This chapter overcomes this issue using the global games equilibrium refinement. This preserves the "instability" properties of sun-spots in a unique equilibrium setting. In section 3.3 I follow Frankel et al. (2003) to resolve equilibrium multiplicity by introducing a third friction: lack of common knowledge. In line with the theory developed in the global games literature, some exogenous variation is required to pin down equilibrium, though I show that this variation can be of negligible size. The equilibrium in my model is a strategy (function) that prescribes a level of production to any possible observed signal about economic fundamentals. The latter are defined as the probability of extreme events: states of the world in which agents' optimal strategies are determined solely by fundamentals, independent of the actions of others.

Two stochastic events play an important role in my model: a 'good' event, in which agents' strictly dominant strategy is to produce levels that accord with a classical Pareto equilibrium, regardless of what other firms do; and a 'bad' event, in which agents dominant strategy is to produce the lowest positive levels available to them. This is the 'dominance regions' assumption. A closely related model was provided by Schaal and Taschereau-Dumouchel (2015), where non-convexities in the production function give rise to strategic

³In Blanchard and Kiyotaki (1987), the occurrence of multiple equilibria depends on the elasticity of aggregate demand to real money balances. If the latter exceeds unity, the phenomenon of multiple equilibria is precluded. Moreover, in Blanchard and Kiyotaki's model if the conditions for strategic complementarities are fulfilled, there is room for two equilibria, as opposed to the continuum of equilibria that results from a nominal wage rigidity. See Appendix C.7 for additional discussion on the macro literature.

⁴See Benhabib and Farmer (1999) for a survey of the early literature. Further discussion of more recent papers is provided below.

complementarities in the choice of capacity utilisation. In their model, dominance regions arise naturally from their setup: high (low) enough realisations of TFP shocks imply that high (low) capacity is a dominant strategy. Instead of this, I assume that there are external rare events that may hit the economy, as in Barro (2006). These events could be interpreted as an oil embargo, a natural disaster, a technological revolution or a free-trade deal.⁵

The sense in which exogenous variation is negligible is that the unconditional probability of extreme events is close to zero,⁶ while the potential payoffs/losses from these events are finite. Figure 3.1 provides a graphical example. In panel (c) I break down expected profits to their endogenous and exogenous components, showing that they differ by several orders of magnitude.⁷ In my model, the mere possibility of realization of such extreme states of the world pins down equilibrium in all other states. How does this occur?

Consider a simplified example with two agents, i and j . Each agent decides her production level (and therefore employment) in stage 1, but can observe demand for her product only at stage 2. Their decision is based on information about a state variable $\theta \in \Theta$, which is related to the probability of the two extreme events. The probability of event A (B) is an increasing (decreasing) function of θ , such that at some level $\bar{\theta} - 2\varepsilon$ ($\underline{\theta} + 2\varepsilon$) the probability of event A (B) is simply 1. If none of the events occurred in stage 2, all agents proceed to sell their goods in a monopolistic competition to consumer-workers. The information agent i possesses at stage 1 about the real state variable is captured by a noisy private signal $\theta_i = \theta + \epsilon_i$, with ϵ_i randomly drawn from a distribution with support $[-\varepsilon, \varepsilon]$. At extremely high (low) levels of the signal $\bar{\theta}$ ($\underline{\theta}$), agent i is sure that event A (B) will occur, and thus sets her production level at $\bar{Q}$ ($\underline{Q}$) independently of what j does, and vice versa for j . However, at intermediate signals, i 's decision depends on j 's action, which in turn depends on i 's action. Agent i would have to consider not only what j believes about the state θ , but also what j believes that i believes about it, and what j believes that i believes that j believes, and so on. Consider the thought process of agent i observing an intermediate signal $\theta_i = \theta^*$, and contemplating what can be inferred from this signal about j 's action. Since θ^* is observed, i knows that the real θ is distributed, conditional on θ_i , between $[\theta^* - \varepsilon, \theta^* + \varepsilon]$. It is possible, therefore, that agent j observed $\theta_j = \theta^* + 2\varepsilon$; how will j behave in that case? j needs to consider any $\theta \in [\theta^* + \varepsilon, \theta^* + 3\varepsilon]$; then j would think it is possible that agent i observed $\theta_i = \theta^* + 4\varepsilon$. In that case agent i would need to consider the distribution of $\theta \in [\theta^* + 3\varepsilon, \theta^* + 5\varepsilon]$.

⁵Barro also considers the great depression as a rare event. In my model, defining this as an external event would miss the point of the endogenous nature of the business cycle.

⁶This effectively means the probability of both events remain close to zero for most of the state-space, but increases very fast at the extremes.

⁷Observe that the gain/loss attributed to exogenous events does not rise even when their probability increases (extreme realizations of the state variable θ), because the endogenous response is already in sync with minimizing losses or maximizing gains from those events.

The pattern in the above is readily visible. Because agents don't have common knowledge, i is not certain about j 's perception of fundamentals, which means she has to consider j 's higher order beliefs. Agent i would, at some point, 'hit the boundary' $\bar{\theta}$ – in which case i knows j 's action with certainty. In this way, the extreme values $\bar{Q}$ and $\underline{Q}$ exert 'remote influence' on all agents' strategies at intermediate states (Morris et al., 1995). Section 3.3 follows Frankel et al. (2003) to show that given certain assumptions, informally sketched above, there is a unique symmetric equilibrium in which both agents adopt a strategy function $Q(\theta_i)$, $Q : \Theta \rightarrow [\underline{Q}, \bar{Q}]$; given that both agents use this strategy, it is optimal for each agent to use this strategy, i.e. it is a Nash equilibrium. In section 3.4 I show that each value $Q(\theta_i)$ is a weighted average of $\bar{Q}$ and $\underline{Q}$, with weights depending on structural parameters such as those related to competition and production technology, and the function $P_z(\theta)$ for $z \in \{A, B\}$.

Figure 3.1 Graphical example

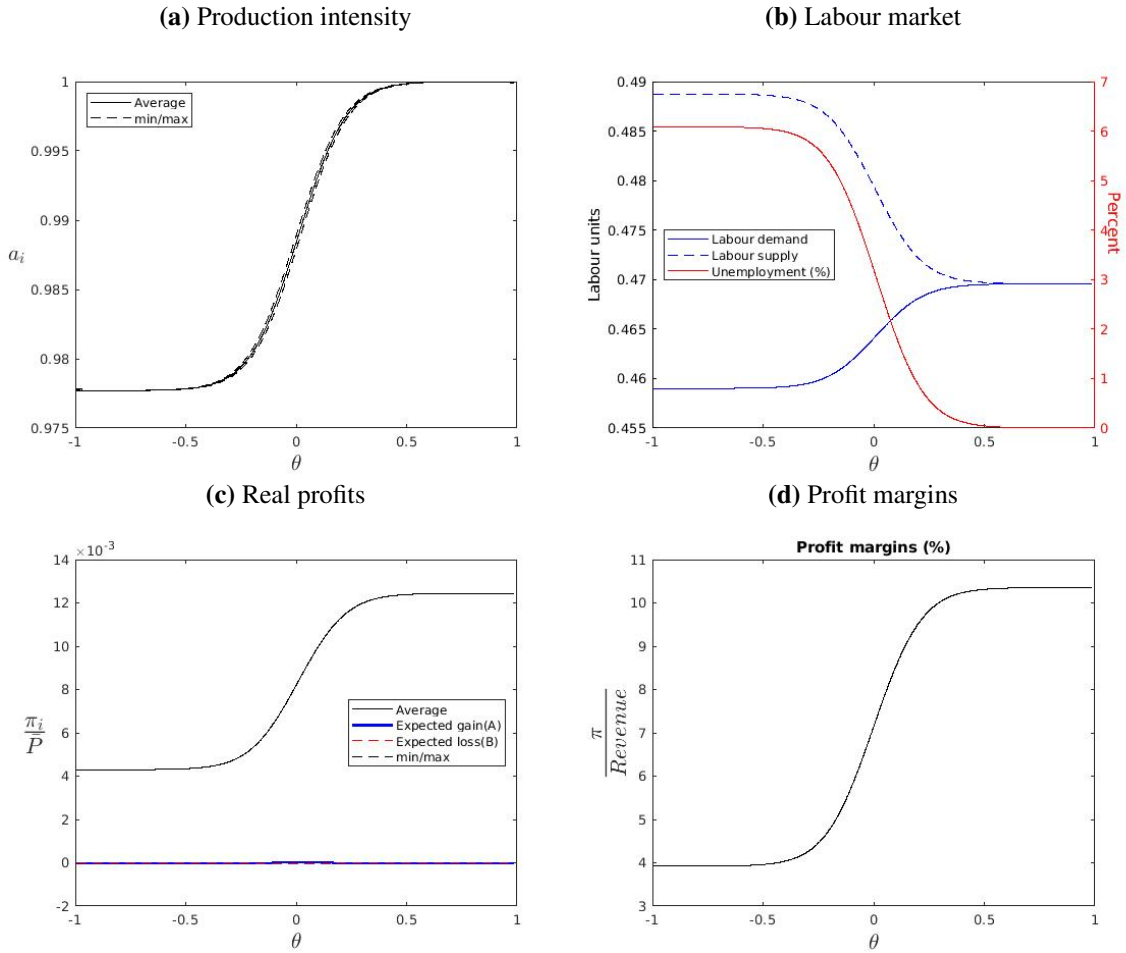


Figure 3.1 provides a graphical example of an equilibrium, calibrated to match the fall in GDP (2.33%) and rise in unemployment (6%) experienced during the last business cycle trough in the United States. It uses an exponential functional form for $\mathbb{P}_A$, increasing from 0

to 1 on $\Theta = [-1, 1]$ and rate of increase of 5. Returns to scale are assumed to be 0.99, and private signals about the fundamental distributed uniformly around the fundamental value θ , $\theta_i \sim U[\theta - 0.01, \theta + 0.01]$. Even though a value of 0.1 is a crude approximation, the shape of the function of equilibrium actions is very close to its analytical counterpart.

In 3.1(a) I plot production intensity ($\frac{Q(\theta)}{\bar{Q}}$) as a function of the state variable θ ; production intensity rises with θ , with values chosen to match the drop in (detrended) GDP experienced in the last business cycle in the U.S. (2.33%).⁸ In panel (b) I plot labour market demand and supply curves as they vary with the state of nature; by construction, labour supply intersects with labour demand only when production is at $\bar{Q}$, consistent with my assumption of downward wage rigidity. Preferences over consumption and labour are assumed to be additively separable, with the Frisch elasticity set at 1.4. Unemployment increases with θ , rising by 6% over the state-space to match the spike in U.S. unemployment rate between October 2006 and October 2009. Panels (c)-(d) show real profits and profit margins respectively; these increase with θ (which is achieved via fixed costs to production); profit margins are calibrated to match S&P500 profit margins, which dropped from around 10% to 4% during the 2007-9 financial crisis.

In section 3.4 I provide conditions under which negligible disturbances to the probabilities of exogenous events can shift equilibrium output by a non-negligible magnitude. I assume returns to scale are constant and the probability of exogenous events are symmetric. In the limit as the deviation from common knowledge goes to zero, equilibrium actions are proportional to the relative probabilities of extreme events: $Q(\theta) = \underline{Q} + (\bar{Q} - \underline{Q}) \frac{\mathbb{P}_A(\theta)}{\mathbb{P}_A(\theta) + \mathbb{P}_B(\theta)}$. The sharp increase in output around $\theta = 0$ is a direct consequence of the curvature of $\mathbb{P}_A$ and $\mathbb{P}_B$, because it also varies the proportion of $\mathbb{P}_A$ to $\mathbb{P}_A + \mathbb{P}_B$. If these assumptions are not satisfied it is difficult to provide an explicit expression for $Q(\theta_i)$, but it is easy to simulate a fixed map iteration that converges to equilibrium rather quickly. As we increase the deviation from common knowledge (larger ϵ), equilibrium output becomes a smooth straight line.

Even without knowing the exact shape of $Q(\theta_i)$, we do know that the unconditional variance of output, $Var[Q(\theta)]$, remains non-negligible even if the unconditional probability of extreme events is negligible. This is what I call an infinite signal-to-ratio amplification mechanism, which is a result of a coordination problem between producers. Demand is ex-post⁹ low because agents expect it to be so ex-ante – at that level of fundamentals – and so they respond optimally to their expectations. In other words, there is an element of self-fulfilling prophecy to the equilibrium.

⁸The value $\underline{Q}$ was chosen exogenously, but C.5 shows that this value is related to the severity of wage rigidity, so in principle it could be an endogenous variable: it would depend on the proportion of workers subject to rigid wages.

⁹After agents observe their private signals.

Two recent papers from the sun-spots literature are relevant in this respect: Angeletos and La'O (2013) (AL) and Farmer (2013). AL study a yeoman farmer economy, i.e. inhabited by agents who are both producers and consumers. Their economy exhibits a unique equilibrium in which non-fundamental shocks ('sentiment shocks': extrinsic variation in expectations that are self-fulfilling) can cause variation in aggregate production. Similarly, Farmer presents a model in which beliefs about stock market valuation (current value of the stream of dividends of firms) follow an exogenous random process, and goes on to show these beliefs are self-fulfilling and affect output. Although there are similarities in the results of these papers and mine, my model differs in several respects. A first difference is that I focus on extrinsic variation in fundamentals that is negligible, while AL and Farmer focus on non-fundamental shocks to beliefs. In this respect both papers are more in the vein of Diamond and Dybvig (1983), where beliefs could follow, or be affected by, an external process that is independent of fundamentals – yet it does affect the equilibrium. However, this means that without specifying beliefs externally the equilibrium is not unique. Instead, I follow the global-games approach of Goldstein and Pauzner (2005) and Schaal and Taschereau-Dumouchel (2015) where a unique equilibrium does not require specifying an arbitrary process for beliefs. Rather, beliefs are pinned down by specifying a parsimonious exogenous process of fundamentals.

AL are very much aware that there is a way to map sentiment shocks to fundamentals using global games, though they do not do it in their paper. The reason their focal point is a 'deviation from rationality' is due to the following claim: trying to explain business cycle fluctuations through an amplification mechanism (of fundamental shocks) which results from a coordination problem is "likely to be counter-productive", because in the class of models they are considering (Angeletos and Pavan, 2007; Morris and Shin, 2002), the volatility of technology shocks constitutes an upper bound on the volatility of output. However, my model is not nested by the class of models AL consider; in it, a coordination problem between producers does generate non-negligible fluctuations from fundamentals that have negligible volatility. This is because in my model there exist multiple equilibria under complete information, while in theirs there is a unique equilibrium. A second difference regards the labour market. In AL, Farmer and in my model, agents' production strategies are prevented from being tied by the economy-wide marginal substitution between labour and consumption. A general equilibrium model in which there is a frictionless, non-segmented labour market with flexible nominal wages, would always equate marginal revenue of firms to the nominal wage. Since the latter is a common variable across firms, their production is uniquely pinned down by the real wage and technology. That is, the labour market imposes an additional condition that rules out multiple equilibria. This is so even with nominal

price rigidities, as in the New Keynesian models. There, there could be underemployment – through there could not be unemployment.

Schaal and Taschereau-Dumouchel (2015) avoid this by assuming non-convexity in the production function; output is controlled by two decision variables, adding a degree of freedom to firms' decision. In Farmer, self-fulfilling prophecies are not due to nominal rigidities from which “prices or wages are prevented from adjusting to their equilibrium levels”. Instead, search costs play the role of disentangling the real wage from the marginal substitution between labour and consumption. In Angeletos and La'O, the choice to model a yeoman farmer economy means that the labour market is segmented, and firms' production strategy is pinned down by *local* marginal substitution between labour and consumption. The latter is different between the different segments of the labour market because consumers themselves have different information sets. Modelling the economy as consisting of yeoman farmers effectively shifts the focus to consumer confidence. In contrast to this, I show that wage rigidity gives rise to a producer-confidence problem that is able to effect significant fluctuations regardless of any consumer confidence problems. Empirically, there is mixed evidence regarding the relative importance of consumer / producer confidence.

This chapter is also related to the global games literature. Early theoretical contributions were made by Aumann (1976); Carlsson and Van Damme (1993); Monderer and Samet (1989); Morris and Shin (2002); Rubinstein (1989). For many years economists have understood that strategic complementarities are important in many real world economic interactions (Cooper and John, 1988). One of the problems in working with models that have this property is that they have multiple equilibria, thus posing a selection problem and precluding policy analysis. Carlsson and Van Damme (1993) were the first to show how relaxing common knowledge can result in selection of a unique equilibrium, provided there are ‘extreme’ states in which agents' action is independent of others' actions. In an interesting application for financial stability, Goldstein and Pauzner (2005) show that the probability measure of these extreme states can be negligible, but nevertheless equilibrium is still unique. I contribute to the theoretical literature by extending Goldstein and Pauzner's finding, and show that exogenous variation in general can be of negligible size, and yet the fluctuations that result from the coordination problem persist in equilibrium.

The rest of the chapter is arranged as follows. The next two sections flesh out the single period model – first in the common knowledge case (Section 3.2), and then without common knowledge (Section 3.3). Section 3.4 uses this model to show that endogenous fluctuations in output persist in equilibrium even though exogenous risk is assumed to be negligible. Section 3.5 extends the single period model without common knowledge to an infinite horizon OLG, and discusses asset pricing implications of the coordination problem. Section 3.6 concludes.

3.2 Model: Complete Information, Multiple Equilibria

A closed economy exists for a single period. A unit continuum of consumer-workers consume a unit continuum of differentiated goods¹⁰ $(Q_i)_{i=0}^1$ with a Dixit-Stiglitz preferences on consumption

$$\frac{1}{1-\gamma} \left(\left[\int_{i=0}^1 Q_i^{1/\rho} di \right]^\rho \right)^{1-\gamma}$$

where $\rho \in (0, 1)$ is the taste-for-variety parameter and γ the coefficient of relative risk aversion (CRRA). They supply undifferentiated labour elastically, taking prices and wages as given. In all they maximize

$$U(Q, N) = \frac{1}{1-\gamma} \left(\left[\int_{i=0}^1 Q_i^{1/\rho} di \right]^\rho \right)^{1-\gamma} - \frac{1}{1+\frac{1}{v}} N^{1+\frac{1}{v}}$$

$$\text{s.t. } \int_{i=0}^1 P_i Q_i \leq W N + \int_{i=0}^1 \Pi_i$$

where v is the Frisch Elasticity of labour supply, P_i is the nominal price of good i , N is labour, Π_i are the profits of firm i (producing good i) and W the nominal wage. Throughout this section I maintain the assumption that the nominal wage has a minimum value $\hat{W}$, to which it equals when the real wage is higher than the ratio of marginal substitution between labour and consumption $MRS_{NQ}(Q)$

Assumption 2 (Minimum nominal wage).

$$W = \begin{cases} \hat{W} & \text{if } \hat{W}/P > MRS_{NQ}(Q) \\ P \cdot MRS_{NQ}(Q) & \text{else} \end{cases} \quad (3.1)$$

$$MRS_{NQ}(Q) = -\frac{U_N}{U_Q} = A^{-\frac{1}{v}} Q^{\frac{1}{\rho v} + \gamma}$$

It is a well established empirical fact that wages exhibit little variation even when faced with significant changes in the demand for labour (Hall et al., 1980). Hall (2005) provides micro-foundations for wage rigidity. Although this assumption is rather strong, I maintain it to keep the complexity of the model at a minimum; all the results in this chapter are consistent with much weaker assumptions about wage rigidity. In Appendix C.5 I show that allowing a fraction of workers to accept lower wages doesn't change the results of this

¹⁰The assumption that there is a unit of goods can be achieved by assuming fixed costs of innovation of new products and free entry.

section qualitatively. I also show that expanding the current framework to have two periods, a first period in which the nominal wage is constrained in the above manner, and a second period in which nominal wages are flexible – while allowing firms to store goods produced in the first period at no cost and sell them in the second period – does not alter the results of this section.

A unit continuum of firms produce differentiated goods using labor as the only input, and compete in a monopolistic competition on selling it to consumer-workers. Unsold goods are worth nothing to firms. Each firm can produce Q_i with labour inputs endogenously, where Q_i is constrained to be at least $\underline{Q} > 0$, which can be arbitrarily small but positive. The production function of Q_i is given by

$$Q_i = f(N_i) = \left(AN_i\right)^\alpha \quad (3.2)$$

where A is a technology parameter and α determines returns to scale¹¹. Firm i 's profits are given by Π_i

$$\begin{aligned} \max \Pi_i &= P_i Q_i - \text{Cost}(Q_i) \\ \text{s.t. } Q_i &= \left(\frac{P_i}{P}\right)^{-\sigma} \frac{\int_{j=0}^1 P_j Q_j dj}{P} \\ P &\equiv \left(\int_{j=0}^1 P_j^{1-\sigma} dj\right)^{\frac{1}{1-\sigma}} \quad \sigma \equiv \frac{1}{1-\rho} > 1 \end{aligned}$$

where Q_i and P_i exhibit the well known relation between demand for the good and the price, and σ the elasticity of substitution between goods. Strictly speaking, a firm can choose both price and quantity, but it is clear that as long as $P_i > W f'(N_i)$ firms' optimal choice is a pair P_i, Q_i that satisfies the demand relation (Grossman and Barro, 1976, pp. 42). Thus we can write the firm's problem as reduced to maximizing nominal profits

$$\max \Pi_i = P_i \left(\frac{P_i}{P}\right)^{-\sigma} \frac{I}{P} - \text{Cost}\left(\left(\frac{P_i}{P}\right)^{-\sigma} \frac{I}{P}\right)$$

¹¹Note that for $\alpha = 1$ the firm exhibits constant returns to scale in terms of labor inputs. This is an assumption that is often made in the literature, though in this model we will assume $\alpha \neq 1$, and look at the different implications of DRS/IRS. The CRS assumption makes sense inasmuch as the production unit (or factory) reaches its most efficient method of production for a very small scale compared to overall production. This is an extreme assumption, that does not often hold in reality. Moreover, one can distinguish between returns to scale from a technological point of view, which pertains to certain conditions in the production unit, and returns to scale from a financial point of view, which could depend on fixed costs. For $FN > 0$ we could have a region in which total returns are increasing with scale, while marginal returns (marginal product of labor) are decreasing.

where $I = \int_{j=0}^1 P_j Q_j dj$ is the total nominal income in the economy. The first and second order conditions of this maximization problem are respectively

$$\begin{aligned} (1-\sigma) \left(\frac{P_i}{P}\right)^{-\sigma} \frac{I}{P} + \frac{\sigma W}{\alpha A} \left[\left(\frac{P_i}{P}\right)^{-\sigma} \frac{I}{P} \right]^{\frac{1}{\alpha}-1} \frac{1}{P_i} \left(\frac{P_i}{P}\right)^{-\sigma} \frac{I}{P} &= 0 \\ -\sigma(\sigma+1) \frac{1}{P_i^2} \left(\frac{P_i}{P}\right)^{-\sigma} \frac{I}{P} - \frac{1}{P_i} \frac{\sigma}{\alpha} \left(\frac{1}{P_i} \frac{\sigma}{\alpha} + 1 \right) \frac{\sigma W}{\alpha A} \left[\left(\frac{P_i}{P}\right)^{-\sigma} \frac{I}{P} \right]^{\frac{1}{\alpha}-1} &< 0 \end{aligned}$$

A symmetric equilibrium is one in which all firms choose identical pairs (P_i, Q_i) , so that $\forall i \ P_i = P, Q_i = Q$

$$P = \frac{\sigma}{\sigma-1} \frac{1}{\alpha} \frac{1}{A} Q^{\frac{1}{\alpha}-1} W \quad (3.3)$$

Now, the fact nominal wage is possibly stuck at a too high value implies a multiplicity of equilibria: any pair (P, Q) that satisfies eq. 3.3 is an equilibrium in which all firms behave optimally. The significance of the wage rigidity from a modelling perspective is that we are missing a condition to determine actual quantities and the real wage. Instead, equilibrium is determined by agents beliefs about other agents' actions, which then become self-fulfilling. Moreover, coordination on $Q_1 < Q_2$ for any $Q_2 \leq \bar{Q}$ is sub-optimal from a social perspective. Both firms' profits *and* consumers-workers welfare is higher at Q_2 . In other words, the continuum of equilibria can be Pareto ranked.

If technological returns to scale are weakly decreasing ($\alpha \leq 1$), there is a single equilibrium which is Pareto optimal, a pair $(P, \bar{Q})$ that equates $MRS_{NQ}(Q)$ to the real wage:

$$\underbrace{A^{-\frac{1}{\alpha}} \bar{Q}^{\frac{1}{\alpha} + \gamma}}_{MRS_{NQ}(\bar{Q})} = \alpha(1-\sigma^{-1}) A \bar{Q}^{1-\frac{1}{\alpha}} \quad (3.4)$$

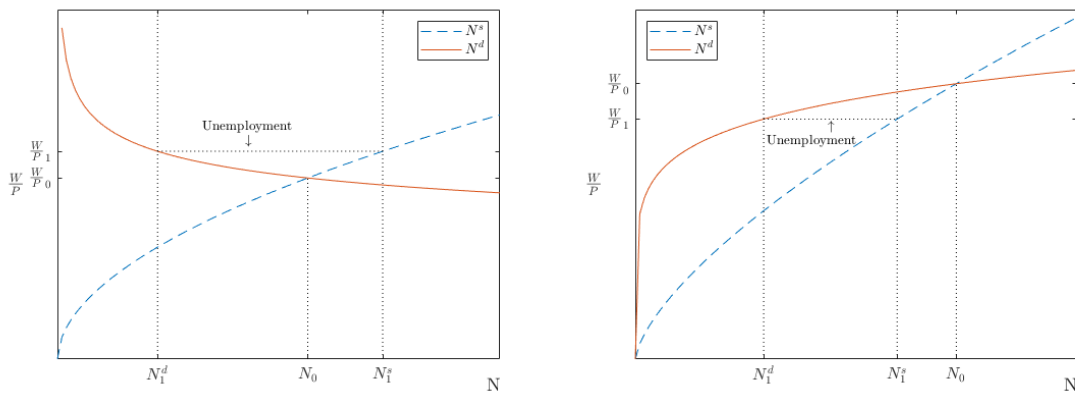
This is because if we exclude increasing returns to scale, optimal pricing is a non-increasing function of expected overall production, but the real wage is decreasing with expected production. Thus, they meet only once. This also means that, in case returns to scale are not increasing, the real wage is either a-cyclical or counter-cyclical.

Barro and Grossman (1971) study a disequilibrium model in which quantities are constrained by aggregate demand, where involuntary unemployment “does not require a rise in the real wage above the level consistent with full employment equilibrium.” In monopolistic competition models, firms face a downward demand function, and so they choose a pair (P, Q) to maximize profits. In Barro and Grossman, *both quantities and prices are constrained exogenously* (since the real wage is fixed) – which reduces the firm's decision

to a trivial one – it prefers to supply as many goods as it can sell at that level of real wages. Barro and Grossman’s case is accommodated in my model for constant returns to scale. However, in my model, firms are constrained by quantities in a more realistic way, through monopolistic competition. In my model, if returns to scale are decreasing, the real wage is required to rise above the level consistent with full employment – otherwise this would imply sub-optimal behaviour by firms. The flip-side of that is that workers on the labour market are left in the exact same place that firms are in Barro and Grossman: preferring to supply as much labour as can be absorbed by firms at that level of real wages. In Barro and Grossman, this is true more generally: both workers and firms are left with trivial choices.

Figure 3.1 illustrates the labour market demand and supply curves under the assumption of a nominal wage rigidity. $(N_0, \frac{W}{P}_0)$ designate values for a full employment equilibrium, corresponding to production of $\bar{Q}$, which is also Pareto efficient. In panel (a), firms’ technology exhibits decreasing returns, so labour demand N^d is decreasing in the real wage. If firms expect all other firms to produce at level N_1^d , then optimal pricing corresponds to a real wage which is higher than the marginal substitution between labour and consumption $MRS_{NQ}(Q)$. Effective employment (labour demand N^d) sets the real wage: the higher is employment, the lower is the real wage – i.e. the real wage is counter-cyclical. In panel (b) firms’ technology exhibits increasing returns to scale (IRS), in which case the real wage is pro-cyclical. With constant returns to scale (CRS), optimal pricing implies a fixed real wage, rationalising the example given in Barro and Grossman. Observe that beliefs about Q that are higher than $\bar{Q}$ are not self-fulfilling, because they are inconsistent with our assumption of downward nominal wage rigidity. Importantly, nominal wage rigidity implies that the labour market doesn’t clear, so for any other belief aside of $\bar{Q}$, there is involuntary unemployment.

Figure 3.1 The labour market with nominal wage rigidity, comparative statics
 (a) Decreasing returns to scale (b) Increasing returns to scale



We saw that this model admits a multiplicity of equilibria, determined by expectations. This multiplicity is a result of the nominal wage rigidity. Nominal wage flexibility, which is assumed in classical economics, is the same as imposing labour market clearing – i.e. an additional condition that determines the real wage. Keynes (2018) departs from this assumption (the second postulate), and instead makes the assumption that nominal wages are sticky, so that expectations play a role in determining equilibrium.

There is an extensive literature, post-Keynes, that discusses how wage rigidities can imply insufficient aggregate demand, dating back to Patinkin (1948) and Clower (1965). This was termed ‘disequilibrium theory’ by Barro and Grossman (1971) and Grossman and Barro (1976), to reflect the fact that not all markets clear. In retrospect, Barro (1979) expressed the significance of the role of “expectations in macro analysis”, but he downplays the effects of rational departure from the second postulate of classical economics, saying that it “does not imply that levels of employment would differ significantly from the (efficient) values that would have been attained under flexible wages.”¹² He saw the role for expectations as important in other contexts.¹³

However, Barro has in mind a model in which the reasons for output being too low are external. The role of expectations in determining equilibrium is absent, and equilibrium multiplicity is not treated. In my model, downward nominal wage rigidity implies a continuum of equilibria with significantly different levels of output and employment. The reason output is too low is internal – a coordination problem that arises due to the nominal wage rigidity.

It is important to emphasize that nominal wage rigidity is not meant to be taken literally.¹⁴ It is not the aim of this work to provide an explanation of wage rigidities; other works have considered micro-foundations for this friction (see e.g. Hall, 2005). Moreover, it is a custom to assume nominal price rigidities in New-Keynesian models, to which there is mixed evidence (Baudry and Le Bihan, 2004). What would happen in the current model if we replace the nominal wage rigidity with a nominal price rigidity?

Let us assume that P is fixed at some arbitrary value. Firms can freely decide quantities, and workers able to undercut each other freely. The question we are interested in is: are

¹²And also: “Notably, the approach suggests that such market features as sticky wages or the apparent non-price, quantity rationing associated with layoffs would be of secondary interest in analyses of business cycles.”

¹³“Thanks especially to the work of Robert Lucas, we have a much better idea of the significance of well-informed private expectations in macro analysis... This significance arises in at least three areas: positive analyses of the effects of monetary and other shocks on economic activity; analyses of the role of government policies; and evaluations and carrying out of econometric estimation...”

¹⁴As in Barro: “Presumably, sticky wages or prices are not intended to be taken literally as the source of private sector inefficiency. The underlying problem must reflect some deeper economic elements, such as imperfect information about the present or future, factor mobility costs, or some types of significant transaction costs.”

there self-fulfilling rational-expectation beliefs about quantities that are lower than $\bar{Q}$? Can firms have an ex-ante belief that other firms will produce less than $\bar{Q}$ which is consistent with the optimality of that action given workers' supply of labour? The answer is no. If the nominal wage is flexible, the real wage will always be equal to the marginal substitution between labour and consumption, due to workers' optimality. Given the real wage, optimal production decisions are to produce $\bar{Q}$.

It is also important to note that wage rigidities are not necessary to achieve these results. The important thing is to find a consistent way to close the model with self-fulfilling beliefs when the labour market clearing condition is abandoned. For example, Schaal and Taschereau-Dumouchel (2015) assume non-convexities in the production function, adding a degree of freedom to firms' decisions. Angeletos and La'O (2013) study a yeoman farmer economy, i.e. inhabited by agents who are both producers and consumers. They analyze the case with monopolistic competition¹⁵ (agents want to consume goods produced by other agents), flexible prices / wages and no common knowledge. In stage 1 agents form expectations and choose production. In stage 2 each agent is matched with another, and the pair trades. In their model, the equilibrium is affected by non-fundamental shocks ('sentiment shocks': extrinsic variation in expectations that are self-fulfilling) that cause variation in aggregate production. Farmer (2013) attains similar results by assuming search costs in the labour market. In his model, search costs play the role of disentangling the real wage from the marginal substitution between labour and consumption. My model differs from the latter two papers in that I do not assume that the evolution of beliefs is exogenous. Instead, I apply global games theory to resolve equilibrium multiplicity, replacing the assumption of exogenous beliefs with an exogenous process for payoff relevant variables (fundamentals).

In sum, in this section I presented a macroeconomic model in which a nominal wage rigidity results in a coordination problem between producers. In the next section I impose additional structure on the information available to agents, which implies beliefs are not arbitrary. This significantly weakens the assumption about beliefs. I apply global games methodology to resolve equilibrium, and show that output, employment and asset prices fluctuate even in face of negligible external risk.

3.3 Model: Incomplete Information, Unique Equilibrium

In this section I introduce standard global games assumptions to the model above, and prove that in this setup equilibrium exists and is unique. In equilibrium, optimal strategies vary

¹⁵Monopolistic competition is already implied by saying that an economy consists of yeoman farmers, because perfect competition in such an economy is equivalent to the trivial case of autarky.

between the minimum $\underline{Q}$ and the maximum $\bar{Q}$, and depend on the probability of exogenous events. In the next section I show that even when the unconditional probability of exogenous events is negligible, there still remains non-negligible variation in the optimal production strategy over the state space.

Consider the single period model of a closed economy, introduced in the last section. In that model, I showed that a wage rigidity and monopolistic competition open the door to self-fulfilling beliefs: if all producers believe that all the others will produce some $Q \in [\underline{Q}, \bar{Q}]$ and price it in accordance with equation 3.3, then it is indeed optimal for an individual producer to produce Q .

The global games literature taught us that the above phenomenon is an artefact of a very strong assumption, namely that of common knowledge. If we deviate even slightly from such perfect agreement, equilibrium uniqueness is restored. The standard global games approach to building a complex information structure parsimoniously is to introduce heterogeneity in agents' information sets via private noisy signals about the economy's fundamentals. The latter are defined as (possibly stochastic) payoff relevant parameters. In my model, all exogenous variation enters the economy in the form of two (mutually exclusive) stochastic events: A and B . In event A , the economy hit by a demand shock: it is integrated with a neighbouring economy that has flexible wages. This resolves the insufficient demand potentially faced by the firm, guaranteeing a market for the firm's goods regardless of other firms' actions.¹⁶ In state B , a disaster hits the economy, destroying all produced goods above $\underline{Q}$.

Events A and B are stochastic, mutually exclusive, and may arrive after firms choose production but before the goods-market opens. Their probability depends on an imperfectly observed state variable $\theta \sim g$, where g is a continuous probability density function (pdf) with a bounded support $\Theta \subset (-\infty, \infty)$, symmetric about 0. Formally, event $z \in \{A, B\}$ depends on a Bernoulli random variable 1_z with pdf $\mathbb{P}_z(\theta)$, which is a continuous and increasing (decreasing) function of the state variable θ for $z = A$ ($z = B$).

Firms are run by entrepreneurs who maximize expected real profits. At stage 1, they observe a noisy signal $\theta_i = \theta + \varepsilon_i$, and then they choose production simultaneously. ε_i is distributed uniformly, $\varepsilon_i \sim U[-\epsilon, \epsilon]$, and is independent of θ :

$$\mathbb{P}[\theta = x] = \mathbb{P}[\theta = y] \quad \forall x, y \in [\theta_i - \epsilon, \theta_i + \epsilon]$$

All ε_i are pairwise independent. The assumption about noisy signals furnishes the model with imperfect common knowledge. I assume the departure from common knowledge is

¹⁶There are many possible ways to define event A, e.g. adding a period between 1 and 2 in which firms can produce additional goods at flexible wages, but assuming the market for goods opens only at stage 2.

small: $\epsilon \rightarrow 0$. Thus, entrepreneurs choose $Q_i(\theta_i)$ to maximize:

$$\mathbb{E} \left[\frac{\Pi_i(\theta)}{P(\theta)} \mid \theta_i \right] = \frac{1}{2\epsilon} \int_{\theta_i-\epsilon}^{\theta_i+\epsilon} \left((1 - \mathbb{P}_A(\tilde{\theta}) - \mathbb{P}_B(\tilde{\theta})) Q_i(\theta_i) \left(\frac{Q_i(\theta_i)}{I(\tilde{\theta})/P(\theta)} \right)^{-\frac{1}{\sigma}} \right. \quad (3.5)$$

$$\left. + \mathbb{P}_A(\tilde{\theta}) Q_i(\theta_i) \left(\frac{Q_i(\theta_i)}{Q} \right)^{-\frac{1}{\sigma}} + \mathbb{P}_B(\tilde{\theta}) Q_i(\theta_i) \left(\frac{Q_i(\theta_i)}{Q} \right)^{-\frac{1}{\sigma}} - \frac{Cost(Q_i)}{P(\theta)} \right) d\tilde{\theta} \\ \frac{I(\theta)}{P(\theta)} = \left[\frac{1}{2\epsilon} \int_{\tilde{\theta}=\theta-\epsilon}^{\theta+\epsilon} Q(\tilde{\theta})^{1-\sigma^{-1}} d\tilde{\theta} \right]^{\frac{1}{1-\sigma^{-1}}} \quad (3.6)$$

where Q_i designates the amount of goods produced by firm i , $P_i(Q_i, \theta, P) = \left(\frac{Q_i(\theta_i)}{I(\tilde{\theta})/P(\theta)} \right)^{-\frac{1}{\sigma}}$ the inverse demand function, P the price index and I/P the real income index. Firms choose Q_i upon observing the signal, and then choose the highest price at which they can sell that produce – *given all other agents' production*. That is, like in Angeletos and La'O (2013), prices are chosen at a second stage so that all goods are sold and all money tokens used.¹⁷

Unlike the former section, production intensity and prices vary in the cross-section. In the case where there is no common knowledge, the average price is set as part of the fixed mapping of prices and quantities. That is, there is a pair of functions $(P(\theta), Q(\theta))$ that constitute a fixed map into themselves. Since Dixit-Stiglitz demand functions depends on the ratio of individual firm's price and the price index – the price index is indeterminate, in the same way that nominal variables are indeterminate in classical models – only quantities and the real wage matter. This was not a problem in the case of common knowledge because there was no cross-section variation. In order to simplify this, I make the following assumption:

Assumption 3. *The price index $P(\theta)$ follows the common knowledge benchmark, and is given by*

$$P(\theta) = \frac{\sigma}{\sigma-1} \frac{1}{\alpha} \frac{W}{A} \left(I(\theta)/P(\theta) \right)^{\frac{1}{\alpha}-1}$$

This assumption is equivalent to making an assumption about the velocity of money (see Appendix C.6.3 for further discussion of this point). Given this price index, firm i 's optimal quantity at signal θ_i satisfies the first order condition:

$$\pi_i(\theta_i) = \quad (3.7) \\ (1 - \sigma^{-1}) \frac{1}{2\epsilon} \int_{\tilde{\theta}=\theta_i-\epsilon}^{\theta_i+\epsilon} \left((1 - \mathbb{P}_A(\tilde{\theta}) - \mathbb{P}_B(\tilde{\theta})) \left(\frac{Q_i}{I(\tilde{\theta})/P(\tilde{\theta})} \right)^{-\frac{1}{\sigma}} + \right.$$

¹⁷Recall that firms never wish to withhold goods from the market, because the price elasticity of demand σ is larger than unity.

$$\mathbb{P}_A(\tilde{\theta}) \left(\frac{Q_i(\theta_i)}{\bar{Q}} \right)^{-\frac{1}{\sigma}} + \mathbb{P}_B(\tilde{\theta}) \left(\frac{Q_i(\theta_i)}{\underline{Q}} \right)^{-\frac{1}{\sigma}} - \left(\frac{Q_i}{I(\tilde{\theta}/P(\tilde{\theta}))} \right)^{\frac{1}{\alpha}-1} d\tilde{\theta} = 0$$

On the one hand, for high enough $\mathbb{E}[\mathbb{P}_B(\theta)|\theta_i]$, firm i 's optimal quantity is $\underline{Q}$ independently from other firms' actions, since expected revenues plummet while costs are unaffected by the event. On the other hand, for a high enough $\mathbb{E}[\mathbb{P}_A(\theta)|\theta_i]$, firm i 's optimal strategy is to produce $\bar{Q}$ even if all other firms produce at minimum intensity. At intermediate signals, however, the firm's strategy depends on other firms' strategies. This in turn means a symmetric equilibrium is one in which all firms follow $Q(\theta)$, a strategy that prescribes playing $Q(\theta)$ at signal θ , and $\pi_i(\theta) = 0$ for any possible θ . We can write the value of optimal $Q(\theta)$ as an implicit expression, i.e. depending on the real income index (which in turn depends on production decisions of firms):

$$Q(\theta_i) = \left[\frac{\int \left((1 - \mathbb{P}_A(\tilde{\theta}) - \mathbb{P}_B(\tilde{\theta})) \left(\frac{1}{I(\tilde{\theta}/P(\tilde{\theta}))} \right)^{-\frac{1}{\sigma}} + \mathbb{P}_A(\tilde{\theta}) \left(\frac{1}{\bar{Q}} \right)^{-\frac{1}{\sigma}} + \mathbb{P}_B(\tilde{\theta}) \left(\frac{1}{\underline{Q}} \right)^{-\frac{1}{\sigma}} d\tilde{\theta} \right)^{\frac{1}{1/\sigma+1/\alpha-1}}}{\int \left(\frac{1}{I(\tilde{\theta}/P(\tilde{\theta}))} \right)^{\frac{1}{\alpha}-1} d\tilde{\theta}} \right]$$

I now go on to show existence and uniqueness of equilibrium in this setting. At a given signal θ_i , player i faces a distribution of possible states $\tilde{\theta}|\theta_i \sim U(\theta_i - \epsilon, \theta_i + \epsilon)$. At each state $\tilde{\theta}$, player i faces a distribution of other players' actions, given by $\left(Q(\tilde{\theta}) \right)_{\tilde{\theta} \in (\tilde{\theta}-\epsilon, \tilde{\theta}+\epsilon)}$. Thus, at each state $\tilde{\theta}$ the composite good (real income) is determinant, with a trivial CDF $F_{\tilde{\theta}}(q)$

$$F_{\tilde{\theta}}(q) = \mathbb{P}(I(\tilde{\theta})/P(\tilde{\theta}) \leq q) = 1_{I(\tilde{\theta})/P(\tilde{\theta}) \leq q} \quad (3.8)$$

This then means that at signal θ_i , firm i faces the following CDF about real income:

$$F_{\theta_i}(q) = \mathbb{P}(I(\theta)/P(\theta) \leq q) = \frac{1}{2\epsilon} \int_{\tilde{\theta}=\theta_i-\epsilon}^{\theta_i+\epsilon} F_{\tilde{\theta}}(q) d\tilde{\theta} \quad (3.9)$$

The first property we require in order to show uniqueness of equilibrium, is provided in the following Lemma.

Lemma 1 (Strategic complementarities). *Given returns to scale are decreasing or constant ($\alpha \leq 1$), firms have strategic complements. Let $a_i : \Theta \rightarrow [0, 1]$ be the strategy of firm i , specifying for each signal a number $a_i(\theta) = \frac{1}{\bar{Q}-\underline{Q}} \left(Q_i(\theta) - \underline{Q} \right) \in [0, 1]$. Let $a = (a_i)_{i=0}^1$ be the strategy profile of all players, and write $a \geq \hat{a}$ if $F_{\theta_i}(q) \leq \hat{F}_{\theta_i}(q)$, for all θ_i and for all*

q .¹⁸ Then we have

$$\frac{\partial \Pi_i(\theta_i, a)}{\partial Q_i} \geq \frac{\partial \Pi_i(\theta_i, \hat{a})}{\partial Q_i} \quad (3.10)$$

i.e. when firms face higher functions a , their incentive to increase their action increases.

The proof is provided in Appendix C.1. Strategic complementarities is the key condition for equilibrium uniqueness. In addition, firms' payoff functions are continuous in all arguments (their signal, other players' and their own strategy), increasing in the state θ , and are Lipschitz continuous in their own action while the incentive to raise their action is Lipschitz continuous in others' actions. Finally, the assumption about extreme events prescribes that firms' optimal strategy is the highest (lowest) possible at the high (low) enough signals, regardless of other firms' strategies. Following Frankel et al. (2003) we have

Proposition 3 (Unique equilibrium). *There exists a unique symmetric equilibrium in which all firms play the same strategy profile $a_i : \Theta \rightarrow [0, 1]$; it is (effectively) the only one surviving iterated deletion of strictly dominated strategies. Moreover, this strategy profile is continuous and monotonically increasing in the signal θ .*

The proof is provided in Appendix C.2. In this section I have extended the model from the case of common-knowledge to one in which there is no common knowledge about the state variable θ , and showed that in this case there is a unique symmetric equilibrium. In the next section I assume the unconditional probability of exogenous events is negligible, and show that there still remains non-negligible variation in production strategies.

3.4 Almost Purely Endogenous

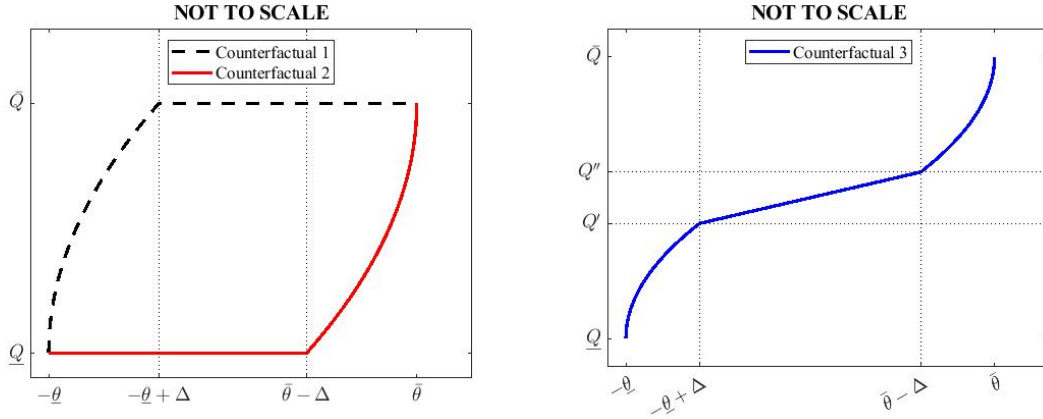
This section provides an analytical example of an equilibrium in which exogenous risk is negligible, but nevertheless risk persists in the economy through the endogenous actions of agents. Recall that Proposition 3 tells us that there exists a unique function $Q : \Theta \rightarrow [\underline{Q}, \bar{Q}]$, mapping every possible signal to a single production plan. Proposition 3 tells us that this function must be increasing and continuous, but does not tell us what is the exact shape of the function; this can only be ascertained in simulation of the fixed map. In general it could be the case that $Var[Q]$ is negligible, and therefore, in fact, there might be no fluctuations in

¹⁸It is evident that when $a(\theta) \geq \hat{a}(\theta) \ \forall \theta$, the latter property is satisfied.

Figure 3.1 Potential equilibria with negligible variance of $Q(\theta)$

(a) Examples 1-2

(b) Example 3



equilibrium.¹⁹

$$Var[Q] = \int_{\theta \in \Theta} \left(Q(\theta) - \mathbb{E}Q(\theta) \right)^2 d\theta \quad (3.11)$$

That would be the case if $Q(\theta) \approx \mathbb{E}Q(\theta)$ everywhere but a negligible subset of the state-space Θ . Examples are given in Figure 3.1. (a) shows cases where the optimal strategy is one in which agents produce the maximum (minimum) everywhere but the interval $[\bar{\theta}, \bar{\theta} + \Delta]$ ($[\bar{\theta}, \bar{\theta} - \Delta]$). If Δ is small, $Var[Q]$ would be negligible. Similarly, (b) shows a more general case where the optimal strategy increases very fast up to Q' , then increases very slowly to Q'' and finally increases to $\bar{Q}$ very fast. If both Δ and $Q'' - Q'$ are small, the variance would be negligible.

The way I exclude anomalies 1-2 is by imposing symmetry on the exogenous risk, which would imply the optimal strategy has to pass through the mid-point $\frac{\bar{Q} - \underline{Q}}{2}$. Recall that exogenous risk enters the model through events A and B, corresponding to a ‘good’ and ‘bad’ events respectively. In addition to taking exogenous risk to zero, I assume symmetry of the probabilities of events A and B. Together, these assumptions may be called ‘Symmetric tail events’:

Assumption 4 (Symmetric tail events). *The following conditions are satisfied:*

1. *Symmetry:* $\mathbb{P}_A(\theta) = \mathbb{P}_B(-\theta) \quad \forall \theta \in \Theta$
2. $\mathbb{P}_A$ is continuous, twice differentiable and strictly increasing for $\theta \in [\bar{\theta}, \bar{\theta}]$. It receives zero at $\underline{\theta}$, $\mathbb{P}_A(\underline{\theta}) = 0$ and 1 at $\bar{\theta}$, $\mathbb{P}_A(\bar{\theta}) = 1$.

¹⁹Note that I didn’t specify the density function of θ . My results follow through for any arbitrary process with a proper density $g(\theta) > 0 \quad \forall \theta \in \Theta$.

3. *Rate of increase:* $\frac{\partial P_A}{\partial \theta}(\tilde{\theta}) \geq \frac{\partial^2 P_A}{\partial \theta^2}(-\tilde{\theta})$ for any $\tilde{\theta} \in \Theta$
4. $\int_{\tilde{\theta} \in \Theta} \mathbb{P}_A(\tilde{\theta})g(\tilde{\theta})d\tilde{\theta} \rightarrow 0$ and $\int_{\tilde{\theta} \in \Theta} \mathbb{P}_A(\tilde{\theta})d\tilde{\theta} \rightarrow 0$

The probability of event A is increasing in θ (Assumption 4.2): higher states of the world imply higher probability of A occurring. This means, due to assumed symmetry (Assumption 4.1), that the probability of event B is decreasing in θ . Increasing probability of event A is standard for global games. It means that, all else equal, agents' payoffs increase with θ . Assumption 4.3 means that the rate of increase of P_A is at least as large as the rate of increase of its derivative, for any two states $\tilde{\theta}, -\tilde{\theta}$. This assumption is a sufficient condition for the equilibrium fixed map to be convex for states $\theta < 0$ and concave $\theta > 0$. Assumption 4.4 amounts to setting the unconditional probability of each of the events to zero, while making sure that the way in which these events distort agents' endogenous strategies is symmetric.

I make a further assumption about technology that imply linear best-responses, which is needed to establish the symmetry of the equilibrium strategy (Lemma 3).

Assumption 5 (Technology). *Technology exhibits constant returns to scale (CRS), $\alpha = 1$.*

Lemma 2 (Linear best-responses). *Given assumptions 4-5, agents' best responses are linear, and are given by*

$$Q(\theta) = \int_{\tilde{\theta}=\theta-\epsilon}^{\theta+\epsilon} \left((1 - \mathbb{P}_A(\tilde{\theta}) - \mathbb{P}_B(\tilde{\theta}))Q(\tilde{\theta}) + \mathbb{P}_B(\tilde{\theta})\bar{Q} + \mathbb{P}_A(\tilde{\theta})\underline{Q} \right) d\tilde{\theta} \quad (3.12)$$

All proofs are provided in Appendix C.3. Linearity of best responses, together with Assumption 4.1 on symmetry, imply that the equilibrium strategy is symmetric about $\theta = 0$, and specifically that at $\theta = 0$ agents optimal decision in equilibrium would be to produce the average of the maximum quantity and minimum quantities $\frac{\bar{Q}-Q}{2}$. This excludes the possibility of anomalies 1-2.

Lemma 3 (Symmetry of the equilibrium strategy). *The optimal strategy $Q(\theta)$ satisfies $Q(\theta = 0) = \frac{\bar{Q}-Q}{2}$.*

With this result, anomalies 1-2 are excluded. But this result is theoretically consistent with anomaly 3: the optimal strategy would increase very fast very close to $\frac{\bar{Q}-Q}{2}$ at some $\underline{\theta} + \Delta$ and then increase slowly until $\bar{\theta} + \Delta$, passing exactly through the mid-point at $\theta = 0$, only to increase very fast again at $\bar{\theta} - \Delta$. In order to exclude this scenario I assume that the probability of event A it is not decelerating (non-negative second derivative).²⁰

²⁰In fact, a finite upper bound would be sufficient.

Assumption 6 (Non-deceleration of tail risk). $\mathbb{P}_A$ is weakly accelerating (non-negative second derivative), $\frac{\partial^2 \mathbb{P}_A}{\partial \theta^2} \geq 0$. Let θ_0, θ_1 be states such that $\theta_1 > \theta_0$. Then this then implies that $\forall \Delta > 0$:

$$\mathbb{E}[\mathbb{P}_A \mid \theta_1 + \Delta] - \mathbb{E}[\mathbb{P}_A \mid \theta_1] \geq \mathbb{E}[\mathbb{P}_A \mid \theta_0 + \Delta] - \mathbb{E}[\mathbb{P}_A \mid \theta_0] \quad (3.13)$$

Lemma 4 (Acceleration of the fixed-map). The equilibrium strategy $Q : \Theta \rightarrow [\underline{Q}, \bar{Q}]$ defined by equation 3.12 is (weakly) accelerating $\forall \theta \in (\underline{\theta}, 0)$. Let $\theta_0, \theta_1 \in (\underline{\theta}, 0)$ be two states such that $\theta_1 > \theta_0$. Then we have

$$\frac{\partial Q}{\partial \theta}(\theta_1) \geq \frac{\partial Q}{\partial \theta}(\theta_0) \quad (3.14)$$

Non-decelerating tail risk implies that the equilibrium strategy cannot decelerate $\forall \theta \in (\underline{\theta}, 0)$. Together with the fact that it has to pass through $\frac{\bar{Q} - \underline{Q}}{2}$ at $\theta = 0$, this ensures volatility of production is strictly positive.

Proposition 4 (Volatility of output and employment). In equilibrium, aggregate output is volatile. The volatility of employment is a consequence of that. More formally, let $\chi(x)$ be the set of all states such that $\mathbb{P}_A(\theta) > x$, $\chi(x) = \{\theta : \mathbb{P}_A(\theta) > x\}$, and denote $|\chi(x)|$ the measure of this set. Since $\mathbb{P}_A$ is a monotonic function, $|\chi(x)| = 1 - \frac{\mathbb{P}_A^{-1}(x) - \underline{\theta}}{|\Theta|}$, where $\mathbb{P}_A^{-1}$ is the inverse of $\mathbb{P}_A$. Then in the limit as $|\chi(x)| \rightarrow 0$ for any $x > 0$,

$$\lim_{|\chi(x)| \rightarrow 0} \text{Var}[Q] = \int_{\theta \in \Theta} \left(Q(\theta) - \mathbb{E}Q(\theta) \right)^2 d\theta \geq \frac{1}{12} [\bar{Q} - \underline{Q}]^2 \quad (3.15)$$

This shows that it's not necessary to resort to deviations from rationality in order to obtain non-negligible fluctuations in macro variables in face of negligible variation in external risk. This contrasts with AL's assertion that, as the volatility of fundamentals vanish, so will the volatility in output. What accounts for the difference in results?

As prefaced above, AL consider the class of models studied in Morris and Shin (2002) in which there is a unique equilibrium in complete information, whereas my model has a continuum of equilibria. In my model, the expected strategic complementarity varies with the fundamentals: it is close to 1 everywhere apart from a negligible part of the state-space Θ . That is, its unconditional mean is just about 1, while its unconditional volatility is negligible. To compare my model and Morris and Shin (2002), I shall abuse my notation to match theirs. The degree of strategic complementarity, $r(\theta)$ is proportional to the probabilities of extreme

events A and B

$$r(\theta) = 1 - \mathbb{P}_A(\theta) - \mathbb{P}_B(\theta) \quad (3.16)$$

Firms' payoff in my model (with CRS and close to perfect competition), conditional on state θ , is

$$u_i(a, \theta) = -r(\theta) \left[a_i - \underbrace{\left(\int_j a_j dj \right)}_{\bar{a}}(\theta) \right]^2 - (1 - r(\theta)) \left(a_i - Z(\theta) \right)^2 |Z(\theta)| \quad (3.17)$$

This loss function is a reduced form that may arise in many economic environments. It has two terms: one reflecting a motive to cooperate (if $r(\theta) > 0$) and one reflecting external motive – wanting to be close to the fundamentals $Z(\theta)$, which is a trinomial variable

$$Z(\theta) = \begin{cases} 1 & \text{w.p. } P_A(\theta) \\ -1 & \text{w.p. } P_B(\theta) \\ 0 & \text{w.p. } 1 - P_A(\theta) - P_B(\theta) \end{cases} \quad (3.18)$$

Firms minimize expected loss. The resulting optimal strategy in my model is effectively

$$a_i(\theta_i) = \mathbb{E}_i[r(\theta)\bar{a}(\theta)] + \mathbb{E}_i[P_A(\theta)] - \mathbb{E}_i[P_B(\theta)] \quad (3.19)$$

In this formulation it is no longer the case that $Var(\bar{a}) \leq Var(Z(\theta)\theta)$, as in AL's model: a change in θ almost always implies an infinitesimal change in the payoff relevant variable, but due to the coordination motive, a non-negligible change in agents' actions $\bar{a}$ could occur, if it is sufficiently close to the 'centre' of the support of θ . The way fundamentals $Z(\theta)$ enter the loss function in a special way, facilitating dominance regions. This is obviously a stark case, and is presented in this way to be consistent with the macro model of this chapter and to drive the point home. By allowing a_i to take the form of an S -shape, it starkly shows the effect of the self-fulfilling prophecy as an amplification mechanism. However, even if we wrote $Z(\theta) = \theta$ as in Morris and Shin (2002), and dispensed with $|Z(\theta)|$ in eq. 3.17, we would have gotten excess volatility *over* fundamentals, whether these be taken as $(1 - r(\theta))\theta$ or simply as θ . In the latter case, the amplification mechanism would not be infinite, but nevertheless the upper bound on the variance of equilibrium actions would not hold.

I now go on to discuss asset pricing implications of my model. So far we've looked at a single period model in which claims to firms' dividends is the only asset. In this case, the

asset price at state θ would simply be the value of dividends (profits) in that state. Since profits vary over the state space, so will asset prices. However, one of the main insights of asset pricing theory in the last few decades is that it is variation in the discount factor – rather than dividends – which accounts for the large swings observed in asset prices. A single period model is, therefore, not compatible with this type of explanation, because there is no room for a role for the discount factor in a single period model. In the next section I extend the single period model to one with infinite horizon overlapping generations, and study asset prices in a multi-period context when facing a coordination problem.

3.5 Infinite Horizon OLG

In this section I extend the single period model introduced in sections 3.2-3.4 to a multi-period dynamic model. I do this in order to study asset pricing implications in an economy destabilized by a coordination problem.

The world exists for an infinite number of periods. Agents who inhabit the economy live for two periods. In each period a unit measure of new agents is born, so in each period the world is inhabited by two groups of agents: young and old. Other than generational heterogeneity there is no difference between agents in the model, so each generation can be fully characterized by a representative agent.

Agents work when young but consume throughout their lives. The old generation holds an equity claim on a mutual fund who owns all the economy's firms, which they sell to the young in order to purchase goods on the market. The young work, and allocate the proceeds of their labour between current consumption and the (cum-dividend) equity claim. Period t young agents maximize expected utility from consumption in both periods

$$\frac{1}{1-\gamma} \mathbb{E}_t[Q_t^{1-\gamma} + \delta Q_{t+1}^{1-\gamma}] - \frac{1}{1+\frac{1}{v}} N_t^{1+\frac{1}{v}} \quad (3.20)$$

$$Q_t = \left(\int_{i=0}^1 Q_{i,t}^\rho di \right)^{\frac{1}{\rho}}$$

subject to the budget constraint

$$\frac{\int_{i=0}^1 P_{i,t} Q_{i,t} di}{P_t} + \phi_t p_t \leq N_t \frac{W}{P_t} + \phi_t d_t \quad (3.21)$$

$$\frac{\int_{i=0}^1 P_{i,t+1} Q_{i,t+1} di}{P_{t+1}} \leq \phi_t p_{t+1} \quad (3.22)$$

where δ is the time-discount, ϕ are holdings of the mutual funds shares (in percent), and p_t and d_t are the real price (cum dividend) and the real dividend of the mutual fund respectively:

$$d_t(\theta) = \frac{\int_{i=0}^1 \Pi_i(\theta) di}{P(\theta)} \quad (3.23)$$

The nominal wage is rigid in the sense of assumption 2. This means the nominal wage at time t , W_t cannot go below some arbitrary value $\hat{W}_t$; it equals that value whenever labour demand is lesser or equal to labor supply.

$$W_t(\theta) = \begin{cases} \hat{W}_t & \text{if } \underbrace{\frac{1}{A} \frac{1}{2\epsilon} \int Q(\tilde{\theta})^{\frac{1}{\alpha}} d\tilde{\theta}}_{\text{Labour demand}} \leq \underbrace{\left(\frac{\int p_{t+1}(\tilde{\theta})^\gamma d\tilde{\theta}}{p_t - d_t} (1 - \sigma^{-1}) \alpha A \left[\frac{1}{2\epsilon} \int (Q(\tilde{\theta}))^{1-\sigma^{-1}} d\tilde{\theta} \right]^{\frac{1-\alpha^{-1}}{1-\sigma^{-1}}} \right)^v}_{\text{Labour supply}} \\ P(\theta) \left(\frac{1}{A} \frac{1}{2\epsilon} \int Q(\tilde{\theta})^{\frac{1}{\alpha}} d\tilde{\theta} \right)^{-\frac{1}{v}-1} \frac{p_t - d_t}{p_{t+1}(\bar{\theta})^\gamma d\bar{\theta}} & \text{else} \end{cases}$$

Unlike the former sections, labour supply depends on asset prices and so it also varies on the state space Θ .²¹ We will see below that labour supply decreases as θ increases, ensuring that labour supply and labour demand meet exactly once. As before, total output in the economy varies between $Q(\theta) \in [\underline{Q}, \bar{Q}]$. $\underline{Q} > 0$ can be arbitrarily low, while $\bar{Q}$ is defined by the equality

$$\frac{1}{A} \bar{Q}^{\frac{1}{\alpha}} = \left(\frac{\int_{\bar{\theta}=0}^1 p_{t+1}(\bar{\theta})^\gamma d\bar{\theta}}{p_t(\bar{\theta}) - d_t(\bar{\theta})} (1 - \sigma^{-1}) \alpha A [\bar{Q}]^{1-\alpha^{-1}} \right)^v \quad (3.24)$$

The latter in turn depends on asset prices (through the marginal substitution between labour and consumption). The above equation does not mean a redefinition of $\bar{\theta}$, which is a parameter in the model. At state $\bar{\theta}$, high production is ensured by assumption. However, that does not preclude that high production occurs at lower states endogenously. But even if this was the case, the asset price would be the same as in the state $\bar{\theta}$. What is, then, the asset price?

Assets in this economy are priced in the same way as in a Lucas tree economy with stochastic endowment. This is because labour supply is always less than or equal to labour demand, and thus it is equivalent to a case in which consumers have no decision about labour supply, but rather take their real wage as exogenous – just as is the case with an endowment (Grossman and Barro, 1976, pp. 45).

²¹This could be avoided by assuming GHH preferences, that would make the labour supply curve depend only on the real wage.

Consider the pricing of an asset with a stochastic cash flow $d_t(\theta)$ at time t , in an OLG model with endowment $e_t(\theta)$ to the young generation, and the old generation holds a claim to the asset. Recall that the young generation discounts time by factor δ . Let $p_t(\theta)$ be the price of the asset in period t at state θ . How is $p_t(\theta)$ priced? If the asset is cum dividend²², generation t optimality yields²³

$$p_t(\theta) = d_t(\theta) + \frac{\delta\psi}{u'(e_t + d_t - p_t)} \left(\frac{1}{1 - \delta\chi} \right) \quad (3.27)$$

$$\psi = \mathbb{E}_s [d_{s+1} u'(p_{s+1})] \quad \forall s \geq t \quad \chi = \mathbb{E}_s \left[\frac{u'(p_{s+1})}{u'(e_{s+1} + d_{s+1} - p_{s+1})} \right] \quad \forall s \geq t$$

The RHS of equations 3.27 decreases with $p(\theta')$ $\forall \theta' \in \Theta$, so that given prices in all other states, the price in state θ is unique. Assuming θ is iid, the variance of asset prices is given by:

$$\begin{aligned} Var[p] = Var[d] + \left(\frac{\delta\psi}{1 - \delta\chi} \right)^2 Var[u'(e + d - p)] \\ + 2 \left(\frac{\delta\psi}{1 - \delta\chi} \right) Cov[d, [u'(e + d - p)]^{-1}] \end{aligned} \quad (3.28)$$

In the cum dividend case, either dividend volatility, or aggregate income volatility are sufficient for asset prices to be volatile.²⁴ As long as output is volatile, then asset prices are volatile as well. Importantly, this could be the case even when exogenous risk factors are negligible. This result is essentially an extension of Proposition 4. Asset prices are

²²If the asset would be sold ex dividend, we would instead get

$$p_t(\theta) = \frac{\delta\psi}{u'(e_t - p_t)} \left(\frac{1}{1 - \delta\chi} \right) \quad (3.25)$$

$$\begin{aligned} \psi &= \mathbb{E}_{t+s} [d_{s+1} u'(p_{s+1} + d_{s+1})] \\ \chi &= \mathbb{E}_s \left[\frac{u'(p_{s+1} + d_{s+1})}{u'(e_{s+1} - p_{s+1})} \right] \end{aligned}$$

$$Var[p] = \left(\frac{\delta\psi}{1 - \delta\chi} \right)^2 Var[u'(e - p)] \quad (3.26)$$

²³The full derivation is given in Appendix C.4

²⁴In the ex dividend case, as long as there is volatility in the endowment, there must be volatility of asset prices.

volatile directly due to dividend variance, and indirectly due to discount-factor variance and its covariance with dividends, just as standard asset pricing theory suggests.

It is important also to emphasize what the statement above doesn't say – that asset prices are volatile while dividends aren't. In the asset pricing literature it is common to take the dividend series directly as the fundamental. Asset prices are uniquely determined in equilibrium. They are not directly determined by a self-fulfilling prophecy, as in OLG models of fiat money for example. The self-fulfilling prophecy in the production economy affects asset prices only through the discount factor and through determining the dividend process.²⁵

In sum, I've shown that at states where firms choose to produce $\bar{Q}$, the marginal substitution between labour and consumption equals the marginal utility from consuming an additional unit of consumption times the real wage. Firms face the exact same maximization problem as in the last section, so that Lemma 1 follows through to an OLG world, and thus there is a unique equilibrium in which, although exogenous risk factors are negligible, there still remains risk in the economy due to the coordination problem and consequently asset prices are volatile. Excess volatility of prices as compared with the dividend process does not arise.

3.6 Discussion

In this chapter I propose a stylized macro model in which a coordination problem between producers generates endogenous fluctuations in output, employment and asset prices – even if exogenous variation is negligible. This chapter's main contributions are twofold: to the systemic risk literature and to macroeconomics. As to the first, I show that it is theoretically possible to have an economy in which exogenous risk factors are negligible, yet endogenous risk persists in equilibrium. I term this an amplification mechanism with an infinite signal-to-noise ratio. The reason this could occur is because of multiple equilibria. If economic outcomes are not uniquely determined in the complete information case, there is an opening for extreme amplification in the incomplete information case. In the model, demand is ex-post low for a given level of fundamentals because agents expect this outcome ex-ante. In other words, there is an element of self-fulfilling prophecy to the equilibrium. However, unlike sun-spot models, it is not the case that 'anything goes', nor is it the case that beliefs follow an autonomous process (as perhaps hinted at in Keynes's 'animal spirits' concept). What

²⁵Allen et al. (2006) study a model in which there is an element of self-fulfilling prophecy to asset prices, and as a consequence, prices deviate from expected dividends. However, in their model asset prices are less volatile due this property, not more.

distinguishes the global games literature is that by relaxing the assumption about common knowledge, beliefs about other agents' actions are pinned down by imposing a parsimonious structure on exogenous variation in fundamentals (payoff-relevant exogenous variables).

The second contribution is to the macroeconomics literature. I emphasize an insight that is often overlooked in Keynes General Theory, that once the labour market is precluded from imposing on all firms the condition that marginal substitution between labour and consumption is equal to the real wage, the door is opened for multiple equilibria. I further contribute to that literature by applying modern game theory²⁶ in order to resolve equilibrium multiplicity. Resolving equilibrium multiplicity is important in view of the difficulty to conduct comparative statics, and as a result making the model unable to evaluate economic policy. This is one reason economists have a distaste for the role that notions like 'animal-spirits' play in classical Keynesian models.

The way we conceptualize risk, and the cause of economic fluctuations, matters. If, on the one hand, we hold that risk is largely exogenous, and thereby not in our control, we may view recessions as a natural phenomenon. This would lead to a normative prescription on how to handle such events; e.g. we might implement contractionary fiscal policy because we think we can no longer afford current expenditure when revenue from taxation decreases during a recession. On the other hand, if we view recessions as a market failure, we may wish to implement different policies to try to prevent recessions from occurring, and to alleviate their costs. If indeed the problem in recessions is a lack of effective demand, counter-cyclical fiscal policy (expansionary during recessions) may be preferable.

Economists often intuitively accept the hypothesis that recessions are a result of some kind of market failure, perhaps several, although there is by no means agreement on that front. The contribution of this dissertation pertains succinctly to self-fulfilling nature of economic beliefs, showing that macroeconomic downturns need not involve significant variation in economic fundamentals, because these could be used almost purely as a coordination device. The quote from Hall et al. (1980) brought in the beginning of this chapter, best captures this. My thesis conceptualizes formally why, in Robert Hall's words, there is sometimes distinctly less work to do in the whole U.S. economy.

²⁶That was not available a few decades ago, when discussion about wage rigidity were more active.

Conclusions

This dissertation contains three studies focusing on systemic risk. It contributes to several fields in economics and finance: banking, financial stability, macroeconomics and applied microeconomics. It makes contributions both to the theory and to the empirical body of knowledge in those fields. In the next few paragraphs I briefly summarize the main contributions of each chapter, shedding light on the connecting theme between them.

The first chapter applies the seminal work by Goldstein and Pauzner to study the question: how does the connectivity between banks affect the stability of the financial system? Contrary to previous literature, it shows that reducing connectivity (credit exposures) between banks will not necessarily increase the network's stability. Previous studies have ignored the endogenous response of economic agents when studying financial stability, focusing instead the mechanics of contagion. They compared financial crises to a pandemic: in order to stop a crisis from spreading, one has to 'quarantine' individual banks from one another. This makes sense in the models that are typically being used in the literature, because they are systems that do not preserve value. By having (non-) agents that could passively absorb losses, a shock that enters the system can quickly evaporate through those 'black holes'. This is indeed what would happen to a virus under quarantine – it would simply die. Under this logic, it makes sense to sever links between banks, in order to avoid the spread of losses across the financial system. However, the financial system is a closed system: every creditor has a debtor. All else equals, reducing credit risk only shifts the real risk from some agents to others. Using this logic, I show that credit exposures in fact also have a positive effect on stability, making it impossible to determine ex-ante whether reducing interbank connectivity would be good for financial stability.

The first chapter developed a theory that puts into question the usefulness of several policies that were introduced in the past 20 years, that were aimed at stabilizing the financial system by reducing interbank credit risk. One such policy, called insolvency netting, is studied empirically in the second chapter. There, I develop a new way to identify structural changes in systemic risk, and apply it to evaluate whether insolvency netting reduced systemic

risk. A large part of the literature on measuring systemic risk focused on the cross-section covariance of tail risks. In contrast, I propose to study the covariance between a economic fundamentals and default risk at the firm level. For a given level of fundamentals, a lower correlation indicates that the threshold of fundamentals at which bond holders refuse to roll over their claims has gone down – i.e. the company will fail at worse fundamentals. I show that a two year trend in falling correlations for American non-financial companies suddenly slowed in April 2005, the same month when the legislation was passed. The trend for European companies continued uninterrupted. As to financial companies, I find that the pre-financial crisis default insurance had very little correlation with firms' fundamentals. In the U.S. it is of the opposite sign to economic intuition. The R^2 on those regressions is low, indicating a potential spurious correlation. This finding echoes Kelly et al. (2016); such different dynamics between financial and non-financial sectors' default-insurance could indicate that market participants priced in implicit government guarantees.

In the third chapter I take a step back into theory, studying how a coordination problem could give rise to a new type of amplification mechanism, one that allows even minute innovations to economic fundamentals to bring about large fluctuations in economic outcomes. The key insight from this chapter is that in economic environments in which there are multiple equilibria in complete information, the unique equilibrium in incomplete information has a knife-edge property: economic outcomes are radically different below and above a threshold of fundamentals. The reason for this is that fundamentals are almost purely used as a coordination device. In the chapter I apply similar machinery as in the first chapter to a general equilibrium model with a nominal wage rigidity. I show that this type of friction implies that we are missing a condition for labour market clearing, thus giving rise to a coordination problem between producers. In complete information there are multiple equilibria that can be Pareto ranked. I apply global games methods to select a unique equilibrium, and study its stability properties.

References

- Acemoglu, D., Carvalho, V. M., Ozdaglar, A., and Tahbaz-Salehi, A. (2012). The network origins of aggregate fluctuations. *Econometrica*, 80(5):1977–2016.
- Acemoglu, D., Ozdaglar, A., and Tahbaz-Salehi, A. (2015). Systemic risk and stability in financial networks. *The american economic review*, 105(2):564–608.
- Acharya, V. V., Pedersen, L. H., Philippon, T., and Richardson, M. (2017). Measuring systemic risk. *The review of financial studies*, 30(1):2–47.
- Adrian, T. and Brunnermeier, M. K. (2016). Covar. *The American Economic Review*, 106(7):1705.
- Adrian, T., Etula, E., and Groen, J. J. (2011). Financial amplification of foreign exchange risk premia. *European economic review*, 55(3):354–370.
- Adrian, T., Etula, E., and Muir, T. (2014). Financial intermediaries and the cross-section of asset returns. *The Journal of Finance*, 69(6):2557–2596.
- Adrian, T., Moench, E., and Shin, H. S. (2016). Dynamic leverage asset pricing. *Federal Reserve Bank of New York, Staff Reports*.
- Agell, J. and Lundborg, P. (2003). Survey evidence on wage rigidity and unemployment: Sweden in the 1990s. *Scandinavian Journal of Economics*, 105(1):15–30.
- Allen, F. and Carletti, E. (2013). What is systemic risk? *Journal of Money, Credit and Banking*, 45(s1):121–127.
- Allen, F., Carletti, E., Goldstein, I., and Leonello, A. (2018). Government guarantees and financial stability. *Journal of Economic Theory*, 177:518–557.
- Allen, F. and Gale, D. (2000). Financial contagion. *Journal of political economy*, 108(1):1–33.
- Allen, F. and Gale, D. (2004). Competition and financial stability. *Journal of money, credit and banking*, pages 453–480.
- Allen, F., Morris, S., and Shin, H. S. (2006). Beauty contests and iterated expectations in asset markets. *Review of financial Studies*, 19(3):719–752.
- Altman, E. I. (1968). Financial ratios, discriminant analysis and the prediction of corporate bankruptcy. *The journal of finance*, 23(4):589–609.

- Altonji, J. G. and Devereux, P. J. (2000). The extent and consequences of downward nominal wage rigidity. In *Research in labor economics*, pages 383–431. Emerald Group Publishing Limited.
- Angeletos, G.-M. and La'O, J. (2013). Sentiments. *Econometrica*, 81(2):739–779.
- Angeletos, G.-M. and Lian, C. (2019). Confidence and the propagation of demand shocks. *Unpublished*.
- Angeletos, G.-M. and Pavan, A. (2007). Efficient use of information and social value of information. *Econometrica*, 75(4):1103–1142.
- Asquith, P., Gertner, R., and Scharfstein, D. (1994). Anatomy of financial distress: An examination of junk-bond issuers. *The quarterly journal of economics*, 109(3):625–658.
- Aumann, R. J. (1976). Agreeing to disagree. *The annals of statistics*, pages 1236–1239.
- Azariadis, C. (1981). Self-fulfilling prophecies. *Journal of Economic theory*, 25(3):380–396.
- Azariadis, C. and Stiglitz, J. E. (1983). Implicit contracts and fixed price equilibria. *The Quarterly Journal of Economics*, pages 2–22.
- Baird, D. G. (2017). Pari passu clauses and the skeuomorph problem in contract law. *Duke LJ Online*, 67:84.
- Ball, L. and Romer, D. (1991). Sticky prices as coordination failure. *The American Economic Review*, Vol. 81:539–552.
- Ball, L. M. (2018). *The Fed and Lehman Brothers: setting the record straight on a financial disaster*. Cambridge University Press.
- Barro, R. J. (1979). Second thoughts on keynesian economics. *The American Economic Review*, 69(2):54–59.
- Barro, R. J. (2006). Rare disasters and asset markets in the twentieth century. *The Quarterly Journal of Economics*, 121(3):823–866.
- Barro, R. J. and Grossman, H. I. (1971). A general disequilibrium model of income and employment. *The American Economic Review*, 61(1):82–93.
- Baudry, L. and Le Bihan, H. (2004). Price rigidity: Evidence from the french cpi macro-data. *SSRN Electronic Journal*.
- Benhabib, J. and Farmer, R. E. (1999). Indeterminacy and sunspots in macroeconomics. *Handbook of macroeconomics*, 1:387–448.
- Benhabib, J. and Schmitt-Grohé (2001). Monetary policy and multiple equilibria. *American Economic Review*, 91(1):167–186.
- Benos, A. and Papanastasopoulos, G. (2007). Extending the merton model: A hybrid approach to assessing credit quality. *Mathematical and computer modelling*, 46(1-2):47–68.

- Bernanke, B., Gertler, M., and Gilchrist, S. (1994). The financial accelerator and the flight to quality. Technical report, National Bureau of Economic Research.
- Bernanke, B. S. (2008). Remarks on the economic outlook. In *Speech prepared for the International Monetary Conference, Barcelona, Spain*, volume 3.
- Bernanke, B. S. et al. (2007). The subprime mortgage market. In *Federal Reserve Bank of Chicago Proceedings*, number 1036.
- Bewley, T. F. (2007). Fairness, reciprocity, and wage rigidity. *Behavioral Economics and Its Applications*, pages 157–188.
- Bharath, S. T. and Shumway, T. (2008). Forecasting default with the merton distance to default model. *The Review of Financial Studies*, 21(3):1339–1369.
- Billio, M., Getmansky, M., Lo, A. W., and Pelizzon, L. (2012). Econometric measures of connectedness and systemic risk in the finance and insurance sectors. *Journal of financial economics*, 104(3):535–559.
- Blanchard, O. J. and Kiyotaki, N. (1987). Monopolistic competition and the effects of aggregate demand. *The American Economic Review*, Vol. 77:647–666.
- Bliss, R. R. and Kaufmann, G. (2004). Derivatives and systemic risk: Netting, stays, and closeout. Technical report, Working paper, Wake Forest University and Federal Reserve Bank of Chicago.
- Bolton, P. and Oehmke, M. (2015). Should derivatives be privileged in bankruptcy? *The Journal of Finance*, 70(6):2353–2394.
- Borio, C. E. and Drehmann, M. (2009). Assessing the risk of banking crises—revisited. *BIS Quarterly Review*, March.
- Borovička, J., Hansen, L. P., and Scheinkman, J. A. (2016). Misspecified recovery. *The Journal of Finance*, 71(6):2493–2544.
- Brunnermeier, M. K. and Sannikov, Y. (2014). A macroeconomic model with a financial sector. *American Economic Review*, 104(2):379–421.
- Calomiris, C. W. and Kahn, C. M. (1991). The role of demandable debt in structuring optimal banking arrangements. *The American Economic Review*, pages 497–513.
- Campbell, J. Y. (1987). Stock returns and the term structure. *Journal of financial economics*, 18(2):373–399.
- Carlsson, H. and Van Damme, E. (1993). Global games and equilibrium selection. *Econometrica: Journal of the Econometric Society*, pages 989–1018.
- Chapman, J. S. C. (2016). Memorandum decision on omnibus motion of the noteholder defendants to dismiss the fourth amended complaint.
- Chari, V. V. and Jagannathan, R. (1988). Banking panics, information, and rational expectations equilibrium. *The Journal of Finance*, 43(3):749–761.

- Clower, R. W. (1965). A keynesian counter revolution: A theoretical appraisal. *The Theory of Interests*.
- Cochrane, J. H. (1994). Shocks. In *Carnegie-Rochester Conference series on public policy*, volume 41, pages 295–364. Elsevier.
- Cooper, R. and John, A. (1988). Coordinating coordination failures in keynesian models. *The Quarterly Journal of Economics*, Vol. 103(3):441–463.
- De Jonghe, O. (2010). Back to the basics in banking? a micro-analysis of banking system stability. *Journal of financial intermediation*, 19(3):387–417.
- Denis, D. J. and Denis, D. K. (1995). Causes of financial distress following leveraged recapitalizations. *Journal of financial economics*, 37(2):129–157.
- Denison, E., Sarkar, A., and Fleming, M. J. (2019). Creditor recovery in lehman’s bankruptcy. Technical report, Liberty Street Economics (blog, 14 January 2019), Federal Reserve Bank of New York; <https://libertystreeteconomics.newyorkfed.org/2019/01/creditor-recovery-in-lehmans-bankruptcy.html>.
- Diamond, D. W. and Dybvig, P. H. (1983). Bank runs, deposit insurance, and liquidity. *Journal of political economy*, 91(3):401–419.
- Diamond, P. A. (1982). Aggregate demand management in search equilibrium. *Journal of political Economy*, 90(5):881–894.
- Dichev, I. D. (1998). Is the risk of bankruptcy a systematic risk? *the Journal of Finance*, 53(3):1131–1147.
- Dichev, I. D. and Piotroski, J. D. (2001). The long-run stock returns following bond ratings changes. *The Journal of Finance*, 56(1):173–203.
- Du, Y. and Suo, W. (2004). Assessing credit quality from equity markets: Is a structural approach a better approach. *Canadian Journal of Administrative Studies (forthcoming)*.
- Duffie, D., Saita, L., and Wang, K. (2007). Multi-period corporate default prediction with stochastic covariates. *Journal of financial economics*, 83(3):635–665.
- Duffie, D. and Skeel, D. A. (2012). A dialogue on the costs and benefits of automatic stays for derivatives and repurchase agreements. *Faculty Scholarship. Paper 386*.
- Elton, E. J., Gruber, M. J., Agrawal, D., and Mann, C. (2001). Explaining the rate spread on corporate bonds. *the journal of finance*, 56(1):247–277.
- Enders, Z. and Kollmann, R. (2010). Global banks and international business cycle. *Unpublished manuscript ECARES, Universite Libre de Bruxelles*.
- Erol, S. and Vohra, R. (2018). Network formation and systemic risk. *Available at SSRN 2546310*.
- Fama, E. F. and French, K. R. (1989). Business conditions and expected returns on stocks and bonds. *Journal of financial economics*, 25(1):23–49.

- Fama, E. F. and Schwert, G. W. (1977). Asset returns and inflation. *Journal of financial economics*, 5(2):115–146.
- Farmer, R. and Benhabib, J. (1994). Indeterminacy and increasing returns. *Journal of Economic Theory*, 63:19–41.
- Farmer, R. E. (2013). Animal spirits, financial crises and persistent unemployment. *The Economic Journal*, 123(568):317–340.
- Fehr, E. and Falk, A. (1999). Wage rigidity in a competitive incomplete contract market. *Journal of political Economy*, 107(1):106–134.
- Fehr, E. and Goette, L. (2005). Robustness and real consequences of nominal wage rigidity. *Journal of Monetary Economics*, 52(4):779–804.
- Frankel, D. M., Morris, S., and Pauzner, A. (2003). Equilibrium selection in global games with strategic complementarities. *Journal of Economic Theory*, 108(1):1–44.
- Gabaix, X., Krishnamurthy, A., and Vigneron, O. (2007). Limits of arbitrage: theory and evidence from the mortgage-backed securities market. *The Journal of Finance*, 62(2):557–595.
- Gai, P. and Kapadia, S. (2010). Contagion in financial networks. In *Proceedings of the Royal Society of London A: Mathematical, Physical and Engineering Sciences*, page rspa20090410. The Royal Society.
- Glasserman, P. and Young, H. P. (2015). How likely is contagion in financial networks? *Journal of Banking & Finance*, 50:383–399.
- Goldstein, I., Kopytov, A., Shen, L., and Xiang, H. (2020). Bank heterogeneity and financial stability. Technical report, National Bureau of Economic Research.
- Goldstein, I. and Pauzner, A. (2005). Demand–deposit contracts and the probability of bank runs. *the Journal of Finance*, 60(3):1293–1327.
- Goodhart, C. A., Basurto, M. A. S., et al. (2009). *Banking stability measures*, volume 627. Citeseer.
- Gorton, G. (1988). Banking panics and business cycles. *Oxford economic papers*, 40(4):751–781.
- Griffin, J. M. and Lemmon, M. L. (2002). Book-to-market equity, distress risk, and stock returns. *The Journal of Finance*, 57(5):2317–2336.
- Grossman, H. I. (1972). Was keynes a" keynesian"? a review article.
- Grossman, H. I. and Barro, R. J. (1976). *Money, employment and inflation*. Cambridge University Press.
- Hall, R. E. (2005). Employment fluctuations with equilibrium wage stickiness. *American economic review*, 95(1):50–65.

- Hall, R. E. et al. (1980). Employment fluctuations and wage rigidity. *Brookings Papers on economic activity*, 1(1980):91–123.
- Hand, J. R., Holthausen, R. W., and Leftwich, R. W. (1992). The effect of bond rating agency announcements on bond and stock prices. *The journal of finance*, 47(2):733–752.
- He, Z., Kelly, B., and Manela, A. (2017). Intermediary asset pricing: New evidence from many asset classes. *Journal of Financial Economics*, 126(1):1–35.
- Holmstrom, B. and Tirole, J. (1997). Financial intermediation, loanable funds, and the real sector. *the Quarterly Journal of economics*, 112(3):663–691.
- Holthausen, R. W. and Leftwich, R. W. (1986). The effect of bond rating changes on common stock prices. *Journal of Financial Economics*, 17(1):57–89.
- Huang, X., Zhou, H., and Zhu, H. (2009). A framework for assessing the systemic risk of major financial institutions. *Journal of Banking & Finance*, 33(11):2036–2049.
- Kaiser, E. (2007). Subprime losses could hit \$100 billion: Bernanke. *Reuters*.
- Kealhofer, S., Kwok, S., and Weng, W. (1998). Uses and abuses of bond default rates, kmv corporation.
- Keim, D. B. and Stambaugh, R. F. (1986). Predicting returns in the stock and bond markets. *Journal of financial Economics*, 17(2):357–390.
- Kelly, B., Lustig, H., and Van Nieuwerburgh, S. (2016). Too-systemic-to-fail: What option markets imply about sector-wide government guarantees. *American Economic Review*, 106(6):1278–1319.
- Keynes, J. M. (2018). *The general theory of employment, interest, and money*. Springer.
- Kindleberger, C. P. and Aliber, R. Z. (1978). Manias. *Panics and Crashes: A History of*.
- Kiyotaki, N. (1988). Multiple expectational equilibria under monopolistic competition. *The Quarterly Journal of Economics*, 103(4):695–713.
- Kiyotaki, N. and Moore, J. (1997). Credit cycles. *Journal of political economy*, 105(2):211–248.
- Kocherlakota, N. et al. (2000). Creating business cycles through credit constraints. *Federal Reserve Bank of Minneapolis Quarterly Review*, 24(3):2–10.
- Krause, M. U. and Lubik, T. A. (2007). The (ir) relevance of real wage rigidity in the new keynesian model with search frictions. *Journal of Monetary Economics*, 54(3):706–727.
- Ladley, D. (2013). Contagion and risk-sharing on the inter-bank market. *Journal of Economic Dynamics and Control*, 37(7):1384–1400.
- Liu, X. (2016). Interbank market freezes and creditor runs. *The Review of financial studies*, 29(7):1860–1910.

- Longstaff, F. A., Mithal, S., and Neis, E. (2005). Corporate yield spreads: Default risk or liquidity? new evidence from the credit default swap market. *The journal of finance*, 60(5):2213–2253.
- Martin, I. and Gao, C. (2021). Volatility, valuation ratios, and bubbles: An empirical measure of market sentiment.
- Matta, R. and Perotti, E. C. (2016). Liquidity runs. *Tinbergen Institute Discussion Paper*.
- Mengle, D. (2010). The importance of close-out netting. *ISDA Research Notes*, 1.
- Merton, R. C. (1974). On the pricing of corporate debt: The risk structure of interest rates. *The Journal of finance*, 29(2):449–470.
- Monderer, D. and Samet, D. (1989). Approximating common knowledge with common beliefs. *Games and Economic Behavior*, 1(2):170–190.
- Mooney Jr, C. W. (2014). The bankruptcy code's safe harbors for settlement payments and securities contracts: When is safe too safe. *Tex. Int'l LJ*, 49:245.
- Morris, S., Rob, R., and Shin, H. S. (1995). p-dominance and belief potential. *Econometrica: Journal of the Econometric Society*, pages 145–157.
- Morris, S. and Shin, H. S. (1998). Unique equilibrium in a model of self-fulfilling currency attacks. *American Economic Review*, pages 587–597.
- Morris, S. and Shin, H. S. (2002). Social value of public information. *The American Economic Review*, 92(5):1521–1534.
- Ohlson, J. A. (1980). Financial ratios and the probabilistic prediction of bankruptcy. *Journal of accounting research*, pages 109–131.
- Opler, T. C. and Titman, S. (1994). Financial distress and corporate performance. *The Journal of finance*, 49(3):1015–1040.
- Patinkin, D. (1948). Price flexibility and full employment. *American Economic Review*, 38(4):543–564.
- Perotti, E. et al. (2010). Systemic liquidity risk and bankruptcy exceptions. *VoxEU.org*, 13.
- Roe, M. J. (2013). Derivatives markets in bankruptcy. *Comparative Economic Studies*, 55(3):519–534.
- Ross, S. (2015). The recovery theorem. *The Journal of Finance*, 70(2):615–648.
- Rotemberg, J. J. and Saloner, G. (1986). A supergame-theoretic model of price wars during booms. *The American Economic Review*, Vol. 76:390–407.
- Rubinstein, A. (1989). The electronic mail game: Strategic behavior under "almost common knowledge". *The American Economic Review*, pages 385–391.
- Schaal, E. and Taschereau-Dumouchel, M. (2015). Coordinating business cycles. *Available at SSRN 2567344*.

- Shapiro, C. and Stiglitz, J. E. (1984). Equilibrium unemployment as a worker discipline device. *The American Economic Review*, 74(3):433–444.
- Siriwardane, E. N. (2015). Concentrated capital losses and the pricing of corporate credit risk. *Harvard Business School working paper series# 16-007*.
- Stiglitz, J. E. (1984). Theories of wage rigidity.
- Tett, G. (2009). *Fool's gold: How unrestrained greed corrupted a dream, shattered global markets and unleashed a catastrophe*. Hachette UK.
- Valukas, A. (2010). Report of the examiner in the chapter 11 proceedings of lehman brothers holdings inc. Retrieved October, 23:2012.
- Vassalou, M. and Xing, Y. (2004). Default risk in equity returns. *The journal of finance*, 59(2):831–868.
- Veronesi, P. and Zingales, L. (2010). Paulson's gift. *Journal of Financial Economics*, 97(3):339–368.
- Wood, P. R. (2007). *Principles of international insolvency*. Sweet & Maxwell.
- Woodford, M. (1986). Stationary sunspot equilibria in a finance constrained economy. *Journal of Economic Theory*, 40(1):128–137.
- Yankov, V. et al. (2020). The liquidity coverage ratio and corporate liquidity management. Technical report, Board of Governors of the Federal Reserve System (US).
- Zhu, H. (2006). An empirical comparison of credit spreads between the bond market and the credit default swap market. *Journal of Financial Services Research*, 29(3):211–235.

A

A.1 Proof of Proposition 1

Theorem 1 in Goldstein and Pauzner (2005) shows that a unique threshold-equilibrium exists, and is the only equilibrium possible. Their result follows from two separate properties of the differential utility function v : (1) single crossing; (2) one-sided strategic complementarities. The latter means that v is decreasing in n (proportion of agents who withdraw early) whenever v is positive. With single crossing, given that all agents use a threshold-equilibrium, then there exists a unique threshold-equilibrium. With one-sided strategic complementarities, any equilibrium must be a threshold-equilibrium.

A.1.1 Definitions and Preliminary Results

A mixed strategy for a patient agent i is a function $s_i : [-\epsilon, 1 + \epsilon] \rightarrow [0, 1]$, for each possible signal, the agent withdraws early with probability s_i . Let $\tilde{n}(\theta)$ be a random variable with support $[0, 1]$ as the total number of agents that withdraw early in state θ . Let $\tilde{n}_k = \tilde{n}(\theta_k)$ be the number of agents who withdraw early in island k . $\tilde{n}_k$ is defined by its CDF $F_{\theta,k}(n)$:

$$F_{\theta}(n) = \mathbb{P}[\tilde{n}(\theta) \leq n] = \mathbb{P}\left[\lambda + (1 - \lambda) \int_{i=0}^1 s_i(\theta + \epsilon_i) di \leq n\right]$$

Let $n^x(\theta)$ be the inverse CDF

$$n^x(\theta) = \inf\{n : F_{\theta}(n) \geq x\}$$

A patient agent decides to withdraw early with probability 1, $s_i = 1$, if $\Delta(\cdot, \theta_i) < 0$, withdraw late with probability 1, $s_i = 0$, if $\Delta(\cdot, \theta_i) > 0$, and is indifferent between the two

(and therefore could use any mix of pure strategies) if $\Delta(\cdot, \theta_i) = 0$. Recall the definition of Δ (eq. 1.11-1.12):

$$\Delta(c, \theta_i, s_{-i}) = \frac{1}{2\epsilon} \int_{\theta_i - \epsilon}^{\theta_i + \epsilon} \underbrace{\left(\int_{n=\lambda}^1 v(\boldsymbol{\theta}, n) dF_{\boldsymbol{\theta}}(n) \right)}_{\mathbb{E}[v(\boldsymbol{\theta}, \tilde{n}(\boldsymbol{\theta})) | \boldsymbol{\theta}]} d\theta$$

$$v(c, \boldsymbol{\theta}, n) = u(c_2) - [qu(c) + (1 - q)u(c_2)]$$

A threshold strategy is characterized by a single number, θ^* , that a patient agent – if she observes a signal higher than that number – will withdraw late, and otherwise she will withdraw early. A threshold-equilibrium is one in which all agents' follow threshold strategies. If all agents follow the same strategy, the proportion of agents who withdraw early in each island n is deterministic.

There are three main differences between my model and that of GP: an additional island, asset returns and bankruptcy costs. First, in GP's model there is a single island, which precludes any heterogeneity in period 2. Second, GP have a fixed return with varying probability, while in my model returns vary with θ_k , but are certain for a given θ_k . Third, my model has bankruptcy/illiquidity costs whereas GP's model doesn't. In both models there are no costs to liquidating the long asset up to a proportion x , at which point first period claims are paid on a first-come-first-served basis. However, in GP $x = 1$ so there are no costs at all. Conversely, in my model liquidation of the long asset is cost-free up to a proportion $x \leq 1$. Not all assets are liquidated to satisfy first period withdrawals, which means there is always some value left to agents who did not run and/or didn't manage to withdraw their deposit in time.

Limiting liquidation of the asset at $x \leq 1$ implies there is always some value left in period 2. Period 2 consumption (eq. 1.18-1.21) is given by

$$c_2(\boldsymbol{\theta}) = \begin{cases} \frac{A_k}{1 - \min[n, x/c]} & \text{if } A_j \geq D \ \forall j \in \{k, -k\} \\ \frac{A_k - B + \psi[B + A_{-k}]}{1 - \min[n, x/c]} & \text{if } A_{-k} < D \ \wedge \ A_k - B + \psi[B + A_{-k}] > D \\ (1 - \psi)(A_k + B) & \text{if } A_k < D \ \wedge \ A_{-k} - B + \psi[B + A_k] > D \\ \frac{A_k + \psi A_{-k}}{1 + \psi} & \text{else} \end{cases}$$

$$A_k = (1 - \min[nc, x])R(\theta_k)(1 - \gamma_1 \times 1_{n > x/c})(1 - \gamma_2 \times 1_{k \text{ insolvent}})$$

where A_k is the bank's external assets, which may suffer losses due to illiquidity (insufficient period 1 value, $n > x/c$) and insolvency (insufficient period 2 value). Lemma 5 (as in GP) states basic facts about the function $\Delta(c, \theta_i)$. Let $(\tilde{n} + a)(\theta) = \tilde{n}(\theta + a)$.

Lemma 5 (Lemma 1). (i) $\Delta(\theta_i, \tilde{n}(\cdot))$ is continuous in the threshold strategy θ_i ; (ii) $\Delta(\theta_i + a, (\tilde{n} + a)(\cdot))$ is continuous and non-decreasing in a ; (iii) $\Delta(\theta_i + a, (\tilde{n} + a)(\cdot))$ is strictly increasing in a if two conditions are satisfied: $\theta_i + a < \bar{\theta} + \epsilon$; and $\tilde{n}(\theta) < 1/c$ with positive probability for $\theta \in [\theta_i + a - \epsilon, \theta_i + a + \epsilon]$.

Proof of Lemma 5. Note that (i) means Δ is continuous in θ_i given a certain distribution $\tilde{n}(\cdot)$ that doesn't change with θ_i ; a change in θ_i only changes the limits of the integration. The integrand $\int_{n=\lambda}^1 v(\theta, n) dF_\theta(n)$ is bounded because v is bounded. Continuity follows from the fact that $\int_{n=\lambda}^1 v(\theta, n) dF_\theta(n)$ is Riemann integrable over its domain $[\lambda, 1]$. Riemann integrability follows from the fact that $\int_{n=\lambda}^1 v(\theta, n) dF_\theta(n)$ can be partitioned into monotone functions, e.g. in case $\gamma_2 = 0$:

$$\int_{n=\lambda_L}^1 v(\theta, n) dF_\theta(n) = \int_{n=\lambda_L}^{x/c} v(\theta, n) dF_\theta(n) + \int_{n=x/c}^1 v(\theta, n) dF_\theta(n)$$

and each of those is Riemann integrable (because it is monotone, as we are not changing the distribution $\tilde{n}(\cdot)$, and v is monotone on each sub-interval), and thus each is continuous in θ_i . In case $\gamma_2 > 0$, partition is into potentially 9 parts, each of which is monotone in θ . The sum of two or more continuous functions is continuous. (ii) Continuity with respect to a follows because v is bounded and Δ is an integral over a segment $[\theta_i - \epsilon, \theta_i + \epsilon]$. Note that the distribution $\tilde{n}$ is only shifted by a constant a . Thus if we compare any θ with $\theta + a$, $\tilde{n}$ is identical, but at $\theta + a$ fundamentals are higher, so $R(\theta)$ is higher and consequently v . Since v is, ceteris paribus, non-decreasing in θ , so is Δ . (iii) In continuation to (ii), if the two conditions are kept then v is strictly increasing for at least part of the interval $[\theta_i - \epsilon, \theta_i + \epsilon]$, thus its integral is strictly increasing. This completes the proof. $\square$

Definition 5 (Strategic Complementarities). A game is said to exhibit strategic complementarities if the incentive of agent i to increase s_i is increasing in s_{-i} . More precisely, if $s_i \geq \hat{s}_i$ and $s_{-i} \geq \hat{s}_{-i}$ – which implies $F_\theta(n) \leq \hat{F}_\theta(n)$ – we have

$$u(s_i, s_{-i}, \theta) - u(\hat{s}_i, s_{-i}, \theta) \geq u(s_i, \hat{s}_{-i}, \theta) - u(\hat{s}_i, \hat{s}_{-i}, \theta)$$

A game is said to exhibit one-sided strategic complementarities if the incentive of agent i to increase s_i is increasing in s_{-i} whenever that incentive is positive.

Lemma 6 (Strategic Complementarities). The game played by patient agents exhibits strategic complementarities.

Proof of Lemma 6. Simplifying v yields $v(n, \cdot) = q(u(c_2) - u(c))$. Observe that the probability of consumption in period 1 is certain ($q = 1$) for $n \leq x/c$, and is decreasing from 1 to $\frac{x}{c}$

for $n \in (x/c, 1]$. Moreover, due to bankruptcy costs, for a fixed θ , A_k may have at most two discontinuities: one exactly at $n = x/c$; the other potentially at a point $n < x/c$ (since A_k decreases in n and may hit insolvency while still liquid). Thus, when $n > x/c$, A_k is fixed. $\frac{\partial v}{\partial n} < 0$ for $n \in [\lambda, x/c]$, and $\frac{\partial v}{\partial n} > 0 (< 0)$ in n for $n \in (x/c, 1]$ when $u(c_2) - u(c) < 0$ (> 0). In case $v(n, \cdot) > 0$ for $n \in (x/c, 1)$, then v is decreasing because $c_2 > c$ and because q is decreasing.

This proves strategic complementarities in a single island k , and for a fixed n . Observe that $\frac{\partial A_{-k}}{\partial n} \leq 0$, and that $\frac{\partial c_2}{\partial A_{-k}} \geq 0$. Thus due to the chain rule: $\frac{\partial c_2}{\partial n} \leq 0$, increasing the incentive to run on the bank. $\square$

Lemma 7 (Single Crossing). *Keeping θ fixed, there exists $n' \in (\lambda, 1)$ such that $v(n, \theta) \geq 0 \forall n < n'$ and $v(n, \theta) \leq 0 \forall n \geq n'$.*

Proof of Lemma 7. Single crossing is implied by strategic complementarities. If $v(n, \cdot) < 0$ for some $n \in (x/c, 1)$, it is negative throughout, so that v may cross (up to discontinuity) from positive to negative at most once. $\square$

The statement in Lemma 7 is for a fixed θ . Observe that if all agents follow a threshold strategy, $\frac{\partial n}{\partial \theta_k} \leq 0$, and therefore $\frac{\partial c_2}{\partial \theta_j} \geq 0 \forall j \in \{k, -k\}$: better fundamentals in either island imply (weakly) higher consumption in period 2. Thus varying θ over the interval $[\theta_i - \epsilon, \theta_i + \epsilon]$, v could cross zero (up to discontinuity) at most once.

A.1.2 Unique Threshold-Equilibrium

Proof of Proposition 1. The proof follows GP Proposition 1 closely. The only difference is that in my model v crosses zero only once up to a discontinuity, which doesn't alter the details of the proof much.

A threshold-equilibrium with threshold θ^* exists iff, given that all other agents use threshold strategies θ^* , each agent finds it optimal to withdraw early when she observes a signal below θ^* , and late if above it:

$$\Delta(c, \theta_i, \tilde{n}(\cdot, \theta^*)) < 0 \forall \theta_i < \theta^* \quad (\text{A.1})$$

$$\Delta(c, \theta_i, \tilde{n}(\cdot, \theta^*)) > 0 \forall \theta_i > \theta^* \quad (\text{A.2})$$

By continuity (Lemma 5i), a patient agent is indifferent when she observes θ^*

$$\Delta(c, \theta^*, \tilde{n}(\cdot, \theta^*)) = 0 \quad (\text{A.3})$$

I now show given that eq. A.3 holds, eqs. A.1-A.2 hold as well. Assume $\theta_i < \theta^*$, and partition each interval $[\theta^* - \epsilon, \theta^* + \epsilon]$ and $[\theta_i - \epsilon, \theta_i + \epsilon]$ to two complementary intervals: the intersection between the two, and its complement for each interval.

Let $Y = [\theta_i - \epsilon, \theta_i + \epsilon] \cap [\theta^* - \epsilon, \theta^* + \epsilon]$ be the intersection (possibly empty) of the interval over which $\Delta(c, \theta_i, \tilde{n}(\cdot, \theta^*))$ and $\Delta(c, \theta^*, \tilde{n}(\cdot, \theta^*))$ are evaluated, and two disjoint intervals $Z^i = [\theta_i - \epsilon, \theta_i + \epsilon] \setminus Y$ and $Z^* = [\theta^* - \epsilon, \theta^* + \epsilon] \setminus Y$. This yields

$$\Delta(c, \theta^*, \tilde{n}(\theta)) = \frac{1}{2\epsilon} \int_{\theta \in Y} v(\theta, n(\theta)) d\theta + \frac{1}{2\epsilon} \int_{\theta \in Z^*} v(\theta, n(\theta)) d\theta \quad (\text{A.4})$$

$$\Delta(c, \theta_i, \tilde{n}(\theta)) = \frac{1}{2\epsilon} \int_{\theta \in Y} v(\theta, n(\theta)) d\theta + \frac{1}{2\epsilon} \int_{\theta \in Z^i} v(\theta, n(\theta)) d\theta \quad (\text{A.5})$$

Equation A.4 must equal zero by assumption. Because $\theta_i < \theta^*$, the part integrated over the complement interval is non-positive (due to single crossing, Lemma 7). When comparing the second part in each equation (the complement of intersection), we must have

$$\frac{1}{2\epsilon} \int_{\theta \in Z^*} v(\theta, n(\theta)) d\theta > \frac{1}{2\epsilon} \int_{\theta \in Z^i} v(\theta, n(\theta)) d\theta$$

This implies eq. A.1 must be negative. A diametrical argument holds for eq. A.2. Due to strategic complementarities (Lemma 6), GP's argument follows through in full. $\square$

A.2 Proof of Proposition 2

Proof. We are interested in the effect of netting on the probability of runs. The proof utilizes the implicit function theorem on this indifference condition to show the effect of the interbank mutual claim B on the probability of runs θ^* . The result states that if $\gamma_2 = 0$ then $\frac{\partial \theta^*}{\partial B} < 0$. Since netting decreases B , the result shows netting increases the probability of bank runs.

Equation 1.11 is the indifference condition that determines the optimal threshold signal θ^* . In equilibrium, all agents use the threshold strategy θ^* , and their indifference condition is equal to zero.

$$\Delta(\theta^*, \cdot) = \frac{1}{2\epsilon} \int_{\theta^* - \epsilon}^{\theta^* + \epsilon} \underbrace{\left(\int_{n=\lambda}^1 v(\theta, n) dF_{\theta}(n) \right)}_{\mathbb{E}[v(\theta, n) | \theta_i = \theta^*]} d\theta = 0$$

Recall that $v = q(u(c_2) - u(c))$. Thus we can write

$$\Delta(\theta^*, \cdot) = \mathbb{E}[q(u(c_2) - u(c)) | \theta_i = \theta^*] \quad (\text{A.6})$$

It is clear from eq. A.6 that $\frac{\partial \Delta(\theta^*, \cdot)}{\partial \mathbb{E}[u(c_2) | \theta_i = \theta^*]} > 0$. Increasing expected utility from consumption in period 2 increases the incentive to withdraw late. Note also that $\frac{\partial \Delta(\theta^*, \cdot)}{\partial \theta^*} \geq 0$ because changing θ^* varies the limits of the integration leaving the distribution of early withdrawals on the interval $[\theta^* - \varepsilon, \theta^* + \varepsilon]$ unchanged. Due to the implicit function theorem, we have

$$\frac{\partial \theta^*}{\partial \mathbb{E}[u(c_2)]} < 0 \quad (\text{A.7})$$

Now consider the effect of B on $\mathbb{E}[u(c_2) | \theta_i = \theta^*]$, given that $\gamma_2 = 0$. In equilibrium, every state θ is associated with c_2 in island k and $-k$. Denote by $t = t(\theta) \in \mathbb{R}$ the transfer from k to $-k$ that is due to the interbank connection B ; it is either:

- $t > 0$ $-k$ insolvent; if k is also insolvent, it has a higher return than $-k$;
- $t < 0$ k insolvent; if $-k$ is also insolvent, it has a higher return than k

Observe that, due to symmetry of f_θ , for every state $\theta^+ = (\theta, \theta')$ there exists $\theta^- = (\theta', \theta)$ such that: (1) $f(\theta^+) = f(\theta^-)$; (2) $t(\theta^+) = -t(\theta^-)$. Generally, whenever $t > 0$ (< 0) period 2 consumption in island k ($-k$) is higher than that of island $-k$ (k).

Let Θ^+ be a subspace of Θ such that $(\theta, \theta') \in \Theta^+$ iff $t(\theta', \theta) > 0$; Θ^+ is a set of unique states of the world (up to island names). Denote by Θ^- its mirror image sub-space $\Theta^- = \{(\theta', \theta) : (\theta, \theta') \in \Theta^+\}$, and write c_2 as the consumption in island k in case $B = 0$. When computing expected utility differential in island k we have

$$\begin{aligned} \Delta(\theta^*, \cdot) &= \int_{\theta^+ \in \Theta^+} q(u(c_2 + t) - u(c)) d\theta^+ f(\theta^+ | \theta_i = \theta^*) \\ &+ \int_{\theta^- \in \Theta^-} q(u(c_2 - t) - u(c)) f(\theta^- | \theta_i = \theta^*) d\theta^- \\ &+ \int_{\theta \in \Theta \setminus (\Theta^+ \cup \Theta^-)} q(u(c_2) - u(c)) f(\theta | \theta_i = \theta^*) d\theta \\ &= \int_{\theta^+ \in \Theta^+} \left(q(\theta^+) (u(c_2(\theta^+) + t) - u(c)) \right. \\ &+ \left. q(\theta^-) (u(c_2(\theta^-) - t) - u(c)) \right) f(\theta^+ | \theta_i = \theta^*) d\theta^+ \\ &+ \int_{\theta \in \Theta \setminus (\Theta^+ \cup \Theta^-)} q(u(c_2) - u(c)) f(\theta | \theta_i = \theta^*) d\theta \end{aligned} \quad (\text{A.8})$$

Note that when $t > 0$, the fundamentals in island k are lower than in $-k$. Since agents observe a signal informing them on the common factor of θ_k , $n_k = n_{-k}$ in all states, therefore

the probability of consumption in period 1 is always the same in both islands. Also, observe that varying B can only change q via θ^* , as it doesn't affect liquid assets x . Thus, keeping θ^* fixed, q is unchanged when varying B . We need to show

$$q(\theta^+)u(c_2(\theta^+) + t) + q(\theta^-)u(c_2(\theta^-) - t) \geq q(\theta^+)u(c_2(\theta^+)) + q(\theta^-)u(c_2(\theta^-))$$

which is true because: (1) $q(\theta^+) = q(\theta^-)$; (2) $c_2(\theta^+) + t \leq c_2(\theta^-) - t$. Moreover, in the limit as $B \rightarrow \infty$, both banks are in effect consolidated, so we have

$$\lim_{B \rightarrow \infty} t = \frac{c_2(\theta^-) - c_2(\theta^+)}{2}$$

Because agents are risk averse,

$$u(c_2(\theta^+) + t) + u(c_2(\theta^-) - t) \geq u(c_2(\theta^+)) + u(c_2(\theta^-))$$

Finally, $\frac{\partial \mathbb{E}u(c_2)}{\partial B} > 0$, so together with eq. A.7 we get the desired result. □

A.2.1 Implicit Function Theorem, Conditions

The main condition for using the implicit function theorem is continuous non-zero first derivatives of Δ in a small open ball around the values in the parameter space that make $\Delta(\theta^*, \cdot) = 0$. This condition holds for $\frac{\partial \Delta}{\partial \theta^*}$ due to Lemma 5. Varying B doesn't change the limits of the integration of the integral in Δ . The partial derivative $\frac{\partial \Delta}{\partial B}$ is continuous because $\frac{\partial v}{\partial B}$ is continuous. The latter is the case because varying B doesn't affect outside assets A_k, A_{-k} ; rather it varies only the interbank transfer t .

A.3 Technical Appendix

This appendix describes in detail my simulation methodology. I make assumptions about the utility, technology and the distribution of θ and ν_k to simplify expressions of the function Δ (expected differential utility):

1. The utility function exhibits constant relative risk aversion:

$$u(c_t) = 1 - \exp(-\alpha c_t)$$

2. Returns on total assets are assumed to be linear in θ :

$$R_k = R(\theta_k, x) = xR_f + (1-x)\theta_k$$

where R_f is the gross risk free rate and θ the return on the long asset.

3. The common and idiosyncratic components of returns are all pairwise independent and distributed normally

$$\begin{pmatrix} \theta \\ \nu_k \\ \nu_{-k} \end{pmatrix} = N \left(\begin{bmatrix} \mu_\theta \\ 0 \\ 0 \end{bmatrix}, \begin{bmatrix} \sigma_\theta^2 & 0 & 0 \\ 0 & \sigma_\nu^2 & 0 \\ 0 & 0 & \sigma_\nu^2 \end{bmatrix} \right)$$

Recall $\theta_k = \theta + \nu_k$. Let ϕ and Φ be (respectively) the pdf and CDF of a normal random variable. Recall that (in island k) the return θ_k and the level of period 2 assets A_k are given by:

$$\begin{aligned} A_k(\theta_k) &= (1 - \min[nc, x])R(\theta_k)\Gamma \\ R_k &= R(\theta_k, x) = xR_f + (1-x)\theta_k \end{aligned}$$

It is useful to work with their inverse:

$$\begin{aligned} A^{-1} = \theta_k(A, x, c, \theta) &= \begin{cases} \left(\frac{A}{1-\gamma_1} \frac{1}{1-\min[nc, x]} - xR_f \right) \frac{1}{1-x} & \text{if } A \geq D \\ \left(\frac{A}{\Gamma} \frac{1}{1-\min[nc, x]} - xR_f \right) \frac{1}{1-x} & \text{if } A < D \end{cases} \\ R^{-1} = \theta(R, x) &= (R - xR_f)/(1-x) \end{aligned}$$

For a given θ , the state-space is two dimensional: one axis for ν_k and another for ν_{-k} . I partition the state-space into 4 regions, based on whether k and $-k$ are bankrupt in period 2. Region IV is sub-partioned to 3, depending on whether k ($-k$) is insolvent due to the interbank connection (see Figure 1.1).

The boundaries between regions are presented in Table A.1. These are used to partition the integral $\int_{\theta_k \in \mathbb{R}} \int_{\theta_{-k} \in \mathbb{R}} \dots$ into sums of integrals on mutually exclusive sub-spaces – one sum for each region:

$$\Delta(\theta_i = \theta^*, \theta^*) = \sum_{m \in \{I, II, III, IVa, IVb, IVc\}} \Delta_m \quad (\text{A.9})$$

$$\Delta(\theta^*, \theta_i, \cdot) = \int_{\theta=\theta_i-\epsilon}^{\theta_i+\epsilon} \int_{\theta_k \in \mathbb{R}} \int_{\theta_{-k} \in \mathbb{R}} q(u(c_2) - u(c)) f_{\theta}(\theta_k, \theta_{-k} | \theta) d\theta_k d\theta_{-k}$$

A.3.1 Region I

Both banks are solvent, $A_k \geq D \wedge A_{-k} \geq D$. Bank $-k$ is solvent, given θ , with probability $(1 - \Phi(\hat{\theta}_1, \theta, \sigma_\nu))$. Otherwise bank $-k$ has no effect on (period 2) expected utility in island k (given θ), which is given by

$$q(\theta) \int_{\theta_k=\hat{\theta}_1}^{\infty} u\left(\frac{A_k}{1 - \min[n, x/c]}\right) \phi(\theta_k, \theta, \sigma_\nu) d\theta_k$$

Due to CARA utility we have

$$\begin{aligned} &= -\frac{1}{\sigma_\nu \sqrt{2\pi}} \int_{\theta_k=\hat{\theta}_1}^{\infty} e^{-\alpha \frac{A_k}{1 - \min[n, x/c]}} e^{-\frac{1}{2\sigma_\nu^2}(\theta_k - \theta)} d\theta_k \\ &= -\frac{1}{\sigma_\nu \sqrt{2\pi}} \int_{\theta_k=\hat{\theta}_1}^{\infty} e^{-\alpha \left(c - \frac{c-1}{1 - \min[n, x/c]}\right) (xR_f + (1-x)\theta_k)} (1 - \tilde{\gamma}_1) e^{-\frac{1}{2\sigma_\nu^2}(\theta_k - \theta)} d\theta_k \end{aligned}$$

The assumption of CARA utility with normal distribution, together with the linear (in θ_k) functional form of returns R_k , means we can complete the square to get a shifted mean normal distribution (because u is exponential function linear in θ_k). Let b_1 be the coefficient on θ_k due to the utility function, and a_1 the free coefficient:

$$\begin{aligned} a_1 &= -\alpha \left(c - \frac{c-1}{1 - \min[n, x/c]}\right) xR_f (1 - \tilde{\gamma}_1) \\ b_1 &= -\alpha \left(c - \frac{c-1}{1 - \min[n, x/c]}\right) (1-x)(1 - \tilde{\gamma}_1) \end{aligned}$$

This results in

$$\begin{aligned} &= -e^{a_1 + b_1\theta + \sigma_\nu^2(b_1)^2/2} \frac{1}{\sigma_\nu \sqrt{2\pi}} \int_{\theta_k \in [\hat{\theta}_1, \infty)} e^{-\frac{1}{2\sigma_\nu^2}(\theta_k - (\theta + b_1\sigma_\nu^2))} d\theta_k \\ &= -e^{a_1 + b_1\theta + \sigma_\nu^2(b_1)^2/2} \left(1 - \Phi(\hat{\theta}_1, \theta + b_1\sigma_\nu^2, \sigma_\nu)\right) \end{aligned}$$

Thus we have

$$Z_1(\theta) = \left[1 - \Phi(\hat{\theta}_1, \theta, \sigma_\nu) - e^{a_1 + b_1\theta + (b_1\sigma_\nu)^2/2} \left(1 - \Phi(\hat{\theta}_1, \theta + b_1\sigma_\nu^2, \sigma_\nu)\right)\right] \left(1 - \Phi(\hat{\theta}_1, \theta, \sigma_\nu)\right)$$

$$\begin{aligned}
&= \left\{ 1 - \Phi(\hat{\theta}_1, \theta, \sigma_\nu) - \exp \left(-\alpha \left(c - \frac{c-1}{1 - \min[n, x/c]} \right) (xR_f + (1-x)\theta) (1 - \tilde{\gamma}_1) \right. \right. \\
&\quad \left. \left. + \left[\alpha \left(c - \frac{c-1}{1 - \min[n, x/c]} \right) (1-x)(1 - \tilde{\gamma}_1) \sigma_\nu \right]^2 / 2 \right) \times \right. \\
&\quad \left. \left(1 - \Phi \left(\hat{\theta}_1, \theta - \alpha(1-x) \left(c - \frac{c-1}{1 - \min[n, x/c]} \right) (1 - \tilde{\gamma}_1) \sigma_\nu, \sigma_\nu \right) \right) \right\} \times \left(1 - \Phi(\hat{\theta}_1, \theta, \sigma_\nu) \right) \\
&\Delta_1(\theta^*) = \int_{\theta=\theta^*-\epsilon}^{\theta^*+\epsilon} q(\theta) Z_1(\theta) d\theta
\end{aligned}$$

A.3.2 Region II

Bank k solvent, $-k$ insolvent, $A_k + \psi(A_{-k} + B) \geq D + B \wedge A_{-k} < D$. The lower boundary of the integral on θ_{-k} , $A^{-1}(\frac{1}{\psi}(B(1-\psi) + D - A_k))$ is derived from the condition that k stays solvent, i.e. the inequality $A_k + \psi(A_{-k} + B) \geq D + B$. Period 2 consumption in island k depends on how deep is $-k$'s insolvency, and is given by

$$\begin{aligned}
c_2(\boldsymbol{\theta} | \boldsymbol{\theta} \in \Theta_2) &= \left(c - \frac{c-1}{1 - \min[n, x/c]} \right) (xR_f(1 - \tilde{\gamma}_1 + \psi\Gamma) + \\
&\quad (1-x)((1 - \tilde{\gamma}_1)\theta_k + \psi\Gamma\theta_{-k})) - \frac{(1-\psi)B}{1 - \min[n, x/c]}
\end{aligned}$$

We can disentangle the component due to the interbank connection in $u(c_2(\boldsymbol{\theta} | \boldsymbol{\theta} \in \Theta_2))$, with the coefficients (for completing the square) being:

$$\begin{aligned}
a_2^k &= -\alpha \left(\left(c - \frac{c-1}{1 - \min[n, x/c]} \right) xR_f(1 - \tilde{\gamma}_1) - \frac{(1-\psi)B}{1 - \min[n, x/c]} \right) \\
a_2^{-k} &= -\alpha \psi \left(c - \frac{c-1}{1 - \min[n, x/c]} \right) xR_f\Gamma \\
b_2^k &= -\alpha \left(c - \frac{c-1}{1 - \min[n, x/c]} \right) (1-x)(1 - \tilde{\gamma}_1) \\
b_2^{-k} &= -\alpha \psi \left(c - \frac{c-1}{1 - \min[n, x/c]} \right) (1-x)\Gamma
\end{aligned}$$

When completing the square on the non-constant component (belonging to k), this yields

$$Z_2(\theta, A_{-k}) = e^{a_2^k + b_2^k \theta + (b_2^k \sigma_\nu)^2 / 2}$$

$$\begin{aligned} & \frac{1}{\sigma_\nu \sqrt{2\pi}} \int_{\theta_k = \hat{\theta}_2(A_{-k})}^{\infty} e^{-\frac{1}{2\sigma_\nu^2}(\theta_k - (\theta + \sigma_\nu^2 b_2^k))^2} d\theta_k \\ & = e^{a_2^k + b_2^k \theta + (b_2^k \sigma_\nu)^2 / 2} \end{aligned}$$

$$\left(1 - \Phi\left(\hat{\theta}_2(A_{-k}), \theta + \sigma_\nu^2 b_2^k, \sigma_\nu\right)\right)$$

(period 2) expected utility (given θ) in island k is given by

$$\begin{aligned} \Delta_2(\theta^*) &= \int_{\theta = \theta^* - \epsilon}^{\theta^* + \epsilon} q(\theta) \int_{\theta_{-k} = -\infty}^{\hat{\theta}_1} \left[1 - \Phi(\hat{\theta}_2(A_{-k}), \theta, \sigma_\nu) - Z_2(\theta, A_{-k}) \times \right. \\ & \quad \left. \exp(a_2^{-k} + b_2^{-k} \theta_{-k})\right] \phi(\theta_k, \theta, \sigma_\nu) d\theta_k d\theta \\ &= \int_{\theta = \theta^* - \epsilon}^{\theta^* + \epsilon} q(\theta) \left\{ \int_{\theta_{-k} = -\frac{x}{1-x}}^{\hat{\theta}_1} \left[1 - \Phi(\hat{\theta}_2(A_{-k}), \theta, \sigma_\nu) - Z_2(\theta, A_{-k}) \times \right. \right. \\ & \quad \left. \exp(a_2^{-k} + b_2^{-k} \theta_{-k})\right] \phi(\theta_k, \theta, \sigma_\nu) d\theta_k + \\ & \quad \left. \left[1 - \Phi(\hat{\theta}_2(0), \theta, \sigma_\nu) - Z_2(\theta, 0)\right] \Phi\left(-\frac{x}{1-x}, \theta, \sigma_\nu\right) \right\} d\theta \end{aligned}$$

In this region we cannot use completion to square because of the expression Z_2 inside the integral is a function of θ_k . Moreover, we cannot use integral2 because the boundary of the integral on θ_k depends on θ , thus we have to use nested integrals.

A.3.3 Region III

Bank k insolvent, $-k$ solvent, $A_k < D \wedge A_{-k} + \psi(A_k + B) \geq D + B$. The lower boundary of the integral on θ_{-k} , $A^{-1}(B(1 - \psi) + D - \psi A_k)$ is derived from the condition that $-k$ stays solvent, i.e. the inequality $A_{-k} + \psi(A_k + B) \geq D + B$. Period 2 consumption in island k doesn't depend on $-k$'s assets, but the boundary of the integral does depend on it, as well as the probability that $-k$ remains solvent. Those are given by

$$\begin{aligned} c_2(\theta | \theta \in \Theta_3) &= (1 - \psi) \left(\left(c - \frac{c-1}{1 - \min[n, x/c]} \right) (xR_f + (1-x)\theta_k) \Gamma + \frac{B}{1 - \min[n, x/c]} \right) \\ Z_3(A_k) &= 1 - \Phi\left(\hat{\theta}_2(A_k), \theta, \sigma_\nu\right) \end{aligned}$$

The coefficients in this region are

$$a_3^k = -\alpha(1 - \psi) \left(c - \frac{c-1}{1 - \min[n, x/c]} \right) xR_f \Gamma$$

$$a_3^{-k} = -\alpha(1-\psi) \frac{B}{1 - \min[n, x/c]}$$

$$b_3^k = -\alpha(1-\psi) \left(c - \frac{c-1}{1 - \min[n, x/c]} \right) (1-x)\Gamma$$

(period 2) expected utility (given θ) in island k is given by

$$\begin{aligned} \Delta_3(\theta^*) &= \int_{\theta=\theta^*-\epsilon}^{\theta^*+\epsilon} q(\theta) \int_{\theta_k=-\infty}^{\hat{\theta}_1} Z_3(A_k) \times u(c_2(\theta | \theta \in \Theta_2)) \phi(\theta_k, \theta, \sigma_\nu) d\theta_k d\theta \\ &= \int_{\theta=\theta^*-\epsilon}^{\theta^*+\epsilon} q(\theta) \left\{ \int_{\theta_k=-\frac{x}{1-x}}^{\hat{\theta}_1} Z_3(A_k) \times u(c_2(\theta | \theta \in \Theta_2)) \phi(\theta_k, \theta, \sigma_\nu) d\theta_k \right. \\ &\quad \left. + Z_3(0) \left(1 - \exp(a_3^{-k}) \right) \Phi\left(-\frac{x}{1-x}, \theta, \sigma_\nu\right) \right\} d\theta \end{aligned}$$

In this region we cannot use completion to square because of the expression Z_3 inside the integral on θ_k . Moreover, we cannot use reduce further the integral because the boundary depends on θ_k , thus we have to use nested integrals.

Region IVa

Both banks are insolvent, and both would have been insolvent even when the other bank is solvent: $A_k < D \wedge A_{-k} < D$. Period 2 consumption in island k is given by:

$$c_2(\theta | \theta \in \Theta_{IVa}) = \frac{A_k + \psi A_{-k}}{(1+\psi)(1 - \min[n, x/c])}$$

One can decompose expected utility from withdrawing late in this region to two components independent from one another,

$$\begin{aligned} \int_{\theta=\theta_i-\epsilon}^{\theta_i+\epsilon} q(\theta) \int_{\theta_k=-\infty}^{\hat{\theta}_1} \int_{\theta_{-k}=-\infty}^{\hat{\theta}_1} u(c_2(\theta | \theta \in \Theta_{IVa})) \phi(\theta_k, \theta, \sigma_\nu) \phi(\theta_{-k}, \theta, \sigma_\nu) d\theta_{-k} d\theta_k d\theta = \\ \int_{\theta_k=-\frac{x}{1-x}}^{\hat{\theta}_1} e^{-\frac{\alpha A_k}{(1+\psi)(1 - \min[n, x/c])}} \phi(\theta_k, \theta, \sigma_\nu) d\theta_k \int_{\theta_{-k}=-\frac{x}{1-x}}^{\hat{\theta}_1} e^{-\frac{\alpha \psi A_{-k}}{(1+\psi)(1 - \min[n, x/c])}} \phi(\theta_{-k}, \theta, \sigma_\nu) d\theta_{-k} \int_{\theta=\theta_i-\epsilon}^{\theta_i+\epsilon} q(\theta) \left(\Phi(\hat{\theta}_1, \theta, \sigma_\nu)^2 - \right. \\ \left. \Phi\left(-\frac{x}{1-x}, \theta, \sigma_\nu\right)^2 \right) d\theta \end{aligned}$$

and then complete the square for each. The expression $\Phi(\hat{\theta}_1, \theta, \sigma_\nu)^2$ comes from adding 1 to the utility function. The lower boundary of the integral is $-\infty$ for the constant component of utility, but is $-\frac{x}{1-x}$ for the one depending on assets, since in that case assets are simply zero.

The coefficients for completing the square are:

$$\begin{aligned} a_4^k &= -\alpha \left(c - \frac{c-1}{1 - \min[n, x/c]} \right) x \Gamma \times \frac{1}{1 + \psi} \\ a_4^{-k} &= a_4^k \times \psi \\ b_4^k &= -\alpha \left(c - \frac{c-1}{1 - \min[n, x/c]} \right) (1-x) \Gamma \frac{1}{1 + \psi} \\ b_4^{-k} &= b_4^k \times \psi \end{aligned}$$

This yields

$$\begin{aligned} \Delta_{IVa} &= \int_{\theta=\theta_i-\epsilon}^{\theta_i+\epsilon} q(\theta) \left\{ \Phi(\hat{\theta}_1, \theta, \sigma_\nu)^2 - \right. \\ &\left[\left(\Phi(\hat{\theta}_1, \theta + b_4^k \sigma_\nu^2, \sigma_\nu) - \Phi\left(-\frac{x}{1-x}, \theta + b_4^k \sigma_\nu^2, \sigma_\nu\right) \right) \times e^{a_4^k + b_4^k \theta [(b_4^k)^2 \sigma_\nu^2]^2 / 2} + \Phi\left(-\frac{x}{1-x}, \theta, \sigma_\nu\right) \right] \times \\ &\left. \left[\left(\Phi(\hat{\theta}_1, \theta + b_4^{-k} \sigma_\nu^2, \sigma_\nu) - \Phi\left(-\frac{x}{1-x}, \theta + b_4^{-k} \sigma_\nu^2, \sigma_\nu\right) \right) \times e^{a_4^{-k} + b_4^{-k} \theta [(b_4^{-k})^2 \sigma_\nu^2]^2 / 2} + \Phi\left(-\frac{x}{1-x}, \theta, \sigma_\nu\right) \right] \right\} d\theta \end{aligned}$$

This expression depends only on θ , reducing the triple to a single integral.

Region IVb

Both banks are insolvent, but k 's insolvency is due to $-k$ and not vice versa. The condition for this is: $A_k + \psi(A_{-k} + B) < D + B \wedge A_{-k} < D$ (mirror of region II in the first inequality). Since both banks insolvent, period 2 consumption is the same as in region IVa:

$$c_2(\theta | \theta \in \Theta_{IVb}) = \frac{A_k + \psi A_{-k}}{(1 + \psi)(1 - \min[n, x/c])}$$

so that k 's integral (the non-constant component) has closed form as before in region IVa (with different boundaries)

$$\Phi(\hat{\theta}_2(A_{-k}), \theta + b_4^k \sigma_\nu^2, \sigma_\nu) - \Phi(\hat{\theta}_1, \theta + b_4^k \sigma_\nu^2, \sigma_\nu)$$

The component depending on $-k$ doesn't have closed form because the boundaries of k 's integral depend on $-k$. We are left with

$$\begin{aligned} \Delta_{IVb} &= \int_{\theta=\theta_i-\epsilon}^{\theta_i+\epsilon} q(\theta) \int_{\theta_{-k}=-\infty}^{\hat{\theta}_1} \\ &\left[\left(\Phi(\hat{\theta}_2(A_{-k}), \theta, \sigma_\nu) - \Phi(\hat{\theta}_1, \theta, \sigma_\nu) \right) - \right. \end{aligned}$$

$$\begin{aligned}
& \left(\Phi(\hat{\theta}_2(A_{-k}), \theta + b_4^k \sigma_\nu^2, \sigma_\nu) - \Phi(\hat{\theta}_1, \theta + b_4^k \sigma_\nu^2, \sigma_\nu) \right) \times e^{a_4^k + b_4^k \theta + (b_4^k)^2 \sigma_\nu^2 / 2} \times \\
& \exp(a_4^{-k} + b_4^{-k} \theta_{-k}) \Big] \phi(\theta_{-k}, \theta, \sigma_\nu) d\theta_{-k} d\theta = \\
& \int_{\theta=\theta_i-\epsilon}^{\theta_i+\epsilon} q(\theta) \left\{ \int_{\theta_{-k}=-\frac{x}{1-x}}^{\hat{\theta}_1} \right. \\
& \left[\left(\Phi(\hat{\theta}_2(A_{-k}), \theta, \sigma_\nu) - \Phi(\hat{\theta}_1, \theta, \sigma_\nu) \right) - \right. \\
& \left. \left(\Phi(\hat{\theta}_2(A_{-k}), \theta + b_4^k \sigma_\nu^2, \sigma_\nu) - \Phi(\hat{\theta}_1, \theta + b_4^k \sigma_\nu^2, \sigma_\nu) \right) \times e^{a_4^k + b_4^k \theta + (b_4^k)^2 \sigma_\nu^2 / 2} \times \right. \\
& \left. \exp(a_4^{-k} + b_4^{-k} \theta_{-k}) \Big] \phi(\theta_{-k}, \theta, \sigma_\nu) d\theta_{-k} + \\
& \left. \Phi\left(-\frac{x}{1-x}, \theta, \sigma_\nu\right) \left[\left(\Phi(\hat{\theta}_2(0), \theta, \sigma_\nu) - \Phi(\hat{\theta}_1, \theta, \sigma_\nu) \right) - \right. \right. \\
& \left. \left. e^{a_4^k + b_4^k \theta + (b_4^k)^2 \sigma_\nu^2 / 2} \times \left(\Phi(\hat{\theta}_2(0), \theta + b_4^k \sigma_\nu^2, \sigma_\nu) - \Phi(\hat{\theta}_1, \theta + b_4^k \sigma_\nu^2, \sigma_\nu) \right) \right] \right\} d\theta
\end{aligned}$$

Region IVc

Both banks are insolvent, but $-k$'s insolvency is due to k and not vice versa. The condition for this is: $A_k < D \wedge A_{-k} + \psi(A_k + B) < D + B$ (mirror image of region III for the second inequality). Since both banks are insolvent, period 2 consumption is the same as in region IVa:

$$c_2(\theta | \theta \in \Theta_{IVc}) = \frac{A_k + \psi A_{-k}}{(1 + \psi)(1 - \min[n, x/c])}$$

so that $-k$'s integral has closed form as before in region IVa (with different boundaries)

$$\Phi(\hat{\theta}_2(A_k), \theta + b_4^{-k} \sigma_\nu^2, \sigma_\nu) - \Phi(\hat{\theta}_1, \theta + b_4^{-k} \sigma_\nu^2, \sigma_\nu)$$

The component depending on k doesn't have closed form because the boundaries of $-k$'s integral depend on k . We are left with

$$\begin{aligned}
\Delta_{IVc} = & \int_{\theta=\theta_i-\epsilon}^{\theta_i+\epsilon} q(\theta) \left\{ \int_{\theta_k=-\frac{x}{1-x}}^{\hat{\theta}_1} \right. \\
& \left[\left(\Phi(\hat{\theta}_2(A_k), \theta, \sigma_\nu) - \Phi(\hat{\theta}_1, \theta, \sigma_\nu) \right) - \right. \\
& e^{a_4^{-k} + b_4^{-k} \theta + (b_4^{-k})^2 \sigma_\nu^2 / 2} \times \left(\Phi(\hat{\theta}_2(A_k), \theta + b_4^{-k} \sigma_\nu^2, \sigma_\nu) - \Phi(\hat{\theta}_1, \theta + b_4^{-k} \sigma_\nu^2, \sigma_\nu) \right) \times \\
& \left. \left. \exp(a_4^k + b_4^k \theta_k) \phi(\theta_k, \theta, \sigma_\nu) d\theta_k \right] + \right.
\end{aligned}$$

$$\Phi\left(-\frac{x}{1-x}, \theta, \sigma_\nu\right) \left[\left(\Phi(\hat{\theta}_2(0), \theta, \sigma_\nu) - \Phi(\hat{\theta}_1, \theta, \sigma_\nu) \right) - e^{a_4^{-k} + b_4^{-k}\theta + (b_4^{-k})^2 \sigma_\nu^2 / 2} \times \left(\Phi(\hat{\theta}_2(0), \theta + b_4^{-k} \sigma_\nu^2, \sigma_\nu) - \Phi(\hat{\theta}_1, \theta + b_4^{-k} \sigma_\nu^2, \sigma_\nu) \right) \right] d\theta$$

A.3.4 Simulation: Graphs

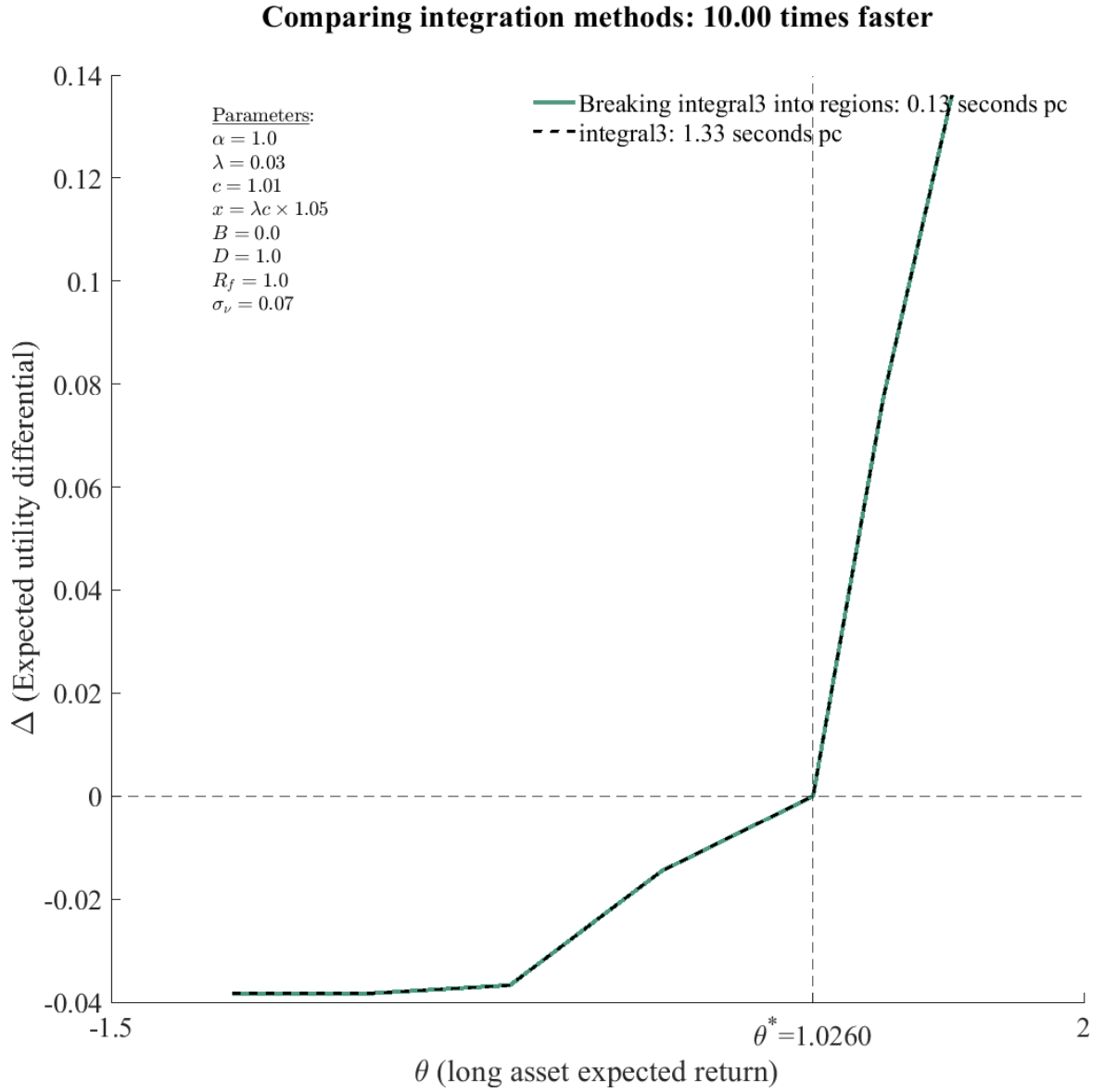


Figure A.1 Comparing integration methods: expected differential utility $a\Delta$

Region	θ_k	θ_{-k}
Θ_I	$\theta_j \geq \overbrace{\left(\frac{D}{1-\tilde{\gamma}_1} \frac{1}{1-\min[nc, x]} - xR_f \right) \frac{1}{1-x}}^{\hat{\theta}_1} \quad \text{for } j \in \{k, -k\}$	
Θ_{II}	$\theta_k \geq \overbrace{\left(\frac{D+B(1-\psi)-\psi A_{-k}}{1-\tilde{\gamma}_1} \frac{1}{1-\min[nc, x]} - xR_f \right) \frac{1}{1-x}}^{\hat{\theta}_2(A_{-k})} \quad \theta_{-k} < \hat{\theta}_1$	
Θ_{III}	$\theta_k < \hat{\theta}_1$	$\theta_{-k} \geq \hat{\theta}_2(A_k)$
Θ_{IVa}	$\theta_j < \hat{\theta}_1 \quad \text{for } j \in \{k, -k\}$	
Θ_{IVb}	$\theta_k < \hat{\theta}_2(A_{-k})$	$\theta_{-k} < \hat{\theta}_1$
Θ_{IVc}	$\theta_k < \hat{\theta}_1$	$\theta_{-k} < \hat{\theta}_2(A_k)$

Table A.1 State-space partition: boundaries

B

B.1 Equity Information

To compute equity returns and market capitalization for each firm, I used Compustat to get closing prices (prccd) and shares outstanding (cshoc) for each security which participates in a firm's earnings' report, as well as a few other auxiliary variables (ajexdi, trfd, curcdd, epf, mkvalincl). Market capitalization (E) in US dollars is computed by multiplying the number of shares outstanding by the dollar price prccd_usd:

$$E = \text{prccd_usd} \times \text{cshoc} \quad \text{prccd_usd} = \text{prccd} \times E_{\text{USD/CCY}}$$

where $E_{\text{USD/CCY}}$ is the exchange rate between the currency CCY the security is traded in, and USD. In order to compute equity daily returns (retd) accurately one has to take account of stock splits, dividend payouts and share buybacks; information on those is contained in ajexdi and trfd; they are used to compute the adjusted closing price (prcadj):

$$\text{retd}_t = \Delta t \ln(\text{prcadj}_t / \text{prcadj}_{t-1}) \quad \text{prcadj}_t = \text{prccd}_t \times \text{trfd}_t / \text{ajexdi}_t$$

where Δt is the number of business days between period t and $t - 1$. The above is performed for all securities that participate in earnings report as documented by Compustat (mkvalincl for North American firms and epf for all others). To arrive at the total market value of a firm, the market capitalization of these securities is summed. The total return of a firm is the weighted average of its individual securities' return. In case a firm has multiple securities participating in its earnings' report, dates when data is available on each security may not always overlap, which results in noisy measurement. To overcome this issue, I completed the time-series of each security so that all relevant securities for computing market capitalization have identical dates, using a linear completion method.

B.2 Term Structure of the Risk Free Rate

To compute expected interest rates, I used information on interest rate swaps¹ and LIBOR² from FRED. Since swap rates are paid by a fixed rate payer in return for receiving three month LIBOR, they contain a direct estimate of market expectations of the LIBOR. These are available for maturities between 1-30 years.

LIBOR rates with maturities up to 12 months³ are used to compute the 3M-LIBOR forward rates; the forward in m months is computed via the relation:

$$\text{FORWARD}(m, 3) = \left(\frac{\text{LIBOR}(m+3)^{(m+3)/12}}{\text{LIBOR}(3)^{3/12}} \right)^{\frac{12}{m}} \quad (\text{B.1})$$

Then the expected average 3M-LIBOR rate is computed as the average forward; for example the expected average 3M-LIBOR in the coming 9 months is computed as

$$\frac{\text{FORWARD}(1/22, 3) + \text{FORWARD}(3, 3) + \text{FORWARD}(6, 3)}{3}$$

The final result is a vector of expected 3-months LIBOR rates at sufficient time-steps, which could then be used in equation 2.6 to discount default intensities.

B.3 Matching Procedure

In order to correlate equity and CDS market information at the company level, I matched Markit's Ticker (CDS) with Compustat gvkey (equity and balance sheet data). The information I had available from Markit was: company name and country. The process of matching was manual, because I learned that automated methods based on company name and little else were unreliable. Markit's Red Code would have been extremely useful in the matching process, as it would have allowed to automating this process. I began the matching procedure from Markit's variable ShortName; Markit abbreviated certain words, e.g. AMERICAN is written AMERN or 'UNITED' is written UTD. I collected a dictionary of all abbreviations as I went along. I then searched the company name in a Compustat database that had

¹ICE Benchmark Administration Limited (IBA), ICE Swap Rates, 11:00 A.M. (London Time), Based on U.S. Dollar, 3 Year Tenor [ICERATES1100USD3Y], retrieved from FRED, Federal Reserve Bank of St. Louis; <https://fred.stlouisfed.org/series/ICERATES1100USD3Y>, July 27, 2021.

²ICE Benchmark Administration Limited (IBA), 3-Month London Interbank Offered Rate (LIBOR), based on U.S. Dollar [USD3MTD156N], retrieved from FRED, Federal Reserve Bank of St. Louis; <https://fred.stlouisfed.org/series/USD3MTD156N>, July 27, 2021

³Rates are available at overnight maturity, 1-2 weeks, 1-12 months.

Ticker	ShortName	Country	gvkey
AMR-AmAirlines	AMERN AIRLS INC	United States	001045
AMR	AMR CORP	United States	001045
AMERAI	AMERN AIRLS GROUP INC	United States	001045
PNW-ARPubSvc	AZ PUB SVC CO	United States	001075
PNW	PINNACLE WEST CAP CORP	United States	001075
ABT	ABBOTT LABS	United States	001078
ABY	ABITIBI CONSOL INC	Canada	001081
AMD	ADVANCED MICRO DEVICES INC	United States	001161
AET	AETNA INC.	United States	001177
APD	AIR PRODS & CHEMS INC	United States	001209
ACV	ALBERTO CULVER CO	United States	001239
AL	ALCAN INC.	Canada	001243
RIOALC	RIO TINTO ALCAN INC	Canada	001243
RICOUSA	RICOH USA INC	United States	001246
IKN	IKON OFFICE SOLUTIONS INC	United States	001246
AYE	ALLEGHENY ENGY INC	United States	001279
HON-Inc	HONEYWELL INC	United States	001300
HON	HONEYWELL INTL INC	United States	001300
ALTEL	ALLTEL CORP	United States	001318
AA	ALCOA INC.	United States	001356
HESS	HESS CORP	United States	001380
FO	FORTUNE BRANDS INC	United States	001408
FORTBRA	FORTUNE BRANDS HOME SEC INC	United States	001408
BEAMINC	BEAM INC	United States	001408
BEAMSUN	BEAM SUNTORY INC	United States	001408
AEP-ResInc	AEP RES INC	United States	001440
AEP	AMERN ELEC PWR CO INC	United States	001440

Table B.1 List of participating companies (excerpt), Markit Ticker to Compustat gvkey

companies' name, legal name, business description amongst other. I recorded all gvkeys I found (for many Tickers I found multiple matching gvkeys), and later searched the internet for those companies, tracing name changes and mergers & acquisitions. In some cases I used Bloomberg to get additional information about the latter. Compustat often contains multiple gvkeys for the same company. I tracked those by observing records that had little to no data on one or more variables, or if there was no overlap between Markit's data and Compustat. An excerpt from the full list is provided in Table B.1. I will make the full matching table available to referees upon request.

Table B.2 List of large banks

Bank name	Country	Markit	Bloomberg	gvkey
Australia and New Zealand	Australia	ANZ	ANZ.AX	015889
Commonwealth Bank	Australia	CBA	CBA.AX	024512
Macquarie Group Limited	Australia	MAGRL	MQG.AX	212856
National Australia Bank	Australia	NAB	NAB.AX	014802
Westpac	Australia	WSTP	WBK	015362
BAWAG Group AG	Austria	BWGPSK	BG.VI	325711
BKS Bank AG	Austria	NaN	BKS.VI	NaN
Erste Group Bank AG	Austria	ERGBA	EBS.VI	214659
Oberbank AG	Austria	NaN	OBS.VI	015730
Raiffeisen Bank Internati	Austria	RFIAG	RBI.VI	272817
Volksbank Vorarlberg e. G	Austria	NaN	VVPS.VI	224229
Wiener Privatbank SE	Austria	NaN	WPB.VI	202851
KBC Group NV	Belgium	KRCP-BankNV	KBC.BR	015703
Banco do Brasil S.A.	Brazil	BANBRA	BBAS3.SA	030581
Itaú Unibanco	Brazil	ITUNI	ITUB	030580
Bank of Montreal	Canada	BMO	BMO	015580
Bank of Nova Scotia (Scot	Canada	BNS	BNS	015582
Canadian Imperial Bank of	Canada	CM	CM	015581
Royal Bank of Canada	Canada	RY	RY	015633
Toronto-Dominion Bank	Canada	TD	TD	015706
Agricultural Bank of Chin	China	AGRBLT	1288.HK	269783
Bank of Beijing	China	NaN	601169.CH	286109
Bank of China (BOC)	China	BCHINL	3988.HK	267461
Bank of Communications	China	BKCOMM	3328.HK	267462
Bank of Jiangsu Co., Ltd.	China	NaN	600919.SS	322055
Bank of Shanghai Co., Ltd	China	NaN	601229.SS	322886
China CITIC Bank	China	CHINCIT	0998.HK	284458
China Construction Bank (	China	CCBC	0939.HK	274364
China Everbright Bank	China	CHEVE	6818.HK	269784
China Merchants Bank	China	CHINMB	3968.HK	254023
China Minsheng Bank	China	NaN	1988.HK	249482
China Zheshang Bank Co.,	China	NaN	2016.HK	321431
Hang Seng Bank Limited	China	HANSEN	0011.HK	015498
Hua Xia Bank	China	NaN	600015.SS	260010

Industrial Bank (China)	China	NaN	601166.SS	282483
Industrial and Commercial	China	ICBCHN	ICBC	279378
Ping An Bank	China	NaN	000001.SZ	213380
Postal Savings Bank of Ch	China	NaN	1658.HK	322494
Shanghai Pudong Developme	China	SHANPUD	600000.SS	243829
Alm. Brand A/S	Denmark	NaN	ALMB.CO	015513
Danske Bank	Denmark	DANBNK	DANSKE.CO	015552
Jutlander Bank A/S	Denmark	NaN	JUTBK.CO	281731
Jyske Bank A/S	Denmark	JYBC	JYSK.CO	015600
Ringkjobing Landbobank A/	Denmark	NaN	RILBA.CO	213403
Spar Nord Bank A/S	Denmark	NaN	SPNO.CO	211463
Sparekassen Sjælland-Fyn	Denmark	NaN	SPKSJF.CO	320825
Sydbank A/S	Denmark	SYDBK	SYDB.CO	209212
Vestjysk Bank A/S	Denmark	NaN	VJBA.CO	250960
Aktia Pankki Oyj	Finland	NaN	AKTIA.HE	NaN
Evli Pankki Oyj	Finland	NaN	EVLI.HE	320802
Nordea	Finland	NORDBAAD	NDA-FI.HE	015863
BNP Paribas	France	BNP	BNP	015532
Caisse Regionale de Credi	France	NaN	CRLA.PA	211483
Caisse Regionale de Credi	France	NaN	CNF.PA	254058
Caisse Regionale de Credi	France	NaN	CRAP.PA	257078
Caisse regionale de Credi	France	NaN	CRLO.PA	219615
Caisse regionale de Credi	France	NaN	CRAV.PA	269176
Courtois S.A.	France	NaN	COUR.PA	222254
Crédit Agricole	France	ACAFP	ACA.PA	024563
Crédit Mutuel	France	BFCM	227143Z FP	NaN
Groupe BPCE	France	BPCE	3692594Z FP	NaN
Natixis S.A.	France	CCBP-NATIXI	KN.PA	015547
Société Générale	France	SOCGEN	GLE.PA	015784
Allianz SE	Germany	ALZSE	ALV.DE	015724
Commerzbank	Germany	CMZB	CBK.DE	015575
DZ Bank	Germany	DZBK	DZBK GR	267620
Deutsche Bank	Germany	DB	DB	015576
Deutsche Pfandbriefbank A	Germany	HREH-DPFBNK	PBB.DE	213367
Wirecard AG	Germany	NaN	WDI.DE	242158
comdirect bank AG	Germany	NaN	COM.DE	237656

Alpha Bank A.E.	Greece	ALPHBK	ALPHA.AT	062410
Attica Bank S.A.	Greece	NaN	TATT.AT	216582
Bank of Greece	Greece	BKGREE	TELL.AT	200569
Eurobank Ergasias Service	Greece	EUROERG	EUROB.AT	227537
National Bank of Greece S	Greece	NATGRE	ETE.AT	030582
Piraeus Bank S.A.	Greece	PIRB	TPEIR.AT	212811
BOC Hong Kong (Holdings)	Hong Kong	NaN	2388.HK	031648
Axis Bank Limited	India	AXISB	AXISBANK.BO	252278
Bank of Baroda	India	BOB	BANKBARODA.NS	205654
Bank of India Limited	India	BNKIND	BANKINDIA.NS	216028
HDFC Bank Limited	India	NaN	HDB	144535
Punjab National Bank	India	NaN	PNB.NS	255702
State Bank of India	India	SBIIN-StateBkIn	SBIN.NS	203666
Union Bank of India	India	UBIND	UNIONBANK.NS	257156
Bank Negara Indonesia	Indonesia	NaN	BBNI.JK	212970
PT Bank Central Asia Tbk	Indonesia	NaN	BBCA.JK	242865
PT Bank Mandiri (Persero)	Indonesia	NaN	BMRI.JK	257421
PT Bank Rakyat Indonesia	Indonesia	NaN	BBRI.JK	157304
Bank of Ireland Group plc	Ireland	NaN	BIRG.L	063590
Permanent TSB Group Holdi	Ireland	NaN	IL0A.IR	223500
Bank Leumi le- Israel B.M	Israel	BNKLLI	LUMI.TA	002018
Assicurazioni Generali S.	Italy	ASSGEN	G.MI	015804
BPER Banca S.p.A.	Italy	NaN	BPE.MI	200647
Banca Generali S.p.A.	Italy	NaN	BGN.MI	281525
Banca Intesa	Italy	NaN	CABOT	NaN
Banca Mediolanum S.p.A.	Italy	NaN	BMED.MI	212955
Banca Monte dei Paschi di	Italy	MONTE	BMPS.MI	024584
Banca Popolare di Sondrio	Italy	NaN	BPSO.MI	200646
Banco BPM Societa per Azi	Italy	BANCBPM	BAMI.MI	225081
Banco di Desio e della Br	Italy	NaN	BDB.MI	216662
Credito Emiliano S.p.A.	Italy	NaN	CE.MI	015745
Credito Valtellinese S.p.	Italy	NaN	CVAL.MI	201329
Intesa Sanpaolo	Italy	SANPAO	ISP.MI	016348
Mediobanca Banca di Credi	Italy	BACRED	MB.MI	015593
SanPaolo IMI SpA	Italy	NaN	SPI	024589
UniCredit	Italy	UNICI	UCG IM	015549

Unione di Banche Italiane	Italy	UDBI	UBI.MI	270266
Daiwa Securities Group In	Japan	DAIWA	DAIWA	015483
Japan Post Bank	Japan	NaN	7182 JT	320609
Mitsubishi UFJ Financial	Japan	MUFJ	MUFG	252940
Mizuho Financial Group	Japan	MIZUHO-GR	MFG	248136
Nomura Holdings	Japan	NOMURA	NMR	015613
Norinchukin Bank	Japan	NORBK	NORBK	024582
Resona Holdings	Japan	RESONA-Bank	8308.T	015662
Sumitomo Mitsui Financial	Japan	SUMIBK	8316 JT	010137
Sumitomo Mitsui Trust Hol	Japan	SUMITS	8309.T	015651
Hana Financial Group Inc.	Korea	HANABK	086790.KS	274764
Industrial Bank of Korea	Korea	INDKOR	024110.KS	236556
KB Financial Group Inc	Korea	CITNAT	KB	030501
Nonghyup Bank	Korea	NaN	0214764D KS	NaN
Shinhan Bank	Korea	SHNFIN	SHG	025714
Woori Financial Group Inc	Korea	WOORI	WF	252819
CIMB Group Holdings Berha	Malaysia	BCHB-CIMBB	1023.KL	200932
Hong Leong Bank Berhad	Malaysia	HLBKMK	5819.KL	206381
Malayan Banking Berhad	Malaysia	MAYMK	1155.KL	015688
Public Bank Berhad	Malaysia	NaN	1295.KL	015632
RHB Capital Berhad	Malaysia	RHBC-Bk	1066.KL	015742
AmBank	Malaysia	AMMMK-AmBank	1015.KL	NaN
Banco del Bajio, S.A., In	Mexico	NaN	BBAJIOO.MX	291854
Fibra Mty, S.A.P.I. de C.	Mexico	NaN	FMTY14.MX	324337
Fibra Plus	Mexico	NaN	FPLUS16.MX	NaN
Gentera, S.A.B. de C.V.	Mexico	NaN	GENTERA.MX	296458
Grupo Financiero Banorte,	Mexico	NaN	GFNORTEO.MX	211934
Grupo Financiero Inbursa,	Mexico	NaN	GFINBURO.MX	202553
Grupo Financiero Multiva,	Mexico	NaN	GFMULTIO.MX	279454
ABN AMRO Group	Netherl.	AAB	ABN.AS	015504
Aegon N.V.	Netherl.	AEGON	AGN.AS	015598
ING Group	Netherl.	INTNED	ING	015617
Rabobank	Netherl.	COOERAB	RABOBK	252120
Van Lanschot Kempen N.V.	Netherl.	VANLBK-FVanBk	VLK.AS	235457
ABG Sundal Collier Holdin	Norway	NaN	ASC.OL	221787
Aurskog Sparebank	Norway	NaN	AURG.OL	282360

DNB ASA	Norway	DNBAS	DNB.OL	015538
Gjensidige Forsikring ASA	Norway	NaN	GJF.OL	296091
Grong Sparebank	Norway	NaN	GRONG-ME.OL	324762
Helgeland Sparebank	Norway	NaN	HELG.OL	252661
Holand og Setskog Spareba	Norway	NaN	HSPG.OL	252663
Pareto Bank ASA	Norway	NaN	PARB.OL	317286
Sandnes Sparebank	Norway	NaN	SADG.OL	232143
Skue Sparebank	Norway	NaN	SKUE.OL	252667
SpareBank 1 BV	Norway	NaN	SBVG.OL	252539
SpareBank 1 Nord-Norge	Norway	NaN	NONG.OL	210996
SpareBank 1 Ostfold Akers	Norway	NaN	SOAG.OL	274992
SpareBank 1 Ringerike Had	Norway	NaN	RING.OL	252577
SpareBank 1 SMN	Norway	NaN	MING.OL	235916
Sparebanken More	Norway	NaN	MORG.OL	247911
Sparebanken Ost	Norway	NaN	SPOG.OL	NaN
Sparebanken Sor	Norway	NaN	SOR.OL	243378
Sparebanken Vest	Norway	NaN	SVEG.OL	247906
Storebrand ASA	Norway	NaN	STB.OL	015647
Totens Sparebank	Norway	NaN	TOTG.OL	247706
Voss Veksel- og Landmands	Norway	NaN	VVL.OL	250901
BDO Unibank, Inc.	Philippines	NaN	BDOUY	252238
Bank of the Philippine Is	Philippines	NaN	BPHLF	200772
Banco Comercial Portugues	Portugal	NaN	BCP.LS	023848
Qatar Islamic Bank (Q.P.S	Qatar	NaN	QIBK.QA	251223
Qatar National Bank (Q.P.	Qatar	NaN	QNBK.QA	251221
Credit Bank of Moscow (pu	Russia	NaN	CBOM.ME	311549
Public joint stock compan	Russia	NaN	ROSB.ME	278833
Sberbank of Russia	Russia	SBERBANK	SBER.ME	221682
VTB Bank	Russia	NaN	VTBR.ME	177236
DBS Bank	Singapore	DBSSP	MU7.SI	015743
Oversea-Chinese Banking C	Singapore	NaN	O39.SI	015753
United Overseas Bank Limi	Singapore	NaN	U11.SI	015679
Investec Group	Sth Africa	NaN	INVPL	209985
Banco Bilbao Vizcaya Arge	Spain	BBVSM	BBVA	015181
Banco Santander	Spain	SANTNDR	SAN	014140
Banco de Sabadell, S.A.	Spain	BANSAB	SAB.MC	245436

Bankia, S.A.	Spain	BANKSA	BKIA.MC	297957
Bankinter, S.A.	Spain	BKTSM	BKTSM	015846
CaixaBank	Spain	CAXBK	CAIXAB	286879
Liberbank, S.A.	Spain	NaN	LBK.MC	315365
Unicaja Banco, S.A.	Spain	NaN	UNI.MC	324887
Avanza Bank Holding AB (p	Sweden	NaN	AZA.ST	222362
Resurs Holding AB	Sweden	NaN	RESURS.ST	321561
Skandinaviska Enskilda Ba	Sweden	SEB	SEB-A.ST	015671
Svenska Handelsbanken AB	Sweden	SVSKHB	SHB-A.ST	015654
Swedbank AB	Sweden	SWEDBK	SWED-A.ST	024578
Banque Cantonale Vaudoise	Switzerl.	NaN	BCVN.SW	015531
Banque Cantonale de Genève	Switzerl.	NaN	BCGE.SW	211774
Banque Cantonale du Jura	Switzerl.	NaN	BCJ.SW	015844
Basellandschaftliche Kant	Switzerl.	NaN	BLKB.SW	015590
Basler Kantonalbank	Switzerl.	NaN	BSKP.SW	015535
Berner Kantonalbank AG	Switzerl.	NaN	BEKN.SW	016215
Credit Suisse	Switzerl.	CSGAG	CSGN	028838
EFG International AG	Switzerl.	NaN	EFGN.SW	274229
Glarner Kantonalbank	Switzerl.	NaN	GLKBN.SW	317860
Julius Baer Gruppe AG	Switzerl.	NaN	BAER.SW	292854
Luzerner Kantonalbank AG	Switzerl.	NaN	LUKN.SW	015678
St. Galler Kantonalbank A	Switzerl.	NaN	SGKN.SW	245119
Thurgauer Kantonalbank	Switzerl.	NaN	TKBP.SW	317379
UBS Group	Switzerl.	UBS	UBS	144496
VP Bank AG	Switzerl.	NaN	VPBN.SW	016310
Vontobel Holding AG	Switzerl.	NaN	VONN.SW	021481
Walliser Kantonalbank	Switzerl.	NaN	WKBN.SW	NaN
Walliser Kantonalbank	Switzerl.	NaN	ORE6.L	NaN
Bankgkok Bank PCL	Thailand	NaN	BBL	NaN
Kasikornbank Public Compa	Thailand	NaN	KBANK.BK	027833
Krung Thai Bank Public Co	Thailand	NaN	KTB.BK	025607
TMB Bank	Thailand	NaN	TMB	031161
Thanachart Bank	Thailand	NaN	TCAP.BK	201773
The Siam Commercial Bank	Thailand	NaN	SCB.BK	025608
Alliance Trust PLC	U.K.	NaN	ATST.L	200128
Arbuthnot Banking Group P	U.K.	NaN	ARBB.L	208274

Bank of Scotland PLC	U.K.	HBOS	HBOS PLC	015528
Barclays	U.K.	BACR	BCS	012673
Close Brothers Group plc	U.K.	NaN	CBG.L	200883
HSBC Holdings PLC	U.K.	HSBC	HSBC	015509
Lloyds Banking Group	U.K.	LLOYDS-Bank	LLOY LN	015929
Metro Bank PLC	U.K.	NaN	MTRO.L	321365
Nationwide Building Socie	U.K.	ANGLIA	NBS.L	030515
Natwest	U.K.	NATWGRO	NWG US	013294
OneSavings Bank Plc	U.K.	NaN	OSB.L	NaN
PCF Group plc	U.K.	NaN	PCF.L	231878
Paragon Banking Group PLC	U.K.	NaN	PAG.L	221845
Provident Financial plc	U.K.	NaN	PFG.L	015631
Rathbone Brothers Plc	U.K.	NaN	RAT.L	207053
Royal Bank of Scotland Gr	U.K.	RBOS	RBSPLC	015634
Schroders plc	U.K.	NaN	SDR.L	015638
Secure Trust Bank Plc	U.K.	NaN	STB.L	306390
Standard Chartered	U.K.	STAN	STAN.L	015645
Virgin Money UK PLC	U.K.	NaN	VMUK.L	321217
Ambac Financial Group Inc	U.S.	ABK	ABK AA	024287
Bank of America	U.S.	BACF-BankNA	BAC	007647
Capital One	U.S.	COF	COF	030990
Citigroup Inc.	U.S.	C	C	003243
Comerica	U.S.	CMA-Bank	CMA	003231
Commerce Bancshares	U.S.	NaN	CBSH	003238
Cullen/Frost Bankers	U.S.	NaN	CFR	003643
Fifth Third	U.S.	FITB	FITB	004640
Goldman Sachs	U.S.	GS	GS	114628
Huntington Bancshares	U.S.	HBAN-NatBank	HBAN	005786
JPMorgan Chase	U.S.	JPM	JPM	002968
Jefferies	U.S.	JEF	JEF	006682
KeyCorp	U.S.	KEY	KEY	009783
Lehman Brothers	U.S.	LEH	LEH	030128
M&T Bank	U.S.	NaN	MTB	004699
Merrill Lynch	U.S.	MER	MER	007267
Morgan Stanley	U.S.	MWD	MS	012124
Northern Trust	U.S.	NTRS	NTRS	007982

PNC Financial Services	U.S.	PNC	PNC	008245
People's United	U.S.	NaN	PBCT	016245
Regions Financial	U.S.	RGNF	RF	004674
State Street	U.S.	STT	STT	010035
The Bank of New York Mell	U.S.	BNYMEL	BK	002019
Truist Financial	U.S.	TRUIFIN	TFC	011856
U.S. Bancorp	U.S.	USB	USB	004723
Wachovia	U.S.	WB	WBCORP	004739
Wells Fargo	U.S.	WFC	WFC	008007
Zions	U.S.	ZION	ZION	011687

B.4 Charts

Figure B.1 T-years risk neutral PoD (indices), by location

(a) 1-year

(b) 5-years

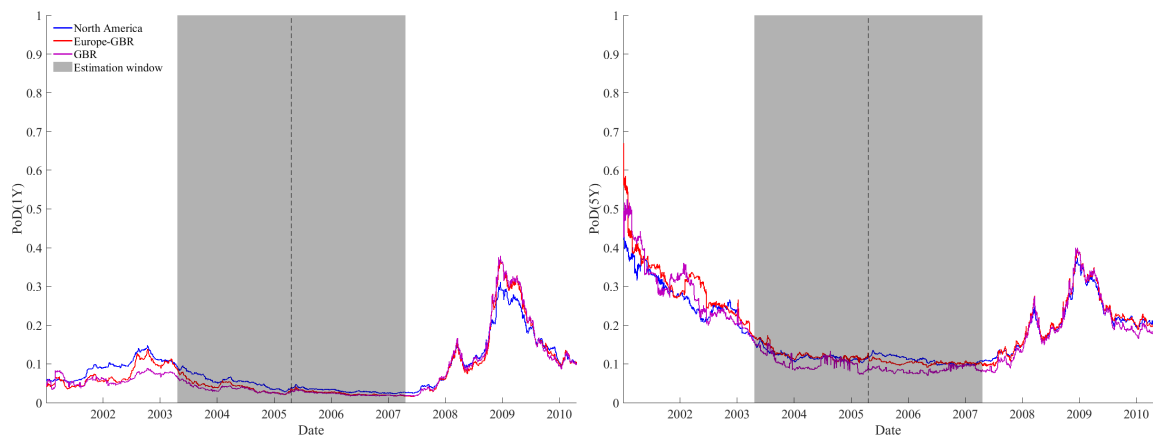


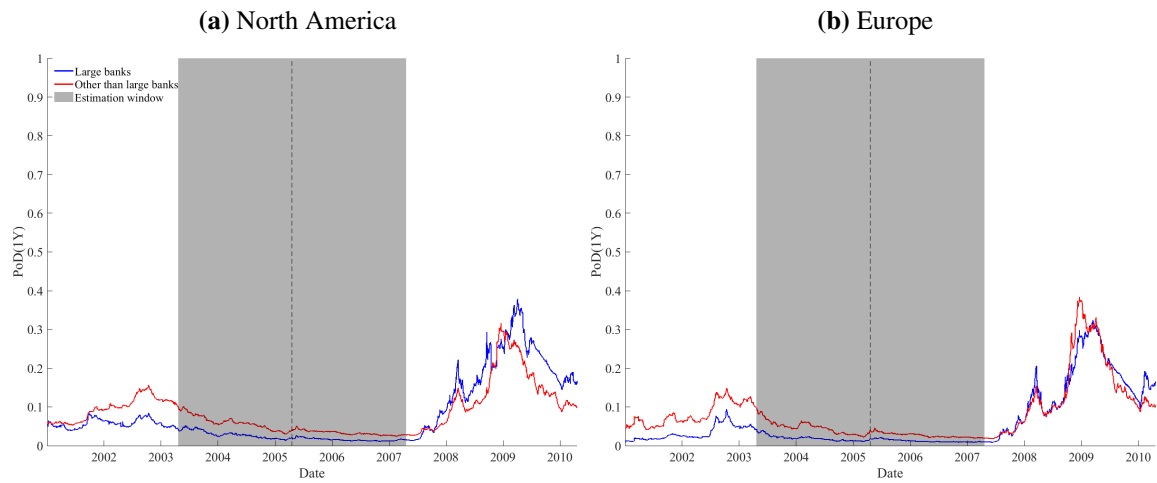
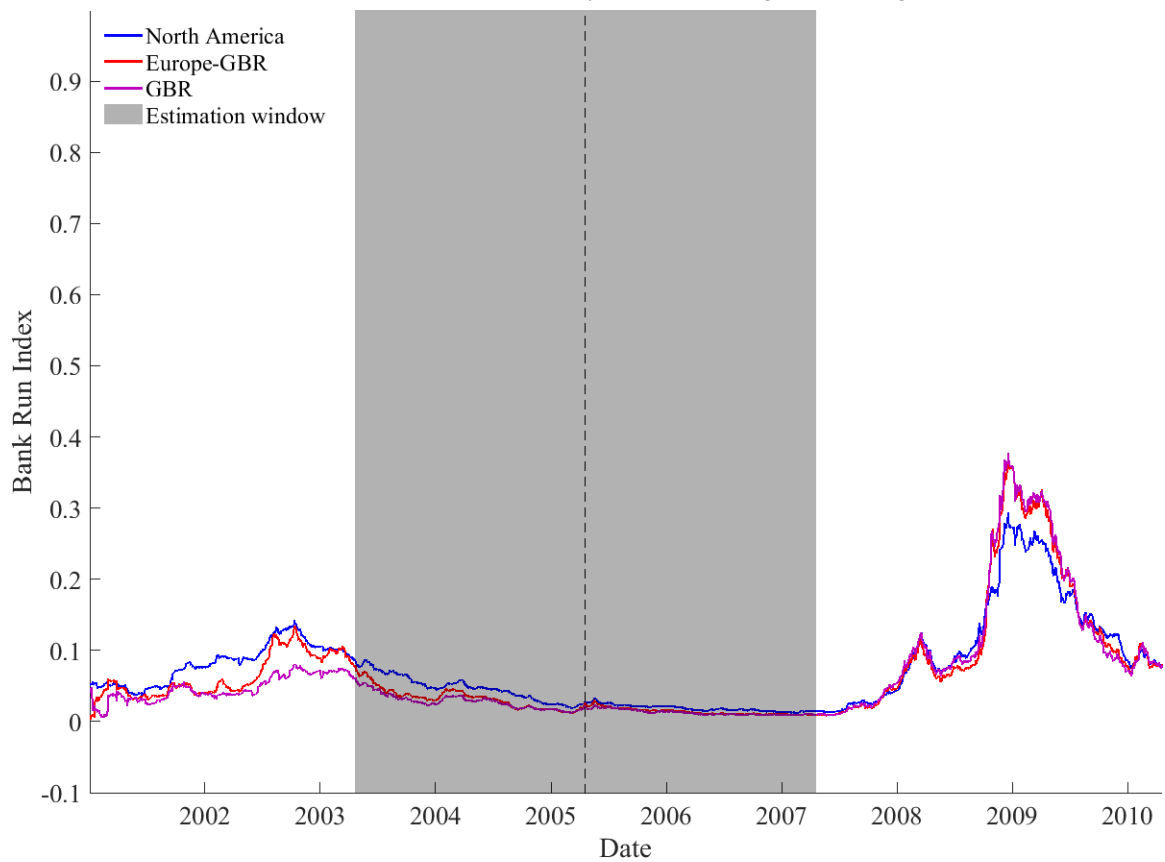
Figure B.2 1-year PoD, large banks and non-large banks**Figure B.3** Bank run index, by location (weighted average)

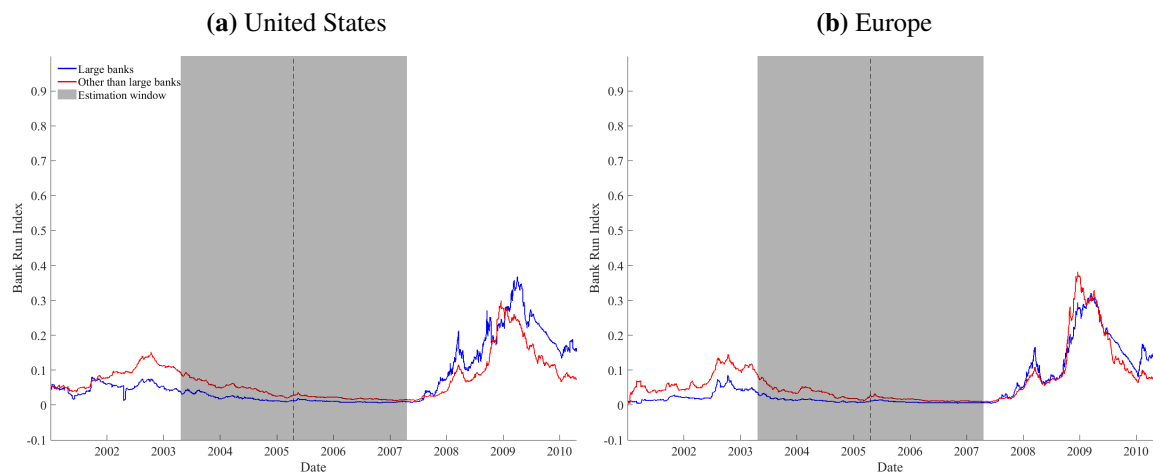
Figure B.4 Bank run index, large banks and non-large banks**Figure B.5** Distance to default, by location (index)

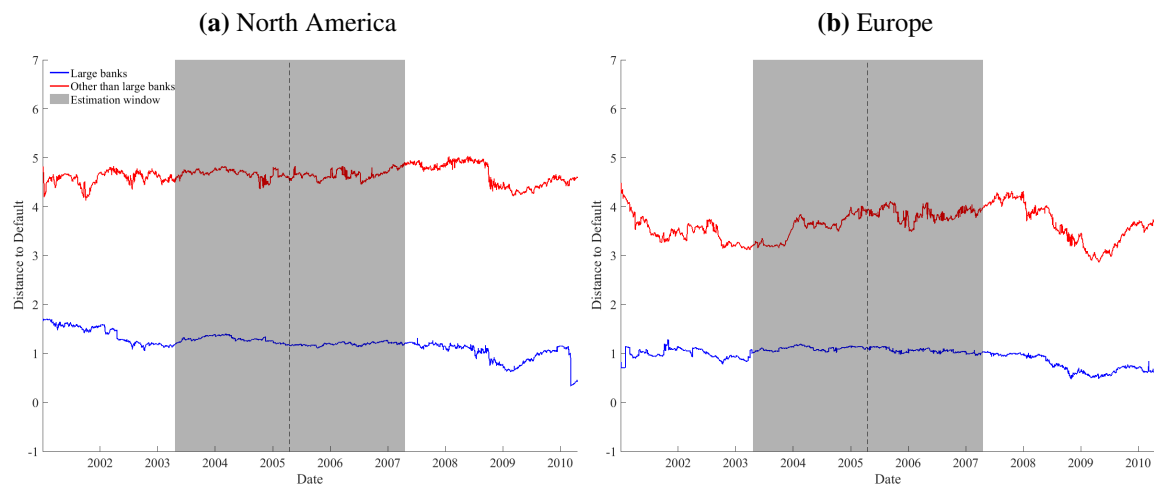
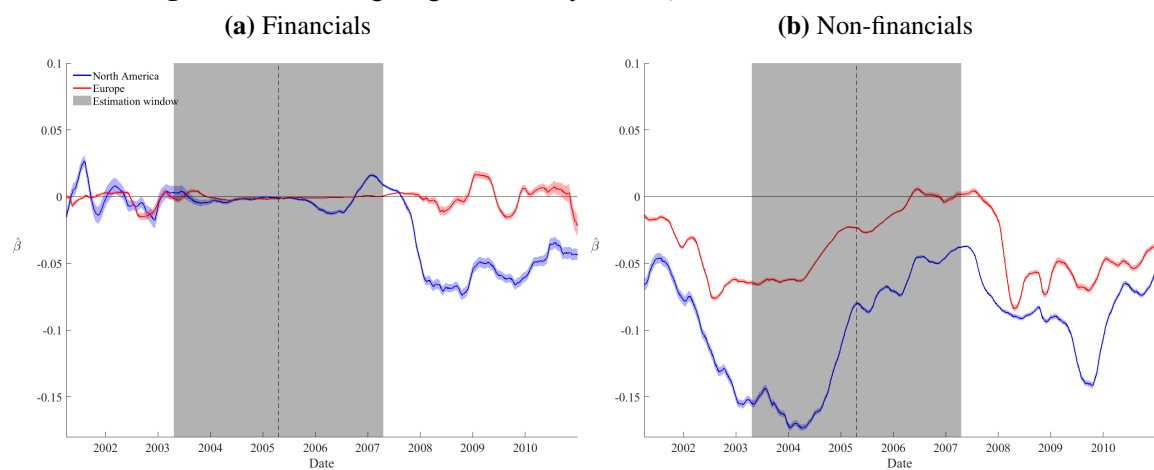
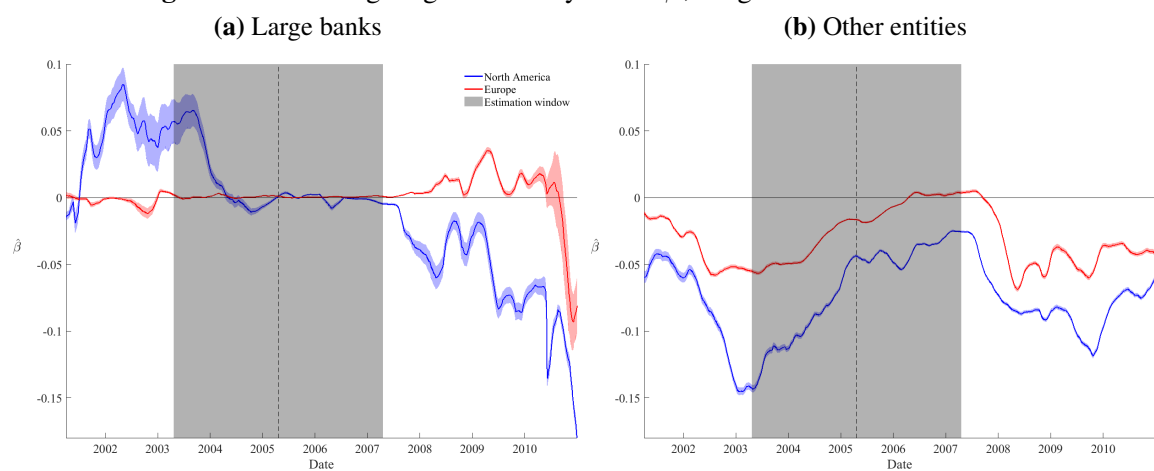
Figure B.6 Distance to default, large banks and non-large banks**Figure B.7** First stage regressions: Systemic- β , Financials vs Non-financials**Figure B.8** First stage regressions: Systemic- β , Large banks vs Other entities

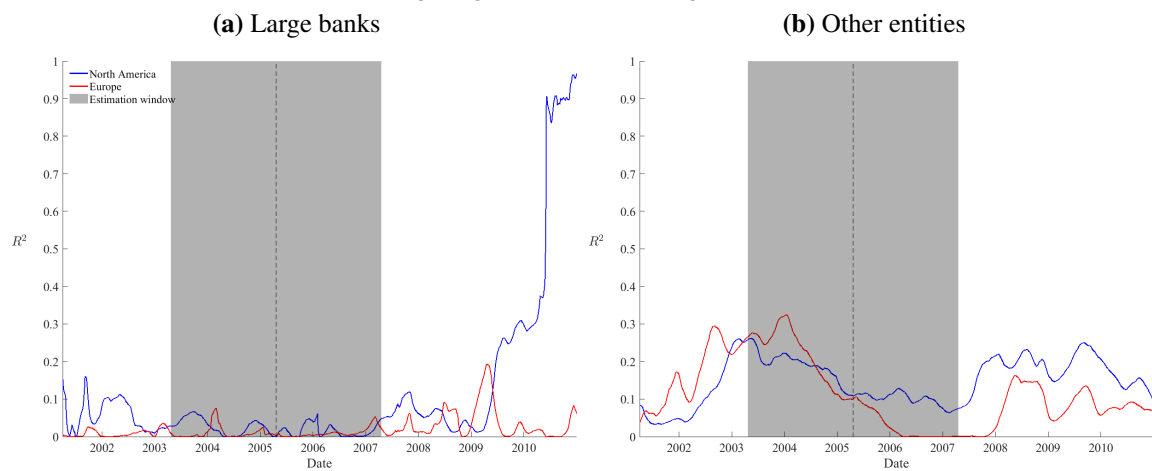
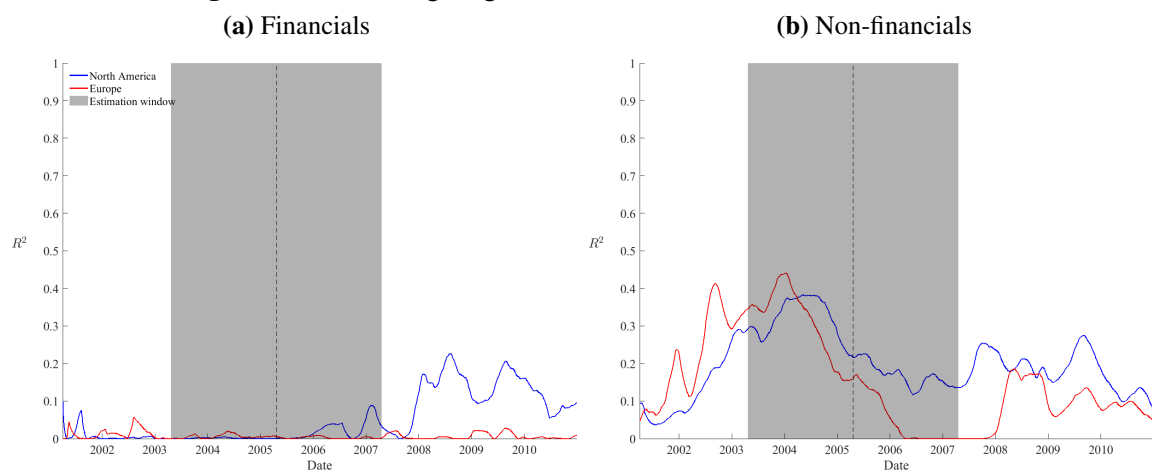
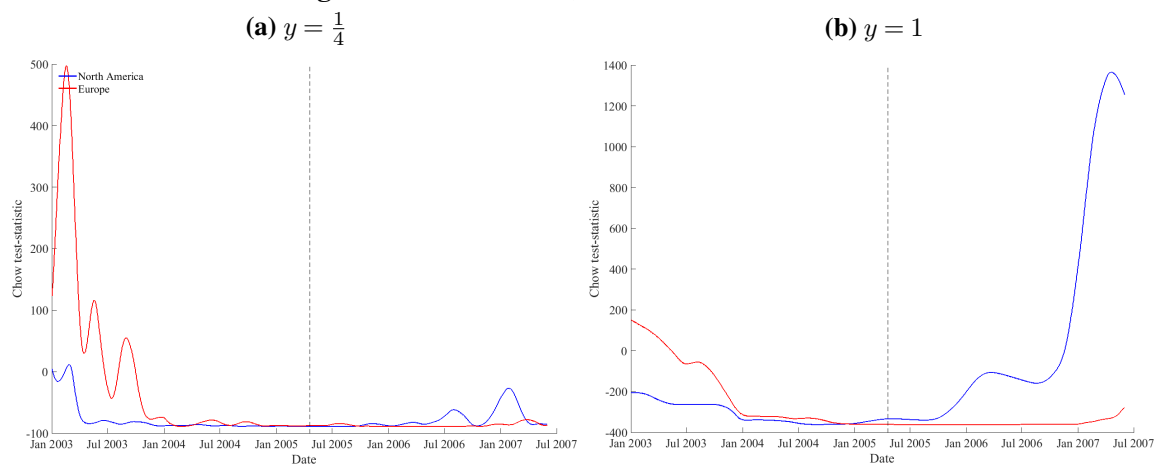
Figure B.9 First stage regressions: R^2 , Large banks vs Other entities**Figure B.10** First stage regressions: R^2 , Financials vs Non-financials**Figure B.11** Chow tests, fixed window: Financials

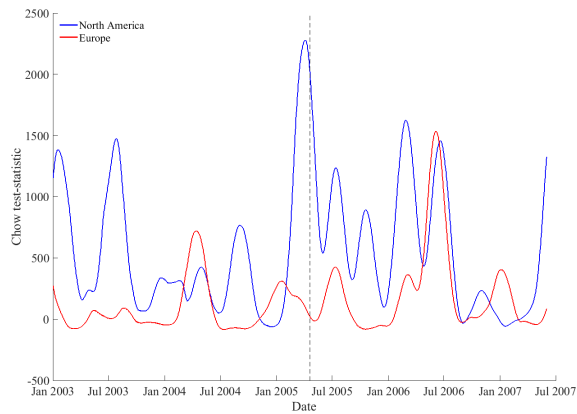
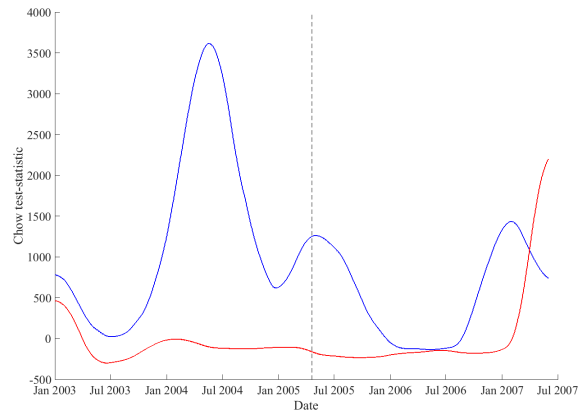
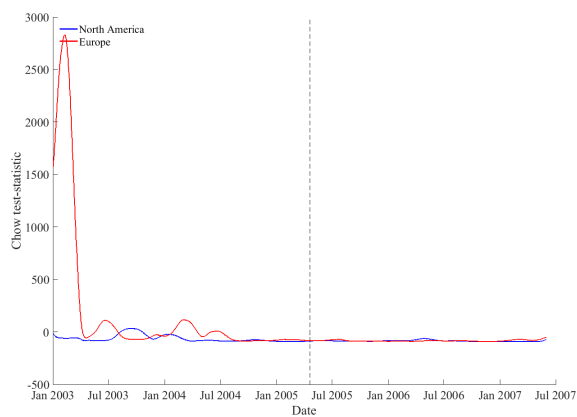
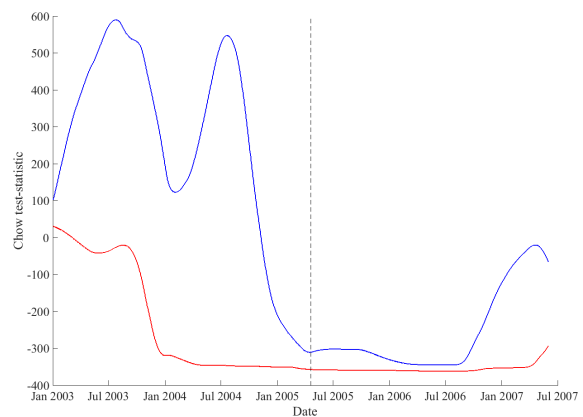
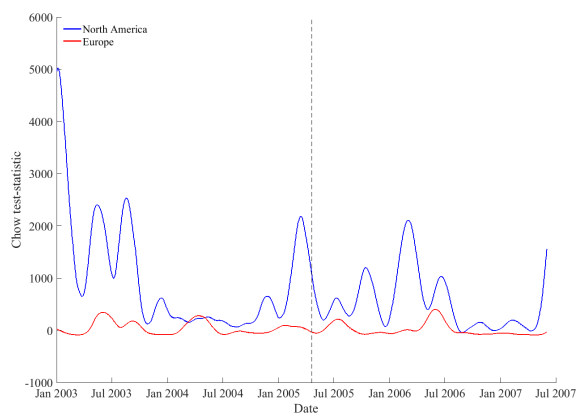
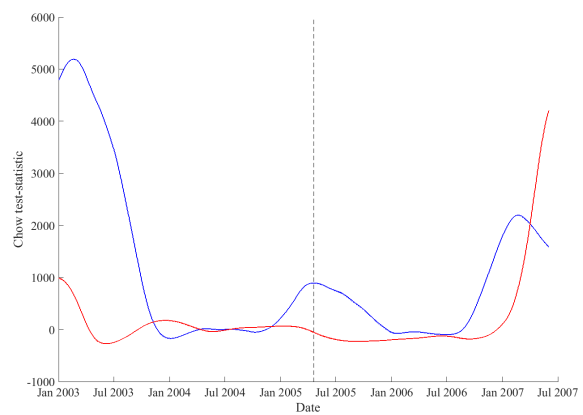
Figure B.12 Chow tests, fixed window: Non-Financials**(a)** $y = \frac{1}{4}$ **(b)** $y = 1$ **Figure B.13** Chow tests, fixed window: Large banks**(a)** $y = \frac{1}{4}$ **(b)** $y = 1$ **Figure B.14** Chow tests, fixed window: Other than large banks**(a)** $y = \frac{1}{4}$ **(b)** $y = 1$ 

Figure B.15 Chow tests, moving window ($y = 1$): Financials vs Non-financials

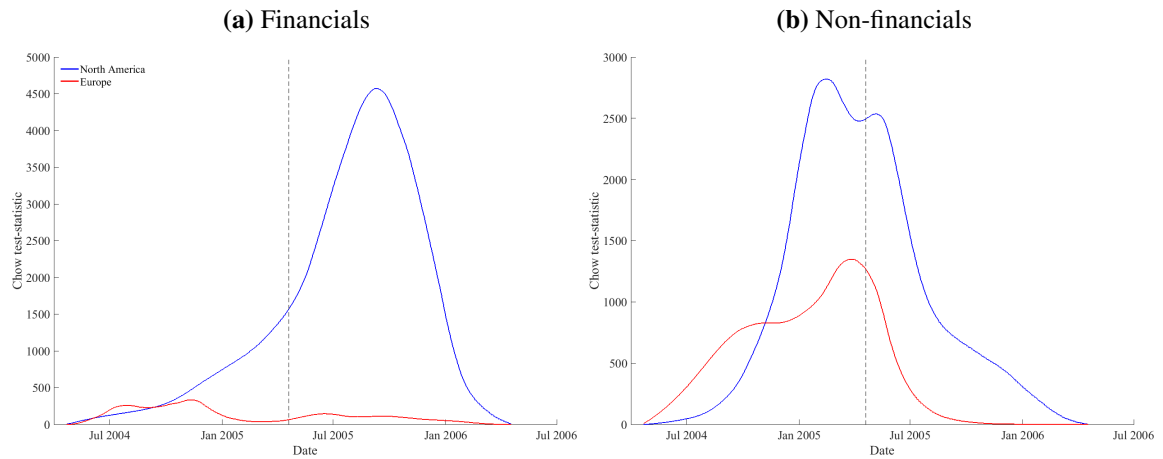
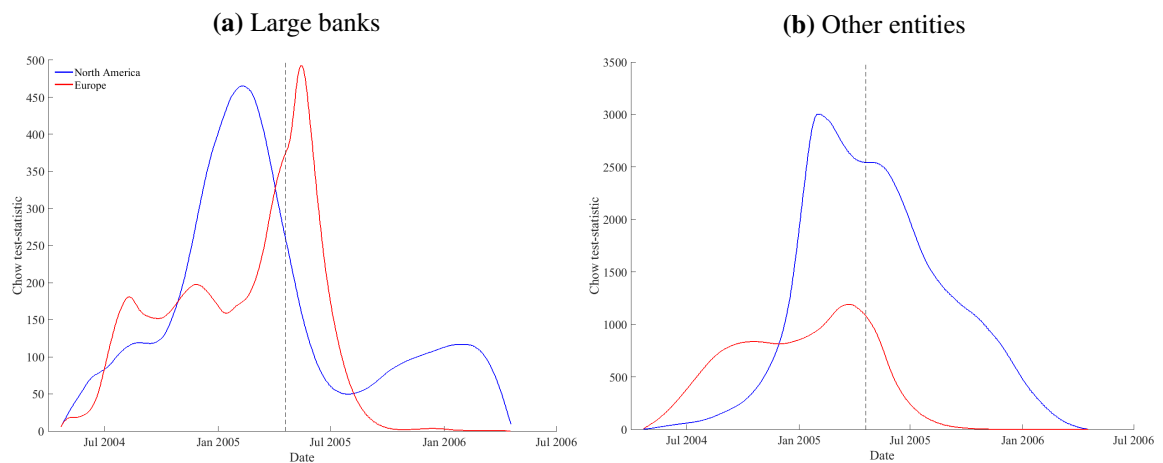


Figure B.16 Chow tests, moving window ($y = 1$): Large banks vs Other entities



C

C.1 Proof of Lemma 1

Lemma 1 shows that firms have strategic complements in quantities. This is a key property for the proof of unique equilibrium.

Proof. Consider equation 3.7 reproduced below

$$\begin{aligned} \pi_i(\theta_i, a) = & (1 - \sigma^{-1}) \frac{1}{2\epsilon} \int_{\tilde{\theta}=\theta_i-\epsilon}^{\theta_i+\epsilon} \left((1 - \mathbb{P}_A(\tilde{\theta}) - \mathbb{P}_B(\tilde{\theta})) \left(\frac{Q_i}{I(\tilde{\theta})/P(\tilde{\theta})} \right)^{-\frac{1}{\sigma}} + \right. \\ & \left. \mathbb{P}_A(\tilde{\theta}) \left(\frac{Q_i(\theta_i)}{\bar{Q}} \right)^{-\frac{1}{\sigma}} + \mathbb{P}_B(\tilde{\theta}) \left(\frac{Q_i(\theta_i)}{\underline{Q}} \right)^{-\frac{1}{\sigma}} - \left(\frac{Q_i}{I(\tilde{\theta})/P(\tilde{\theta})} \right)^{\frac{1}{\alpha}-1} \right) d\tilde{\theta} = 0 \end{aligned}$$

This equation is decreasing in Q_i , $\frac{\partial \pi_i}{\partial Q_i} < 0$; a sufficient condition for it to be ‘increasing’ in real income I/P is that $\alpha \leq 1$. More formally, consider $\hat{\pi} = \frac{\partial \Pi(\theta_i, \hat{a})}{\partial Q_i}$ as the first order condition under an alternative strategy profile $\hat{a}$. Recall that real income is equal to the composite good and $\frac{\partial I/P}{\partial Q(\theta)} \geq 0$. Let $\theta(Q)$ be the inverse of the optimal strategy $Q(\theta)$. Then we have

$$\begin{aligned} \int_{\tilde{Q}=\hat{Q}(\theta_i-\epsilon)}^{Q(\theta_i+\epsilon)} \frac{1 - \mathbb{P}_A(\theta(\tilde{Q})) - \mathbb{P}_B(\theta(\tilde{Q}))}{\frac{I}{P}(\theta(\tilde{Q}))^{-1/\sigma}} dF_{\theta_i}(\tilde{Q}) &\geq \int_{\tilde{Q}=\hat{Q}(\theta_i-\epsilon)}^{Q(\theta_i+\epsilon)} \frac{1 - \mathbb{P}_A(\theta(\tilde{Q})) - \mathbb{P}_B(\theta(\tilde{Q}))}{\frac{I}{P}(\theta(\tilde{Q}))^{-1/\sigma}} d\hat{F}_{\theta_i}(\tilde{Q}) \\ \int_{\tilde{Q}=\hat{Q}(\theta_i-\epsilon)}^{Q(\theta_i+\epsilon)} \frac{1}{\frac{I}{P}(\theta(\tilde{Q}))^{\frac{1}{\alpha}-1}} dF_{\theta_i}(\tilde{Q}) &\leq \int_{\tilde{Q}=\hat{Q}(\theta_i-\epsilon)}^{Q(\theta_i+\epsilon)} \frac{1}{\frac{I}{P}(\theta(\tilde{Q}))^{\frac{1}{\alpha}-1}} d\hat{F}_{\theta_i}(\tilde{Q}) \end{aligned}$$

The result $\pi(\theta_i) \geq \hat{\pi}(\theta_i)$ follows directly from these two inequalities.

□

C.2 Proof of Proposition 3

Proposition 3 follows Frankel et al., (FMP) to show that a unique symmetric equilibrium exists in this model.

Proof. This proposition follows from Theorem 5 in FMP. In what follows I show that each of the five assumptions made by FMP to show unique equilibrium is kept in my model. I start with Lipschitz continuity.

Recall $\pi(\theta_i)$ from equation 3.7, the first derivative of firm i 's payoff in her own action.

$$\begin{aligned} \pi_i(\theta_i) = & (1 - \sigma^{-1}) \frac{1}{2\epsilon} \int_{\tilde{\theta}=\theta_i-\epsilon}^{\theta_i+\epsilon} \left((1 - \mathbb{P}_A(\tilde{\theta}) - \mathbb{P}_B(\tilde{\theta})) \left(\frac{Q_i}{I(\tilde{\theta})/P(\tilde{\theta})} \right)^{-\frac{1}{\sigma}} + \right. \\ & \left. \mathbb{P}_A(\tilde{\theta}) \left(\frac{Q_i(\theta_i)}{\bar{Q}} \right)^{-\frac{1}{\sigma}} + \mathbb{P}_B(\tilde{\theta}) \left(\frac{Q_i(\theta_i)}{\underline{Q}} \right)^{-\frac{1}{\sigma}} - \left(\frac{Q_i}{I(\tilde{\theta})/P(\tilde{\theta})} \right)^{\frac{1}{\alpha}-1} \right) d\tilde{\theta} = 0 \end{aligned}$$

To show Lipschitz continuity it is sufficient to show that the first derivative is bounded. The firm's marginal real revenue is increasing in $\frac{1}{2\epsilon} \int_{\tilde{\theta}=\theta_i-\epsilon}^{\theta_i+\epsilon} \frac{I}{P}(\tilde{\theta}) d\tilde{\theta}$, and the marginal cost is decreasing in it because the average price is increasing in total production. Expected marginal real revenue is bounded from above by 1 and from below by 0. The marginal real cost is bounded from below by 0, and from below by

$$(1 - \sigma^{-1}) \left(\frac{\bar{Q}}{\underline{Q}} \right)^{\frac{1}{\alpha}-1}$$

All of the bounds above are finite, therefore firms' payoffs are Lipschitz continuous in their own action.

For the second Lipschitz condition it is sufficient to show that the derivative of π_i wrt real income is bounded. Consider two alternative strategy profiles $a, \hat{a}$, and the first derivative $\pi_i, \hat{\pi}_i$. Without loss of generality, assume $a \geq \hat{a}$. We need to show that, at some arbitrary $\theta_i \in \Theta$, the following is bounded

$$\int_{\tilde{Q}=\hat{Q}(\theta_i-\epsilon)}^{Q(\theta_i+\epsilon)} \frac{1 - \mathbb{P}_A(\theta(\tilde{Q})) - \mathbb{P}_B(\theta(\tilde{Q}))}{\frac{I}{P}(\theta(\tilde{Q}))^{-1/\sigma}} dF_{\theta_i}(\tilde{Q}) - \int_{\tilde{Q}=\hat{Q}(\theta_i-\epsilon)}^{Q(\theta_i+\epsilon)} \frac{1 - \mathbb{P}_A(\theta(\tilde{Q})) - \mathbb{P}_B(\theta(\tilde{Q}))}{\frac{I}{P}(\theta(\tilde{Q}))^{-1/\sigma}} d\hat{F}_{\theta_i}(\tilde{Q}) -$$

$$\int_{\tilde{Q}=\hat{Q}(\theta_i-\epsilon)}^{Q(\theta_i+\epsilon)} \frac{1}{\frac{I}{P}(\theta(\tilde{Q}))^{\frac{1}{\alpha}-1}} dF_{\theta_i}(\tilde{Q}) + \int_{\tilde{Q}=\hat{Q}(\theta_i-\epsilon)}^{Q(\theta_i+\epsilon)} \frac{1}{\frac{I}{P}(\theta(\tilde{Q}))^{\frac{1}{\alpha}-1}} d\hat{F}_{\theta_i}(\tilde{Q})$$

From Lemma 1 we know that 0 is a lower bound for this expression. An upper bound is given by $\bar{Q}^{1/\sigma} + \frac{1}{Q^{1/\alpha-1}}$. This concludes the part of the Lipschitz conditions. Continuity with respect to own action and to total production is easy to show, since both the average price and the inverse demand function are continuous in these variables. Continuity with respect to the signal follows from continuity of the functions $\mathbb{P}_z$ for $z \in \{A, B\}$. This concludes the proof. $\square$

C.3 Proof of Proposition 4

In this section I prove Proposition 4, that says that the variation in equilibrium strategy $Q(\theta)$ is non-negligible – even when the variation in fundamentals (exogenous relevant payoff variables) is negligible. In other words, there exist states θ_0, θ_1 such that the ratio between the change in $Q(\theta_1) - Q(\theta_0)$ to the change in a payoff relevant variable $\mathbb{P}_z(\theta_1) - \mathbb{P}_z(\theta_0)$ – the probability of an event that entails a finite payoff – that this ratio is infinite: a ‘small’ change in the probability of this event implies a ‘large’ change in equilibrium strategy. Formally, define the set $\chi(x)$, in which the probability of $\mathbb{P}_z(\theta)$ is larger than x , for any $x > 0$,

$$\chi(x) = \{\theta : \mathbb{P}_A(\theta) > x\}$$

In the limit in which the measure of that set, $|\chi(x)|$, is taken to zero, $Var[Q]$ is bounded from below by an expression that is proportional only to $\bar{Q} - \underline{Q}$ (it does not depend on x):

$$\lim_{|\chi(x)| \rightarrow 0} Var[Q] \geq \frac{1}{12} [\bar{Q} - \underline{Q}]^2$$

The proof is written in several steps. First, I study the properties of the fixed point map $Q(\theta)$, and show that under conditions of symmetry between the probabilities of events A and B, this map has to go through the mid-point at $\theta = 0$, $Q(0) = [\bar{Q} - \underline{Q}]/2$. I do this by representing the convergence of $Q(\theta)$ as a fictitious Markov process with two absorbing states, A and B. Transition between states can be thought of as the reasoning-process of an agent contemplating higher order beliefs of other agents.

$Q(\theta)$ is closely related to the conditional probability that this process ‘stops’ at event z , $\psi^z(\theta)$ given that it starts in state θ ; this is shown in Lemma 2. Under conditions of symmetry, ψ^z goes through 1/2 exactly at zero, which then implies $Q(0) = [\bar{Q} - \underline{Q}]/2$.

Lastly, I show that the second derivative of $\psi^z(\theta)$, and therefore of $Q(\theta)$, is larger than zero for $\theta < 0$ and lower than zero for $\theta > 0$. This then implies that the variance $Var[Q]$ is bounded from below by the one implied if $Q(\theta)$ was a straight line on the state-space between zero and one.

Proof of Proposition 4. Consider the following fictitious Markov process, describing how higher order beliefs determine equilibrium outcomes. Studying this convergence process will assist us in making statements about the fixed map. The process is defined for the 1-dimensional state-space Θ . Given an arbitrary starting point at time t , $\theta^t \in \Theta$ in which the process was not absorbed, the process evolves in discrete time: it transitions into another transitory state in accordance with some transition density $p(\theta^t, \theta^{t+1})$, and then is either ‘absorbed’ into state $z \in \{A, B\}$ with probability $\mathbb{P}_z(\theta^{t+1})$ – in which case the process stops – or it is not-absorbed with probability $1 - \sum_z \mathbb{P}_z(\theta^t)$ – in which case the process continues afresh from θ^{t+1} . States A and B can be thought of simply as the extremities of Θ if $\frac{\partial p(\theta, A)}{\partial \theta} \geq 0$ and $\frac{\partial p(\theta, B)}{\partial \theta} \leq 0$, thus A and B are names for collapsing the extremities of Θ into single points.¹

Let z^t be the event of the process being absorbed to state z at time t ; the probability of being absorbed at time t depends only on θ^t . Thus the probability of the joint event that the process will transition from θ^t to $\theta^{t+1} = \tilde{\theta}$ and won’t be absorbed is

$$\mathbb{P}(\theta^{t+1} = \tilde{\theta} \cap \neg(A^{t+1} \cup B^{t+1}) \mid \theta^t) = p(\theta^t, \tilde{\theta}) \mathbb{P}(\neg(A^{t+1} \cup B^{t+1}) \mid \theta^{t+1} = \tilde{\theta})$$

This statement is closely related to the assumption that the noise in agents’ signals is independent from the stochastic process governing fundamentals. Let $\psi^z(\theta)$ be the conditional probability that the process stops in state z given it begins in state θ at time zero. Let $\mathcal{G}$ be the set of all possible densities on the interval Θ , and $\psi^z \in \mathcal{G}$. Define the operator $G : \mathcal{G} \rightarrow \mathcal{G}$

$$(\psi^z G)(\theta) = \int_{\tilde{\theta} \in \Theta} p(\theta, \tilde{\theta}) \psi^z(\tilde{\theta}) d\tilde{\theta} \quad (\text{C.1})$$

This operator describes the evolution of the stochastic process in the following way: given the current state is θ , what is the probability that in the next period the process will transition to state z .

$$\psi^z(\theta) = \int_{\tilde{\theta} \in \Theta} p(\theta, \tilde{\theta}) \mathbb{P}_z(\tilde{\theta}) d\tilde{\theta} + \quad (\text{C.2})$$

¹Whether the probability of event z is everywhere non-zero is not a necessary condition, though makes it easier to prove that the function is convex.

$$\begin{aligned} & \int_{\tilde{\theta} \in \Theta} p(\theta, \tilde{\theta}) \left[1 - \sum_z \mathbb{P}_z(\tilde{\theta}) \right] \int_{\tilde{\theta} \in \Theta} p(\tilde{\theta}, \tilde{\theta}) \mathbb{P}_z(\tilde{\theta}) d\tilde{\theta} d\tilde{\theta} + \dots \\ \psi^z(\theta) = & \mathbb{E}[\mathbb{P}_z(\tilde{\theta}) \mid \theta] + \mathbb{E} \left[\left(1 - \sum_{u \in A, B} \mathbb{P}_u(\tilde{\theta}) \right) (\psi^z G)(\tilde{\theta}) \mid \theta \right] + \dots \end{aligned} \quad (\text{C.3})$$

$$\psi_0^z(\theta) = (\mathbb{P}_z G) \mathbb{P}_z(\tilde{\theta}) d\tilde{\theta}, \quad \psi_{t+1}^z(\theta) = \left(\left[1 - \sum_z \mathbb{P}_z(\tilde{\theta}) \right] \psi_t^z(\tilde{\theta}) G \right)(\theta) \quad (\text{C.4})$$

$$\psi^z(\theta) = \sum_{t=0}^{\infty} \psi_t^z(\theta) \quad (\text{C.5})$$

Lemma 8 (Properties of ψ^z , ψ_t^z and p). 1. $\sum_{z \in \{A, B\}} \psi^z(\theta) = 1 \quad \forall \theta$,

2. $p(\theta, \tilde{\theta}) = p(-\theta, -\tilde{\theta})$
3. Let θ_0 and θ_1 be two states such that $\theta_1 \geq \theta_0$. Then we have $p(\theta_0, \theta_0 + \delta) = p(\theta_1, \theta_1 + \delta) \quad \forall \delta \in [\underline{\theta} - \theta_0, \bar{\theta} - \theta_1]$.
4. $\psi_t^A(\theta) = \psi_t^B(-\theta) \Rightarrow \psi^A(\theta) = \psi^B(-\theta)$,
5. $\psi^A(0) = \psi^B(0) = 0.5$
6. $\frac{\partial \psi^A(\theta)}{\partial \theta} \geq 0$

Proof of Lemma 8. 1. The process only stops when it reaches either A or B. Given that the probability $\mathbb{P}_z(\theta) > 0$ for some θ , given enough time, the process will stop. Therefore the probability of stopping at either A or B is 1.

2. This is a result of uniformly distributed noise. In effect, $p(\theta, \tilde{\theta}) = 1 \quad \forall \tilde{\theta} \in [\theta - \epsilon, \theta + \epsilon]$ and $p(\theta, \tilde{\theta}) = 0 \quad \forall \tilde{\theta} \notin [\theta - \epsilon, \theta + \epsilon]$.
3. This is a result of uniformly distributed noise. Then if the ‘distance’ between the signal θ_0 and a state is δ , the posterior probability that state $\theta_0 + \delta$ occurred is as likely as any other state $\theta_1 + \delta$ when the signal is θ_1 .
4. By induction. Verify that the statement is true for $t = 0$ by assumption 4.1, $\psi_0^A = \mathbb{E}[\mathbb{P}_A(\tilde{\theta}) \mid \theta] = \mathbb{E}[\mathbb{P}_B(\tilde{\theta}) \mid -\theta] = \psi_0^B$. Assume that the statement is true for t : $\psi_t^A(\theta) = \psi_t^B(-\theta)$. And now, at step $t + 1$ we have

$$\begin{aligned} \psi_{t+1}^A(\theta) = & \left(\left[1 - \sum_z \mathbb{P}_z(\tilde{\theta}) \right] \psi_t^z(\tilde{\theta}) G \right)(\theta) = \int_{\tilde{\theta} \in \Theta} p(\theta, \tilde{\theta}) \left[1 - \sum_z \mathbb{P}_z(\tilde{\theta}) \right] \psi_t^A(\tilde{\theta}) d\tilde{\theta} = (\text{C.6}) \\ & \int_{\tilde{\theta} \in \Theta} p(-\theta, -\tilde{\theta}) \left[1 - \sum_z \mathbb{P}_z(-\tilde{\theta}) \right] \psi_t^B(-\tilde{\theta}) d\tilde{\theta} = \psi_{t+1}^B(-\theta) \end{aligned}$$

which is true because of 2.

5. Combining 1 and 4.

6. This part says that the probability of being absorbed in state A starting from state θ , is increasing in θ . This is simply a result assumption 4.2 which says that the probability of being absorbed in event A (B) is increasing (decreasing) in θ , in addition to result 1, which says that the sum of the probabilities equals 1. Then $\psi^A(\theta) - \psi^B(\theta) \geq 0 \forall \theta \geq 0$, with exact equality only at $\theta = 0$. In the limit as $\epsilon \rightarrow 0$, $\mathbb{E}_\theta[f(\mathbb{P}_A, \mathbb{P}_B)] = f(\mathbb{P}_A, \mathbb{P}_B)$. This then implies that

$$\psi^z(\theta) = \frac{\mathbb{P}_z(\theta)}{\sum_u \mathbb{P}_u} \quad (\text{C.7})$$

$$\frac{\partial \psi^z}{\partial \theta}(\tilde{\theta}) = \frac{\frac{\partial \mathbb{P}_z}{\partial \theta}(\tilde{\theta}) \mathbb{P}_{-z}(\tilde{\theta}) - \frac{\partial \mathbb{P}_{-z}}{\partial \theta}(\tilde{\theta}) \mathbb{P}_z(\tilde{\theta})}{[\sum_u \mathbb{P}_u(\tilde{\theta})]^2} \quad (\text{C.8})$$

which is greater than zero for A and smaller for B. Even without the limit we can prove this, by using induction argument.

□

Let $\mathcal{H}$ be the set of all possible maps from Θ to $[0, \bar{Q}]$. Consider a candidate equilibrium $Q \in \mathcal{H}$. Let $H : \mathcal{H} \rightarrow \mathcal{H}$ be the fixed map operator, defined by Lemma 2

$$(QH)(\theta) = \int_{\tilde{\theta}=\theta-\epsilon}^{\theta+\epsilon} \left((1 - \mathbb{P}_A(\tilde{\theta}) - \mathbb{P}_B(\tilde{\theta}))Q(\tilde{\theta}) + \mathbb{P}_B(\tilde{\theta}) \times \underline{Q} + \mathbb{P}_A(\tilde{\theta})\bar{Q} \right) d\tilde{\theta}$$

An equilibrium satisfies

$$Q(\theta) = (QH)(\theta) \quad \forall \theta \quad (\text{C.9})$$

Observe that $p(\theta_i, \tilde{\theta}) = 1$ for $\tilde{\theta} \in [\theta_i - \epsilon, \theta_i + \epsilon]$ because the noise is uniform. This renders the connection between H and G visible. Using the definition of equilibrium and of the operator to get

$$Q(\theta) = \underline{Q} \times \sum_{t=0}^{\infty} \psi_t^B + \bar{Q} \times \sum_{t=0}^{\infty} \psi_t^A = \psi^B(\theta) \times \underline{Q} + \psi^A(\theta) \times \bar{Q} \quad (\text{C.10})$$

Using Lemma 8, we see that $Q(\theta = 0) = \bar{Q}/2 - \underline{Q}/2$. Thus, we have established that the equilibrium unique optimal strategy implies producing exactly the average of the maximum and minimum levels of production. $\square$

Proof of Lemma 2. The proof relies on the limiting case, $\epsilon \rightarrow 0$, CRS technology and almost perfect competition. Observe that with constant returns to scale, and $\sigma \downarrow 1$, and using equation 3.7, we have

$$Q(\theta) = \int_{\tilde{\theta}=\theta-\epsilon}^{\theta+\epsilon} \left((1 - \mathbb{P}_A(\tilde{\theta}) - \mathbb{P}_B(\tilde{\theta})) \frac{I(\tilde{\theta})}{P(\tilde{\theta})} + \mathbb{P}_A(\tilde{\theta})\bar{Q} + \mathbb{P}_B(\tilde{\theta})\underline{Q} \right) d\tilde{\theta} \quad (\text{C.11})$$

Next, consider equation 3.6 that describes real income as a function of produced goods, reproduced here for convenience:

$$\frac{I(\theta)}{P(\theta)} = \left[\frac{1}{2\epsilon} \int_{\tilde{\theta}=\theta-\epsilon}^{\theta+\epsilon} Q(\tilde{\theta})^{1-\sigma^{-1}} d\tilde{\theta} \right]^{\frac{1}{1-\sigma^{-1}}}$$

We see that real income at state θ is a power-mean of Q in states $\tilde{\theta} \in [\theta - \epsilon, \theta + \epsilon]$. Because ϵ is small, this is a power mean of values that are arbitrarily close to one another – because we know that, in equilibrium, Q is a continuous function). Formally, we have

$$\lim_{\epsilon \rightarrow 0} \left[\frac{1}{2\epsilon} \int_{\tilde{\theta}=\theta-\epsilon}^{\theta+\epsilon} Q(\tilde{\theta})^{1-\sigma^{-1}} d\tilde{\theta} \right]^{\frac{1}{1-\sigma^{-1}}} = Q(\theta) \quad (\text{C.12})$$

Substituting-in this result into equation C.11 we get the desired result. $\square$

Proof of Lemma 4. First, recall that the equilibrium fixed map Q is unique, continuous and increasing (Proposition 3). By Lemma 2 we know that at each state θ , $Q(\theta)$ is a weighted linear combination of three terms: $\bar{Q}$ with weight $\mathbb{E}[\mathbb{P}_A(\tilde{\theta}) \mid \theta]$, $\underline{Q}$ with weight $\mathbb{E}[\mathbb{P}_B(\tilde{\theta}) \mid \theta]$, and a weighted arithmetic average $\int_{\tilde{\theta}=\theta-\epsilon}^{\theta+\epsilon} \left((1 - \mathbb{P}_A(\tilde{\theta}) - \mathbb{P}_B(\tilde{\theta})) Q(\tilde{\theta}) d\tilde{\theta} \right)$ of the values of Q in the neighbourhood of $Q(\theta)$. Note that for any θ_0, θ_1 such that $\theta_1 > \theta_0$ we have $\mathbb{E}[\mathbb{P}_A(\tilde{\theta}) \mid \theta_1] \geq \mathbb{E}[\mathbb{P}_A(\tilde{\theta}) \mid \theta_0]$ and $\mathbb{E}[\mathbb{P}_B(\tilde{\theta}) \mid \theta_0] \geq \mathbb{E}[\mathbb{P}_B(\tilde{\theta}) \mid \theta_1]$, so that the $Q(\theta_1)$ puts a higher weight on the maximum, and a lower weight on the minimum, values of Q – as compared with $Q(\theta_0)$. Let $\Delta > 0$; $\frac{\partial Q}{\partial \theta}(\theta_i)$ is defined as

$$\frac{\partial Q}{\partial \theta}(\theta_i) = \lim_{\Delta \rightarrow 0} \frac{Q(\theta_i + \Delta) - Q(\theta_i)}{\Delta}$$

$$\frac{\partial^2 Q}{\partial \theta^2}(\theta_i) = \lim_{\Delta \rightarrow 0} \frac{\frac{\partial Q}{\partial \theta}(\theta_i + \Delta) - \frac{\partial Q}{\partial \theta}(\theta_i)}{\Delta}$$

We need to show that

$$\frac{\partial^2 Q}{\partial \theta^2}(\theta_i) \geq 0 \quad \forall \theta < 0 \quad (\text{C.13})$$

Recall that in equilibrium Q is proportional to ψ^z (eq. C.10)

$$Q(\theta) = \psi^B(\theta) \times \underline{Q} + \psi^A(\theta) \times \bar{Q} = \underline{Q} + (\bar{Q} - \underline{Q})\psi^A(\theta)$$

because $\sum_z \psi^z(\theta) = 1 \quad \forall \theta$ (Lemma 8.1). Then to prove equation C.13 it is sufficient to show that $\frac{\partial^2 \psi^A}{\partial \theta^2}(\tilde{\theta}) \geq 0 \quad \forall \tilde{\theta} < 0$:

$$\begin{aligned} \frac{\partial^2 \psi^A}{\partial \theta^2}(\tilde{\theta}) = & \quad (\text{C.14}) \\ & \frac{\left[\frac{\partial^2 \mathbb{P}_A}{\partial \theta^2}(\tilde{\theta}) \mathbb{P}_B(\tilde{\theta}) - \frac{\partial^2 \mathbb{P}_B}{\partial \theta^2}(\tilde{\theta}) \mathbb{P}_A(\tilde{\theta}) - 2 \frac{\partial \mathbb{P}_A}{\partial \theta}(\tilde{\theta}) \frac{\partial \mathbb{P}_B}{\partial \theta}(\tilde{\theta}) \right]}{[\sum_u \mathbb{P}_u(\tilde{\theta})]^2} \\ & - \frac{2 \left[\frac{\partial \mathbb{P}_A}{\partial \theta}(\tilde{\theta}) + \frac{\partial \mathbb{P}_B}{\partial \theta}(\tilde{\theta}) \right] \left[\frac{\partial \mathbb{P}_A}{\partial \theta}(\tilde{\theta}) \mathbb{P}_B(\tilde{\theta}) - \frac{\partial \mathbb{P}_B}{\partial \theta}(\tilde{\theta}) \mathbb{P}_A(\tilde{\theta}) \right]}{[\sum_u \mathbb{P}_u(\tilde{\theta})]^3} \geq 0 \quad \forall \tilde{\theta} < 0 \end{aligned}$$

This is true because $\frac{\partial \mathbb{P}_B}{\partial \theta}(\tilde{\theta}) < 0$, because $\frac{\partial \mathbb{P}_A}{\partial \theta}(\tilde{\theta}) + \frac{\partial \mathbb{P}_B}{\partial \theta}(\tilde{\theta}) > 0$ for $\tilde{\theta} < 0$, and because of assumption 4 about symmetry and the rate of increase P_A , which implies

$$-\frac{\partial^2 \mathbb{P}_B}{\partial \theta^2}(\tilde{\theta}) \mathbb{P}_A(\tilde{\theta}) - \frac{\partial \mathbb{P}_A}{\partial \theta}(\tilde{\theta}) \frac{\partial \mathbb{P}_B}{\partial \theta}(\tilde{\theta}) = \mathbb{P}_A(\tilde{\theta}) \frac{\partial \mathbb{P}_B}{\partial \theta}(\tilde{\theta}) \left[-\frac{\frac{\partial^2 \mathbb{P}_B}{\partial \theta^2}(\tilde{\theta})}{\frac{\partial \mathbb{P}_B}{\partial \theta}(\tilde{\theta})} - \frac{\frac{\partial \mathbb{P}_A}{\partial \theta}(\tilde{\theta})}{\mathbb{P}_A(\tilde{\theta})} \right] \geq 0 \quad (\text{C.15})$$

To see this, observe that symmetry of P_A and P_B means that their derivatives are also symmetric. Specifically, their n^{th} derivatives satisfy: $\frac{\partial^n \mathbb{P}_A}{\partial \theta^n}(\tilde{\theta}) = (-1)^n \frac{\partial^n \mathbb{P}_B}{\partial \theta^n}(-\tilde{\theta})$. For our case, $\frac{\partial^2 \mathbb{P}_B}{\partial \theta^2}(\tilde{\theta}) = \frac{\partial^2 \mathbb{P}_A}{\partial \theta^2}(-\tilde{\theta})$ and $\frac{\partial \mathbb{P}_B}{\partial \theta}(\tilde{\theta}) = -\frac{\partial \mathbb{P}_A}{\partial \theta}(-\tilde{\theta})$. Assumption 4.3 then says

$$\left[-\frac{\frac{\partial^2 \mathbb{P}_B}{\partial \theta^2}(\tilde{\theta})}{\frac{\partial \mathbb{P}_B}{\partial \theta}(\tilde{\theta})} - \frac{\frac{\partial \mathbb{P}_A}{\partial \theta}(\tilde{\theta})}{\mathbb{P}_A(\tilde{\theta})} \right] \leq 0$$

Again, this expression relies on $\mathbb{P}_z$ strictly being positive everywhere, and on the fact that the noise ϵ is small. A proof using induction can be made even without the requirement that $\mathbb{P}_z$ is strictly positive everywhere (where the expression above may not be defined. This

would still rely on the noise being small, and thus that the probability of no event happening $1 - \sum \mathbb{P}_z$ is close to zero for neighbouring states.

That the second derivative of $Q(\theta)$ is positive for $\theta < 0$ implies, due to symmetry, that it is negative for $\theta > 0$. This then means that a lower bound for $Var[Q(\theta)]$ is given by the variance if Q was a straight line from $\underline{Q}$ to $\bar{Q}$:

$$\underline{Q} + (\bar{Q} - \underline{Q}) \frac{\theta - \underline{\theta}}{\bar{\theta} - \underline{\theta}} \quad (\text{C.16})$$

This is because the unconditional expectation of Q , $\mathbb{E}Q$ is simply $Q(0) = (\bar{Q} + \underline{Q})/2$. A convex function on $\theta < 0$ is always lower than the straight line, and conversely for $\theta > 0$. Thus

$$|Q(\theta) - Q(0)| > |\underline{Q} + (\bar{Q} - \underline{Q}) \frac{\theta - \underline{\theta}}{\bar{\theta} - \underline{\theta}} - (\bar{Q} + \underline{Q})/2| \quad (\text{C.17})$$

This completes the proof. $\square$

C.3.1 Numerical Example

It is perhaps of use to provide a concrete example of the above mathematical construction. Consider the following composition of the exponential function, $P_z : [-1, 1] \rightarrow [0, 1]$

$$P_A(\theta) = \left(1 + \exp(-2\eta) \frac{1}{1 - \exp(-2\eta)} \right) \exp(\eta(\theta - 1)) - \exp(-2\eta) \frac{1}{1 - \exp(-2\eta)}$$

$$P_B(\theta) = 1 + \exp(-2\eta) \frac{1}{1 - \exp(-2\eta)} - \left(1 + \exp(-2\eta) \frac{1}{1 - \exp(-2\eta)} \right) \exp(\eta(\theta - 1))$$

Note that I defined $\bar{\theta} = 1$ and $\underline{\theta} = -1$. This function satisfies all of the assumptions I've made in the paper. It increases from zero to one along the interval as defined, it is convex, and since it's higher derivatives all have constant rate of increase η , while the function P_A itself has a rate of increase at least as large as that, assumption 4.3 is also satisfied.

η is a curvature parameter. Figure C.1 illustrates two choices for η : the higher it is, the lower is the measure of the set $\chi(x)$ for any x , while the slope of S -shape becomes steeper.

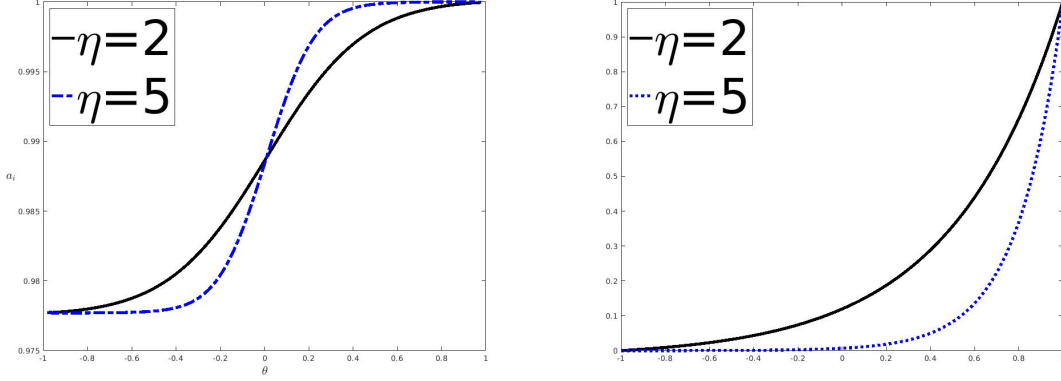
C.4 Optimal Asset Prices, Derivation

In the case the asset is cum dividend, the young generation consumes the endowment (labour income), plus the dividend minus the asset price. Since both young and old have the same

Figure C.1 Graphical example, varying η

(a) Production intensity

(b) Probability of event A



taste for variety, their consumption of the different goods is identical up to multiplication by a constant (as they differ in incomes). I therefore treat asset prices in terms of the composite good. I also make the assumption that realizations of θ are iid over time, so that the future always looks the same (i.e. not contingent on the state). Period t young agents maximize expected utility from consumption in both periods

$$\begin{aligned} \max_{\phi_t} \mathbb{E}[u(Q_t) + \delta u(Q_{t+1})] \quad \text{s.t.} \\ Q_t + \phi_t p_t \leq e_t + \phi_t d_t \\ Q_{t+1} \leq \phi_t p_{t+1} \end{aligned}$$

Recall that the decision over N_t is a trivial one, and is therefore omitted. Optimality of the young generation and asset market clearing ($\phi_t = 1$) implies

$$u'(e_t(\theta) - p_t(\theta) + d_t(\theta))[p_t(\theta) - d_t(\theta)] = \delta \mathbb{E}_t[p_{t+1}(\theta)u'(p_{t+1}(\theta))] \quad (\text{C.18})$$

$$p_t(\theta) = d_t(\theta) + \delta \mathbb{E}_t \left[d_{t+1} \frac{u'(p_{t+1})}{u'(e_t + d_t - p_t)} + \delta \frac{u'(p_{t+1})}{u'(e_t + d_t - p_t)} \mathbb{E}_{t+1} \left(d_{t+2} \frac{u'(p_{t+2})}{u'(e_{t+1} + d_{t+1} - p_{t+1})} + \dots \right) \right]$$

$$p_t(\theta) = d_t(\theta) + \frac{\delta}{u'(e_t + d_t - p_t)} \left(\mathbb{E}_t [d_{t+1} u'(p_{t+1})] + \delta \mathbb{E}_t \left[\frac{u'(p_{t+1})}{u'(e_{t+1} + d_{t+1} - p_{t+1})} \mathbb{E}_{t+1} [d_{t+2} u'(p_{t+2})] \right] + \dots \right)$$

$$\psi = \mathbb{E}_s [d_{s+1} u'(p_{s+1})] \quad \forall s \geq t$$

$$\chi = \mathbb{E}_s \left[\frac{u'(p_{s+1})}{u'(e_{s+1} + d_{s+1} - p_{s+1})} \right] \quad \forall s \geq t$$

$$p_t(\theta) = d_t(\theta) + \frac{\delta}{u'(e_t + d_t - p_t)} \left(\psi + \delta\chi\psi + \delta^2\chi^2\psi + \dots \right)$$

$$p_t(\theta) = d_t(\theta) + \frac{\delta\psi}{u'(e_t + d_t - p_t)} \left(\frac{1}{1 - \delta\chi} \right)$$

Notice that volatility in asset prices comes both from the volatility of the dividend and from the discount factor. If the asset is not cum dividend, there would still be volatility in asset prices. In this case the young generation consumes the endowment (labour income) minus the asset price. Period t young agents maximize expected utility from consumption in both periods

$$\begin{aligned} \max_{\phi_t} \mathbb{E}[u(Q_t) + \delta u(Q_{t+1})] \quad \text{s.t.} \\ Q_t + \phi_t p_t \leq e_t \\ Q_{t+1} \leq \phi_t (p_{t+1} + d_{t+1}) \end{aligned}$$

Which would result in

$$u'(e_t(\theta) - p_t(\theta)) [p_t(\theta)] = \delta \mathbb{E}_t [(p_{t+1} + d_{t+1}) u'(p_{t+1} + d_{t+1})]$$

$$p_t(\theta) = \delta \mathbb{E}_t \left[d_{t+1} \frac{u'(p_{t+1} + d_{t+1})}{u'(e_t - p_t)} + \delta \frac{u'(p_{t+1} + d_{t+1})}{u'(e_t - p_t)} \mathbb{E}_{t+1} \left(d_{t+2} \frac{u'(p_{t+2} + d_{t+2})}{u'(e_{t+1} - p_{t+1})} + \dots \right) \right]$$

$$p_t(\theta) = \frac{\delta}{u'(e_t - p_t)} \left(\mathbb{E}_t [d_{t+1} u'(p_{t+1} + d_{t+1})] + \delta \mathbb{E}_t \left[\frac{u'(p_{t+1} + d_{t+1})}{u'(e_{t+1} - p_{t+1})} \mathbb{E}_{t+1} [d_{t+2} u'(p_{t+2} + d_{t+2})] \right] + \dots \right)$$

$$\psi = \mathbb{E}_s [d_{s+1} u'(p_{s+1} + d_{s+1})] \quad \forall s \geq t$$

$$\chi = \mathbb{E}_s \left[\frac{u'(p_{s+1} + d_{s+1})}{u'(e_{s+1} - p_{s+1})} \right] \quad \forall s \geq t$$

$$p_t(\theta) = \frac{\delta}{u'(e_t(\theta) - p_t(\theta))} (\psi + \delta\chi\psi + \delta^2\chi^2\psi + \dots)$$

$$p_t(\theta) = \frac{\delta\psi}{u'(e_t(\theta) - p_t(\theta))} \left(\frac{1}{1 - \delta\chi} \right)$$

Thus, insofar as $e_t(\theta) - p_t(\theta)$ is not constant, asset prices are volatile also in the ex-dividend case. However, $e_t(\theta) - p_t(\theta)$ cannot be a constant. If it were, this would imply $p_t(\theta)$ to be a constant – but then this would contradict the premise that $e_t(\theta) - p_t(\theta)$ is constant since $e_t(\theta)$ is volatile.

The variance in the ex-dividend case is

$$Var[p] = \left(\frac{\delta\psi}{1 - \delta\chi} \right)^2 Var[u'(e - p)]^{-1}$$

while in the cum-dividend case it is

$$Var[p] = Var[d] + \left(\frac{\delta\psi}{1 - \delta\chi} \right)^2 Var[u'(e + d - p)]^{-1} + 2 \left(\frac{\delta\psi}{1 - \delta\chi} \right) Cov[d, [u'(e + d - p)]^{-1}]$$

In both cases, insofar as there is variation in aggregate income, endowment (labour income) and the dividend, there is volatility in asset prices.

C.5 Relaxing Wage Rigidity

C.5.1 Fraction of Workers have Sticky Wages

Let λ be the fraction of workers that are able to offer labour services at a completely flexible wage: the nominal wage that these workers accept always equates their marginal substitution between labour and consumption to the real wage. How does the economy behave under this weaker assumption? In the former case we ruled out producers' expectations of total goods on the market that were above $\bar{Q}$, and below $\underline{Q}$, where the former was defined endogenously (by equating the marginal substitution between labour and consumption with the real wage) and the latter exogenously. If some fraction of workers can offer labour services at a flexible wage, and the labour market is perfectly competitive, we are able to define $\underline{Q}$ endogenously as well.

In the case of perfect competition in the labour market, the nominal wage is set at $\hat{W}$ for $Q \in [\underline{Q}, \bar{Q}]$, where $\underline{Q}$ equates labour demand N^d to the labour supply of flexible workers

– so that at $\underline{Q}$ only flexible workers work. $Q < \underline{Q}$ is then inadmissible, because it implies the market for flexible labour doesn't clear. In other words, expectations of Q below $\underline{Q}$ are inadmissible. Expectations of Q above $\bar{Q}$ are inadmissible for the same reasons as we had before. Expectations of $Q \in [\underline{Q}, \bar{Q}]$ are admissible, and operate as a self fulfilling prophecy as before, without the need for an exogenous endowment of goods $\underline{Q}$. Nominal wages below $\hat{W}$ when production is expected to be lower than $\bar{Q}$ but above $\underline{Q}$ are not consistent with the assumption about the wage rigidity.

In sum, the assumption about a minimum nominal wage can be relaxed without affecting the qualitative results of this section. As long as there is a positive fraction of workers (and perfect competition in the labour market) that are constrained by a minimum nominal wage, there arises a continuum of equilibria in which production is lower than the Pareto efficient one at $\bar{Q}$. The larger the fraction of workers that cannot offer labour services at a nominal wage lower than some arbitrary minimum, the larger that set of admissible equilibria (from the downside, i.e. $\underline{Q} > 0$ is closer to zero).

C.5.2 Short-run Sticky Wage

One could argue that the assumption about sticky wages is strong, and that in the long run wages will equal the marginal substitution between labour and consumption. This subsection deals with the case of two periods, one in which the nominal wage is rigid, and another in which it is flexible. Firms are assumed to be able to sell in the second period goods produced in the first period with no storage costs. We will see that this does not alter the results of the former section.

A firm's choice consists of choosing three variables: how much to produce in each period ($\hat{Q}_{i,t}$ for $t = 1, 2$), and how much of their produce to sell in each period ($Q_{i,t}$ for $t = 1, 2$). Of course, given total production of the firm, the latter is reduced to a single choice variable ($Q_{i,1} + Q_{i,2} = \hat{Q}_{i,1} + \hat{Q}_{i,2}$). The firm is assumed to be constrained by its demand function in each period ($\frac{P_{i,t}}{P_t} \leq \left(\frac{Q_{i,t}}{I_t/P_t}\right)^{-\frac{1}{\sigma}} I_t/P_t$, in each period Q_t for $t = 1, 2$), and to be able to sell in the first period at most what it produced in the first period ($Q_{i,1} \leq \hat{Q}_{i,1}$). I assume the firm maximizes the sum of real profits in both periods: $\pi_1/P_1 + \pi_2/P_2$, where P_t is the usual price index at period t and $P_{i,t}$ is the inverse demand function of firm i given all firms' production at period t :

$$\sum_t \pi_t/P_t = \sum_t \left(Q_{i,t} P_{i,t}(Q_{i,t}, Q_t, P_t)/P_t - \frac{1}{A} \frac{W_t}{P_t} \hat{Q}_{i,t}^{\frac{1}{\alpha}} \right) \quad (\text{C.19})$$

I now establish a symmetric equilibrium exists in which all firms produce $(Q_t)_{t=1,2}$ and market all of their produce in the period they are produced. First, observe that, given a decision about how much to sell in each period, producing optimally implies²

$$\hat{Q}_{i,1} = \frac{\hat{Q}_{i,1} + \hat{Q}_{i,2}}{1 + \left(\frac{W_1/P_1}{W_2/P_2}\right)^{\frac{\alpha}{1-\alpha}}} \quad (\text{C.20})$$

That is, the optimal production plan depends on the ratio of real wages of both periods.³ Next, observe that given some decision over total production in both periods $\sum_t \hat{Q}_{i,t}$, optimal marketing decision of that produce is given by⁴

$$Q_{i,t} = \frac{I_t/P_t}{\sum_t I_t/P_t} \sum_t \hat{Q}_{i,t} \quad (\text{C.21})$$

Now, observe that given the optimal decision about how to produce and how much to sell, the optimal decision of a firm is to produce⁵

$$\sum_t \hat{Q}_{i,t} = \sum_t I_t/P_t \quad (\text{C.22})$$

²I ignore for the moment the question of whether this production plan satisfies the constraint $Q_{i,1} \leq \hat{Q}_{i,1}$. We will later see that the optimal decision does satisfy it.

³This results from the first order condition of eq. C.19 wrt $\hat{Q}_{i,1}$ keeping $\sum_t \hat{Q}_{i,t}$ constant:

$$\frac{1}{A} \frac{1}{\alpha} \frac{W_1}{P_1} \hat{Q}_{i,1}^{\frac{1}{\alpha}-1} = \frac{1}{A} \frac{1}{\alpha} \frac{W_2}{P_2} \hat{Q}_{i,2}^{\frac{1}{\alpha}-1}$$

⁴This is a result of the first order condition of eq. C.19 wrt $Q_{i,1}$.

⁵Let $Q_i = \sum_t \hat{Q}_{i,t}$ and $I/P = \sum_t I_t/P_t$. The full derivation is given by taking the derivative of

$$\sum_t \left\{ \frac{I_t/P_t}{I/P} Q_i \left(\frac{I_t/P_t Q_i}{I/P} \right)^{-\frac{1}{\sigma}} - \frac{1}{A} \frac{W_t}{P_t} \left[1 + \left(\frac{W_1/P_1}{W_2/P_2} \right)^{\frac{\alpha}{1-\alpha}} \right]^{-\frac{1}{\alpha}} \left(\frac{W_1/P_1}{W_t/P_t} \right)^{\frac{1}{1-\alpha}} Q_i^{\frac{1}{\alpha}} \right\}$$

wrt Q_i , which is given by

$$(1 - \sigma^{-1}) \sum_t \frac{I_t/P_t}{I/P} \left(\frac{Q_i}{I/P} \right)^{-\frac{1}{\sigma}} - \frac{1}{A} \frac{1}{\alpha} \left[1 + \left(\frac{W_1/P_1}{W_2/P_2} \right)^{\frac{\alpha}{1-\alpha}} \right]^{-\frac{1}{\alpha}} \sum_t \frac{W_t}{P_t} \left(\frac{W_1/P_1}{W_t/P_t} \right)^{\frac{1}{1-\alpha}} Q_i^{\frac{1}{\alpha}-1} = 0$$

and noticing that in a symmetric equilibrium the ratio of real wages is given by

$$\frac{W_1/P_1}{W_2/P_2} = \left(\frac{I_1/P_1}{I_2/P_2} \right)^{1-\frac{1}{\alpha}}$$

If returns to scale are decreasing, this first order condition is decreasing in Q_i , thus there is a unique intersection with zero, which is given by $Q_i = I/P$.

That is, the total amount that the firm chooses to sell is equal to the aggregate real income in both periods. Given these results it follows that $Q_{i,t} = I_t/P_t$, so that optimal behaviour of the firm requires it to operate in the same way as other firms do. Thus if all firms produce Q_2 in the second period, and $Q_1 < Q_2$ in the first period, firm i will do the same. Q_2 is pinned down by the condition of labour market clearing in the second period at the Pareto efficient level. However, even when allowing firms to sell their first period produce in the second period does not resolve the multiplicity of equilibria that exists in period 1.

C.6 Additional Points for Discussion

C.6.1 Nominal Wage Rigidity

Nominal wage rigidity is an important assumption in my model, and it therefore deserves additional discussion. My model does not provide an explanation of wage rigidity. Instead, it studies the macro consequences for an economy should it exist. Hall (2005) provides micro-foundations for wage rigidity in a costly search-and-match labour market. In his model, wage rigidity is an equilibrium, though there might be other possible equilibria. For explaining why this equilibrium is selected he refers to Bewley (2007), who argues that in practice nominal wage cuts are extremely rare because they have a negative effective on employee morale.

The significance of wage rigidity from a theoretical perspective, is that it enables the disentanglement of firms' marginal (real) revenue from the marginal substitution between labour and consumption. This, however, is not the only enabler. Angeletos and La'O (2013) achieves a similar result without any nominal rigidities, by studying segmented labour markets. The reason I use wage rigidity is that it is simple enough in order to not obscure the model's central feature in the case of common knowledge, namely multiple equilibria.

Moreover, as I note in section 3.2, wage rigidity is a well-established empirical fact. Keynes (2018) documents wage rigidity during the great depression. Hall et al. (1980) documents cyclical movements of output, employment and wages for the US economy in the 1970s. He shows that labour input (in terms of working hours) and the unemployment rate varied significantly over that decade, while the rate of growth in nominal wages is fairly stable. He says that: "There is no point any longer in pretending that the labor market is an auction market cleared by the observed average hourly wage... The traditional idea of sticky nominal wages and unilateral profit maximization by employers has hardly been overturned."

Bewley (2007) contains a survey of the literature on the existence of wage rigidity. He writes that there is ambiguous evidence when wage rigidity is gauged by surveying firms,

though even then he mostly presents evidence of downward wage rigidity. There is less ambiguous evidence of wage rigidity when it is gauged by observing company pay records (which are presumably more reliable because they are objective). Fehr and Falk (1999) provide evidence supporting Shapiro and Stiglitz's shirking theory. They conduct double-auction experiments to check the hypothesis that, faced with incomplete labour contracts, higher wages induce workers to exert higher effort. They conclude that workers often underbid each other, but employers refuse to accept such offers when faced with incomplete labour contracts – while they do accept them when faced with complete labour contracts. Moreover, they confirm that workers effort is positively correlated with wage levels. Altonji and Devereux (2000) present evidence that wage cuts are rare, and almost always associated with changes in full time status. Agell and Lundborg (2003) report results from a repeat survey of Swedish manufacturing firms in 1990s. They find no evidence that increase in unemployment rate mitigated the mechanisms generating wage rigidity. Fehr and Goette (2005) find robustness of wage-rigidity to a low growth environment in Switzerland in the 1990s.

In a discussion of Hall et al., LH Summers writes that the fact we observe *nominal* wage rigidity is an embarrassment to most contract theories, and that nominal wage rigidity is the major unsolved problem in understanding the labor market. His reading of Hall et al. suggests that there are in fact a multiplicity of labour markets, but only some have unemployment – namely the low-skilled labour markets. This would be consistent with a role for the minimum wage as an explanatory factor, though, as Stiglitz writes, persistent unemployment is a phenomenon predating the welfare institutions of the 20th century.

In the same discussion, William Nordhaus said that writing contracts that fix the real wage involve high costs, so the nominal wage can be used reliably in stable inflation environment. Azariadis and Stiglitz (1983) conceptualize wage rigidity as an information-processing failure, and bring the example of the custom of collective bargaining to predetermine money wages several years in advance.⁶

Standard NK models feature a frictionless labour market (Krause and Lubik, 2007), even though it is widely recognized that rigid wages are central in Keynes's explanation of the persistence of unemployment (Stiglitz, 1984). Far less widely recognized is the connection between wage rigidity and multiple equilibria. For example, Hall et al. (1980) says that

⁶Bewley (2007) writes: "John Maynard Keynes (1936) proposed that downward wage rigidity is explained by employees' preoccupation with pay differentials among workers in similar jobs at different firms. I found, however, that such external pay differentials are not an issue, except in highly unionized industries. In most companies, employees know so little about pay rates at other firms that they do not know whether they are underpaid. Although labor unions do try to keep their members informed of pay rates at other companies, unions are weak in the United States." Bewley's cite studies conducted in the 1990s, it is possible that the amount of information employees today have about salaries in other companies is greater.

“It is clear why employees in one firm work harder when there is more work to do in that firm, but not why there is more total work to do in the aggregate economy in some years than in others. Institutional arrangements in the labor market like temporary and permanent layoffs and unresponsive wages make good economic sense for individual firms dealing with their own fluctuations in demand, but it is not known why they sometimes operate in unison to depress employment throughout the economy.”

In the discussion of Hall et al., Robert Gordon said that “there is no puzzle in explaining fluctuations in aggregate employment if nominal wage stickiness is accepted and nominal demand is given, although this leads to a disequilibrium description of total output and employment.” In that same discussion, William Nordhaus noted that “such a disequilibrium view of macro fluctuations explained the actual facts of the typical cycle considerably better than the equilibrium business cycle view associated with Robert Lucas and Thomas Sargent.”

In the general disequilibrium model developed by Grossman and Barro (1976), both money wages and prices are fixed, and firms only choose quantities. In effect, the determination of effective demand is exogenous therefore equilibrium is unique. My model focuses on how expectations determine effective demand, and a unique equilibrium is a result of relaxing the common knowledge assumption. Demand is not simply given, as Robert Gordon suggests. Both demand and supply of goods (effective demand and supply) are determined endogenously (the goods market clears), and can vary with producers expectations. Nominal wage rigidity enables this phenomenon by disentangling firms’ actions from the economy-wide marginal disutility from labour.

C.6.2 Neutrality of Money

Is money neutral in this economy? Since the nominal wage is fixed, the total money in the economy is fixed for a given state of the world. To vary the amount of money in this model at a given state θ where production is already given, is to change one of it’s founding assumptions that the nominal wage is fixed. In this model, the amount of nominal money tokens is not fixed, it is determined by firm’s optimal behaviour provided the wage is fixed. Thus in the strict sense, yes, money is neutral – if we change the value of W but not the fact that it is fixed at W , there would be no change whatsoever. But for a given W , the amount of money in the model is perfectly correlated with total production, and is determined endogenously. To vary it ex-ante is to change the nominal wage.

One could think about an extension of the OLG model in section 3.5 to incorporate a central-bank that issues a safe bond denominated nominally. By doing so, a central bank

would be able to implement monetary policy that may change equilibrium. In that sense, money could be non-neutral in this type of economy.

C.6.3 Price Index Indeterminacy

What determines the price index in case there is no common knowledge? The price index is a function of the state $P(\theta)$, and is given only implicitly in a similar way to the fact that $Q(\theta_i)$ given only implicitly. The difference between the two is that, given $P(\theta)$, $Q(\theta_i)$ is a decision variable. This decision function is given only implicitly: the decision rule at signal θ_i depends on the rules at other signals. An admissible decision rule is a fixed path such that

$$Q(\theta_i) = \Omega_Q(Q(\theta), P(\theta))$$

at each and every signal θ_i . FMP provide the conditions of existence and uniqueness of this fixed decision rule. $P(\theta)$ on the other hand is not a decision variable, but is determined as part of this decision rule. In fact, we should say that the pair $Q(\theta), P(\theta)$ are to be the fixed function. The issue is that even when $Q(\theta)$ is given, $P(\theta)$ still depends on itself.

$$P(\theta) = \Omega_P(P(\theta), Q(\theta))$$

If there is a self-consistent average price function $P(\theta)$ for every decision rule $Q(\theta)$ then prices are indeterminant. Prices are determinant when there is no price dispersion because of the assumption about the fixed wage: For realizations of θ in which all agents follow the same the strategy, all prices are the same, and the average price is determined by the common knowledge benchmark in equation 3.3. We know that for the dominance regions these price are different.

Velocity of Money

The total money supply in the economy is determined by $M(\theta)$

$$M(\theta) = \frac{W}{A} \frac{1}{2\epsilon} \int_{\tilde{\theta}=\theta-\epsilon}^{\tilde{\theta}+\epsilon} Q(\tilde{\theta})^{1/\alpha} d\tilde{\theta} + FN$$

Since profits are positive, the same money token is used more than once (the velocity of money $V(\theta) > 1$). The velocity of money and the average price are co-determined:

$$V(\theta) = \frac{\bar{P}(\theta)}{M(\theta)} \frac{1}{2\epsilon} \int_{\tilde{\theta}=\theta-\epsilon}^{\tilde{\theta}+\epsilon} Q(\tilde{\theta}) \left(1 + \eta^{-1} \left[1 - \frac{Q(\tilde{\theta})}{\frac{1}{2\epsilon} \int_{\tilde{\theta}=\theta-\epsilon}^{\tilde{\theta}+\epsilon} Q(\tilde{\theta}) d\tilde{\theta}} \right] \right) d\tilde{\theta} - 1$$

Thus, different average prices are consistent with the same division of the real revenue to real profits and real cost, as long as we can vary the velocity of money accordingly.

C.6.4 Trading Mechanism

In what follows I provide an account of the trading mechanism of the economy. This is done for the purpose of internal completeness – to tie in assumptions required to deliver a unique equilibrium.

The economy's trade is facilitated in the following manner. An authority called the central bank is in charge of the issuance of fiat money, an IOU of the central bank. The value of money is derived solely from the fact that people trust they can pass it on to others for the procurement of goods and services, and is determined in equilibrium in accordance with agents' expectations of aggregate production and prices.

Firms are fully owned by consumer/workers, who hold equity claims on firms' profits, denominated in money. In order to produce, firms need to pay money to workers, who in turn use it to purchase goods from firms. Firms observe their private signal θ_i , and then approach the central bank asking for a loan at a zero interest. The size of the loan depends on the amount of labour they wish to employ and on the nominal wage in equilibrium. Thus, the amount of money in the economy grows when the firms wish to produce more intensively.

In turn, firms seek to hire labor at the price W_t . Recall that W_t has a minimum value of W . They pay workers and produce the amount of goods. After the production process is concluded, the goods' market is opened. Firms then sell their merchandise at the highest price they can sell all goods. After an initial round of sales all firms pay a dividend, and a second round of sales commences – using dividends to purchase any goods left. Finally, firms return the loan to central bank.

C.7 Related Literature, Macroeconomics

New-Keynesian (NK) models combine monopolistic competition in goods' markets with price rigidities ('stickiness') to deliver non-neutrality of money. The latter is then used as a rationale for the efficacy of monetary policy, the main tool by which central banks attain their goal of price stability. The 'Keynesian feature' of NK models is related to the monopolistic competition effect, which results in underemployment; this is termed by Blanchard and Kiyotaki (1987) as aggregate demand externality.⁷

⁷“The results of this paper are tantalizingly close to those of traditional Keynesian models: under monopolistic competition, output is too low, because of an aggregate demand externality.”

In those models there are two conditions that pin down a unique equilibrium: firm level optimization pins down production as a function of the real wage; labour markets clearing pins down the real wage by equating it to the marginal substitution between labour and consumption. Instead of this, I follow Keynes and explore an otherwise standard macroeconomic model, with a nominal wage rigidity: the nominal wage can go up, but not down. This does away with the condition of labour markets clearing, and in doing so gives rise to a coordination problem between producers. In contrast, NK models (such as Blanchard and Kiyotaki (1987)) typically assume a rigidity in prices for consumption goods rather than wages, which is less in the spirit of Keynes (Grossman, 1972, Footnote 11), and prevents self-fulfilling beliefs from playing a significant role.

Yet aggregate demand externality is insufficient, on its own, to produce an aggregate demand effect – that is a purely monetary effect, from the quantity of money to a real effect on output and employment. An additional friction in the form of price rigidities can deliver the sought for aggregate demand effect. Although Blanchard and Kiyotaki don't relate aggregate demand effect to its Keynesian features, supply-side vs demand-side is readily considered as the opposition between the two leading schools of economic thought in the decades that preceded their publication. Demand effects are very well known to be in 'Keynesian tradition'. Moreover, price rigidities are often referred to as a Keynesian idea (Rotemberg and Saloner, 1986).

There is, however, an important difference between Keynes's and New-Keynesian models. The former considers a short-term effect where the nominal wage is fixed (strictly speaking, cannot decrease), while the latter use a general form of nominal rigidities, but do so at the level of goods' prices rather than the wage. Blanchard and Kiyotaki [BK] confound their monetary effect with Keynes's effective demand, which is a central concept to his theory. In Keynes's model, labour markets don't clear, which opens the way for pure demand effects through a coordination of producers on a different equilibrium. The emphasis in BK is quite different. In the same vein as Keynes, Grossman and Barro (1976) (Chapter 2) distinguish between notional quantities, resulting from a Walrasian model in which all markets clear, to effective quantities, which result from a model in which at least some markets don't clear.

In BK, multiple equilibria and a coordination problem can exist, though their model does not produce this phenomenon robustly. The occurrence of multiple equilibria depends on the elasticity of aggregate demand to real money balances, if the latter exceeds unity, multiple equilibria is precluded. Moreover, in their model if the conditions for strategic complementarities are fulfilled, there is room for two equilibria. We will see that in my model there is a continuum of possible equilibrium. BK require menu-costs (price rigidities) to produce this multiple equilibria, whereas in my model there are no menu-costs. Finally,

their treatment of multiple equilibria is limited. Instead, they reference the reader to other models such as: Cooper and John (1988), Rotemberg and Saloner (1986), Ball and Romer (1991).

Models that admit equilibrium multiplicity abound, but although many have argued that it could play a role in macro environments, to the best of my knowledge they don't point out how nominal wage rigidity gives rise to a coordination problem. Ball and Romer (1991) model an economy of yeoman farmers subject to monetary shocks. In their model, multiple equilibria is a result of strategic complementarities if monetary shocks are of intermediate size. In case of intermediate shocks, there are three possible equilibria: one in which all farmers adjust prices, one in where no farmer adjusts prices, and finally one in which some adjust, some don't adjust, and all are indifferent to switching. However, the latter equilibrium is unstable. Moreover, their analytical result of multiple equilibria depends on a Taylor approximation, without which they rely on simulations.

Their model is similar to mine in that "strategic complementarity in price adjustment is tied to a simpler kind of strategic complementarity: the positive dependence of a farmer's utility-maximizing price in the absence of menu costs on the prices of others." Later in their paper they describe a continuum of equilibria in cut-off values of the monetary shock: if the value of the monetary shock is greater than some value, all farmers adjust their prices. That is, the multiplicity of equilibrium is not so much about the price itself (which is a reduced strategically to a binary variable), but rather about the cut-off value of the monetary shock. This is akin to what BK describe for their model.

Rotemberg and Saloner (1986) study oligopolistic behaviour, first in partial equilibrium, and then in a general equilibrium model with two sectors, one oligopolistic and another competitive.⁸ In their model, firms tend to collude less when demand shifts towards their sector; thus business cycles are a result of a relative shift in demand for the oligopolistic sector, and are bad not only because output is low but also because of greater microeconomic distortions. They leave the cause of shifts in sectoral demand unexplained, though contend that shifts in nominal money supply are correlated with both business cycles and demand for durable goods (which they describe as oligopolistic). My model is closer to what they describe as an alternative view of business cycles, in which aggregate demand changes are not reflected in the nominal wage (in my model because the nominal wage is too high). Finally, their treatment of multiple equilibria is minimal and not related to any strategic

⁸In their setup labour supply is horizontal at the price of the oligopolistic sector P_1 , meaning that at a given real wage labour supply is completely elastic. This is somewhat close to assuming fixed nominal wage like in model, in that as long as quantities are below the second best, labour wishes to supply at least as much labour as firms wish to employ. However, they say that this assumption doesn't play a significant role in their results.

complementarities. Rather, the advent of a large set of punishment-strategies that are needed to sustain collusion is a potential source of multiplicity.

Keynes's model is different from New Keynesian models in several respects. The most important aspect is perhaps the assumption that Keynes maintains almost throughout the General Theory, namely that workers might not have the expedient by which to reduce the nominal wage.⁹ Moreover, a closer look at Keynes's definitions of aggregate demand¹⁰, we can see how this can be confounded with the notion of the quantity of money in the economy, and thereby with pure monetary effects. However, for Keynes's theory, the important concept is effective demand.¹¹¹²

Keynes's full setup is a more elaborate macroeconomic model than the one I propose in this paper, involving capital, investment and saving. My model is highly stylized, aiming to capture a crucial phenomenon in his theory: that employers' expectations of other employers' actions results in a coordination problem and could result in economy-wide recessions.¹³ In a discussion of effective demand and Say's law, he clearly points out that the classical theory

⁹See e.g. chapter 2, section II:

"But the other, more fundamental, objection, which we shall develop in the ensuing chapters, flows from our disputing the assumption that the general level of real wages is directly determined by the character of the wage bargain. In assuming that the wage bargain determines the real wage the classical school have slipped in an illicit assumption. For there may be no method available to labour as a whole whereby it can bring the wage-goods equivalent of the general level of money wages into conformity with the marginal disutility of the current volume of employment. There may exist no expedient by which labour as a whole can reduce its real wage to a given figure by making revised money bargains with the entrepreneurs."

Chapter 2, section III:

"the struggle about money-wages primarily affects the distribution of the aggregate real wage between different labour-groups, and not its average amount per unit of employment, which depends, as we shall see, on a different set of forces."

¹⁰The proceeds of a level of employment – aggregate [nominal] income (i.e. factor cost plus profit) resulting from a given amount of employment.

¹¹"the substance of the General Theory of Employment...".

¹²The value of aggregate demand at the intersection with aggregate supply – expectation of proceeds which will just make it worth the while of the entrepreneurs to give that employment. Aggregate supply being the function that relates the supply price – the expectation of proceeds which will just make it worth the while of the entrepreneurs to give that employment – to a given amount of employment. Thus when the two intersect, there is an equilibrium in which all producers behave optimally.

¹³See Chapter 5, section I:

"All production is for the purpose of ultimately satisfying a consumer. Time usually elapses, however, and sometimes much time, between the incurring of costs by the producer (with the consumer in view) and the purchase of the output by the ultimate consumer. Meanwhile the entrepreneur (including both the producer and the investor in this description) has to form the best expectations he can as to what the consumers will be prepared to pay when he is ready to supply them (directly or indirectly) after the elapse of what may be a lengthy period; and he has no choice but to be guided by these expectations, if he is to produce at all by processes which occupy time."

assumes that

“...effective demand, instead of having a unique equilibrium value, is an infinite range of values all equally admissible; and the amount of employment is indeterminate except in so far as the marginal disutility of labour sets an upper limit.”

But in Keynes’s model the necessity of labour markets’ clearing is one of the major assumptions of classical theory which he breaks with – resulting in non-frictional involuntary unemployment, a feature that is completely absent from the literature cited above. In fact, he devotes a considerable part of his magnum opus to this question – which he refers to as the second postulate of classical economics, the first one being that the real wage equals the marginal product of labour.

Grossman and Barro (1976) develop a general disequilibrium model, though they neglect the determination of effective demand (they take it as exogenous), whereas the current model focuses how expectations determine effective demand. They consider a thought experiment (pp. 55-56) in which they demonstrate the effect of a Keynesian multiplier. Their model is very different from mine. They actually get a unique equilibrium – they have no self-fulfilling prophecy element to their model. I think the reason for that is that both wages and prices are fixed, and firms only choose quantities.

Other works have considered problems of insufficient demand that are not a result of a wage rigidity. In Clower (1965), households have a choice between consuming and saving, thus his approach is explicitly choice theoretic. Angeletos and Lian (2019) is modern micro-foundation of the same notion. Kiyotaki (1988) locates discoordination of producers in a model with capital, though he requires increasing returns to scale. Unlike these models, my model misses the inter-temporal aspect of cutting down production/consumption now to save more for the future.

In sum, macroeconomic models have hitherto followed Keynes only in using monopolistic competition to facilitate an equilibrium in which the marginal product of labour does not equal the real wage (the first postulate), though they refrain from following Keynes in abandoning the second postulate, that the labour market clears. If we take a standard monopolistic competition model, and fix the nominal wage, as indeed Keynes does in most of his book, we get a multiplicity of equilibria. This by itself might be a good reason for economists to adopt the second postulate, though it need not be the only way to achieve a unique equilibrium.

