



LUND UNIVERSITY

Networks, Information and Economic Volatility

Cokayne, Graeme

2016

Document Version:

Publisher's PDF, also known as Version of record

[Link to publication](#)

Citation for published version (APA):

Cokayne, G. (2016). *Networks, Information and Economic Volatility*. [Doctoral Thesis (monograph), Lund University School of Economics and Management, LUSEM]. MediaTryck Lund.

Total number of authors:

1

General rights

Unless other specific re-use rights are stated the following general rights apply:

Copyright and moral rights for the publications made accessible in the public portal are retained by the authors and/or other copyright owners and it is a condition of accessing publications that users recognise and abide by the legal requirements associated with these rights.

- Users may download and print one copy of any publication from the public portal for the purpose of private study or research.
- You may not further distribute the material or use it for any profit-making activity or commercial gain
- You may freely distribute the URL identifying the publication in the public portal

Read more about Creative commons licenses: <https://creativecommons.org/licenses/>

Take down policy

If you believe that this document breaches copyright please contact us providing details, and we will remove access to the work immediately and investigate your claim.

LUND UNIVERSITY

PO Box 117
221 00 Lund
+46 46-222 00 00

Networks, Information and Economic Volatility

Graeme Cokayne



LUND
UNIVERSITY

DOCTORAL DISSERTATION

by due permission of the School of Economics and Management,
Lund University, Sweden.

To be defended at Holger Crafoord EC3:211, Lund
on October 6, 2016 at 14:15.

Faculty opponent

Paolo Pin, University of Bocconi

Distributed by Department of Economics, P.O. Box 7082, S-220 07 Lund, Sweden

I, the undersigned, being the copyright owner of the abstract of the above-mentioned dissertation, hereby grant to all reference sources permission to publish and disseminate the abstract of the above-mentioned dissertation.

Signature _____ Date 2016-08-30 _____

Organization Lund University Department of Economics P.O. Box 7082 S-220 07 Lund, Sweden	Document name Doctoral dissertation	
	Date of issue	
Author Graeme Cokayne	Sponsoring organization	
Networks, Information and Economic Volatility Networks, Information and Economic Volatility:		
Abstract This thesis makes a contribution to network theory and how it applies to economics. It consists of three self-contained papers. It predominantly considers how connections between economic entities can affect economic outcomes. In particular, the first two papers examine social learning, in the one case by applying it to portfolio choice, and in the other by conducting an experiment to determine how people incorporate information from others into their beliefs to achieve economic outcomes. The final paper looks at how sectoral shocks can be transmitted through the economy through the connections between industries. The first paper, <i>Social learning and financial markets: Can informed neighbors make up for a lack of financial acumen?</i>, allows financially- informed and uninformed agents to consult with each other on different social networks and examines what the benefits are to the two different groups of agents on the different networks in terms of financial investment. The second paper, <i>Incorporating information into beliefs on networks: An experiment</i>, studies how agents incorporate into their beliefs information they receive from other agents. In particular, it considers whether agents use the DeGroot model or a Bayesian approach. Further, it considers how the density of network connections affects the results of consultation. The third paper, <i>Sectoral shocks and aggregate volatility</i>, devises a demand-side measure of industry influence and shows that such a demand-side measure is needed to fully understand the influence of industries with a case study of the Australian mining industry.		
Keywords Social Networks, Learning, Portfolio Choice, Information, DeGroot, Bayes, Macroeconomics		
Classification system and/or index terms (if any) JEL Classification: D85, D83, G11 C67, E32, D03		
Supplementary bibliographical information	Language English	
ISSN and key title 0460-0029 Lund Economic Studies no. 196	ISBN 978-91-7623-965-0 (print) 978-91-7623-966-7 (pdf)	
Recipient's notes	Number of pages	Price
	Security classification	

Networks, Information and Economic Volatility

Graeme Cokayne



LUND
UNIVERSITY

LUND ECONOMIC STUDIES NUMBER 196

Copyright © Graeme Cokayne 2016

Distributed by the Department of Economics

Lund University

P.O. Box 7082

S-220 07 Lund

SWEDEN

ISBN 978-91-7623-965-0 (print)

ISBN 978-91-7623-966-7 (pdf)

ISSN 0460-0029

Printed in Sweden by Media-Tryck, Lund University

Lund 2016

Contents

Abstract	i
Acknowledgments	iii
Introduction	1
1 Background	1
2 Social Learning and Financial Markets: Can informed neighbors make up for a lack of financial acumen?	2
3 Incorporating Information into Beliefs on Networks: An experiment	3
4 Sectoral Shocks and Aggregate Volatility	4
Paper I: Social Learning and Financial Markets: Can informed neighbors make up for a lack of financial acumen?	11
1 Introduction	12
2 General Model	14
3 Specific Model	15
3.1 Network Structures and Scenarios	16
3.2 Agents' Beliefs	19
3.3 Signals	19
3.4 Updating beliefs through Consultation	21
3.5 Utility maximization	23
4 Analysis	25
4.1 Expected Beliefs and Variance in Beliefs after Consultation	26
4.1.1 Scenario 1: Perfect Knowledge of Neighbors	27
4.1.2 Scenario 2: No Knowledge of Neighbors	31
4.2 Consulting on the Three Specific Networks	32
4.2.1 Hub Networks	32
4.2.2 Random Wheel Networks	33
4.2.3 Non-random Wheel Networks	36
4.2.4 Summary	40
4.3 Certainty Equivalents, Expected Returns and Variance in Returns	40

4.3.1 Summary	46
5 Conclusion	47
A Appendix	51
A.1 Proof of Proposition 1: Beliefs that are closer to correctly certain (CC) beliefs lead to higher expected returns	51
A.2 Proof of Proposition 2: Analysts' post-signal beliefs will be closer to CC beliefs than their prior beliefs regardless of the state of nature or prior beliefs.	52
A.3 Derivations of expected beliefs and variance in beliefs under no neighbor knowledge	52
A.4 Expected variance of an analyst's post-consultation beliefs with perfect neighbor knowledge on a random-wheel network.	55
A.5 An analyst's expected post-consultation beliefs with no neighbor knowledge on a random-wheel network.	55
A.6 The expected variance in an analyst's post-consultation beliefs with no neighbor knowledge on a random-wheel network.	56
A.7 The expected post-consultation variance of uninformed agents on random wheel networks with no neighbor knowledge.	57
A.8 Derivations of the Neighborhood Mix of Agents in Non-random Networks (Propositions 3 - 4)	57
A.8.1 Proof of Proposition 3	58
A.8.2 Proof of Proposition 4	60
A.9 The expected variance of analysts' post-consultation beliefs on a non-random wheel network under perfect neighbor knowledge	61
A.10 The expected variance of uninformed agents' post-consultation beliefs on a non-random wheel network under perfect neighbor knowledge	62
A.11 The expected post-consultation beliefs of analysts on a non-random wheel network under no neighbor knowledge	63
A.12 The expected post-consultation variance of analysts on a non-random wheel network under no neighbor knowledge	63
 Paper II: Incorporating Information into Beliefs on Networks: An Experiment	 67
1 Introduction	68

2 DeGroot's Model and Network Theory	71
3 Experiment Design	73
4 Hypotheses	77
4.1 Marginal Benefit of Consultation	79
4.2 Costs of consulting	82
5 Experiment Details and Subjects	84
6 Main Results	86
6.1 Bayes vs DeGroot: Testing Hypothesis 1	86
6.2 Consultations: Testing Hypothesis 2	89
6.3 Convergence to Consensus	91
6.4 Accuracy of beliefs	92
7 Subsidiary Results	94
7.1 Non-consulters	94
7.2 Initial Estimates	95
8 Conclusions	97
A Appendix	102
 Paper III: Sectoral Shocks and Aggregate Volatility	 121
1 Introduction	122
2 Literature Review	124
3 Economic Model	126
4 Estimation	128
4.1 Level of Disaggregation	128
4.2 Data	129
4.3 Cross-country comparisons	130
4.3.1 Broad disaggregation	130
4.3.2 Fine disaggregation	135
5 Case study: Australian mining industry	137
5.1 Recent history of the Australian mining industry	137
5.2 The mining industry in the Australian supply and demand influ- ence vectors	138
6 Conclusion	142

A	Supply shocks derivation	146
B	Demand shocks derivation	147
C	Influence Measures and Network Theory	149
C.1	Supply Influence	150
C.2	Demand Influence	151

Abstract

This thesis makes a contribution to network theory and how it applies to economics. It consists of three self-contained papers. It predominantly considers how connections between economic entities can affect economic outcomes. In particular, the first two papers examine social learning, in the one case by applying it to portfolio choice, and in the other by conducting an experiment to determine how people incorporate information from others into their beliefs to achieve economic outcomes. The final paper looks at how sectoral shocks can be transmitted through the economy through the connections between industries.

The first paper, *Social learning and financial markets: Can informed neighbors make up for a lack of financial acumen?*, allows financially- informed and uninformed agents to consult with each other on different social networks and examines what the benefits are to the two different groups of agents on the different networks in terms of financial investment.

The second paper, *Incorporating information into beliefs on networks: An experiment*, studies how agents incorporate into their beliefs information they receive from other agents. In particular, it considers whether agents use the DeGroot model or a Bayesian approach. Further, it considers how the density of network connections affects the results of consultation.

The third paper, *Sectoral shocks and aggregate volatility*, explores how sectoral shocks can propagate through the economy through the input-output networks. It devises a demand-side measure of industry influence and shows that such a demand-side measure is needed to fully understand the influence of industries with a case study of the Australian mining industry.

Keywords: Social Networks, Learning, Portfolio Choice, Information, DeGroot, Bayes, Macroeconomics

JEL Classification: D85, D83, G11 C67, E32, D03

Acknowledgments

I shall be telling this with a sigh
Somewhere ages and ages hence:
Two roads diverged in a wood, and I —
I took the one less traveled by,
And that has made all the difference.

The road not taken – Robert Frost

Through the course of my PhD studies I have often wondered whether Robert Frost was correct when he implicitly praised taking the road less traveled by. I was recruited to the PhD program as a macroeconomics student. However, just prior to starting, I read Duncan Watts' book "6 Degrees" which outlined how the study of networks could revolutionize economics. I was intrigued by this and immediately was inspired by the various ways in which the study of networks could give insights into macroeconomic questions.

With perhaps more confidence than was justified, I turned up to the PhD introductory meetings and declared that I wanted to study the less traveled by road of network theory rather than the well-trodden path of macroeconomics. Never mind that I had not taken a course in network theory, and never mind that none of the faculty at the department of economics specialized in network theory, or that none of the other students were studying network theory, I was going to do this. Many times through the course of my PhD studies I have wondered at the wisdom of that decision. Ultimately though, pursuing this path has given me a great deal of fulfillment and now that it is nearly over I would not change it.

There are many people I have to thank for helping me along the way. Certainly far too many than I can thank here, so what follows is in no way an exhaustive list.

First of all, to my two supervisors, Frederik Lundtofte and Erik Wengström, thank you for all of the help that you have given me. Even though neither of you specialize in network theory, you have always been a great help whenever I have needed it. Frederik, you regularly forced me to be more rigorous in my theoretical analysis (which I very much needed!!), and Erik, you helped me greatly with designing the experiment and adding to the microeconomic side of my research.

I would like to especially thank the members of the Knut Wicksell Centre

for all of their helpful suggestions and questions. Also for expressing interest in what I was doing, you helped enormously in giving me inspiration to continue my research. Furthermore, thank you for the collegial atmosphere that you created that let me know that there were others who were going through the same experience as me and who cared about my results. In no particular order I would like to thank Hossein Asgharian, Martin Strieborny, Lu Liu, Björn Hansson, Valeriia Dzhamaeva, Hans Byström, Anders Vilhelmsson, Jens Forssbäck, Naciye Sekerci, Anamaria Cociorva, Veronika Lunina and Hassan Sabzevari.

There are various other faculty members at Lund University who have helped me along the way. Whether it was through teaching interesting courses, or with answers to my questions, or asking me questions about my research, giving me suggestions, or just generally showing an interest in what I was doing, you encouraged me greatly. Moreover, you have made the Department of Economics at Lund University a pleasant yet challenging place to work. So for all of that I would like to thank David Edgerton, Fredrik NG Andersson, Pontus Hansson, Tommy Andersson, Fredrik Sjöholm, Joakim Gullstrand, Jerker Holm, Natalia Montinari, Karin Olofsdotter, Maria Persson, Dag Rydorff, Margareta Dackehag, Alessandro Martinello, Kaveh Majlesi and Joakim Westerlund.

To the administration staff at the Department of Economics, thank you for always being ready to answer my questions and help me sort out the myriad issues that cropped up but that never occur to a researcher until they need to be solved. So thank you to Jenny Wate, Nathalie Stenbeck, Mariana Anton, Rikke Barthélemy, Karin Wandér and Azra Padjan.

To my PhD cohort, Katarzyna Burzynska, Mingfa Ding, Anna Andersson, Jens Gudmundsson and Cecilia Hammarlund: we started this journey together and, though you all finished well before me, I greatly enjoyed the times we had along the way. I would have liked to have been able to socialize more with you if only the restrictions of a young family and living in a different town had not got in the way. Nonetheless, I fondly look back on our many Friday lunches and conversations about life, research, university and everything else.

Life at the department would, of course, not have been the same without the many conversations both large and small that I had with the other PhD students. To all of you, you have enriched my time studying for my PhD. In that vein I would like to thank Daniel Ekeblom, Lina Marie Ellegård, Hilda Ralsmark, Albin Erlandson, Gustav Kjellsson, Elvira Andersson, Caren Nilsson, Mahtemeab Tadesse, Jens Dietrichson, Sofie Gustavsson, Thomas Eriksson, Karin Bergman,

Johan Blomquist, Emanuel Alfranseder, Milda Norkute, Claes Ek, Usman Khalid, Karl McShane, Margaret Samahita, Anna Welander, Lina Ahlin, Björn Thor Arnarson, Ross Wilson, Yana Pryymachenko, Jim Ingebretsen Carlson, Sara Moricz, Hjördis Hardardottir, Mehmet Caglar Kaya, Dominika Krygier, Kristoffer Persson, Dominice Goodwin, Yana Petrova, and Pol Campos.

Monday evenings through most of my PhD career were taken up with playing the very Swedish sport of innebandy. It was a great release that I looked forward to and was rather sad to have to give it up over the last couple of years of my studies. To Henrik Lundgren, Mats Alveesson, Birger Nilsson and all of the other regular and irregular innebandy crew, thank you for the games.

Finally, I must thank my family for their support and providing a side of my life away from the study of economics that has given my so much joy and fulfillment. To my partner, Jessica, who was with me at the start of the PhD journey, and to my children, Hannsi and Carita, who joined us along the way, I love you very much and thank you for everything you have given to my life.

INTRODUCTION

Introduction

1 Background

This dissertation investigates various ways in which networks can influence economic outcomes, specifically through both transferring information between individual economic agents and more broadly linking industries within the economy. Focusing on the connections between economic agents and how their interactions can affect economic outcomes is a comparatively new but rapidly growing area of economic research.

Rational choice theory assumes that consumers are utility-maximizing rational agents making individual decisions based on their individual preferences. Furthermore, rational choice theory assumes that agents are fully informed about their preferences, their alternative possible courses of action and other relevant information about the world. However, humans are not independent of their fellow human beings. We often influence each other's decisions. Nor are we fully informed about all aspects of our lives. We often must take cues from others' behavior and learn from each other (Jackson 2008).

In order to become more informed about the world, people often turn to each other, asking friends and other social contacts for information on which to base their consumption and investment decisions. In this way people can become more informed about the world, and so closer to one of the assumptions of rational choice theory, but at the same time, losing their independence. Learning from those to whom an agent is socially connected is called social learning (see, for example, Goyal 2007, Bala & Goyal 1998, 2001). Chapters 2 and 3 contribute to this area of research by investigating how people gain information and how they might incorporate this information into their beliefs and hence actions.

Another example in which economics has ignored the connections within an economy is how macroeconomics has considered sectoral shocks. Traditionally, macroeconomics has ignored the effects of sectoral shocks as it has been assumed that industries are independent and so in a diversified economy sectoral shocks will tend to cancel each other out (Gabaix 2011). However, industries are not independent of each other but rather fit into a network of interrelationships. Chapter 4 explores how the interdependencies between industries might alter our conclusions about economic volatility.

2 Social Learning and Financial Markets: Can informed neighbors make up for a lack of financial acumen?

Chapter 2 investigates how social learning can affect financial decisions with particular reference to portfolio choice. People often need to make financial decisions. However, many people are not particularly well-informed when it comes to financial matters. As a result people will often turn to their social networks to gain information to make financially beneficial decisions. Chapter 2 investigates whether this is a reasonable strategy for people to use, where the benefits lie for people who are not financially well-informed and indeed for those who are, and on which networks they might best gain knowledge about financial matters.

Applying social learning to portfolio choice theory is fairly thin area of research, which appears somewhat strange given how many people turn to their social contacts for advice when it comes to financial matters. Various papers consider how social neighbors can affect agents' choice to participate in the stock market but not how they might choose their portfolio (see, for example, Hong et al. 2004, Ivkovich & Weisbenner 2007, Hong et al. 2005, Cohen et al. 2008). The few papers that do exist in this area tend to assume that all agents have equal access to information, even if their information does not always turn out to be so useful. These models might be considered to be how consultations between financial experts might affect financial outcomes (Acemoglu et al. 2011, Ozsoylev & Walden 2011).

In contrast, this is the first paper that has treated financial agents as categorically different in term of their access to information and how this might affect their portfolio choices in response to social learning. That is, how consultation might affect experts and non-experts differently.

The model presented considers three specific networks, designed to represent different social networks on which people might seek financial advice. The networks represented are finance-related blogs, Facebook and LinkedIn. Agents can choose how many people they wish to consult with and on which network they can consult to improve their beliefs. In this model, consulting is defined as learning their neighbors' beliefs about the likely state of nature. Following consultation, the agents incorporate their neighbors' beliefs into their own and then choose their optimal portfolio based on their beliefs about the likely state of nature.

The paper analyses how consultation affects the expected beliefs and variance in beliefs for the two types of agents. It then analyses how their beliefs flow through to expected returns and variance in returns and, ultimately, to the certainty equivalents of their investments.

The analysis shows that for informed agents the benefits of consultation do not lie in increases in their expected returns but rather in a reduction in the variance in their returns. These benefits are realized on finance-related blogs and LinkedIn, regardless of how well they know how informed their neighbors are. On Facebook, these gains are only realized if they can tell how well-informed their neighbors are. For uninformed agents the benefits lie in an improvement in their expected returns with only a relatively small gain through a reduced variance. Again the greatest benefit is on finance-related blogs. However, if they can tell how informed their neighbors are, networks similar to Facebook can also be beneficial.

3 Incorporating Information into Beliefs on Networks: An experiment

Within the model of chapter 2 agents used the DeGroot (1974) model to incorporate the information from their neighbors into their beliefs about the state of nature. This is a simple model in which people weight the beliefs of their social contacts and then using these weights, they update their beliefs and a weighted average of their neighbors' beliefs. Given the simplicity of this model of information incorporation, it is very attractive to use and assumes a low cognitive load on the part of the agents, and, thus, has been regularly used in the economic literature (see, for example, Jadbabaie et al. 2012, Golub & Jackson 2010, Eyster & Rabin 2009). An alternative model of information incorporation that has also regularly been used in the economic literature is based on Bayes rule (Acemoglu et al. 2011, Holt & Smith 2009, Charness et al. 2007). The benefit of this model is that the updated beliefs of agents will be use information in much more sophisticated way. The downside is that it assumes a high degree of sophistication on the part of the agents. So which do people actually use when they incorporate information from other people into their beliefs?

Chapter 3 outlines an experiment I conducted to try to tease out which model is most reasonable model to use to describe how people incorporate information into their beliefs. In general, it finds that the DeGroot model is a reasonable

one to describe how people come to their beliefs. However, there was a significant group who would follow what might be described as a Bayes rule model at least some of the time. This group was larger than that has been found through previous research.

Another prediction of the DeGroot model, and indeed also those models that rely more on Bayes rule, is that people on networks will come to consensus if they consult with each other enough times. Given that many networks of people do not end up at consensus a reasonable question is why not? What stops them from reaching consensus? Chapter 3 also investigates whether it is the perceived costs and benefits of consulting that stops people from consulting to consensus, and how the network structure might affect this.

In general, I found that the perceived benefits of consulting were the most important factors in whether people decided to consult with their neighbors an extra round, with little evidence that the costs of consulting played much role. People on more sparsely-connected networks tended to consult more times than those on more densely-connected networks. However, this extra consultation did not lead to better results in terms of uncovering information.

4 Sectoral Shocks and Aggregate Volatility

While chapters 2 and 3 focused on individual actors, networks can affect economic outcomes at a broader perspective. Chapter 4 considers how interrelationships between industries can affect macroeconomic volatility.

Traditionally macroeconomics has ignored the effects of sectoral shocks. This has largely been because it was felt that in a well-diversified economy positive shocks to one sector will be counterbalanced by negative shocks to another and so they will have no aggregate effect (Gabaix 2011). This theory relies on the assumption that industries are independent of each other. Of course, industries are not independent of each other because they are connected through input-output relationships. Therefore, if an industry has a productivity shock this will be passed on to the industries that it supplies inputs to, and then through to the industries that those industries supply inputs to and so on. Looking the other way, if an industry is hit by a negative demand shock that industry will not demand so many inputs from its supplying industries, and so the shock will be passed on to those industries that supply inputs to it. The supplying industries will then not demand so many inputs from their supplying industries and so

on. In this manner demand and supply shocks will be passed on beyond the original industries that receive those shocks.

Some research has looked at this process. In particular, Acemoglu et al. (2012, 2013) looked at supply shocks and the influence that industries have over each other through that channel. I construct a similar instrument to measure the influence of industries through the demand channel and show how a demand-side measure is needed with a case study of the Australian mining industry.

The Australian mining industry is a good case study here because mining has been considered one of the keys to Australia's strong economic performance over the past couple of decades. However its influence over the Australian economy is almost completely through the demand side. While it produces a large amount of exports and demands inputs from a number of other industries, it supplies very few inputs to other industries within Australia. Indeed, using my measures of industry influence, the mining industry is one of the least influential industries on the supply side but one of the most influential on the demand-side.

References

- Acemoglu, D., Carvalho, V. M., Ozdaglar, A., & Tahbaz-Salehi, A. (2012). The network origins of aggregate fluctuations. *Econometrica*, 80(5), 1977–2016.
- Acemoglu, D., Dahleh, M. A., Lobel, I., & Ozdaglar, A. (2011). Bayesian learning in social networks. *The Review of Economic Studies*, 78(4), 1201–1236.
- Acemoglu, D., Ozdaglar, A., & Tahbaz-Salehi, A. (2013). The network origins of large economic downturns. Working Paper 19230, National Bureau of Economic Research.
- Bala, V. & Goyal, S. (1998). Learning from neighbours. *Review of Economic Studies*, 65(3), 595–621.
- Bala, V. & Goyal, S. (2001). Conformism and diversity under social learning. *Economic Theory*, 17, 101–120.
- Charness, G., Karni, E., & Levin, D. (2007). Individual and group decision making under risk: An experimental study of bayesian updating and violations

- of first-order stochastic dominance. *Journal of Risk and Uncertainty*, 35(2), 129–148.
- Cohen, L., Frazzini, A., & Malloy, C. (2008). The small world of investing: Board connections and mutual fund returns. *Journal of Political Economy*, 116(5), 951–979.
- DeGroot, M. H. (1974). Reaching a consensus. *Journal of the American Statistical Association*, 69(345), 118–121.
- Eyster, E. & Rabin, M. (2009). Rational and naive herding. CEPR Discussion Papers 7351, C.E.P.R. Discussion Papers.
- Gabaix, X. (2011). Granular origins of aggregate fluctuations. *Econometrica*, 79(3), 733–772.
- Golub, B. & Jackson, M. O. (2010). Naïve learning in social networks and the wisdom of crowds. *American Economic Journal: Microeconomics*, 2(1), 112–49.
- Goyal, S. (2007). *Connections: An Introduction to the Economics of Networks*. Princeton: Princeton University Press.
- Holt, C. A. & Smith, A. M. (2009). An update on bayesian updating. *Journal of Economic Behavior & Organization*, 69(2), 125 – 134.
- Hong, H., Kubik, J. D., & Stein, J. C. (2004). Social interaction and stock-market participation. *The Journal of Finance*, 59(1), 137–163.
- Hong, H., Kubik, J. D., & Stein, J. C. (2005). Thy neighbor's portfolio: Word-of-mouth effects in the holdings and trades of money managers. *The Journal of Finance*, 60(6), 2801–2824.
- Ivkovich, Z. & Weisbenner, S. (2007). Information diffusion effects in individual investors' common stock purchases covet thy neighbors' investment choices. NBER Working Papers 13201, National Bureau of Economic Research, Inc.
- Jackson, M. O. (2008). *Social and Economic Networks*. Princeton, New Jersey: Princeton University Press.
- Jadbabaie, A., Molavi, P., Sandroni, A., & Tahbaz-Salehi, A. (2012). Non-bayesian social learning. *Games and Economic Behavior*, 76(1), 210 – 225.

Ozsoylev, H. N. & Walden, J. (2011). Asset pricing in large information networks.
Journal of Economic Theory, 146(6), 2252 – 2280.

PAPER I

Social Learning and Financial Markets: Can informed neighbors make up for a lack of financial acumen?

Abstract

Much of the literature on social learning examines equally-informed agents and determines that the full benefits of social learning only occur asymptotically. This paper investigates what the benefits of social learning are for agents with different levels of financial knowledge and how quickly those benefits accrue on different networks designed to reflect real-world social networks. It finds that the main benefit to financially knowledgeable agents is a reduction in the variance of their returns, while uninformed agents benefit mainly from an improvement in their expected returns. When consulting on hub networks most of the gains can be achieved through consulting only a few other agents, while to gain similar benefits from consulting on wheel networks far larger neighborhoods are generally required.

Keywords: Social Networks, Learning, Social Influence, Bounded Rationality, Portfolio Choice.

JEL Classification: D85, D83, G11

I would like to thank Frederik Lundtofte and Erik Wengström for their many helpful comments and suggestions. Also, thank you to all of the people at the Knut Wicksell Centre for their encouragement and helpful feedback.

1 Introduction

In considering the question of choice among investment options, many investors resort to asking people in their social network for advice. For example, is it currently a good time to invest in real estate? Should agents ask their pension funds to focus their investment on stocks or bonds? This naturally prompts questions about whether this is a reasonable strategy for uncovering information about the current merits of investing in different asset classes. If investment analysis is beyond the abilities of some agents who still must make investment decisions, can they make up for this lack of financial knowledge by asking more-informed people to whom they are socially connected? Moreover, even if an agent is able to conduct investment analysis themselves, might they still gain from consulting other informed investors in their social neighborhood?

Studies have shown that participation in financial markets is significantly influenced by agents' social contacts. This is true not only among finance professionals (Hong et al. 2005, Cohen et al. 2008) but also across the general public (Hong et al. 2004, Ivkovich & Weisbenner 2007). Therefore, it is worth considering what the benefits of social learning are for finance professionals and the general public, and how these benefits might differ between the two groups. This paper contributes to the literature by applying a social learning model to a portfolio choice setting and studying what the gains from consultation are in terms of certainty equivalents for different agents in a network in which there are both informed and uninformed agents, and in which the signals received by the informed agents are neither fully informative of the state of nature nor of the returns to investment. Furthermore, it considers how quickly gains can be realized as the neighborhood size increases, on three specific networks designed to reflect real world networks through which people can gain financial advice. These networks are: a network with a central core of informed investors with uninformed investors observing the core from outside, to represent finance-related blogs; a decentralized network with informed investors randomly distributed throughout the network, to represent Facebook; and finally, a decentralized network with informed investors grouped together, to represent LinkedIn.

In this study, agents choose a portfolio consisting of a risk-free asset and a risky asset. The expected returns on the risky asset changes with an unknown state of nature. A subset of agents receives private signals regarding the unknown state of nature. All agents can then consult with a subset of other agents (their "neighbors") to come to final beliefs about the state of nature. In the

context of this paper “consulting” refers to agents observing the beliefs of their neighbors about the state of nature. Having observed their neighbors’ beliefs agents then incorporate those beliefs into their own beliefs by constructing a weighted average of their neighbors’ beliefs (including their own beliefs). This method is based on the model produced by DeGroot (1974). After agents have updated their beliefs following consultation, they choose a portfolio of assets based on their beliefs about the state of nature.

Those agents that receive a signal about the state of nature (and hence are more informed) gain in certainty equivalents on hub networks (reflecting finance blogs), and non-random wheel networks (representing LinkedIn) with most of the gains realized through comparatively small neighborhoods. If they can determine whether their neighbors are informed or not, informed agents can also gain through consulting on random wheel networks (representing Facebook), but it requires far larger neighborhoods than for the other networks. While uninformed agents gain in certainty equivalents on all of the networks, the gains are far greater on the hub networks, even for small neighborhoods. If they can determine the type of their neighbors, uninformed agents can also gain through consulting random wheel networks, though, as with informed agents, they require large neighborhoods to achieve this. There appears to be little benefits for uninformed agents to consult on non-random wheel networks.

Various social learning models include agents receiving a private informative signal which they can share with other agents on their network (for example, Acemoglu et al. (2011) and Ozsoylev & Walden (2011)), though few include uninformed agents in their models. In a finance context these studies could be thought of as dealing with financial market professionals, as they can be expected to be more informed than the general public in terms of financial investment. Acemoglu et al. (2011) consider a model in which agents receive a return based on whether they correctly work out what the state of nature is. Prior to making their decisions, agents received a private signal and can observe a stochastically determined subset of the agents that have previously acted. The study determines certain conditions under which asymptotic learning occurs (that is, when the probability of agents making the correct decision approaches one as the number of agents goes to infinity). In the model of Ozsoylev & Walden (2011), agents receive a noisy signal about the return on a risky asset. As agents become more connected to other agents, their beliefs about the return on the asset become more accurate. Again all agents are equally likely to have an accu-

rate signal.

While these models give interesting insights into the nature of social learning, when it comes to portfolio choice it seems more likely that agents will share their beliefs about general asset classes to invest in rather than their portfolios (as per Acemoglu et al. (2011)) or their beliefs about the retruns of those assets (as per Ozsoylev & Walden (2011)). In this study's model, observations are made directly of other agent's beliefs and, even if an agent is certain about the state of nature, they will still not be certain about the returns on the assets. Furthermore, in contrast to these two models, this model includes agents who do not receive a private signal as well as those who do, and so can distinguish the gains that might accrue to both types of agent given the structure of the network.

A model of social learning that includes agents with heterogeneously informative signals is that of Jadbabaie et al. (2013). In this model, agents receive signals that can distinguish between some but not all potential states of nature. Given that only one state of nature is drawn in any period only some agents' signals will be informative. The authors then analyze how quickly the agents learn the true state of nature through consultation. From a finance perspective it is not clear that it is possible to completely distinguish the true state of nature (particularly if there are many potential states of nature). Moreover, the heterogeneity of the agents' information in the model of Jadbabaie et al. (2013) is not determined categorically but rather only after the state of nature is determined. In my model, no agent's signal is completely informative, but it is possible for an agent to know categorically how informed they are. This appears to be more relevant to a real world setting in which agents should know how informed they are likely to be, even if they cannot be certain that their information is correct.

Sections 2 and 3 introduce the general and specific models to be investigated in this paper. Section 4 presents the analysis of the results of consultation on different networks and conclusions are provided in Section 6.

2 General Model

In the model, there is a set, $N = \{1, \dots, n\}$ with $n \geq 3$, of individuals in a society who choose an action from a finite set of alternatives, A_i . It is assumed that all individuals have the same set of alternatives, i.e. $A_i = A_j = A$ for all individuals. Denote by a_i the action taken by individual i . The pay-offs associated with an action depend on the state of nature θ , which belongs to a finite set, Θ . The state

of nature is exogenously determined.

If θ is the true state of nature and individual i chooses action $a_i \in A$, then she will observe outcome $y \in Y$ with conditional density $\phi(y, a; \theta)$ and obtains a reward $r(a, y)$.

Individuals do not know the true state of nature but their private information is summarized in beliefs over the set of states. For individual i this belief is denoted π_i . The set of beliefs is denoted $\mathcal{P}(\Theta)$. It is assumed that $0 < \pi_i(\theta) < 1, \forall \theta$ and $i \in N$. Given belief π , an individual's expected utility function from action a is:

$$u(a, \pi) = \sum_{\theta \in \Theta} \pi(\theta) \int_Y r(a, y) \phi(y, a; \theta) dy$$

Given a belief, π , an agent chooses an action that maximizes the expected pay-offs. Formally, $B : \mathcal{P}(\Theta) \rightarrow A$ is the optimality correspondence such that:

$$B(\pi) = \{a \in A \mid u(a, \pi) \geq u(a', \pi), \forall a' \in A\}$$

For each $i \in N$, let $b_i : \mathcal{P}(\Theta) \rightarrow A$ be a selection from the optimality correspondence B . Let δ_θ represent a point mass belief on the state θ ; then $B(\delta_\theta)$ denotes the set of optimal actions if the true state is θ .

3 Specific Model

There are two assets, a risk-free asset and a risky asset. The agents' action set is $A = \{a_0, a_1\}$ and $a_i \in \mathbb{R}$, where a_0 is the share of wealth invested in the risk-free asset and a_1 is the share of wealth invested in the risky asset. By assumption, $a_0 + a_1 = 1$.¹ The states of nature are $\Theta = \{\theta_0, \theta_1\}$. The outcome observed by all agents (y) is the return on the risky asset and the reward received by the agents ($r(a, y)$) is the return on their investment portfolio.

The return on the risk-free asset is R_f in both states of nature, while the return on the risky asset (y) is either R_h or R_l , with $R_l < R_f < R_h$. The return on the agent's portfolio is then $r = [ay + (1 - a)R_f]$. In state θ_1 , the probability that the return on the risky asset is R_h is $P(y = R_h \mid \theta = \theta_1) = P_1$, and in state θ_0 , $P(y = R_h \mid \theta = \theta_0) = P_0$. It is assumed that $P_0 < P_1$. Therefore, the conditional expected return on the risky asset is:

¹ Given this assumption, $a_0 = 1 - a_1$ and so the agents' action set becomes $A = \{a_1\}$. For convenience sake, we drop the subscript on a from here.

$$\begin{aligned}
\tilde{R} &= P_1 R_h + (1 - P_1) R_l & \text{if state} = \theta_1 \\
\hat{R} &= P_0 R_h + (1 - P_0) R_l & \text{if state} = \theta_0 \\
\text{and} \quad \tilde{R} &> R_f > \hat{R}
\end{aligned}$$

It is assumed that the possible returns and the probability of those returns are known to the agents.

The agents have a prior belief about the probability that the state of nature is state 1, π_i .² Some agents (called “analysts”) then obtain a private signal as to the state of nature. After the analysts receive their signals, agents on the networks can consult with the other agents in their neighborhood to update their beliefs about the state of nature. The agents then decide a portfolio and the outcome and pay-offs are determined based on the state of nature and the agents’ actions.

3.1 Network Structures and Scenarios

Each individual is located as a node on a network. They are connected to a subset of other agents on the network with whom they can consult (share beliefs). It is assumed that agents know their own type (i.e. analyst or uninformed agent), the general nature of the network on which they are located, in terms of the general structure of the network and the distribution of the agents between informed and uninformed agents, and how many other agents they will consult. It is assumed that, before consultation, agents do not know the type of the agents to whom they are connected.

Two scenarios are considered. Under one scenario, through the process of consultation, agents are able to perfectly determine the type of their neighbors, while under the other scenario they cannot determine their neighbors type at all and so assume that the probability that their neighbors are analysts is equal for all of their neighbors. These two scenarios are designed as the extremes of knowledge that an agent can gain about the knowledgeability of their neighbors. In reality, people can probably discern with some degree of accuracy how knowledgeable their neighbors are. The scenarios were chosen in order to abstract from modeling how accurately people could discern how knowledgeable

²Given that there are only two states of nature, the agent’s belief about the probability that the state of nature is state 0 is $1 - \pi_i$, and so their beliefs are fully described with reference to state one.

their neighbors are and instead see these scenarios as border cases with reality somewhere in between

The potential gains from consultations were analyzed using three different network structures. These network structures were selected to represent three different social networks through which people might gain information about financial markets. The first was a hub network, in which the informed agents formed a central hub (which was itself a wheel network), while the uninformed agents were situated outside this hub and could observe some of those within the hub (Figure 1, Top Panel). This network could be thought of as uninformed investors reading finance-related discussion boards and blogs on the Internet, or gaining advice from finance professionals who have consulted with other finance professionals. The second network was a wheel network in which the informed agents (that is, those agents that received a signal) were randomly distributed throughout the network (Figure 1, Middle Panel). This could be thought of as resembling Facebook, where people are grouped without specific reference to their knowledge of financial matters. The final network was also a wheel network but on this network the informed agents were grouped together with links to the uninformed agents (those who did not receive a private signal) only at the margins (Figure 1, Bottom Panel). This can be thought of as resembling LinkedIn, where people are grouped with some reference to their occupation and so finance workers might be more likely to be grouped together than on Facebook.

It should be noted that these network structures are only stylized renderings of the networks that they are supposed to represent. For one thing the networks are unlikely to be balanced. On these networks, some agents will have more connections than others (in some cases many more connections). The number of connections agents have will affect how much influence they have over the network and potentially how much weight each agent to whom they are connected will place on their opinions. However, this influence will only affect their opinions after at least one round of consultation has been completed. In this analysis, agents only consult with each other for one round, which will only include the opinions of their neighbors based on their neighbor's signals, not on their neighbors consultations. Therefore, this is not a major consideration in this analysis.

Another difference between these network structures and their real-world analogs is that the financially- informed agents are unlikely to be completely

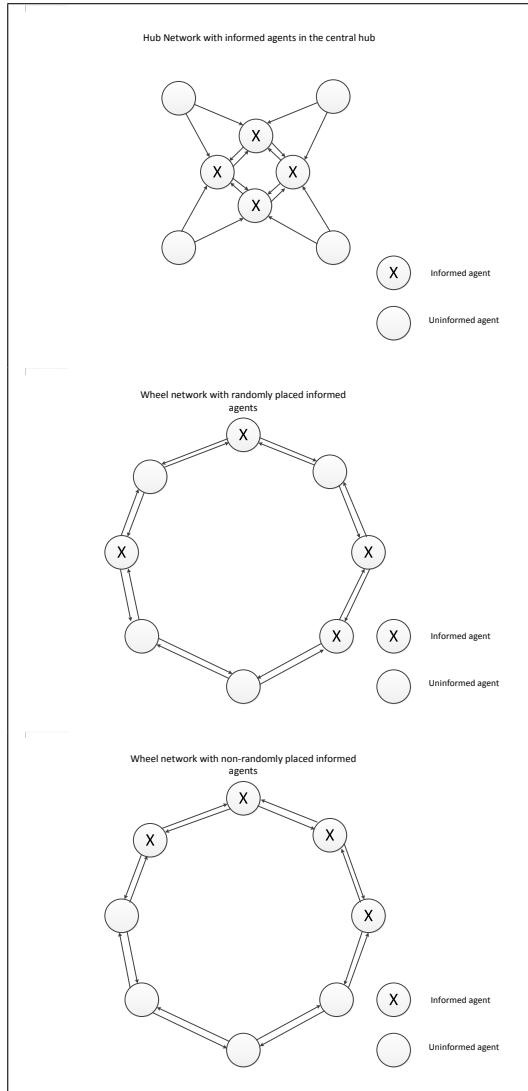


Figure 1: Network Structures: The three basic network structures that were analyzed were: a hub network with the informed agents forming a central hub, while the uninformed agents are situated outside this hub observing some of those within the hub (top panel); a wheel network with informed agents randomly placed around the network (middle panel); a wheel network with informed agents grouped together with links to uninformed agents only at the margins (bottom panel); In each of the diagrams the arrows indicate the direction of observation.

randomly spread over Facebook nor completely clustered on LinkedIn. On Facebook people connect with people they work with and so finance professionals are likely to be more clustered than a completely random assignment would suggest, while people on LinkedIn also connect with people that they don't work with. Nonetheless, it is still likely that the clustering of financially-knowledgeable people is lower on Facebook than on LinkedIn, as LinkedIn is more specifically geared to people's employment. Therefore, the interpretation should be that these structures are probably more extreme than in the real world and that the real world conclusions are probably somewhat less extreme than suggested here. Nonetheless, these structures should give some idea of the relative benefits of consulting on each of these networks to glean financial information.

3.2 Agents' Beliefs

Each agent has beliefs about the probability that the state of nature is state 1. This belief is updated after the agent receives a signal as to the state of nature (if the agent does receive a signal); and during consultation with other agents. Therefore, we can define the beliefs of agent i as:

π_{1i} belief before receiving a signal

π_{2i} belief after receiving a signal, before consultation

π_{3i} belief after consultation

The agent's belief, π_{1i} , will summarize all relevant information known to the agent before receiving a signal. This might include, past returns on the risky asset, any public information about the assets and any consultation that the agent has previously engaged in. It is assumed that all agents on the network have independent and identically distributed beliefs, prior to agents receiving their signals.

3.3 Signals

Some agents receive a private signal of the state of nature. The agents that receive this signal are called "informed agents" or "analysts". The agents that do not receive this signal are called "uninformed agents". The signal received by agent i is θ_i^m . The accuracy of this signal, η , is assumed to be identical for all analysts and is defined as $\eta = Pr(\theta_i^m = \theta_j | \theta = \theta_j), \eta \in [0.5, 1]$. This means that it is

likely that the analysts will not all receive the same signal but that each analyst has an equal probability of receiving a correct signal. If an agent has a belief, prior to obtaining the signal, that $Pr(\theta = \theta_1) = \pi_{1i}$, then according to Bayes' Theorem, the probability that the state of nature is state 1, given that the signal indicated that it is state 1 is:

$$\begin{aligned} P(\theta = \theta_1 | \theta_i^m = \theta_1) &= \frac{\eta \pi_{1i}}{\eta \pi_{1i} + (1 - \eta)(1 - \pi_{1i})} \\ &= \pi_{B1i} \end{aligned}$$

The probability that the state of nature is state 1 given that the signal indicated that it will be state 0 will be:

$$\begin{aligned} P(\theta = \theta_1 | \theta_i^m = \theta_0) &= \frac{(1 - \eta) \pi_{1i}}{(1 - \eta) \pi_{1i} + \eta(1 - \pi_{1i})} \\ &= \pi_{B0i} \end{aligned}$$

Therefore, the agent will set

$$\pi_{2i} = \begin{cases} \pi_{B1i} & \text{if } \theta_i^m = \theta_1 \\ \pi_{B0i} & \text{if } \theta_i^m = \theta_0 \end{cases}$$

Uninformed agents do not receive a signal and so for them $\pi_{2i} = \pi_{1i}$.

Here we define the concept of “correctly certain” (CC) beliefs. CC beliefs are defined as $\pi = 1$, given that the state is state one, and $\pi = 0$ given that the state is state zero.

Proposition 1 *Beliefs that are closer to correctly certain (CC) beliefs lead to higher expected returns*

See Section A.1 in the Appendix for the proof of this proposition.

Proposition 2 *Analysts' post-signal beliefs will be closer to CC beliefs than their prior beliefs regardless of the state of nature or prior beliefs.*

See Section A.2 in the Appendix for the proof of this proposition. Informally, this follows from the fact that analysts have more information on which to base their beliefs, as long as the signal is informative.

These two propositions combine to show that analysts will expect to have higher returns than the uninformed agents in the absence of consulting. Through the course of consulting, those agents that place greater weight on the beliefs of analysts should also expect that their post-consultation beliefs will be closer to CC beliefs and hence will have higher expected returns on their portfolios, *ceteris paribus*.

3.4 Updating beliefs through Consultation

After the analysts have received their signals and updated their beliefs based on those signals, all agents are able to consult with their neighbors; that is, they can observe their neighbors' beliefs about the state of nature. The method through which agents incorporate their neighbors' beliefs into their own updated beliefs is based on the DeGroot (1974) model.

In the DeGroot model, each agent is located as a node on a network. The set $N = \{1, \dots, n\}$ is the set of nodes (agents). A real-valued $n \times n$ matrix g describes the relationships between each of the nodes. The relationship between two agents, i and j , is denoted by g_{ij} , where $g_{ij} \in \{0, 1\}$. $g_{ij} = 1$ denotes an information flow from agent j to agent i , namely agent i can observe something about agent j . In the context of this model, the observation is the agent's belief about the state of nature (i.e. the probability that agent j places on it being state 1, π_{2j}). This study does not explore reasons why agents might wish to share their information with other agent, beyond the possibility of increasing the certainty equivalent of their investment. However, other researchers have investigated reasons for exchanging information including; extrinsic and intrinsic motivations (Lin 2007), mediating the effects of social capital on competitiveness (Wu 2008), and increased profits from increased precision from the pooled information (Vives 2007).

The network is directed so that $g_{ij} = 1$ does not imply that $g_{ji} = 1$, that is, if agent i can observe agent j this does not necessarily imply that agent j can observe agent i ³. The set of agents that agent i can observe is called the *neighborhood* of i , denoted $N_i^d(g) = \{k \in N | g_{ik} = 1\}$. The number of agents that an agent is linked to (that is, the cardinality of the agent's neighborhood) is called that agent's *degree*, and is denoted $d_i(g) = \#\{j | g_{ji} = 1\} = \#N_i^d(g)$ (Jackson 2008, p. 29).

³By assumption, $g_{ii} = 1$.

Each agent assigns a weight to the information that they receive from each person in their neighborhood. This weight is represented by a number, t_{ij} , where:

$$t_{ij} = \begin{cases} x_{ij} \in [0, 1] & \text{if } g_{ij} = 1 \\ 0 & \text{if } g_{ij} = 0 \end{cases}$$

and

$$\sum_{j=1}^n t_{ij} = 1 \quad \forall i$$

The information transfer matrix, T , is defined as:

$$T_{n,n} = \begin{pmatrix} t_{11} & t_{12} & \cdots & t_{1n} \\ t_{21} & t_{22} & \cdots & t_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ t_{n1} & t_{n2} & \cdots & t_{nn} \end{pmatrix}$$

Each of the rows in T will sum to 1. However, the columns will not necessarily also sum to 1.

To calculate how information is passed around the network we can set up the vector of initial beliefs about the state of nature (i.e. the vector of beliefs held by each of the agents immediately prior to consultation). Namely:

$$\Pi_2 = \begin{pmatrix} \pi_{21} \\ \pi_{22} \\ \vdots \\ \pi_{2n} \end{pmatrix}$$

After information sharing the beliefs matrix will be:

$$\Pi_3 = T\Pi_2$$

The belief of agent i after they have finished consulting with their neighbors is π_{3i} . The beliefs of agents prior to investment will therefore be based on their prior beliefs about the state of nature, any private signals they might receive, the beliefs of their neighbors and the weights that they place on their neighbors' beliefs. Apart from the weight that they place on their neighbors' beliefs,

agents are, therefore, assumed to use only relevant information and to use this information correctly when coming to their beliefs. The psychological literature suggests that, in certain groups, the members do not always act rationally in coming to their beliefs and opinions about the world.⁴ However, in trying to abstract from such concerns to examine the economics of the decision making, I have not included such possibilities in this model.

A note here should be made about why the DeGroot model was chosen to incorporate neighbors' beliefs into agents updated beliefs rather than another model; for example, one based on Bayes Rule. There were three main reasons for this choice. The first was that, to use Bayes Rule, agents need to know the evidence on which their neighbors' beliefs are based. In this model, with agents having a distribution of priors, it is not clear that an agent could infer from their neighbors' posterior beliefs what their signals were and so it would not be possible to use Bayes Rule.

Secondly, various studies have shown that people generally do not use Bayes Rule but rather simpler methods of incorporating information, such as the DeGroot model (see for example: Acemoglu et al. (2011), Holt & Smith (2009), Charness et al. (2007), Tenenbaum et al. (2006), Charness & Levin (2005), Gale & Kariv (2003), El-Gamal & Grether (1995), Grether (1992)). Most of these studies suggest that people sometimes use Bayes rule but usually only under certain specific conditions, such as that they are sophisticated agents in large networks and can distinguish between new and repeated information, but even then Bayes rule is an incomplete explanation for how people incorporate information. Finally, the model becomes far more tractable if a simpler model is used.

The use of Bayes rule to incorporate personal experience into beliefs and the DeGroot model to incorporate others' beliefs is also consistent with previous studies such as those of Jadbabaie et al. (2012) and Jackson & Golub (2007).

3.5 Utility maximization

After consultation with their neighbors, each agent chooses the share of wealth that she will invest in the risky asset, a_i , with the balance invested in the risk-free asset. All agents are assumed to be small, price takers so that their choice of portfolio does not affect the returns on the assets. This leads to the utility maximization problem of:

⁴One well known example of this is groupthink. For a review of the literature concerning groupthink see Turner & Pratkanis (1998).

$$\begin{aligned}
Max_{a_i} \quad E(U_i) = & \pi_{3i} \left[P_1 U[a_i W_i R_h + (1 - a_i) W_i R_f] \right. \\
& \left. + (1 - P_1) U[a_i W_i R_l + (1 - a_i) W_i R_f] \right] \\
& + (1 - \pi_{3i}) \left[P_0 U[a_i W_i R_h + (1 - a_i) W_i R_f] \right. \\
& \left. + (1 - P_0) U[a_i W_i R_l + (1 - a_i) W_i R_f] \right]
\end{aligned}$$

where W_i is the agent's wealth before investment.

By assumption, agents have an iso-elastic (power) utility function:

$$U(W_i) = \frac{W_i^{1-\gamma} - 1}{1-\gamma} \quad \gamma > 0$$

where γ is a measure of risk aversion. Therefore, the utility maximization problem becomes:

$$\begin{aligned}
Max_{a_i} \quad E(U) = & \frac{\pi_{3i} P_1}{1-\gamma} \left[[a_i W_i R_h + (1 - a_i) W_i R_f]^{1-\gamma} - 1 \right] \\
& + \frac{\pi_{3i} (1 - P_1)}{1-\gamma} \left[[a_i W_i R_l + (1 - a_i) W_i R_f]^{1-\gamma} - 1 \right] \\
& + \frac{(1 - \pi_{3i}) P_0}{1-\gamma} \left[[a_i W_i R_h + (1 - a_i) W_i R_f]^{1-\gamma} - 1 \right] \\
& + \frac{(1 - \pi_{3i}) (1 - P_0)}{1-\gamma} \left[[a_i W_i R_l + (1 - a_i) W_i R_f]^{1-\gamma} - 1 \right]
\end{aligned}$$

This maximization problem is solved by setting a_i to the following value:

$$a_i^* = \frac{R_f (g_i^{\frac{-1}{\gamma}} - 1)}{(R_f - R_l) (g_i^{\frac{-1}{\gamma}} - 1) + R_h - R_l}$$

where,

$$g_i = \frac{-(R_l - R_f) \left[\pi_{3i} (1 - P_1) + (1 - \pi_{3i}) (1 - P_0) \right]}{(R_h - R_l) \left[\pi_{3i} P_1 + (1 - \pi_{3i}) P_0 \right]}$$

The optimal value of a_i , a_i^* , will be continuously increasing in π_{3i} .

Following the choice of portfolios, the return on the risky asset is determined and so are the returns on the agents' portfolios.

4 Analysis

In analyzing this model, a number of simplifying assumptions are made.

Assumption 1 *Agents know the structure of the network and the shares of the types of agents on the network but not their own position on the network.*

Assumption 2 *Agents' pre-signal beliefs are independent and identically distributed.*

Assumption 3 *Agents do not know their neighbors' types before consultation. But they might be able to determine their neighbors' type through the process of consultation.*

Assumption 4 *Agents know the network on which they are consulting and the number of agents they are consulting (k).*

Assumption 5 *The number of analysts in agent i 's neighborhood (ψ_i) is independent of all agents pre-consultation beliefs (π_{2j}).*

The analysis conducted in this paper attempts to answer the question of which network it is best to consult to gain knowledge about financial markets. Therefore, for individual agents the question is; through consultation on which networks is their certainty equivalent expected to be greatest and under which circumstances (such as the size of their neighborhood). A first step is to consider what the agents expect their beliefs will be after consultation, given the size of their neighborhood and the network on which they will consult, and what they expect the variance in those beliefs will be.

Under Assumption 2 (that the agents pre-signal beliefs are iid), the pre-consultation beliefs of uninformed agents will also be iid (as their pre-consultation beliefs will simply be their pre-signal beliefs). Given that the signals that analysts receive are independent of each other, the analysts' pre-consultation beliefs will also be iid (though likely with a different distribution to the uninformed agents' beliefs). Therefore, the agents' pre-consultation beliefs then can be thought of as independent random variables drawn from two populations; one for the analysts and one for the uninformed agents.

As shown in the previous section, agents' beliefs after consultation will be a convex combination of the pre-consultation beliefs of their neighbors. Therefore, given expectations about the number of analysts in agents' neighborhoods,

we can determine the expected post-consultation beliefs and variance in agent's beliefs as convex combinations of the pre-consultation expected beliefs and variance in beliefs of the two types of agent, and hence compare the different networks under the two scenarios.

4.1 Expected Beliefs and Variance in Beliefs after Consultation

Through consulting their neighborhoods, agents update their beliefs as weighted averages of their neighbors' beliefs. The absolute weights that agent i places on one of their neighbor's, j , beliefs is t_{ij} . As will be shown the effects of consultation on agent's beliefs and hence on their optimal portfolios will be dictated by how much overall weight is placed on analysts beliefs in the agent's updated beliefs. That is, the sum of the absolute weights placed on the analysts in the agent's neighborhood.

In analyzing the model, however, it is easier to work with relative weights, as these are more separable. The relative weight that agent i places on one of their neighbor's, j , beliefs is w_{ij} . The difference between w_{ij} and t_{ij} is that the sum of the t_{ij} across all j for each agent must be equal to 1, while there is no such restriction on the relative weights, w_{ij} . The relationship between the relative weights w_{ij} and the absolute weights, t_{ij} , is:

$$t_{ij} = \frac{w_{ij}}{\sum_{j=1}^{k+1} w_{ij}}$$

Where k is the total number of agents in an agent's neighborhood (excluding the agent herself), w_{ij} is the relative weight agent i puts on the beliefs of agent j in consultation. In this model, the agent herself is the $(k+1)^{th}$ agent in the neighborhood.

The beliefs of an agent after one round of consultation will then be:

$$\pi_{3i} = \frac{\sum_{j=1}^{\psi_i} w_{ij} \pi_{2j} + \sum_{j=\psi+1}^k w_{ij} \pi_{2j} + w_{ii} \pi_{2i}}{\sum_{j=1}^{k+1} w_{ij}}$$

Where ψ_i is the number of analysts in agent i 's neighborhood (not including the agent herself if she is an analyst), π_{2j} is the pre-consultation belief of agent j , and w_{ii} is the relative weight an agent puts on her own beliefs in consultation.

Therefore, the first expression in the numerator is the contribution to the agent's post-consultation beliefs that comes from the analysts in her neighborhood, the second term is the contribution from uninformed agents and the fi-

nal term in the numerator is the contribution from her own pre-consultation beliefs.

The number of analysts in an agent's neighborhood (ψ_i) is unknown to the agent before she consults with her neighborhood. However, prior to consultation she can have some expectations over ψ_i given the network on which she is consulting, how large her neighborhood is and the relative numbers of analysts and uninformed agents on the network.

The signals received by analysts are independent, and so the pre-consultation beliefs of agents will be independent. Therefore, conditional on ψ , the variance in agent i 's beliefs post-consultation will be:

$$\begin{aligned}
 Var(\pi_{3i}|\psi) &= Var\left(\sum_{j=1}^{k+1} \frac{w_{ij}}{\sum_{j=1}^{k+1} w_{ij}} \pi_{2j}\right) \\
 &= \sum_{j=1}^{k+1} \left(\frac{w_{ij}}{\sum_{j=1}^{k+1} w_{ij}}\right)^2 Var(\pi_{2j}) \\
 &= \sum_{j=1}^{\psi_i} \left(\frac{w_{ij}}{\sum_{j=1}^{k+1} w_{ij}}\right)^2 Var(\pi_{2j}) + \sum_{j=\psi_i+1}^k \left(\frac{w_{ij}}{\sum_{j=1}^{k+1} w_{ij}}\right)^2 Var(\pi_{2j}) \\
 &\quad + \left(\frac{w_{ii}}{\sum_{j=1}^{k+1} w_{ij}}\right)^2 Var(\pi_{2i})
 \end{aligned}$$

The first term is the contribution to the variance in the agent's beliefs from the analysts in her neighborhood, the second term is the contribution from the uninformed agents and the final term is the contribution from the variance in her own pre-consultation beliefs.

Based on the general expressions for the expected beliefs and variance in those beliefs of agents post-consultation, we can consider different scenarios based on how much knowledge the agent can gain about the other people in their neighborhood.

4.1.1 Scenario 1: Perfect Knowledge of Neighbors

Under the first scenario of perfect neighbor knowledge, each agent can determine the type of the agents in their neighborhood through the process of consultation. Therefore, while agents do not know the types of agents in their neigh-

borhoods before consultation, they can determine the types of their neighbors perfectly after consultation. Because each agent knows the type of each of the agents in her neighborhood after consultation, she gives zero weight to the uninformed agents and equal weight to the analysts in their neighborhood, including the agent herself if she is an analyst. Therefore, under this scenario $w_{ij} = 1$ for all neighbors who are analysts, $w_{ij} = 0$ for all neighbors who are uninformed agents, $w_{ii} = 1$ if the agent is an analyst and $w_{ii} = 0$ if the agent is an uninformed agent (and $\psi_i \geq 1$). If the agent is an uninformed agent and there are no analysts in her neighborhood, the agent puts equal weight on each of the agents in her neighborhood including her own beliefs.

In this scenario, an analyst's beliefs after consultation (given ψ_i) will be:

$$\pi_{3i} = \frac{\sum_{j=1}^{\psi_i} \pi_{2j} + \pi_{2i}}{\psi_i + 1}$$

Prior to consultation, ψ_i and π_{2j} will be unknown to the agent. Nonetheless, she can have expectations over these variables and, under the assumption that ψ_i and π_{2j} are independent, the expectations over the post-consultation beliefs will be:

$$\begin{aligned} E[\pi_{3i}] &= E\left[\frac{\psi_i + 1}{\psi_i + 1}\right] \pi_a \\ &= \pi_a \end{aligned}$$

Where π_a is the expected pre-consultation beliefs of analysts.

An uninformed agent's beliefs after consultation (when $\psi_i \geq 1$) will be:

$$\pi_{3i} = \frac{\sum_{j=1}^{\psi_i} \pi_{2j}}{\psi_i}$$

Taking expectations over this gives:

$$\begin{aligned} E[\pi_{3i} | \psi_i \geq 1] &= E\left[\frac{\psi_i}{\psi_i}\right] \pi_a \\ &= \pi_a \end{aligned}$$

When $\psi_i = 0$, an uninformed agent's beliefs after consultation will be:

$$\pi_{3i} = \frac{\sum_{j=1}^{k+1} \pi_{2j}}{k+1}$$

Taking expectations over this gives:

$$\begin{aligned} E[\pi_{3i} | \psi_i = 0] &= \frac{k+1}{k+1} \pi_u \\ &= \pi_u \end{aligned}$$

Where π_u is the expected pre-consultation beliefs of an uninformed agent. Therefore, an uninformed agent's expected post-consultation beliefs will be:

$$E[\pi_{3i}] = P(\psi_i = 0)\pi_u + P(\psi_i \geq 1)\pi_a$$

Therefore, if agents can perfectly determine the type of the agents in their neighborhoods then analysts will expect that their beliefs after consultation will be the same as the expected pre-consultation beliefs of the analysts. Uninformed agents will expect that their beliefs will approach those of analysts as the probability that there are no analysts in their neighborhood decreases.

Under perfect neighbor knowledge the variance of an analyst's beliefs post-consultation given the number of analysts in their neighborhood, ψ_i , will be:

$$\begin{aligned} Var(\pi_{3i} | \psi_i) &= \sum_{j=1}^{\psi_i} \left(\frac{1}{\psi_i + 1} \right)^2 Var(\pi_{2j}) + \left(\frac{1}{\psi_i + 1} \right)^2 Var(\pi_{2i}) \\ &= \frac{1}{(\psi_i + 1)^2} \left(\sum_{j=1}^{\psi_i} Var(\pi_{2j}) + Var(\pi_{2i}) \right) \end{aligned}$$

Given that the agent is an analyst,

$$Var(\pi_{2j}) = Var(\pi_{2i}) \quad \forall j$$

and so

$$Var(\pi_{3i}|\psi_i) = \frac{1}{\psi_i + 1} Var(\pi_2^a)$$

Where $Var(\pi_2^a)$ is the variance of analysts' pre-consultation beliefs. Taking expectations over ψ_i gives:

$$E[Var(\pi_{3i})] = E\left[\frac{1}{\psi_i + 1}\right] Var(\pi_2^a)$$

If $\psi_i \geq 1$, the variance of an uninformed agent's beliefs will be:

$$Var(\pi_{3i}|\psi_i \geq 1) = \sum_{j=1}^{\psi_i} \left(\frac{1}{\psi_i}\right)^2 Var(\pi_2^a)$$

Taking expectations over ψ_i gives:

$$E[Var(\pi_{3i}|\psi_i \geq 1)] = E\left[\frac{1}{\psi_i}|\psi_i \geq 1\right] Var(\pi_2^a)$$

If $\psi_i = 0$, the variance of an uninformed agent's beliefs will be:

$$\begin{aligned} Var(\pi_{3i}|\psi_i = 0) &= \sum_{j=1}^{k+1} \left(\frac{1}{k+1}\right)^2 Var(\pi_{2j}) \\ &= \frac{1}{k+1} Var(\pi_2^u) \end{aligned}$$

where $Var(\pi_2^u)$ is the variance in uninformed agents' pre-consultation beliefs.

Therefore, the expected variance in post-consultation beliefs for an uninformed agent with perfect neighbor knowledge is:

$$E[Var(\pi_{3i})] = P(\psi_i = 0) \frac{1}{k+1} Var(\pi_2^u) + P(\psi_i \geq 1) E\left[\frac{1}{\psi_i}|\psi_i \geq 1\right] Var(\pi_2^a)$$

Thus, under perfect neighbor knowledge, both analysts and uninformed agents will expect that the variance in their beliefs will decrease as the number of analysts in their neighborhood increases, though uninformed agents might have greater variance after consultation compared with before if the variance of analysts' beliefs is greater than that of uninformed agents.

4.1.2 Scenario 2: No Knowledge of Neighbors

Under the second scenario, agents cannot determine the type of their neighbors through consultation, and so they have no knowledge of the type of the agents in their neighborhood. Therefore, agents give equal weight to all of their neighbors. The relative weight that they give to their neighbors is based on the probability that any individual neighbor is an analyst, that is $w_{ij} = v_i, \forall j, j \neq i$, where v_i is the probability that an agent in agent i 's neighborhood is an analyst. The expected number of analysts in an agent's neighborhood will then be $E[\psi_i] = kv_i$. It should be noted that k and v_i might not be independent. The agent knows their own type so $w_{ii} = 1$ if the agent is an analyst and $w_{ii} = 0$ if the agent is an uninformed agent. See Section A.3 in the Appendix for the derivation of the following results.

Under this scenario, an analyst's expected beliefs after consultation will be:

$$E[\pi_{3i}] = \frac{(kv_i^2 + 1)\pi_a + kv_i(1 - v_i)\pi_u}{kv_i + 1}$$

An uninformed agent's beliefs after consultation will be:

$$E[\pi_{3i}] = v_i\pi_a + (1 - v_i)\pi_u$$

The expected variance of an analyst's beliefs post-consultation will be:

$$E[Var(\pi_{3i})] = \frac{kv_i^3 + 1}{(kv_i + 1)^2} Var(\pi_2^a) + \frac{kv_i^2(1 - v_i)}{(kv_i + 1)^2} Var(\pi_2^u)$$

The expected variance of an uninformed agent's beliefs will be:

$$E[Var(\pi_{3i})] = \frac{v_i}{k} Var(\pi_2^a) + \frac{(1 - v_i)}{k} Var(\pi_2^u)$$

Therefore, under the scenario of no knowledge of neighbors' type, given that the absolute weight that analysts will place on the beliefs of analysts post-consultation is expected to be less than 1, analysts expect that their beliefs will move away from CC beliefs as their neighborhood increases. For uninformed

agents, however, they expect that their beliefs will be closer to CC beliefs after consultation compared with before consultation, but the size of their neighborhood will not affect how close to CC beliefs they are as long as the probability that a neighbor is an analyst, ν_i , is unaffected by the size of the neighborhood.

Analysts expect that the variance in their beliefs will decrease with the size of their neighborhood (as long as the variance of analysts' beliefs is greater than the variance of uninformed agents beliefs). Uninformed agents expect that the variance in their beliefs will decrease with the size of their neighborhood as long as the probability of a neighbor being an analyst is unaffected by the size of the neighborhood. However, the variance of uninformed agents beliefs might increase with consultation compared with no consultation if the variance of analysts' beliefs is greater than that of uninformed agents.

4.2 Consulting on the Three Specific Networks

Having considered what will happen to the beliefs and the variance of beliefs of agents after consultation, and what is expected to happen to those beliefs and variance in beliefs, we can now consider what will happen on the specific networks that were chosen for this study.

4.2.1 Hub Networks

When agents consult on a hub network they know that all of their neighbors will be analysts simply because of the structure of the network, that is, $\psi_i = k$ for all agents. Therefore, the perfect neighbor knowledge scenario is the only one that applies. Furthermore, the uninformed agents will certainly have at least one analyst in their neighborhood. Therefore, we can say that the expected beliefs of both analysts and uninformed agents will be π_a . This means that the analysts' beliefs are not expected to move any closer to CC beliefs after consultation, while the uninformed agents' beliefs will (though no closer than the analysts' beliefs). This result follows from the fact that analysts in the hub networks in this model are essentially identical and that the DeGroot (1974) model of consultation uses a weighted average of beliefs rather than a Bayesian analysis. Therefore, given all publicly available information, the expected beliefs of any two (or more) agents before consultation will be identical. Their expected beliefs after consulting with each other will be a weighted average of their (identical) pre-consultation expected beliefs, and so will be unchanged.

If the uninformed agents on a hub network fully adopt the beliefs of the analysts with whom they consult (that is, they place no weight on their own pre-consultation beliefs) their expected post-consultation beliefs will be identical to those of the analysts. The uninformed agents will then have expected beliefs identical to that of the analysts if they consult with at least one analyst in each period.

The variance of analysts' beliefs will be:

$$\begin{aligned} Var(\pi_{3i}) &= \frac{1}{\psi_i + 1} Var(\pi_a) \\ &= \frac{1}{k + 1} Var(\pi_a) \end{aligned}$$

and the variance of uninformed agents' beliefs will be:

$$\begin{aligned} Var(\pi_{3i}) &= \frac{1}{\psi_i} Var(\pi_a) \\ &= \frac{1}{k} Var(\pi_a) \end{aligned}$$

Therefore, on a hub network, the variance of analysts' beliefs will fall post-consultation and will be less than that of uninformed agents who have the same size neighborhood. The benefits of consultation on a hub network for an analyst are not that they expect to improve their beliefs but rather that they will reduce the variance in their beliefs. For uninformed agents, the expected benefits are an improvement in their expected beliefs, and if they increase their neighborhood beyond a minimal one, a reduction in the variance in their beliefs.

4.2.2 Random Wheel Networks

Assumption 6 *On a random wheel network, agents pre-consultation beliefs about the number of analysts in their neighborhoods follows a binomial distribution.*

That is, agents assume that the number of analysts in their neighborhood is a binomial random variable.⁵ Under this assumption, each neighbor is effec-

⁵Strictly speaking this should be a hyper-geometric random variable as the neighbors are sampled without replacement. However, as long as the neighborhood is small compared with the size of the network and the probability of a neighbor being an analyst is sufficiently far away from zero or one the binomial distribution is a reasonable approximation.

tively a separate trial, and so the number of trials is k , the size of the neighborhood (not including the agent herself), and the probability of “success” (that a neighbor is an analyst) is v_i . Therefore, the expected number of analysts in agent i ’s neighborhood, ψ_i , is kv_i .

On random wheel networks, the probability of that a neighbor is an analyst, v_i , is the share of analysts on the network and is independent of the size of the neighborhood. That is:⁶

$$v_i = \frac{x}{x+y}$$

Where x is the number of analysts on the network as a whole, and y is the number of uninformed agents on the network as a whole. This probability is the same for analysts and uninformed agents. Therefore, the expected belief of an analyst after consultation on a random wheel network with **perfect neighbor knowledge** is π_a , while that of an uninformed agent will be:

$$\begin{aligned} E[\pi_{3i}] &= P(\psi_i = 0)\pi_u + P(\psi_i \geq 1)\pi_a \\ &= (1 - v_i)^k \pi_u + (1 - (1 - v_i)^k) \pi_a \\ &= \left(1 - \frac{x}{x+y}\right)^k \pi_u + \left(1 - \left(1 - \frac{x}{x+y}\right)^k\right) \pi_a \\ &= \left(\frac{y}{x+y}\right)^k \pi_u + \left(1 - \left(\frac{y}{x+y}\right)^k\right) \pi_a \end{aligned}$$

Therefore, as the number of neighbors, k , increases on a random wheel network with perfect neighbor knowledge, the expected post-consultation beliefs of uninformed agents will approach that of analysts.

We turn now to the variance of agents’ beliefs on a random wheel network with perfect neighbor knowledge (See Section A.4 in the Appendix for this derivation). The expected variance of an analyst’s post-consultation beliefs with perfect neighbor knowledge on a random wheel network will be:

⁶Strictly the probability that a neighbor is an analyst is $v_i = \frac{x-1}{x+y-1}$ if the agent is an analyst, and $v_i = \frac{x}{x+y-1}$ if the agent is an uninformed agent. However, for a large enough network these two expressions will be approximately equal to $\frac{x}{x+y}$.

$$E[Var(\pi_{3i})] = \frac{x+y}{(k+1)x} \left[1 - \left(\frac{y}{x+y} \right)^{k+1} \right] Var(\pi_2^a)$$

The expected variance of an uninformed agent's beliefs on a random wheel network with perfect neighbor knowledge will be:

$$E[Var(\pi_{3i})] = \left(\frac{y}{x+y} \right)^k \frac{1}{k+1} Var(\pi_2^u) + \left(1 - \left(\frac{y}{x+y} \right)^k \right) E \left[\frac{1}{\psi_i} | \psi_i \geq 1 \right] Var(\pi_2^a)$$

Therefore, on a random wheel network with perfect neighbor knowledge the expected post-consultation variance of analysts' beliefs will decrease as the neighborhood increases. For uninformed agents as the neighborhood increases the post-consultation variance of their beliefs is expected to fall. However, if the pre-consultation variance of analysts' beliefs is greater than that of uninformed agents' beliefs, the variance of uninformed agents' beliefs after consultation might be greater than compared with no consultation.

When agents have **no neighbor knowledge** on a random wheel network, the expected beliefs of an analyst will be (see Section A.5 in the Appendix for a derivation of this result):

$$E[\pi_{3i}] = \frac{[kx^2 + (x+y)^2]\pi_a + kxy\pi_u}{kx^2 + (x+y)^2 + kxy}$$

The expected beliefs of an uninformed agent will be:

$$E[\pi_{3i}] = \frac{x}{x+y} \pi_a + \frac{y}{x+y} \pi_u$$

Therefore, the expected weight on the pre-consultation beliefs of analysts will always be greater in the expected post-consultation beliefs of analysts compared with that of uninformed agents. This means that on random wheel networks with no neighbor knowledge the expected post-consultation beliefs of analysts will be closer to CC beliefs than those of uninformed agents regardless of the size of the neighborhood. However, the expected post-consultation beliefs of analysts will be further away from CC beliefs than they were before consultation, while those of uninformed agents will be closer to CC beliefs than before consultation.

The expected variance of analysts' beliefs on random wheel networks with no neighbor knowledge is (see Section A.6 in the Appendix for a derivation of this result):

$$E[Var(\pi_{3i})] = \frac{kx^3 + (x+y)^3}{(x+y)(kx+x+y)^2} Var(\pi_2^a) + \frac{kx^2y}{(x+y)(kx+x+y)^2} Var(\pi_2^u)$$

The expected post-consultation variance of uninformed agents on random wheel networks with no neighbor knowledge will be (see Section A.7 in the Appendix for a derivation of this result):

$$E[Var(\pi_{3i})] = \frac{x}{k(x+y)} Var(\pi_2^a) + \frac{y}{k(x+y)} Var(\pi_2^u)$$

Therefore, the variance of both analysts' and uninformed agents' beliefs will fall as the neighborhood size increases, though there might be an initial increase in variance for one of the types of agents if the difference in variance between analysts' and uninformed agents' pre-consultation beliefs is large.

4.2.3 Non-random Wheel Networks

Assumption 7 *On a non-random wheel network, it is assumed that k is an even number. That is, under consultation agents consult with an equal number ($\frac{k}{2}$) of their nearest neighbors on each side on the network.*

On non-random wheel networks the expected numbers of analysts in the neighborhoods of analysts will be different to those of uninformed agents. The make-up of analysts' and uninformed agents' neighborhoods can be determined if we know their position on the network. Restricting the analysis to situations where the neighborhoods are "small" compared with the network as a whole, that is, $k < \min(2x; 2y)$, we can determine the expected number of analysts in both analysts' and uninformed agents' neighborhoods.

Proposition 3 *The expected number of analysts in analysts' neighborhoods on a non-random wheel network is $E(\psi_a) = \frac{k}{x} \left(x - \frac{k}{4} - \frac{1}{2} \right)$*

Proposition 4 *The expected number of analysts in uninformed agents' neighborhoods on a non-random wheel network is $E(\psi_u) = \frac{k(k+2)}{4y}$*

For proofs of Propositions 3 and 4 see Sections A.8.1 and A.8.2 in the Appendix.

Agents, by assumption, do not know their position on the network. Therefore, on non-random wheel networks, they assume that it is equally likely that

they are situated in any of the possible positions on the network for their type. There are y possible positions for uninformed agents on the network. Of these y positions, $y-k$ positions will have no analysts in their neighborhood. Therefore, if an uninformed agent considers that all positions on the networks are equally likely the probability that they have no analysts in their neighborhood will be $P(\psi_i = 0) = \frac{y-k}{y}$.

Therefore, on non-random wheel networks with **perfect neighbor knowledge**, the expected beliefs of analysts will be π_a and the expected beliefs of uninformed agents will be:

$$\begin{aligned} E[\pi_{3i}] &= P(\psi_i = 0)\pi_u + P(\psi_i \geq 1)\pi_a \\ &= \frac{y-k}{y}\pi_u + \frac{k}{y}\pi_a \end{aligned}$$

Therefore, under perfect neighbor knowledge the expected post-consultation beliefs of uninformed agents on a non-random wheel network will approach those of analysts as the size of the neighborhood increases but at a much slower rate than that of the random wheel network.

To determine the expected variance of their post-consultation beliefs, agents need to consider the expected number of analysts in their neighborhoods. According to Proposition 3, the expected number of analysts in analysts' neighborhoods is:

$$E(\psi_a) = \frac{k}{x} \left[x - \frac{k}{4} - \frac{1}{2} \right]$$

Given that $E(\psi_a) = k\nu_a$, where ν_a is the probability that any particular neighbor in an analyst's neighborhood is themselves an analyst. Then:

$$\begin{aligned} \nu_a &= \frac{E(\psi_a)}{k} \\ &= 1 - \frac{k+2}{4x} \end{aligned}$$

This means that the probability that any particular neighbor on an analyst's neighborhood is an analyst decreases as the neighborhood increases.

The expected number of analysts in uninformed agents' neighborhoods, according to Proposition 4, is:

$$E(\psi_u) = \frac{k(k+2)}{4y}$$

Given that $E(\psi_u) = k\nu_u$, where ν_u is the probability that any particular neighbor in an uninformed agent's neighborhood is an analyst. Then:

$$\begin{aligned}\nu_u &= \frac{E(\psi_u)}{k} \\ &= \frac{(k+2)}{4y}\end{aligned}$$

This means that the probability that any particular neighbor on an uninformed agent's neighborhood is an analyst increases as the neighborhood increases.

The expected variance of analysts' post-consultation beliefs on a non-random wheel network under perfect neighbor knowledge will be:

$$E[Var(\pi_{3i})] = \left[\frac{x-k}{x(k+1)} + \frac{2}{x} \sum_{i=1}^{\frac{k}{2}} \left(\frac{1}{\frac{k}{2} + i} \right) \right] Var(\pi_2^a)$$

(See Section A.9 in the Appendix for the derivation of this result.)

The expected variance of uninformed agents' post-consultation beliefs on a non-random wheel network under perfect neighbor knowledge will be:

$$E[Var(\pi_{3i})] = \frac{y-k}{y(k+1)} Var(\pi_2^u) + \frac{2}{y} \sum_{i=1}^{\frac{k}{2}} \frac{1}{i} Var(\pi_2^a)$$

(See Section A.10 in the Appendix for the derivation of this result.)

Therefore, the expected variance of analysts' post-consultation beliefs will fall as the neighborhood increases under perfect neighbor knowledge on a non-random wheel network, and, while the dynamics of the variance of uninformed agents beliefs are not clear, as their neighborhood increases it becomes more likely that the variance of their beliefs will fall with consultation.

Under **no neighbor knowledge** the expected post-consultation beliefs of analysts on a non-random wheel network are:

$$E[\pi_{3i}] = \frac{k(4x - k - 2)^2 + 16x^2}{4xk(4x - k - 2) + 16x^2}\pi_a + \frac{k(k + 2)(4x - k - 2)}{4xk(4x - k - 2) + 16x^2}\pi_u$$

(See Section A.11 in the Appendix for the derivation of this result.)

Under no neighbor knowledge the expected post-consultation beliefs of uninformed agents on a non-random wheel network are:

$$E[\pi_{3i}] = v_i\pi_a + (1 - v_i)\pi_u$$

With $v_i = \frac{(k+2)}{4y}$

$$= \frac{(k + 2)}{4y}\pi_a + \left(\frac{4y - k - 2}{4y}\right)\pi_u$$

Therefore, the expected post-consultation beliefs of analysts will move away from CC beliefs as the size of the neighborhood increases but at a slower rate than on a random wheel network. The post-consultation beliefs of uninformed agents are expected to be closer to CC beliefs than before consultation but not as close as on a random wheel network unless the neighborhood is extremely large.

Under no neighbor knowledge the expected post-consultation variance of analysts' beliefs on a non-random wheel network is:

$$E[Var(\pi_{3i})] = \frac{k(4x - k - 2)^3 + 64x^3}{4x[k(4x - k - 2) + 4x]^2}Var(\pi_2^a) + \frac{k(k + 2)(4x - k - 2)^2}{4x[k(4x - k - 2) + 4x]^2}Var(\pi_2^u)$$

(See Section A.12 in the Appendix for a derivation of this result.)

Under no neighbor knowledge the expected post-consultation variance of uninformed agents' beliefs on a non-random wheel network is:

$$E[Var(\pi_{3i})] = \frac{v_i}{k}Var(\pi_2^a) + \frac{(1 - v_i)}{k}Var(\pi_2^u)$$

With $v_i = \frac{(k+2)}{4y}$

$$\begin{aligned} E[Var(\pi_{3i})] &= \frac{\frac{(k+2)}{4y}}{k} Var(\pi_2^a) + \frac{(1 - \frac{(k+2)}{4y})}{k} Var(\pi_2^u) \\ &= \frac{k+2}{4yk} Var(\pi_2^a) + \frac{4y-k-2}{4yk} Var(\pi_2^u) \end{aligned}$$

4.2.4 Summary

The expected beliefs of analysts do not get any closer to CC beliefs on any of the networks under either scenario. If the analysts do not know the type of their neighbors their beliefs will move away from CC beliefs after consultation. For uninformed agents, however, their beliefs will move closer to CC beliefs on all networks under both scenarios. The variance of analysts beliefs will fall on all networks under all scenarios. The dynamics of the variance in uninformed agents' post-consultation beliefs are not clear for all networks though in general the variance will tend to be smaller on larger neighborhoods compared with smaller neighborhoods.

4.3 Certainty Equivalents, Expected Returns and Variance in Returns

Having established the effects on expected beliefs and the variance in those beliefs from consulting on the different networks under the two scenarios for both analysts and uninformed agents, we can now turn to how those beliefs affect the certainty equivalents of the agents' investments to assess the potential gains to be made from consultation.

Expected beliefs and variance in beliefs affect the expected returns and variance in returns through the choice of portfolio. Without making assumptions about the distribution of beliefs and the parameters of the model it is difficult to be too conclusive about the effects of the consultation on the choice of portfolio and then on expected returns, variance in returns and ultimately on the certainty equivalents of the agents' investments. However, under certain reasonable assumptions we can make some general conclusions about what is likely to happen.

As the degree of risk aversion increases beyond a low level, the optimal portfolio function becomes near linear. In this case, changes in the expected optimal

portfolio are mainly affected by changes in the expected beliefs with only a relatively small effect coming from changes in the variance in beliefs. Therefore, in cases where consultation affects both the expected beliefs of agents and the variance in their beliefs, the effect on the optimal portfolio from the expected beliefs is likely to dominate. It is only in the cases where consultation only has an effect on the variance in beliefs should we expect the expected optimal portfolio to be significantly affected by the variance in beliefs. The relationship between the expected beliefs and the optimal portfolio is positive so as the agents becomes more certain that the state is state one, the more the optimal portfolio will be weighted towards the risky asset. There is likely to be a positive relationship between changes in the expected variance of beliefs and the variance of the optimal portfolio, that is, increases (decreases) in the variance in beliefs will likely lead to increases (decreases) in the variance of the optimal portfolio.

The function mapping the optimal weight allocated to the risky asset to the expected returns is linear in the optimal portfolio so we can say that the expected returns will be affected by the optimal portfolio, and the variance in returns will be affected by the variance in the optimal portfolio. In state zero, a greater allocation to the risky asset will reduce the expected returns. The reverse will be the case in state one. Given that the agents are risk averse, we can say that increases in expected returns will increase the certainty equivalents, while increases in the variance in returns will decrease the certainty equivalents. It is possible that higher order moments might also affect the certainty equivalents, though here we abstract from such concerns.

Bringing all of this together, we can say that when consultation moves the expected beliefs towards CC beliefs this will work to increase the certainty equivalents. Likewise, when consultation reduces the variance in beliefs this will also work to increase the certainty equivalents. Therefore, cases where consultation moves expected beliefs towards CC beliefs and reduces the variance in beliefs will increase the certainty equivalents, while cases in which consultation moves beliefs away from CC beliefs and increases the variance of beliefs will reduce the certainty equivalents. When the effects of consultation on expected beliefs and the variance in beliefs have opposite effects on the certainty equivalents, it is not a simple matter to determine which effect will dominate in general.

To try to tease out the effects on certainty equivalents from consultation, a numerical example was run. In this numerical example, the returns were calculated for all three networks under both scenarios for neighborhood sizes rang-

R_f	1.02	P_1	0.7
R_h	1.10	P_0	0.3
R_l	0.95	γ	6
η	0.9		

Table 1: Parameters: The parameters used for the analysis of the certainty equivalents of the different networks for each agent were set according to those in this table.

ing from zero to 20 neighbors. Within each of these neighborhood sizes, the results were calculated for each possible number of analysts in each neighborhood and the number of those analysts receiving a correct signal under both states of nature. Certainty equivalents, expected returns and variance in returns were calculated using the probabilities of each of these possible neighborhoods given the networks and the priors of the agents. The parameters assumed in the example were as shown in Table 1.

Assumption 8 *To simplify the calculations it was assumed that all agents had the same beliefs at the start of the period.*

This assumption means that the effects on certainty equivalents presented here are likely to be an upper bound given the parameters, as agents will have a minimum variance in their beliefs and hence in their returns. The expected beliefs and variance in beliefs for each of the agents were calculated for the two agent types under the two scenarios for consultation on the three different networks with neighborhoods sizes varying from zero (no consultation) to 20 neighbors.

The results of the certainty equivalents, expected returns and variance in returns on each of the three networks and under the two scenarios for different sizes of neighborhoods are shown in Figure 2 for the analysts and in Figure 3 for the uninformed agents. The general results are robust to changes in the parameters.

For analysts, as shown in Figure 2, consulting will increase the certainty equivalents of their investment on hub networks and on non-random wheel networks, under both scenarios when their neighborhoods contain up to 20 neighbors and on random wheel networks if they can discover the type of their neighbors. If analysts cannot discover the types of agent that their neighbors are, then the certainty equivalents of their investment will fall when they con-

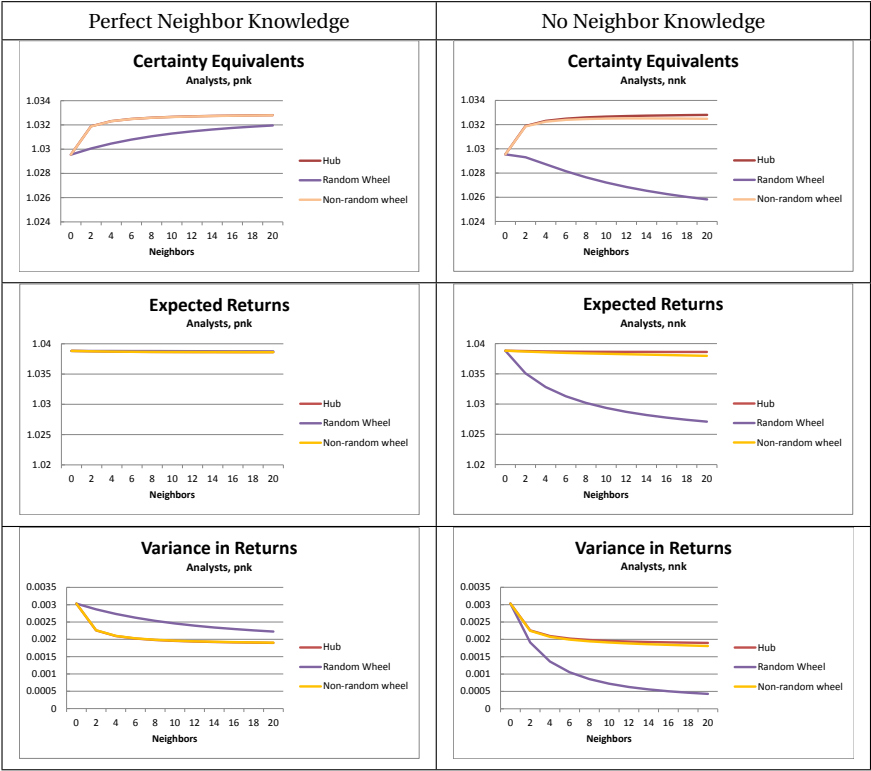


Figure 2: Analysts: The figure shows the effects of consultation among neighborhoods of various sizes for analysts in terms of certainty equivalents, expected returns and the variance in returns under the scenarios when the agent can tell perfectly the type of their neighbors (Perfect neighbor knowledge) and when they cannot tell the type of their neighbors at all (No neighbor knowledge).

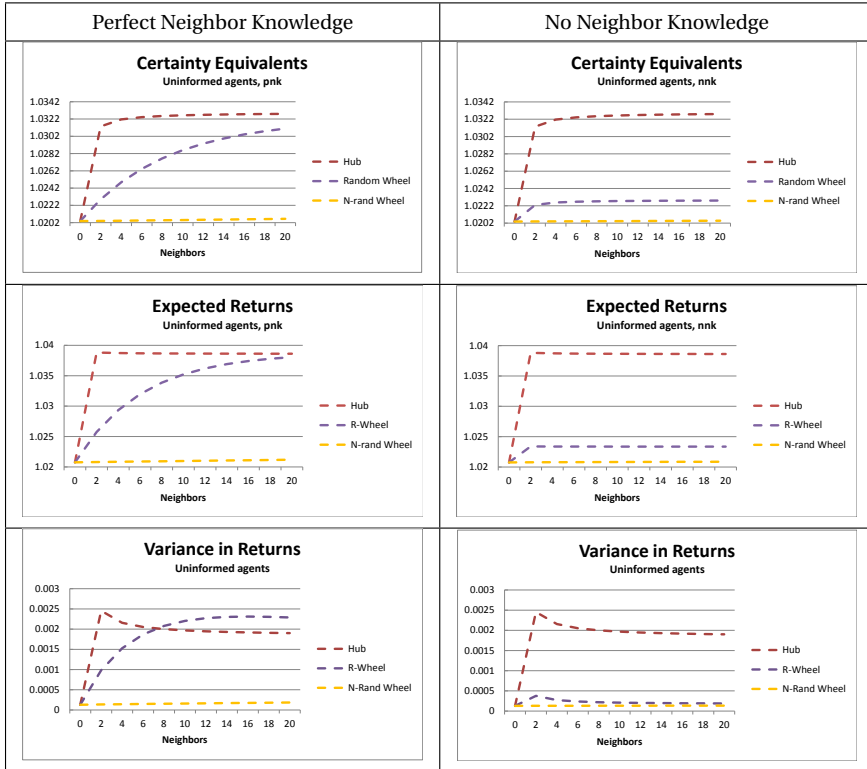


Figure 3: Uninformed Agents: The figure shows the effects of consultation among neighborhoods of various sizes for uninformed agents in terms of certainty equivalents, expected returns and the variance in returns under the scenarios when the agent can tell perfectly the type of their neighbors (Perfect neighbor knowledge) and when they cannot tell the type of their neighbors at all (No neighbor knowledge).

sult on a random wheel network. Consulting on a hub network gives the greatest increase but the difference between the certainty equivalents on a hub network compared with those on a non-random wheel network are extremely small. While the difference is slightly larger when the analyst cannot tell which type of agent her neighbors are, even then the difference is very small.

For analysts, the source of the improvement in certainty equivalents (when they improve) is entirely due to a fall in the variance of returns. On no networks under either scenario do the expected returns in investment improve with consulting. On most networks expected returns fall marginally, though on a random wheel network under the scenario of no neighbor knowledge the expected returns on investment fall significantly. This fall is due to analysts giving weight to many uninformed agents in their beliefs and so their expected returns worsen significantly. This fall in expected returns is the cause of the fall in certainty equivalents on random wheel networks with non neighbor knowledge and indicates that analysts would be better off not consulting at all than consulting on a random-wheel networks if they cannot determine how knowledgeable their neighbors are.

On each of the networks under both scenarios the variance in returns falls for analysts. Similar falls are recorded on the hub and non-random wheel networks under both of the scenarios. On a random wheel network when there is perfect neighbor knowledge the falls in variance are less than those on the other networks. When there is no neighbor knowledge, the falls in variance of returns are far greater on the random wheel network than on the other networks. This fall is due to the analysts including more uninformed agents (who have lower variance in beliefs than analysts) in their post-consultation beliefs.

Figure 3 shows that for uninformed agents, consulting increases the certainty equivalents of their investment on each of the three networks under both of the scenarios. However, these gains are comparatively very small on the non-random wheel network under both scenarios. On a hub network, almost all of the gains are made on a small neighborhood with much smaller marginal gains to be made by increasing their neighborhood beyond this small amount. This suggests that uninformed agents can receive almost the full benefits of consulting by only reading a few finance blogs. The greatest difference between the two scenarios in terms of certainty equivalents is found on the random wheel network. While the uninformed agents do gain on the random wheel networks when they cannot tell the type of their neighbors, far greater gains are made

when they are able to determine whether their neighbors are informed or not. Even with perfect knowledge of their neighbors the gains on a random wheel network only approach those of the hub network when the neighborhood becomes quite large.

It is clear from Figure 3 that, for uninformed agents, the main source of the gains in certainty equivalents is through an increase in expected returns. On each of the networks under both of the scenarios the expected returns of the uninformed agents increases in similar proportions to that of the certainty equivalents. On hub networks the expected returns increase significantly with the smallest amount of consultation. On random wheel networks, the expected returns approach those of the hub networks but only on larger neighborhoods under the perfect neighbor knowledge scenario. Under no neighbor knowledge the gains are far more modest. There are few gains on the non-random wheel network under either of the scenarios.

The variance in returns for uninformed agents increases on each of the networks under both scenarios when they consult compared with not consulting. This would normally tend to reduce the certainty equivalents of their investments. However, the increase in the certainty equivalents resulting from the improvements in expected returns outweighs any decrease from the increased variance.

4.3.1 Summary

For analysts, the gains from consulting derive exclusively from a reduction in the variance of their returns with no gains in their expected returns. For analysts, consulting on hub networks or on non-random wheel networks provide almost equivalent gains whether or not they can tell the type of agent that is their neighbor, given that they are likely to be among other financially-knowledgeable people on either platform. Nonetheless, consulting on random wheel networks can also be beneficial for analysts, provided they can tell which of their neighbors are financially-informed and which are not and they have a large neighborhood.

On the other hand, for uninformed agents, the benefits of consulting derive almost exclusively from improvement in their expected returns. Of the network types that have been examined here, the greatest gains are to be made through consulting hub networks, and they only need consult comparatively few neighbors to receive almost all of the benefit. However, similar gains can be made on random wheel networks if they can tell which of their neighbors are ana-

lysts and which are uninformed, though they need to have large neighborhoods to achieve these gains. For uninformed agents there appears to be little to be gained from consulting on non-random wheel networks as they are likely to be too far away from financially-informed neighbors for this to be of great use.

5 Conclusion

This study examines where the benefits from consultation accrue to financially informed and uninformed agents. It compares three networks, a hub network, a random wheel network and a nonrandom wheel network. These networks are specifically designed to represent real-life social networks on which people might consult to gain financial knowledge; specifically, finance-related blogs, Facebook and LinkedIn. Informed and uninformed agents could share their beliefs with other agents to whom they are socially connected and incorporate their neighbors' views into their own. Using these updated beliefs about the state of nature, they then choose a portfolio consisting of a risk-free asset and a risky asset. Certainty equivalents, expected returns and the expected variance in returns are calculated for each of these agents on each of these networks under two scenarios reflecting how well an agent could determine how knowledgeable their neighbors are.

The study suggests that for informed agents the best networks on which to consult are hub networks where they are linked to other informed agents, such as finance-related blogs, or on networks where they are closely connected to other financially knowledgeable people, such as LinkedIn. However, if they can determine how financially knowledgeable their neighbors are, Facebook can yield similar gains, though the size of the neighborhood required for such gains to be realized is significantly larger than for finance blogs or LinkedIn. This might not be a great problem given the size of many neighborhoods on Facebook. The gains to be made for financially informed agents is exclusively through a reduction in the variance of their returns.

For uninformed agents, a hub network is also the best network to consult. However, in contrast to informed agents, LinkedIn is the worst network to consult, because most of the uninformed agents are simply too far away from the informed agents for it to be a helpful source of information. Facebook is better than LinkedIn for uninformed agents if they can determine the type of their neighbors. The gains to be made for uninformed agents are largely through

improvement in their expected returns, with only modest gains to be made through a reduction in the variance of their returns.

References

- Acemoglu, D., Dahleh, M. A., Lobel, I. & Ozdaglar, A. (2011), 'Bayesian learning in social networks', *The Review of Economic Studies* **78**(4), 1201–1236.
- Charness, G., Karni, E. & Levin, D. (2007), 'Individual and group decision making under risk: An experimental study of bayesian updating and violations of first-order stochastic dominance', *Journal of Risk and Uncertainty* **35**(2), 129–148.
- Charness, G. & Levin, D. (2005), 'When optimal choices feel wrong: A laboratory study of bayesian updating, complexity, and affect', *The American Economic Review* **95**(4), 1300–1309.
- Cohen, L., Frazzini, A. & Malloy, C. (2008), 'The small world of investing: Board connections and mutual fund returns', *Journal of Political Economy* **116**(5), 951–979.
- DeGroot, M. H. (1974), 'Reaching a consensus', *Journal of the American Statistical Association* **69**(345), 118–121.
- El-Gamal, M. A. & Grether, D. M. (1995), 'Are people bayesian? uncovering behavioral strategies', *Journal of the American Statistical Association* **90**(432), 1137–1145.
- Gale, D. & Kariv, S. (2003), 'Bayesian learning in social networks', *Games and Economic Behaviour* **45**(2), 329–346.
- Grether, D. M. (1992), 'Testing bayes rule and the representativeness heuristic: Some experimental evidence', *Journal of Economic Behavior & Organization* **17**(1), 31 – 57.
- Holt, C. A. & Smith, A. M. (2009), 'An update on bayesian updating', *Journal of Economic Behavior & Organization* **69**(2), 125 – 134.
- Hong, H., Kubik, J. D. & Stein, J. C. (2004), 'Social interaction and stock-market participation', *The Journal of Finance* **59**(1), 137–163.
- Hong, H., Kubik, J. D. & Stein, J. C. (2005), 'Thy neighbor's portfolio: Word-of-mouth effects in the holdings and trades of money managers', *The Journal of Finance* **60**(6), 2801–2824.

- Ivkovich, Z. & Weisbenner, S. (2007), Information diffusion effects in individual investors' common stock purchases covet thy neighbors' investment choices, NBER Working Papers 13201, National Bureau of Economic Research, Inc.
- Jackson, M. O. (2008), *Social and Economic Networks*, Princeton University Press, Princeton, New Jersey.
- Jackson, M. O. & Golub, B. (2007), Naïve Learning in Social Networks: Convergence, Influence and Wisdom of Crowds, Working Papers 2007.64, Fondazione Eni Enrico Mattei.
- Jadbabaie, A., Molavi, P., Sandroni, A. & Tahbaz-Salehi, A. (2012), 'Non-bayesian social learning', *Games and Economic Behavior* **76**(1), 210 – 225.
- Jadbabaie, A., Molavi, P., Sandroni, A. & Tahbaz-Salehi, A. (2013), Information heterogeneity and the speed of learning in social networks, Research Paper 13-28, Columbia Business School.
- Lin, H.-F. (2007), 'Effects of extrinsic and intrinsic motivation on employee knowledge sharing intentions', *J. Information Science* **33**(2), 135–149.
- Ozsoylev, H. N. & Walden, J. (2011), 'Asset pricing in large information networks', *Journal of Economic Theory* **146**(6), 2252 – 2280.
- Tenenbaum, J. B., Griffiths, T. L. & Kemp, C. (2006), 'Theory-based bayesian models of inductive learning and reasoning', *Trends in Cognitive Sciences* **10**(7), 309 – 318. Special issue: Probabilistic models of cognition.
- Turner, M. E. & Pratkanis, A. R. (1998), 'Twenty-five years of groupthink theory and research: Lessons from the evaluation of a theory', *Organizational Behavior and Human Decision Processes* **73**(2-3), 105 – 115.
- Vives, X. (2007), Information sharing: Economics and anti-trust, Occasional Paper 07/11, IESE Business School.
- Wu, W.-p. (2008), 'Dimensions of social capital and firm competitiveness improvement: The mediating role of information sharing', *Journal of Management Studies* **45**(1), 122–146.

A Appendix

A.1 Proof of Proposition 1: Beliefs that are closer to correctly certain (CC) beliefs lead to higher expected returns

Assumption 9 $P_1 R_h + (1 - P_1) R_l > R_f > P_0 R_h + (1 - P_0) R_l$

If the state is state one, CC beliefs are $\pi = 1$, which is the maximal value for beliefs. Therefore, in state one, beliefs that are closer to CC beliefs will have a higher value. The optimum portfolio a^* is increasing in beliefs, so in state one beliefs that are closer to CC beliefs will have a higher value for a^* . The expected returns on a portfolio in state one are:

$$E(R) = P_1 [a^* R_h + (1 - a^*) R_f] + (1 - P_1) [a^* R_l + (1 - a^*) R_f]$$

The derivative of the expected return with respect to a^* is:

$$\frac{\partial E(R)}{\partial a^*} = P_1 R_h + (1 - P_1) R_l - R_f$$

This derivative is positive by assumption. Therefore, in state one, the closer beliefs are to CC beliefs the greater a^* will be and the greater the expected returns will be.

If the state is state 0, CC beliefs are $\pi = 0$, which is the minimal value for beliefs. Therefore, in state zero, beliefs that are closer to CC beliefs will have a lower value. The optimum portfolio a^* is increasing in beliefs, so in state zero beliefs that are closer to CC beliefs will have a lower value for a^* . The expected returns on a portfolio in state zero are:

$$E(R) = P_0 [a^* R_h + (1 - a^*) R_f] + (1 - P_0) [a^* R_l + (1 - a^*) R_f]$$

The derivative of the expected return with respect to a^* is:

$$\frac{\partial E(R)}{\partial a^*} = P_0 R_h + (1 - P_0) R_l - R_f$$

This derivative is negative by assumption. Therefore, in state zero, the closer beliefs are to CC beliefs the lower a^* will be and the greater the expected returns will be.

Therefore, in either state the closer beliefs are to CC beliefs the greater the expected returns will be.

A.2 Proof of Proposition 2: Analysts' post-signal beliefs will be closer to CC beliefs than their prior beliefs regardless of the state of nature or prior beliefs.

If an analyst receives signal $\theta^m = 1$ their belief will be:

$$\pi_2 = \frac{\eta\pi_1}{\eta\pi_1 + (1-\eta)(1-\pi_1)}$$

If an analyst receives signal $\theta^m = 0$ their belief will be:

$$\pi_2 = \frac{(1-\eta)\pi_1}{(1-\eta)\pi_1 + \eta(1-\pi_1)}$$

$$\begin{aligned} E(\pi|\theta = 1) - \pi_1 &= \frac{\pi_1(1-\pi_1)^2(1-2\eta)^2}{[1 - (\pi_1 + \eta - 2\eta\pi_1)](\pi_1 + \eta - 2\eta\pi_1)} \\ &> 0 \quad \text{If } \pi_1 \neq 0, 1 \text{ and } \eta \neq 0.5 \end{aligned}$$

$$\begin{aligned} \pi_1 - E(\pi|\theta = 0) &= \frac{\pi_1^2(1-\pi_1)(1-2\eta)^2}{[1 - (\pi_1 + \eta - 2\eta\pi_1)](\pi_1 + \eta - 2\eta\pi_1)} \\ &> 0 \quad \text{If } \pi_1 \neq 0, 1 \text{ and } \eta \neq 0.5 \end{aligned}$$

Therefore in either state of nature the expected post-signal beliefs of the analysts will be closer to CC beliefs than the prior beliefs, for all prior beliefs. Furthermore, regardless of the state of nature the difference between the expected post-signal beliefs and the prior beliefs will be greater the greater is η .

A.3 Derivations of expected beliefs and variance in beliefs under no neighbor knowledge

Under no neighbor knowledge:

$$w_{ij} = v_i \quad \forall j, j \neq i$$

where v_i is the probability that an agent in i 's neighborhood is an analyst.

$$w_{ii} = \begin{cases} 1, & \text{if the agent is an analyst,} \\ 0, & \text{if the agent is an uninformed agent,} \end{cases}$$

Under this scenario, an analyst's beliefs after consultation (given ψ_i) will be:

$$\pi_{3i} = \frac{\sum_{j=1}^{\psi_i} v_i \pi_{2j} + \sum_{j=\psi_i+1}^k v_i \pi_{2j} + \pi_{2i}}{k v_i + 1}$$

Under Assumption 5, that agents' pre-consultation beliefs (π_2) and the number of analysts in their neighborhood (ψ_i) are independent of each other, taking expectations over π_2 and ψ_i gives:

$$\begin{aligned} E[\pi_{3i}] &= \frac{E[\psi_i] v_i \pi_a + (k - E[\psi_i]) v_i \pi_u + \pi_a}{k v_i + 1} \\ &= \frac{(k v_i^2 + 1) \pi_a + k v_i (1 - v_i) \pi_u}{k v_i + 1} \end{aligned}$$

An uninformed agent's beliefs after consultation (given ψ_i) will be:

$$\pi_{3i} = \frac{\sum_{j=1}^{\psi_i} v_i \pi_{2j} + \sum_{j=\psi_i+1}^k v_i \pi_{2j}}{k v_i}$$

Given Assumption 5, taking expectations over π_2 and ψ_i gives:

$$\begin{aligned} E[\pi_{3i}] &= \frac{E[\psi_i] \pi_a + (k - E[\psi_i]) \pi_u}{k} \\ &= \frac{k v_i \pi_a + (k - k v_i) \pi_u}{k} \\ &= v_i \pi_a + (1 - v_i) \pi_u \end{aligned}$$

Turning now to the variance in agent's beliefs, under no neighbor knowledge, the variance of an analyst's beliefs post-consultation will be:

$$\begin{aligned} Var(\pi_{3i}) &= \sum_{j=1}^{\psi_i} \left(\frac{v_i}{k v_i + 1} \right)^2 Var(\pi_2^a) + \sum_{j=\psi_i+1}^k \left(\frac{v_i}{k v_i + 1} \right)^2 Var(\pi_2^u) \\ &\quad + \left(\frac{1}{k v_i + 1} \right)^2 Var(\pi_2^a) \end{aligned}$$

Taking expectations over ψ_i gives:

$$\begin{aligned} E[Var(\pi_{3i})] &= \frac{E[\psi_i]v_i^2 + 1}{(kv_i + 1)^2} Var(\pi_2^a) + \frac{(k - E[\psi_i])v_i^2}{(kv_i + 1)^2} Var(\pi_2^u) \\ &= \frac{kv_i^3 + 1}{(kv_i + 1)^2} Var(\pi_2^a) + \frac{kv_i^2(1 - v_i)}{(kv_i + 1)^2} Var(\pi_2^u) \end{aligned}$$

The variance of an uninformed agent's beliefs will be:

$$Var(\pi_{3i}) = \sum_{j=1}^{\psi_i} \left(\frac{v_i}{kv_i} \right)^2 Var(\pi_2^a) + \sum_{j=\psi_i+1}^k \left(\frac{v_i}{kv_i} \right)^2 Var(\pi_2^u)$$

Taking expectations over ψ_i gives:

$$\begin{aligned} E[Var(\pi_{3i})] &= \frac{E[\psi_i]}{k^2} Var(\pi_2^a) + \frac{k - E[\psi_i]}{k^2} Var(\pi_2^u) \\ &= \frac{v_i}{k} Var(\pi_2^a) + \frac{(1 - v_i)}{k} Var(\pi_2^u) \end{aligned}$$

Therefore, under the scenario of no knowledge of neighbors' type, analysts expect that their beliefs will move away from CC beliefs as their neighborhood increases. For uninformed agents, however, they expect that their beliefs will be closer to CC beliefs after consultation compared with before consultation, but the size of their neighborhood will not affect how close to CC beliefs they are as long as the probability that a neighbor is an analyst, v_i , is unaffected by the size of the neighborhood.

Analysts expect that the variance in their beliefs will decrease with the size of their neighborhood (as long as the variance of analysts' beliefs is greater than the variance of uninformed agents beliefs). Uninformed agents expect that the variance in their beliefs will decrease with the size of their neighborhood as long as the probability of a neighbor being an analyst is unaffected by the size of the neighborhood. However, the variance of uninformed agents beliefs might increase with consultation compared with no consultation if the variance of analysts' beliefs is greater than that of uninformed agents.

A.4 Expected variance of an analyst's post-consultation beliefs with perfect neighbor knowledge on a random-wheel network.

The expected variance of an analyst's post-consultation beliefs on a random wheel network will be:

$$E[Var(\pi_{3i})] = E\left[\frac{1}{\psi_i + 1}\right] Var(\pi_2^a)$$

On a random-wheel network it is assumed that ψ_i follows a binomial distribution with trials k and probability of success v_i . The expected value of $\frac{1}{X+1}$ when X is a binomial variable with n trials and p chance of success is:

$$\begin{aligned} E\left[\frac{1}{X+1}\right] &= \sum_{k=0}^n \frac{1}{k+1} \binom{n}{k} p^k (1-p)^{n-k} \\ &= \frac{1}{(n+1)p} \cdot \sum_{k=0}^n \binom{n+1}{k+1} \cdot p^{k+1} (1-p)^{n-k} \\ &= \frac{1}{(n+1)p} \cdot (1 - (1-p)^{n+1}) \end{aligned}$$

Therefore:

$$E\left[\frac{1}{\psi_i + 1}\right] = \frac{1}{(k+1)v_i} \cdot (1 - (1-v_i)^{k+1})$$

and so

$$E[Var(\pi_{3i})] = \frac{1}{(k+1)v_i} \cdot (1 - (1-v_i)^{k+1}) Var(\pi_2^a)$$

with $v_i = \frac{x}{x+y}$

$$E[Var(\pi_{3i})] = \frac{x+y}{(k+1)x} \cdot \left[1 - \left(\frac{y}{x+y}\right)^{k+1}\right] Var(\pi_2^a)$$

A.5 An analyst's expected post-consultation beliefs with no neighbor knowledge on a random-wheel network.

When agents have **no neighbor knowledge** on a random wheel network, the expected beliefs of an analyst will be:

$$E[\pi_{3i}] = \frac{(kv_i^2 + 1)\pi_a + kv_i(1 - v_i)\pi_u}{kv_i + 1}$$

with $v_i = \frac{x}{x+y}$

$$\begin{aligned} &= \frac{\left(k\left(\frac{x}{x+y}\right)^2 + 1\right)\pi_a + k\frac{x}{x+y}\left(1 - \frac{x}{x+y}\right)\pi_u}{k\frac{x}{x+y} + 1} \\ &= \frac{\left(\frac{kx^2}{(x+y)^2} + 1\right)\pi_a + k\frac{x}{x+y}\left(\frac{y}{x+y}\right)\pi_u}{k\frac{x}{x+y} + 1} \\ &= \frac{\left(\frac{kx^2 + (x+y)^2}{(x+y)^2}\right)\pi_a + \frac{kxy}{(x+y)^2}\pi_u}{\frac{kx+x+y}{x+y}} \\ &= \left(\frac{[kx^2 + (x+y)^2]\pi_a + kxy\pi_u}{(x+y)^2}\right) \cdot \frac{x+y}{kx+x+y} \\ &= \frac{[kx^2 + (x+y)^2]\pi_a + kxy\pi_u}{(x+y)(kx+x+y)} \\ &= \frac{[kx^2 + (x+y)^2]\pi_a + kxy\pi_u}{kx^2 + (x+y)^2 + kxy} \end{aligned}$$

A.6 The expected variance in an analyst's post-consultation beliefs with no neighbor knowledge on a random-wheel network.

On a random wheel network the variance of an analyst's beliefs will be:

$$E[Var(\pi_{3i})] = \frac{kv_i^3 + 1}{(kv_i + 1)^2} Var(\pi_2^a) + \frac{kv_i^2(1 - v_i)}{(kv_i + 1)^2} Var(\pi_2^u)$$

with $v_i = \frac{x}{x+y}$

$$\begin{aligned}
&= \frac{k \left(\frac{x}{x+y} \right)^3 + 1}{\left(k \frac{x}{x+y} + 1 \right)^2} \text{Var}(\pi_2^a) + \frac{k \left(\frac{x}{x+y} \right)^2 \left(1 - \frac{x}{x+y} \right)}{\left(k \frac{x}{x+y} + 1 \right)^2} \text{Var}(\pi_2^u) \\
&= \frac{\left(\frac{kx^3 + (x+y)^3}{(x+y)^3} \right)}{\frac{(kx+x+y)^2}{(x+y)^2}} \text{Var}(\pi_2^a) + \frac{\left(\frac{kx^2}{(x+y)^2} \right) \left(\frac{y}{x+y} \right)}{\frac{(kx+x+y)^2}{(x+y)^2}} \text{Var}(\pi_2^u) \\
&= \left(\frac{(kx^3 + (x+y)^3)(x+y)^2}{(x+y)^3(kx+x+y)^2} \right) \text{Var}(\pi_2^a) + \left(\frac{kx^2 y (x+y)^2}{(x+y)^3(kx+x+y)^2} \right) \text{Var}(\pi_2^u) \\
&= \frac{kx^3 + (x+y)^3}{(x+y)(kx+x+y)^2} \text{Var}(\pi_2^a) + \frac{kx^2 y}{(x+y)(kx+x+y)^2} \text{Var}(\pi_2^u)
\end{aligned}$$

A.7 The expected post-consultation variance of uninformed agents on random wheel networks with no neighbor knowledge.

The expected post-consultation variance of uninformed agents on networks with no neighbor knowledge will be:

$$E[\text{Var}(\pi_{3i})] = \frac{v_i}{k} \text{Var}(\pi_2^a) + \frac{(1-v_i)}{k} \text{Var}(\pi_2^u)$$

with $v_i = \frac{x}{x+y}$

$$= \frac{x}{k(x+y)} \text{Var}(\pi_2^a) + \frac{y}{k(x+y)} \text{Var}(\pi_2^u)$$

A.8 Derivations of the Neighborhood Mix of Agents in Non-random Networks (Propositions 3 - 4)

For Propositions 3 - 4, agent a_1 is the analyst who has an uninformed agent as their immediate neighbor in the anti-clockwise direction. The analyst immediately next to a_1 in the clockwise direction is agent a_2 , and so on until agent a_x is reached. Similarly, for uninformed agents, agent b_1 is the uninformed agent whose immediate neighbor in the anti-clockwise direction is an analyst. Their neighbor in the clockwise direction is agent b_2 and so on until agent b_y is reached. In all of these derivations it is assumed that the size of the neighborhood is small. More specifically it is assumed that $k \leq \min(2x, 2y)$, where

k is the number of agents in each neighborhood other than the agent herself, x is the number of analysts in the network, and y is the number of uninformed agents in the network.

A.8.1 Proof of Proposition 3

The number of analysts in each analyst's neighborhood in an anti-clockwise direction are:

$$\begin{aligned}
 a_1 &= 0 \\
 a_2 &= 1 \\
 a_3 &= 2 \\
 &\vdots \\
 a_{\frac{k}{2}-1} &= \frac{k}{2} - 2 \\
 a_{\frac{k}{2}} &= \frac{k}{2} - 1 \\
 a_{\frac{k}{2}+1} &= \frac{k}{2} \\
 a_{\frac{k}{2}+2} &= \frac{k}{2} \\
 &\vdots \\
 a_{x-1} &= \frac{k}{2} \\
 a_x &= \frac{k}{2}
 \end{aligned}$$

The sum of these agents' neighbors is:

$$\begin{aligned}
 Sum &= 0 + 1 + 2 + \dots + \underbrace{\left(\frac{k}{2} - 2\right) + \left(\frac{k}{2} - 1\right) + \left(\frac{k}{2}\right)}_{\left(\frac{k}{2} + 1\right) \text{ agents}} + \underbrace{\left(\frac{k}{2}\right) + \dots + \left(\frac{k}{2}\right)}_{\left(x - \left(\frac{k}{2} + 1\right)\right) \text{ agents}} \\
 &= \sum_{i=1}^{\frac{k}{2}} i + \frac{k}{2} \left(x - \left(\frac{k}{2} + 1\right)\right) \\
 &= \left(\frac{k}{2} + 1\right) \left(\frac{k}{4}\right) + \frac{k}{2} \left(x - \frac{k}{2} - 1\right) \\
 &= \frac{k}{2} \left(x - \frac{k}{4} - \frac{1}{2}\right)
 \end{aligned}$$

There will be an equal number of analysts in analysts' neighborhoods going in a clockwise direction from the agents, therefore the total number of analysts in analysts' neighborhoods (not including the analyst herself) will be:

$$Total = k \left(x - \frac{k}{4} - \frac{1}{2}\right)$$

To get the expected number of analysts in any analyst's neighborhood this total must be divided by the number of analysts (namely by x) to give:

$$E(\psi_a) = \frac{k}{x} \left(x - \frac{k}{4} - \frac{1}{2}\right)$$

A.8.2 Proof of Proposition 4

The number of analysts in each uninformed agent's neighborhood in an anti-clockwise direction are:

$$\begin{aligned}
 b_1 &= \frac{k}{2} \\
 b_2 &= \frac{k}{2} - 1 \\
 b_3 &= \frac{k}{2} - 2 \\
 &\vdots \\
 b_{\frac{k}{2}-1} &= 2 \\
 b_{\frac{k}{2}} &= 1 \\
 b_{\frac{k}{2}+1} &= 0 \\
 b_{\frac{k}{2}+2} &= 0 \\
 &\vdots \\
 b_y &= 0
 \end{aligned}$$

The sum of these agents' neighbors is:

$$\begin{aligned}
 Sum &= \underbrace{\left(\frac{k}{2}\right) + \left(\frac{k}{2} - 1\right) + \left(\frac{k}{2} - 2\right) + \dots + 2 + 1}_{\left(\frac{k}{2}\right) \text{ agents}} + \underbrace{0 + \dots + 0}_{\left(y - \frac{k}{2}\right) \text{ agents}} \\
 &= \sum_{i=1}^{\frac{k}{2}} i \\
 &= \frac{\frac{k}{2} \left(\frac{k}{2} + 1\right)}{2} \\
 &= \frac{k(k+2)}{8}
 \end{aligned}$$

There will be an equal number of analysts in uninformed agents' neighborhoods going in a clockwise direction from the agents. Therefore, the total number of analysts in uninformed agents' neighborhoods will be:

$$Total = \frac{k(k+2)}{4}$$

To get the expected number of analysts in any uninformed agents' neighborhood this total must be divided by the number of uninformed agents (namely by y) to give:

$$E(\psi_u) = \frac{k(k+2)}{4y}$$

A.9 The expected variance of analysts' post-consultation beliefs on a non-random wheel network under perfect neighbor knowledge

The expected variance of analysts' post-consultation beliefs under perfect neighbor knowledge will be:

$$E[Var(\pi_{3i})] = E\left[\frac{1}{\psi_i + 1}\right] Var(\pi_2^a)$$

On a non-random wheel network, there will be $x - k$ analysts that have $k + 1$ analysts in their neighborhood (including themselves), 2 analysts with k analysts in their neighborhood, 2 analysts with $k - 1$ analysts in their neighborhoods, and so on until there are 2 analysts with $\frac{k}{2} + 1$ analysts in their neighborhood.

Given that each position is equally likely, the probability that an analyst consulting on a non-random wheel network has $k + 1$ analysts in their neighborhood is $\frac{x-k}{x}$, and the probability of any other possible number of analysts in their neighborhood is $\frac{2}{x}$. Therefore,

$$\begin{aligned} E\left[\frac{1}{\psi_i + 1}\right] &= \frac{x-k}{x} \left[\frac{1}{k+1}\right] + \frac{2}{x} \left[\frac{1}{\frac{k}{2}+1} + \frac{1}{\frac{k}{2}+2} + \dots + \frac{1}{k}\right] \\ &= \frac{x-k}{x} \left[\frac{1}{k+1}\right] + \frac{2}{x} \sum_{i=1}^{\frac{k}{2}} \left(\frac{1}{\frac{k}{2}+i}\right) \end{aligned}$$

and so:

$$E[Var(\pi_{3i})] = \left[\frac{x-k}{x(k+1)} + \frac{2}{x} \sum_{i=1}^{\frac{k}{2}} \left(\frac{1}{\frac{k}{2} + i} \right) \right] Var(\pi_2^a)$$

A.10 The expected variance of uninformed agents' post-consultation beliefs on a non-random wheel network under perfect neighbor knowledge

The expected variance of uninformed agents' post-consultation beliefs under perfect neighbor knowledge will be:

$$E[Var(\pi_{3i})] = P(\psi_i = 0) \frac{1}{k+1} Var(\pi_2^u) + P(\psi_i \geq 1) E \left[\frac{1}{\psi_i} | \psi_i \geq 1 \right] Var(\pi_2^a)$$

On a non-random wheel network, there will be $y-k$ uninformed agents with no analysts in their neighborhood. Given that each position is assumed to be equally likely, the probability that an uninformed agent has no analysts in their neighborhood on a non-random wheel network will be: $P(\psi_i = 0) = \frac{y-k}{y}$. Therefore, $P(\psi_i \geq 1) = \frac{k}{y}$.

Of the positions which have at least one analyst in their neighborhood, there will be 2 positions with $\psi_i = 1$, 2 positions with $\psi_i = 2, \dots, 2$ positions with $\psi_i = \frac{k}{2}$. Therefore:

$$\begin{aligned} E \left[\frac{1}{\psi_i} | \psi_i \geq 1 \right] &= \frac{2}{k} \left[1 + \frac{1}{2} + \dots + \frac{1}{\frac{k}{2}} \right] \\ &= \frac{2}{k} \left[\sum_{i=1}^{\frac{k}{2}} \frac{1}{i} \right] \end{aligned}$$

and so

$$\begin{aligned} E[Var(\pi_{3i})] &= \frac{y-k}{y} \frac{1}{k+1} Var(\pi_2^u) + \frac{k}{y} \frac{2}{k} \left[\sum_{i=1}^{\frac{k}{2}} \frac{1}{i} \right] Var(\pi_2^a) \\ &= \frac{y-k}{y(k+1)} Var(\pi_2^u) + \frac{2}{y} \sum_{i=1}^{\frac{k}{2}} \frac{1}{i} Var(\pi_2^a) \end{aligned}$$

A.11 The expected post-consultation beliefs of analysts on a non-random wheel network under no neighbor knowledge

Under **no neighbor knowledge** the expected post-consultation beliefs of analysts are:

$$E[\pi_{3i}] = \frac{(kv_i^2 + 1)\pi_a + kv_i(1 - v_i)\pi_u}{kv_i + 1}$$

With $v_i = 1 - \frac{k+2}{4x}$

$$\begin{aligned} E[\pi_{3i}] &= \frac{(k(1 - \frac{k+2}{4x})^2 + 1)\pi_a + k(1 - \frac{k+2}{4x})(1 - (1 - \frac{k+2}{4x}))\pi_u}{k(1 - \frac{k+2}{4x}) + 1} \\ &= \frac{4x}{k(4x - k - 2) + 4x} \left[\frac{k(4x - k - 2)^2 + 16x^2}{16x^2} \pi_a + \frac{k(4x - k - 2)(k + 2)}{16x^2} \pi_u \right] \\ &= \frac{k(4x - k - 2)^2 + 16x^2}{4xk(4x - k - 2) + 16x^2} \pi_a + \frac{k(k + 2)(4x - k - 2)}{4xk(4x - k - 2) + 16x^2} \pi_u \end{aligned}$$

A.12 The expected post-consultation variance of analysts on a non-random wheel network under no neighbor knowledge

Under no neighbor knowledge the expected post-consultation variance of analysts' beliefs is:

$$E[Var(\pi_{3i})] = \frac{kv_i^3 + 1}{(kv_i + 1)^2} Var(\pi_2^a) + \frac{kv_i^2(1 - v_i)}{(kv_i + 1)^2} Var(\pi_2^u)$$

With $v_i = 1 - \frac{k+2}{4x}$

$$\begin{aligned} E[Var(\pi_{3i})] &= \frac{k(1 - \frac{k+2}{4x})^3 + 1}{(k(1 - \frac{k+2}{4x}) + 1)^2} Var(\pi_2^a) + \frac{k(1 - \frac{k+2}{4x})^2(1 - (1 - \frac{k+2}{4x}))}{(k(1 - \frac{k+2}{4x}) + 1)^2} Var(\pi_2^u) \\ &= \frac{\frac{k(4x - k - 2)^3 + 64x^3}{64x^3}}{\left[\frac{k(4x - k - 2) + 4x}{4x} \right]^2} Var(\pi_2^a) + \frac{\left[\frac{k(4x - k - 2)^2}{16x^2} \cdot \frac{(k + 2)}{4x} \right]}{\left[\frac{k(4x - k - 2) + 4x}{4x} \right]^2} Var(\pi_2^u) \\ &= \frac{k(4x - k - 2)^3 + 64x^3}{4x[k(4x - k - 2) + 4x]^2} Var(\pi_2^a) + \frac{k(k + 2)(4x - k - 2)^2}{4x[k(4x - k - 2) + 4x]^2} Var(\pi_2^u) \end{aligned}$$

PAPER II

Incorporating Information into Beliefs on Networks: An Experiment

Abstract

In many applications in economic theory a model is needed to explain how people incorporate information into their beliefs about the world. Two of the main models of information incorporation are Bayes rule and DeGroot's model. This study compares these two models on networks in an experimental setting. While the DeGroot model provides the most accurate estimates of agents' behaviors, there was a large group, possibly as large as half of subjects, who sometimes use a Bayesian approach. Further, consultation was found to help overcome individual uncertainty about the state of nature but could not overcome more general uncertainty, and more consultation beyond a minimal amount did not improve beliefs. Consultation generally led to a convergence in beliefs but rarely to the point of consensus.

Keywords: Social Networks, Learning, Information, DeGroot, Bayes, Experiment.

JEL Classification: D83, D03

I would like to thank Frederik Lundtofte, Erik Wengström, Natalia Montinari, Jerker Holm, Margaret Samahita, Claes Ek, Jim Ingebretsen Carlson and Pol Campos for their many helpful comments and suggestions. Thank you also to Jenny Wate for all of her help with the administrative side of the experiment. This experiment was funded by the Crafoord Foundation, whose help is greatly appreciated.

1 Introduction

In order to gain knowledge about the world beyond their own personal experience people consult with others to whom they are socially connected; their neighbors. They then must incorporate the knowledge that they gain into their beliefs. This paper investigates how people incorporate other people's beliefs about the world into their own beliefs. It contributes to the literature by directly comparing two of the main models of information incorporation (namely the DeGroot (1974) model and one using Bayes rule) in an experimental setting. It finds that, while the DeGroot model more accurately reflects the information incorporation process, over half of subjects will deviate from the DeGroot model's predictions when Bayes rule suggest that they should. Further, this study finds that convergence to group consensus is unrelated to people's connectivity or the number of consultations people make with their neighbors, despite these both being predictions of the DeGroot and Bayes models. Nonetheless, connectivity does help people gain more accurate information about the world.

In the economic literature two main models tend to be used to model the incorporation of information into beliefs. These models are the Degroot model and one based on Bayes rule. The Degroot model is a bounded rationality model in which people assign weights to the beliefs of the people they consult, and so update their beliefs as a convex combination of their neighbors' beliefs. The benefits of this model is that it imposes a far lighter cognitive load on agents compared with Bayes rule, requires far less knowledge of how neighbors came to their beliefs, and, hence, how to incorporate agents' neighbors' beliefs into their own. Therefore, it is expected to be used in a wider variety of contexts. As such it is regularly used in the economic literature to model information incorporation (see for example: Jadbabaie et al. (2012), Golub & Jackson (2010), Eyster & Rabin (2009, 2010), DeMarzo et al. (2003), Ellison & Fudenberg (1993, 1995)).

In contrast to the DeGroot model, information incorporation according to Bayes rule requires agents to use information beyond the beliefs of their neighbors such as beliefs about the evidence on which their neighbors gained their knowledge, and whether their neighbors' beliefs are based on their neighbors' personal experience or if their neighbors gained their knowledge through others (see for example: Acemoglu et al. (2011), Holt & Smith (2009), Charness et al. (2007), Tenenbaum et al. (2006), Charness & Levin (2005), Gale & Kariv (2003), El-Gamal & Grether (1995), Grether (1992)). This requires significantly greater sophistication on the part of agents than the DeGroot model. However, having

just one agent using Bayes rule can significantly improve the knowledge of a network compared with only having agents using the DeGroot model (Mueller-Frank 2014).

The current study evaluates these models in an experimental setting. Subjects were asked to assess the likelihood that there were at least 50 white balls in a hypothetical urn of 100 balls. To assess this probability, they each received an independent sample of three balls from the urn and could then view other subjects' estimates in order to update their own estimate. The subjects could then choose to consult as many times as they liked to get more accurate beliefs about the number of white balls in the urn. In this study, the subjects could have used the DeGroot model to update their estimates, or they could have used a more Bayesian process, as they could potentially infer the evidence on which their neighbors have derived their beliefs. Overall, a DeGroot model predicted the subjects estimates better than a prediction based on Bayes rule. However, over half of the subjects, at least some of the time, deviated from the DeGroot model when Bayes rule suggested that they should. This suggests that more people are Bayesian than suggested by previous research such as Chandrasekhar et al. (2015).

In an experiment comparable to the current one, Chandrasekhar et al. (2015) directly compared the use of DeGroot and Bayes rule. Chandrasekhar et al. (2015) provided subjects with binary information based on which they took binary actions. In viewing neighbors' actions, subjects could then change their course of action if they wished. They found that the best model to describe the way agents incorporated information was one in which none of the subjects were Bayesian. They show that when Bayesian and DeGroot models have diverging predictions only 17 per cent of agents take the Bayesian option. The current study approaches the question from a different perspective to that of Chandrasekhar et al. (2015), with subjects provided with more nuanced information from which they could derive beliefs across a range, and hence the data aligned more closely with that of the original DeGroot formulation. Taken together these studies suggest that, while the DeGroot model is the best to describe how people incorporate information into their beliefs, there is a significant minority of people who will use a Bayesian approach if they can.

The DeGroot model also predicts that if agents consult with each other enough times they will eventually come to consensus, even though this is not a goal of their interactions, as long as there is an information path from at least one agent

to all other agents on the network (Berger 1981). Nonetheless, consensus only happens when the rounds of consultation go to infinity and this convergence tends to be quite slow (Golub & Jackson 2010). In reality, agents do not continue to consult indefinitely and so consensus is unlikely to result from the DeGroot model in practice.

There have been numerous experiments studying consensus formation on networks. Most of these experiments involve agents consciously attempting to reach consensus as the goal of their interactions (see for example, Judd et al. (2010), Kearns et al. (2009), Kearns et al. (2006)). In contrast, comparatively little experimental research has looked at agents reaching consensus passively on networks. That is, agents reaching consensus even though consensus is not a goal of their interaction. In one study on passive consensus formation, Mueller-Frank & Neri (2015) found consensus was hard to achieve when it was not a direct goal of interaction.

Therefore this experiment also investigates whether agents reach consensus, how often agents choose to consult, and how accurate their beliefs are after consultation. Subjects were allowed to consult as often as they wanted to, and, in general, it was found that subjects would continue to consult as long as they saw a marginal benefit in doing so, while the costs of consultation did not appear to affect their decisions to consult. Consensus was rarely reached in this experiment. Subjects consulted more often, and their beliefs also converged more often, on sparse networks than on more densely connected networks. Nonetheless, despite consulting more often on sparse networks, the accuracy of subjects' final beliefs tended to be worse than on more densely-connected networks. Furthermore, given the type of network on which subjects were consulting, consulting more times did not lead to more accurate final beliefs. This suggests that consulting a minimal number of times might give as good information as consulting more times, but consulting with a broader range of people can improve the accuracy of beliefs.

In Section 2, I will outline Degroot's model with reference to network theory. I will describe the experimental design in Section 3, and Section 4 will provide the main hypotheses to be tested. Section 5 gives details of the experiment and the subjects. Section 6 shows the main results, with subsidiary results presented in Section 7, and conclusions are provided in Section 8.

2 DeGroot's Model and Network Theory

The DeGroot (1974) model of consensus formation continues to be used by researchers to model the effects of information sharing on a social network (see for example, Jadbabaie et al. (2012), Golub & Jackson (2010), Eyster & Rabin (2009, 2010), DeMarzo et al. (2003), Ellison & Fudenberg (1993, 1995)). This study is designed to test some of the predictions that come out of the Degroot model. Specifically, the experiment will test whether or not agents updated beliefs can be considered to be a weighted average of their neighborhoods' previous beliefs. Furthermore it will test whether or not agents' beliefs converge towards a consensus, and the reasons behind how often agents choose to consult with each other.

In the DeGroot model, each agent is located as a node on a network. The set $N = \{1, \dots, n\}$ is the set of nodes. A real-valued $n \times n$ matrix g describes the relationships between each of the nodes. The relationship between two agents, i and j , is denoted by g_{ij} , where $g_{ij} \in \{0, 1\}$. $g_{ij} = 1$ denotes an information flow from agent j to agent i , namely agent i can observe something about agent j .

The network is directed so that $g_{ij} = 1$ does not imply that $g_{ji} = 1$. That is, if agent i can observe agent j this does not necessarily imply that agent j can observe agent i ¹. The set of agents that agent i can observe is called the *neighborhood* of i , denoted $N_i^d(g) = \{k \in N | g_{ik} = 1\}$. All agents in agent i 's neighborhood are agent i 's *neighbors*. The number of agents that an agent is linked to (that is, the cardinality of the agent's neighborhood) is called that agent's *degree*, and is denoted $d_i(g) = \#\{j | g_{ji} = 1\} = \#N_i^d(g)$ (Jackson 2008, p. 29). A path in a network between nodes i and j is defined as a sequence of links $i_1 i_2, i_2 i_3, \dots, i_{K-1} i_K$ such that $i_k i_{k+1} \in \{g | g_{i_k i_{k+1}} = 1\}$ for each $k \in \{1, \dots, K-1\}$, with $i_1 = i$ and $i_K = j$, and such that each node in $i_1, \dots, i_K$ is distinct (Jackson 2008, p. 23). A network in which there exists a directed path between any two nodes is called a *strongly connected* network.

Each agent assigns a weight to the information that they receive from each person in their neighborhood and, as such, their updated beliefs will be a convex combination of the beliefs of their neighbors. This weight is represented by a number, t_{ij} , where:

¹By assumption, $g_{ii} = 1$.

$$t_{ij} = \begin{cases} x_{ij} \in [0, 1] & \text{if } g_{ij} = 1 \\ 0 & \text{if } g_{ij} = 0 \end{cases}$$

and

$$\sum_{j=1}^n t_{ij} = 1 \quad \forall i$$

The information transfer matrix, T , is defined as:

$$T_{n,n} = \begin{pmatrix} t_{11} & t_{12} & \cdots & t_{1n} \\ t_{21} & t_{22} & \cdots & t_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ t_{n1} & t_{n2} & \cdots & t_{nn} \end{pmatrix}$$

Each of the rows in T will sum to 1. However, the columns will not necessarily also sum to 1. To calculate how information is passed around the network we can set up the matrix of initial beliefs about the state of nature. Namely:

$$\Pi^0 = \begin{pmatrix} \pi_1 \\ \pi_2 \\ \vdots \\ \pi_n \end{pmatrix}$$

After one round of information sharing the beliefs matrix will be:

$$\Pi^1 = T\Pi^0$$

After n rounds of information sharing the beliefs matrix will be:

$$\Pi^n = T^n\Pi^0$$

Given that all of the elements of T are between zero and one, the elements in each of the subsequent beliefs matrix will be weighted averages of the elements in the initial beliefs matrix (DeMarzo et al. 2003, p. 10).

If, and only if, there exists a positive integer, n , such that every element in at least one column of T^n is strictly positive then the beliefs will converge to a consensus (Berger 1981). This means that if there is an information path from

at least one agent to every other agent on the network (that is, it is a strongly-connected network), a consensus will form as the rounds of consultation go to infinity, and this consensus will be a weighted average of the agents' initial beliefs.

3 Experiment Design

The experiment was conducted entirely through computers, using students drawn from the Lund University School of Economics and Management. The experiment was conducted over 4 phases of 3 rounds each for a total of 12 rounds. At the start of each round the subjects were randomly assigned to networks of five people. Each of the phases involved a different treatment and the order of the presentation of the treatments changed between the sessions. The treatments were: a sparsely-connected network with a 'baseline' treatment; a highly-connected network with a 'baseline' treatment; a highly-connected network with an 'information' treatment; and a non-consultation treatment. These treatments will be explained below.

In the experiment there was a hypothetical urn with 100 balls in it. Each of these balls was either white or black. The number of white balls in the urn was drawn from a uniform distribution across all possible outcomes. Each of the subjects was asked to assess the probability that the number of white balls in the urn was greater than or equal to 50. To assist them in assessing this probability each subject received their own independent sample of three balls from the urn.

Having initially estimated the probability that the number of white balls in the urn was greater than or equal to 50, the subjects were then able to consult with their neighbors. This consultation involved viewing their neighbors' estimates of the probability that the number of white balls was greater than or equal to 50. Based on this consultation, the subjects could then update their estimate of the probability. They could then elect to consult with their neighbors again to view their neighbors' updated estimates. Consultation continued until no-one on the network wanted to consult again. The subjects' final estimates were used to determine their payoffs.

The subjects' payoffs were calculated as follows:

$$\pi_i(x_i) = \begin{cases} 100 [1 - (x_i - 1)^2] & \text{if } W \geq 50 \\ 100 [1 - x_i^2] & \text{if } W < 50. \end{cases}$$

Where $\pi_i(x_i)$ is the payoff to subject i , x_i is subject i 's final estimate of the probability that the number of white balls is greater than or equal to 50, W is the number of white balls in the urn.

Under this payoff structure subjects maximize their expected payoff by setting their final estimate equal to their actual belief about the probability that the number of white balls is at least 50. Therefore, if the subjects are risk neutral that is their best strategy.²

In this experiment, two balanced networks were tested. These networks are depicted in Figures 1 and 2. Both networks contain five subjects, with one network sparsely-connected, and the other network more highly-connected. In consultation on the sparsely-connected network, each subject was able to observe one other subject's estimates, and in turn had one other subject observe their estimates (a different subject to that which they could observe). In the highly-connected network, each subject was able to observe three other subjects.

Balanced networks were chosen because balanced networks give the most straightforward predictions in the DeGroot model. That is, with each person connected to an equal number of people and each person being as informed as everyone else, an equal weighting under the DeGroot model is an obvious prediction. If some subjects were more connected than others or had more or better information it was not always clear what a reasonable weighting scheme would be. The Bayesian updating was also more clear with balanced networks. Therefore, with these networks, there should be less noise in the data.

Two treatments were used on the highly-connected network; a 'baseline' treatment and an 'information' treatment. Under the 'baseline' treatment the subjects only received the estimates of their neighbors and were allowed to incorporate the information from their neighbors in whatever way they wanted to. Under the 'information' treatment, in addition to their neighbors' estimates, subjects also received a simple average of their neighbors' current estimates excluding their own previous estimate, and a simple average of their neighbors' current estimates including their own previous estimate. Subjects could then use this information in whatever way they wished to come to their updated estimates. Assuming that subjects wished to use the DeGroot model of consulting

²If the subjects are risk averse then their best strategy is to set their final estimate closer to 0.5 than what they actually assess the probability to be. The more risk averse they are, the closer to 0.5 they should set their final estimate.

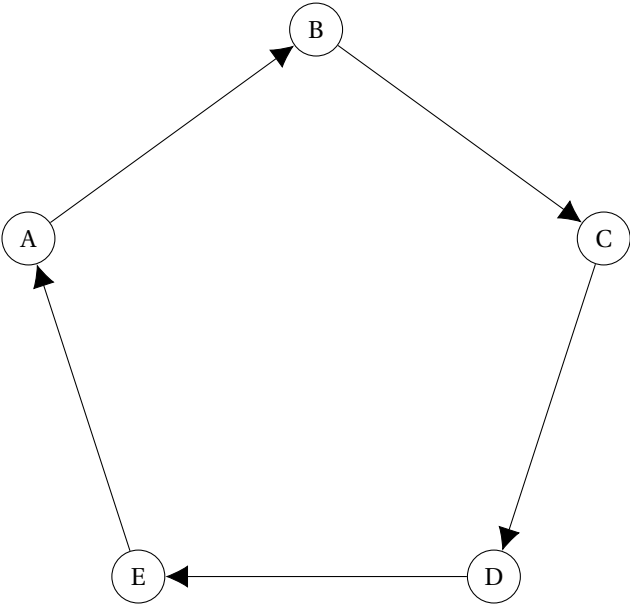


Figure 1: Sparsely-connected Consultation Network. The circles represent subjects and the arrows represent the direction of viewing. Each subject can view one of the other four people’s estimates on the network. Each subject’s estimates can be viewed by one of the other four people on the network.

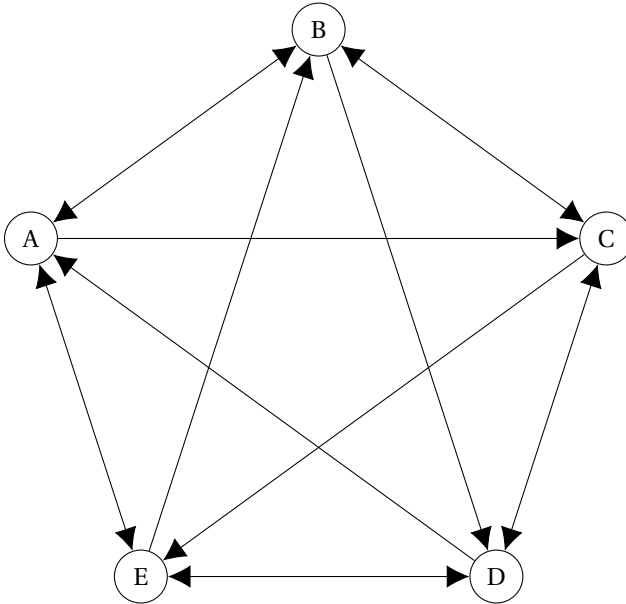


Figure 2: Highly-connected Consultation Network. The circles represent subjects and the arrows represent the direction of viewing. Each subject can view three of the other four people's estimates on the network. Each subject's estimates can be viewed by three of the other four people on the network.

with equal weights, then, under the ‘information’ treatment, subjects were not required to estimate these averages themselves and so incurred fewer costs to consulting. This treatment was, therefore, included to investigate the impact of the costs of consulting on subjects’ decisions to consult. The ‘information’ treatment was only used on the more densely-connected network.

4 Hypotheses

While Bayes rule is often used in theoretical papers, research suggests that people are more likely to use the DeGroot model in practice. Therefore, the first hypothesis that will be tested in this study is that subjects use the Degroot model.

Hypothesis 1 *Subjects use the DeGroot model*

Under the DeGroot model subjects’ post-consultation estimates will be a weighted average of their neighbors’ estimates. Therefore, regardless of how they weight their neighbors’ estimates, the subjects’ post-consultation beliefs will be within the range of estimates on their neighborhood. A simple assessment of the possibility that subjects are using the DeGroot model is, therefore, how often their post-consultation estimates fall outside the range of neighborhood estimates. The more often the estimates fall outside of the range of neighborhood estimates the less likely is it that they are using the DeGroot model.

If subjects are able to make an educated guess as to their neighbors’ samples based on their neighbors’ first round estimates, subjects can get a more accurate estimate of the probability that the number of white balls is at least 50 using Bayes’ rule compared with using an average of their neighbors’ estimates. For certain combinations of samples the estimate based on Bayes’ rule will fall outside of the range of the neighborhood’s estimates. Therefore, one predictor of whether subjects are using the De Groot model or Bayes’ rule is to measure how often the post-consultation beliefs fall outside of the range of neighborhood estimates when Bayes’ rule predicts that they should.

Given that no subject has more information than any other, under the DeGroot model an equal weighting of the neighborhood estimates would seem reasonable. Therefore, a final test can be done by comparing whether a simple average of the neighborhood estimates is a better predictor of the post-consultation estimates than that predicted by Bayes’ rule. This comparison can

be done through comparing root mean square errors (RMSE) of the two predictors.

Under the information treatment subjects will receive the averages of their neighborhood's estimates and so should be primed to use the averages. This will possibly make them less likely to use Bayes Rule in favor of DeGroot. Furthermore, having seen the information treatment in a previous phase, subjects could be primed to use the DeGroot model in future phases making them more likely to use the DeGroot model in phases subsequent to the information treatment phase.

Predictions based on this hypothesis, are that:

1. if subjects use the DeGroot model
 - (a) their updated estimates after the first round of consultation will not be outside the range of their neighborhood's estimates
 - (b) $RMSE(DeGroot) < RMSE(Bayes)$
2. if subjects use Bayes Rule
 - (a) their updated estimates after the first round of consultation will be outside the range of their neighborhood's estimates when the Bayesian update is outside the range
 - (b) $RMSE(DeGroot) > RMSE(Bayes)$
3. The information treatment should reduce the likelihood of the subjects using Bayes Rule
4. Subjects should be less likely to use Bayes rule after they have seen the information treatment

Under both the Degroot model and Bayes rule, networks that continue to consult should end up at consensus. However, many networks do not reach consensus. This prompts questions about why networks do not reach consensus. One explanation could be that agents on networks stop consulting on networks because they do not see any net benefits to reaching consensus. That is, that they believe that the costs of consulting an extra time outweigh the benefits of doing so.

Hypothesis 2 *Costs of consulting cause subjects to stop consulting before consensus is reached*

If there is a cost to consulting, people will only continue to consult as long as they expected that the marginal benefit to another round of consulting will be greater than the cost. First, I look at the marginal benefits to consulting and then I will consider the costs involved.

4.1 Marginal Benefit of Consultation

If it is assumed that the subjects give correct estimates based on their individual samples, the estimates can be thought of as independent random variables, with a mean, μ_x , and standard deviation, σ_x , determined by the number of white balls in the urn. If we define the estimate of subject i after n rounds of consultation as x_{ni} then the expected value of the estimate before consultation will be $E(x_{0i}) = \mu_x$ with variance $Var(x_{0i}) = \sigma_x^2$. If the subjects use the DeGroot model, then the estimate of subject i after one round of consultation will be:

$$x_{1i} = \sum_{j=1}^k w_j x_{0j}$$

where w_j is the weight that subject i places on subject j 's estimate (with $w_j \in [0, 1]$ and $\sum_{j=1}^k w_j = 1$) and k is the number of subjects in agent i 's neighborhood. Assuming that the weights assigned to neighbors' beliefs are independent of the neighbors' estimates, the expected value of the estimate after one round of consultation then will be:

$$\begin{aligned} E(x_{1i}) &= E\left(\sum_{j=1}^k w_j x_{0j}\right) \\ &= \sum_{j=1}^k w_j E(x_{0j}) \\ &= \mu_x \sum_{j=1}^k w_j \\ &= \mu_x \end{aligned}$$

Similar calculations can be made for subsequent rounds of consultation to show that the expected value of the estimate is unaffected by the number of rounds of consultation.

Considering now the variance of the estimates, the variance before consultation is $Var(x_{0i}) = \sigma_x^2$. After one round of consultation the variance of the estimates will be:

$$\begin{aligned} Var(x_{1i}) &= Var\left(\sum_{j=1}^k w_j x_{0j}\right) \\ &= \sum_{j=1}^k w_j^2 Var(x_{0j}) \\ &= \sigma_x^2 \sum_{j=1}^k w_j^2 \end{aligned}$$

Assuming that the agents do not change the weights that they assign to their neighbors' beliefs, after two rounds of consultation the variance of the estimates will be:

$$\begin{aligned} Var(x_{2i}) &= Var\left(\sum_{j=1}^k w_j x_{1j}\right) \\ &= \sum_{j=1}^k w_j^2 Var(x_{1j}) \\ &= \sigma_x^2 \left(\sum_{j=1}^k w_j^2\right)^2 \end{aligned}$$

Similarly, after n rounds of consultation the variance of the estimates will be:

$$Var(x_{ni}) = \sigma_x^2 \left(\sum_{j=1}^k w_j^2\right)^n$$

The conditions that $w_j \in [0, 1]$ and $\sum_{j=1}^k w_j = 1$, imply that $\frac{1}{k} \leq \sum_{j=1}^k w_j^2 \leq 1$. Therefore, the variance in the estimates will be a weakly convex decreasing function of the number of rounds of consultation. $\sum_{j=1}^k w_j^2 = 1$ if one, and only one, subject's estimate is given weight in the consultation (presumably the subject themselves). Therefore, as long as at least one other subject's estimate is given weight, and hence that there actually is consultation, the variance in the estimates will be a strictly convex decreasing function of the number of rounds of

consultation.³

Given that the payoff function is quadratic, we can determine that the expected value of the payoff is:

$$E[\pi(x_{ni})] = \begin{cases} 100 [1 - (\mu_x - 1)^2] - 100\sigma_x^2 \left(\sum_{j=1}^k w_j^2 \right)^n & \text{if } W \geq 50 \\ 100 [1 - \mu_x^2] - 100\sigma_x^2 \left(\sum_{j=1}^k w_j^2 \right)^n & \text{if } W < 50 \end{cases}$$

The expected payoff, therefore, is a concave increasing function of the rounds of consultation, regardless of the number of white balls in the urn. The marginal benefit to consulting (here defined as the first derivative of the expected payoff function with respect to the number of rounds of consultation) is:

$$\text{Marginal benefit} = -100\sigma_x^2 \left(\sum_{j=1}^k w_j^2 \right)^n \ln \left(\sum_{j=1}^k w_j^2 \right)$$

This is a convex decreasing function of the number of rounds of consultation, regardless of the number of white balls in the urn.

If $b = \sum_{j=1}^k w_j^2$ and it is reasonably assumed that b is decreasing in k (that is, $\sum_{j=1}^k w_j^2$ is lower in larger neighborhoods, which will be the case, for example, with an equal weighting of neighbors' estimates), then, $b_H < b_S < 1$, where b_H is b on more highly-connected networks and b_S is b on more sparsely-connected networks. The marginal benefit of consultation then will be greater on highly-connected networks compared with that on sparsely-connected networks, when:

$$\left(\frac{b_S}{b_H} \right)^n < \frac{\ln(b_H)}{\ln(b_S)}$$

If, as assumed, $b_H < b_S < 1$ then the marginal benefit to consultation will initially be greater in more-connected networks, but as the rounds of consultation increase the marginal benefit to consultation on highly-connected networks will fall below that on sparsely-connected networks.

Predictions based on the marginal benefits of consulting include:

1. More subjects should choose to consult on densely-connected networks than on sparsely-connected networks

³From now it is assumed that at least one other subject's estimate is given weight in the consultation, and hence $\sum_{j=1}^k w_j^2 < 1$.

2. Those that choose to consult on densely-connected networks should stop consulting earlier than on sparsely-connected networks.

4.2 Costs of consulting

In this study, in considering the costs that a subject might believe they will incur with an extra round of consultation, there appear to be three different potential sources of costs of consulting:

1. Costs of integrating multiple sources of information.

This cost implies that it is harder for people to integrate more pieces of information than fewer into their beliefs. In the context of this study, it suggests that people in the larger neighborhoods should face greater costs to consulting (other things being equal) compared with those in the smaller neighborhoods because they have more sources of information to integrate. Under the DeGroot model, in which the subjects are estimating averages of their neighbors estimates, it says that it is harder for the subjects (and, hence, involves a greater cost) to calculate an average of four estimates than two estimates.

2. Costs of integrating diverse information.

This cost implies that it is harder for people to integrate more diverse pieces of information, compared with less diverse information, into their beliefs. If the subjects are using the DeGroot model, this cost suggests that it will be more difficult for the subjects to calculate the average of estimates with a greater variance than those with a smaller variance.

3. Fatigue from having previously consulted.

As subjects consult it is conceivable that fatigue builds up and so the costs of consulting in later rounds are greater than the costs in earlier rounds.

If subjects consult until consensus is reached ($Var(x_{ni}) = 0, \forall i$), then it is clear that the costs of consulting are so low as to pose no obstacle to reaching consensus. Even if costs are effectively zero, subjects are unlikely to consult until absolute consensus is reached. Rather they are likely to consult until they believe that further consultation will not increase their return. If subjects consult until this point we can conclude that costs of consultation were not a factor in inhibiting consultation.

Under the ‘information’ treatment, each of these costs will be effectively zero (assuming that subjects wish to use the DeGroot model, with equal weighting). Therefore, a test of whether costs of consulting stop subjects from reaching consensus is if networks under the ‘information’ treatment get closer to consensus than those under the ‘baseline’ treatment. If the difference between the final variance between the two treatments is greater on networks with greater initial variance would be an indication that cost type 2 is a factor, and finally if networks consult for more rounds under the ‘information’ treatment compared with the ‘baseline’ treatment this would be an indication that fatigue is a factor in consultation.

Based on the costs of consulting, predictions for the experiment include:

1. In general, if costs are a factor:
 - (a) More subjects will choose to consult under the ‘Information’ treatment compared with the ‘Baseline’ treatment.
 - (b) Those that consult will have more consultations under the ‘Information’ treatment compared with the ‘Baseline’ treatment.
2. If cost type 1 is a factor:
 - (a) Fewer subjects will choose to consult on dense networks compared with sparse networks.
 - (b) Those that consult will have fewer consultations on dense networks compared with sparse networks.
3. If cost type 2 is a factor:
 - (a) There will be fewer consultations on networks with greater initial variance.
 - (b) Networks with greater initial variance will have greater final variance.
4. If cost type 3 is a factor:
 - (a) Fewer subjects will choose to consult later in the session.
 - (b) Those that consult will have fewer consultations later in the session.

5 Experiment Details and Subjects

The experiment was designed using the software program z-tree (Fischbacher 2007) and conducted at the Lund University School of Economics and Management (LUSEM). The experiment was conducted across nine sessions, with a total of 120 subjects. The subjects were drawn from the student body at LUSEM. In each session the subjects completed four phases of three rounds for a total of twelve rounds. The subjects were paid for three of the twelve rounds randomly selected by the computer. In each of the phases a different treatment was presented. The four treatments used were:

1. a sparsely-connected network, baseline treatment (T1),
2. a highly-connected network, baseline treatment (T2),
3. a highly-connected network, information treatment (T3),
4. a non-consulting treatment (T4).

Under the ‘baseline’ treatment subjects only received their neighbors’ most recent estimates. Under the ‘information’ treatment subjects received their neighbors’ most recent estimates as well as the averages of their neighbors’ most recent estimates and the average of their neighborhoods’ most recent estimates (that is, the average of their neighbors’ estimates and their own most recent estimate). In the non-consulting treatment the subjects were not able to consult with other subjects and their payment was based on their initial estimates. Table 1 shows the details of the sessions. In the first six sessions, the order in which the consultation treatments (T1, T2, T3) were presented was rotated so that every permutation of the order of the three treatments was used. The final phase in each session was always the non-consulting treatment. The final three sessions repeated the order of sessions 1, 4 and 6 as these were the least-attended sessions of the first six.

After the instructions for the experiment were read to the subjects and before they started the experiment the subjects completed control questions (CQ) in which they were asked to estimate the probability that there were at least 50 white balls in the urn based on each of the four possible samples. After the subjects had submitted their estimates they were each provided with the true probabilities based on the different samples. These true probabilities are shown in Table 2. They were allowed to refer to these true probabilities during the course

Table 1: Session Information

Session	Number of Subjects	Order of Treatments
1	10	T1, T2, T3, T4
2	15	T2, T3, T1, T4
3	15	T3, T2, T1, T4
4	10	T1, T3, T2, T4
5	15	T2, T1, T3, T4
6	10	T3, T1, T2, T4
7	20	T1, T2, T3, T4
8	10	T1, T3, T2, T4
9	15	T3, T1, T2, T4

Table 2: The true estimates of the probability of there being 50 or more white balls in the urn based on a single sample of three balls randomly drawn (without replacement) from the urn

Number of white balls in sample	0	1	2	3
True estimate	0.06	0.32	0.70	0.94

of the experiment. The subjects were thus induced to have correct priors about the number of white balls in the urn (that is, priors reflecting that the number of white balls in the urn were drawn from a uniform distribution across all possible outcomes).

Following the experiment, to gauge their level of risk aversion, the subjects were asked to select from eight different gambles the one gamble that they preferred. The gambles were modified versions of those in Eckel & Grossman (2002) and are presented in Table 4 in the Appendix. The gambles were in order of increasing riskiness, with gamble 1 a risk-free gamble and gamble 8 the most risky gamble. The subjects were not paid based on their choice of gamble. Following the experiment subjects also completed a questionnaire on general demographic information and a cognitive reflection task (CRT). The CRT consisted of five questions in which the intuitive answers were not the correct answers. The CRT was an extended version of the CRT in Frederick (2005) and the questions are in Table 5 in the Appendix.

Of the subjects, 69 were men and 51 women. The ages of subjects ranged from 18 to 49 with an average age of 24.9 and a median age of 24. 46 subjects

were undergraduate students, 72 were postgraduate students, with a further two not studying. 78 subjects were majoring in economics, 24 in business administration and 18 had other majors. 66 subjects had attended a Swedish high school and 54 had not.

6 Main Results

In the regressions shown in the Appendix and commented on in this section, session dummies and demographic information on the subjects' age, gender, undergraduate/postgraduate status, major, and whether they went to a Swedish high school or not were included but are not presented in the tables. The data was clustered at the subject level where appropriate.

6.1 Bayes vs DeGroot: Testing Hypothesis 1

Under the DeGroot model, subjects' updated estimates are a weighted average of their neighborhoods' beliefs. One of the consequences of this is that, if subjects are using the DeGroot model, their updated estimates should never be outside the range of their neighborhood's estimates. Therefore, if a subject updates their estimate to an estimate that is outside the range of their neighborhood, they clearly are not using the DeGroot model. However, even if their updated estimate is outside their neighborhood's range this does not mean that they are necessarily using a Bayesian updating framework, as the Bayesian update might not be outside the range of their neighborhood.

To calculate the Bayesian update it is necessary to infer the evidence on which the original estimate was made. While subjects can credibly infer the evidence on which their neighbors' initial estimates are based (particularly if their neighbors' estimates coincide with the true probabilities), inferring the evidence on which their neighbors' subsequently updated estimates are based requires significant assumptions about how their neighbors updated their estimates. Given these complexities, this analysis was only conducted on the first updates after the initial estimates were made. To make inferences about the neighbors' samples, I have treated a neighbor's estimate as uninformative if their estimate was 0.5. If the neighbor's estimate was not 0.5, I assumed that the subjects would assess their neighbor's sample to be whichever sample corresponds to the true estimate that is closest to the estimate of the neighbor's

estimate. That is, if, for example, the neighbor gave an estimate of 0.75, I have assumed that the subject assessed the neighbor's sample as being two white balls and one black ball as that corresponds to the nearest true estimate (0.699) to 0.75, regardless of what the neighbor's actual sample was.

The calculation of the Bayesian update was based on the inferred sample of the subject's neighbors and used the true probability based on the subject's own sample as their prior belief (rather than their initial estimate). The true probability was used because the initial estimate could be affected by risk aversion, while the subject would know what the true probability was as this had been given to them. Nonetheless, the results did not change significantly when the Bayesian update was assessed using the subjects' own initial estimates.

Across observations in which the subject chose to consult at least once (882 observations), the neighborhood's average estimate (including the subject's own estimate) proved to be a better predictor of subjects' updated estimates ($RMSE=0.168$) compared with the Bayesian update ($RMSE=0.255$), indicating that the DeGroot model was a better predictor of subject behavior. Nonetheless, Table 3 shows that, of the times in which a subject chose to consult at least once, 17.7 per cent of the time (156 observations) the subject's updated estimate was outside the range of their neighborhood's estimates and so they could not have been using the DeGroot model. Furthermore, when the Bayesian estimate was outside the range of the neighborhood's estimates (447 observations), 30 per cent of the time (134 observations) the subject also updated their estimate outside the range of their neighborhood's estimates.

Of the 120 subjects, 111 had, when they chose to consult, at least one occasion when the Bayesian update was outside the range of their neighborhoods' estimates. Of these 111 subjects, 65 updated their estimate outside their neighborhoods' range on at least one occasion when the Bayesian update suggested that they should. This suggests that 59 per cent of subjects who had a chance to "show that they were Bayesian", took at least one of those opportunities. This implies a far higher rate of Bayesian updating on the part of subjects than has been shown by previous research, such as that of Chandrasekhar et al. (2015) who found that only 17 per cent of subjects would choose the Bayesian approach when the predictions of the Bayesian and DeGroot models diverged.

A probit model was run to investigate the factors that caused a person to update their estimate in line with the Bayesian update. The sample for this model was all observations in which a subject chose to consult at least once

Table 3: Bayes vs Degroot Shares. The sample used for the table is all observations in which a subject chose to consult at least once. The columns show the number of observations where the Bayesian update is outside the range of the subject's neighborhood. The rows show the number of observations when the subject's updated estimate following the initial consultation was outside the range of their neighborhood's estimates.

		Bayesian update outside range		Total
		Yes	No	
Actual update outside range	Yes	134	22	156
	No	313	413	726
	Total	447	435	882

and the Bayesian update was outside the range of the subject's neighborhood's estimates. Table 7 in the Appendix shows the results of this probit model. Subjects were more likely to update their estimate outside the range of their neighbors' estimates: the further the neighborhood average was from 0.5; the further the Bayesian estimate was from the neighborhood's range; or if they received a 'moderate' sample⁴.

That the subjects were more likely to update outside the neighborhood range the further the neighborhood average was from 0.5, perhaps indicates that, when the subjects were more certain that the number of white balls was on one side of 50 than the other, they felt more secure giving an update further from 0.5 and hence more likely outside their range. Assuming that subjects that used Bayesian updating had some error around the Bayesian estimates, it is not unexpected that the further the Bayesian update was outside the neighborhood range the more likely they were to also update their estimates outside the range of their neighborhoods. The final factor that increased the likelihood of a subject updating her estimate outside the range of her neighbors was if the subject received a moderate sample. This perhaps suggests that subjects with a more extreme sample had more extreme estimates and so did not have so far to update their estimates to outside the range.

Interestingly, personal factors such as, CRT scores, risk aversion and scores on the control questions had no effect on the likelihood of following the Bayesian

⁴ A 'moderate' sample is defined as one that contains two balls of one color and one ball of the other color.

update outside the neighborhood range. Receiving the averages of the neighborhood also did not affect the likelihood of subjects updating outside the range, suggesting that if subjects were Bayesian they could ignore a signal not to be. Being more or less connected also had no effect on the likelihood of being Bayesian.

These results suggest that, while the best model for predicting the updated estimates of subjects was the Degroot model, there was a significant minority of subjects who did not use the DeGroot model, probably in favor of being Bayesian. Moreover, personal factors and factors relating to the network conditions did not seem to affect the likelihood of being Bayesian. Instead conditions relating to the information received, and how confident this information made subjects, caused them to follow the Bayesian estimate.

6.2 Consultations: Testing Hypothesis 2

To test Hypothesis 2, regressions were run at the group level and at the individual level to investigate the influences on the number of times agents consulted. Table 8 shows the results for regressions at the group level, while Table 9 shows the results from two regressions looking at the decisions of individuals to consult.

Under the DeGroot model, groups reach consensus if allowed to consult an infinite number of times. In reality people do not consult infinitely so the first group-level regression considered the factors that determined the number of times groups consulted. When subjects were connected on sparse networks, the average number of consultations was significantly greater than when subjects were connected on more highly connected networks (treatment 2). However, when subjects were on highly-connected networks with additional information (treatment 3), the average consultations was not significantly different to those under treatments one or two. This is consistent with the analysis provided previously that the marginal benefit of consulting falls more quickly on highly connected networks than on sparsely connected networks and so groups on the sparsely connected network consult more times. That subjects were no more likely to consult under the ‘information’ treatment than under the ‘baseline’ treatment from the densely-connected network indicates that the costs of consulting were not a factor. Fatigue does not appear to be a factor in deciding whether or not to consult, given that the number of consultations groups undertook was not affected by how late in the session it was.

Another result was that greater initial variance in a group’s estimates led the

group members to consult more times. As shown in Section 4, there is greater marginal benefit to consulting if the group has greater initial variance and so, this evidence suggests that the main factor causing people to consult longer is the perceived marginal benefits of consultation.

Given that decisions around consulting occurred at the individual level it is worthwhile analyzing the individual level data so I now look at the factors influencing their individual decisions regarding how many times to consult.

Each of the regressions in Table 9 has the number of times the subjects chose to consult as the dependent variable. The first regression includes all subjects in all consultation rounds (1080 observations). The second regression restricts the sample to those times when the subject chose to consult at least once.

The results of the first regression (that included those subjects that chose to consult and those that chose not to consult) indicate that subjects chose to consult more times if: they were less risk averse; they received a 'moderate' sample; their initial estimate was further from 0.5; or their initial estimate was further from their neighborhood's average. Subjects tended to choose to consult fewer times if: their neighborhood's average was further from 0.5 or if their neighborhood's variance was greater. Of the treatments, subjects consulted more times if they were under treatment one than treatments two or three, while there was no statistically significant difference between treatments two and three.

As will be shown later, the dominant group among the subjects who chose not to consult were subjects who were risk averse and chose to estimate 0.5 regardless of other factors. This explains why these factors became insignificant when the sample is restricted to just those who chose to consult. Furthermore, receiving a moderate sample does not affect the number of consultations once the sample is restricted to those who choose to consult at least once, suggesting the receiving a moderate sample makes subjects more likely to choose to consult than not consult, but does not affect the number of times they choose to consult after that.

Looking at those factors that remain significant when the sample is restricted to those who chose to consult at least once, less risk averse subjects consulted more than more risk averse subjects. This result is somewhat surprising as choosing to consult more did not involve any reduction in the payoff and so would seem to be a less risky strategy. Potentially the risk averse subjects were less likely to stray too far from an estimate of 0.5 and so saw less gain in consulting.

Having fewer connections lead to more consultations than having more con-

nections. This is consistent with the analysis that suggested that the marginal benefits of consultation would reduce more quickly with more connections than fewer. Receiving the averages of neighbors' beliefs did not affect the number of consultations, again suggesting that the costs involved in updating beliefs was not a significant factor in deciding how many times to consult.

Once subjects had chosen to consult, measures of more general uncertainty tended to increase the number of consultations subjects chose, but measures of more individual uncertainty did not. The closer the number of white balls were to fifty (and hence the closer to the border between the two payoff states) and the closer was the neighborhood average to 0.5 the more consultations subjects chose. However, receiving a moderate sample and having an initial estimate closer to 0.5 did not affect the number of consultations chosen once the subject had chosen to consult.

Given the neighborhood variance, the further subjects' initial estimates were from the neighborhood average the more they would consult. This suggest that if subjects were initially more in line with their neighborhood they felt there was less to gain through consultation and so consulted less.

Taken together these results suggest that subjects' decisions to consult were based primarily on the perceived benefits of consulting rather than the costs of doing so. That is, subjects continued to consult until they felt that there was no more marginal benefit to consulting. Furthermore, uncertainty about that state of nature induced subjects to choose to consult and consult more times.

6.3 Convergence to Consensus

One of the predictions of the DeGroot model is that strongly-connected networks will converge to consensus if left to continue to consult. To test this prediction regressions were run at the group level for all groups in the consultations rounds (216 observations). Table 8 shows the results for these regressions.

In this experiment, consensus can be defined as all members of the group having the same beliefs. This will entail that there is no variance in group members' beliefs after consultation is finished. Only six times in the experiment was the final variance of a groups' beliefs equal to zero. Of these instances, five times the consensus was at zero or one, and once the consensus was at 0.5. Consensus occurred once under treatment one, four times under treatment two and once under treatment three. The consensus at 0.5 was under treatment two and took two consultations to reach. Of the other instances of consensus, three were

reached with a maximum of two consultations within the group and the other two instances took three consultations within the groups.

Given how few groups reached consensus, it is difficult to draw many conclusions from those instances. Instead I looked at the factors that influence the final variance of the groups and finally the factors that cause groups members' beliefs to converge. Table 8 shows that groups that had a larger initial variance tended also to have a larger final variance and that groups' final variance tended to be smaller later in the sessions. Interestingly, the number of consultations did not affect the final variance nor did the number of connections. These results suggest that while greater initial variance led group members to consult more often, this extra consultation did not reduce the final variance to the same levels as those groups with smaller initial variance. So perhaps while groups with greater variance saw greater initial benefit in consulting, this extra consulting did not ultimately lead to closer beliefs due to the cost of integrating more diverse information. That the treatments did not affect the final variance suggests that the costs of integrating more information sources was not a factor in consultations.

Groups' members' beliefs converged (that is, the final variance of the group's beliefs was less than the initial variance) 73 per cent of the time (158 groups out of 216). Convergence is predicted both if it is assumed that the subjects follow the DeGroot model or that they are Bayesian (Gale & Kariv (2003), Choi et al. (2005), Choi et al. (2012)). Therefore, it is interesting that over a quarter of groups' beliefs diverged. Groups were more likely to converge if they had a greater initial variance (which possibly reflects that they had more scope to converge) or if it was later in the session. On the whole, the treatments had no effect on the likelihood of groups converging, nor did the number of consultations. These are somewhat strange results as the DeGroot model predicts that convergence should occur more quickly on more connected networks, and that convergence should be more likely the more times a group consults.

6.4 Accuracy of beliefs

Ultimately, agents seek information about the world in order to improve their beliefs about the world and their confidence in the correct beliefs. Therefore, it is worthwhile considering which factors led subjects to become more confident of the correct answers. Two regressions were run to investigate the efficacy of the consultations. The results of these regressions are shown in Table 10. In each

of these regressions, the dependent variable is the absolute difference between the subjects' final estimates and the Correctly Certain (CC) beliefs. CC beliefs are defined as estimates that reflected subjects that were certain of the correct state of nature. That is, if the number of white balls is at least 50, the CC belief is 1 and if the number of white balls is less than 50, the CC belief is 0. This measure was used rather than the payoffs themselves as the non-linear nature of the payoffs would overweight incorrect answers. The first regression's sample is all subjects in all consultation rounds (1080 observations). The sample for the second regression is the observations in which subjects chose to consult at least once (882 observations).

In the regression that includes the non-consulters, subjects' final estimates were closer to the CC beliefs if: they were less risk averse; they consulted a larger neighborhood; the number of white balls was further from 50; their initial estimate was further from 0.5; their neighborhood average was further from 0.5; their neighborhood had a smaller variance; or their initial estimate was closer to their neighborhood's average. When the sample is restricted to those subjects who chose to consult at least once most of the same factors remain significant. However, those factors relating to subjects' initial estimates are no longer significant.

It is reasonable that factors relating to initial estimates should have less impact when excluding non-consulters as, for non-consulters, their initial estimates are their final estimates and so will have greater impact. Nonetheless, it is interesting that for people who choose to consult, individual initial conditions appear not to have any impact on their final result. Having an initial estimate close to 0.5 or further from the neighborhood average can be overcome. What cannot be overcome, however, are external initial conditions. That is, if the number of white balls is near 50, the initial neighborhood average is near 0.5 or greater initial neighborhood variance all negatively impact the final result.

Subjects do significantly worse under treatment one than under treatment two or three, with no significant difference between treatments two and three. This suggests that consulting more widely improves outcomes. On the other hand, reducing the costs of incorporating information does not improve outcomes, possibly because subjects did not use the averages much anyway. Interestingly, neither consulting more nor revising estimates more improved the outcomes, suggesting that consulting widely is important but consulting more often among your neighbors does not improve results. Alternatively, this shows

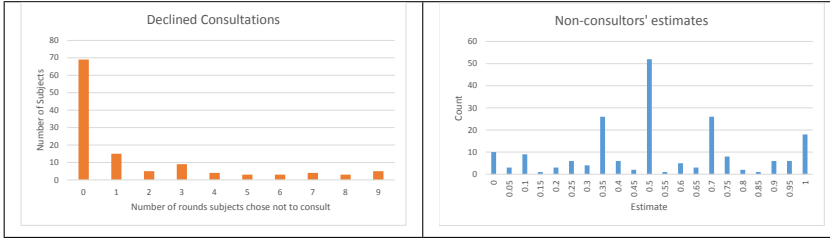


Figure 3: Declined consultations and estimates. The chart on the left shows the number of subjects who chose not to consult the relevant number of times. For example, 15 subjects chose not to consult once, while 9 subjects chose not to consult three times. The chart on the right is a histogram of the estimates made by subjects who chose not to consult with their neighbors in a particular consultation round. The modal estimate is 0.5 with smaller peaks at the true estimates associated with the ‘moderate’ samples.

that information travels more easily around a highly connected networks than a less-connected network as the information path around a less-connected network is more easily broken by an agent stopping consulting.

7 Subsidiary Results

7.1 Non-consulters

Of the 1080 times that subjects could consult (120 subjects by 9 consultation rounds), subjects declined to consult with their neighbors 198 times, or nearly 20 per cent of the time. Given that consulting did not reduce the subject’s pay-off, it seems somewhat strange that, so many times, subjects should choose not to consult at all. As shown in Figure 3, over 40 per cent of the subjects (51 of the 120 subjects) chose not to consult at least once, with five subjects choosing not to consult at all. Figure 3 also shows a histogram of the estimates given by subjects when they did not consult. The modal estimate is 0.5. Smaller peaks can be seen around the true estimates associated with the moderate samples and smaller peaks at 0 and 1. The smaller peaks at 0 and 1 compared with the true estimates associated with the moderate samples, suggest that non-consulters who received ‘extreme’ samples split their estimates between the true estimates associated with those samples and the corner solutions.

To investigate the decision not to consult, a probit model was run on a dummy variable that took the value of one if the subject chose not to consult in a particular consultation round, and zero otherwise. The results of this model are shown in Table 11. Subjects were more likely to choose not to consult if: they were more risk averse; it was earlier in the session; they had not received a 'moderate' sample; or their estimate was closer to 0.5. These results suggest that the dominant group among the non-consulters were subjects who were more risk averse, and chose to estimate 0.5 knowing that they would then be guaranteed to receive 75 kronor regardless of the number of white balls in the urn.

As mentioned in the previous section, when subjects received a moderate sample they were presumably less certain about the state of nature and so desired more information before making their final estimate and so chose to consult at least once. That subjects were less likely to choose not to consult later in the session possibly reflects that they felt more comfortable with the process and so were more willing to consult, or alternatively they might have felt that they were not doing so much and so boredom led them to consult.

To further investigate the group of non-consulters, an additional regressions was run to investigate the factors that led subjects who chose not to consult to estimate closer to the correctly certain beliefs (CC beliefs). As shown in Table 11, subjects were more likely to make an estimate closer to the CC beliefs if they: were less risk averse; did better on the control questions; or did not get a moderate sample. These results again confirm that the dominant group among the non-consulters were risk averse, with a less intuitive grasp of probability theory, who chose to estimate 0.5.

7.2 Initial Estimates

In order for subjects to use Bayesian updating they must be able to rely on the estimates of their neighbors to provide accurate information about the evidence on which those estimates are based. Therefore, an investigation into the accuracy of the beliefs as expressed by the estimates is warranted.

The subjects were provided with the actual probabilities of there being at least fifty white balls in the urn based on the different possible samples. Figure 4 shows the share of subjects who accurately reported those probabilities in their initial estimates. Over the first nine rounds of the experiment, there appears to be an increasing share that report those probabilities accurately (and those who are within five percentage points of the correct estimates). Up to 70 per cent of

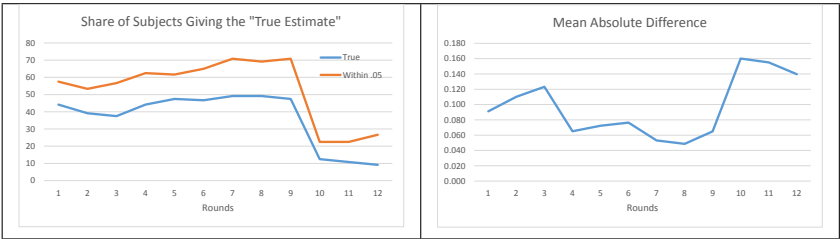


Figure 4: Initial estimates. The figure on the left shows the share of subjects whose initial estimate was the true estimate of the probability that there were at least 50 white balls in the urn, as well as the share whose initial estimate was within five percentage points of the true estimate. The figure on the right shows the mean absolute difference between the subjects’ initial estimates and the true probabilities based on their samples.

subjects gave initial estimates that were within 5 per cent of the true probability. In the final three rounds, however, this share drops dramatically. Figure 4 also shows the mean absolute difference between the initial estimates and the true estimates. Again there appears to be an improvement in the accuracy of initial estimates over the course of the first nine rounds, which then reverses over the final three rounds. This suggests that over the course of the consultation rounds (rounds 1-9) subjects learnt to give more accurate estimates based on their samples. However, once they got to the non-consultation rounds (rounds 10-12) they felt less need to be accurate in their estimates.

Possibly this was because they found that they preferred to have accurate information from other subjects and so due to some form of cooperative motivation tried to help other subjects as well by giving accurate information. However, once the consultation rounds were finished, there was no cooperative need to give the true estimates anymore. Therefore, subjects instead gave estimates more in line with predictions based on their risk aversion. This explanation is borne out by regressions conducted on the subjects initial estimates.

Table 12 shows the results of two regressions. The first regression shows the factors contributing to the accuracy of the initial estimates, measured by the absolute difference between the initial estimates and the true probabilities based on the relevant samples. As the figures suggested, subjects got more accurate as the round number increased but were much less accurate in the non-consultation rounds (T4). Furthermore, this regression suggests that subjects

were more accurate in their initial estimates if they got a 'moderate' sample.

Giving the correct estimate of the probability based on the subject's beliefs is only the utility maximizing solution if the subject is risk neutral. If the subject is risk averse they should give an estimate that is closer to 0.5 than their beliefs about the probability that there are at least 50 white balls, and if they are risk loving they should give an estimate that is further away from 0.5 than their beliefs. The second regression in Table 12 tests this proposition. The dependent variable is the absolute difference between the initial estimate and 0.5, minus the absolute difference between the true estimate and 0.5. This measures whether or not the subject's initial estimate is closer to 0.5 than the true probability. The results show that a subject is more likely to give an estimate in the non-consultation rounds that is closer to 0.5 than the true estimate if they did better on the CRT or if they received a moderate sample. However, the subjects' risk aversion did not make them more likely to give an initial estimate further away from 0.5 than the true estimate.

These results suggest that, as subjects went through the experiment, they gave progressively more accurate initial estimates, possibly in an act of reciprocity as they appreciated receiving accurate estimates from their neighbors. However, once they got to the non-consultation rounds they felt less obligated to give accurate estimates. Nonetheless, during the consultation rounds, with up to 70 per cent of subjects within 5 per cent of the true estimate, it can be considered that it was reasonable for subjects to rely on other subjects' estimates to come to conclusions about their neighbors' samples.

8 Conclusions

For many applications in economics it is important to have an accurate model of how people incorporate information into their beliefs. This study compares two of the main theories of information incorporation. I investigated this topic through conducting an experiment in which agents had to estimate the probability that the number of white balls in an urn of 100 balls was greater than 50, using a private sample and information about other agents' estimates of the probability. The specific questions that I sought to answer were whether agents used a Bayesian or DeGrootian method to incorporate other agents' beliefs into their own; what influenced agents decisions to consult with their neighbors; what caused groups to approach consensus and how accurate agents' final esti-

mates were.

Overall, the study shows that over half of agents would take a Bayesian approach at least some of the time, which is a significantly greater share than previous research has suggested. Nonetheless, an estimate based on the DeGroot model outperforms a Bayesian estimator in predicting the updates of agents' beliefs. The size of the group that will use a Bayesian approach gives some support to the use of this model in economic theories, though it should be maintained that the DeGroot approach will possibly still provide more accurate estimates overall.

Decisions of agents appear to be focused mostly on the perceived benefits of consulting rather than any costs involved. Agents tend to consult more times with their neighbors when they have fewer neighbors or if there is more general uncertainty about the state of nature. This confirms that the perceived marginal benefits from consulting leads people to consult more when they are less connected. However, this increase in consultation does not lead to more convergence to others' beliefs or to better final beliefs. The average number of consultations and the connectedness of the networks have no effect on how close to consensus groups finish, despite the DeGroot model suggesting that they should. Agents have better final estimates the more connected they are. However, the number of times they consult with their neighbors has no effect on the accuracy of their final estimates.

This is not to say that consultation cannot improve beliefs. Consultation can improve beliefs by overcoming individual uncertainty about the state of the world, but it will not overcome more general uncertainty. While consulting more widely can improve beliefs, consulting the same neighbors more often does not improve beliefs. This possibly suggested that information is more likely to be passed around a highly connected network than a poorly connected one, as it is far easier for an information path to be broken in a poorly connected network.

References

- Acemoglu, D., Dahleh, M. A., Lobel, I. & Ozdaglar, A. (2011), 'Bayesian learning in social networks', *The Review of Economic Studies* **78**(4), 1201–1236.
- Berger, R. L. (1981), 'A necessary and sufficient condition for reaching a consensus using DeGroot's method', *Journal of the American Statistical Association* **76**(374), 415–418.
- Chandrasekhar, A. G., Larreguy, H. & Xandri, J. P. (2015), Testing models of social learning on networks: Evidence from a lab experiment in the field, Working Paper 21468, National Bureau of Economic Research.
- Charness, G., Karni, E. & Levin, D. (2007), 'Individual and group decision making under risk: An experimental study of Bayesian updating and violations of first-order stochastic dominance', *Journal of Risk and Uncertainty* **35**(2), 129–148.
- Charness, G. & Levin, D. (2005), 'When optimal choices feel wrong: A laboratory study of Bayesian updating, complexity, and affect', *The American Economic Review* **95**(4), 1300–1309.
- Choi, S., Gale, D. & Kariv, S. (2005), *Experimental and Behavioral Economics (Advances in Applied Microeconomics, Volume 13)*, Emerald Group Publishing Limited, chapter Behavioral Aspects of Learning in Social Networks: An Experimental Study, pp. 25–61.
- Choi, S., Gale, D. & Kariv, S. (2012), 'Social learning in networks: a quantal response equilibrium analysis of experimental data', *Review of Economic Design* **16**(2-3), 135–157.
- DeGroot, M. H. (1974), 'Reaching a consensus', *Journal of the American Statistical Association* **69**(345), 118–121.
- DeMarzo, P. M., Vayanos, D. & Zwiebel, J. (2003), 'Persuasion bias, social influence, and unidimensional opinions', *Quarterly Journal of Economics* **118**(3), 909–968.
- Eckel, C. C. & Grossman, P. J. (2002), 'Sex differences and statistical stereotyping in attitudes toward financial risk', *Evolution and Human Behavior* **23**(4), 281–295.

- El-Gamal, M. A. & Grether, D. M. (1995), 'Are people Bayesian? Uncovering behavioral strategies', *Journal of the American Statistical Association* **90**(432), 1137–1145.
- Ellison, G. & Fudenberg, D. (1993), 'Rules of Thumb for Social Learning', *Journal of Political Economy* **101**(4), 612–43.
- Ellison, G. & Fudenberg, D. (1995), 'Word-of-mouth communication and social learning', *The Quarterly Journal of Economics* **110**(1), 93–125.
- Eyster, E. & Rabin, M. (2009), Rational and naive herding, CEPR Discussion Papers 7351, C.E.P.R. Discussion Papers.
- Eyster, E. & Rabin, M. (2010), 'Naïve herding in rich-information settings', *American Economic Journal: Microeconomics* **2**(4), 221–43.
- Fischbacher, U. (2007), 'z-Tree: Zurich toolbox for ready-made economic experiments', *Experimental Economics* **10**(2), 171–178.
- Frederick, S. (2005), 'Cognitive reflection and decision making', *Journal of Economic Perspectives* **19**(4), 25–42.
- Gale, D. & Kariv, S. (2003), 'Bayesian learning in social networks', *Games and Economic Behaviour* **45**(2), 329–346.
- Golub, B. & Jackson, M. O. (2010), 'Naïve learning in social networks and the wisdom of crowds', *American Economic Journal: Microeconomics* **2**(1), 112–49.
- Grether, D. M. (1992), 'Testing Bayes rule and the representativeness heuristic: Some experimental evidence', *Journal of Economic Behavior & Organization* **17**(1), 31 – 57.
- Holt, C. A. & Smith, A. M. (2009), 'An update on Bayesian updating', *Journal of Economic Behavior & Organization* **69**(2), 125 – 134.
- Jackson, M. O. (2008), *Social and Economic Networks*, Princeton University Press, Princeton, New Jersey.
- Jadbabaie, A., Molavi, P., Sandroni, A. & Tahbaz-Salehi, A. (2012), 'Non-Bayesian social learning', *Games and Economic Behavior* **76**(1), 210–225.

- Judd, S., Kearns, M. & Vorobeychik, Y. (2010), 'Behavioral dynamics and influence in networked coloring and consensus', *Proceedings of the National Academy of Sciences*.
- Kearns, M., Judd, S., Tan, J. & Wortman, J. (2009), 'Behavioral experiments on biased voting in networks', *Proceedings of the National Academy of Sciences*.
- Kearns, M., Suri, S. & Montfort, N. (2006), 'An experimental study of the coloring problem on human subject networks', *Science* **313**(5788), 824–827.
- Mueller-Frank, M. (2014), 'Does one Bayesian make a difference?', *Journal of Economic Theory* **154**, 423 – 452.
- Mueller-Frank, M. & Neri, C. (2015), A General Model of Boundedly Rational Observational Learning: Theory and Experiment, IESE Research Papers D/1120, IESE Business School.
- Tenenbaum, J. B., Griffiths, T. L. & Kemp, C. (2006), 'Theory-based Bayesian models of inductive learning and reasoning', *Trends in Cognitive Sciences* **10**(7), 309 – 318. Special issue: Probabilistic models of cognition.

A Appendix

Table 4: Risk choices. Subjects were asked to choose one of the gambles described below. In each gamble, each of the two possible outcomes had a 0.5 probability of occurring, as though a coin was tossed and the subject would be awarded the amount of money corresponding to the result of the coin toss for the gamble they had chosen. The subjects were not actually paid the amount shown but rather asked to choose the gamble that they would most prefer if they were to be paid for their choice. This is a modified version of the risk choices in Eckel & Grossman (2002).

Gamble number	1	2	3	4	5	6	7	8
Heads	240	210	190	170	150	110	70	10
Tails	240	300	340	380	420	500	540	570

Instructions

February 19, 2016

1 General Information

- In this experiment, there will be 4 phases of 3 rounds for a total of 12 rounds.
- Your total earnings for the experiment will be based on your actions in 3 of the 12 rounds randomly selected by the computer.
- At the beginning of each round, you will be randomly assigned to a group of five people.
- In each round, you will be asked to say how certain you are that there are at least 50 white balls in an urn of 100 balls.
- To assess this you will be provided with a random sample of balls from the urn and you can interact with other people in your group.
- Your payoff will increase the more certain you are of the correct answer. That is,
 - If there are at least 50 white balls in the urn, your payoff increases the more certain you are that this is so.
 - If there are less than 50 white balls in the urn, your payoff increases the more certain you are that this is so.
- You will interact exclusively within your group through your computer for the duration of each round without knowing the identity of the other members of your group.
- Your payment after the experiment will be based on your actions taken during the experiment.
- If you have questions during the experiment please raise your hand.

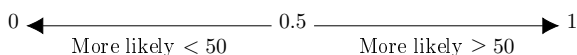
2 The Experiment

- In the experiment, there is a hypothetical urn with 100 balls in it. Each of these balls is either white or black.
- There will be a separate urn for each group.
- The number of white balls in the urn will be randomly generated at the start of each round.
- You will be asked to assess the likelihood that the number of white balls in the urn is at least 50.
- To assist in assessing this likelihood each person will each receive their own random sample of three balls from the urn.
- Your sample will appear on your computer screen as a set of 3 letters. Each of the letters will be either “w” or “b”, depending on how many of the balls in your sample are white or black.
- Having initially assessed the likelihood that the number of white balls in the urn is at least 50, based on your sample, you will then be able to consult with one of the other members of your group (your neighbor).
- This consultation will involve viewing your neighbor’s estimate, based on their sample. Your neighbor will be able to consult with their neighbor at the same time. Their neighbor will be another member of your group.
- You can then update your estimate, and your neighbor can update their estimate as well.
- You can then elect to consult with your neighbor again to view their updated estimate.
- Consultation continues until no-one on the network wants to consult again.
- Your final estimate will be used to determine your earnings.

3 Estimate

- Your estimate of the likelihood that the number of white balls in the urn is at least 50 will be a number between zero and one.
- The estimate can be made up to an accuracy of 3 decimal places.
- The closer the estimate is to one, the more certain you are that the number of white balls is at least 50.

- The closer the estimate is to zero, the more certain you are that the number of white balls is less than 50.
 - If your estimate is 1, this means that you are certain that there are at least 50 white balls in the urn.
 - If your estimate is zero, this means that you are certain that there are less than 50 white balls in the urn.
 - If your estimate is 0.500, this means that you think it is equally likely that there are more or less than 50 white balls in the urn.



4 Earnings

Your earnings each round will depend on your final estimate and the number of white balls in the urn.

- Your earnings each round will be between zero and 100 kr.
- You will be paid for 3 rounds randomly selected by the computer. Therefore, your maximum earnings for the experiment will be 300 kr.
- If the sum of your earnings for the three rounds selected by the computer is less than 50 kr, you will be paid 50 kr. Therefore, your minimum total earnings for the experiment will be 50 kr.
- If the number of white balls in the urn is at least 50, your earnings will be greater, the closer your final estimate is to 1.
- If the number of white balls in the urn is less than 50, your earnings will be greater, the closer your final estimate is to zero.
- Examples:
 - if the number of white balls in the urn is at least 50; and:
 - * your estimate is 1.000, your earnings are 100 kr.
 - * your estimate is 0.750, your earnings are 93,75 kr.
 - * your estimate is 0.500, your earnings are 75 kr.
 - * your estimate is 0.250, your earnings are 43,75 kr.
 - * your estimate is 0.000, your earnings are 0 kr.

- if the number of white balls in the urn is less than 50; and:
 - * your estimate is 1.000, your earnings are 0 kr.
 - * your estimate is 0.750, your earnings are 43,75 kr.
 - * your estimate is 0.500, your earnings are 75 kr.
 - * your estimate is 0.250, your earnings are 93,75 kr.
 - * your estimate is 0.000, your earnings are 100 kr.

For those interested, the mathematical formula for your earnings is as follows:

$$\pi(x) = \begin{cases} 100 [1 - (x - 1)^2] & \text{if } W \geq 50 \\ 100 [1 - x^2] & \text{if } W < 50. \end{cases}$$

Where $\pi(x)$ is your earnings, x is your final estimate of the probability that the number of white balls is at least 50, W is the number of white balls in the urn.

5 The Groups

The groups in which you will consult are as shown in Figure 1.

- There are five people in each group, represented by the circles in the figure.
- Each of the five people assigned to each group will be randomly and anonymously assigned to one of five roles (A, B, C, D, E).
- Each of these roles are essentially identical.
- The arrows show the directions of consultation.
- Each subject can view one of the other four people's estimates on the network.
- Each subject's estimates can be viewed by one of the other four people on the network.
- During consultation:
 - subject A will see the estimates of subject B
 - subject B will see the estimates of subject C
 - subject C will see the estimates of subject D
 - subject D will see the estimates of subject E
 - subject E will see the estimates of subject A

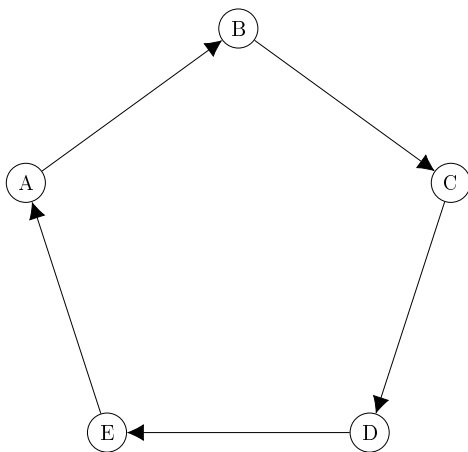


Figure 1: Consultation Network. The circles represent subjects and the arrows represent the direction of viewing. Each subject can view one of the other four people's estimates on the network. Each subject's estimates can be viewed by one of the other four people on the network.

6 Feedback and Payment

- At the end of each round the computer will show you: the number of white balls in the urn, your final estimate and how much you will earn for that round if it is selected for payment.
- At the end of the experiment your total earnings will be the sum of your earnings from each of your selected rounds, rounded to the nearest krona.

Table 5: Cognitive Reflection Task (CRT). Subjects were asked to answer the following questions as part of the final questionnaire. The number of questions they got correct was used as data. This is an extended version of the CRT in Frederick (2005).

Question	
1	A house contains a living room and a kitchen that are perfectly square. The living room has four times the area of the kitchen. If the walls of the kitchen are four meters long, how long are the walls in the living room?
2	A store owner reduced the price of a pair of SEK 1000 shoes by 10%. The next week, he reduced it by a further 10%. How much do the shoes cost now?
3	If it takes 5 machines 5 minutes to make 5 forks how long (in minutes) would it take 100 machines to make 100 forks?
4	In a lake, there is a patch of lily pads. Every day, the patch doubles in size. If it takes 48 days for the patch to cover the entire lake, how long (in days) would it take for the patch to cover half of the lake?
5	A tennis racket and a ball cost SEK 110 in total. The racket costs SEK 100 more than the ball. How much does the ball cost (in SEK)?

Table 6: Data description

Variable	Description
ABS(Bayes - range)	The absolute distance from the Bayesian update to the range of the subject's neighborhood's initial estimates.
ABS(final - CC)	The absolute difference between the subject's final estimate and the CC belief.
ABS(init - nhood ave)	The absolute difference between the subject's initial estimate and the average of their neighborhood's initial estimates.
ABS(init - true)	The absolute difference between the subject's initial estimate and the true probability based on their sample.
ABS(init - 0.5)	The absolute difference between the subject's initial estimate and 0.5
ABS(init - 0.5) - ABS(true - 0.5)	The absolute difference between the subject's initial estimate and 0.5 minus the absolute difference between the true estimate based on the subject's sample and 0.5.
ABS(init ave - 0.5)	The absolute difference between the average of the group's initial estimates and 0.5
ABS(Nhood ave - 0.5)	The absolute difference between the average of subject's neighborhood's initial estimates and 0.5.
ABS(Wballs - 50)	The absolute difference between the number of white balls in the urn and 50.
Consultations	The number of times the subject consulted their neighbors' estimates in a round.
CRT Correct	The number of questions the subject got correct on the CRT.
MAE(CQ)	The mean absolute error of the subject's answers to the control questions.
N10	A dummy variable: 1= there were 10 people in the session; 0 otherwise.
N20	A dummy variable: 1= there were 20 people in the session; 0 otherwise.
N'hood variance	The variance in the initial estimates of the subject's neighborhood.
Revisions	The number of times the subject revised their estimates in a round.
Risk choice	The number of the gamble that the subject chose. The higher the number, the more risky was the gamble.
Round	The number of the round in the session.
Smod	A dummy variable: 1= the subject received a 'moderate' sample; 0 otherwise. (A moderate sample is a sample with two balls of one color and one of the other color.)
T1	Treatment 1: Sparse network, baseline treatment
T2	Treatment 2: Dense network, baseline treatment
T3	Treatment 3: Dense network, information treatment
T4	Treatment 4: No consultation treatment
T3 previous	A dummy variable: 1=the subject had performed T3 in a previous phase; 0 otherwise.

Table 7: Bayes vs Degroot Probit model. The dependent variable is a dummy variable that is 1 if the subject's first updated estimate is outside the range of their neighborhood's estimates and zero otherwise. Positive coefficient implies that the subject's estimate is more likely to be outside the range of their neighborhood's estimates. The sample is all observations when the subject chooses to consult at least once and the Bayesian update is outside the range of the subject's neighborhood's estimates (447 observations). ***=(p-value<0.01); **=(p-value<0.05); *=(p-value<0.1). Marginal effects were calculated holding other variables at their medians.

Dependent variable	Dummy (actual update outside range)	
Sample	Consultations>0 and Bayes outside range	
No. of Observations	447	
Variable	Coefficient	Marginal effect
Constant	-2.8560** (1.1202)	
CRT correct	0.0563 (0.00670)	0.0167 (0.0204)
Risk Choice	-0.0094 (0.0467)	-0.0028 (0.0140)
MAE(CQ)	0.3091 (0.6206)	0.0917 (0.1815)
T2	0.1817 (0.1776)	0.0570 (0.0536)
T3	-0.1042 (0.2710)	-0.0298 (0.0795)
T3 Previous	-0.0399 (0.4322)	-0.0118 (0.1297)
Round	0.0056 (0.0492)	0.0016 (0.0148)
ABS(Wballs - 50)	-0.0095 (0.0078)	-0.0028 (0.0023)
Smod	0.4354** (0.2159)	0.1291* (0.0715)
ABS(init - 0.5)	1.2080 (0.9163)	0.3582 (0.2835)
ABS(Nhood avge - 0.5)	5.6309*** (0.8830)	1.6698*** (0.3486)
ABS(Bayes - range)	5.0924*** (1.2237)	1.5101*** (0.4186)
Pseudo R^2	0.1648	

Table 8: Group analysis. The dependent variable for the first regression is the average number of consultations in each group. The dependent variable for the second regression is the final variance of estimates for each group. The dependent variable for the final probit model is a dummy variable that records 1 if the group converged (final variance less than initial variance) and zero otherwise. Marginal effects were calculated holding other variables at their medians. The sample for each model is all groups in the consultation rounds (216 observations). ***=(p-value<0.01); **=(p-value<0.05); *=(p-value<0.1)

Dependent variable	Average Consultations	Final Variance	Convergence	
Sample	all groups	all groups	all groups	
No. of Observations	216	216	216	
Variable	Coefficient	Coefficient	Coefficient	Marginal effect
Constant	2.6682*** (0.1801)	0.0224 (0.0126)	-0.3478 (0.7471)	
T2	-0.1421* (0.0843)	-0.0009 (0.0058)	-0.2276 (0.2483)	-0.0499 (0.2483)
T3	-0.1985 (0.1372)	-0.0007 (0.0095)	-0.5901 (0.4034)	-0.1295*** (0.0488)
T3 Previous	-0.1676 (0.2110)	0.0153 (0.0145)	-0.5711 (0.6180)	-0.1253 (0.0969)
Round	0.0155 (0.0244)	-0.0032* (0.0017)	0.1553** (0.0723)	0.0341*** (0.0100)
ABS(Wballs-50)	-0.0049 (0.0044)	0.0003 (0.0003)	-0.0064 (0.0131)	-0.0014 (0.0029)
ABS(init avge - 0.5)	-0.2537 (0.4347)	-0.0481 (0.0299)	-0.5029 (1.2762)	-0.1103 (0.2783)
Initial variance	1.7393* (0.8856)	0.4098*** (0.0615)	9.2773*** (2.9822)	2.0353** (1.0171)
Average Consultations		0.0004 (0.0049)	0.1151 (0.2043)	0.0253 (0.0466)
adj R^2	0.3526	0.2904	0.1307	

Table 9: Consultations. The dependent variable is the number of consultations. The sample for the first regression is all subjects in all consultation rounds (120 subjects by 9 consultation rounds = 1080 observations). The sample for the second regression is subjects in consultation rounds when the subject chooses to consult at least once (882 observations). Standard errors are in parentheses. ***=(p-value<0.01); **=(p-value<0.05); *=(p-value<0.1)

Dependent variable	Consultations	Consultations
Sample	all consultation rounds	Consultations > 0
No. of Observations	1080	882
Variable	Coefficient	Coefficient
Constant	3.119*** (0.6715)	3.5328*** (0.6621)
CRT correct	0.0110 (0.0530)	-0.0017 (0.0425)
Risk Choice	0.0786** (0.0340)	0.0467 (0.0290)
MAE(CQ)	0.0981 (0.4384)	0.2784 (0.3657)
T2	-0.1930** (0.0816)	-0.2184*** (0.0712)
T3	-0.2689** (0.1232)	-0.2638** (0.1056)
T3 Previous	-0.2224 (0.1862)	-0.0829 (0.1642)
Round	0.0212 (0.0206)	-0.0123 (0.0185)
ABS(Wballs-50)	-0.0041 (0.0037)	-0.0069** (0.0032)
Smod	0.3219** (0.1489)	0.0679 (0.1106)
ABS(init-0.5)	1.8835*** (0.6610)	0.3366 (0.4287)
ABS(N'hood avge-0.5)	-1.1425** (0.4421)	-0.7944** (0.3878)
N'hood variance	-2.1094* (1.1793)	-1.7566 (1.1746)
ABS(init-nhood avge)	0.8949** (0.3584)	0.9265** (0.3689)
R^2	0.2363	0.1508

Table 10: Absolute final error. The dependent variable is the absolute difference between the subjects' final estimates and the Correctly Certain (CC) belief. A negative coefficient implies a more accurate estimate. The sample for the first regression is all subjects in all consultation rounds (120 subjects by 9 consultation rounds = 1080 observations). The sample for the second regression is subjects in consultation rounds when the subject chooses to consult at least once (882 observations). Standard errors are in parentheses. ***=(p-value<0.01); **=(p-value<0.05); *=(p-value<0.1)

Dependent variable	ABS(final - CC)	ABS(final - CC)
Sample	All consultation rounds	Consultations > 0
No. of Observations	1080	882
Variable	Coefficient	Coefficient
Constant	0.6760*** (0.0553)	0.6674*** (0.0875)
CRT correct	-0.0058 (0.0053)	-0.0096 (0.0066)
Risk Choice	-0.0080** (0.0031)	-0.0075* (0.0040)
MAE(CQ)	0.0547 (0.0512)	0.0156 (0.0631)
T2	-0.0693*** (0.0200)	-0.0773*** (0.0210)
T3	-0.0952*** (0.0280)	-0.1070*** (0.0305)
T3 Previous	-0.0388 (0.0476)	-0.0412 (0.0539)
Round	0.0055 (0.0053)	0.0064 (0.0059)
ABS(Wballs-50)	-0.0052*** (0.0010)	-0.0060*** (0.0011)
Smod	0.0231 (0.0214)	0.0301 (0.0297)
ABS(init-0.5)	-0.2354*** (0.0803)	-0.0979 (0.1136)
ABS(N'hood avge-0.5)	-0.6697*** (0.0718)	-0.6954*** (0.0887)
N'hood variance	0.6567*** (0.2611)	0.8252*** (0.2910))
ABS(init-nhood avge)	0.1533* (0.0897)	0.0589 (0.1003)
Consultations	0.0034 (0.0076)	0.0083 (0.0095)
Revisions	-0.0145 (0.0117)	-0.0155 (0.0123)
R^2	0.3151	0.3045

Table 11: Non-consulters. The dependent variable in the first probit model is a dummy variable that records 1 if the subject chose not to consult at all in a round, and zero otherwise. A positive coefficient implies a higher likelihood of choosing not to consult. Marginal effects were calculated holding other variables at their medians. The sample is all subjects in the consultation rounds (120 subjects by 9 consultation rounds = 1080 observations). The dependent variable for the second regression is the absolute difference between the estimate and 0.5. The dependent variable for the third regression is the absolute difference between the estimate and the CC belief. The samples for the final two regressions is observations where the subject chose not to consult in the consultations rounds. ***=(p-value<0.01); **=(p-value<0.05); *=(p-value<0.1)

Dependent variable	Non-consultation dummy		ABS(estimate - 0.5)	ABS(estimate - CC)
Sample	All subjects in consultation rounds		Consultations=0	Consultations=0
No. of Observations	1080		198	198
Variable	Coefficient	Marginal effect	Coefficient	Coefficient
Constant	-2.2484** (0.9067)		0.3043*** (0.0971)	0.4096*** (0.1353)
CRT correct	-0.0267 (0.0710)	-0.0052 (0.0140)	-0.0163 (0.0116)	-0.0040 (0.0130)
Risk Choice	-0.1103** (0.0529)	-0.0216 (0.0132)	0.0073 (0.0096)	-0.0327*** (0.0113)
MAE(CQ)	0.2729 (0.5512)	0.0535 (0.1085)	-0.2563** (0.1117)	0.3224** (0.1288)
T2	-0.0708 (0.1021)	-0.0139 (0.0216)	-0.0060 (0.0228)	-0.0412 (0.0390)
T3	0.0781 (0.1682)	-0.0153 (0.0307)	-0.0286 (0.0333)	0.0001 (0.0670)
T3 Previous	0.3746 (0.2700)	0.0735* (0.0438)	0.0245 (0.0607)	0.0055 (0.0895)
Round	-0.0682** (0.0309)	-0.0134** (0.0052)	-0.0045 (0.0067)	0.0037 (0.0114)
ABS(Wballs-50)	-0.0009 (0.0067)	-0.0002 (0.0013)	-0.0002 (0.0012)	-0.0030 (0.0019)
Smod	-0.4522** (0.1968)	-0.0887** (0.0439)	-0.1500*** (0.0234)	0.1719*** (0.0383)
ABS(init-0.5)	-2.7474*** (0.8513)	-0.5389** (0.2352)		
R^2	0.2262		0.5606	0.3644

Table 12: Initial Estimates. The dependent variable in the first regression is the absolute difference between the initial estimate and the actual probability based on the subject's sample. Negative coefficient implies the initial estimate is closer to the true estimate than average. The sample is all subjects' initial estimates across all rounds (120 subjects by 12 rounds = 1440 observations). The dependent variable in the second regression is the absolute difference between the initial estimate and 0.5, minus the absolute difference between the actual probability based on the subject's sample and 0.5. Negative coefficient implies the initial estimate is closer to 0.5 than the true estimate. The sample is all subjects' initial estimates in the non-consultation rounds (120 subjects by 3 non-consultation rounds = 360 observations). ***=(p-value<0.01); **=(p-value<0.05); *=(p-value<0.1)

Dependent variable	ABS(init-true)	ABS(init-0.5)-ABS(true-0.5)
Sample	All experimental rounds	Non-consultation rounds
No. of Observations	1440	360
Variable	Coefficient	Coefficient
Constant	0.2441*** (0.0520)	0.0701 (0.1489)
CRT correct	-0.0044 (0.0047)	-0.0323** (0.0141)
Risk Choice	-0.0007 (0.0036)	-0.0124 (0.0082)
MAE(CQ)	0.0433 (0.0403)	-0.0516 (0.1041)
T2	0.0087 (0.0076)	
T3	0.0203 (0.0152)	
T4	0.1298*** (0.0134)	
T3 Previous	-0.0082 (0.0196)	
Smod	-0.0277*** (0.0104)	-0.1262*** (0.0322)
Round	-0.0068*** (0.0022)	
Round 11		0.0145 (0.0188)
Round 12		0.0359* (0.0206)
R^2	0.1550	0.2200

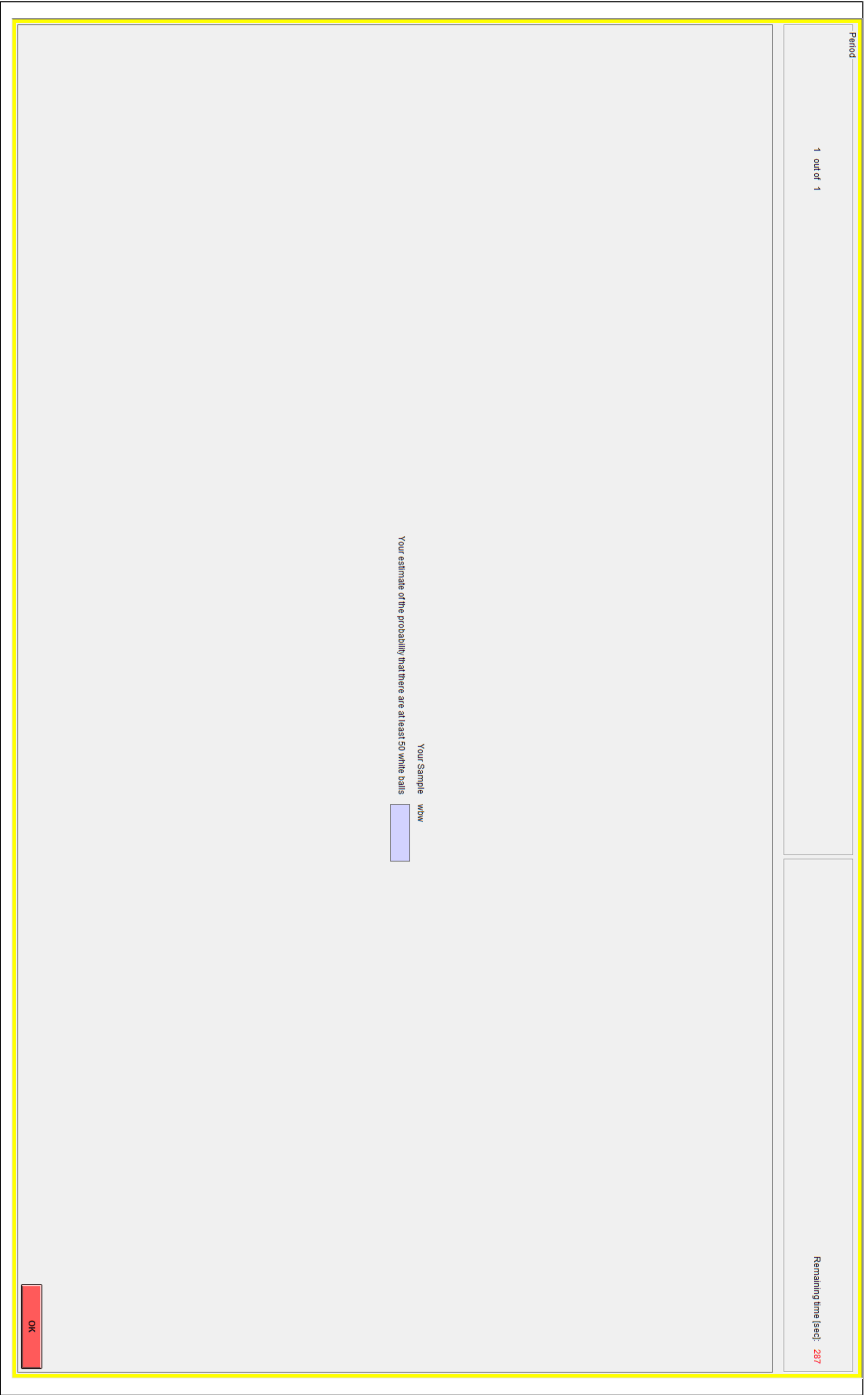


Figure 5: Screenshot from the experiment: Input for the subject's initial estimate based on the subject's sample.

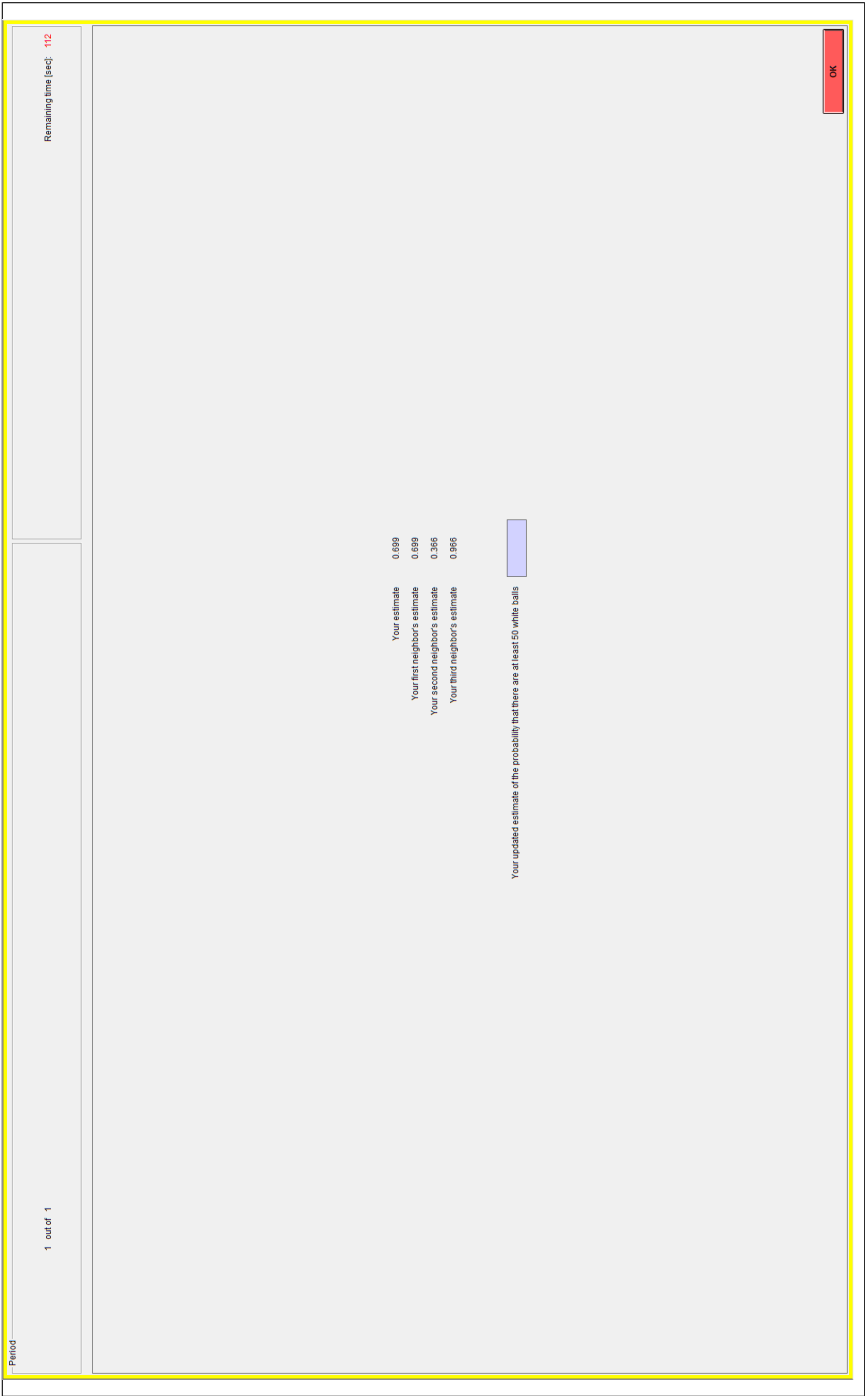


Figure 6: Screenshot from the experiment: Input for the subject's updated estimate based on the subject's neighborhood's estimates in Treatment 2.

PAPER III

Sectoral Shocks and Aggregate Volatility

Abstract

Macroeconomics has generally downplayed the effects of sectoral shocks on aggregate volatility because it was considered that, in a well-diversified economy, sectoral shocks would tend to cancel out at the aggregate level. However, Acemoglu et al. (2012, 2013) have shown that if the input-output network is unbalanced, the effects of sectoral shocks will not decay as fast as the diversification argument contends. In this paper, I extend the model of Acemoglu et al. (2012) by including a demand-side measure of industry influence. Applying this measure to various economies shows that including the demand-side influence of industries can capture important sectoral sources of aggregate volatility. I also present a specific case study of the Australian mining industry to demonstrate this point.

Keywords: Sectoral shocks, macroeconomics, input-output networks, aggregate output

JEL Classification: C67, E32

1 Introduction

Macroeconomics has typically taken the position that sectoral shocks have little effect on the macroeconomy because, in a well-diversified economy, negative shocks to some sectors will be balanced by positive shocks to others. Therefore, shocks at the sectoral level are unlikely to cause fluctuations large enough to be significant at the macro level. Recently, however, researchers have been looking more closely at the arguments behind this theory and have found that there are mechanisms through which sectoral shocks can propagate through the economy to result in aggregate fluctuations. Furthermore, there is some evidence that the importance of sectoral shocks on aggregate volatility has been increasing (Foerster et al. 2011). Moreover, monetary policy can have stronger effects on some industries compared with others, and so policymakers need to be aware of the influence of the industries that they will be affecting (Dixon et al. 2014, Lawson & Rees 2008).

In this paper, I have constructed a measure of demand-side industry influence which expands on the work of Acemoglu et al. (2012, 2013) who constructed supply-side measures of industry influence. These measures of industry influence consider the influence that industries have on other industries through the use of input-output tables and can be used to identify whether an economy is susceptible to sectoral shocks and to rank industries in order of their relative influence on the economy.

We can consider that one industry has influence on another if the first industry supplies inputs to the second industry. Furthermore, this influence is greater the larger the share of the second industry's inputs are sourced from the first industry. The influence of the first industry is further enhanced if the second industry itself has large influence over other industries. Therefore, it is not only the direct effects that an industry has on industries to which it supplies inputs that matter, but also indirect effects on industries that source their inputs from the industries that an industry supplies.

This supply influence is, however, not the only way in which an industry can be influential in an economy. Another source of influence is through the demand side. If an industry demands a large share of another industry's output as an input then the first industry will have influence over the second industry. This paper shows that only focusing on the supply-side influences can lead to incomplete conclusions about whether an economy is susceptible to sectoral shocks which can create a misleading picture of which industries are the most

influential in an economy.

The economic model used in this paper is based on the work of Long & Plosser (1983), in which a representative agent consumes goods that are produced using Cobb-Douglas production technology, in which all goods are potentially used as inputs. By solving for the competitive equilibrium from both the consumers' and the producers' viewpoints I construct a measure of the supply of inputs from one industry to other industries and another measure of the demand for inputs from industries. These measures describe the influence that each industry has over other industries through the supply of and demand for inputs. These measures can also be considered through network theory in terms of measures of centrality where the input-output tables define the network of industry linkages in an economy.

Having constructed measures of industry influence for both the supply- and demand-side, I then estimate the influence of industries in various countries. I find that including the demand-side measures of industry influence alters the measured susceptibility of countries to sectoral shocks. This suggests that only focusing on the supply-side measure potentially ignores important sources of sectoral influences on aggregate volatility. In terms of the more influential industries, the results show that there are strong similarities between countries. For example, the most influential industries in Sweden, the US and Australia tend to be service industries.

A strong example of how incorporating the demand side into measures of influence gives a more complete picture of the industry influences in an economy is the case of Australia's mining industry. In recent years, Australia's economy has performed extremely well, having not had a recession since the early 1990's. One of the main reasons for this strong performance, especially since 2000, has been the strength of the mining industry and its production of exports, particularly to China, but also its demand for inputs as it has expanded its capacity over the past decade or so (Connolly & Orsmond 2011). This indicates that the influence of the mining industry on the Australian economy is greater than its contribution to GDP growth from its gross value added (GVA). However, using Acemoglu et al's measure of industry influence suggests that the mining industry is one of the least influential industries in the Australian economy as it does not produce many inputs for other Australian industries. Rather its influence on the Australian economy is through its demand for other industries' outputs and its production of final output. I confirm this in a case study of Australia, showing

that, despite the recent strong influence of the mining industry on the fortunes of the Australian economy, using only a supply-side measure of industry influence indicates that the mining industry has little importance in the Australian economy. Adding the demand-side measure corrects this misleading picture, confirming that the demand-side measure is important when considering the influences of industries on an economy as a whole.

In Section 2, I review the macroeconomic literature relating to sectoral shocks. In Section 3, I outline my economic model. Section 4 applies the model to a number of economies and Section 5 details a case study of the Australian mining industry. Conclusions are provided in Section 6.

2 Literature Review

The diversification argument that has typically been used in macroeconomics to downplay the importance of sectoral shocks says that in a well-diversified economy the influence of any one industry goes to zero as the number of industries increase (Gabaix 2011). Therefore, shocks to any one sector will be balanced by opposite shocks to other industries and so sectoral shocks will tend to have little aggregate effect. Nonetheless, sectoral shocks have been estimated to contribute significant volatility to aggregate output (Atalay 2014, Roson & Sartori 2014, Mehrotra & Sergeyev 2013, Storer 1996), suggesting that the diversification argument does not always hold. The argument rests on the assumptions that industries are independent and identical, and that the distribution of shocks to those industries are independent and identical. Clearly in most economies industries are not identical, meaning that shocks to larger industries are likely to have a greater impact on the aggregate economy than shocks to smaller ones and so can result in aggregate volatility despite diversification (Gabaix 2011). Moreover, the shocks to industries can be greater in some industries than others and industries can have common cycles or trends, and hence are not independent of each other (Harvey & Mills 2002). Indeed, even with completely independent shocks, if sectors in the economy are not independent these shocks can result in aggregate volatility (Jovanovic 1987).

One of the main channels through which industries are not independent is the input-output linkages between industries (Aroche Reyes & Marquez Mendoza 2012). Input-output linkages between industries result in spillovers between industries so that shocks to one industry can propagate through the econ-

omy far more easily, causing greater aggregate volatility than the diversification argument would suggest (Shea 2002). Therefore, a productivity shock to an industry will not only be felt by that industry but also by those industries that the first industry supplies with inputs (Bak et al. 1992). Furthermore, the second industries will provide inputs to third industries and so, the productivity shock to the first industry will be felt by more industries than just their direct customer industries.

Building on the seminal work of Long & Plosser (1983), and using a measure of industry influence based on the provision of inputs by industries, Acemoglu et al. (2012, 2013) show that certain arrangements of the input-output matrix can result in the influence of industries not falling to zero as quickly as the diversification argument suggests. The supply influence vector outlined in Section 3 broadly replicates Acemoglu et al's work. They show that as the input-output matrix becomes more unbalanced, some industries will retain significant influence over the economy despite an increase in industry diversification. This unbalancedness means that productivity shocks of the same sign are more likely to (and hence will more often) land on industries with significant influence, causing aggregate volatility. Therefore, if the input-output matrix is unbalanced, aggregate fluctuations can result from distributions of productivity shocks with much thinner tails than the diversification argument would suggest.

The main extension to the approach of Acemoglu et al. (2012) that this paper formulates is to add a demand-side measure of influence. Acemoglu et al only consider the influence that industries have over other industries through the supply of inputs, but industries also have influence over other industries through the demand for inputs. If one industry demands a significant share of another industry's output as an input to production, then if the first industry suffers a negative shock, the second industry will have reduced demand for its products. In turn this reduction in demand for the second industry's products will then affect the industries that provide inputs to it, further propagating the shock upstream. Therefore, including a measure based on demand-side links between industries can provide a more complete idea of the influences of industries in an economy.

A consideration of demand-side shocks was included in Acemoglu et al. (2015), though this consideration is different to that presented in this paper. Acemoglu et al. (2015) include a government sector, and so represent demand shocks by exogenous changes to government consumption. Furthermore, they do not in-

clude a measure of industry influence through the demand side. In this paper, I instead represent demand shocks as changes in a “taste” parameter. This “taste” parameter alters the demand for the products of each industries in final consumption. Therefore, this gives a more general depiction of demand shocks than those which work exclusively through changes in government consumption. I also include a demand-influence vector, which is comparable to the supply-side measure, to give an overall idea of the industries through which an economy is most susceptible to demand shocks.

3 Economic Model

The economic model used in this paper is based on that used by Shea (2002), which is itself based on that of Long & Plosser (1983). There is a representative agent with preferences

$$U = \sum_{i=1}^N \delta_i \log(C_i) - L; \quad \delta_i = a_i \exp(d_i); \quad \sum_{i=1}^N a_i = 1. \quad (1)$$

Consumption depends on the consumption, C_i , of N goods, and hours worked L . The goods have stochastic preference weights, δ_i , with mean zero taste or demand shocks d_i .

Each good is produced through a Cobb-Douglas production process that utilizes labor and capital inputs.¹ Each of the goods are potentially used as inputs for the production of itself and the other goods.

$$Q_i = \lambda_i^{\alpha \gamma_i} L_i^{\alpha \gamma_i} \left[\prod_{k=1}^N (X_{ki})^{(1-\alpha)\beta_{ki}} \right]; \quad \alpha \gamma_i + \sum_{k=1}^N (1-\alpha)\beta_{ki} = 1. \quad (2)$$

Here, Q_i is the amount of good i produced, X_{ki} is the amount of good k used in the production of good i , α is the labor share in economy, L_i is the hours employed in the production of good i and $\lambda_i = \exp(s_i)$, where s_i are mean-zero supply shocks. Market clearing requires that:

$$Q_i = \sum_{k=1}^N X_{ik} + C_i; \quad L = L_i. \quad (3)$$

¹In theory each good is homogeneous and so should be considered to be produced by a single firm. Later, when I apply data to the model, the goods will be output produced by industries. Therefore, in this paper the terms ‘firm’ and ‘industry’ are used interchangeably.

Solving for the competitive equilibrium of the economy (see Appendix A for the derivation of this result), as per Acemoglu et al. (2012, p. 2005), gives the deviations from steady-state log GDP due to supply shocks as:

$$y_s = v_s^T \mathbf{s} + K; \quad \mathbf{v}_s = (\mathbf{I} - (1 - \alpha)\mathbf{W}^T)^{-1} \alpha \mathbf{l} \quad (4)$$

where v_s is the supply influence vector of the economy, $\mathbf{s}$ is the vector of supply shocks to each industry, K is a constant independent of the vector of shocks, α is the economy's labor share of production, $\mathbf{W}$ is the input-output matrix of the economy (normalized so that element w_{ji} is the share of industry j 's inputs provided by industry i), $\mathbf{I}$ is an $n \times n$ identity matrix (n being the number of industries) and $\mathbf{l}$ is the vector of the shares of the workforce employed in each industry.

This supply-side measure is fairly similar to that produced by Acemoglu et al. (2012, 2013).

Turning to the demand side (see Appendix B for the derivation of this result), the deviations from steady-state log GDP due to demand shocks will be:

$$y_d = v_d^T \mathbf{d} + M; \quad \mathbf{v}_d = (\mathbf{I} - (1 - Z)\mathbf{G}^T)^{-1} Z \mathbf{e} \quad (5)$$

where v_d is the demand influence vector of the economy, $\mathbf{d}$ is the vector of demand shocks to each industry, M is a constant independent of the vector of shocks, Z is the share of the economy's total production that is final production, $\mathbf{G}$ is the input-output matrix of the economy (normalized so that element g_{ji} is the share of total intermediate output produced by industry j and used by industry i as an input), $\mathbf{I}$ is an $n \times n$ identity matrix and $\mathbf{e}$ is the vector of the shares of the economy's final production produced by each industry.

Combining equations 4 and 5, gives the overall log deviations from steady-state GDP:

$$y = v_s^T \mathbf{s} + v_d^T \mathbf{d} + K + M; \quad (6)$$

This result is analogous to that produced by Shea (2002). Similarly to Shea's result, in this model shocks only propagate in one direction; supply shocks propagate to industries downstream, while demand shocks only affect industries upstream.

The total influence of an industry on the economy, as shown by Equation 6, is related to the size of the shocks that hit an industry as well as how those shocks

are propagated to other industries. For example, an industry might have a large influence coefficient (as measured by its value in the influence vectors, v_s and v_d) but only ever get hit by comparatively small shocks and so have fairly small influence on the deviations from steady state GDP. Alternatively, another industry's shocks might not propagate very far through the economy (small influence coefficient in v_s and v_d) but get regularly hit by very large shocks and so have a larger overall influence on the economy than that suggested by the influence coefficient. In this way, both the influence measures and the size of the shocks need to be considered together.

The influence measures presented in this paper can also be thought of as measures of centrality from network theory. In network theory, a node has greater centrality if it has greater influence on the other nodes in the networks and, hence, on the network as a whole. Under this interpretation, the input-output tables are descriptions of the input-output network of the economy and the influence measures are used to rank the industries (or nodes) as to how great an influence they have over other industries (nodes) and hence the economy as a whole. Specifically the influence measures used in this paper can be thought of as centrality measures based on the measures produced by Bonacich (1987) and Bonacich & Lloyd (2001) and are derivations based on network theory are presented in Section C in the Appendix.

4 Estimation

4.1 Level of Disaggregation

In estimating the measures of industry influence a question that must be addressed is what level of disaggregation is best. Data for input-output tables are provided at various levels of disaggregation depending on the country and the institution providing the data. Acemoglu et al. (2012, p. 1997) argue that, for their purposes, even the finest level of disaggregation available through the Bureau of Economic Analysis (BEA) tables (approximately four-digit SIC definition) is not fine enough. In their model, they require that the level of disaggregation is such that each sector (or firm) produces a homogeneous product and hence require a finer level of disaggregation than is generally provided by statistical institutions.

However, Acemoglu et al's model is based on industries being subject to

shocks that are independent of those hitting other industries. As some sectoral shocks are likely to directly affect more than one homogeneous product, an argument can be made for the level of disaggregation to be at the level at which shocks will generally hit a single industry or firm. This would potentially argue for a broader level of disaggregation and so, the data should be considered at a broader range such as at the industry or sub-industry level.

A further complication is when the interest is in cross-country comparisons. Readily-available, comparable data across countries tend to only be found at fairly broad level, such as that provided by the OECD. Therefore, it is at this level that such comparisons must be made, due to data limitations. In the estimations that follow I provide estimates at both a broad and a fine degree of disaggregation to try to explore the results based on these different purposes.

As will be shown, the discussion around the level of disaggregation is not trivial. Choosing a different level of disaggregation can lead to significantly different results when computing the influence vectors and determining how susceptible an economy is to shocks of a certain type.

4.2 Data

For the broad disaggregation estimates, I took the input-output data from the OECD StatExtracts online library, dated to “the mid 2000s”. The estimates of the labor share in production for Australia, US and Sweden were also taken from the OECD StatExtracts online library.

For the estimation of the Australian economy at the fine disaggregation level, I used the input-output data from the Australian Bureau of Statistics (ABS) Input-output Tables for 2012/13 (catalogue no. 5209.0.55.001), and the data used to estimate the labor shares by industry from the “ABS Labour Force (Detailed)” publication (catalogue no. 6291.0.55.003).

The Bureau of Economic Analysis (BEA) Input-output tables for 2007 provided the input-output data for the US, with the labor share by industry data from the Bureau of Labor Statistics.

The Swedish input-output data came from the Symmetric Input-output tables for 2010 from Statistics Sweden. The Labor share by industry data for Sweden came from the OECD iLibrary National Accounts for Sweden.

4.3 Cross-country comparisons

4.3.1 Broad disaggregation

The demand-side and supply-side influence of industries were estimated for Australia, US and Sweden. For ease of comparison, the data were taken from a common source, namely the OECD, which provides input-output tables for countries with the economies divided into 36 industries.² As stated previously, this level of disaggregation might be considered to be too broad (Acemoglu et al. 2012, p. 1997) but, nonetheless, this allows for more direct comparisons across countries.

The three countries were chosen because of their influence of the mining industry on their economies. As will be shown later the Australian mining industry gives a strong demonstration of the need for a demand-side measure to supplement the supply-side measure of industry influence, and so the demand and supply influence vectors were estimated for Australia. Sweden was chosen because like Australia it is a developed country with a large mining industry. As will be shown, Sweden's mining industry has more influence on the supply side than the demand side, showing that mining does not only influence a country's economy through the demand side. Thereby, Sweden gives a good contrast to the example of Australia. The US was included for completeness.

Summary statistics for the supply influence and demand influence vectors for Australia, US and Sweden are shown in Table 1. In Australia and the US, the demand influence vector is more skewed than the supply influence vector, while in Sweden the reverse is true. This suggests that Australia and the US are more susceptible to demand shocks, while Sweden is more susceptible to supply shocks (Acemoglu et al. 2012). Furthermore, if we were to only consider the supply influence vector, we would be downplaying the level of unbalancedness in the Australian and US economies and overplaying the unbalancedness in the Swedish economy. Nonetheless, even if we only considered the supply influence vector, we would still conclude that the Swedish economy is less unbalanced than the Australian or US economies, though not by as much as if we

²The data for Australia and Sweden are missing some industries. The Australian data only have 32 industries and the Swedish data have 34 industries, while the US data have the full complement of 36 industries. The missing industries in the Australian data are Office, Accounting & Computing Machinery; Radio, Television & Communication Equipment; Renting of Machinery & Equipment; and Research & Development. The missing industries in the Swedish data are Radio, Television & Communication Equipment; and Research & Development.

Table 1: Summary statistics of the Demand Influence vectors and Supply Influence vectors for Australia, US and Sweden at the broad level of disaggregation. The skewness statistic is the adjusted Fisher-Pearson standardized moment coefficient. Data: OECD. (Total industries: Australia 32, US 36, Sweden 34.)

	Demand Influence			Supply Influence		
	Aus	USA	Swe	Aus	USA	Swe
Mean (1/n)	0.031	0.028	0.029	0.031	0.028	0.029
Median	0.022	0.013	0.022	0.021	0.014	0.018
Std dev	0.031	0.034	0.023	0.027	0.031	0.026
Skewness	2.20	2.55	0.90	1.58	1.99	1.30
Maximum	0.139	0.167	0.082	0.127	0.127	0.096
Minimum	0.003	0.002	0.003	0.004	0.003	0.002

had the fuller picture.

The ten most influential industries in each country according to the demand-side measure are shown in Table 2, and according to the supply-side measure in Table 3. For Australia, the most influential industries appear to be Construction and Wholesale & Retail Trade, given that they are very influential on both the demand- and supply-side. For the US, the most influential industries on both the demand- and supply-side are Public Administration & Defense, and Wholesale & Retail trade and for Sweden the most influential industry is Health & Social Work. In Australia, six industries appear in both the demand-side and supply-side lists. For the US eight industries are in both list and for Sweden seven industries do.

There is a strong similarity across these countries in the industries that are most influential, which is consistent with other research (Jones 2011), with more similarities on the supply-side. While the order of industry influence varies among the countries, on the demand-side four industries appear in all three countries' lists, six industries appear in two of the lists with a further six industries only appearing in one of the lists. On the supply-side, five industries appear in all three countries lists, six appear in two of the lists and a further three appear in only one of the lists. With the similarities between these countries' most influential industries, it is also worth noting how central service industries are to these economies. On the demand side, seven of the ten most influential

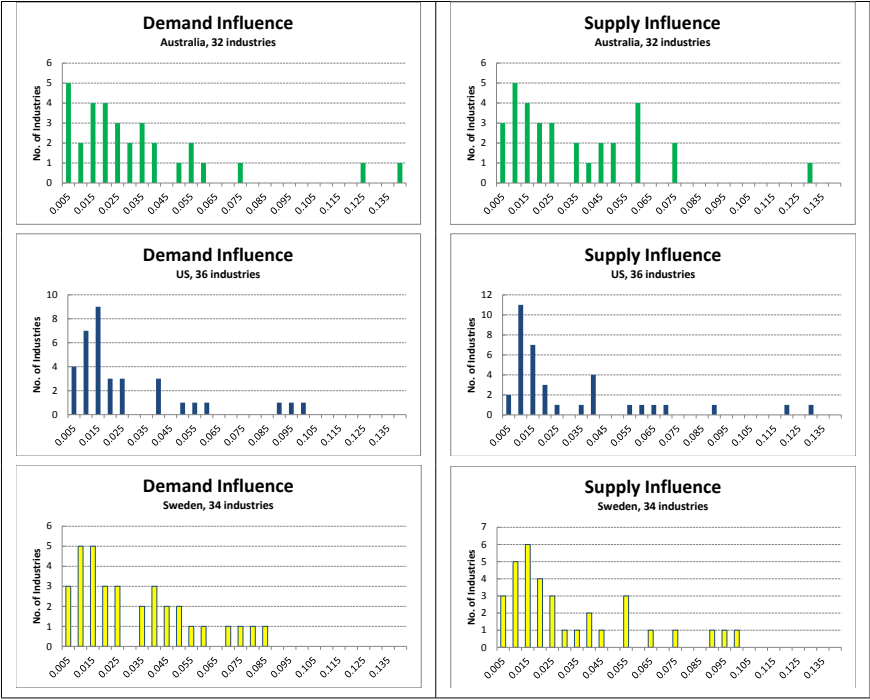


Figure 1: Histograms of the Demand and Supply influence vectors (broad disaggregation) in Australia, US and Sweden. Source: OECD.

Table 2: The ten most influential industries in Australia, the US and Sweden according to the Demand Influence Vector (Broad disaggregation). The number in brackets is the measure of demand-side influence of the industry. Data: OECD. (Total industries: Australia 32, US 36, Sweden 34.)

Rank	Australia	US	Sweden
1	Construction (0.139)	Public Administration, Defense & Compulsory Social Security (0.105)	Health & Social Work (0.082)
2	Wholesale & Retail trade and Repairs (0.123)	Wholesale & Retail trade and Repairs (0.073)	Wholesale & Retail trade and Repairs (0.079)
3	Real Estate Activities (0.074)	Real Estate Activities (0.062)	Real Estate Activities (0.072)
4	Transport & Storage (0.059)	Health & Social Work (0.059)	Motor Vehicles, Trailers & Semi-trailers (0.068)
5	Public Administration, Defense & Compulsory Social Security (0.055)	Construction (0.046)	Other Machinery & Equipment Manufacturing (0.056)
6	Food Products, Beverages & Tobacco (0.054)	Finance & Insurance (0.038)	Transport & Storage (0.054)
7	Health & Social Work (0.048)	Other Community, Social & Personal Services (0.029)	Public Administration, Defense & Compulsory Social Security (0.049)
8	Hotels & Restaurants (0.039)	Food Products, Beverages & Tobacco (0.023)	Education (0.047)
9	Other Business Activities (0.038)	Motor Vehicles, Trailers & Semi-trailers (0.023)	Other Electrical Machinery & Apparatus Manufacturing (0.042)
10	Mining & Quarrying (0.035)	Hotels & Restaurants (0.023)	Other Business Activities (0.042)

Table 3: The ten most influential industries in Australia, the US and Sweden according to the Supply Influence Vector (Broad disaggregation). The number in brackets is the measure of supply-side influence of the industry. Data: OECD. (Total industries: Australia 32, US 36, Sweden 34.)

Rank	Australia	US	Sweden
1	Wholesale & Retail trade and Repairs (0.127)	Wholesale & Retail trade and Repairs (0.121)	Transport & Storage (0.096)
2	Other Business Activities (0.073)	Public Administration, Defense & Compulsory Social Security (0.112)	Health & Social Work (0.090)
3	Construction (0.070)	Other Business Activities (0.084)	Other Business Activities (0.088)
4	Computer & Related Activities (0.060)	Health & Social Work (0.064)	Wholesale & Retail trade and Repairs (0.070)
5	Finance & Insurance (0.060)	Other Community, Social & Personal Services (0.057)	Public Administration, Defense & Compulsory Social Security (0.063)
6	Health & Social Work (0.058)	Finance & Insurance (0.055)	Electricity, Gas & Water Supply (0.053)
7	Transport & Storage (0.057)	Hotels & Restaurants (0.052)	Education (0.051)
8	Hotels & Restaurants (0.047)	Construction (0.036)	Construction (0.051)
9	Other Community, Social & Personal Services (0.046)	Real Estate Activities (0.034)	Real Estate Activities (0.042)
10	Education (0.044)	Transport & Storage (0.033)	Mining & Quarrying (0.040)

Table 4: Summary statistics of the Demand Influence vectors and Supply Influence vectors for Australia, US and Sweden at the fine level of disaggregation. The skewness statistic is the adjusted Fisher-Pearson standardized moment coefficient. Data: ABS, BEA, Statistics Sweden. (Total industries: Australia 114, US 382, Sweden 59.)

	Demand Influence			Supply Influence		
	Aus	USA	Swe	Aus	USA	Swe
Mean	0.009	0.002	0.017	0.009	0.003	0.017
Median	0.003	0.001	0.009	0.004	0.001	0.011
Std dev	0.014	0.006	0.017	0.011	0.006	0.017
Skewness	3.17	7.04	1.68	2.26	4.82	1.59
Maximum	0.090	0.062	0.082	0.060	0.050	0.073
Minimum	0.000	0.000	0.001	0.000	0.000	0.000

industries are service industries in each of the countries; while on the supply-side, in Australia and the US, nine of the ten most influential industries are service industries and in Sweden, the seven most influential industries are service industries.

Therefore, these directly comparable data show that the industries that are most central to these three economies are mostly service industries, with Australia and the US more susceptible to demand shocks, while Sweden is more susceptible to supply shocks.

4.3.2 Fine disaggregation

The influence vectors were then estimated using the more disaggregated data available from the countries' statistical bureaus. These data are not directly comparable across countries because the countries use different definitions for their industries and have different levels of disaggregation relative to each other. However, this is the finest level of disaggregation possible with publicly available data and so is closest to the homogeneous-product level of disaggregation recommended by Acemoglu et al. (2012).

Summary statistics for the influence vectors at the finer level of disaggregation are provided in Table 4 and their histograms are shown in Figure 2.

At the finer levels of disaggregation the influence vectors become more skewed

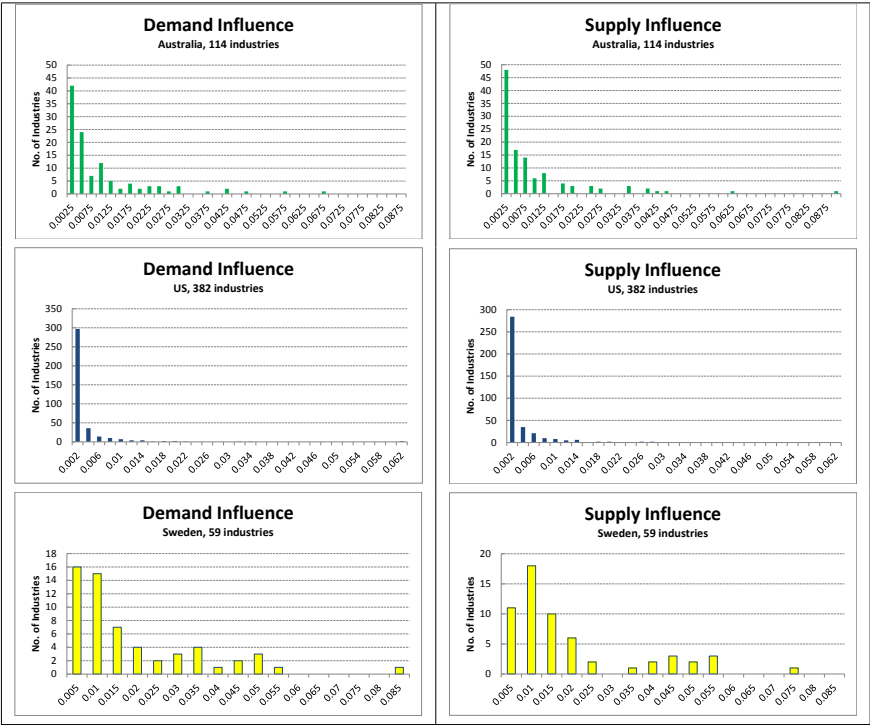


Figure 2: Histograms of the Demand and Supply influence vectors (fine disaggregation) in Australia, US and Sweden. Source: ABS, BEA, Statistics Sweden.

for each of the countries. For the US and Sweden the demand influence vector is more skewed than the supply influence vector, while in Australia the reverse is the case. This suggests that the US and Sweden are somewhat more susceptible to sectoral demand shocks than supply shocks. Furthermore, it suggest that if we only considered the supply influence vector, we would conclude that the US and Sweden are less susceptible to sectoral shocks causing aggregate volatility than is actually the case. This result contrasts somewhat with that found at the broad disaggregation level where in Australia the demand influence vector was more skewed than the supply influence vector while in Sweden the reverse was the case. Therefore, how we interpret the likelihood of demand or supply shocks causing aggregate volatility can be affected simply by the level of disaggregation used in the estimations.

This paper does not provide a model for determining the ‘ideal’ level of disaggregation for constructing the influence vectors. However, given the model’s specification that sectoral shocks should be independent of each other, disaggregating at the level at which sectoral shocks are independent would seem a reasonable solution. Naturally, data constraints might well have more sway over this decision than theory.

5 Case study: Australian mining industry

The mining industry in Australia over recent years provides a stark example of why there is a need in the framework of the model of this paper to include the demand side influence of industries. This case study is included to show how an industry can have significant influence on the aggregate outcomes of an economy even though it rates quite low in terms of influence from the supply side model of Acemoglu et al. (2012).

5.1 Recent history of the Australian mining industry

In the early 2000s, the mining industry in Australia was largely derided as an example that the Australian economy was being held back because it was too dependent on the ‘old economy’ industries of mining, agriculture and related industrial manufacturing, rather than the ‘new economy’ of hi-tech manufacturing and services. Nonetheless, over the course of the following decade the mining industry experienced one of the strongest booms in Australian economic

history as the prices of key commodities rose to historically high levels causing a surge in mining investment and production, particularly in coal, iron ore and liquefied natural gas (Connolly & Orsmond 2011). As a result of the resource boom the mining industry is now the largest industry in Australia in terms of real GVA (Figure 3). This growth has been marked by sustained strong growth in the GVA of iron ore mining since the late 1990s, while more recently the growth in oil & gas extraction and coal mining have also contributed to growth.

The turnaround in resource commodity prices, which had been at their lowest inflation-adjusted level for decades at the start of the 2000s, has largely been attributed to emerging economies increasing demand for inputs to energy and steel production (Connolly & Orsmond 2011, p. 6). The increase in commodity prices caused extremely strong growth in mining investment as mining companies sought to expand their capacity. This increase in mining investment and the resulting increase in capacity saw revenue from the mining industry increase from 6 per cent of nominal GDP in 2000 to 14 per cent by 2010, and mining investment rise from $1\frac{1}{2}$ per cent of GDP to over 4 per cent over the same period (Connolly & Orsmond 2011, pp. 12-13). Furthermore, the mining industry has contributed around 12 per cent of the growth in Australian GDP since 2000.

While the mining industry is now the largest industry in the Australia economy, its influence is generally considered to be even greater than its share of GDP would suggest due to its use of inputs from other industries. The resource economy broadly defined, that is including activities supporting the mining industry, has been estimated at around 18 per cent of gross value added in 2011/12, double its share in 2003/04 (Rayner & Bishop 2013).

5.2 The mining industry in the Australian supply and demand influence vectors

The recent history of the Australian mining industry indicates that it has had a significant and growing influence on Australia's economic fortunes in recent years. However, when the supply influence vector is calculated for Australia, the mining industry is estimated to be one of the least influential industries in the Australian economy. Using the ABS data and aggregating the data into 19 broad industry categories, mining is estimated to be the 16th most influential industry on the supply side (Table 5). At the finer level of disaggregation (114 industries), the mining sub-industries are again estimated to be compara-

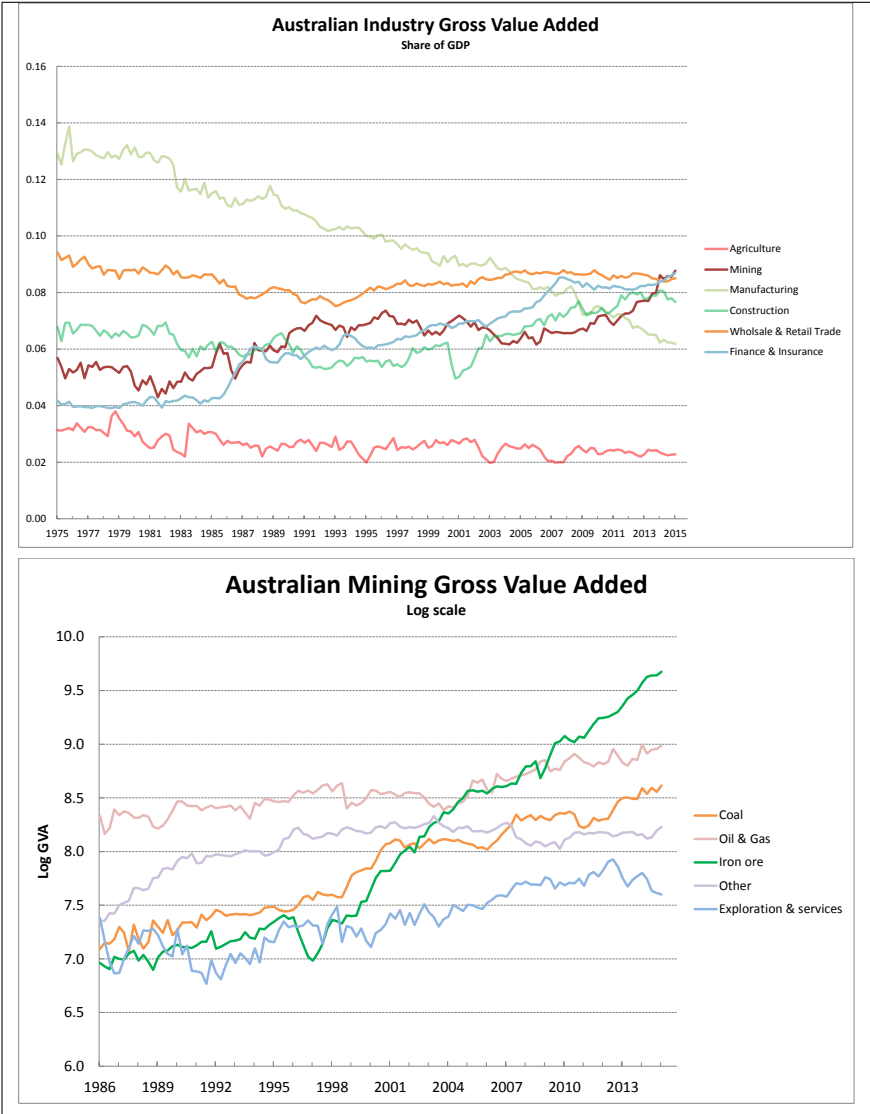


Figure 3: Selected Australian industries' real gross value added as a share of real GDP. Log real GVA of Australian mining sub-industries. Source: ABS.

tively lowly-ranked in terms of influence (Oil & Gas extraction 38th; Non-ferrous metal ore mining 50th; Coal mining 57th; Iron ore mining 67th; Exploration & mining support services 69th; and Non-metallic mineral mining 95th). This low ranking of the mining industries compared with their recent performance suggests that something is missing in this measure of industry influence.

In contrast, when the demand influence vectors are estimated, the mining industry is estimated to be the fourth most influential industry at the broad disaggregation level (behind Construction; Manufacturing; and Rental, Hiring & Real Estate services), while most of the sub-industries are also estimated to be far more influential than the supply influence vector would indicate (Coal mining 14th; Iron ore mining 19th; Oil & gas extraction 25th; Non-ferrous metal ore mining 27th; Exploration & mining support services 51st; and Non-metallic mineral mining 95th). While, according to these measures, the mining industry is still not the most influential industry in Australia, these results indicate that the demand-side measure is capturing an important source of influence that should not be ignored.

Interestingly, the situation in Australia where the mining industry is more influential on the demand-side than the supply side contrasts with that in the US and Sweden where the reverse is the case. Using the OECD data, Mining & Quarrying is the least influential industry on the demand side in Sweden but the tenth most influential on the supply side. In the US, it is the 28th most influential on the demand side but the 21st on the supply side. This suggests that, while the mining industry in Australia is not a large source of inputs to the Australian economy, it is in Sweden. This is perhaps reflected in the generally higher ranking of industrial manufacturing industries in the Swedish economy (such as Motor Vehicle manufacturing and Other Machinery & Equipment manufacturing) compared with Australia.

The importance of considering the demand-side influence measure is shown starkly when viewing the evolution of the influence measures of the mining industry's sub-industries during the time of the mining boom. Figure 4 shows how the different measure of influence have changed through the years from 1999 to 2012. Not only do the sub-industries have greater absolute influence on the demand side, the demand measures also show a strong increase in the influence of the mining industries over this time. This increase in influence is especially strong for Coal mining and Iron Ore mining, with a smaller increases in Oil & Gas extraction and Non-ferrous Metal mining. In comparison, growth

Table 5: Australian industries and their rankings according to the demand-side and the supply-side measures at a broad disaggregation level. Source: ABS, 19 industries.

Rank	Demand Influence	Supply Influence
1	Construction	Manufacturing
2	Manufacturing	Professional, Scientific & Technical Services
3	Rental, Hiring & Real Estate Services	Construction
4	Mining	Administrative & Support Services
5	Public Administration & Safety	Health Care & Social Assistance
6	Health Care & Social Assistance	Financial & Insurance Services
7	Transport, Postal & Warehousing	Retail Trade
8	Financial & Insurance Services	Transport, Postal & Warehousing
9	Retail Trade	Rental, Hiring & Real Estate Services
10	Education & Training	Education & Training
11	Wholesale Trade	Accommodation & Food Services
12	Professional, Scientific & Technical Services	Public Administration & Safety
13	Accommodation & Food Services	Wholesale Trade
14	Information, Media & Telecommunications	Information, Media & Telecommunications
15	Administrative & Support Services	Other Services
16	Agriculture, Forestry & Fishing	Mining
17	Other Services	Agriculture, Forestry & Fishing
18	Electricity, Gas, Water & Waste Services	Electricity, Gas, Water & Waste Services
19	Arts & Recreation Services	Arts & Recreation Services

in the supply- influence measures over this time are far more muted, with only Oil & Gas extraction showing a significant increase in influence. While Iron Ore mining's influence grew over this time, it remained at a very low-level despite that growth.

This adds more weight to the argument that including the demand-side measure adds a significant source of economic influence and so should be included alongside the supply-side measure.

6 Conclusion

In this paper, I add to the literature on how sectoral shocks can affect aggregate volatility by constructing a measure of industry influence through demand channels to upstream suppliers. While previous work has attributed influence to industries that provided significant shares of other industries' inputs, the demand-side measure that I use considers the influence that derives from industries using the inputs provided by other industries. By including this channel to the supply side effects of sectoral shocks, a more complete picture can be generated as to how susceptible economies are to sectoral shocks and which industries are key to those economies. What this analysis suggests is that, by only focusing on the supply-side influences of industries, we will likely miss significant sectoral influences on aggregate volatility. A stark example of this shown in this article is that of the mining industry in Australia, which has had a very large influence on the economic fortunes of Australia over recent years. Using only a supply-side measure of influence would significantly downplay the influence of the mining industry in Australia. Adding the demand-side measure gives the mining industry a more prominent position in the ranking of Australian industries.

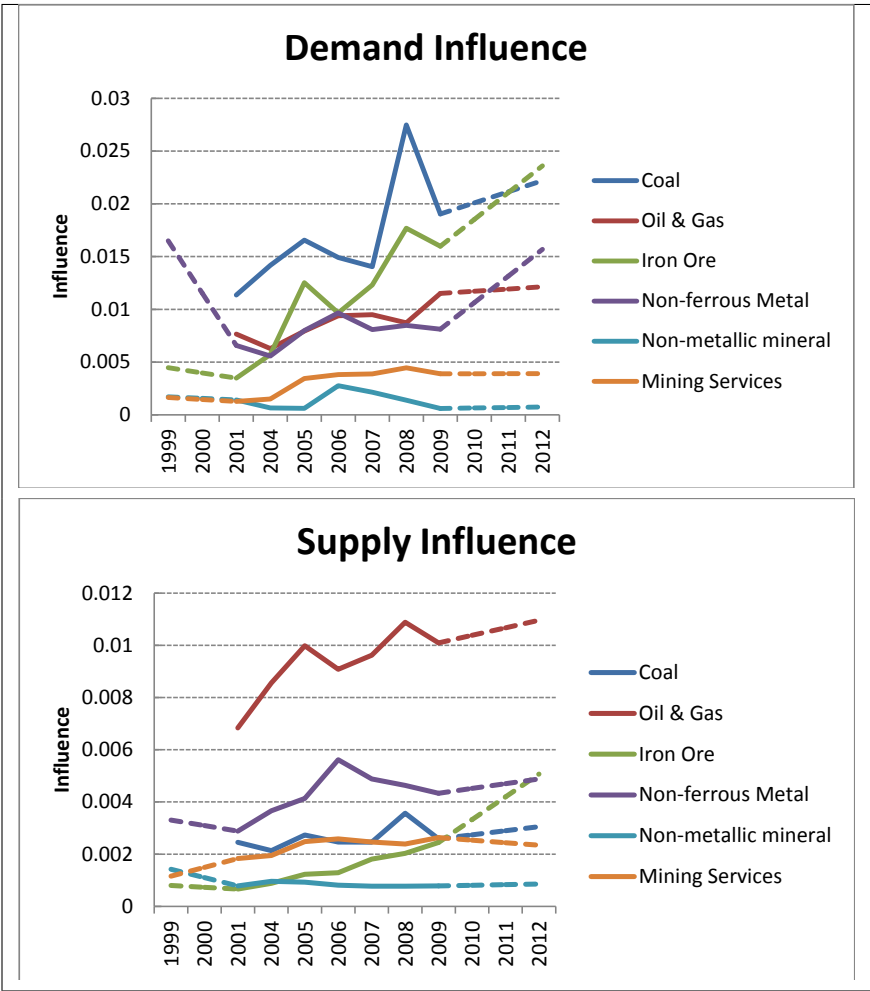


Figure 4: Influence measures of the mining sub-industries in Australia for the period 1999-2012. Data for 2000, 2010 and 2011 interpolated as Input-Output tables were not published for these years. In 1999, the Coal and Oil & Gas industries were combined in the Input-Output tables and so their separate data only begins in 2001. Source: ABS.

References

- Acemoglu, D., Akcigit, U. & Kerr, W. (2015), Networks and the macroeconomy: An empirical exploration, Working Paper 21344, National Bureau of Economic Research.
- Acemoglu, D., Carvalho, V. M., Ozdaglar, A. & Tahbaz-Salehi, A. (2012), 'The network origins of aggregate fluctuations', *Econometrica* **80**(5), 1977–2016.
- Acemoglu, D., Ozdaglar, A. & Tahbaz-Salehi, A. (2013), The network origins of large economic downturns, Working Paper 19230, National Bureau of Economic Research.
- Aroche Reyes, F. & Marquez Mendoza, M. A. (2012), An economic network in North America, MPRA Paper 61391, University Library of Munich, Germany.
- Atalay, E. (2014), How Important Are Sectoral Shocks, Working Papers 14-31, Center for Economic Studies, U.S. Census Bureau.
- Bak, P., Chen, K., Scheinkman, J. & Woodford, M. (1992), Aggregate fluctuations from independent sectoral shocks: Self-organized criticality in a model of production and inventory dynamics, Working Paper 4241, National Bureau of Economic Research.
- Bonacich, P. (1987), 'Power and Centrality: A Family of Measures', *American Journal of Sociology* **92**(5), 1170–1182.
- Bonacich, P. & Lloyd, P. (2001), 'Eigenvector-like measures of centrality for asymmetric relations', *Social Networks* **23**, 191–201.
- Connolly, E. & Orsmond, D. (2011), The Mining Industry: From Bust to Boom, Research Discussion Paper 2011-08, Reserve Bank of Australia.
- Dixon, H., Franklin, J. & Millard, S. (2014), Sectoral shocks and monetary policy in the United Kingdom, Working Paper 499, Bank of England.
- Foerster, A., Sarte, P. D. & Watson, M. (2011), 'Sectoral versus aggregate shocks: A structural factor analysis of industrial production', *Journal of Political Economy* **119**(1), 1 – 38.
- Gabaix, X. (2011), 'Granular origins of aggregate fluctuations', *Econometrica* **79**(3), 733–772.

- Harvey, D. I. & Mills, T. C. (2002), 'Common features in UK sectoral output', *Economic Modelling* **19**(1), 91 – 104.
- Jackson, M. O. (2008), *Social and Economic Networks*, Princeton University Press, Princeton, New Jersey.
- Jones, C. I. (2011), Misallocation, Economic Growth, and Input-Output Economics, NBER Working Papers 16742, National Bureau of Economic Research.
- Jovanovic, B. (1987), 'Micro shocks and aggregate risk', *The Quarterly Journal of Economics* **102**(2), 395–409.
- Lawson, J. & Rees, D. (2008), A Sectoral Model of the Australian Economy, Research Discussion Paper 2008-01, Reserve Bank of Australia.
- Long, J. B. & Plosser, C. I. (1983), 'Real Business Cycles', *Journal of Political Economy* **91**(1), 39–69.
- Mehrotra, N. & Sergeyev, D. (2013), Sectoral Shocks, the Beveridge Curve and Monetary Policy, 2013 Meeting Papers 919, Society for Economic Dynamics.
- Rayner, V. & Bishop, J. (2013), Industry dimensions of the resource boom: An input-output analysis, Research Discussion Paper 2013-02, Reserve Bank of Australia.
- Roson, R. & Sartori, M. (2014), Input-output linkages and the propagation of domestic productivity shocks: Assessing alternative theories with stochastic simulation, MPRA Paper 59884, University Library of Munich, Germany.
- Shea, J. (2002), 'Complementarities and Comovements', *Journal of Money, Credit and Banking* **34**(2), 412–433.
- Storer, P. (1996), 'Separating the effects of aggregate and sectoral shocks with estimates from a Markov-switching search model', *Journal of Economic Dynamics and Control* **20**(1-3), 93–121.

A Supply shocks derivation

The derivation of the supply influence vector that follows is based largely on that in Acemoglu et al. (2012, p. 2005) with the only difference being that in Acemoglu et al's model $\gamma_i = 1, \forall i$.

The first order conditions for industry i 's production function (Equation 2) for labor, L_i , and inputs, X_{ki} , gives $L_i^* = \frac{P_i Q_i \alpha \gamma_i}{h}$ and $X_{ki}^* = \frac{P_i Q_i (1-\alpha) \beta_{ki}}{P_k}$, where h is the market wage, and P_j is the price of good j . Substituting these values into industry i 's production function and taking logs yields:

$$\alpha \gamma_i \log(h) = \alpha \gamma_i s_i + B + \log(P_i) - (1-\alpha) \sum_{k=1}^N \beta_{ki} \log(P_k) + (1-\alpha) \sum_{k=1}^N \beta_{ki} \log(\beta_{ki})$$

where $B = \alpha \gamma_i \log(\alpha \gamma_i) + (1-\alpha) \log(1-\alpha)$. Dividing the above equation through by $\alpha \gamma_i$, multiplying by the i th element of the supply influence vector $v_s^T = \alpha \mathbf{I}^T (\mathbf{I} - (1-\alpha)\mathbf{W})^{-1}$ and summing over all sectors gives:

$$y_s = \log(h) = v_s^T \mathbf{s} + K$$

where K is a constant independent of the vector of shocks equal to:

$$K = \sum_{i=1}^N \frac{l_i}{\gamma_i} \log(P_i) + \frac{B}{\alpha} \sum_{i=1}^N \frac{v_i}{\gamma_i} + \frac{1-\alpha}{\alpha} \sum_{i=1}^N \sum_{k=1}^N \frac{v_i}{\gamma_i} \beta_{ki} \log(\beta_{ki})$$

B Demand shocks derivation

The first order condition for the consumer's utility function for consumption of good i yields, $C_i^* = \frac{h\delta_i}{P_i}$. Substituting this value and the optimum demand for industry i ' output as inputs, $X_{ik}^* = \frac{P_k Q_k (1-\alpha) \beta_{ik}}{P_i}$, into the resource constraint gives:

$$\begin{aligned} Q_i &= \sum_{k=1}^N X_{ik} + C_i \\ &= \sum_{k=1}^N \left[\frac{P_k Q_k (1-\alpha) \beta_{ik}}{P_i} \right] + \frac{h\delta_i}{P_i} \end{aligned}$$

Sum over all i firms

$$Q = \sum_{i=1}^N \sum_{k=1}^N \left[\frac{P_k Q_k (1-\alpha) \beta_{ik}}{P_i} \right] + \sum_{i=1}^N \frac{h\delta_i}{P_i}$$

A log linearization of this result yields:

$$q_d \approx \sum_{i=1}^N \sum_{k=1}^N \left[\frac{P_k Q_k (1-\alpha) \beta_{ik}}{P_i} \right] \frac{1}{Q} q_k + \sum_{i=1}^N \frac{h a_i}{P_i} \frac{1}{Q} d_i + R$$

where R is a constant independent of the demand shocks. The deviation from steady-state output attributable to deviations in individual industry i 's output is then:

$$q_{di} \approx \sum_{k=1}^N \left[\frac{P_k Q_k (1-\alpha) \beta_{ik}}{P_i} \right] \frac{1}{Q} q_k + \frac{h a_i}{P_i} \frac{1}{Q} d_i + R_i$$

At this point we can define some terms. The share of total output that is final output in the steady state

$$Z = \sum_{i=1}^N \frac{ha_i}{P_i} \frac{1}{Q}$$

The share of total output that is intermediate output in the steady state

$$1 - Z = \sum_{i=1}^N \sum_{k=1}^N \left[\frac{P_k Q_k (1 - \alpha) \beta_{ik}}{P_i} \right] \frac{1}{Q}$$

The share of total final output that is produced by industry i

$$e_i = \frac{\frac{ha_i}{P_i}}{\sum_{i=1}^N \frac{ha_i}{P_i}}$$

The share of total intermediate output produced by industry i that is demanded by industry k

$$g_{ik} = \frac{\frac{P_k Q_k (1 - \alpha) \beta_{ik}}{P_i}}{\sum_{k=1}^N \frac{P_k Q_k (1 - \alpha) \beta_{ik}}{P_i}}$$

Substituting these expressions into the equation for industry i 's log deviations from steady state output due to demand shocks gives:

$$q_{di} \approx \sum_{k=1}^N (1 - Z) g_{ik} q_k + e_i Z d_i + R_i$$

Solving this for the log linearization of deviations from steady state total output due to demand shocks gives:

$$y_d = v_d^T \mathbf{d} + M; \quad \mathbf{v}_d = (\mathbf{I} - (1 - Z)\mathbf{G}^T)^{-1} \mathbf{Z} \mathbf{e}$$

C Influence Measures and Network Theory

In network theory, a network is a collection of connected nodes and, in general, a node is considered more central to a network if it has greater influence on, or importance to, the network as a whole. The various ways to measure which node in a network is most central to the network can be categorized into four broad groups, based on the node's characteristics that are mainly used in the measurement. These groups are: degree (how many other nodes the node is connected to); closeness (how near the node is to other nodes); betweenness (how important the node is for connecting other nodes), and neighbors' characteristics (how important a node's neighbors are) (Jackson 2008, p. 37). The measures used in this paper to estimate industries' influence on the economy are based on those produced by Bonacich (1987) and Bonacich & Lloyd (2001), which are measures based on the characteristics of a node's neighbors. That is, a node is central to the network if it is connected to, and has influence over, other important nodes.

The general measure of centrality of Bonacich & Lloyd (2001) gives a node greater centrality based on the influence that the node has from exogenous sources and its endogenously-derived influence from connections to other important nodes. A relative weighting is then applied between the exogenous and endogenous sources of influence to arrive at the final measure of centrality. The general form of the Bonacich measure of centrality is:

$$x = \alpha A^T x + e \quad (7)$$

where x is the vector of node centrality; α is a scalar parameter that reflects the relative importance of endogenous versus exogenous factors in a node's centrality; A is an adjacency matrix reflecting the direct influence of nodes on each other, where the element a_{ij} is the influence that i has over j ; and e is a vector of exogenous sources of influence. Solving Equation 7 for the vector of node centrality, x , yields:

$$x = (I - \alpha A^T)^{-1} e \quad (8)$$

where I is the identity matrix.

In the model presented in this paper, industries are nodes in the input-output matrix. The exogenous influence of an industry is their direct influence

on the economy, while the endogenous source of influence is their indirect influence on the economy through their effect on other industries. The exogenous influences that industries have on the economy are their use of labor in production and their production of final goods. The effects that industries have on other industries in this model are through their provision of inputs to, or demands for inputs from, other industries. An industry can, therefore, have influence on the economy through the supply-side or the demand-side and, as such, two measures of influence are produced for each industry.

C.1 Supply Influence

The influence vector in the article by Acemoglu et al. (2012, p. 1985) is:

$$\mathbf{v} \equiv \frac{\alpha}{n} [\mathbf{I} - (1 - \alpha)\mathbf{W}^T]^{-1} \mathbf{1} \quad (9)$$

where α is the economy's labor share of production, n is the number of industries, $\mathbf{I}$ is an $n \times n$ identity matrix, $\mathbf{W}$ is the input-output matrix of the economy (normalized so that element w_{ji} is the share of industry j 's inputs provided by industry i), and $\mathbf{1}$ is an n -dimensional vector of ones. This measure differs slightly from that of Bonacich in that the weights on the exogenous and endogenous sources of influence (α and $1 - \alpha$) sum to one.

The measure of an industry's supply-influence presented in this article is conceptually similar to that produced by Acemoglu et al. (2012). Under this measure an industry is considered more influential to the economy, if it employs a large share of the economy's workforce or if it provides a large share of the inputs to other industries' production. These two supply influences of an industry are weighted according to the labor share in production of the economy to give the supply influence of industry i :

$$v_{si} = (1 - \alpha)\mathbf{W}^T \mathbf{v}_s + \alpha l_i \quad (10)$$

where α is the economy's labor share of production, $\mathbf{W}$ is the input-output matrix of the economy (normalized so that element w_{ji} is the share of industry j 's inputs provided by industry i), and l_i is the share of the workforce employed in industry i . Solving Equation 10 for the supply-influence vector of the economy gives:

$$\mathbf{v}_s = (\mathbf{I} - (1 - \alpha)\mathbf{W}^T)^{-1} \alpha \mathbf{l} \quad (11)$$

where $\mathbf{I}$ is an $n \times n$ identity matrix (n being the number of industries) and $\mathbf{l}$ is the vector of the shares of the workforce employed in each industry. The difference between the supply-influence vector here and that of Acemoglu et al is that Acemoglu et al implicitly assume that every industry employs the same share of the labor force, while in my model the shares of the labor force used by each industry are allowed to differ.

C.2 Demand Influence

An industry has a large influence on the economy through the demand side if it demands a large share of other industries' output as inputs to its own production, or if it produces a large share of the economy's final production. The exogenous and endogenous sources of demand influence are weighted by the share of final production in the economy's total production. The demand influence of industry i is:

$$v_{di} = (1 - Z)\mathbf{G}^T v_d + Z e_i \quad (12)$$

where Z is the share of the economy's total production that is final production, $\mathbf{G}$ is the input-output matrix of the economy (normalized so that element g_{ji} is the share of total intermediate output produced by industry j and used by industry i as an input), and e_i is the share of the economy's final production produced by industry i .

Solving Equation 12 for the demand influence vector of the economy gives:

$$\mathbf{v}_d = (\mathbf{I} - (1 - Z)\mathbf{G}^T)^{-1} Z \mathbf{e} \quad (13)$$

where $\mathbf{I}$ is an $n \times n$ identity matrix and $\mathbf{e}$ is the vector of the shares of the economy's final production produced by each industry.